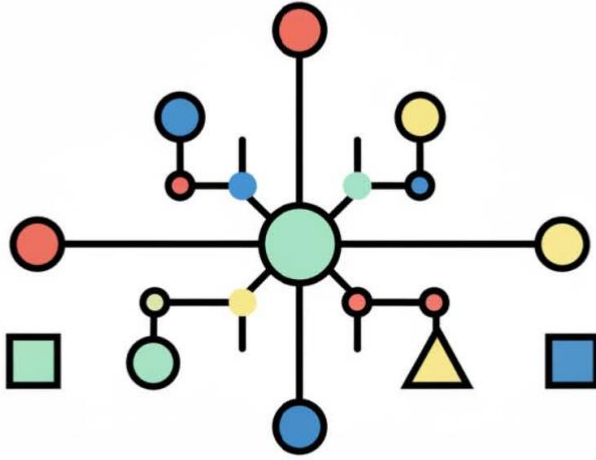


ISBN: 978-625-93538-0-7

Çağdaş Dilbilim Yazıları I Contemporary Studies in Linguistics I

İMÜ Dilbilimi 15. Yıl Armağanı
IMU Linguistics 15th Anniversary Commemorative Volume



Editörler / Edited by
Oktay ÇINAR, Fırat BAŞBUÇ, Hakan AYDEMİR





Çağdaş Dilbilim Yazıları I
Contemporary Studies in Linguistics I

İMÜ Dilbilimi 15. Yıl Armağanı
IMU Linguistics 15th Anniversary Commemorative
Volume

Editörler / Edited by
Oktay ÇINAR, Fırat BAŞBUĞ, Hakan AYDEMİR



Yayın Kurulu / Editorial Board:

Prof. Dr. Marcel Erdal - Goethe University Frankfurt, Germany
Prof. Dr. Christoph Schroeder - University of Potsdam, Germany
Prof. Dr. Jeanine Treffers-Daller - University of Reading, UK
Dr. Serkan Uygun – Bahçeşehir University, Türkiye

Yayıncı/Publisher: **Artsürem Bilim Sanat Danışmanlık Bilişim AŞ.**
Oğulbey Mahallesi, K.Evleri, 571, Gölbaşı-Ankara, Türkiye
Tel: +90 530 922 90 90 bilgi@artsurem.tr www.artsurem.tr

T.C. Kültür ve Turizm Bakanlığı Yayıncı Sertifika Numarası /
Publisher Certificate Number: **46183**

ISBN: **978-625-93538-0-7**

Açık Erişim Tarihi / Open Access Date: **31 Aralık 2025**

Açık Erişim Adresi / Open Access URL:
<https://artsurem.tr/index.php/books/catalog/book/2>

DOI: **10.7816/imuling-15-2025-cl56**

Kapak Tasarımı / Cover Design: **Fırat BAŞBUĞ**

Kapak Görseli / Cover Image: **Nano Banana**

Series / Dizi: **Contemporary Studies in Linguistics.**

Çınar, O., Başbuğ, F., & Aydemir, H. (Edl.). (2025). Çağdaş Dilbilim
Yazıları I: İMÜ Dilbilimi 15. Yıl Armağanı. Artsürem.

Telif Hakkı ve Açık Eriřim Lisansı: © 2025 Yazarlar ve Editörler.

Copyright & Open Access License: © 2025 The Authors and Editors. Published by Artsürem Publishing.

Artsürem Yayıncılık tarafından yayımlanmıştır. Bu eser, yazarın ve kaynağın usulüne uygun olarak belirtilmesi kaydıyla, herhangi bir ortamda ticari olmayan kullanıma, dağıtıma ve çoğaltmaya izin veren Creative Commons Atıf-GayriTicari 4.0 Uluslararası Lisansı (CC BY-NC 4.0) koşulları altında dağıtılan açık erişimli bir kitaptır. Önceden izin alınmaksızın eserde deęişiklik yapılamaz veya türetilmiş eserler oluşturulamaz.

This is an open-access book distributed under the terms of the Creative Commons Attribution-NonCommercial 4.0 International License (CC BY-NC 4.0), which permits non-commercial use, distribution, and reproduction in any medium, provided the original author and source are credited. No derivatives or modifications are allowed without prior permission.

Yayın Etięi: Artsürem Yayıncılık, Yayın Etięi Komitesi (COPE) tarafından belirlenen etik ilkelere ve en iyi uygulamalara sıkı sıkıya baęlıdır.

Publishing Ethics: Artsürem Publishing adheres strictly to the ethical guidelines and best practices set forth by the Committee on Publication Ethics (COPE).

Hakem Deęerlendirme Beyanı: Bu kitap, yayına kabul edilmeden önce ilgili alandaki iki bağımsız, dış akademik uzman tarafından çift kör hakemlik sürecinden geçirilmiştir.

Peer Review Statement: This edited volume has been subjected to a double-blind peer review process by two independent, external academic experts in the relevant field prior to publication acceptance.

İÇİNDEKİLER / TABLE OF CONTENTS

Önsöz / Preface	1
Douglas Q. Adams <i>Indeterminacy in etymology: Reconstructing the PIE word for ‘sheep’ with notes also on Tocharian B ās ‘goat’ and PIE *h₂éwis ‘bird’</i>	5
Hakan Aydemir <i>Morphophonological, morphosyntactic, and iconic constraints governing irreversible nominal and verbal bicoordinatives in Turkish</i>	13
Fırat Başbuğ <i>Grasping grammar: The haptic turn in language documentation</i>	65
Oktay Çınar <i>Çevrim içi dilbilim araştırmaları için PCIBex Farm: Tasarım ve uygulama</i>	81
Marcel Erdal <i>Distinctive diacritics in the Atil-Turkic epitaphs and in the Yarkand documents</i>	109
Zeynep Erdemir, Ümit Atlamaz, Ömer Demirok <i>Revisiting verbal reflexivity: On the morphosyntax and argument structure of IN derived verbs in Turkish</i>	117
Ruhan Güçlü <i>Does Chatgpt argue through structure or stance? A metadiscourse analysis of English argumentative essays across topics</i>	147
Ali Çağan Kaya, Emre Yağlı <i>Türkçede ünlü formantlarının meta-analitik incelenmesi: Yöntembilimsel farklılıklar ve etkileri</i>	170
Mehmet Akif Kılıç <i>Fonetik–fonolojik süreklilik ve ünlü dörtgeni bağlamında Türkçedeki /a/ ünlüsünün kavramsal konumu</i>	193

Michael Knüppel <i>Altaic studies – What has not been compared so far</i>	200
Hülya Kocagül <i>Yapay zekâ çağında adli dilbilimde yazar analizi</i>	212
Nurbanu Korkmaz <i>Idiom studies in the last decade: A bibliometric mapping of research traces based on web of science</i>	244
Emel Kökpınar Kaya <i>Ya dışındadır dilbilimin ya da içinde</i>	263
Tai Ma <i>Remarks on Turkish interrogative complement clauses and verb subcategorization</i>	273
Engin Evrim Önem <i>Self-paced reading tests with psychopy: A versatile tool for linguistic data collection</i>	289
Alexander V. Savelyev, Yuriy M. Svoyskiy, Anastasiya A. Chernukhina <i>Yenisey Türkçesi runik yazıtlarının belgelenmesi: 2024 yılı epigrafik saha çalışması</i>	305
Mürselin Taşan, Serkan Uygün <i>Morphological processing of Turkish derived words: Does bilingualism affect the processing route?</i>	311
Abdullah Topraksoy <i>Understanding humor in Temel jokes: Insights from conversation analysis, conversational maxims, and speech act theory</i>	335
Serkan Uygün <i>Optional subject-verb agreement in heritage Turkish</i>	350
Kardelen Tuana Yılmaz, Murat Özgen <i>Phasal diagnostics: A critical review</i>	372

IMU Contemporary Studies in Linguistics I

Abstract

Contemporary Studies in Linguistics I advances a clear editorial thesis: contemporary linguistics is unified not by a single method, language family, or theoretical school, but by a common explanatory task—showing how linguistic patterns become observable, interpretable, and defensible through different forms of evidence. Rather than presenting diversity as an end in itself, this edited volume treats historical records, manuscripts and inscriptions, corpora, experiments, discourse data, bibliometric mappings, and AI-mediated texts as complementary evidential domains for linguistic analysis.

Across twenty chapters, the volume moves from historical-comparative reconstruction and lexical history to morphosyntax, phonetics and phonology, semantics, discourse and metadiscourse, bilingual and heritage-language processing, bibliometric research, and emerging interfaces between linguistics and artificial intelligence. What binds these contributions is a shared set of questions: how are linguistic patterns constrained, processed, documented, and transformed; what counts as adequate evidence for linguistic analysis; and how do new tools reshape the empirical foundations of the field?

A major strength of the volume lies in its combination of theoretically informed argumentation with methodologically explicit studies on Turkish and other languages. Its central contribution is to show that methodological pluralism is not miscellany but cumulative argument: different kinds of data illuminate different dimensions of the same object—language as historical record, structured system, cognitive process, and situated use. By bringing historical depth, formal analysis, empirical measurement, and digital innovation into a single frame, the volume offers a coherent account of what contemporary linguistics is, how it proceeds, and how cumulative linguistic knowledge is built across heterogeneous materials and methods.

Keywords: contemporary linguistics; linguistic evidence; linguistic explanation; methodological pluralism; historical and comparative linguistics; experimental linguistics; discourse studies; digital methods; artificial intelligence

Önsöz

Çağdaş Dilbilim Yazıları I: İMÜ Dilbilimi 15. Yıl Armağan Kitabı, İstanbul Medeniyet Üniversitesi Dilbilimi Bölümünün kuruluşunun on beşinci yılı dolayısıyla hazırlanan bir anı kitabıdır. Kitap, bölümün akademik gelişim sürecini ve bu sürece eşlik eden kurumsal hafızayı görünür kılmayı amaçlamakta ve aynı zamanda Türkiye’den ve dünyanın farklı üniversitelerinden dilbilimcilerin katkılarıyla oluşan bir derleme niteliği de taşımaktadır.

İstanbul Medeniyet Üniversitesi 2010 yılında kurulduğunda, Dilbilimi Bölümü de bu sürecin kurucu bir unsuru olmuştur. Üniversitenin kuruluşu, 14 Temmuz 2010’da TBMM’de kabul edilen 6005 sayılı kanunun, 21 Temmuz 2010 tarihli Resmî Gazetede yayımlanmasıyla yürürlüğe girmiştir. Bölümün ilk yapılanması sayılan 1416 sayılı kanun kapsamındaki görevlendirmelerin 7 Ekim 2010 tarihinde yapılması, buna karşın kurucu rektör atamasının 10 Aralık 2010 tarihinde gerçekleşmesi bölümümüzün tarihi açısından oldukça ilginçtir. Bu süreç, bölümümüzün İstanbul Üniversitesi Dilbilimi Bölümüyle birlikte İstanbul’daki ilk bağımsız dilbilimi bölümü olarak kabul edilmesini de mümkün kılmaktadır. Henüz üniversitenin fiziki mekânlarının ve idari yapısının tam olarak oluşmadığı bir dönemde hayata geçen bölümümüz; öğrenci kabulünden önce nitelikli akademik kadronun yetiştirilmesine ve bilimsel altyapının oluşturulmasına odaklanmıştır. Bu vizyoner yaklaşım, ilerleyen yıllarda bölümün ulusal ve uluslararası akademik ağlara entegre olabilmesinin de zeminini hazırlamıştır.

Bölümün akademik çekirdeğinin oluşumunda, merkezi sınavlar ve insan kaynağına yapılan stratejik yatırımlar belirleyici olmuştur. Bu çerçevede, 2010 yılında 1416 sayılı kanun kapsamında MEB bursuyla yurt dışına gönderilen Fırat Başbuğ ve Hülya Kocagül ile 2011 yılında Öğretim Üyesi Yetiştirme Programı (ÖYP) aracılığıyla doktora eğitimini tamamlamak üzere Hacettepe Üniversitesinde görevlendirilen Oktay Çınar, bölümün kurumsal kimliğini şekillendirmede önemli rol oynamışlardır. 2013 yılı ise İMÜ Dilbilimi Bölümü açısından önemli bir dönüm noktasıdır. Bu yıl, Hakan Aydemir bölümde fiilen görev yapan ilk akademik personel olarak atanmıştır. Uzun bir yapılanma sonrasında, 2024 yılında araştırma görevlisi Zeynep Feyza Erdemir’in kadroya katılmasıyla bölümün akademik yapısı güçlenmeye devam etmiştir.

Her ne kadar Dilbilimi Bölümü 2010 yılında kurulmuş olsa da lisans düzeyinde ilk öğrenci kabulü 2022–2023 eğitim-öğretim döneminde gerçekleşmiştir. Bu süre, geriye dönüp bakıldığında, güçlü bir akademik altyapının ve araştırma kültürünün oluşturulması için bir hazırlık dönemi olarak değerlendirilebilir. Bölüm, ilk öğrencilerini kabul ettiğinde, farklı alt alanlarda üretim yapan, ulusal ve uluslararası alan yazını yakından takip eden ve çağdaş yöntemleri kullanan bir akademik çevre haline gelmiştir.

Bu armağan kitabı, söz konusu akademik birikimin yalnızca bölüm içinden değil, daha geniş bir akademik çevreden gelen katkılarla somutlaşmış halidir. Kitapta yer alan bölümler, İstanbul Medeniyet Üniversitesi Dilbilimi Bölümü öğretim elemanlarının çalışmalarının yanı sıra, Türkiye’deki ve yurt dışındaki farklı üniversitelerde görev yapan dilbilimcilerin özgün araştırmalarını da içermektedir. Bu yönüyle kitap, bölümün akademik üretimini çevreleyen geniş bir bilimsel ağın da bir göstergesidir. Kitapta yer alan çalışmalar, dilbilimin farklı alt alanlarına yayılan geniş bir konu çeşitliliğini temsil etmektedir. Tarihsel ve karşılaştırmalı dilbilim, biçimbilim, sesbilim, sözdizim ve anlambilim, söylem ve metin incelemeleri, deneysel dilbilim çalışmaları ile dijital araştırma yöntemleri ve yapay zekâ temelli yaklaşımlar bu çeşitliliğin başlıca örnekleri arasında yer almaktadır. Bu yönüyle kitap, yalnızca bir “anı kitabı” olmanın ötesinde, güncel dilbilim araştırmalarının yöntemsel ve kuramsal çeşitliliğini de bir araya getiren bir derleme niteliği taşımaktadır. Farklı kuramsal yaklaşımların, veri türlerinin ve araştırma geleneklerinin aynı çatı altında buluşması, kitabın temel hedeflerinden birini oluşturmaktadır. Bu armağan kitabı, İstanbul Medeniyet Üniversitesi Dilbilimi Bölümünün kuruluşundan bugüne emeği geçen tüm akademik personele; kitaba katkı sunan Türkiye’den ve dünyadan dilbilimcilere; bölümün gelişimine katkı sağlayan öğrencilere ve bu kurumsal hafızayı paylaşan tüm paydaşlara ithaf edilmiştir. Geçmişten alınan güçle, dilin karmaşık yapısını anlamaya ve anlatmaya yönelik bilimsel yolculuğumuzun uzun yıllar boyunca sürmesi dileğiyle...

Editörler

Oktay ÇINAR, Fırat BAŞBUĞ, Hakan AYDEMİR
31 Aralık 2025. Kadıköy, İstanbul

Preface

Contemporary Studies in Linguistics I: IMU Linguistics 15th Anniversary Commemorative Volume is a book prepared on the occasion of the fifteenth anniversary of the establishment of the Department of Linguistics at Istanbul Medeniyet University. The volume aims to make visible both the academic development process of the department and the institutional memory that has accompanied this process, while at the same time constituting an edited collection brought together through contributions from linguists working at universities in Türkiye and around the world. When Istanbul Medeniyet University was founded in 2010, the Department of Linguistics was established as one of its founding units. The establishment of the university entered into force with the publication of Law No. 6005, adopted by the Grand National Assembly of Türkiye on 14 July 2010, in the Official Gazette on 21 July 2010. From the perspective of the department's history, it is particularly noteworthy that the first appointments under Law No. 1416—considered the initial structuring of the department—were made on 7 October 2010, whereas the appointment of the founding rector took place later, on 10 December 2010. This sequence of events also allows our department to be regarded, alongside the Department of Linguistics at Istanbul University, as one of the first independent linguistics departments in Istanbul. Established at a time when the university's physical infrastructure and administrative structure had not yet been fully formed, the department prioritized the training of a qualified academic staff and the development of a solid scientific foundation prior to admitting students. This visionary approach later laid the groundwork for the department's integration into national and international academic networks. The formation of the department's academic core was shaped by centralized examinations and strategic investments in human resources. Within this framework, Fırat Başbuğ and Hülya Kocagül, who were sent abroad in 2010 under the Ministry of National Education scholarship within the scope of Law No. 1416, as well as Oktay Çınar, who was appointed to Hacettepe University in 2011 to complete his doctoral studies through the Faculty Member Training Program (ÖYP), played significant roles in shaping the department's institutional identity. The year 2013 represents an important milestone for the Department of Linguistics at IMU, as Hakan Aydemir was appointed as the first academic staff member actively serving in the department. Following a long period of institutional development, the department's academic structure continued to be strengthened in 2024 with the appointment of research assistant Zeynep Feyza Erdemir.

Although the Department of Linguistics was established in 2010, its first undergraduate student intake took place in the 2022–2023 academic year. In retrospect, this period can be regarded as a preparatory phase dedicated to building a strong academic infrastructure and research culture. By the time the department admitted its first students, it had already become an academic environment characterized by scholarly production across different subfields, close engagement with national and international literature, and the use of contemporary methodologies. This book represents the tangible outcome of this academic accumulation, enriched not only by contributions from within the department but also by a broader academic community. In addition to the works of faculty members from the Department of Linguistics at Istanbul Medeniyet University, the volume includes original research by linguists affiliated with various universities in Türkiye and abroad. In this respect, the book also stands as an indicator of the extensive scholarly network surrounding the department’s academic production. The contributions included in the volume reflect a broad thematic diversity spanning multiple subfields of linguistics. Historical and comparative linguistics, morphology, phonetics and phonology, syntax and semantics, discourse and text studies, experimental linguistics, as well as digital research methods and artificial intelligence–based approaches are among the primary examples of this diversity. In this sense, the volume goes beyond being merely a commemorative book and also serves as an edited collection that brings together the methodological and theoretical diversity of contemporary linguistic research. One of the fundamental aims of the book is to create a shared space in which different theoretical perspectives, data types, and research traditions converge under a single framework. This book is dedicated to all academic staff who have contributed to the Department of Linguistics at Istanbul Medeniyet University since its establishment; to the linguists from Türkiye and around the world who have contributed to the volume; to the students who have supported the development of the department; and to all stakeholders who share this institutional memory. With the strength drawn from the past, we hope that our scholarly journey toward understanding and explaining the complex structure of language will continue for many years to come.

Editors

Oktay ÇINAR, Fırat BAŞBUĞ, Hakan AYDEMİR
31 December 2025, Kadıköy, Istanbul

INDETERMINACY IN ETYMOLOGY: RECONSTRUCTING THE PIE WORD FOR ‘SHEEP,’ WITH NOTES ALSO ON TOCHARIAN B *ās* ‘GOAT’ AND PIE **h₂éwis* ‘BIRD’

Douglas Q. ADAMS

Prof., University of Idaho, dqadams@uidaho.edu

Adams, D. Q. (2025). Indeterminacy in etymology: Reconstructing the PIE word for ‘sheep,’ with notes also on Tocharian B *ās* ‘goat’ and PIE *h₂éwis* ‘bird’. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 05–12). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X001>

ABSTRACT

The Indo-European word for ‘sheep’ is well-represented in most Indo-European groups. However, its exact shape in the proto-language is not as well-established as usually thought. The initial laryngeal has been reconstructed as **h₁-*, **h₂-*, or **h₃-* by various investigators and the noun’s accentual pattern often left undefined. The argument here is that we cannot reconstruct with certainty. It seems perhaps most likely that it was an acrostatic **h₂ówi-/h₂éwi-* but Tocharian might also reflect either **h₂ōwi-* and **h₃ōwi-* and, if the latter, **h₃ówi-/h₃éwi-* would certainly be possible for other Indo-European groups. Knowing if PIE **h₃-* became *χ-* in Lycian would help delimit the number of indeterminacies but not solve the problem decisively. We will be left with multiple possibilities.

Keywords: Indo-European etymology, PIE laryngeals (**h₂*, **h₃*), ‘Sheep’ etymon (**h₂ówi-* ~ **h₃ówi-*), Acrostatic nouns, Lengthened-grade Athematic formations (Hitt. *ās-*, Skr. *āsāt*), Brugmann’s Law, Tocharian B (*ās-*), Anatolian evidence, Comparative method.

ÖZ

Hint-Avrupa kökenli ‘koyun’ anlamındaki sözcük, Hint-Avrupa dil gruplarının çoğunda iyi tanıklanmıştır. Ne var ki, bu sözcüğün Ana Hint-Avrupa dilindeki (PIE) tam biçimi, genelde düşünüldüğü kadar iyi belirlenememiştir. Önsesteki gırtlaksız, farklı araştırmacılar tarafından **h₁-*, **h₂-* ya da **h₃-* olarak kurgulanmış; sözcüğün vurgu örüntüsü ise çoğu kez tanımlanmadan bırakılmıştır. Bu çalışmadaki temel sav, bunun kesin bir biçimde kurgulanamayacağıdır. En olası görünen çözüm, bunun kök vurgulu (acrostatic) bir **h₂ówi-/h₂éwi-* biçiminde olduğudur; ancak Toharca, **h₂ōwi-* ya da **h₃ōwi-*

biçimlerinden birini de yansıtır olabilir. İkinci durum söz konusuysa, diğer Hint-Avrupa grupları için **h₃ówi-/h₃éwi-* biçimi de kesinlikle olanaklıdır. PIE **h₃-*’nin Likçede *χ-*’ye dönüşüp dönüşmediğinin bilinmesi, belirsizliklerin sayısını sınırlamaya yardımcı olabilir; ancak sorunu kesin biçimde çözmez. Birden fazla olasılıkla karşı karşıya kalırız.

Anahtar Sözcükler: Hint-Avrupa etimolojisi, Ana Hint-Avrupa dili (PIE) gırtlaksızları (**h₂, *h₃*), ‘Koyun’ etimonu (**h₂ówi- ~ *h₃ówi-*), Kök vurgulu (acrostatic) adlar, Uzun dereceli atematik yapılar (Hit. *ās-*, Skr. *āsāt*), Brugmann Yasası, Toharca B (*ās-*), Anadolu dilleri kanıtları, Karşılaştırmalı yöntem

Though one of the most iconic and well-attested lexical items in the reconstructed PIE vocabulary, the exact shape of the word for ‘sheep/ewe’ (in Brugmann terms **owis*) is not as firmly established as one might think. Of the twelve well-attested Indo-European stocks Albanian and Iranian show no traces of the PIE word. And, leaving aside Tocharian for the moment, there are nine other well-attested stocks of which eight have been provided with dictionaries by members of the “Leiden School”: Kloekhorst for Anatolian (2008:337-338), De Vaan for Italic (2008:437-438), Derksen for Slavic (2008:384), Matasović for Celtic (2009:301), Beekes for Greek (2010:1061-1062), Martirosyan for Armenian (2010:419), Kroonen for Germanic (2013:45), Derksen for Baltic (2015:74). To this set we can add Puhvel for Hittite (1991:279-280) and Melchert (2004) for Lycian. Of these four give no Tocharian data at all (Puhvel, Melchert, Derksen [2008], and Beekes). The others mention the Tocharian word for ‘ewe,’ giving more commonly the nominative plural *awi* (De Vaan, Matasović, Derksen [2015]); Kroonen gives the nominative singular *āw*,¹ Kloekhorst mentions *ā*-forms in Tocharian, but he, like everyone but Matasović, proceeds to ignore the Tocharian data. All but Matasović reconstruct **h₃owī-* (or most commonly **h₃ewī-*) tout simple.² Kloekhorst dismisses the Tocharian data as “outside his competence,” Kroonen admits that the Tocharian shows “unexpected vocalism.” On the other hand, Beekes claims positive support for **h₃ewī-* because of a lack Brugmann’s Law in this word in Sanskrit.³

Of the Leiden School Matasović alone (2009), following Pinault (1997:191ff.), reconstructs an acrostatic paradigm **h₂ówis/h₂éwyos* (which would be phonetically [**h₂áwyos*] already in PIE), with generalization of the *-o-* in Armenian, Greek, Italic, and Celtic; the *-a-* in Germanic, Baltic, and Slavic being ambiguous (**o* and **a* having fallen together). Melchert (2004:81) also reconstructs PIE **h₂owis*, specifically rejecting **h₃ewis*,⁴ but makes no comment about its morphological structure.

¹ Aside from the nominative singular and plural the only other attested member of the Tocharian B paradigm is the obviously analogical genitive singular *awantse*.

² Well, not quite: Puhvel gives **h₃ewī-* (in the notation used here) or, quite surprisingly, **h₁owī-*.

³ Going outside the universe of dictionaries, one might note that Olsen’s discussion (2018:70) of the word gives either **h₂owis* or **h₃owis* as possible proto-forms. As we will see, I think she’s on the right track.

⁴ Just as he did ten years earlier in his 1993 Lycian dictionary.

But, even side from Tocharian, there are issues of various sorts with the almost universal preference for **h₃ewi-*, issues with both the Indic and Anatolian. The complication with Indic involves the intersection of this word with Brugmann’s Law. For those researchers, and there are many, who reject the law there actually is no problem: any initial laryngeal followed by any short mid vowel would yield Sanskrit *a-*, such as we find in *ávis*. However, for those who do believe, and there are many, including Beekes, in some formulation of Brugmann’s Law, whereby a PIE **-o-* in an open syllable followed by a resonant appears in Sanskrit as *-ā*⁵ except where that **-o-* is “non-apophonic,” i.e., when it does not alternate with **-e-*. An **-e-* being “colored” by laryngeal **h₃-* to **-o-* is assumed to make it non-apophonic because **-e-* and **-o-* have fallen together in this environment. The defenders of an initial **h₃-* in this word for ‘sheep’ are unanimous in saying we do not have lengthening in this word in Sanskrit because it was the *-o-* of **h₃owi-* was non-apophonic in PIE.

Though it’s not universally accepted, I myself think that Brugmann’s Law seems a safe bet for Indo-Iranian, though the original distribution has been blurred/disturbed by various analogical changes (e.g., the extension of the lengthened grade to derived causatives [like *sādayati* ‘seats’] or the extension, outside the realm of Brugmann’s Law, of the lengthened grade of the nominative singular of root nouns to the accusative singular [**pōds/podm̐* > **pōds/pōdm̐*]). Nevertheless, the notion of “non-apophonic *o*” seems very shaky to me. Laryngeal-coloring seems clearly to be of PIE date so we should have had **h₃owi-* already in this putative word for sheep, rather than **h₃ewi-*, with no suggestion that its **-o-* was phonetically distinct from any other **-o-* in that variety of late PIE that gave rise to Indo-Iranian. Thus, there would seem to be no reason why Brugmann’s Law should not apply. So why not **ávi-*? Labeling the **-o-* of **h₃owi-* as “non-apophonic” hardly satisfies any *phonetic* criteria.⁶ Moreover, there’s no phonological reason why the initial vowel of Sanskrit *ávis* could not reflect the vowel of the weak cases of an acrostatic **h₂ówis/h₂áwyos*.⁷ And thus, it turns out that, whether there was a Brugmann’s Law or not, with or without the “non-apophonic clause,” Sanskrit offers no deciding vote on whether the word for sheep had an initial **h₂-* or **h₃-* and we are no further along than we were with the first seven Indo-European languages we first

⁵ That is Brugmann’s Law in the Kleinbans-Pedersen formulation of 1900.

⁶ It would be possible, I suppose, to construct an argument that **h₃-* colored and **-e-* to **-o-*, say, distinct from **-o-* (and thus still an allophone of **-e-*), and thus not subject to Brugmann’s Law and that this allophone of **-e-* only fell together (independently?) with **-o-* or **-a-* later in all branches of Indo-European. Possible I suppose, but it seems quite unlikely (and convoluted) and I don’t recall seeing that argument ever having been made. Certainly, none of the lexicologists listed above asserting **h₃ewi-* as the ancestral form of ‘sheep’ have made that argument.

⁷ That the acrostatic paradigm seen in Sanskrit is inherited from Proto-Indo-European is supported by the same acrostatic paradigm seen in Greek, *óis/óios*. Note, too, the parallel generalization of the vocalism of the weak cases in Sanskrit *vis/vés* ‘bird.’ Compare the generalization of the vocalism of the strong cases in Latin *avis*. Both descend from a proterodynamic paradigm **h₂ówis/h₂wéis*. For ‘bird’ see further Blažek and Adams (2023 [with previous literature]).

considered. As far as the eight “Brugmannian” stocks are concerned the initial laryngeal could be any laryngeal.⁸

Turning to the situation in Anatolian, we see it is not as straightforward as one might wish, but for different reasons. In most Anatolian languages (e.g., Hittite and Luwian) PIE **-a-* and **-o-* fall together as *-a-* (or **-ā-* when stressed in an open syllable). But Lycian is different; PIE **-a-* is Lycian *-a-*, but PIE **-o-* is Lycian *-e-*. So, obviously, the merger had not taken place in Proto-Anatolian. Moreover, while Hittite and Luwian reflect both initial PIE **h₂-* and **h₃-* as *h-*, Lycian may be different. PIE **h₂-* is definitely reflected in Lycian as *χ*⁹ but the Lycian reflex of PIE **h₃-* is disputed: some see it as zero (Melchert, 1994:305), others as *χ* (like **h₂-*) (Kloekhorst, 2008:219-220). The primary datum is provided by *χawa-* ‘sheep.’ All are agreed that this word owes its declensional class *-a-* to the influence of *wawa-* ‘cow’ (< **g^wow-ah₂-*). The **χewa-* thus produced becomes *χawa-* by regular *a*-umlaut (i.e., **-e-...-a- > -a-...-a-*). As noted, the *χ*- is taken by some as evidence that **h₃-* was preserved in Lycian and, thus that **h₂-* and **h₃-* had not merged in Proto-Anatolian. However, examples of either putative development (*χ* or zero) are very few (one each) and deriving Lycian *χawa-*, and the cognates in Hittite and Luwian, from **h₂owis* is altogether possible.¹⁰

χawa- is the single lexical item that provides what positive evidence there is that **h₃-* appears as *χ*- in Lycian. The negative evidence, i.e., evidence that suggests initial **h₃-* was not preserved in Lycian, is Lycian *epirijeti*. This word has long been equated with Hittite *happirije/a-* ‘trade, sell, deliver,’ both taken as reflecting a putative PIE **h₃operye/o-* (see most authoritatively Melchert, 1993:17 and 2004:15 with previous literature). However, Rasmussen (1992:56-59) has argued that the meaning ‘sells’ given to the Lycian word is not assured and given mostly because of its formal resemblance with Hittite *happirije/a-*. For Rasmussen it might possibly be ‘sells’ but might be several other things instead. Following Rasmussen’s lead but even more strongly rejecting the equation with Hittite *happerije/a-* is Kloekhorst’s discussion, 2008:296. If it should mean ‘sells’ then an equation with *happirije/a-* would be certain, but, if not, then not. Moreover, even if it should be ‘sells,’ the scarcity of evidence allows it to be the case that initial **h₂-* and **h₃-* remained differentiated in Proto-Anatolian and pre-Lycian and, in a development reminiscent of Saussure’s Law where a laryngeal disappears in the presence of an *o*-grade (e.g., *pórnē* ‘prostitute, whore,’ a derivative of **perh₂-* ‘sell’), Proto-Anatolian **h₃-* disappeared before the retained **-o-* of pre-Lycian (before it became the *-e-* seen in attested Lycian) but remained elsewhere (or at least before *-a-*). Only later did initial **h₂-* and **h₃-* fall together as *χ*- in Lycian. Alterna-

⁸ As noted above, Puhvel (1991:279-280) allows either **h₃owis* (my modernization of his transcription) or **h₁owis* (!).

⁹ Cf. *Xaha-* ‘altar’ and Hittite *hassa-* ‘fireplace, hearth’ (compare Latin *āra* ‘altar’) and *χba(i)-* ‘irrigate’ and Hittite *hapā(i)-* ‘pour water over’ (compare further Tocharian *āp* ‘water’).

¹⁰ This is, for instance, Melchert’s position (2004:81) who is firm in reconstructing **h₂owi-* on the basis of the Anatolian evidence.

tively, the original **h₃-* here was “assimilated” to **h₂-* when the following vowel became *-a-* and later **h₃-* disappears.

Ultimately, though **h₁-* is excluded (contra Puhvel), Anatolian, like the other Indo-European stocks examined so far, provides no compelling evidence for either **h₃-* or **h₂-* as the initial laryngeal of the word for ‘sheep.’¹¹

But Tocharian is possibly a game-changer. As in Lycian, PIE **-o-* remains distinct from PIE **-a-* in most cases.¹² The Tocharian B word for ‘ewe’ is *ā(u)w* (nom. sg.), *awi* (nom. pl.) < **āwī*. Obviously, it belongs with the other words for ‘sheep’ and/or ‘ewe’ listed above. But how? If there had been a (PIE) **-o-* in the initial syllable, it would have appeared in Tocharian B as *-e-*. But if we start with an acrostatic paradigm **h₂ówis/h₂áwyos*, the actual Tocharian result would be realized if, like the one possibility outlined for Sanskrit, the one where the initial laryngeal had been **h₂-* and the paradigm acrostatic, the vocalism of the weak cases had been extended to the strong ones. If this is so, we have the interesting situation where the roots of two PIE words are phonologically identical (in their early PIE guise): **h₂ew-* ‘bird’ and **h₂ew-* ‘sheep.’ They are united too by their stem class, **h₂ew-i-* for both, but they are differentiated by their ablaut/accent pattern. **h₂ewi-* ‘bird’ was apparently a proterokinetic noun (strong cases **h₂éwi-*, weak cases, illustrated by the genitive singular, **h₂wéis*) as shown by Sanskrit *ví-* ‘bird,’¹³ while **h₂ewi-* ‘sheep’ was acrostatic **h₂ówi-/h₂éwi-*.¹⁴ Here we would have a morphophonemic “minimal pair.”

However, there is another possible explanation the *ā-* of Tocharian *āw*. That *ā-* could be a reflex of PIE **-ō-*. That is, we might have a lengthened grade formation **h_{2/3}ōwis*. Other Tocharian words that have been seen as having lengthened grades are *āntse* ‘shoulder’ (< PIE **h_xōmso-*, cf. Latin *humerus*), *sāle* ‘ground; basis’ (cf. Latin *solum*

¹¹ Data from other Anatolian languages, e.g., Palaic or Lydian, would probably not be helpful here. In Palaic both initial **h₂-* and **h₃-* are at least graphically identical (for the latter, compare Palaic *hapari(ya)-* ‘hand over’). In Lydian there is no evidence for initial **h₃-* but **h₂-* has definitely disappeared and that means **h₃-* is almost surely gone too (for both languages see Melchert, 1994).

¹² Notably they fall together when original PIE **-o-*, in Proto-Tocharian **-e-* (accidentally aping the outcome in Lycian), is followed in the next syllable by Proto-Tocharian **-ā-*, i.e., **-e-...-ā- > -ā-...-ā-*, again showing the same outcome as Lycian. (This process is normally referred to as *ā*-umlaut by Tocharianists.) In Tocharian A *ā*-umlaut occurs only when the original **-o-* was unstressed while in Tocharian B it affected both stressed and unstressed **-e-*. The putative Proto-Tocharian **-e-* of ‘sheep’ was not followed by an **-ā-*, so it would not have been subject to *ā*-umlaut.

¹³ The Tocharian cognate also shows the reflex of the weak cases. It is *wiñce** ‘nestling,’ known only in the derived adjective *wiñcaññe* ‘pertaining to a nestling’ (Blažek and Adams, 2022 [2023]:341-343). This must be a putative PIE **h₂wi-nt-en-* (adjective as if **h₂wi-nt-en-yo-*). The ‘small animal suffix’ **-nt-* seen in Tocharian is the post-vocalic equivalent of the post-consonantal **-nt-* which gives Slavic *-et-* in names for juvenile animals, e.g., OCS *telēt-* ‘calf,’ etc. It should be re-iterated that this diminutive/juvenile formation for animals is a rather significant isogloss common to Tocharian and Slavic (Blažek and Adams, 2022[2023]:343).

¹⁴ Somewhat like Greek *ō(i)ā* ‘sheepskin,’ in showing a long *-ō-* after an initial **h_{2/3}-* in a derivative, is **h₂ōw(i)yom* ‘egg’ from **h₂awi-* (< **h₂ewi-*) ‘bird.’ In Germanic there is also the pair **mari-* ‘sea’ (> Modern English *mere*) and **mōra-* ‘moor.’ My intent here is not to explore the presence of lengthened grade formations in Indo-European, merely show that what I’m postulating as a possibility in Tocharian is paralleled elsewhere in Tocharian (see next paragraph) and elsewhere in the Indo-European world.

‘ground’), *āsta* ‘bones’ (cf. Greek *ostéon*) *āšce* ‘head’¹⁵ (cf. Greek *ózdos* ‘branch’ or *ostéon* ‘bone’) (cf. Darms, 1978:325, with previous literature). We find *i*-stems with **-ō-* in Germanic for instance (e.g., in Anglian varieties of Old English there is *dæg* ‘day’ (< Proto-Germanic **dōgi-* beside **daga-* ‘day’)) and a few others of similar formation.¹⁶ As already pointed out, this Proto-Indo-European word for ‘sheep’ shows a lengthened grade derivative in Greek *ō(i)ā* ‘sheepskin.’

Of course, in this latter case, both **h₂ōwi-* and **h₃ōwi-* are possible antecedents of the Tocharian form. What makes the possibility of a lengthened grade in the history of *āw* attractive is that we could then derive Tocharian B *ās-* ‘she-goat’¹⁷ from a form with similar lengthened grade **h₁ōsi-* (cf. the similarly derived *awantaññe* ‘pertaining to ewes’ and *asantaññe* ‘pertaining to she-goats’¹⁸), a derivative of **h₁es-* ‘be, sit,’ and thus originally ‘goods, chattels.’¹⁹ It seems obvious that these two words have influenced one another morphologically in some respects, so an influence on the root vowel as well is certainly not out of the question. Of course, the influence might as well have been from ‘she-goat’ to ‘ewe’ as the other way around. Or, indeed, the “accidental” resemblance of the Tocharian reflexes of **h₂awi-* and **h₁ōsi-*, where the initial had become in both cases Tocharian **ā-* may have set the stage for further mutual influence.

So, we are in a sense back where we started. The question is, “was the initial consonant of the PIE word for ‘sheep, ewe’ an **h₂-* or an **h₃-*?” The answer would seem to be, “we can’t be sure.” The Sanskrit reflex *ávis* was taken as definitive evidence of **h₃-* by Beekes (and others) but, as we have seen above, it is not. In the past Pinault (1997) and I (Adams, 1999) have seen Tocharian B *āw* as definitive evidence for **h₂-* but again, as we have seen, it is not. (Though these latter accounts have the virtue of taking the Tocharian data into account.) Lycian²⁰ might be decisive, if it can be

¹⁵ Historically an *i*-stem, nom. sg. *āšce*, acc. sg. *āšc*, nom. pl. *āscī*, acc. pl. *āstām*, whether **ōzd-i-* or **ōst-i-*.

¹⁶ Very similar to **dagaz* ‘day’ and **dōgiz* ‘day’ is **dalaz* ‘valley’ (> NHG *Tal*, Modern English *dale*) and **dōliz* ‘valley,’ the latter seen in ON *dæl* ‘small valley’ and *dæll* (< **dōljaz*) ‘valley dweller.’

¹⁷ Also called ‘doe’ or ‘nanny goat’ (as a ‘he-goat’ is also a ‘buck’ or ‘billy goat’), though not in my variety of English.

¹⁸ Both from secondary plurals, i.e., *āwántā** and *āsántā** (not attested but implied by *asantaññe*)? *-nta* and *-nma* are the two most productive plural formations in Tocharian B.

¹⁹ Hittite *āssu* ‘goods, chattels’ may be as if from **h₁osus* with geminate *-ss-* (thus as if **h₁ossu*) generalized from the underlying verb **h₁es-* when the verb occurred before consonant-initial endings (a la Melchert, 1994:151-152 [though he doesn’t give this precise explanation for *āssu*]). Kloekhorst (2008:223-224) makes the same semantic connection (and a connection with Greek *eús* ‘good, brave, strong [in war]’ and Sanskrit *su-*) but starts from an extremely unlikely reduplicated **h₁o-h₁s-* with an unexplained **-o-* in the reduplicated syllable. Puhvel’s derivation (1984:199ff.) from **ans-u* from **ans-* (Hittite *ass-/assiya-*) ‘be favored’ is phonetically plausible but semantically strained. However, that an adjective, ‘dear,’ derived from *ass-/assiya-* and a noun ‘goods’ derived from **h₁es-* should have secondarily become associated/conflated is quite possible. Further it should be noted that a root-noun **h₁os-* or **h₁ōs-*, connected to **h₁ēs-* ‘sit,’ in my view a long-vowel iterative-intensive of **h₁es-* ‘be,’ by both Puhvel (1984:296) and Kloekhorst (2008:220), appears in Hittite *as-āwar/as-aun-* ‘sheep-pen’ (< **h₁ōs-ōwar* or, Kloekhorst’s preferred proto-form, **h₁s-ōwar*).

²⁰ Or other Anatolian languages where ‘sheep’ is not yet attested.

determined that **h₃-* does not show up as *χ-*, but at present the evidence is altogether too exiguous.

The lesson learned, or perhaps relearned, is that our reconstructions are utterly dependent on our data and the access of new data will change, or potentially change, our reconstructions and, no, one cannot simply declare Tocharian to be “outside my competence” and ignore it—no more than one can ignore Greek or Germanic, etc. And, perhaps a bit unsettling, there is more than one possible, even plausible reconstruction: ultimately our reconstructions, not just this one, are to varying degrees and in varying ways, indeterminate.²¹

References

- Adams, D. Q. (2013). *A dictionary of Tocharian B* (2nd ed., revised and greatly enlarged). (Leiden Studies in Indo-European, 10).
- Blažek, V., & Adams, D. Q. (2022 [2023]). Indo-European “bird.” *Journal of Indo-European Studies*, 50, 335–360.
- Darms, G. (1978). *Schwäher und Schwager, Han und Huhn: Die Vrddhi-Ableitung im Germanischen*. Kitzinger.
- Derksen, R. (2008). *Etymological dictionary of the Slavic inherited lexicon* (Leiden Indo-European Dictionary Series, 4). Brill.
- Derksen, R. (2015). *Etymological dictionary of the Baltic inherited lexicon* (Leiden Indo-European Dictionary Series, 13). Brill.
- Hansen, B. S. S. (2017). Layers of root nouns in Germanic. In B. S. S. Hansen, B. N. Whitehead, T. Olander, & B. A. Olsen (Eds.), *Etymology and the Indo-European lexicon: Proceedings of the 14th Fachtagung der Indogermanischen Gesellschaft, 17–22 September 2012, Copenhagen* (pp. 169–182). Reichert.
- Kloekhorst, A. (2008). *Etymological lexicon of the Hittite inherited lexicon* (Leiden Indo-European Dictionary Series, 5). Brill.
- Kroonen, G. (2013). *Etymological dictionary of Proto-Germanic* (Leiden Indo-European Dictionary Series, 11). Brill.
- Matasović, R. (2009). *Etymological dictionary of Proto-Celtic* (Leiden Indo-European Etymological Dictionary Series, 9). Brill.
- Melchert, H. C. (1993). *Lycian lexicon* (2nd rev. ed.). (Studia Anatolica, 1). Self-published.
- Melchert, H. C. (2004). *A dictionary of the Lycian language*. Beech Stave Press.
- Nussbaum, A. J. (1998). Severe problems. In J. Jasanoff, H. C. Melchert, & L. Oliver (Eds.), *Mir Curad: Studies in honor of Calvert Watkins* (pp. 521–538). Innsbrucker Beiträge zur Sprachwissenschaft.

²¹ My current best guess is that the very real mutual influence of *ās* ‘she-goat’ and *āw* ‘ewe’ results from the **h₂a-* of ‘ewe’ and the **h₁ō-* of ‘she-goat’ falling together as **ā-*, accidentally as it were, in Proto-Tocharian. Naturally, the semantic proximity of the two words would certainly reinforce that mutual influence.

- Olsen, B. A. (2017). Armenian. In M. Kapović (Ed.), *The Indo-European languages* (2nd ed., pp. 421–451). Routledge.
- Olsen, B. A. (2018). Notes on the Indo-European vocabulary of sheep, wool and textile production. In R. Iversen & G. Kroonen (Eds.), *Digging for words: Archaeolinguistic case studies from the XV Nordic TAG Conference held at the University of Copenhagen, 16–18 April 2015* (pp. 69–77). (BAR International Series 2888). BAR Publishing.
- Pinault, G.-J. (1997). Terminologie du petit bétail en tokharien. *Studia Etymologica Cracoviensia*, 2, 175–219.
- Puhvel, J. (1984). *Hittite etymological dictionary* (Vols. 1–2). (Trends in Linguistics: Documentation, 1). Mouton.
- Puhvel, J. (1991). *Hittite etymological dictionary* (Vol. 3). (Trends in Linguistics: Documentation, 5). Mouton de Gruyter.
- Rasmussen, J. E. (1992). Initial *h₃* in Anatolian: A vote for chaos. *Copenhagen Working Papers in Linguistics*, 2, 53–61.

MORPHOPHONOLOGICAL, MORPHOSYNTACTIC, AND ICONIC CONSTRAINTS GOVERNING IRREVERSIBLE NOMINAL AND VERBAL BICOORDINATIVES IN TURKISH

Hakan AYDEMİR

Assoc. Prof. Istanbul Medeniyet University, Department of Linguistics, aydemirhaakan@gmail.com,
ORCID: 0000-0002-2368-7103

Aydemir, H. (2025). Morphophonological, morphosyntactic, and iconic constraints governing irreversible nominal and verbal bicoordinatives in Turkish. In O. Çınar, F. Başbuğ, & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 13–64). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X002>

ABSTRACT

This study investigates the structural principles governing irreversible nominal and verbal binary coordinatives in Turkish, referred to here as *bicoordinatives* (Turkish: *ikileme*). While such binary coordinative formations have long been studied in descriptive work, the rules governing such binary structures have not yet been fully established. Drawing on a relatively large dataset from both Modern Turkish and Old Turkic, the present study argues that these formations are regulated by a systematic interaction of morphophonological, morphosyntactic, and iconic constraints, formulated here as *Ordering Principles* (OPri). The analysis identifies seven hierarchically interacting rules that determine the fixed order of *asymmetric bicoordinatives*. These rules range from segmental and syllabic properties to higher-level iconic relations reflecting event structure and semantic sequencing. It is shown that, when the balance of the syllable pattern is disrupted, asymmetries in syllable structure and syllable count override segmental preferences. The study further demonstrates that these ordering principles are not restricted to Modern Turkish but are already attested in Old Turkic too, indicating long-term structural stability within the language. Beyond bicoordinatives, the paper situates these findings within a broader *block-recursive model of coordination* (BiCo–TriCo–MulCo), showing how binary coordinative units serve as the building blocks for larger coordinative structures. The proposed theoretical framework contributes to Turkish linguistics, coordination theory, and typological research by providing a unified and empirically grounded account of *irreversible coordination*.

Keywords: Turkish linguistics, Bicoordinatives, Irreversible binomials, Coordination, Ordering principles, Sonority hierarchy, Morphophonology, Iconicity, Typology

ÖZ

Çalışma, burada *bicoordinative* olarak adlandırılan Türkçedeki kalıplaşmış ad ve eylem ikilemelerini düzenleyen yapısal ilkeleri incelemektedir. Bu tür eşbağımlı ikili yapılar betimleyici çalışmalarda uzun süreden beri incelense de, bu ikili yapıları düzenleyen kurallar henüz tam olarak belirlenememiştir. Çağdaş Türkçe ve Eski Türkçeden görece geniş bir veri kümesine dayanan bu çalışma, söz konusu yapıların aslında kurallı olduğunu, biçimsesbilimsel, biçimdizimsel ve görüntüsel kısıtlamaların düzenli etkileşimiyle ortaya çıktığını savunmaktadır. Bu kısıtlamalar, çalışmada *Dizilim İlkeleri* (OPri) olarak ele alınmaktadır. İncelemede, *bakımsız ikilemelerin* kalıplaşmış dizilimini belirleyen ve öncelikle olarak etkileşim halinde olan yedi kural tanımlanmaktadır. Bu kurallar, sesbirimsel ve seslem yapısına dayalı özelliklerden, olay yapısı ve anlamsal sıralamayı yansıtan daha üst düzey görüntüsel ilişkilere kadar uzanmaktadır. Seslem örüntüsündeki dengenin bozulduğu durumlarda, seslem yapısı ve seslem sayısındaki farklılıkların sesbirimsel tercihleri geçersiz kıldığı gösterilmektedir. Çalışma, bu ilkelerin yalnızca Çağdaş Türkçeye özgü olmadığını, Eski Türkçede de aynı biçimde işlediğini ortaya koymaktadır. Çalışma ayrıca, ikilemeleri daha geniş bir *öbek yinelemeli eşbağımlama modeli* (BiCo–TriCo–MulCo) içinde ele alarak, ikilemelerin daha büyük eşbağımlı dizilerin temel yapıtaşları olduğunu göstermektedir. Önerilen kuramsal çerçeve, *kalıplaşmış eşbağımlamanın* bütüncül ve veriye dayalı bir açıklamasını sunarak Türk dilbilimine, eşbağımlama kuramına ve tipolojik araştırmalara katkıda bulunmaktadır.

Anahtar Sözcükler: Türk dilbilimi, İkileme, Eşbağımlama, Dizilim ilkeleri, Titreşimsel önceleme, Biçimsesbilim, Görüntüsellik, Tipoloji

1. Introduction

The structure of irreversible binomials and biverbals — i.e. irreversible coordinations, referred to in this study as *bicoordinatives* (e.g. *alt üst*, *bıkmak usanmak*, etc.) — has received increasing attention in both theoretical and Turkish linguistics.¹ However, this attention has so far focused primarily on nominal formations, while their verbal counterparts have remained largely unexplored, particularly in Turkish linguistics. Despite this growing interest, the structural constraints governing these formations in Turkish remain poorly understood.

These constructions, characterized by the coordination of two lexemes, play a significant role in the morphophonological, morphosyntactic, and semantic structuring of the Turkish lexicon. However, a fundamental and unresolved question persists: according to what principles do these formations emerge and undergo lexicalization? Despite their frequency and apparent regularity, no comprehensive or satisfying account has yet been provided. This study aims to address that gap by revisiting the theoretical framework that I

¹ For a detailed, annotated, and critical survey of relevant research on Turkish bicoordinatives from 1899 to 2007, see Aydemir (2007), an unpublished MA thesis that is publicly and online available via the Turkish Council of Higher Education (YÖK) National Thesis Center. For more recent synchronic discussions, see Gök & Canalis (2023) and Özkan (2019: 2). The present article therefore does not repeat the historical overview.

first proposed in 2007,² while also offering a typologically informed and formally explicit analysis of bicoordinative structures in Turkish, both synchronic and diachronic.

Recent research in phonological typology and morphophonology has shown that ordering asymmetries in linguistic constructions — including coordination and compounding — are systematically shaped by several universal forces. These include the sonority hierarchy (Clements 1990; Parker 2002, 2008), principles of articulatory ease and perceptual salience (Lindblom 1990), and iconicity (Haiman 1985; Croft 2003). Some of the earliest observations in this domain can already be traced back to Jespersen (1904).

Within this framework, Turkish bicoordinatives provide a particularly revealing testing ground, since they exhibit a complex interplay of phonological, morphosyntactic, and iconic factors that determine their irreversible order. While the descriptive literature abounds with examples of Turkish binomials, a coherent model that unifies these different dimensions under a single explanatory framework has not yet been proposed. The present study seeks to fill this gap by systematically identifying the *morphophonological*, *morphosyntactic*, and *iconic constraints* as *Ordering Principles* (OPri) that govern the order of nominal and verbal bicoordinatives in Turkish and by clarifying their functional and diachronic motivation.

2. Conceptual and Terminological Framework

2.1 Definition and scope of the term *bicoordinative*

The analytical distinctions required by this study make it necessary to introduce a partly refined terminology, since the structural, morphophonological, morphosyntactic, and iconic properties of Turkish bicoordinatives can only be adequately described, evaluated, and precisely delimited through an appropriately differentiated set of terms.

In this study, therefore, I propose the terms *bicoordination* (as a process) and *bicoordinative* (as its product) as a typologically grounded and morphophonologically — morphosyntactically motivated category that captures a specific and structurally marked class of *irreversible* binomial and biverbal coordinations in Turkish. Thus, these notions refer to binary coordination structures that exhibit fixed internal ordering.³ The label is a translation and adaptation of the term *Bi-Koordinativ* (cf. “*co-ordinative*”, Bloomfield 1933), which I first introduced in my German-language master’s

² The basic theoretical framework of this study was first developed in my German-language master’s thesis (Aydemir 2007). While the original observations have remained robust, the present version expands and refines those initial insights, situating them more explicitly within a broader typological and theoretical framework.

³ Although *irreversibility* and *lexicalization* are not synonymous concepts in general linguistic theory, they strongly converge in Turkish. Virtually all irreversible bicoordinatives in Turkish are also lexicalized and idiomatized, behaving as fixed lexical units. Nevertheless, in this study the two notions are kept conceptually distinct: *irreversibility* refers to a structurally mandated ordering constraint, whereas *lexicalization* denotes the diachronic and semantic consolidation of the construction as a stable lexical item.

thesis (Aydemir 2007), alongside the related categories *Tri-Koordinativ* (*tricoordinative*; see §2.4) and *Multi-Koordinativ* (*multicoordinative*; see §2.5), the latter denoting formations consisting of more than two successive irreversible bicoordinatives.

Until that earlier work, these formations had not been systematically classified, and existing terminology in the literature — terms such as *hendiadys*, *word pairs*, *reduplication*, or *binomial* — tended to capture only certain subtypes within a much broader systemic phenomenon, both structurally and functionally, observable across Turkic.⁴

In contrast, the notion *bicoordinative* allows for a more comprehensive typological, morphophonological, morphosyntactic, and iconic account (i.e. *iconicity*, see §5.3) that also enables cross-linguistic comparison. It captures (a) the internal morphophonological, morphosyntactic, and iconic alignment patterns, (b) the contrast between *symmetric* and *asymmetric* formations (see §2.2), and (c) the specific coordination properties shared by binomial and biverbal constructions.

Among these, *asymmetric bicoordinatives* (AsBi) are characterized by a fixed precedence relation between their components (e.g. *alt üst* ‘upside-down’, *dere tepe* ‘over hill and dale’), which makes the order [A] [B] obligatory and the reverse [B] [A] impossible. This contrasts with *symmetric bicoordinatives* (SyBi), which exhibit formal mirroring or full repetition and therefore do not encode any internal precedence hierarchy, even though they are likewise fixed, inseparable, and typically lexicalized (e.g. *bile bile* ‘knowingly’ ← [know-CONV] [know-CONV], *güle güle* ‘good bye’ ← [smile-CONV] [smile-CONV]).

This typological distinction between symmetric and asymmetric formations forms the structural basis of my analysis and motivates the formal ordering constraints captured in Rules 1–7. It also links these constraints to broader functional, synchronic, and diachronic tendencies observed in Turkic.

A further distinction essential to Turkish is that between *syndetic* and *asyndetic coordinations* (see §2.3). This distinction is essential for understanding the structural variation within Turkish bicoordinatives. *Asyndetic bicoordinatives* — e.g. *ev bark*, *dal budak*, *alt üst*, etc. — lack any overt conjunction and show tighter morphophonological integration. *Syndetic bicoordinatives* employ overt coordinators such as *ve*, *ile*, *ya*, *ya da*, etc., in Modern Turkish (see §2.3 for details). In Old Turkic, syndetic coordination is attested both (a) through the genuinely coordinative suffixal pattern [...+II] [...+II],

⁴ In the German (and partly also English) Turkological literature — and to a large extent also in the Turkish linguistic literature written in English — the terminology used for this phenomenon is neither unified nor conceptually stable. A wide range of labels in German — such as *Zwillingsformel*, *Zwillingswort*, *Hendiadyoin*, *Reduplikation*, *Wortreduplikation*, *Doppelausdruck*, *Paarformel*, *zweigliedrige Komposita*, *Wortpaar*, *Paarwort*, *Dopplung*, *Verdopplung*, *Synonymkompositum*, *Antonymkompositum*, *Synonymenpaar*, *Aufzählung*, *Binom*, *Biverb*, *Reimpaar*, *Worthäufung*, as well as English terms like *compound expression*, *word-pair*, *hendiadys*, *word-combination*, *duplication*, *doubling*, *reduplication*, or *fixed two-word sequence*, etc. — illustrate this lack of terminological consensus. Although many of these labels accurately capture specific subtypes within the wider system, they are frequently used as general designations, which has further contributed to the conceptual ambiguity surrounding the phenomenon.

and (b) through overt conjunctions such as *hām* ‘and; both ... and’, *ya* ‘or’, which likewise introduce explicitly marked coordination (see §2.3 for details). This corresponds to the traditional typological opposition between syndetic and asyndetic coordination. As shown schematically in Figure 1, these two types constitute the highest division of the Turkish irreversible coordination system. The primary division between asyndetic and syndetic coordination corresponds to the absence or presence of explicit conjunctions or coordinating morphology.

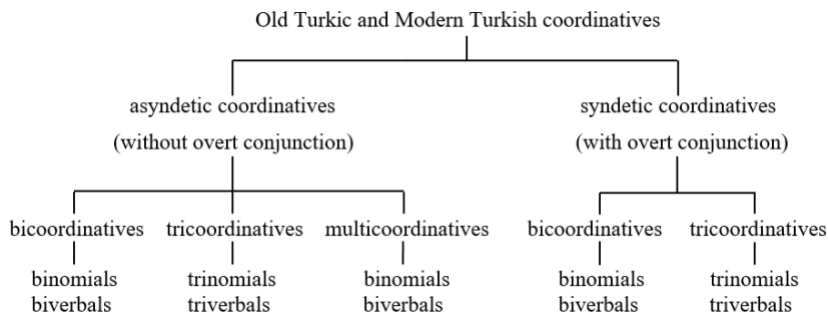


Figure 1. Revised schema of Old Turkic and Modern Turkish irreversible coordinatives (adapted from Aydemir 2007)

While Figure 1 provides a schematic overview of the major coordinative types, concrete realizations of tricoordinatives and multicoordinatives are illustrated and formally modelled in §3.1 (see especially the generalized structural formulas and examples therein).

To accurately assess, characterize, and classify the Turkish cases of irreversible coordinative structures, it is essential to distinguish these major structural types. This differentiation is specific to Turkish and has not been systematically employed in previous research, which is why a precise and comprehensive account of the phenomenon has so far remained lacking.

This conceptual framework may provide a unified basis for analyzing bicoordinatives not only in Turkic and other Altaic languages but also in structurally comparable systems across other language families, including Uralic (e.g. Samoyedic, Hungarian) and Indo-European (e.g. Tocharian, Sogdian), as well as further typologically diverse languages, both synchronically and diachronically.

It should be noted, however, that despite the general typological schema summarized above, the present study focuses exclusively on the formation rules and structural constraints governing irreversible *binary coordinations* (bicoordinatives) in Old and Modern Turkish. Ternary (tricoordinative) formations and multicoordinative chains — while typologically relevant and briefly illustrated where necessary — are not examined in detail in the present work.

The descriptive model developed in the following sections builds on these conceptual distinctions and provides a unified typological and structural account of bicoordinatives in Turkish.

2.2 Symmetry and asymmetry in bicoordinatives

Symmetry and *asymmetry* in ordering refer to two structural subtypes of lexicalized binary coordinative constructions, which are classified here under the broader term *bicoordinative*. These concepts clarify the internal alignment and structural organization of binary patterns governed by morphophonological (see §5.2), morphosyntactic (see §5.2), and iconic principles (see §5.3), and provide a principled basis for their classification.

In this framework, *asymmetry* refers to a structural hierarchy or preference in which the first slot [A] holds a position of priority over the second [B], despite their apparent coordination (e.g. *deli dolu* ‘zappy, high-spirited’ ← [mad / wild] + [full / overflowing]). Thus, (1) in a *symmetric bicoordinative* (SyBi) the two coordinated elements show full formal repetition or structural mirroring with no detectable precedence or internal hierarchy, typically through formal identity, mirroring, or tight formal parallelism (e.g. *yavaş yavaş* ‘gradually’ ← [slow(ly)] + [slow(ly)]; *gide gide* ‘more and more, gradually’ ← [go-CONV] + [go-CONV]), (2) while an *asymmetric bicoordinative* (AsBi) involves a fixed, irreversible order (i.e. [A] + [B] but not [B] + [A]) with structural and/or iconic integration (e.g. *gide gele* ‘constantly coming and going’ ← [go-CONV] + [come-CONV]).

In this framework, asymmetric bicoordinatives involve irreversible ordering between the two coordinated elements and the first element, i.e. the slot [A], occupies a structurally prioritized position. This priority may be grounded in morphophonological constraints (e.g. segmental markedness, sonority, or syllable balance), morphosyntactic structure (e.g. event structure or argument hierarchy), or iconic principles reflecting conceptual sequencing. In other words, asymmetric bicoordinatives encode a structural or iconic hierarchy, which makes the ordering constraint an inherent part of their lexicalization. Thus, the internal hierarchy is a defining feature of lexicalized and irreversible bicoordinatives (e.g. *gide gele* but not *gele gide*).

To sum up, asymmetry, in this framework, entails four key features: (1) the ordering of the slots [A] and [B] is irreversible; (2) there is an implicit structural or semantic hierarchy between the two slots; (3) this hierarchy is morphophonologically, morphosyntactically, or iconically encoded; and (4) the construction is lexicalized and behaves as a fixed coordinative unit. This typological distinction provides a more coherent basis for classification and helps clarify several long-standing descriptive and terminological ambiguities in the literature. Thus, the symmetric vs. asymmetric distinction proposed here offers a more systematic and encompassing framework, within which the various reduplicated subtypes listed above represent internally coherent patterns of symmetric bicoordination.

Within the broader category of symmetric bicoordinatives — defined by full formal mirroring, strict morphological parallelism, and the absence of

an internal precedence — the following formally and functionally distinct subtypes can be identified. The classification proposed here is form-based, i.e. part-of-speech-based (see also Appendix), distinguishing bicoordinatives according to the lexical category of their base elements, even though these constructions may exhibit functional flexibility and frequently occur in adverbial contexts:⁵

- a) *reduplicated adverbial bicoordinatives* (OTu. *yana yana*; MTu. *seke seke, koşa koşa*):⁶ encode manner, intensity, and a range of event-structural or affective nuances associated with verbal action.
- b) *reduplicated adjectival bicoordinatives* (OTu. *öñi öñi*; MTu. *iri iri, koca koca*): encode intensification or augmentative emphasis of a property.
- c) *reduplicated nominal bicoordinatives* (OTu. *ew ew*; MTu. *yer yer, adım adım*): encode distributive or iterative spreading across loci or entities.
- d) *reduplicated distributive bicoordinatives* (MTu. *teker teker, birer birer, azar azar, parça parça*): encode partitive–distributive plurality or stepwise iteration.
- e) *reduplicated pronominal bicoordinatives*:
 - reflexive: OTu. *öz öz, kántü kántü*
 - demonstrative: OTu. *ol ol, bo bo*
 - interrogative-indefinite: OTu. *kayu kayu, kim kim*
 encode partitive-distributive reference or emphatic generalization.
- f) *reduplicated cardinal bicoordinatives* (OTu. *bir bir*; MTu. *bir bir*): encode incremental iteration or distributive quantification.
- g) *reduplicated onomatopoeic bicoordinatives* (MTu. *cik cik, şırl şırl, ciyak ciyak*): encode iconic repetition of sound events.
- h) *figura etymologica-based bicoordinatives* (MTu. *inim inim inle-, horul horul horla-*): encode repetitive or continuous verbal action through figura etymologica.

These subtypes form a coherent internal hierarchy based on three dimensions:

- a) *formal* properties (full reduplication)
- b) *functional* domains (adverbial, adjectival, nominal, distributive, pronominal, cardinal, onomatopoeic, figura etymologica), and
- c) *macro-structural* behavior (symmetric coordination).

These reduplicated subtypes exhibit full formal symmetry and lack any internal hierarchy between the two conjuncts. Despite their diverse func-

⁵ For an alternative, function-based classification of comparable reduplicative constructions in Turkish, see Göksel & Kerslake (2005: 90). The present study, by contrast, adopts a form-based approach, focusing on the lexical category of the base elements rather than their contextual function.

⁶ For ease of reference, representative meanings for each symmetric subtype are provided in Appendix; the examples listed in the main text serve primarily to illustrate structural patterning.

tional domains, they all conform to the symmetric bicoordinative structure, characterized by strict parallelism, mirroring, and the absence of any internal ordering asymmetry.

2.3 Syndetic and asyndetic bicoordinatives

The distinction between *syndetic* and *asyndetic coordination* constitutes one of the highest structural divisions within the system of irreversible bicoordinatives in both Turkish and Old Turkic. Although this typological opposition is well established cross-linguistically, its systematic application to Turkic lexicalized and morphologically integrated coordinative formations had not been undertaken in earlier scholarship. The present study therefore adopts this distinction as a foundational categorization and refines it for the purposes of analyzing Turkish bicoordinatives.

This corresponds to the traditional typological opposition between asyndetic and syndetic coordination (Haspelmath 2007). A more systematic application of this distinction to Turkic was first introduced in an earlier work (Aydemir 2007). In the present study, this distinction is applied specifically to Turkish bicoordinatives in order to differentiate between constructions that rely on explicit coordinators (syndetic) and those that occur without any overt connective (asyndetic). These coordinators fall into two major types — *coordinating suffixes* and *coordinating conjunctions* — which are examined in §§2.3.1–2.3.2. These two types of coordinators discussed in this study are henceforth abbreviated as CoS (*coordinating suffix*) and CoC (*coordinating conjunction*). Although the terminology itself is well established crosslinguistically, no previous study had systematically applied this distinction to Turkish lexicalized binomial and biverbal formations before the earlier work in question on the subject.

Within the broader category of syndetic bicoordinatives, two structural subtypes may be distinguished: *syndetic bicoordination* and the relatively rarer *syndetic tricoordination*. The latter is attested only sporadically in Old Turkic (OTu.) sources and is entirely absent from the Modern Turkish (MTu.) data examined here. The Old Turkic example illustrated in (1) displays a polysyndetic pattern (i.e. [...+II ...+II ...+II]) in which three coordinated elements share the same overt coordinator, typically a coordinating suffix. Although typologically noteworthy, such formations fall outside the scope of the present analysis; their structural representations will be briefly discussed and illustrated in the section on *tricoordinatives* (see §2.4), using a CoP-based notation. The discussion in the remainder of this section focuses exclusively on the structural types of syndetic bicoordination, which are both more frequent and more productive in Turkish.

- (1) *süüli aşı kertgünçli*
 army-CoS provision-CoS faith-CoS
 ‘army, provisions, and faith’⁷

Syndetic bicoordination occurs in both binomial (e.g. MTu. *öyle ya da böyle* ‘one way or another’) and biverbal constructions (e.g. MTu. *gitti de gitti* ‘kept going; went on and on’, see §2.3.2.2). However, unlike syndetic bicoordinatives — where a part-of-speech-based (POS) classification is often necessary (see Appendix) — the syndetic types cannot meaningfully be organized by lexical category, since the attested material is restricted to only a handful of structural patterns. In the verbal domain, for instance, syndetic biverbals are attested only with finite verbs (i.e. MTu. *gitti de gitti*).

For analytical clarity, it is therefore more appropriate to classify syndetic bicoordinatives according to the type of coordinator they employ. This yields two principal subtypes: (a) constructions marked by coordinating suffixes (CoS), and (b) constructions marked by coordinating conjunctions (CoC) (see §§2.3.1–2.3.2). In what follows, both CoS and CoC are treated as subcategories of the more general notion of *coordinator*.

In Old Turkic texts, the postpositive coordinating suffix *+II* — the only morpheme that can be analyzed as a true CoS — is attested frequently in the canonical coordinative pattern [...+II] [...+II].⁸ By contrast, the privative suffix *+sXz* is not a coordinating morpheme; it is a derivational marker expressing semantic privation. In other words, although both OTu. *+II* (= *+lı / +li*) and MTu. *+IX* (< *+IXg* = *+lig / +lig / +lug / +lüg*), as well as *+sXz*, are suffixes in morphological terms, only OTu. *+II* functions as a true coordinating morpheme. By contrast, *+sXz* and MTu. *+IX* are purely derivational and do not encode coordination.

Consequently, formations of the type [...+sXz] [...+sXz] (e.g. MTu. *evsiz barksız*, ‘homeless’) or mixed patterns such as [...+IX] [...+sXz] in Modern Turkish do not constitute syndetic bicoordination in the strict sense. The parallelism observed in these expressions results from semantic and pragmatic alignment rather than from a coordinative suffix. For this reason, *+sXz*-based sequences are treated here not as coordinative constructions but as derivationally parallel adjectival chains, whereas only the [...+II] [...+II] pattern is analyzed as genuine syndetic bicoordination.

⁷ TT V B105. A related issue concerns forms such as Old Turkic *uçsuz kıldısız ülgüsüz* [peak+sXz] [border+sXz] [mass+sXz] ‘boundless [...] and immeasurable’ (merits and good deeds) (Clauson 1972: 31), where a string of adjectives marked with *+sXz* (‘privative’) occurs in what superficially resembles a coordinative pattern. However, the suffix *+sXz* is not a coordinating morpheme but a privative derivational suffix and therefore cannot be analyzed as a CoS. The apparent coordinative effect results from semantic and pragmatic parallelism rather than from morphological coordination. In contrast, formations with [...+II] [...+II] in Old Turkic represent genuine CoS-based syndetic coordination. Consequently, mixed patterns of the type [...+IX] [...+sXz] in Modern Turkish do not constitute true cases of syndetic coordination.

⁸ A comparable classification is found in Johannessen (2003: 22), who explicitly employs the term “coordinating suffix” for Old Turkic formations marked with [...+II] [...+II]. Although her analysis is embedded in a broader typology of coordination in natural language, her terminology overlaps in part with the CoS-based approach adopted here. Typologically comparable with the so-called “bisyndetic” coordinators in Haspelmath (2007: 6). For OTu. *+II*, see Erdal 2004: 166–167.

Most expressions in Modern Turkish that contain surface patterns exhibiting [...+*lX*] [...+*lX*] or [...+*sXz*] [...+*sXz*] are fully lexicalized and frequently occurring as fixed idiomatic forms. These include items such as MTu. *evli barklı* ‘married and having a family’ ← [house+*lX*] [property+*lX*] and MTu. *evsiz barksız* ‘homeless, vagrant’ ← [house+*sXz*] [property+*sXz*], whose suffixal marking is synchronically parallel but not demonstrably related to the Old Turkic coordinative [...+*lI*] [...+*lI*] construction. The Old Turkic expression *uçsuz kıldısız* ‘boundless, limitless’, attested in HT IX, is likewise likely to have been lexicalized already in Old Turkic, since parallel forms occur in several other early sources. Its Modern Turkish reflex is *uçsuz bucaksız* ‘immense, endless’ ← [edge/end+*sXz*] [corner+*sXz*].

Expressions of the form [...+*sXz*] [...+*sXz*] represent not true coordinative structures but rather instances of *litotes* — i.e. semantic intensification through double morphological negation. The two suffixes may also co-occur in mixed patterns in Modern Turkish, although in such cases [...+*lX*] invariably precedes [...+*sXz*], as in MTu. *belli belirsiz* ‘vaguely, indistinctly’ ← [sign+*lI*] [certain+*sXz*]. These patterns are derivationally based and semantically parallel, but they do not constitute syndetic coordination in the strict morphological sense (see fn. 6).

Based on the classification presented in Sections §§2.2–2.3 above, the structural types of syndetic bicoordinatives involving an overt coordinator — where [A] and [B] represent distinct components — can be represented as follows:

2.3.1 Types of coordinators: coordinating suffixes (CoS)

The category of CoS is restricted to formations in which the two conjuncts are marked (a) with the *genuinely coordinating* Old Turkic suffix +*lI*, yielding a coordinated structure through parallel morphology (i.e. [A+*lI*] [B+*lI*]). Historically, this suffix is the only productive instance of a true coordinating morpheme in Old Turkic.

In addition to this Old Turkic pattern, Modern Turkish exhibits two further formations that *superficially* resemble CoS-based coordination but are not synchronically or historically related to the Old Turkic suffix; (b) expressions formed with +*lX*, which display a similar surface pattern due to lexicalization rather than to a productive coordinative function; and (c) fully reduplicated adverbial formations of the type [A+*lX*] [B+*lX*], which likewise represent morphological repetition rather than genuine CoS-based coordination:

- a) [A+*lI*] [B+*lI*], e.g. A+CoS B+CoS (Old Turkic type):

This subtype comprises genuine coordinative formations in Old Turkic, where the suffix +*lI* functions as a productive CoS marker. These constructions encode coordinated nominal structures and are attested widely in Old Turkic inscriptions and texts, as illustrated in (2).

- (2) OTu. *inili açılı* (KT/E6)
younger.brother-CoS elder.brother-CoS

‘younger and elder brothers’

These represent structurally symmetrical formations in which both conjuncts carry the same suffixal marking, forming a coordination through morphological parallelism (see also OTu. *tāyri̇li yerli, sāvinçli korkınçlı, aşnuklı amtiklı*, among other examples).

b) [A+*LX*] [B+*LX*] (Modern Turkish *surface pattern*⁹)

Modern Turkish contains numerous lexicalized expressions, including irreversible bicoordinatives, that display the surface pattern [A+*LX*] [B+*LX*], but the suffix +*LX* (historically < +*LXg* = +*lig* / +*lig* / +*lug* / +*lüg*)¹⁰ is not demonstrably related to the Old Turkic coordinative +*II*. These formations *behave* as if they were CoS-based coordinations, yet their coordinative interpretation arises from lexicalization, rather than from a synchronically productive morphosyntactic rule, as illustrated in (3)–(6). Namely — as suggested above too — MTu. +*LX* (= +*lı* / +*li* / +*lu* / +*lü*) is a derivational suffix and not a coordinating morpheme. It therefore does not encode coordination, even though forms derived with this suffix may of course continue to combine productively in new surface patterns at the synchronic level.

- | | |
|--|---|
| <p>(3) <i>evli barklı</i>
house+<i>LX</i> property+<i>LX</i>
‘married and having a family’</p> | <p>(4) <i>senli benli</i>
you+<i>LX</i> I+<i>LX</i>
‘hobnob (with)’</p> |
|--|---|

- | | |
|---|--|
| <p>(5) <i>derli toplu</i>
order(ed)+<i>LX</i> gather(ed)+<i>LX</i>
‘tidy; well-organized’</p> | <p>(6) <i>allı pullu</i>
red+<i>LX</i> flake/scale+<i>LX</i>
‘showily dressed’</p> |
|---|--|

Although synchronically parallel to the Old Turkic type, this suffix +*LX* belongs to a different morphological lineage, and its coordination-like behaviour results from idiomaticization rather than from the presence of a true coordinating suffix.

c) [A+*LX*] [A+*LX*] (full repetition pattern)

A third subtype consists of non-CoS bicoordinatives formed by full formal repetition of an +*LX*-marked base, yielding adverbial expressions in Modern Turkish, as illustrated in (7) and (8) below. Although not coordinative in the strict morphological sense, these expressions share the *reduplicated adverbial bicoordinative template* and typically encode manner, intensity, or habituality (cf. *reduplicated adjectival bicoordinatives*: MTu. *iri iri, koca koca*; and *reduplicated nominal bicoordinatives*: MTu. *yer yer, kapı kapı, ev ev*, etc.; and note that all such reduplicated types fall under the symmetric bicoordinative category outlined in §2.2; cf. also Appendix).

⁹ Here *surface pattern* refers to the identical morphological template shared by two historically distinct suffixal constructions (i.e. OTu. +*II* vs. MTu. +*LX*).

¹⁰ In the Old Turkic inscriptions (8th c.), +*LXg* was a formative meaning ‘possessing the referent of the base’ (Erdal 1991).

- (7) MTu. *uslu uslu*
 sense/mind+*LX* sense/mind+*LX*
 ‘quietly; obediently; in a well-behaved manner’
- (8) MTu. *edali edali*
 manner/affectation+*LX* manner/affectation+*LX*
 ‘gracefully; charmingly’

Although they share the same template of parallel suffixation, these repetitions do not instantiate coordinative structure in the CoS-based sense described under (a); instead, their interpretation arises from adverbialization through lexicalized reduplication rather than from true CoS-based coordination.

The preceding subsection has established the structural and historical profile of suffix-based coordination (CoS) in both Old Turkic and Modern Turkish, distinguishing genuinely coordinative patterns from superficially similar but non-coordinative surface formations. Building on this foundation, the following subsection turns to coordinating conjunctions (CoC), which constitute the second major type of overt coordination in Turkic. Unlike CoS markers, CoC elements outlined in §2.3.2 are free morphemes and encode coordination through syntactic linkage rather than morphological parallelism.

2.3.2 Types of coordinators: Coordinating conjunctions (CoC)

Within the broader class of syndetic bicoordinatives, one major subtype consists of constructions marked by coordinating conjunctions (CoC). Unlike asyndetic bicoordinatives, which show no overt connective, CoC-based formations employ explicit conjunctions — whether Old Turkic (OTu.) or Modern Turkish (MTu.) — to link the two coordinated components [A] and [B] (i.e. [A + CoC + B]). Although syndetic coordination is cross-linguistically common, its systematic application to lexicalized irreversible bicoordinatives in Turkic has received little attention prior to earlier work. The present study adopts this category as a structural subtype of syndetic bicoordination and integrates it into the general typology of Turkish bicoordinatives.

As noted in §2.3, syndetic bicoordination in Turkic can be realized either by coordinating suffixes (CoS) or by coordinating conjunctions (CoC). In contrast to CoS-based formations — which are morphologically highly restricted — CoC-based coordinatives display relatively greater internal variety. Nevertheless, they do not lend themselves to a part-of-speech-based classification, since attested examples represent only a few structural templates. Therefore, the classification used here is based solely on the type of conjunction employed.

The following subsections provide a systematic overview of CoC-based syndetic bicoordinatives, following exclusively the examples and structural types documented in Old Turkic and Modern Turkish.

2.3.2.1 Disjunctive coordinators:

- [A *ya* B] OTu. (KB) *kâliş ya barış* ‘coming or going; movement’ (cf. asyndetic: HT IX *kâliş ya barış*; cf. below: *kâliş hām barış*), *üküş ya az* ‘more or less’, *kiçig ya ulug* ‘small or big’.
- [A *ya da* B] MTu. *er ya da geç* ‘sooner or later’ (cf. asyndetic: *er geç*),¹¹ *öyle ya da böyle* ‘one way or another’ (cf. asyndetic: *öyle böyle değil* ‘it’s wicked’).
- [A *veya* B] MTu. *er veya geç* ‘sooner or later’ (cf. asyndetic: *er geç*), *öyle veya böyle* ‘one way or another’.

2.3.2.2 Copulative coordinators (i.e. ‘and’-type coordination):¹²

- [A *hām* B] OTu. (KB) *açuk hām yaruk* ‘frank and clear’, *alış hām beriş* ‘buying and selling; trade’, *kâliş hām barış* ‘coming and going; movement’ (cf. above *kâliş ya barış*).
- [A *ve* A] MTu. *ancak ve ancak* ‘if and only if’, *asla ve asla* ‘never and ever’.
- [A *da* A] (Finite verbs): MTu. *gitti de gitti* ‘kept going; went on and on’, *ağladı da ağladı* ‘wept and wept’, *çoştı da çoştı* ‘grew increasingly excited’;
(Nominals): MTu. *neler de neler!* ‘all sorts of things; so many things; a great many things’ (cf. asyndetic: *neler neler!* ‘id.’), *aman da aman* ‘oh my goodness!’
- [[A+*la*] B]¹³ MTu. *kaşla göz arasında* ‘in the blinking of an eye; in no time’ (< *kaş ile göz*), *tuzla buz olmak* ‘to shatter into pieces, break to smithereens’ (< *tuz ile buz*); *akla kara(yı seçmek)* ‘meet a lot of difficulties’ (< *ak ile kara*).

2.3.2.3 Adversative coordinators:

- [A [ama A]]¹⁴ MTu. *asla ama asla* ‘never ever’ (cf. above *asla ve asla* ‘never and ever’), *hiç ama hiç* ‘absolutely not’, *hep ama hep* ‘always, all the time’.

¹¹ Note that the element *er* is no longer productive in Modern Turkish and survives only in fossilized expressions such as *er geç* ‘sooner or later’. The syndetic form *er ya da geç* is a later expansion of this older asyndetic template.

¹² By *copulative coordinators*, we refer to coordinative markers expressing additive or joint relations (‘and’-type coordination), typically used in constructions denoting combination, accumulation, or conjoint participation (e.g. ‘X and Y’). This contrasts with disjunctive coordination (‘or’-type), which encodes alternatives. The term is employed here descriptively, independent of any specific syntactic theory.

¹³ Note that the element +*la* (< *ile*) in formations such as *kaşla göz arasında* and *tuzla buz* represents a morphologized reflex of the coordinating conjunction/postposition *ile* ‘with, and’. Although not a productive coordinating suffix in synchrony, its behavior in these expressions is functionally parallel to true CoS-marking in that it links the two conjuncts through a bound coordinating marker.

¹⁴ Here *ama* functions not as an adversative connector but as a scalar intensifier (‘all the more, emphatically’), a well-known pragmatic reanalysis in Modern Turkish.

2.3.2.4 Pseudo-coordinators:

- [[A *mX*] A]¹⁵ MTu. *güzel mi güzel* ‘very beautiful’, *olur mu olur* ‘it may well happen / it might be possible; never say never’, *hiç mi hiç* ‘not at all; not in the least’.
- [*AbeA*]¹⁶ MTu. *anbean* ‘with every moment, moment to moment’, *özbeöz* ‘real; genuine’, *günbegün* ‘day by day’.

To sum up, CoC-based syndetic bicoordinatives in Turkish exhibit a limited set of structural templates, but they play a crucial role in the typology of irreversible coordinations. Because these constructions are heavily habitualized or lexicalized and structurally restricted, the classification used here — based solely on the type of conjunction — provides a more accurate and consistent typology than a part-of-speech-based approach.

This section integrates all relevant types attested in the corpus (Old Turkic and Modern Turkish) insofar as they are pertinent to the present framework, while maintaining full compatibility with the analytical foundations laid out in the preceding subsections.

2.4 Tricoordinatives

Tricoordinatives in Old Turkic and Modern Turkish are coordinative constructions in which three lexemes occur together, some of which form fixed and lexicalized units, while many others — particularly in Old Turkic — are best understood as occasional formations rather than stable components of the lexicon.

As noted in earlier research,¹⁷ these ternary sequences were often created *ad hoc*, either for metrical or rhythmic purposes in literary compositions or as translation-driven attempts to achieve the most accurate or expressive rendering in religious texts. As will be outlined in §2.4.2, such ternary constructions occur in both asyndetic and suffix-marked syndetic forms, although only the former display the morphophonological cohesion that interacts with ordering constraints. In the Karakhanid period (11-12th centuries), and likewise in Old Uyghur manuscripts (9-14th centuries), such tricoordinatives frequently function as rhetorical devices or textual-stylistic techniques, serving to enhance semantic intensity and stylistic ornamentation. Consequently, most Old Turkic tricoordinatives arise through (a) newly coined combinations (see ex. (9)), (b) extensions of already established bicoordinatives (see ex. (10)), or (c) contaminations involving two pre-existing bicoordinatives, as illustrated in (see ex. (11)). Since the present study focuses on irreversible bicoordinatives, these ternary patterns are discussed only briefly to illustrate how larger coordinative sequences develop through the recursive

¹⁵ The inserted interrogative particle *mX* is synchronically a question marker in Turkish, but in these specific bicoordinative constructions it functions exclusively as an intensifier rather than as a marker of interrogation.

¹⁶ The inserted Persian preposition *be* ‘with, for, in order to’ functions here as an intensifier, and all formations containing *be* are treated as fully lexicalized bicoordinatives.

¹⁷ Aydemir 2007.

expansion of binary structures and how they align with the typology of Turkish coordination.

- (9) OTu. *süüli aşı kertgünçli* ‘army, provisions, and faith’
[[army-CoS] [provision-CoS] [faith-CoS]] (see ex. (19) below)
- (10) OTu. *yertinçü yer suw* ‘the earth; the world’
[[earth] + [ground + water]] (see ex. (14) below)
- (11) OTu. *kork- bāz- titrā-* ‘to be afraid, frightened’
[[to be afraid + to shudder] + [to tremble]] (see ex. (18) below)

2.4.1 Structural domains of tricoordinatives

Old Turkic evidence shows that tricoordinatives occur in three structural domains: adjectival, nominal, and verbal. Although they differ in lexical category and semantic function, all instantiate a recursive Coordination Phrase (CoP) structure that can surface either as tightly integrated asyndetic structures (see §2.4.1.1) or as suffix-marked syndetic formations (see §2.4.1.2).

Old Turkic nominal bicoordinatives display both left-branching and right-branching structures, indicating that the CoP architecture in the nominal domain allows recursive coordination in both directions. In addition to the more familiar left-branching patterns (e.g. *yertinçü yer suw* ‘the earth; the world’, see ex. (14) below), the language also exhibits clear cases of right-branching bicoordinatives (e.g. *aş içgü suwsuş* ‘food and drink’, etc., see ex. (15) below), which ultimately surface as flat two-member sequences. These right-branching patterns are particularly important typologically because they provide the structural basis for certain contamination-derived tricoordinatives. For clarity, both types of bicoordinatives are represented below in their flat surface forms. In other words, Old Turkic nominal coordination is not exclusively left-branching; the system licenses right-branching expansions as well, which surface as flat sequences of two elements (and occasionally three).

A comparable two-way recursive pattern is also attested in adjectival bicoordinatives: Old Turkic exhibits both right-branching structures (e.g. *athig küülüğ bālgülüğ* ‘highly renowned’, see ex. (12) below) and left-branching expansions (e.g. *kirsiz arıg süzök* ‘pure and clean’, see ex. (13) below), both of which surface as semantically intensifying ternary clusters.

For clarity, both branching patterns are presented below in parallel, with their full CoP tree structures and their corresponding flat surface representations.

2.4.1.1 Asyndetic tricoordinatives (AsTri)

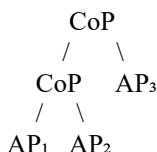
a) Adjectival tricoordinatives (ATri)

Asyndetic adjectival tricoordinatives consist of three adjectives forming a cohesive semantic whole, typically expressing completeness or intensification, as illustrated in (12) and (13). ATri typically produce quality intensification, whereby multiple semantically convergent adjectives combine to yield a cumulative or heightened degree of the property. The adjectives involved are

not required to be fully synonymous; rather, they denote closely related properties within the same semantic domain, whose cumulative combination yields intensification.

- (12) OTu. *atlıg küülig bālgülig*
[[renowned + famous] + [distinguished]] → ‘highly renowned’

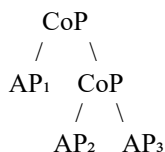
Structural representation:



This right-branching adjectival tricoordinative represents an asyndetic ternary complex consisting of three semantically convergent adjectival items. The first two elements (*atlıg küülig*) form an intermediate coordination, which then combines with the third (*bālgülig*), resulting in a structurally ternary CoP configuration. The overall construction expresses semantic intensification and functions as an asyndetic tricoordinative unit; i.e. [[AP₁ + AP₂] + [AP₃]] (deep) → [AP₁ + AP₂ + AP₃] (surface) = Expansion-type tricoordinative.

- (13) OTu. *kırsız arıg süzök*
[[dirt-free] + [clean + pure]] → ‘pure and clean’

Structural representation:



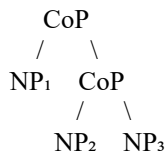
Here, the left-branching *arıg süzök* ‘clean’ is an existing bicoordinative to which *kırsız* is added as a compatible conceptual element. This grouping is motivated by the pre-existing lexicalized bicoordinative *arıg süzök*, which forms a tighter semantic — and likely prosodic — unit, to which *kırsız* is subsequently added; i.e. [AP₁ + [AP₂ + AP₃]] (deep) → [AP₁ + AP₂ + AP₃] (surface) = Expansion-type tricoordinative.

b) Nominal tricoordinatives (NTri)

Asyndetic nominal tricoordinatives typically arise through the lexical expansion of an already established bicoordinative, representing semantic totalities or conceptual domains, as illustrated in (14) and (15) below:

- (14) OTu. *yertinçü yer suw*
[[earth] + [ground + water]] → ‘the earth; the world’

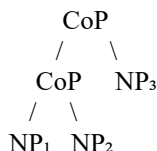
Structural representation:



Here, the left-branching *yer suw* is an existing bicoordinative to which *yertinçü* is added as a compatible conceptual element; i.e. $[[\text{NP}_1] + [\text{NP}_2 + \text{NP}_3]]$ (deep) $\rightarrow [\text{NP}_1 + \text{NP}_2 + \text{NP}_3]$ (surface) = Expansion-type tricoordinative.

- (15) OTu. *aş içgü suwsuş*
 $[[\text{food} + \text{drink}] + [\text{water}]] \rightarrow \text{'food and drink'}$

Structural representation:



The Old Turkic form *aş içgü suwsuş* provides a clear example of a *asyndetic* tricoordinative derived through contamination.¹⁸ The apparent grouping reflects the diachronic derivation of the construction from two overlapping bicoordinatives, rather than a purely synchronic ternary coordination. It originates from two overlapping bicoordinatives — *aş içgü* ‘food and drink’ and *aş suwsuş* ‘food and water’ — in which the shared head (*aş*) is retained only once, while the two second elements (*içgü* and *suwsuş*) are reanalysed as parallel conjuncts (i.e. $aş \text{ içgü} \sim aş \text{ suwsuş} \rightarrow aş \text{ içgü suwsuş}$). This process yields a surface structure consisting of three semantically co-equal nouns, forming a flat ternary CoP that denotes a domain-totality (‘food–drink–water’). Although the underlying formation is of the type $[[\text{NP}_1 + \text{NP}_2] + [\text{NP}_3]]$, the resulting construction functions as a genuine nominal tricoordinative in terms of coordination symmetry and semantic scope; i.e. $[\text{NP}_1 + \text{NP}_2] \sim [\text{NP}_1 + \text{NP}_2] \text{ (contamination)} \rightarrow [\text{NP}_1 + [\text{NP}_2 + \text{NP}_3]]$ (deep) $\rightarrow [\text{NP}_1 + \text{NP}_2 + \text{NP}_3]$ (surface) = Contamination-type tricoordinative.

c) Verbal tricoordinatives (VTri)

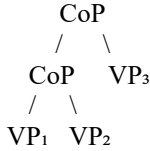
Asyndetic Verbal tricoordinatives arise when three verb stems combine into a cohesive coordinative unit that typically expresses intensified action (quantitative or qualitative), semantic intensification, iterative motion, or a cumulative event sequence, as illustrated in (16)–(18). As in the adjectival and nominal domains, such formations originate either (a) through the expansion of an existing biverbal structure (see ex. (16) and (17)) or (b) through the

¹⁸ In historical linguistics, contamination is a cover term for processes whereby two or more formally or semantically related linguistic forms, coexisting in a speech community, influence each other and merge partially, yielding hybrid outcomes. Such effects may operate at different levels of structure (phonological, morphological, lexical, or constructional). In the present study, the term is used in a more restricted sense, referring specifically to the merging of overlapping coordinative expressions.

contamination of two overlapping bicoordinatives (see ex. (18)), producing ternary event complexes that are semantically integrated, sometimes approximating semantic totality, and morphosyntactically symmetric. The resulting structural patterns closely parallel those observed in ATri and NTri, confirming that all three domains instantiate the same recursive CoP-based coordination mechanism.

- (16) OTu. *agırla- aya- tap-*
[[respect + honour] + [worship]] → ‘to worship; to show reverence’

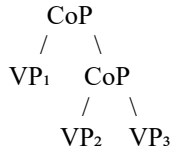
Structural representation:



Here, the right-branching *agırla- aya-* ‘to worship, reverence’ is an existing bicoordinative to which *tap-* ‘to worship’ is added as a compatible conceptual element; i.e. [[VP₁ + VP₂] + [VP₃]] (deep) → [VP₁ + VP₂ + VP₃] (surface) = Expansion-type tricoordinative.

- (17) OTu. *adır- käs- yar-*
[separate + [cut + split]] → ‘to cut and split’

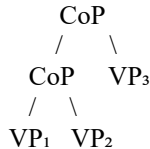
Structural representation:



Here, the left-branching bicoordinative *käs- yar-* ‘to cut and split’ functions as the structural base to which *adır-* ‘to separate’ is added as a semantically compatible element. The three verbal stems belong to the same event domain (‘acts of separation’) and form an asyndetic ternary coordination expressing cumulative and intensified action; i.e. [[VP₁] + [VP₂ + VP₃]] (deep) → [VP₁ + VP₂ + VP₃] (surface) = Expansion-type tricoordinative.

- (18) OTu. *kork- bāz- titrā-*
[[to be afraid + to shudder] + [to tremble]] → ‘to be afraid, frightened’

Structural representation:



The Old Turkic sequence *kork- bāz- titrā-* provides a clear example of a *asyndetic* verbal tricoordinative formed through contamination. As in the nominal contamination case in (15), this structure also originates from two overlapping bicoordinatives — *kork- bāz-* ‘to be afraid; to shudder’ and *kork- titrā-* ‘to be afraid; to tremble’ — in which the shared first element (*kork-*) is retained only once, while the two second elements (*bāz-* and *titrā-*) are reanalysed as parallel verbal conjuncts. The resulting ternary complex forms an *asyndetic* CoP that denotes an intensified fear-related event domain (‘fear–shudder–tremble’). Although the underlying configuration is of the contamination type $[VP_1 + VP_2] \sim [VP_1 + VP_3]$, the reanalysed structure surfaces as a fully integrated *asyndetic* verbal tricoordinative exhibiting coordinative symmetry; i.e. $[VP_1 + VP_2] \sim [VP_1 + VP_3]$ (contamination) $\rightarrow [[VP_1] + [VP_2 + VP_3]]$ (deep) $\rightarrow [VP_1 + VP_2 + VP_3]$ (surface) = Contamination-type tricoordinative.

Numerous further examples of this pattern are attested in Old Turkic, including forms such as *kork- ürk- bālinglä-* ($\leftarrow$ *kork- ürk-* $\sim$ *ürk- bālinglä-*) ‘to be afraid, to be frightened’ and *aç- yad- yarut-* ($\leftarrow$ *aç- yad-* $\sim$ *aç- yarut-*) ‘to expose, reveal, uncover’, among others. These sequences arise through the same mechanism of overlapping bicoordinatives with a shared initial element and a reanalysis of the remaining verbal stems as parallel conjuncts.

From a typological perspective, and for now, it seems that contamination-based *asyndetic* verbal tricoordinatives in Old Turkic predominantly follow right-branching patterns; no unequivocal left-branching contamination type has yet been identified in the available corpus.

Taken together, the *asyndetic* verbal tricoordinatives analysed above demonstrate that Old Turkic employs more than one structural route in building ternary verbal complexes. Both right-branching *asyndetic* expansion patterns (e.g. *agırla- aya- tap-*, (16)) and left-branching *asyndetic* ones (e.g. *adırl- käs-yar-*, (17)) are attested, each reflecting the recursive application of the CoP mechanism to semantically compatible verb stems. In addition, *asyndetic* contamination-type formations (e.g. *kork- bāz- titrā-*, (18)) show that overlapping bicoordinatives can likewise give rise to fully integrated ternary structures. Across these patterns, the resulting sequences form *asyndetic* serial coordinations expressing intensified, cumulative, or multi-step event structures.

To sum up, from a typological standpoint, Old Turkic exhibits two structurally distinct types of expansion-based *asyndetic* verbal tricoordinatives. Some are right-branching, where an existing bicoordinative receives an additional verbal element in final position (as in example (16)), while others are left-branching, where the added verb precedes an already established bicoordinative structure (as in example (17)). Both patterns surface as flat ternary sequences but differ in their internal CoP configuration — $[[VP_1 + VP_2] + [VP_3]]$ versus $[[VP_1] + [VP_2 + VP_3]]$ — demonstrating that the verbal domain allows *asyndetic* recursive coordination in both structural directions. This bidirectional expandability provides the morphosyntactic basis for the full range of *asyndetic* verbal tricoordinatives attested in Old Turkic.

2.4.1.2 Syndetic tricoordinatives (SyTri)

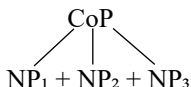
Syndetic tricoordinatives in Old Turkic are extremely rare and, unlike their adjectival, nominal and verbal asyndetic counterparts, do not constitute a productive or structurally diverse category. The only securely attested pattern involves polysyndetic nominal formations in which three coordinated elements are marked by the same coordinating suffix (CoS) — i.e. [...+// ...+// ...+//] — yielding a formally syndetic but semantically unified ternary sequence.

As noted earlier (see §2.3), such formations occur only sporadically in the Old Turkic corpus and have no counterparts in Modern Turkish. Moreover, verbal SyTri structures are unattested, which is expected given that the relevant coordinating suffix (+//) is morphologically incompatible with verbal stems. In principle, a CoC-based (conjunction-marked) tricoordinative would also be typologically possible, but no such example has been identified in Old Turkic textual sources.

The following example — identical to the one introduced in §2.3 — illustrates the only securely documented Old Turkic SyTri configuration, consisting of three nominals, each bearing the coordinating suffix +//:

- (19) OTu. *süüli aşı kertgünçli* (TT V B105)
[[army-CoS] [provision-CoS] [faith-CoS]] → ‘army, provisions, and faith’

Structural representation:



This pattern shows that syndetic tricoordination exists in Old Turkic only as a marginal subtype of nominal coordination, limited to occasional polysyndetic constructions and lacking evidence of broader structural generalization.

2.4.1.3 Idiomatic and formulaic tricoordinatives in modern Turkish

Modern Turkish displays only limited productivity in the formation of tricoordinatives. Unlike Old Turkic — where ternary constructions occur across adjectival, nominal, and verbal domains through both expansion and contamination — the modern language preserves only a small set of habitualised, conventionalised, and stylistically fixed ternary expressions. Although these items do not originate from productive CoP-recursion in the synchronic grammar, they nevertheless satisfy the structural criteria for tricoordinatives: three lexemes forming a cohesive semantic and prosodic unit.

One well-known example is *ham hum şaralop* (or *şorolop*), which originates from the contamination of the bicoordinative *ham hum* ‘mumbling; gibberish’ and the onomatopoetic form *şaralop* (~ *şorolop*), yielding a ternary complex meaning ‘vague muttering; nonsensical talk’. The structure consists of three phonologically and semantically compatible components forming an asyndetic tricoordinative unit.

Other idiomatic expressions likewise surface as fixed ternary patterns in the modern language, including:

- (20) *öldüm öldüm dirildim*
öl-dü-m öl-dü-m diril-di-m
die-PST-1SG die-PST-1SG revive-PST-1SG
‘I was so scared / terrified’ (*lit.* ‘I died, died, and came back to life’)
- (21) *ye iç yat*
ye iç yat
eat-IMP.2SG drink-IMP.2SG lie.down-IMP.2SG
‘eat, drink, and sleep’
- (22) *yandım bittim kül oldum*
yan-dı-m bit-ti-m kül ol-du-m
burn-PST-1SG finish-PST-1SG ash become-PST-1SG
‘I was utterly devastated’ (*lit.* ‘I burned, I was finished, I became ash’)
- (23) *el pençe divan*¹⁹
el pençe divan
hand palm court
‘showing extreme respect; in a deeply reverential posture’

These sequences are synchronically habitualised, functioning as formulaic or idiomatic tricoordinatives whose ternary structure is preserved through frequent usage. Although they do not participate in the productive derivational pathways attested in Old Turkic (i.e. expansion-based or contamination-based CoP recursion), they display the same surface property: three coordinated lexical items forming a cohesive ternary semantic unit.

Overall, Modern Turkish maintains the structural possibility of tricoordination, but its contemporary system relies almost exclusively on habitualised, conventionalised, idiomatized, fixed ternary expressions, rather than the morphologically and stylistically productive patterns characteristic of Old Turkic.

2.4.1.4 Functional interpretation of tricoordinatives (ATri, NTri, VTri)

Across adjectival, nominal, and verbal domains, tricoordinatives contribute a wide range of semantic, aspectual, and stylistic effects. Structurally, they emerge through the recursive expansion or contamination of bicoordinatives; functionally, they yield tightly integrated ternary units that amplify or elaborate the conceptual space encoded by their components. These formations thus operate not only as grammatical coordinations but also as devices of semantic enrichment and stylistic intensification, especially in Old Turkic literary and religious texts. The following functional domains sum-

¹⁹ The expression *el pençe divan* historically derives from an earlier adverbial construction in which the bicoordinative *el pençe* functioned as an independent manner adverb (‘with hands clasped in front, respectfully’) modifying the verb phrase *divan durmak* (‘to bow, to stand in reverence before someone’). Through conventionalization, the adverbial modifier and the verbal base underwent lexical fusion, yielding the fixed idiom *el pençe divan* with the adverbial meaning ‘showing utmost reverence’.

marize the principal interpretive effects associated with ATri, NTri, and VTri structures:

- a) *Iterative event sequencing*: repetitive or cyclic action; OTu. *kork-ürk-bälinglä-* ‘to be afraid, to be frightened’ (see §2.4.1.1(c)).
- b) *Intensified action*: quantitative or qualitative strengthening of the event; OTu. *agırla- aya- tap-* ‘to worship; to show reverence’ (ex. (16)).
- c) *Semantic intensification*: cumulative reinforcement of a conceptual domain; OTu. *atlıg küülüg bälgülüg* ‘highly renowned’ (ex. (12)).
- d) *Semantic totality*: near-exhaustive or domain-wide coverage of related concepts; OTu. *aş içgü suwsuş* ‘food–drink–water’ → ‘food and drink (in general)’ (ex. (15)).
- e) *Aspectual culmination*: cumulative or progressive event structuring; OTu. *adır- käs- yar-* ‘to cut and split’ (ex. (17)).
- f) *Property intensification*: a heightened degree or cumulative buildup of a quality; OTu. *kirsiz arıg süzök* ‘pure and clean’ (ex. (13)).
- g) *Rhetorical or textual–stylistic function*: enhancing rhythm, emphasis, and expressive force in discourse; e.g. *öldüm öldüm dirildim* ‘I was so scared / terrified’ (lit. ‘I died, died, and came back to life’) (ex. (20)).

While tricoordinatives represent the minimal extension of binary coordination beyond two components, Old Turkic and Modern Turkish also exhibit larger coordinative chains consisting of more than three elements. These formations go beyond ternary structures and give rise to what may be termed multicoordinatives, i.e. coordinative constructions involving four or more successive elements. As will be shown in the following section, such multi-coordinative structures emerge either as simple enumerations or as concatenated chains of bicoordinatives, and they play a distinct role in the organization of complex coordinative sequences in Turkic.

2.5 Multicoordinatives (enumerations and concatenations)

Multicoordinatives are coordinative constructions consisting of more than three successive elements. They occur with both nominal and verbal bases and, in principle, allow an unrestricted number of coordinated expressions of varying internal complexity. Under the umbrella term multicoordinative, two structurally and functionally distinct phenomena are subsumed in the present framework:

- a) *simple enumerations* involving more than three coordinated items, and
- b) *concatenations (chainings) of bicoordinatives*, i.e. the recursive coordination of already established bicoordinative units without overt conjunctions.

For the first type, the widely used cross-linguistic term *enumeration* is retained. For the second type, in this study, the term *concatenation* is used in its standard linguistic sense to denote the process — and its result — of the rule-governed, linear placement of linguistic elements, forming strings in which elements stand in a relation of succession (e.g. $X + Y + Z$) (Bussmann 1998; Crystal 2008). This foundational process plays a central role in many approaches to morphological and syntactic theory (see, e.g. Haspelmath & Sims 2010). The present analysis, however, does not engage with theoretical debates on concatenation *per se*, but focuses instead on its empirically observable phonetic and morphophonological correlates in Turkish bicoordinatives. In contrast to simple enumerations, concatenations of bicoordinatives typically display internal structural organization, semantic domain coherence, and rhetorical structuring.

The structures described above as “concatenations (chainings) of bicoordinatives” are henceforth referred to as *multicoordinatives* (MulCo). In the present framework, multicoordinatives are defined as linear coordinative chains formed by the successive alignment of two or more independent bicoordinatives. Unlike tricoordinatives, which involve exactly three coordinated elements within a single recursive CoP domain, multicoordinatives represent higher-order coordinative chains resulting from the iterative linking of bicoordinative units. Depending on the lexical category of their components, three subtypes are distinguished here: NMulCo (nominal multicoordinatives), VMulCo (verbal multicoordinatives), and AMulCo (adjectival multicoordinatives). Structurally, these formations instantiate stepwise recursive coordination at the level of bicoordinative blocks, while functionally they serve as powerful devices of textual structuring, semantic amplification, and rhetorical intensification in Old Turkic.

It should furthermore be noted that a single multicoordinative (MulCo) structure may itself contain more than one NMulCo, AMulCo, or VMulCo segment in succession, yielding mixed nominal, adjectival, and verbal chains within the same linear coordinative sequence, as illustrated in (31) and (32) below.

2.5.1. Enumerations

Most enumerations in Old Turkic are syntactically conditioned and largely context-bound. As already observed by Röhrborn, many of them represent “syntactic chance products” (i.e. *nonce constructions*), especially in translation texts (Röhrborn 1983). In Buddhist Old Turkic, for example, long nominal lists frequently arise through near-literal word-for-word renderings from a source language (most often Chinese). A representative case is the following passage from the Old Turkic *Xuanzang-Biography* (see ex. (24)):

- (24) OTu. *altun, kümüş, tuç, ud, koyn, titir tävä, kızıl tuz, tış tuz, kara tuzta ulatı*
 ‘gold, silver, bronze, cattle, sheep, female and male camels, red salt, white salt, black salt, etc.’

Such enumerations primarily reflect source-text structure and therefore remain of limited structural relevance for coordination typology. They do not normally contain true bicoordinatives or tricoordinatives as internal building blocks and thus lie largely outside the core domain of irreversible coordination.

2.5.2 Concatenations of bicoordinatives

Concatenations are structurally and functionally highly informative, as they are formed by the successive alignment of multiple bicoordinatives. These formations resemble bicoordinatives and tricoordinatives in their role as text-structuring devices and appear to serve expressive, rhetorical, and stylistic purposes such as intensification, variation, elaboration, emphasis, emotional colouring, and semantic specification. Typical examples from Old Turkic illustrate how several bicoordinatives can be aligned consecutively, as illustrated in (25)–(31).

Throughout this section, coordinative formations are classified according to the part-of-speech (POS) category of their constituents, not according to the semantic or functional interpretation they yield in context:

- (25) OTu. *miñ miñ tūmān tūmān*
[[thousand] [thousand]] + [[ten-thousand] [ten-thousand]] = MulCo
'countless, innumerable'
- (26) OTu. *öñi öñi adrok adrok*
[[various] [various]] + [[different] [different]] = AMulCo
'each different in its own way'
- (27) OTu. *eş tuş adaş kudaş*
[[spouse] [associate]] + [[friend] [comrade]] = AMulCo
'friends; kith and kin'
- (28) OTu. *aya- ağırla- tapın- udun-*
[[to honour] [to respect]] + [[to worship] [to reverence]] = VMulCo
'to reverence and venerate'
- (29) OTu. *yaşur- batur- ürt- kizlä-*
[[to hide] [to conceal]] + [[to cover up] [conceal]] = VMulCo
'to hide, conceal'
- (30) OTu. [*yıd yıpar*] [*hwa çäçäk*] *üzä* [*tapıg udug*] *kılsarlar*
incense incense flower flower with offering offering do-COND-3PL = NMulCo
'If they perform offerings with incense and flowers'
- (31) OTu. [*kök kalıg*] *täg* [*uçsuz tüpsüz*] [*öñi öñi*] [*basdaçı täpdäçi*] *ärür*
(Kara & Zieme 1977)
sky heaven like end-PRIV bottom-PRIV different different tread-PTCP step-PTCP be-AOR.3SG → 'Like the sky he is boundless; he treads everything'

In such cases, the individual bicoordinatives remain structurally intact but are arranged into a higher-order coordinative chain. Semantically, these concatenations often cover a coherent conceptual field and evoke effects of hyperbole, rhetorical amplification, and cumulative elaboration. An extreme illustration is provided by extended multicoordinative chains such as the following (see ex. (32)), which consists of thirteen successive bicoordinatives:

- (32) [şlok nom]₁ *ârdini için tol* [ellärin uluşların]₂ [kântlärin suzakların]₃ [ävlärin barkların]₄ [orduların karşılärin]₅ [atların yañaların]₆ [agıların barımlärin]₇ [ädlärin tavarların]₈ [ârdinilärin yinçülärin]₉ [ogulların kızların]₁₀ [kunçuyaların hatunların]₁₁ [amrak isig]₁₂ *özlärinä [tidip ıdalap]₁₃ temin ök nom ârdinig âşidgäli boltılar.* (Zieme 1996)

‘Because these (and other Bodhisattvas), for the sake of the teaching contained in a single śloka,₁ relinquished and abandoned₁₃ their entire lands and realms,₂ their cities and villages,₃ their houses and dwellings,₄ their palaces and royal residences,₅ their horses and elephants,₆ their treasures and possessions,₇ their goods and property,₈ their jewels and pearls,₉ their daughters and sons,₁₀ their wives and concubines,₁₁ (even) their own beloved₁₂ selves, they were then able to hear the Dharma-jewel at once.’

Here the expressive force of the passage (i.e. (32)) largely derives from the cumulative effect of concatenated bicoordinatives rather than from simple syntactic listing. These extended chains are not lexicalized fixed expressions; their expressive effect instead arises from the sequential combination of pre-existing bicoordinatives. In translated Old Turkic passages such as those of the Altun Yaruk (*Suvarṇaprabhāsottamasūtra*, cf. (32)), this ordering is not random but reflects the parallel enumerative structure of the source text.

Finally, Modern Turkish still preserves related multicoordinative behaviour, though with far more limited productivity, as illustrated in (33), where the multicoordinative also based on concatenated bicoordinatives:

- (33) MTu. *gide gele gide gele*
go-CONV come-CONV go-CONV come-CONV
‘coming and going again and again’

Formally, multicoordinatives (MulCo) can be modelled as higher-order recursive coordinations formed through the successive alignment of bicoordinative blocks – which surface as a flat linear sequence – rather than individual lexemes. This process may be represented schematically as $[[Bi_1 + Bi_2] + Bi_3 + \dots + Bi_n]$ (deep) $\rightarrow$ [MulCo] (surface), where “Bi” denotes any nominal, verbal, or adjectival bicoordinative (i.e. NBi, VBi, ABi), and “n” is in principle unbounded (see §2.5.3 below).

Depending on the lexical category of the chained blocks, the resulting structure may be classified as NMulCo (nominal), VMulCo (verbal), or AMulCo (adjectival). Furthermore, mixed multicoordinatives are also possi-

ble, in which nominal, verbal, and adjectival bicoordinatives occur mixed together in any order within the same MulCo chain; e.g. $[[NBi_1 + VBi_2] + ABi_3 + \dots + XBi_n] \rightarrow [MulCo]$ (see §2.5.3 below). A MulCo necessarily requires the successive coordination of at least two bicoordinatives; that is, MulCo structures minimally begin with two bicoordinative units.

2.5.3 Structural schemata of multicoordinatives (MulCo)

For the formal generalization of multicoordinatives as *flat linear alignments* at the surface level and as *chains of recursive bicoordinatives* at the deep level, the following structural schemata are proposed:

- a) General structural schema of multicoordinatives (surface):
 $[[Bi_1 + Bi_2] + Bi_3 + \dots + Bi_n] \rightarrow [MulCo]$,
 where $Bi \in \{NBi, VBi, ABi\}$.

This schema captures the surface realization of MulCo as a flat linear concatenation of independently formed bicoordinative blocks.

- b) Structural schemata of the deep subtypes:
1. Homogeneous multicoordinatives:
 Nominal:
 $[[NBi_1 + NBi_2] + NBi_3 + \dots + NBi_n] \rightarrow [NMulCo]$
 Verbal:
 $[[VBi_1 + VBi_2] + VBi_3 + \dots + VBi_n] \rightarrow [VMulCo]$
 Adjectival:
 $[[ABi_1 + ABi_2] + ABi_3 + \dots + ABi_n] \rightarrow [AMulCo]$
 2. Mixed (heterogeneous) multicoordinatives
 $[[NBi_1 + VBi_2] + ABi_3 + \dots + XBi_n] \rightarrow [MulCo]$,
 where $XBi \in \{NBi, VBi, ABi\}$.

The latter represents mixed multicoordinatives, in which nominal, verbal, and adjectival bicoordinatives may occur successively mixed together in any order within the same coordinative chain.

Adverbial bicoordinative (AdvBi) or tricoordinative (AdvTri) sequences in Turkish are treated here as category-shifted realizations of adjectival or verbal bicoordinatives at the level of syntactic function, and no independent AdvBi or AdvTri classes are therefore postulated.

2.5.4 Typological interpretation of multicoordinatives

Typologically, Turkish coordination displays a strictly hierarchical expansion pathway from BiCo $\rightarrow$ TriCo or BiCo $\rightarrow$ MulCo. While tricoordinatives emerge through single-step recursive expansion or contamination of bicoordinatives, multicoordinatives represent higher-order coordinative chains formed by the iterative concatenation of independent bicoordinative blocks. Unlike simple enumerations, MulCo structures thus instantiate recursion at the level of coordinative units rather than at the level of individual lexical items. A crucial theoretical consequence of the present model is, thus, that in multi-

coordinatives, the basic unit of recursion is not the individual lexeme but the bicoordinative block itself.

The structural schemata above (see §2.5.3) formalize the fact that multicoordinatives are not basic coordinative units, but higher-order constructions derived through iterative recursive coordination of bicoordinative blocks. The model captures:

- a) the flat linear appearance of MulCo at the surface level,
- b) the stepwise recursive structure at the deep level,
- c) the existence of both homogeneous and mixed multicoordinative chains, and
- d) the theoretically unlimited extensibility of coordinative concatenation.

To sum up, from a structural point of view, multicoordinatives constitute the outermost extension of the recursive coordinative system described in this study. While enumerations primarily reflect syntactic and translational mechanisms, chains of bicoordinatives reveal the full stylistic and rhetorical potential of recursive coordination. They show that the same morphophonological and iconic principles governing irreversible bicoordinatives also operate at higher levels of coordinative complexity.

3. Typological Overview of Turkish Coordinative Constructions (BiCo–TriCo–MulCo)

The present BiCo–TriCo–MulCo model intersects with Haspelmath's cross-linguistic typology of coordination (Haspelmath 2004, 2007) in recognising coordination as a structurally complex and typologically variable construction. Haspelmath's framework primarily classifies coordination with respect to linking strategies (conjunctions vs. affixes), symmetry and asymmetry of conjuncts, and clause-level coordination patterns. The present study converges with this approach in treating coordination as an internally structured constructional domain, but diverges in analytical focus by modelling coordination as a block-recursive system operating on coordinative units (*bicoordinatives*, *tricoordinatives*, *multicoordinatives*) rather than on individual conjuncts alone. While Haspelmath's typology is largely conjunction- and clause-oriented, the BiCo–TriCo–MulCo framework foregrounds the recursive expansion of coordinative blocks as higher-order constructional units, thereby providing a complementary perspective on the structural organization of coordination, particularly in morphologically rich languages such as Turkish. While both approaches converge in recognising structurally complex coordination, they diverge in their analytical focus: Haspelmath's typology is primarily clause- and conjunction-oriented, whereas the present framework is block-recursive and construction-based.

3.1 Generalized structural formulas

This subsection presents a set of generalized structural formulas that abstract and schematically capture the recursive architecture of coordinative constructions developed in this study. The formulas are intended as formal summaries of the underlying combinatorial principles governing bicoordinatives, tricoordinatives, and multicoordinatives, highlighting how simple binary coordination serves as the structural base for higher-order recursive extensions.

- a) Bicoordinatives (BiCo)

$$[X_1 + X_2] \rightarrow [\text{BiCo}]$$
- b) Tricoordinatives (TriCo)

Expansion-type:

Right-branching: $[[X_1 + X_2] + [X_3]] \rightarrow [X_1 + X_2 + X_3]$

Left-branching: $[[X_1] + [X_2 + X_3]] \rightarrow [X_1 + X_2 + X_3]$

Contamination-type:

$$[X_1 + X_2] \sim [X_1 + X_3] \rightarrow [X_1 + X_2 + X_3]$$
- c) Multicoordinatives (MulCo)

General recursive schema:

$$[[\text{Bi}_1 + \text{Bi}_2] + \text{Bi}_3 + \dots + \text{Bi}_n] \rightarrow [\text{MulCo}]$$

Mixed (heterogeneous) multicoordinatives

$$[[\text{NBi}_1 + \text{VBi}_2] + \text{ABi}_3 + \dots + \text{XBi}_n] \rightarrow [\text{MulCo}]$$

4. Main Theoretical Contributions of the BiCo–TriCo–MulCo Model

- a) *Block-recursive view of coordination:*

The model reconceptualizes coordination as a recursive system operating on bicoordinative (BiCo) blocks rather than on isolated lexical conjuncts, showing that higher-order coordinatives (TriCo, MulCo) emerge through the iterative coordination of already coordinated units.

- b) *Unified treatment of expansion and contamination:*

It provides a unified structural account of both *expansion-type* and *contamination-type* tricoordinatives within a single CoP-based recursive framework. The model provides the first integrated structural account of both expansion-type and contamination-type tricoordinatives within a single recursive CoP-based framework, demonstrating that these two formation pathways are not independent mechanisms but structurally related outcomes of the same coordinative recursion.

- c) *Formal modelling of multicoordinatives:*

Multicoordinatives are formally characterized for the first time through generalized recursive schemata (e.g. $[[\text{Bi}_1 + \text{Bi}_2] + \text{Bi}_3 + \dots + \text{Bi}_n] \rightarrow [\text{MulCo}]$), capturing both their flat surface appearance and stepwise deep recursive structure.

d) *Partial projection of ordering principles (OPri) across recursive coordinative levels:*

The model shows that the morphophonological, morphosyntactic, and iconic ordering principles (OPri) that strictly govern irreversible bicoordinatives do not transfer mechanically or fully to higher coordinative configurations. These principles remain partially active in tricoordinatives and multicoordinatives, but their governing and ordering force progressively weakens as coordinative blocks expand and are chained sequentially.

In multicoordinatives consisting of two bicoordinative blocks (i.e. $[Bi_1 + Bi_2]$), the irreversible structure of the individual bicoordinatives still allows OPri to operate to a substantial degree across the entire construction. By contrast, once the number of coordinated blocks exceeds two, yielding long multicoordinative chains, the blocks are primarily aligned linearly and sequentially, and no consistent global ordering or iconic priority is observed between adjacent blocks. In such cases, OPri continues to operate internally within each bicoordinative block, but no longer constrains the ordering relations between blocks.

A similar partial projection is observable in contamination-based tricoordinatives formed from the interaction of two pre-existing bicoordinative blocks (e.g. in TriCo contaminated of two bicoordinative blocks, $[Bi_1 + Bi_2]$), where OPri applies locally within the source bicoordinatives but does not always or consistently extend uniformly to the newly formed three-component structure.

e) *A unified synchronic and diachronic framework for BiCo–TriCo–MulCo structures:*

The present study proposes, for the first time, a unified synchronic and diachronic theoretical framework for the systematic classification and formal modelling of irreversible coordinative constructions in Modern Turkish and Old Turkic under the BiCo–TriCo–MulCo architecture. While reduplicative and coordinative formations have long been recognised in the literature, they have typically been treated as isolated lexical or stylistic phenomena under diverse labels. The present model integrates these formations into a single recursive system, explicitly distinguishing bicoordinatives, tricoordinatives, and multicoordinatives as structurally related but hierarchically differentiated construction types, and formalizes their relations by means of generalized structural formulas.

Although the model is developed on the basis of Modern Turkish and Old Turkic data, its block-recursive architecture is formulated at a level of abstraction that may perhaps facilitate cautious cross-linguistic comparison — primarily with other Turkic languages and, secondarily, with certain Altaic and some Uralic languages (e.g. Samoyedic, Hungarian), and perhaps even with structurally compatible patterns observed in extinct languages such as Tocharian and Sogdian.

5. Structural, Constructional, and Iconic Constraints Governing Bicoordinatives in Turkish

The question of whether Turkish bicoordinatives follow systematic structural principles has been raised since the earliest studies on Turkish irreversible coordination,²⁰ yet its full resolution has remained unattained. The view that Turkish bicoordinatives constitute *irregular formations* stems largely from the absence of a coherent model capable of capturing their internal ordering constraints. Yet, as both descriptive and comparative data increasingly suggest, these formations are far from arbitrary.

Early work by Foy (1899), Çağatay (1940–41), and Ağakay (1954) correctly identified several morphophonological and morphosyntactic tendencies governing the internal arrangement of bicoordinatives. Their observations, however, were never integrated into a unified theoretical framework. This has contributed to the assumption by some scholars²¹ — as mentioned above — that Turkish bicoordinatives are irregular or unpredictable constructions.²²

A re-examination of the phenomenon shows, however, that Turkish bicoordinatives exhibit a robust set of *Ordering Principles* (OPri), rooted in their morphophonological structure, morphosyntactic and iconic alignment. These collectively determine the irreversible ordering characteristic of asymmetric bicoordinatives. Some of these principles were implicitly recognised in earlier scholarship, while others emerge clearly only when a sufficiently broad synchronic and diachronic dataset is brought into the analysis. When taken together, these principles reveal that irreversible bicoordinatives form a structurally constrained domain rather than irregular formations.

The following sections therefore reconstruct and refine the formation rules that govern Turkish asymmetric bicoordinatives (AsBi). These rules are based on three sources:

- a) some of the earlier descriptive observations of Foy, Çağatay, and Ağakay,²³
- b) subsequent observations in the typological and phonological literature, and

²⁰ Foy (1899) appears to be the first to address this question.

²¹ Schilling remarks, for example, that “morphophonological criteria apparently play no decisive role in the formation of Old Turkic paired expressions” (Schilling 2001: 159). Müller also considers the rules proposed in the literature to be “vague” (Müller 2004: 46, 50).

²² A detailed critical discussion of earlier descriptive accounts, including those of Foy (1899) and Çağatay (1940–41), is provided in Aydemir (2007) and will not be repeated here.

²³ Several descriptive observations in the earlier literature can be retrospectively related to specific ordering principles formulated in the present study. Thus, Foy (1899) correctly identified a number of tendencies corresponding to Rules 3, 4, 5, and 6, particularly with regard to segmental composition, vocalic onset, syllable count, and syllable structure. Çağatay’s (1940–41) observations can be most directly associated with sonority-based ordering tendencies (Rule 2a), while Ağakay (1954) anticipated aspects of vowel-based preferences (Rule 1) and event-structural sequencing (Rule 7). However, none of these earlier accounts formalised these tendencies as hierarchically organised rules or recognised their systematic interaction across different structural levels. Subsequent studies, often without direct reference to these earlier works, have tended to reiterate individual observations in isolation rather than integrating them into a comprehensive explanatory framework or a unified model.

- c) the results of my own analyses of both Old Turkic and Modern Turkish datasets.

Two clarifications are necessary before introducing these rules: (1) First, they apply exclusively to lexicalised asymmetric bicoordinatives in Modern Turkish. Non-lexicalised, compositional, or contextually generated coordinations represent structurally distinct phenomena and are therefore not subject to OPri-based constraints. (2) Second, while the rules are not constructed with Old Turkic as their primary domain, preliminary comparison shows that a large proportion of Old Turkic examples examined by me conform to the same OPri-based tendencies.

The effectiveness of the criteria proposed here for predicting component order was tested on over 180 Old Turkic and over 400 Modern Turkish bicoordinatives examined by me.²⁴ These numbers do not represent the total size of the corpora, of course, but they are sufficient to reveal stable and recurrent ordering patterns across both diachronic and synchronic periods.

Understanding the structural organisation of Turkish bicoordinatives thus requires acknowledging two foundational observations. (1) First, the irreversible ordering of asymmetric bicoordinatives is governed by a system of morphophonological ordering principles that determine the relative prominence and placement of their components. (2) Second, these principles interact systematically with broader morphosyntactic and iconic factors, giving rise to stable patterns of lexicalised asymmetric binary formations. The following subsections present these principles and prepare the ground for their formal statement in Rules 1–7.

5.1 Phonological and sonority-based constraints (Rules 1, 2a, 2b)

Phonological structure provides the most fundamental layer of asymmetry in Turkish bicoordinatives. Before morphophonological, syntactic, or iconic factors come into play, the ordering of the two components is shaped by basic principles of vowel sonority, consonantal sonority, and articulatory progression. These principles jointly form the phonetic–phonological core of the OPri model, determining the natural direction in which sound sequences tend to unfold within irreversible bicoordinatives.

At this foundational level, the internal order of Turkish bicoordinatives reflects a consistent interaction between acoustic openness, articulatory effort, and segmental mobility.

- a) Rule 1 (see §5.1.1) captures the tendency for pairs to begin with more sonorous (i.e. lower or more open) vowels and proceed toward less

²⁴ I present here a compiled set of Modern Turkish data, partly drawing on materials from Ağakay (1954) and Hatiboğlu (1981), while omitting some items from these sources since many of them are not lexicalised, and adding several others. All Modern Turkish examples included here are well-formed and lexicalised. Italic round brackets “()” indicate frequently occurring inflected forms; non-italic round brackets “()” indicate commonly attested verbal forms or variants; for Old Turkic data, see Aydemir 2007 (earlier list) and Aydemir 2013 (non-list data).

sonorous ones; e.g. *tek tük* ‘here and there’, *çar çur* (et-) ‘to squander’, etc.

- b) Rule 2a (see §5.1.2) extends this logic to consonants, showing that the first member typically contains a segment with higher consonantal sonority, while the second tends to contain one with lower sonority; e.g. *döl döş* ‘descendence, children’, *dere tepe* ‘over hill and dale’, etc.
- c) Rule 2b (see §5.1.3) accounts for cases where vowel and consonant sonority alone do not predict the ordering. These cases follow a systematic [posterior] → [anterior] articulatory progression, reflecting the biomechanical dynamics of speech production. This auxiliary criterion explains otherwise irregular-looking sequences by showing that articulatory movement itself contributes to irreversible ordering; e.g. *kaba saba*, *kanlı canlı*, *kızar bozar*, *kıvrır zıvrır*, etc.

Taken together, these three rules demonstrate that the asymmetry intrinsic to Turkish bicoordinatives originates not from lexical idiosyncrasy, but from universal phonological and articulatory biases governing segment organization. They define the natural “phonological backbone” upon which higher-level morphophonological, morphosyntactic, and iconic constraints subsequently build.

The following sections present these principles in detail, beginning with Rule 1, which establishes the vocalic foundation of phonological asymmetry.

5.1.1 Rule 1 – Vowel sonority hierarchy in Turkish bicoordinatives

In the vowel system of Turkish asymmetric bicoordinatives, ordering tendencies reflect a consistent hierarchy of sonority, governed by degrees of acoustic openness and oral aperture. High vowels, which involve a narrower oral aperture and lower acoustic sonority, tend to follow low vowels whose wider oral aperture and greater sonority facilitate perceptual prominence. This *open-to-close* (*low-to-high*) progression corresponds to the universal vowel sonority scale, in which /a/ [a] shows the highest sonority, followed by /o/ [o], /e/ [e], /ö/ [ø], and /ı/ [ɯ], /i/ [i], /u/ [u], /ü/ [y] at the lowest end. Accordingly, the preferred ordering pattern in Turkish asymmetric bicoordinatives is:

[low / mid vowel] → [high vowel], i.e.
[more sonorous] → [less sonorous]

This pattern is consistent with the widely held view that the sonority hierarchy acts as a fundamental organizing principle in linguistic sequencing (Clements 1990; Parker 2002; Parker 2008). In this respect, Rule 1 forms the vocalic foundation of irreversible ordering and prepares the structural ground for the consonantal and articulatory hierarchies formalised in Rules 2a and 2b.

Rule 1.

The vocalic sonority hierarchy (i.e. *low / mid vowel* → *high vowel*):
the component containing the relatively more sonorous or open vowel
comes first.
[+open] [–open]

Constraint: Rule 1 does not apply when another higher-ranked rule is active, namely:

- a) Rule 4 (vocalic vs. non-vocalic onset) [+vocalic] [–vocalic]: *abuk sabuk, özene bezene*
- b) Rule 6 (open vs. closed syllable) [+open syllabic] [–open syllabic]: *ulu orta, yeme içme*

[+open] [–open]: MTu.: *akıllı uslu, altüst* (< *alt üst*), *az uz, cart curt, cak cuk* (*çiğnemek*), *cambul cumbul, can ciğer, cangıl cungul, çalı çırpı, çanak çömlek, çangıl çungul, çangır çungur, çar çur, çarık çürük, çarpık çurpuk, çayır çimen, dan dun, dangıl dungul, davul zurna, dayamak döşemek, dayalı döşeli, deli dolu, derli toplu, doğru dürüst, efil üfül* (*esmek*), *fakir fukara, falan filan, falan fısıltı, fasa fiso, garç gurç, gelen giden, gelin güvey, gelir gider* (*hesabı*), *gelmek gitmek, hacı hoca, ham hum, kambur kumbur, kara kuru, karı koca, kem küm, langır lungur, namaz(ında) niyaz(ında), oğul uşak, rahat huzur, sağ salım, sağ selamet, sağa sola, sağdan soldan, sağlı sollu, sağda solda, sağı solu, sarmaş dolaş, sayım suyum, sele suya* (*karişmak*), *tek tük, saymak sövmek, saz(ı) söz(ü)* (~ *saz(lı) söz(lü)*), *şangır şungur, tad(ı) tuz(u), taka tuka, takas tukas, tangır tungur, tatlısı tuzlusu, tatsız tuzsuz, tepeden tırnağa, ters tırs, var yok, varı yoğu, vara yoğa* (*üzülmek*), *varsa yoksa, yakmak yıkmak, yalan dolan, yamru yumru, yamuk yumuk, yatak yorgan, yeri yurdu, yer yurt, yeri yurdu, zar zor*.

OTu.: *agır ulug, äzüg igid, kävil- kogşa-, töz tüb*.

5.1.2 Rule 2a – Consonantal sonority hierarchy in Turkish bicoordinatives

The internal ordering of Turkish asymmetric bicoordinatives is not determined by vowels alone. In a large subset of lexicalised formations, the decisive factor is the *relative consonantal sonority* of the two components. Consonantal sonority shows how much a consonant blocks the flow of acoustic energy. Consonants that allow more airflow and have more regular vibration are naturally more sonorous than those that involve tighter constriction. This principle, originating in Jespersen's (1904:186) concept of open articulation and later formalised in phonological sonority theory (Clements 1990; Parker 2002, 2008), predicts that in bicoordinatives the more sonorous consonant tends to occur in the first component, while the second typically contains a less sonorous segment. The Turkish asymmetric bicoordinatives strongly support this universal tendency. Across both Modern Turkish and Old Turkic, bicoordinatives regularly exhibit a descending sonority pattern, whereby the acoustic openness of the initial consonant is greater than that of the following one.

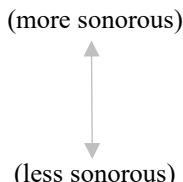
The consonantal hierarchy that emerges from the corpus aligns closely with universal scales proposed in the literature. In Turkish bicoordinatives, the relevant hierarchy can be represented as in (34):

- (34) *rhotic consonants* → *laterals* → *nasals* → *voiced fricatives* → *voiced affricates* → *voiced stops* → *voiceless fricatives* → *voiceless affricates* → *voiceless stops*²⁵

This descending hierarchy captures the preference for the more sonorous consonant to appear in the first component of the bicoordinative, followed by a segment of lower sonority (e.g. *sus pus*, *süs püs*, *soy sop*, *çat pat*, *yırtık pırtık*, etc.; for further examples, see below). In all such formations, the acoustic–articulatory openness decreases from the first member to the second, producing a natural sonority-descending sequence. This tendency indicates that irreversible ordering does not arise from lexical arbitrariness but rather from phonetic and phonological biases favouring sequences that begin with more resonant segments.

While general sonority hierarchies have been discussed since Jespersen (1904), the internal *relative sonority hierarchy of obstruents* had not been clearly established prior to the late 2000s. The ordering proposed here (see ex. (34) and (35a)), originally formulated on the basis of Turkish data in my 2007 MA thesis, coincides exactly with the *relative obstruent hierarchy* later assumed by Parker (2008) on the basis of a broad survey of the phonetic and typological literature (see ex. (35b)). This convergence suggests that the Turkish data reflect a phonetic ordering principle of broader typological relevance:

(35a) relative sonority of obstruents proposed in Turkish	(35b) relative sonority of obstruents proposed by Parker
<i>voiced fricatives</i> <i>voiced affricates</i> <i>voiced stops</i> <i>voiceless fricatives</i> <i>voiceless affricates</i> <i>voiceless stops</i> (Aydemir 2007)	<i>voiced fricatives</i> <i>voiced affricates</i> <i>voiced stops</i> <i>voiceless fricatives</i> <i>voiceless affricates</i> <i>voiceless stops</i> (Parker 2008)



Some bicoordinatives, however, display ordering patterns that cannot be fully accounted for by consonantal sonority alone and require an additional articulatory perspective; e.g. *kaba saba*, *kanlı canlı*, *kızır bozar*, *kıvır zıvır*, etc. (see Rule 2b below). These cases motivate the formulation of a separate articulatory ordering principle, which is developed in Rule 2b below.

²⁵ Aydemir 2007: 56, 58.

5.1.3 Rule 2b – Articulatory progression sequence (auxiliary criterion)

In some bicoordinative pairs (e.g. *tuz buz*, *kaba saba*, *kızar- bozar-*, etc.; for further examples, see below), the ordering appears to run counter to the descending sonority pattern described in Rule 2a. In such cases, the sequence seems to “rise” in sonority rather than descend. Although these bicoordinatives comply with Rule 2b (*posterior* → *anterior* articulatory place sequencing), their initial consonants also display a surface pattern of rising sonority, i.e. an increase in sonority of consonants from left to right ([low sonority] → [higher sonority]), as seen in sequences such as *t–b*, *k–s*, and *k–b*. Thus, these formations do not constitute exceptions to the sonority hierarchy itself; instead, they are governed by an additional principle. This suggests that the natural sound flow in Turkish cannot be explained solely in terms of acoustic-perceptual sonority but must also be considered in relation to the direction of articulatory movement.

Closer inspection shows that some of these pairs follow a systematic progression from *posterior* to *anterior* articulatory positions (e.g. /k/ → /b/ or /k/ → /s/). Although this type of ordering cannot be captured by traditional sonority hierarchies, it is mechanically natural from the perspective of speech production – i.e. the systematic tendency toward *back-to-front* movement of articulators. During articulation, movements of the tongue and other articulators from the back (posterior) toward the front (anterior) of the vocal tract facilitate rhythmic flow, articulatory ease, and coarticulation.

While this directional tendency is not explicitly formulated in the phonetic literature as a “*posterior* → *anterior*” principle, it is compatible with well-known observations that articulatory gestures often proceed from more constricted to less constricted configurations (e.g. Browman & Goldstein 1989; Lindblom 1990).

Therefore, for the patterns observed in this study, the descriptive term *articulatory progression sequence* is proposed. This term differs from classical feature systems (e.g. [dorsal], [coronal], [labial]) in that it characterizes the direction of segmental production, providing an additional, articulation-based measure that complements the phonological interpretation and sonority-driven analysis of bicoordinative structures. This sequence of movement proceeds from less mobile to more mobile articulators, can be summarized as illustrated in (36):

- (36) [laryngeal] [palatal / dental / labial], [palatal] [dental / labial], [dental] [labial], i.e.
tongue root → *tongue body* → *tongue blade* → *tongue front* → *tongue tip* → *lower lip*

It represents the increasing mobility of articulatory organs from the back toward the front of the vocal tract and accounts for the Turkish bicoordinatives exhibiting *rising sonority* as phenomena grounded not in sonority itself but in an articulatory principle of productional naturalness. This articulatory movement from back to front mirrors the natural biomechanical flow

of speech and interacts with the sonority hierarchy to yield a consistent pattern of *descending sonority* combined with *ascending articulatory mobility*.

These two interacting hierarchies — i.e. the *descending sonority hierarchy* (Rule 2a) and the *articulatory progression sequence* (Rule 2b) — together define the phonotactic constraints governing a range of Turkish bicoordinatives. In such cases, *articulatory progression* does not override sonority but interacts with it as an auxiliary ordering principle. While they account for the majority of phonologically regular patterns, other types of coordination obey distinct morpho-semantic or lexical principles discussed in the subsequent rules. In this way, the two hierarchies presented here complement the other rules by introducing the phonotactic dimension of Turkish asymmetric bicoordinative formation.

Accordingly, the tendencies discussed above are formalised in Rules 2a and 2b as follows:

Rule 2a

The consonantal sonority hierarchy (sonority-descending)

rhotic consonants → *laterals* → *nasals* → *voiced fricatives* → *voiced affricates* → *voiced stops* → *voiceless fricatives* → *voiceless affricates* → *voiceless stops*

[fricative] [affricate], [fricative] [stop], [affricate] [stop]

[+sonorant] [–sonorant]

Rule 2b

Relative articulatory mobility hierarchy by place of articulation
(auxiliary criterion)²⁶

tongue root → *tongue body* → *tongue blade* → *tongue front* →
tongue tip → *lower lip*

[laryngeal] [palatal / dental / labial], [palatal] [dental / labial],
[dental] [labial]

[+consonant] [+consonant]:²⁷ MTu. *bağıra çağıra, bağır- çağır-, bağırış çağırış, bakkal çakkal, başı sonu, bata çıka, belli başlı, belli besli, boy pos (/ bos), boylu poslu, bölük pörçük, büyük küçük, börek çörek, canla başla, cıncık*

²⁶ This sequence reflects increasing articulatory mobility from posterior to anterior positions and helps account for bicoordinatives exhibiting surface patterns of rising sonority. The concept of a *relative mobility hierarchy of articulators* is used here as a conceptual tool referring to a phonologically and phonetically motivated phenomenon, rather than to an absolute or universal articulatory scale. This notion aligns with observations in approaches emphasizing *articulatory ease* and *gestural phonology*, and was previously introduced and discussed in earlier work by the author (Aydemir 2007) under the term “*Relative Beweglichkeitshierarchie der Artikulationsorgane (Artikulatoren)*.” The qualifier *relative* therefore underscores its function as a comparative ordering preference specific to coordinative structures. Within this framework, certain ordering preferences in Turkish bicoordinatives reflect a relative articulatory progression from posterior to anterior positions, which is compatible with reduced articulatory cost in sequential production.

²⁷ This feature indicates that the bicoordinatives in this group begin with a consonantal onset.

boncuk, cici bici, cicili bicili, çala çırpa, çalı çırpı, çalmak çırpmak, çalışmak çabalamak, çanak çömlek, çarşı pazar, çat pat, çatal bıçak, çıtı pıtı, çoluk çocuk, dağ taş, dala çıka, dere tepe, derme çatma, dertsiz tasasız, döl döş, düğün bayram (etmek), gezmek tozmak, girdi çıktı, girinti(lı) çıkıntı(lı), giyim kuşam, giyinmek kuşanmak, gizli kapaklı, gizli saklı, hesap kitap, gözüne dizine (durmak), halim selim, halis muhlis, hali vakti, hamal camal (~ hammal cammal), haraç mezat, kaba saba, kader kısmet, kadir kıymet, kan ter, kanlı canlı, kanı canı (kalmamak), kanlı canlı, kapı baca, karda kışta, karda buzda, kargacık burgacık, karman çorman, kızarmak bozarmak, kona göçe, konar göçer, kurt kuş, küçük(lü) büyük(lü), kıvrır zıvrır, saç baş, sarmaş dolaş, senet sepet, senetsiz sepetsiz, senli benli, sere serpe, sıkı fıkı, soy(lu) sop(lu), soyu sopu, sus pus, süklüm püklüm, süs püs, süslü püslü, şöyle böyle, şu bu, şundan bundan, şura bura, şurada burada, şuralı buralı, yağmur çamur, yaka paça, yalan dolan, yaş baş, yaşlı başlı, yata kalka, yatak yorgan, yaz boz, yaz kış, yazar bozar, yılan çıyan, yırtık pırtık.

OTu. aşnukılı amtkılı, barça tükāl, buşuşlık korkınçlık, çiziz tartız, çom- bat-, karlıg buzlug, kaçig köprüg, kalış barış, kâzig tizig, küşüşlüg tilâklig, sävinçli korkınçlı, sı- buz-, sıg tümgä, sıg yuka, sığun muygak, yaruk yaşuk.

5.1.4 Closing remarks to Rules 1, 2a, and 2b

Together, Rules 1, 2a, and 2b establish the phonological foundation of asymmetry in Turkish bicoordinatives. They show that the irreversible ordering observed in these constructions is not accidental or lexically idiosyncratic, but arises from the interaction of vowel sonority, consonantal sonority, and articulatory mechanics. These factors operate at different but complementary levels, shaping ordering preferences that tend to proceed from higher to lower acoustic openness and, independently, from posterior to anterior positions in articulatory space.

Taken together, these interacting hierarchies define a coherent phonological architecture that constrains the formation of asymmetric bicoordinatives in Turkish. This architecture does not by itself account for all aspects of bicoordinative structure, but it provides the necessary phonetic and phonological grounding upon which higher-level asymmetries are built. The following set of morphophonological and morphosyntactic constraints (Rules 3–6) extend this foundation into the structural domain, where phonetic and rhythmic asymmetries become conventionalized through patterns of syllable structure, onset configuration, and morphological alignment.

5.1.5 From phonological asymmetry to structural asymmetry

The phonological principles outlined in Rules 1, 2a, and 2b establish the foundational sound-based conditions — such as vowel and consonantal sonority as well as articulatory directionality — that give rise to asymmetric bicoordinatives in Turkish. Yet the internal organization of these formations cannot be derived from phonological factors alone. Once phonological asymmetry is in place, additional constraints come into effect — constraints that pertain to syllable shape, segmental arrangement, and morphosyntactic align-

ment. In other words, phonological asymmetry constitutes only the initial layer of a more complex system, one whose full structure becomes visible only when higher-level morphophonological and morphosyntactic factors are taken into account. The following rules therefore move beyond the phonological domain and address the morphophonological and morphosyntactic factors that systematically contribute to irreversible ordering.

5.2 Morphophonological and morphosyntactic constraints (Rules 3-6)

Rules 3–6 specify the structural mechanisms through which the phonological patterns identified above become integrated into stable coordinative templates. While Rules 1, 2a, and 2b explain asymmetry in terms of sonority and articulatory directionality, the present section focuses on how segmental composition, syllable configuration, and morphosyntactic organization further shape irreversible ordering. Together, these constraints chart the transition from purely phonological tendencies to fully structuralized patterns characteristic of Turkish bicoordinatives.

5.2.1 Rule 3 – Vowel backness and ordering in bicoordinatives

Rule 3 operates at the level of segmental composition and concerns vowel quality rather than consonantal sonority or articulatory directionality. Unlike Rules 2a and 2b, which are sensitive to consonantal properties, Rule 3 captures an ordering preference based on vowel backness when syllable count is the same. The rule thus introduces a vocalic dimension of asymmetry that interacts with, but is hierarchically subordinate to, the phonological constraints discussed earlier.

Rule 3

With equal syllable count, the component containing a back vowel comes first.
[+back] [–back]

Constraint: Rule 3 does not apply when another higher-ranked rule is active, namely:

- a) Rule 1 (vowel openness) [+open] [–open]: *deli dolu*, *sele suya*,
- b) Rule 2a (sonorant vs. non-sonorant onset) [+sonorant] [–sonorant]: *belli başlı*, *dere tepe*
- c) Rule 4 (vocalic vs. non-vocalic onset) [+vocalic] [–vocalic]: *üst baş*, *iri yarı*, *ipe sapa*
- d) Rule 5 (unequal syllable count) [+fewer syllables] [–fewer syllables]: *el ayak*, *göz kulak*
- e) Rule 6 (open vs. closed syllable) [+open syllabic] [–open syllabic]: *pılı pırtı*, *ulu orta*

[+back] [–back]: MTu. *az öz*, *baş göz (etmek)*, *baş(ı) beyn(i)*, *çarık çürük* (~ *çürük çarık*), *çanak çömlek*, *çayır çimen*, *hayal meyal*, *karışanı görüşeni*, *kaş göz (etmek)*, *kaşı gözü*, *kazma kürek*, *kırık dökük*, *kırıla döküle*, *kırılanı*

döküleni, kıyı(da) köşe(de), mal mülk, paldır küldür, pat küt, yara bere, yazar çizer, yaz- çiz-

OTu. *açın- igid-, aşıl- üstäl-, bıçık tikig, tañdalığ keçälig.*

5.2.2 Rule 4 – Vocalic vs. non-vocalic onset

This rule captures a systematic preference for vowel-initial components in coordinative ordering when syllable number is the same.

Rule 4

With equal syllable count, the component with a vocalic onset comes first.
[+vocalic] [–vocalic]

Constraint: Rule 4 does not apply when another higher-ranked rule is active, namely:

- a) Rule 5 (unequal syllable count) [+fewer syllables] [–fewer syllables]:
aç susuz, ardı arkası

[+vocalic] [–vocalic]: MTu. *abur cubur, abuk sabuk (/ subuk), aç tok, açık saçık, açık seçik, açıl- saçıl-, adı sanı, ağrı sızı, ah vah, ahım şahım, akar kokar, ak(ça) pak(ça), akıl fikir, aksır- tıksır-, alaca bulaca, alacak verecek, alavere dalavere, alıcı verici, alık salık, alım satım, allak bullak, allem kallem, allı pullu, alla- pulla-, aman zaman, ana baba, analı babalı, analı danalı, anasız babasız, anlı şanlı, apar topar, abdestinde namazında, ara sıra, arasor-, arama tarama, ara- tara-, aslı fashı, aslı nesli, as- kes-, aşağı yukarı, aşna fişne, atılmaz satılmaz, atik tetik, ayan beyan, ayda yılda, aygın baygın, ayıla bayıla, ayıl- bayıl-, ayrı gayrı, az buz, az çok, ecir sabır (dilemek), eciş bücüş, eften püften, eğri büğrü, eğri doğru, eksik gedik, el(e) gün(e), el(i) kol(u), el(i) yüz(ü), eni konu, enine boyuna, eninde (/ önünde) sonunda, er geç, esele-besele-, eski püskü, estek köstek, eş dost, et(i) bud(u), etli butlu, etli(ye) sütlü(ye), etme bulma, etten kemikten, etme eyleme, ev bark, evele- gevele-, evir- çevir-, evire çevire (dövmek), evli barklı, ezik büzük, ezil- büzüül-, ezile büzüle, ıcığını cıcığını (çıkarmak), ıkıl- sıkıl-, ıkın- sıkın-, ıkına sıkına, iklim tıklım, ıvrır zıvrır, içi dışı, içli dışlı, ikide birde (~ ikide bir), ileri geri, in cin, incik boncuk, inişli çıkışlı, ipe sapa (gelmemek), ipsiz sapsız, iri yarı, ismi cismi, iyi kötü, ofla- pufla-, okuma yazma, okur yazar, ondan bundan, onu bunu, onun bunun, orada burada, oram buram, öcü böcü, ölç- biç-, ölüm kalım, ölüsü dirisi, öteberi (< öte beri), öteki beriki, ötesine berisine, öyle böyle, özene bezene, özen- bezen-, ufacık tefecik, ufak tefek, ufla- pufla-, üç beş, üçe beşe, utana sıkıla, üfle- püfle-, üfleye püfleye, üfüre püfüre, üst baş.*

OTu. *açıl- yadıl-, alku barça, arığ simäk, arığ turuk, arkurı turkaru, asığ tusu, at kü, atlıg küllüg, ayag çiltäg, ayal- kötrül-, ägsük käreğäk, äñ- çın-, äñinçig tañlançığ, ärän kırkın, äsin bulıt, eligläär bäglär, ilış- tartış-, ödräk küvüz, ög kañ, ögir- sävin-, ögrünç sävinç, öl şı, uçuz yenik, uzun kızka, ülgü kolu, üküş tälüm.*

5.2.3 Rule 5 – Syllable count asymmetry

This rule captures a systematic preference for the component with fewer syllables in coordinative ordering when the two components differ in syllable count. Rule 5 operates at a higher structural level than the phonological constraints discussed in Rules 1, 2a, and 2b, and renders inoperative Rules 3 and 4 whenever syllable count is unequal. In such cases, differences in syllable number take precedence over vowel backness (Rule 3) and onset type (Rule 4), reflecting the strong organizing role of syllabic structure in the stabilization of Turkish bicoordinatives. Thus, segmental preferences (e.g. vowel backness or vocalic onset) become relevant only when the syllabic structure of the two components is balanced.

The abundance and stability of Old Turkic evidence suggest that, like the other rules discussed here, Rule 5 is not a late or peripheral development but has operated consistently since the earliest attested stages of Turkish. Although less frequently discussed than vowel harmony, syllable-count asymmetry appears to form part of a deeper organizing layer of the language, interacting systematically with segmental constraints such as vowel backness and onset type. This constellation of interacting principles points to a phonological system in which harmonic, segmental, and syllabic structures are mutually reinforcing rather than independent, a pattern that may be of particular interest for research at the phonetics–phonology interface.

As illustrated in the following examples, the rules are not mutually exclusive but hierarchically ordered. Rule 5 is ranked above Rule 4 and applies only when syllable count asymmetry is present; when syllable counts are equal, ordering is determined by lower-ranked rules such as Rule 4.

Rule 5

With equal syllable count, the component with fewer syllables comes first.
[+fewer syllables] [–fewer syllables]

Constraint: Rule 5 does not apply when another higher-ranked rule is active, namely:

- a) Rule 1 (vowel openness) [+open] [–open]: *akıllı uslu*
- b) Rule 4 (vocalic vs. non-vocalic onset) [+vocalic] [–vocalic]: *ileri geri, ikide bir*

[+fewer syllables] [–fewer syllables]: MTu. *acı tatlı, aç susuz, al aşağı (olmak / etmek), alet edevat, ar namus, ardı arkası, aslı astarı, az buçuk, bağ bahçe, belli belirsiz, bet beniz, bet bereket, bı- usan-, birlik beraberlik, bit- tüken-, bitmez tükenmez, bul- buluştur-, can ciğer, çift çubuk, çul çaput, dağ bayır (dolaşmak), dal budak, deli divane, dertsiz tasasız, dili damağı (kurumak), din iman, dip bucak, dip doruk, dirlik düzenlik, dışından turnağından (artırmak), don gömlek, don paça, döne dolaşa, dön- dolaş-, dur durak (yok), durmuş oturmuş, dur otur (yok), el alem, el avuç, el ayak, el ense, el etek, el pençe (divan durmak), fakir fukara, falan festekiz, falan feşmekan, gel gelelim, gelin görümce, genci ihtiyarı, gizli kapaklı, görmüş geçirmiş, göz kulak (olmak),*

güç(lü) kuvvet(li), güle oynaya, güllük gülistanlık, günlük güneşlik, haddi hesabı, hak hukuk, hal hatır, han hamam, hanım hanımcık, hısım akraba, ister istemez, it köpek, kan revan, karga tulumba, kapacak, kas- kavur-, kılık kıyafet, kırdı bayırda, kış kıyamet, kıt kanaat, kızı kırsığı, kız kızan, kol kanat, kör topal, kör kütük, kul köle, kul kurban, kulu köpeği (olmak), od ocak, palas pandıras, saç sakal, saç- savur-, sap saman, sar- sarmala-, sesi soluğu, ses seda, sessiz sedasız, seve oynaya, sildi süpürdü, sor- soruştur-, söz sohbet, suya sabuna (dokunmak), sür- sürüştür-, şen şakrak, takım taklavat, tak-takıştır-, taşı tarağı, taşı toprağı, telli duvaklı, top tüfek, toz duman, toz toprak, tuz biber, un ufak, ucu bucağı, ucu ortası, uçsuz bucaksız, uça bucakta, ver-veriştir-, yarı buçuk, yarım yamalak, yerli yabancı, yol yolak, yol yordam, yol yöntem, zehir zemberek, zil zurna.

OTu. *aç okı, aş azuk, aş- üklit-, aşı üstüyü, asdaşı üstädäşi, aya- ağırla-, el uluş, ençgü äsängü, bür budık, çın kertü, çog yalın, çöb evdängü, hwa çäçäk, hua yavişgu, ı ıgaç, irü belgü, iş küdük, iz oruk, ka kadaş, keng alkıg, kog kıçmık, kol- ötün-, kork- titrä-, köl öläng, körk mäniz, köz kulkak, kurt koñuz, midik pirtigçan, münlüg kadaglıg, nom şazın, oñşatıg yöläşürüg (~ yöläşürüg oñşatıg), suk- tündä-, şın süngük, tärkin tavrakın, tıd- bäklä- toyın şınnanç, toz toprak, tüş yemiş, tüşlig utlılıg, uç kızıg, uçsuz kızıgsız, ula- üklit-, ulug kıçig, ürk- bälinglä-, ürküt- bälinglä-, yän ätäk, yeg kodıki, yeg üstünki, yer- münä-, yol oruk, yu- arıt-, yul yulak.*

5.2.4 Rule 6 – Open vs. closed syllable structure

Rule 6 addresses asymmetry at the level of syllable structure by treating the open vs. closed syllable contrast as an ordering principle. While Rules 3–5 operate on segmental composition and syllable count, the present rule targets the internal phonotactic configuration of the syllable itself. When higher-ranked phonological, morphosyntactic, or iconic constraints do not apply, Turkish bicoordinatives show a preference for components containing open syllables to precede those with closed syllables, reflecting a tendency toward syllabic structure simplicity and rhythmic balance. Here, the open vs. closed syllable contrast is understood in relative terms: the ordering reflects a comparison between the overall syllabic structure of the two components, rather than the presence of any particular open or closed syllable within individual words.

Rule 6

With equal syllable count, the component containing an open syllable comes first.

[+open syllable] [–open syllable]

Constraint:

Rule 6 applies only when ordering is not already determined by Rules 1–5 or by iconic constraints (Rule 7).

[+open syllable] [–open syllable]: MTu. *ağır aksak, ana oğul, baba oğul, balta nacak, börtü böcek, çeki düzen, doğru dürüst, dorma dolaş, düğün dernek, gece gündüz, gözü gönü (açılmak), hırlı hırsız, ilim irfan, irili ufaklı, iğne iplik, kazma kürek, kelle kulak, konu komşu, köşe bucak, olur olmaz, pılı pırtı, rica minnet, saçma sapan, sere serpe, sille tokat, tarla tapan, tekme tokat, ulu orta, utanmaz arlanmaz, yalan yanlış, yeme içme.*

OTu. *arta- aşşın-, isig amrak, kaçā küntüz, kurgu yenik, tünlā küntüz, tözün kövşāk, uyat- äymän-, yaru- yaltri-, yaruk yaltrik, yarut- yaltrit-, yerük ägsük.*

5.3 Rule 7 – Iconic and event-based ordering (iconic constraints)

While Rules 1–6 account for asymmetric ordering in Turkish bicoordinatives on phonological, morphophonological, and morphosyntactic grounds, a further major class of constructions is governed by *iconicity*, that is, by semantic and cognitive principles reflecting the linear representation of real-world events. In such cases, the ordering of components mirrors a conceptual sequence grounded in temporal, causal, or scalar relations rather than in segmental or prosodic structure.

In iconic bicoordinatives, the first component represents a state, action, or value that is *conceptually anterior* to the second. This study captures this principle by means of the feature notation [+anterior] → [–anterior], where anteriority is understood in a broad sense, encompassing temporal precedence, causal priority, or scalar progression.

Typical examples include MTu. *bugün yarın* ‘soon, any day now’ ← [‘today’] [‘tomorrow’], *er geç* ‘sooner or later’ ← [early] [late], *üç beş* ‘a small number, a few’ ← [three] [five], as well as OTu. *korkmak titrämäk* ‘to be very afraid’ ← [to fear] [to tremble] and *yu- arit-* ‘to wash and cleanse’ ← [to wash] [to purify]. In all such cases, the fixed order reflects a real-world progression: what occurs earlier, is causally prior, or represents a lower point on a scale precedes what follows.

Reversing the order of these lexicalized bicoordinatives (**yarın bugün, *geç er, beş üç*, etc.)²⁸ results in markedness and semantic degradation, demonstrating that their ordering is not arbitrary but anchored in conceptual structure. The constructions are therefore irreversible not because of phonological well-formedness, but because the reversed sequence no longer aligns with the underlying perceptual or cognitive logic.

In earlier work (2007), I described such constructions as *bikonditionelles Koordinativ* (‘biconditional coordinative’) in order to capture the fact that the two components stand in a logically dependent semantic relation. In the present study, however, I adopt the term *iconic bicoordinative*, aligning the analysis with the broader typological notion of *iconicity*, according to which linguistic form tends to reflect conceptual and experiential structure. This terminological shift preserves the original insight while facilitating comparison

²⁸ Note that *beş üç* remains possible in Turkish, but only as a literal numeral sequence (e.g. a 5–3 score), rather than as a lexicalized or structurally fixed bicoordinative.

with cross-linguistic research on iconic ordering in coordination and binomials.²⁹

Iconic bicoordinatives may thus express, among others:

- a) anterior–posterior relations (e.g. MTu. *eninde sonunda, er geç*),
- b) causal relations (e.g. MTu. *sararmak solmak*, OTu. *korkmak titrämäk*),
- c) scalar progressions (e.g. MTu. *az çok, küçük büyük, üç beş*).

Although these relations differ in surface semantics, they share a directional conceptual structure grounded in anteriority. For this reason, they are uniformly represented here as [+anterior] → [–anterior].

As the ordering in these constructions is determined by conceptual event structure, iconicity constitutes an overriding constraint that is not subordinated to the phonological or morphosyntactic principles described in Rules 1–6.

Rule 7

Iconically motivated obligatory ordering (event hierarchy):³⁰ The component denoting the anterior state or event precedes the posterior one.
[+anterior] [–anterior]

Constraint: none

[+anterior] [–anterior]: MTu. *bugün yarın* ‘in the next few days; sooner or later’ ← [today] [tomorrow], *eninde sonunda* ‘sooner or later; in any case’ ← [before] [after], *kapkaç* ‘snatch theft’ ← *kap kaç!* ← [grab!] [run away!], *sabah akşam* ‘constantly; all the time’ ← [morning] [evening], *sararmak solmak* ‘to become increasingly pale / to fade’ ← [to turn yellow] [to wither]), *üç beş* ‘a few; a small number’ ← [three] [five], etc.

OTu. *açıl- yadıl-* ‘to spread, to expand’ ← [to open] [to spread], *korkmak titrämäk* ‘to be very afraid’ ← [to fear] [to tremble], *yu- arıt-* ‘to wash and cleanse’, *yüğürmäk kaçmak* ‘to flee, to run away’ ← [to run] [to flee], etc.

5.3.1 Bicoordinative compounds: compounding under iconic constraints

A subset of iconic bicoordinatives undergoes further structural reanalysis through compounding and develops into compound nouns in Turkish. The following examples (see (37)), among many others, illustrate how an originally coordinative structure is reanalysed as a single lexical unit through compounding:

- (37) *alaşağı* ‘overthrow’ ← *al aşağı* [bottom] [down],
öteberi ‘various things’ ← *öte beri* [the other side] [this side],

²⁹ This use of iconic ordering is consistent with well-established functional and typological accounts of form–meaning mapping (e.g. Givón 1985, Haiman 1985), though a full discussion of iconicity theory lies beyond the scope of the present study.

³⁰ This corresponds to what I earlier termed (2007) “*Logisch-semantisch orientierte obligatorische Wortfolge*” or “*Inhaltshierarchie*”.

elalem ‘(colloquial) folk’ ← *el alem* [(other) people] [world],
alsat ‘buying an item and selling it quickly’ ← *al sat!* ← [buy!] [sell!],
çekyat ‘sofa bed’ ← *çek yat!* ← [pull out!] [lie down!],
gelgit ‘ebb and tide; going back and forth for nothing’ ← *gel git!* [come!] [go!],
kapkaç ‘snatch theft’ ← *kap kaç!* ← [grab!] [run away!],
oldubitti ‘fait accompli’ ← *oldu bitti* [become-PST] [finish-PST]
vurkaç ‘hit and run’ ← *vur kaç!* ← [hit!] [run!].

Conceptually, these formations preserve the same iconic anterior–posterior structure: the first component denotes the initiating action or state, while the second denotes the resulting action, outcome, or social frame. What changes is their grammatical status: coordination gives way to compounding, while the construction remains morphologically transparent and largely semantically compositional, preserving the core iconic relation between its components, even though some degree of semantic abstraction or conventionalization may occur.

In several cases (e.g. *alaşağı* < *al aşağı et-* ‘take down, overturn’; *çekyat* < *çekyat kaneye* ‘sofa bed’), the apparent semantic shift reflects not a loss of iconic meaning but an ellipsis-based reanalysis; i.e. originally bicoordinative structures are condensed into elliptical nominal forms. In *alaşağı*, for instance, the compound preserves an underlying directional bicoordinative structure (*al* ‘down/below’ + *aşağı* ‘down’), which is morphologically condensed rather than semantically reinterpreted.³¹

This development raises questions of broader typological relevance, particularly with regard to the interaction between iconic sequencing, compounding, and category change. However, a systematic comparison of such formations lies beyond the scope of the present study.³²

While iconic sequencing is widely attested cross-linguistically, the systematic development of compounds from bicoordinative structures under iconic constraints appears to be particularly characteristic of Turkish.

6. Conclusion

This study has examined the internal ordering principles of irreversible nominal and verbal bicoordinatives in Turkish and has shown that their apparent rigidity is not accidental or lexically idiosyncratic. Instead, irreversible ordering emerges from a layered system of interacting constraints

³¹ Note that several bicoordinatives preserve archaic or otherwise non-productive lexical items that no longer function as independent lexemes in Modern Turkish but survive due to lexicalized coordinative templates. A clear example is *alaşağı*, where *al* represents an old directional noun meaning ‘down, below, bottom’ (cf. Hungarian *al* ‘sub, nether, under-’, showing a close formal and semantic parallel), also attested in formations such as *alt* (< *al* + *-t*). Similarly, in *er geç* ‘sooner or later’, *er* is no longer used as a free lexical item in Modern Turkish but is preserved exclusively within the bicoordinative construction. Such cases illustrate that bicoordinatives not only encode iconic ordering but also act as conservative morphological domains that retain archaic elements through formulaic stabilization.

³² The present discussion focuses specifically on compounds derived from bicoordinative structures under iconic constraints (here termed *iconic compounds*). For general treatments of compounding in Turkish, see van Schaik (2002) and Göksel & Haznedar (2007).

operating at phonological, morphophonological, morphosyntactic, and iconic levels.

At the most basic level, phonological asymmetry is established through vowel and consonant sonority relations and articulatory directionality (Rules 1, 2a, and 2b). These principles define preferred sound sequences that favour descending sonority or, in specific cases, rising sonority driven by posterior-to-anterior articulatory progression. Together, they constitute the phonetic–phonological foundation upon which irreversible ordering is built.

Building on this foundation, Rules 3–6 demonstrate how phonological tendencies become structurally stabilized. Segmental composition, onset type, syllable structure, and syllable count interact to yield robust ordering preferences that override lower-level phonological factors when necessary. In particular, syllable count (Rule 5) emerges as a highly stable organizing principle, operative from Old Turkic to Modern Turkish, and capable of neutralizing segmental preferences such as vowel backness or vocalic onset when the balance of syllable structure is disrupted. These findings suggest that structural asymmetries play a central role in the structural stabilization of bicoordinative patterns in Turkish.

Rule 7 extends the analysis beyond form-internal constraints by demonstrating the role of iconic sequencing. In a subset of bicoordinatives, ordering reflects conceptual relations such as temporal succession, causality, or resultative structure, aligning linear form with event structure. These iconic patterns not only motivate irreversible coordination but also provide a bridge to compounding, where coordinative constructions may be reanalysed as single lexical units while often retaining morphological transparency and semantic compositionality.

Taken together, the seven rules outline an integrated architecture of asymmetry in Turkish bicoordinatives. Rather than invoking a single explanatory factor, the study shows that irreversible ordering arises from the cumulative interaction of multiple constraint types, hierarchically organized but mutually reinforcing. This architecture captures both the synchronic regularities of Modern Turkish and their diachronic continuity with Old Turkic, highlighting the depth and stability of coordinative ordering principles in the language.

More broadly, the findings contribute to typological discussions of coordination, phonological asymmetry, and iconicity by demonstrating how language-specific data can reveal general principles of ordering that are neither purely phonological nor purely semantic. Turkish bicoordinatives thus offer a particularly clear window into the interaction between sound structure, constructional organization, and meaning in the emergence of irreversible constructions. Additionally, while iconic sequencing, coordination, and compounding are each widely attested cross-linguistically, Turkish is typologically distinctive in that these three domains converge in a systematic and structurally integrated way within the same class of constructions.

Appendix

A Part-of-Speech–Based Classification of Irreversible Bicoordinatives in Modern Turkish and Old Turkic

This appendix presents the original part-of-speech–based classification of irreversible bicoordinative constructions in Modern Turkish and Old Turkic, covering both binomial and biverbal formations. First developed in my German-language master’s thesis, the taxonomy is reproduced here with only minor adjustments, as it continues to provide an empirically robust and typologically coherent structural framework. The aim of this appendix is strictly classificatory: it organizes the attested bicoordinatives according to their lexical category (noun, adjective, verb, pronoun, numeral, etc.), without proposing semantic or functional generalizations. For reasons of space, only one or two representative examples from Old Turkic (OTu.) and Modern Turkish (MTu.) are given for each subtype, together with their lexical meanings. These meanings help sharpen the semantic contours of each structural pattern and assist readers in interpreting the morphophonological and morphosyntactic types discussed in the main text. This POS-based taxonomy is intended as a comprehensive empirical overview; deeper semantic, pragmatic, or discourse-functional distinctions lie outside the scope of the present constraint-oriented analysis. For additional examples, readers may consult the compilations of Old Turkic and Modern Turkish bicoordinatives.³³

1. Bicoordinatives

1.1. Nominal bicoordinatives (NBi)

1.1.1. Nouns:

- a) Kinship nominals: OTu. *ata ana* ‘parents’; MTu. *ana baba* ‘id.’.
- b) Inanimate nominals: OTu. *öl şı* ‘humidity’; MTu. *kap kacak* ‘pots and pans’.
- c) Oppositional nominals: OTu. *az üküş* ‘more or less’; MTu. *az çok* ‘id.’.
- d) Synonymic nominals: OTu. *asıg tusu* ‘benefit’; MTu. *delik deşik* ‘full of holes’.
- e) Quantifying nominals:
 - Totality: OTu. *alku kamağ* ‘all; the whole’.
 - Multiplicative: OTu. *üküş tälīm* ‘numerous; a great many’; MTu. *onlarca yüzlerce* ‘tens or hundreds; a great many’.

1.1.2. Adjectives: OTu. *uçsız kırıksız* ‘unlimited’; MTu. *uçsuz bucaksız* ‘id.’.

1.1.3. Adverbs: OTu. *keçä küntüz* ‘round the clock’; MTu. *gece gündüz* ‘id.’.

1.1.4. Pronouns:

- a) Personal pronouns: none.
- b) Reflexive pronouns: OTu. *öz öz* ‘each his own’, *käntü käntü* ‘id.’.

³³ Ölmez 1997, 1998; Şen 2003; Ölmez 2017, 2022; Özkan 2019; Akyağın 2024; Karaman 2022; Şahin & Özbek 2025.

c) Demonstrative pronouns: OTu. *ol ol* ‘whichever’; MTu. *şu bu* ‘so-and-so’.

d) Interrogative-indefinite pronouns: OTu. *kim kayu* ‘whoever’, *kim kim* ‘id.’.

1.1.5. Numerals:

a) Cardinals: OTu. (inscriptional) *eki üç* ‘two-three’; MTu. *beş on* ‘around five to ten’.

b) Distributives: OTu. *bir bir* ‘one by one’; MTu. *bir bir* ‘id.’.

1.1.6. Interjections: MTu. *uf puf* ‘huff and puff’, *of of* ‘oh dear; alas’.

1.1.7. Onomatopoeia: MTu. *patır kütür* ‘crash-bang’, *tangır tungur* ‘rattle; clatter’.

1.1.8. Reduplication and reduplicative forms (Nominal Reduplication):

1.1.8.1. Morpheme- or word-level reduplication:³⁴

1.1.8.1.1. Nouns: OTu. *ew ew* ‘from house to house’, *maş maş* ‘step-by-step’; MTu. *yer yer* ‘locally; occasionally’, *adım adım* ‘step-by-step’.

1.1.8.1.2. Adjectives: OTu. *öñi öñi* ‘multifarious’, *başka başka* ‘id.’.

1.1.8.1.3. Adverbs: OTu. *yana yana* ‘again and again’, MTu. *tekrar tekrar* ‘id.’.

1.1.8.1.4. Pronouns:

a) Personal pronouns: none.

b) Reflexive pronouns: OTu. *öz öz* ‘each his own’, *kántü kántü* ‘each his own’.

c) Demonstrative pronouns: OTu. *ol ol* ‘whichever’, *anta anta* ‘here and there’.

d) Interrogative-indefinite pronouns: OTu. *näçä näçä* ‘however many, any number of times’; MTu. *nice nice* ‘a great many’.

1.1.8.1.5. Numerals:

a) Cardinals: OTu. *bir bir* ‘one by one’; MTu. *bir bir* ‘id.’.³⁵

b) Distributives: MTu. *birer birer* ‘one by one’.

1.1.8.1.6. Onomatopoeia:

a) Morpheme reduplication: OTu. (UygW) *kıs kus* ‘smacking sound’; MTu. *cik cik* ‘tweet tweet’.

b) With vowel addition (Paragoge): MTu. *gargara* ‘gargling’

1.1.8.2. m-Reduplication: MTu. *kitap mitap* ‘book or such things’; *at mat* ‘horse and such’.

1.1.8.3. Figura etymologica: MTu. *inim inim inle-* ‘to moan continuously’; *horul horul horla-* ‘to snore loudly’; *oylum oylum oy-*, a colloquial, metaphorical, and often half-joking threat, meaning

³⁴ Note that subsections 1.1.8.1.1–1.1.8.1.6 classify *reduplicated* forms by the part-of-speech (POS) of their base; hence numerals, pronouns, nouns, etc. reappear here.

³⁵ Note that some reduplicated formations (e.g. *bir bir*) exhibit a mismatch between their part-of-speech (POS) classification and their semantic function. In such cases, the items are listed here according to the POS of the base form (e.g. *bir* as a cardinal numeral), while their functional interpretation (e.g. distributive ‘one by one’) is indicated separately. This organization reflects the fact that the classification in this appendix is fundamentally POS-based rather than function-based.

roughly ‘to carve somebody up / to tear somebody to pieces’, and similar figurative expressions of threat; etc.

1.2. Verbal Bicoordinatives (VBi)

Note that the label *Verbal Bicoordinatives (VBi)* is used here in a part-of-speech sense. It designates bicoordinative constructions whose base forms are *verbal in origin*. This category is therefore broader than the narrow functional notion of *biverbals*, which requires the coordination of two finite predicates. Participial, infinitival, and converbial formations included below (see §1.2.2.) are thus *verbal-origin bicoordinatives*, not finite biverbals.

1.2.1. Finite verbs

The finite subtypes listed here constitute *true biverbals*, i.e. coordinations of two inflected verb forms. Non-finite verbal formations, though included in §1.2 due to their verbal origin, do not belong to the biverbal class proper.

1.2.1.1. Indicative:

- a) **Perfect (indefinite / evidential) -mXş:** MTu. *yemiş içmiş* ‘he/she ate and drank’ (cf. adjectival: *gelmiş geçmiş* ‘ever, of all time’).
- b) **Perfect (definite) -DX:** OTu. *bilti uktı* ‘he/she understood’; MTu. *düşündü taşındı* ‘he/she thought and thought again’.
- c) **Present -yor:** —
- d) **Aorist -(X)r, -(A)r, (Neg.) -mAz:** OTu. *ayayur agırlayur* ‘to worship’; MTu. *(ara sıra) gelir gider* ‘comes and goes occasionally’ (cf. binomial: *gelir gider* ‘income and expense’).
- e) **Future -DAçI, -gAy, -(y)AcAk:** —
- f) **Combined Indicative (Mixed TAM Bicoordinatives):**³⁶ A subtype combining two different indicative TAM categories within the same biverbal sequence: MTu. *oldu olacak* ‘on the verge of happening’ (perfect + future) → ‘in/on the cards’, *bitti bitecek* ‘almost finished’ (perfect + future), *gitti gidiyor* ‘on the verge of leaving/going’ (perfect + progressive).

1.2.1.2. Mood:

- a) **Conditional -(y)sA:** MTu. *varsa yoksa* (< *var ise yok ise*) ‘only; nothing but’, *olsa olsa* ‘at most’.
- b) **Optative:** —
- c) **Imperative:** OTu. *yarızun yaltrızun* ‘Let it shine and sparkle!’, *örtbas* ‘cover-up, delitescence’ (< *ört bas*) ← [cover-IMP.2SG] [press-IMP.2SG], *gelgelelim* ‘however, nevertheless’ (< *gel gelelim*) ← [come-IMP.2SG] [come-OPT.1PL] (*lit. come-come-*

³⁶ Mixed TAM refers to the combination of two different tense–aspect–mood (TAM) categories within a single bicoordinative predicate sequence.

let-us), *gelin görün* ‘however’ ← [come-IMP.2PL] [see-IMP.2PL].

1.2.2. Non-finite verbs

1.2.2.1. Infinitive -*mAk*: OTu. *ulamak üklitmäk* ‘to add and multiply’, *ögirmäk sävinmäk* ‘joyousness’; MTu. *bıkmak usanmak* ‘have had fill of something’.

1.2.2.2. Participle:

a) Perfect participle -*DX*, -*mXş* (< -*mIş*), -*dXk* (< -*dOk*):

OTu. *ötgürmiş topulmuş* ‘having penetrated’; MTu. *oldubitti* (< *oldu bitti*) ‘fait accompli’, *girdi çıktı* ‘input-output’, *gelmiş geçmiş* ‘ever, of all time’, *olmuş bitmiş* ‘fait accompli’, *batmış bitmiş* ‘collapsed’, *tanıdık bildik* ‘familiar’.

b) Imperfect participle -*DAçI*, -(*X*)*r*, -(*A*)*r*, (Neg.) -*mAz*, -(*y*)*An*:

OTu. (BT 8) *üzmadäçi käsmädäçi (ärür)* ‘(he is) the one not breaking and not cutting (the lineage and the order)’; lit. both forms are present participles in -*DAçI*; OTu. (UygW) *açtaçı yaddaçı* ‘interpreter’; MTu. *okur yazar* ‘literate person’, *bitmez tükenmez* ‘perpetual, never-ending’, *olur olmaz* ‘unnecessarily’, *olan biten* ‘goings-on’, *gelen geçen* ‘passerby’.

c) Future participle: -(*y*)*AcAk*:

MTu. *alacak verecek* ‘receivables and payables’.

1.2.3. Converb:

a) Converb -(*y*)*A*:

MTu. *tıka basa* ‘chuck-full’, *ite kaka* ‘kicking and screaming’, *güle güle* ‘good bye’, *gitgide* (< *git- git-e*) ‘gradually, more and more’.

b) Converb -(*y*)*U*:

OTu. *kolu ötünü* ‘humbly imploring’, *asa üstäyü* ‘by multiplying and increasing’.

c) Converb -(*y*)*Xp*:

OTu. *ögirip sävinip* ‘by rejoicing’, *tavranıp katıglanıp* ‘by striving’; MTu. *yatıp kalkıp dua et-* ‘to give thanks continually’, *durup durup* ‘repeatedly’.

d) Converb -*mAdAn* (< -*mAtIn* / -*mAdIn*):

OTu. *körmädin äşidmädin* ‘without seeing or hearing’, MTu. *bilip bilmeden konuşmak* ‘shoot off one's mouth, talk through hat’.

Acknowledgements

I wish to thank Assoc. Prof. Aslı Gürer, Assoc. Prof. Emre Yağlı, and Asst. Prof. Oktay Çınar for their collegial review and valuable feedback on an earlier draft of this manuscript.

Abbreviations and Symbols

+ nominal suffix

–	verbal suffix
~	synchronically coexisting overlapping forms
ABi	adjectival bicoordinative
AdvBi	adverbial bicoordinative
AdvTri	adverbial tricoordinative
AMulCo	adjectival multicoordinative
AP	adjective phrase
AsBi	asymmetric bicoordinative
AsTri	asyndetic tricoordinatives
Bi	bicoordinative
BiCo	bicoordinative
BT 8	Kara & Zieme 1977
HT	Old Turkic Xuanzang Biography
HT III	Ölmez & Röhrborn 2001
HT IX	Aydemir 2013
KB	Kutadgu Bilig (11th century)
KT	Kül Tegin inscription (8th century)
MTu.	Modern Turkish
MulCo	multicoordinative
NBi	nominal bicoordinative
NMulCo	nominal multicoordinative
NTri	nominal tricoordinative
OTu.	Old Turkic
SyBi	symmetric bicoordinative
TriCo	tricoordinative
TT V	Bang & Gabain 1931
UygW	Wilkens 2021
VBi	verbal bicoordinative
VMulCo	verbal multicoordinative
VTri	verbal tricoordinative

References

- Ağakay, M. A. (1954). Türkçede Kelime Koşmaları. *Türk Dili Araştırmaları Yıllığı - Belleten*, 2: 97-104.
- Akyalçın, N. (2024). *Türkçe İkilemeler Sözlüğü*. Alfa Basım Yayım Dağıtım.
- Aydemir, H. (2007). *Untersuchung zum Wortschatz der alttürkischen Xuanzang-Biographie (mit besonderer Berücksichtigung von Buch IX)* [Unpublished master's thesis]. Georg-August-Universität Göttingen, Germany.
- Aydemir, H. 2013: *Die alttürkische Xuanzang-Biographie IX. Nach der Handschrift von Paris, Peking und St. Petersburg sowie nach dem Transkript von Annemarie v. Gabain ediert, übersetzt und kommentiert.* (Xuanzangs Leben und Werk. 10. VdSUA. 34). Harrassowitz.
- Bang, W., & von Gabain, A. (1931). *Türkische Turfan-Texte. V.* Sitzungsberichte der Preußischen Akademie der Wissenschaften, Philologisch-Historische Klasse, (pp. 323–356).
- Bloomfield, L. (1935). *Language* (2nd rev. ed.). George Allen & Unwin.
- Browman, C. P., & Goldstein, L. (1989). Articulatory gestures as phonological units. *Phonology*, 6(2), 201–251.

- Bussmann, H. (1998). *Routledge dictionary of language and linguistics*. Routledge.
- Clauson, Sir G. 1972: *An Etymological Dictionary of Pre-Thirteenth-Century Turkish*. Clarendon Press.
- Clements, G. N. (1990). The role of the sonority cycle in core syllabification. In J. Kingston & M. Beckman (Eds.), *Papers in laboratory phonology I: Between the grammar and physics of speech* (pp. 283–333). Cambridge University Press.
- Croft, W. (2003). *Typology and universals* (2nd ed.). Cambridge University Press.
- Crystal, D. (2008). *A dictionary of linguistics and phonetics* (6th ed.). Blackwell.
- Çağatay, S. (1940–41). Uygurcada Hendiadyoinler. İçinde *Türk Lehçeleri Üzerine Denemeler*. Ankara, 1978: 29-66 (reprint of the version published in *Ankara Üniversitesi Dil ve Tarih-Coğrafya Fakültesi Yıllık Araştırmalar Dergisi*, Vol. I, Ankara, pp. 97–145).
- Dik, S. C. (1972). *Coordination: Its implications for the theory of general linguistics*. North-Holland.
- Erdal, M. 1991: *Old Turkic Word Formation. A Functional Approach to the Lexicon*. 1-2. (Turcologica 7). Harrassowitz.
- Erdal, M. 2004: *A Grammar of Old Turkic*. Brill.
- Foy, K. (1899). Studien zur osmanischen Syntax. I. Das Hendiadyoin und die Wortfolge *ana baba*. Mitteilungen des Seminars für Orientalische Sprachen. 2. Berlin. 105–136.
- Givón, T. (1985). Iconicity, isomorphism and non-arbitrary coding in syntax. In J. Haiman (Ed.), *Iconicity in syntax: Proceedings of a symposium on iconicity in syntax, Stanford, June 24–26, 1983* (pp. 187–219). John Benjamins.
- Göksel, A., & Haznedar, B. (2007). *Remarks on compounding in Turkish. MorboComp Project*, University of Bologna.
- Göksel, A. & Kerslake, C. (2005). *Turkish: A Comprehensive Grammar*. Routledge.
- Haiman, J. (1985). *Natural Syntax: Iconicity and Erosion*. Cambridge University Press.
- Haspelmath, M. (2004). Coordinating constructions: An overview. In Haspelmath, Martin (ed.) *Coordinating constructions*. Benjamins, 3–39.
- Haspelmath, M. (2007). Coordination. In T. Shopen (Ed.), *Language typology and syntactic description* (Vol. 2, *Complex constructions*, pp. 1–51). Cambridge University Press.
- Haspelmath, M., & Sims, A. D. (2010). *Understanding morphology* (2nd ed.). Hodder Education.
- Hatiboğlu, V. (1981). *Türk dilinde ikileme*. Türk Dil Kurumu Yayınları.
- Jespersen, O. (1904). *Lehrbuch der Phonetik*. B. G. Teubner.
- Johannessen, J. B. (2003). *Coordination*. Oxford University Press.

- Kara, G. & Zieme, P. (1977). *Die uigurischen Übersetzungen des Guruyogas „Tiefer Weg“ von Sa-skya Paṇḍita und der Mañjuśrīnāmasaṃgīti*. (Berliner Turfantexte 8). Akademie-Verlag.
- Karaman, A. (2022). *Eski Türkçede İkilemeler*. Türk Dil Kurumu.
- Lindblom, B. (1990). Explaining phonetic variation: A sketch of the H&H theory. In W. J. Hardcastle & A. Marchal (Eds.), *Speech production and speech modelling* (pp. 403–439). Kluwer.
- Müller, W. (2004). *Reduplikationen im Türkischen*. Harrassowitz.
- Ölmez, Z. K. (1997). Kutadgu Bilig’de İkilemeler (1). *Türk Dilleri Araştırmaları*, 7: 19–40.
- Ölmez, Z. K. (1998). Kutadgu Bilig’de ikilemeler (2). In J. P. Laut & M. Ölmez (Eds.), *Bahşı Ögdisi: Klaus Röhrborn armağanı* (pp. 235–260). Harrassowitz.
- Ölmez, M. (2017). Eski Uygurca ikilemeler üzerine. *Türk Dili Araştırmaları Yıllığı – Belleten*, 65/2: 243–311.
- Ölmez, M. (2022). Eski Uygurcada İkilemeler Üzerine – 2. *International Journal of Old Uyghur Studies*, 4/1: 38–94.
- Ölmez, M. & Röhrborn, K. 2001: *Die alttürkische Xuanzang-Biographie III. Nach der Handschrift von Paris, Peking und St. Petersburg sowie nach dem Transkript von Annemarie v. Gabain herausgegeben, übersetzt und kommentiert*. (VdSUA 34, 7). Harrassowitz.
- Özkan, B. (2019). *Türkiye Türkçesinde İkili Tekrarlar: Derlem Tabanlı Bir Uygulama*. Pegem Akademi.
- Parker, S. G. (2002). *Quantifying the sonority hierarchy* [Doctoral dissertation]. University of Massachusetts Amherst.
- Parker, S. G. (2008). Sound level protrusions as physical correlates of sonority. *Journal of Phonetics*, 36(1), 55–90.
- Röhrborn, K. (1983). Gruppenflexion und Komposition im Türkischen. In A. von Gabain & W. Veenker (Eds.), *Documenta Barbarorum: Festschrift für Walther Heissig zum 70. Geburtstag* (pp. 317–323). Harrassowitz.
- van Schaik, G. (2002). *The noun in Turkish: Its argument structure and the compounding straitjacket*. Harrassowitz.
- Schilling, U. (2001). Zur Bildung und Bedeutung von Paarformeln im Alt türkischen. *Türk Dilleri Araştırmaları*, 11, 153–170.
- Şahin, H. & Özbek, E. E. (2025). *Türk Dilinde İkilemeler*. Paradigma Akademi Yayınları.
- Şen, S. (2003). *Eski Uygur Türkçesinde İkilemeler* [Unpublished master’s thesis]. Ondokuz Mayıs University, Institute of Social Sciences.
- Uzun, İ. P. (2020). *Seslem Yapısı ve Titreşim Hiyerarşisi*. İçinde I. P. Uzun (Ed.), *Kuramsal ve uygulamalı sesbilim* (ss. 33–72). Seçkin Yayıncılık.
- Wilkens, J. (2021). *Handwörterbuch des Altuigurischen: Altuigurisch–Deutsch–Türkisch*. De Gruyter.
- Zieme, P. (1996). *Altun Yaruq Sudur: Vorworte und das erste Buch. Edition und Übersetzung der altuigurischen Version des Goldglanz-Sūtra (Suvarṇaprabhāsottamasūtra)* (Berliner Turfantexte, 18). Brepols.

GRASPING GRAMMAR: THE HAPTIC TURN IN LANGUAGE DOCUMENTATION

Firat BAŞBUĞ

Asst. Prof., Istanbul Medeniyet University, Department of Linguistics, firatbasbug@gmail.com

ORCID: <https://orcid.org/0000-0001-9824-123X>

Başbuğ, F. (2025). Grasping grammar: The haptic turn in language documentation. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 65–80). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X003>

ABSTRACT

Language documentation has developed robust standards for recording the optical and acoustic dimensions of speech events, but it less often preserves the material conditions under which force-sensitive grammatical contrasts become interactionally available. In domains involving manipulation, fit, resistance, fracture, and effort, visually similar events may differ in ways that are grammatically relevant yet not recoverable from image or waveform alone. This chapter proposes the Haptic Minimal Pair (HMP) as a controlled elicitation method that holds visual geometry constant while varying a single material parameter such as friction, compliance, mass distribution, or tolerated fit. The proposal is methodological rather than conclusive: drawing on exploratory field observations from Sakha, Telengit, and Cilician Arabic, the chapter presents proof-of-concept cases showing how materially calibrated stimuli can make posture predicates, handling verbs, fracture constructions, and ideophones more observable, more comparable across sessions, and more archivally accountable. It also outlines a compact metadata protocol for preserving both digital design specifications and the physical instantiation of stimuli under CARE-aligned, community-governed archival conditions. The chapter argues that documentary adequacy in force-sensitive domains requires not only recording what speakers say and what cameras capture, but also documenting the calibrated material constraints that make particular grammatical choices pragmatically available.

Keywords: documentary linguistics; language documentation; force-sensitive grammar; elicitation design; Haptic Minimal Pair; 3D printing

This research is currently supported by the SADA Institute. Grant Number: SADA-25112

Belgeleme dilbilimi bugüne dek konuşma eyleminin optik ve akustik boyutlarını kaydetmede detaylı standartlar geliştirmiş olsa da kuvvete duyarlı dilbilgisel karşıtlıkların ortaya çıktığı maddi ve etkileşimsel koşulları aynı titizlikle arşivleyememektedir. Nesne manipülasyonu, direnç, esneme veya kırılma gibi fiziksel temas içeren eylemlerde görsel olarak birbirine benzeyen olaylar, dilbilgisel açıdan derin farklılıklar taşıyabilir; ne var ki bu ince ayrımlar standart ses ve görüntü kayıtlarında genellikle kaybolmaktadır. Bu çalışma, görsel geometriyi sabit tutarken sürtünme, esneklik, kütle dağılımı ya da geçme toleransı gibi tek bir maddi parametreyi değiştiren kontrollü bir veri derleme yöntemi olarak Dokunsal Minimal Çift'i (DMÇ) önermektedir. Öneri, yöntemsel niteliktedir: Çalışma; Saha (Yakut), Telengit (Altay) ve Çukurova Arapçası üzerine yürütülen keşif niteliğindeki saha gözlemlerinden hareketle, maddi olarak kalibre edilmiş uyarıların duruş ve kavrama eylemlerini, kırılma yapılarını ve yansıma sözcüklerini nasıl daha gözlemlenebilir, karşılaştırılabilir ve arşivsel bakımdan denetlenabilir kıldığını gösteren yöntemsel gösterim örnekleri sunmaktadır. Bölüm ayrıca, dijital tasarım özellikleri ile uyarıların maddi üretim düzeylerini birlikte koruyan, CARE ilkeleriyle uyumlu ve topluluk-temelli bir üstveri çerçevesi önermektedir. Çalışma, kuvvete duyarlı alanlarda yapılacak nitelikli bir dil belgelemesinin yalnızca kameranın kaydettiği sesi ve görüntüyü değil, o dilbilgisel seçimleri pragmatik düzeyde mümkün kılan fiziksel ve maddi uyarıların da kalibre edilmiş biçimde kayıt altına alması gerektiğini savunmaktadır.

Anahtar Sözcükler: belgeleme dilbilimi; dil belgeleme; kuvvete duyarlı dilbilgisi; veri derleme tasarımı; dokunsal minimal çift; 3D baskı

1. Introduction

Over the last three decades, documentary linguistics has consolidated a rigorous methodological paradigm for producing archives that are ethically grounded, philologically durable, and analytically tractable (Himmelman 1998, 2006; Woodbury 2003, 2011). At the centre of this paradigm lies a commitment to comprehensive, naturalistic corpora: records of language as it is used in the course of ordinary, situated activity. This commitment has reshaped the field's evidential standards. Analyses are now expected to remain accountable to a primary documentary record that can support reuse, reinterpretation, cross-linguistic comparison, and community-oriented forms of access and benefit (Himmelman 1998; Woodbury 2011). These developments have unfolded alongside the standardisation of high-quality audio-visual recording, time-aligned annotation, and increasingly robust digital archiving infrastructures (Austin 2006; Conathan 2011). Yet for all this progress, contemporary documentary practice remains uneven in what it captures well. Sound and image are preserved with increasing fidelity, but the material physics of the communicative encounter are usually left implicit. Documentary records are therefore rich in waveform and pixel while treating physical constraint as background. Variables such as mass distribution, frictional profile, compliance under pressure, and dimensional tolerance are rarely specified, even though they may be crucial to the event being documented. Nor are these properties reducible to visible form alone. They become interactionally and semantically salient in contact: through hands that probe, test, slip, adjust, fail, and try again.

This matters because, in many linguistic systems, grammar is not organised only around what can be seen. It is also organised around what must be felt, resisted, managed, or overcome. When such resistances are absent from elicitation settings, or when they vary unpredictably across sessions without being parameterised in the archive, part of what is grammatically available to speakers may remain pragmatically dormant. A corpus may thus appear comprehensive—especially when supported by careful transcription and glossing—while remaining thin in precisely those domains where language is calibrated to engagement with a resisting material world. What is missing in such cases is not simply more speech, but a replicable account of the biophysical conditions under which particular morphosyntactic resources become relevant and worth deploying (Woodbury 2011).

The issue, then, is methodological rather than merely technological. Many morphosyntactic systems are tuned to manipulation events in which causal relations are not recoverable from static geometry alone. They emerge through effort, obstruction, slippage, support, pressure, and release. Talmy's force-dynamic framework is particularly useful here because it makes explicit a contrast space that languages recurrently lexicalise and grammaticalise: enabling and blocking, overcoming, stable support, and collapse (Talmy 1988). At the same time, it offers a practical diagnostic for documentary design. If an elicitation stimulus cannot generate physical resistance, then the constructions organised around resistance are likely to be under-elicited. In fieldwork terms, these are often the very contexts in which speakers recruit highly specific handling predicates, contact verbs, instrumental morphology, and aspectual packaging to distinguish, for example, a single successful completion from repeated, effortful attempts. They are also contexts that recruit directive sequences and conversational repair, because resistance makes trouble publicly observable and interactionally accountable.

Where the stimulus cannot push back, the communicative ecology that licenses force-dynamic construals is correspondingly weakened. The well-known Atsugewi materials provide a particularly clear illustration—not because they are typologically anomalous, but because they make the dependence of grammatical selection on non-visual, haptic parameters unusually explicit. Atsugewi motion verbs require the selection of dependent roots that conflate motion with figure-type distinctions that are often tactile and material in character: *-caq-* for 'a slimy lumpish object' and *-staq-* for 'runny icky material' (Talmy 1985; cf. Talmy 1972). Beyond figure-type, the system also encodes causal vectors in prefixal satellites, including *uh-* for causation by gravity, *ca-* for motion due to wind, and *cu-* for axial pressure exerted by a linear object. Talmy's attested example brings these dimensions together in a single verbal complex.

(1) s'w-cu-staq-cis-a

1SG-axial.pressure-runny.icky.material-into.fire-FACT

'I prodded the guts into the fire with a stick.' (Talmy 1985: 77)

What matters for present purposes is not simply the descriptive richness of the form, but the conditions that license it. A recording that captures only visible displacement, without a controlled account of the object's resistance and the applied force that organised the action, may be insufficient to reconstruct why this rather than another morphological constellation was selected.

Comparable dependencies on material constraint and haptic feedback are documented in a range of other languages. These include Navajo classificatory handling stems (Young & Morgan 1987), Tzeltal handling verbs sensitive to modes of grasp and placement (Brown 2008), and Korean *kki-ta* as a widely cited tight-fit predicate in which forceful insertion, rather than mere containment, becomes salient (Choi & Bowerman 1991). Such cases suggest that the problem is not marginal. They reveal a broader documentary blind spot: where elicitation relies primarily on images, videos, or objects whose resistive properties remain unspecified, the resulting archive may preserve what moved while under-specifying what had to be overcome for that movement to occur.

This chapter takes that asymmetry as a methodological problem for language documentation. Its central claim is that if we want to document grammars that are sensitive to manipulation, force, and material resistance, then audio-visual recording alone is not enough. What is also needed is a principled way of calibrating and documenting the physical conditions under which these grammatical contrasts become available in interaction. The broader aim is not to displace established documentary practice, but to extend it: from the recording of visible action to the documentation of the resistive material environments within which particular forms become pragmatically and grammatically motivated. Although the demonstrations in this chapter focus on force-sensitive domains, the methodological logic is broader: the HMP framework is in principle applicable wherever linguistically relevant contrasts depend on parameterizable, non-visual properties of the elicitation ecology.

2. Beyond Audio-Visual Adequacy: Material Conditions as Documentary Evidence

Comparative documentary practice has often depended on a disciplined reduction. By stabilising a representational field, pictorial stimuli—ranging from *The Pear Stories* to spatial elicitation kits such as the *Topological Relations Picture Series*—hold constant a set of perceptual cues, thereby making cross-site comparison feasible in the absence of a shared elicitation metalanguage (Chafe 1980; Bowerman & Pederson 1992). That achievement is not in question. The present concern is not documentary linguistics as such, but visually dominated elicitation ecologies in force-sensitive domains. Cross-linguistic research on sensory lexicons shows that the relative codability of perceptual domains varies substantially across languages, and that vision cannot simply be assumed to be universally privileged in lexical encoding (Majid et al. 2018).

In domains where morphosyntax is recruited to track manipulation, causation, and event segmentation, the relevant semantic contrasts are rarely just configurational. They are force-dynamic and contact-sensitive, keyed to properties discovered through resistance rather than ocular inspection. A picture can be exemplary for ensuring recognisability, but it cannot compel a participant to discover that an object jams, that a surface slips, that a material yields, or that a fit must be physically forced. The gap is therefore not a vague call for ‘multimodality’, but the absence of an archivable account of material constraint.

Force dynamics is not an optional enrichment of spatial description; it is a recurrent grammatical route into causation, control, and event segmentation (Talmy 1988). It is also unusually sensitive to elicitation ecology. A static line drawing can effectively elicit descriptions of containment or support. However, even typologically robust dynamic stimuli explicitly designed to capture manipulation and change of state, such as the widely used Cut and Break video clips (Bohnemeyer, Bowerman, & Brown 2001), remain vulnerable to visual conflation: they successfully depict the kinetic trajectory and visible outcome of an action, yet leave the applied force and felt resistance entirely uncalibrated. They cannot generate wedging, sticking, coming loose, or the distributional ecology of constructions deployed when outcomes are contingent on physical pushback.

Where the stimulus cannot resist, grammars of attempt, effort, and failure are pushed to the margins of the record—not because speakers lack these resources, but because the encounter provides no durable reason to deploy them. The audio-visual record is then vulnerable to visual conflation: a tight fit and a merely careful placement may look identical from one camera angle; low friction and poor grip can be visually indistinguishable; compliance under pressure may only be visible as a final outcome, not as a felt gradient. Without an explicit record of the resistance profile that organised the action, the corpus cannot reliably support reproducible inferences about why a particular handling predicate, aspectual packaging, or diathesis alternation was locally licensed. Semantic conflation is a useful analytic precedent for characterising this documentary risk. Haviland, following Talmy, uses the term for cases where a single lexical item bundles distinct semantic components—his standard illustration is English *climb*, which conflates vertical displacement with effort—so that an elicitation context that renders only one component salient can obscure the conditions licensing the full package (Talmy 1985; Haviland 2006: 141–142).

In the present chapter we extend this logic from the lexicon to the stimulus ecology: visual protocols can inadvertently produce documentary conflation, collapsing materially distinct events into a single “same-looking” scene because geometry and outcome are preserved while resistance is erased. Under such conditions, contrasts organised around sticking, jamming, slipping, wedging, or graded fit may be systematically under-elicited—or misanalysed as discourse noise—because the stimulus makes the relevant force-dynamic component pragmatically inert. The methodological point of

the Haptic Minimal Pair (HMP) is therefore de-conflational: by holding visual geometry constant while manipulating a single physical parameter (e.g., compliance, friction, clearance), the protocol forces grammars to resolve distinctions that the visual record would otherwise collapse, making the licensing conditions for handling predicates, attempt/repair formats, and telic packaging reproducible in the archive. Crucially, material constraints do not always target the same grammatical stratum across languages. While Atsugewi makes the dependency affixal, other systems distribute it differently—often through phonosemantic selection (e.g., impact ideophones), diathesis morphology (e.g., medial voice alternations), or large descriptive-verb inventories. The evidential burden is nevertheless uniform: where grammatical retrieval depends on deformation, friction, tightness of fit, or grip-type control, the archive must retain enough of the constraint regime to render token distributions interpretable as grammar rather than as accidental stimulus physics.

A further pressure on elicitation design comes from referential indeterminacy (Haviland 2006). Realia can mitigate this, but exclusive reliance on familiar artefacts introduces a counter-risk: indexical saturation. Because such objects carry stabilised labels and embodied routines, interaction can collapse into rapid naming (“that’s a basket”), bypassing the morphosyntax recruited for part structure, graded fit, attempt, and repair (Levelt 1989). Novelty is an imperfect corrective: unfamiliar local artefacts may inhibit ready-made labels yet also trigger guarded guessing rather than description. The design target is narrower—stimuli that are physically legible while remaining lexically and culturally underdetermined.

3. From Calibrated Material Stimuli (CMS) to Haptic Anchoring

The approach developed here grows out of field observations among the Chulym, Sakha, and Telengit communities, where spatial and motion-related distinctions repeatedly emerged as linguistically salient in naturalistic settings (Başbuğ 2016; 2023). Material constraint belongs in the archive as something described, measured, and made reproducible. In this context, 3D printing emerges not merely as a high-tech novelty, but as a critical methodological tool for creating stimuli with calibrated, 'built-in' physical properties within language documentation practice. Some remarks from iterative prototyping during the development of the MumuLab framework (Turkish Patent and Trademark Office, Trademark App. No. 2025/152264) demonstrate that nominally identical elicitation tools or mascots—derived from a single digital mesh—can generate divergent interactional routines solely as a function of variations in density, compliance, or surface friction. Drawing on pilot field studies conducted for the CODE Model (Andiç & Başbuğ, in review), we distinguish between the stimulus as a physical tool and as a cognitive trigger—a distinction that can be further sharpened through a semiotic lens. In Haviland’s Peircean framing, an index signifies by a physical or causal link to what it stands for—his linguistic illustration is ouch, whose meaning is licensed by an imagined causal coupling between bodily pain and

the utterance (Haviland 2006: 143). I treat calibrated objects as enabling a comparable coupling by design: when a hand encounters a controlled roughness, jam, slippage, or elastic give, the artefact's resistance becomes a haptic index for the very parameter the grammar is recruited to encode. While the concept of 'haptic anchoring' originates in ecological psychology to describe how physical resistance provides a reference frame for bodily postural stability (e.g., Mauerberg-deCastro et al. 2014, p. 303), I repurpose it here as a methodological warrant for language documentation. On this view, the Haptic Anchor (HA) is not a metaphorical "grounding" claim but a strict methodological warrant: linguistic choices (attempt morphology, repair formats, handling predicates, telicity) are made publicly accountable to a reproducible, archivable resistance profile. In short, the CODE protocol implements an engineered indexicality in which contact physics is not backgrounded but converted into an explicit evidential condition of the documentary record.

The Calibrated Material Stimulus (CMS) refers to the engineered object itself, defined by precise physical parameters (e.g., a specific friction profile). In contrast, the Haptic Anchor (HA) refers to the interactional effect: the process by which a speaker's linguistic retrieval is grounded in the object's immediate physical resistance. While the CMS ensures that every participant encounters the same physical challenge, the HA ensures that their grammatical output is a direct response to that calibrated reality rather than a generic label. A rigorous protocol does not require exhaustive fabrication detail; it requires documentary variables. We identify four primary loci that recurrently reorganise morphosyntactic retrieval across the grammars surveyed here:

Mass and its distribution: Reorganises effort profiles, directly shifting event segmentation, completion metrics, and the encoding of agentivity.

Frictional profile: Reorganises grammatical control, yielding structural contrasts between slippage (often construed as diminished control) and dragging (often construed as sustained effort).

Toleranced fit: Reorganises spatial placement, systematically distinguishing loose volumetric containment from forcing, jamming, and the iterative repair sequences that a tight interference fit structurally invites.

Grasp affordances: Reorganises handling predicates, as lexical retrieval frequently depends on biomechanical grip-type control and contact geometry rather than the transfer trajectory.

The point is not to legislate a universal parameter set, but to ensure that whatever parameter regime makes a contrast pragmatically live is retained in the deposit as part of what renders the linguistic record interpretable.

3.1. The Haptic Minimal Pair (HMP) Framework

If additive manufacturing is treated not merely as a fabrication technique but as a parametric system for translating digital specification into constraint-bearing material form, its principal methodological contribution lies in the formalisation of what I term the Haptic Minimal Pair (HMP). While terms like 'haptic phonemes' have been used metaphorically in Human-Computer Interaction (HCI) to describe basic vibrotactile signals for artificial interfaces (e.g., Enriquez et al. 2006), the Haptic Minimal Pair (HMP) proposed here operates within the domain of language documentation. It is not an artificial signaling system, but a methodological framework designed to elicit and isolate natural morphosyntactic responses to physical constraints. By analogy with phonological minimal pair methodology—where tightly controlled acoustic contrasts isolate phonemic distinctions—the HMP framework enables the systematic manipulation of a single physical parameter (e.g., internal density, surface friction, elastic compliance) while preserving constant visual geometry. The analytic objective is controlled variable isolation: resistance is introduced as a distinctive haptic feature under conditions of geometric invariance.

The analogy to phonological minimal pairs is methodological rather than isomorphic. In the physical world, it is difficult to change surface friction without simultaneously altering grip affordances; therefore, the HMP is a heuristic method rather than a tool for perfect isolation. The goal is simply controlled contrast—revealing latent variables without aiming for absolute laboratory sterilization. Ultimately, the HMP is a methodological correction. It helps researchers distinguish between stimuli that look identical on camera but offer entirely different physical affordances to the speaker. When earlier corpora are revisited through this lens, certain morphosyntactic alternations—previously attributed to speaker variability or discourse fluctuation—become interpretable as responses to unrecorded resistance conditions. The so-called “missing physics” of elicitation thus emerges as a latent variable in corpus interpretability. What is at stake is not the replacement of visual protocols, but their supplementation: unless physical constraints are rendered systematically contrastive, a dimension of grammatical conditioning remains structurally underdetermined within the documentary record. High-definition video stabilises optical evidence, but its epistemological reach is limited for force-dynamic variables (Franchetto 2006). Where the licensing of roots or constructions hinges on resistance, a purely visual ecology produces visual conflation (Talmy 1985; Haviland 2006). Accordingly, haptic parameters must be integrated into the elicitation environment to ensure methodological rigor. The HMP is therefore proposed not as an empirical finding but as a methodological correction. Its immediate value lies in making materially conditioned contrasts observable, comparable, and archivally transparent. In the present chapter, this value is demonstrated through force-sensitive domains because they make the evidential role of resistance especially visible. This should not be taken to imply that the framework is intrinsically restricted to those domains. More generally, the HMP provides a way of asking whether

contrasts that appear linguistically unstable, optional, or speaker-specific may in fact depend on unrecorded properties of the elicitation setting.

4. Material constraints in morphosyntax: Case studies

Data note. The discussion in Sections 4.1–4.4 is intended as a methodological proof of concept rather than as a balanced cross-linguistic comparison. Unless otherwise indicated, the examples from Sakha, Telengit, and Cilician Arabic are drawn from the author’s field notes and exploratory elicitation sessions. They are used here to show where materially calibrated stimuli can make linguistically relevant contrasts more observable and more archivally accountable, not to provide exhaustive grammatical descriptions of the languages concerned.

4.1. Spatial volume, kinematic envelope, and posture geometry

Standard visual elicitation protocols typically rely on schematic line drawings or geometric abstractions (e.g., “stick figures”) to elicit posture descriptions. While 2D representations can optically encode volume through perspective, they suppress key physical variables of the encounter—especially haptic feedback, mass distribution, and tissue compliance. In the material considered here, posture descriptions are sensitive to such non-visual affordances as well as to visible geometry. In the Telengit data, sitting events suggest a lexical contrast linked to differences in body volume and posture configuration. When a larger-bodied sitter adopts a seated posture, abdominal and thigh mass can constrain neutral joint angles and favor either a rounded geometry (+Volume) or a base-expansion strategy (+Stability). In the material discussed here, such configurations make converbs like *bolčoy-yp* and *alčay-yp* (or *taltay-yp*) plausible modifiers of the posture verb *otur-* (‘sit’):

(2a) *Semis kiži bolčoy-yp kal-gan otur-dy-Ø.*
 fat person become_spherical-CVB AUX-PTCP sit-PST-3SG
 ‘The fat person sat down rounded (bulging like a sphere).’

(2b) *Semis kiži alčay-yp kal-gan otur-dy-Ø.*
 fat person splay_legs-CVB AUX-PTCP sit-PST-3SG
 ‘The fat person sat with legs splayed wide (due to tissue obstruction).’

These converbs are less natural with low-volume subjects. Where the requisite spatial envelope is absent, speakers may instead prefer a different converbial root describing a folded, angular posture:

(2c) *Aryk kiži korčoy-yp kal-gan otur-dy-Ø.*
 thin person hunch_curl-CVB AUX-PTCP sit-PST-3SG
 ‘The thin person sat hunched/curled up.’

Exploratory elicitation with Sakha-speaking participants suggests that the resultative auxiliary construction (e.g., *-yp kal-gan*) can be used to frame

the posture as a sustained state rather than a momentary action. A related contrast is visible in the Sakha material as well, though the larger-bodied case appears more readily in descriptive phrasing than in a single dedicated converb, while *nikyiy-an olor-or* ('sits with the head drawn into the shoulders / curls inward while sitting') remains a clearer low-volume form. Within a parametric 3D-printing paradigm, the key methodological point is that visual geometry should be separated from latent physical properties. Changing only the internal infill of an articulated mannequin chiefly alters mass, not the external kinematic envelope; to elicit interference-based forms like *alčay*, the stimulus would need either a modified outer shell or modular "tissue" attachments. Otherwise, the relevant biomechanical interference remains under-specified.

4.2. Fracture mechanics, material compliance, and result entailment (Sakha and Cilician Arabic)

When standard visual stimuli depict the destruction of an object, they typically provide only the optical trajectory of separation. Verbs of destruction, however, often encode internal fracture mechanics—the structural reaction of a material to applied stress. Standard elicitation protocols can also import framing bias when the prompt itself presupposes an active/transitive event (e.g., Russian *когда ломаем*, 'when we break it'), making transitive roots more likely regardless of the material event. Haptic elicitation does not eliminate this problem, but it can reduce reliance on translation prompts by letting speakers respond to the material behavior of the object itself. In the Sakha material considered here, spontaneous material failure can be expressed both by relatively general intransitives and by more specific forms. For the contrasts discussed here, the more specific set is analytically the most useful, surfacing either as a simple intransitive or with the completive auxiliary *bar-* ('go'):

(3a) *Taas xampariy-da-Ø.*

stone break.into.pieces-PST-3SG

'The stone shattered/broke into pieces (isotropic brittle failure).'

(3b) *Et tirit-a bar-da-Ø.*

meat tear-CVB

AUX.COMPL-PST-3SG

'The meat tore apart/shredded (yielding compliance).'

(3c) *Mas tostu bar-da-Ø.*

wood snap.ADV AUX.COMPL-PST-3SG

'The wood snapped/broke with a crack (anisotropic fibrous failure).'

Cilician Arabic appears to stratify deformation along a similar axis of material compliance. In the material considered here, brittle destruction of a rigid branch favors *kasar* ('break'), tensile separation of pliable matter favors *qba'* ('pluck/tear'), and viscoplastic deformation favors the mediopassive *in-ma'as* ('get squashed/mashed'). Physical elicitation can also help clarify the relation between telicity and result entailment. In these data, speakers resisted

cancellation framings built directly on brittle-fracture predicates, suggesting that *kasar* and its Form VII counterpart *in-kasar* are relatively strong result predicates. When resistance must be expressed without an achieved break, speakers instead separate the force/manner verb (*ḍarab*, ‘to hit/strike’) from the unattained telic outcome:

- (4) *Ḍarab-t=o* *bas mā* *in-kasar-Ø*.
 hit.PFV-1SG=3SG.M.OBJ but NEG MID-break.PFV-3SG.M
 ‘I hit it (applied force), but it didn’t break (refusal to yield).’
 (**Kasar-t=o bas mā in-kasar-Ø*)

By integrating additive manufacturing, these material differences can at least be made more explicit and reproducible. A rigid PLA print may approximate brittle failure, whereas TPU is more likely to model elastic deformation than the kind of permanent yielding targeted here. To elicit forms associated with squashing or mashing, the stimulus may need to incorporate a more clearly viscoplastic material, such as polymer clay or another deformable insert. Without an object whose mechanical profile is specified and archivable, the conditions governing result entailment remain difficult to reconstruct from the recording alone.

4.3. Surface friction (μ) and grip affordances

The physics of the external surface—especially friction and grip affordances—can reorganise the lexicalisation of handling predicates. Changing the surface texturing of a parametrically designed stimulus alters how readily it slides, catches, or can be held securely. In Cilician Arabic, a controlled slide with manageable traction favors *zaḥlaq-a*, whereas a sudden lubricated slip (e.g., wet soap shooting from the hand) favors *faḥlaṣ-a*. Grip affordances also matter. A rigid object may support a neutral cylindrical hold (*msik-t=o*, ‘I held it’); a minute, solid object may support a precision pincer grip (*niqqay-t*, ‘I precisely picked/selected’). If that rigid volume is replaced by a yielding or granular mass (e.g., heavily hydrated dough or wet sand), speakers may instead prefer a form like *kabbaš-t=o*, associated here with a conformable power grip:

- (5a) *Niqqay-t* *l-ḥajar*.
 pick_precisely.PFV-1SG DEF-stone
 ‘I picked/selected the stone (precision pincer grip).’

- (5b) *Kabbaš-t=o*.
 clutch_squeeze.PFV-1SG=3SG.M.OBJ
 ‘I clutched/squeezed it (conformable power grip for viscoplastic/granular matter).’

4.4. Acoustic-haptic coupling: Impact ideophones (Sakha and Cilician Arabic)

The transduction of kinetic energy into auditory-haptic feedback provides an equally productive domain for sound symbolism. Cross-linguistic research demonstrates that ideophones act as cross-modal transducers, mapping auditory features onto internal density, structural damping, and fluid mechanics (Dingemanse 2012). In the Sakha material considered here, impact ideophones track broad differences in material behavior and impact profile rather than sound alone. For the present contrast, the most defensible pair is *lūs* for the dull fall of a heavy rigid object and *pal* for the wet splat of a soft, saturated mass. The field observations reported here are consistent with an acoustic-material mapping of this kind:

(6a) *Kürbe taas "lūs" gın-a tūs-te-Ø.*
 stone.block IDEO:heavy.thud do-CVB fall-PST-3SG
 ‘The stone block fell with a heavy dull thud.’

(6b) *Inax saaga "pal" gın-a tūs-te-Ø.*
 cow.dung IDEO:splat do-CVB fall-PST-3SG
 ‘The cow dung fell with a wet splat.’ (*pal*, overdamped viscoelastic dissipation).

Cilician Arabic closely parallels this acoustic-material mapping. A rigid, high-frequency impact is robustly encoded as *ṭāq* (*ṭili'-Ø ḥiss ṭāq*, ‘a “tak” sound rose up’). When a heavy, fully saturated substance impacts a surface, the complex fluid mechanics and resulting viscous spread are cross-modally mapped onto heavy labial consonants: *labbāḥ* (*l-waḥīl wāqi'-Ø labbāḥ*, ‘the mud falls with a splat’).

5. Archiving the Autographic: A Metadata Protocol for Material Stimuli

For material constraints to be analytically useful, they must be integrated into ordinary archival practice. However, this should not overwhelm fieldworkers with metadata bloat. We simply need a compact extension for materially calibrated stimuli. I therefore propose a baseline meta-documentary profile for parametrically designed stimuli. Drawing conceptually on Goodman’s (1968) distinction between allographic specification and autographic instantiation, this framework is structured across two primary vectors:

Digital Specification (Design): The underlying geometric files (e.g., .stl, .obj). The deposit should state who may access, reprint, or modify the design under CARE-guided community permissions, preventing culturally sensitive files from becoming “runaway objects” (Carroll et al. 2020, pp. 6-7). While exhaustive physical testing in the field is unrealistic, archiving this digital blueprint should be a minimal standard for future researchers.

Material Instantiation and Physical Condition: The basic fabrication variables that determined the object’s behaviour (e.g., material type like PLA

vs. TPU, infill density) and a brief note on its state during the session. Recording wear, looseness, or repairs helps distinguish genuine grammatical patterning from stimulus drift across sessions.

Archiving this metadata keeps the stimulus auditable. It allows later analysts to ask whether a morphosyntactic gap reflects a boundary of the grammar or simply the absence—or later erosion—of the physical trigger on which the contrast depended.

5.1. Limitations and future directions

Three limitations should be stated explicitly. First, the cases discussed here are exploratory and are not presented as balanced cross-linguistic datasets. Second, an HMP seeks controlled contrast rather than perfect isolation, since altering one material parameter may also affect other perceptual dimensions. Third, the present chapter demonstrates the framework in force-sensitive domains, but the framework itself is not claimed to be limited to them; its extension to other elicitation domains remains to be established empirically, domain by domain. Future work should compare matched visual and materially calibrated stimuli, report fabrication variables systematically, and test whether comparable effects recur across sessions, speakers, materials, and grammatical domains.

6. Conclusion: Toward a Haptically Aware Documentary Praxis

This chapter proposes a limited extension to documentary practice. When grammatical contrasts depend on non-visual properties of an event, the elicitation record should preserve the relevant material conditions alongside audio and video. Visible form alone does not always capture the resistance, fit, compliance, or friction that can make particular contrasts available in interaction. The cases from Sakha, Telengit, and Cilician Arabic are offered as proof-of-concept examples of both the problem and one possible response.

Although the examples discussed here come from force-sensitive domains, the methodological point is broader. Whenever linguistic contrasts depend on parameterizable features of the elicitation environment, those features should be documented rather than left implicit. Additive manufacturing provides one practical way of doing so, because it allows researchers to reproduce stimuli whose friction, compliance, mass distribution, and fit can be varied in controlled ways. Used within community-governed, CARE-aligned archival practice, such stimuli can complement naturalistic corpora by making materially relevant aspects of the event more explicit, comparable, and recoverable.

This is not an argument for replacing naturalistic documentation or for reducing linguistic meaning to physical variables. It is an argument for preserving the material conditions that shape speakers' construals of events when those conditions matter for grammatical choice. If language documentation aims to preserve a durable record of meaning in use, then

archives must retain enough of the elicitation environment to make the linguistic record interpretable.

Abbreviations: ADV adverbial, AUX auxiliary, COMPL completive, CVB converb, DEF definite, FACT factual, IDEO ideophone, M masculine, MID middle, NEG negation, OBJ object, PFV perfective, PST past, PTCP participle, SG singular.

References

- Andiç, S., & Başbuğ, F. (n.d.). *The CODE model: Object-mediated play as a scalable framework for dyadic co-regulation in early childhood* [Manuscript submitted for publication].
- Austin, P. (2006). Data and language documentation. In J. Gippert, N. P. Himmelmann, & U. Mosel (Eds.), *Essentials of language documentation* (pp. 87–112). Mouton de Gruyter. <https://doi.org/10.1515/9783110197730.87>
- Başbuğ, F. (2016). Tehlikedeki dillerin belgelenmesi üzerine: Çulımca ile ilgili bazı tespitler. *idil*, 6(28), 217–232. <https://doi.org/10.7816/idil-06-28-14>
- Başbuğ, F. (2023). Dil, uzam, zaman ve hareket üzerine. *Ulakbilge*, 11(88), 928–940. <https://doi.org/10.7816/ulakbilge-11-88-03>
- Bohnenmeyer, J., Bowerman, M., & Brown, P. (2001). Cut and break clips. In S. C. Levinson & N. J. Enfield (Eds.), *Manual for the field season 2001* (pp. 90–96). Max Planck Institute for Psycholinguistics.
- Bowerman, M., & Pederson, E. (1992). Topological relations picture series. In S. C. Levinson (Ed.), *Space stimuli kit 1.2*. Max Planck Institute for Psycholinguistics.
- Brown, P. (2008). Verb specificity and argument realization in Tzeltal child language. In M. Bowerman & P. Brown (Eds.), *Crosslinguistic perspectives on argument structure: Implications for learnability* (pp. 167–189). Lawrence Erlbaum.
- Carroll, S. R., Garba, I., Figueroa-Rodríguez, O. L., Holbrook, J., Lovett, R., Materechera, S., Parsons, M., Raseroka, K., Rodríguez-Lonebear, D., Rowe, R., Sara, R., Walker, J. D., Anderson, J., & Hudson, M. (2020). The CARE principles for Indigenous data governance. *Data Science Journal*, 19, 43.
- Chafe, W. (Ed.). (1980). *The pear stories: Cognitive, cultural, and linguistic aspects of narrative production*. Ablex.
- Choi, S., & Bowerman, M. (1991). Learning to express motion events in English and Korean: The influence of language-specific

- lexicalisation patterns. *Cognition*, 41, 83–121.
[https://doi.org/10.1016/0010-0277\(91\)90033-z](https://doi.org/10.1016/0010-0277(91)90033-z)
- Conathan, L. (2011). Archiving and language documentation. In P. K. Austin & J. Sallabank (Eds.), *The Cambridge handbook of endangered languages* (pp. 235–254). Cambridge University Press.
- Dingemanse, M. (2012). Advances in the cross-linguistic study of ideophones. *Language and Linguistics Compass*, 6(10), 654–672.
- Enriquez, M., MacLean, K. E., & Chita, C. (2006). Haptic phonemes: Basic building blocks of haptic communication. In *Proceedings of the 8th International Conference on Multimodal Interfaces*.
- Franchetto, B. (2006). Ethnography in language documentation. In J. Gippert, N. P. Himmelmann, & U. Mosel (Eds.), *Essentials of language documentation* (pp. 183–211). Mouton de Gruyter.
- Goodman, N. (1968). *Languages of art: An approach to a theory of symbols*. Bobbs-Merrill.
- Haviland, J. B. (2006). Documenting lexical knowledge. In J. Gippert, N. P. Himmelmann, & U. Mosel (Eds.), *Essentials of language documentation* (pp. 129–162). Mouton de Gruyter.
- Himmelmann, N. P. (1998). Documentary and descriptive linguistics. *Linguistics*, 36(1), 161–195.
- Himmelmann, N. P. (2006). Language documentation: What is it and what is it good for? In J. Gippert, N. P. Himmelmann, & U. Mosel (Eds.), *Essentials of language documentation* (pp. 1–30). Mouton de Gruyter.
- Levelt, W. J. M. (1989). *Speaking: From intention to articulation*. MIT Press.
- Majid, A., Roberts, S. G., Cilissen, L., Emmorey, K., Nicodemus, B., O'Grady, L., ... Levinson, S. C. (2018). Differential coding of perception in the world's languages. *Proceedings of the National Academy of Sciences (PNAS)*, 115(45), 11369–11376.
- Mauerberg-deCastro, E., Moraes, R., Tavares, C. P., Figueiredo, G. A., Pacheco, S. C. M., & Costa, T. D. A. (2014). Haptic anchoring and human postural control. *Psychology & Neuroscience*, 7(3), 301–318.
<https://doi.org/10.3922/j.psns.2014.045>
- Talmy, L. (1972). *Semantic structures in English and Atsugewi* [Unpublished doctoral dissertation]. University of California, Berkeley.
- Talmy, L. (1985). Lexicalization patterns: Semantic structure in lexical forms. In T. Shopen (Ed.), *Language typology and syntactic description: Vol. 3. Grammatical categories and the lexicon* (pp. 57–149). Cambridge University Press.

- Talmy, L. (1988). Force dynamics in language and cognition. *Cognitive Science*, 12(1), 49–100.
- Turkish Patent and Trademark Office. (2025). *MumuLab* (Trademark Application No. 2025/152264).
- Woodbury, A. C. (2003). Defining documentary linguistics. *Language Documentation and Description*, 1, 35–51.
- Woodbury, A. C. (2011). Language documentation. In P. K. Austin & J. Sallabank (Eds.), *The Cambridge handbook of endangered languages* (pp. 159–186). Cambridge University Press.
- Young, R. W., & Morgan, W. (1987). *The Navajo language: A grammar and colloquial dictionary* (2nd ed.). University of New Mexico Press.

ÇEVİRİM İÇİ DİLBİLİM ARAŞTIRMALARI İÇİN PCIBEX FARM: TASARIM ve UYGULAMA

Oktay ÇINAR

Dr. Öğr. Üyesi, İstanbul Medeniyet Üniversitesi, Dilbilimi Bölümü, oktay.cinar@medeniyet.edu.tr, ORCID: 0000-0002-9822-7574

Çınar, O. (2025). Çevrim içi dilbilim araştırmaları için PCibex Farm: Tasarım ve uygulama. İçinde O. Çınar, F. Başbuğ & H. Aydemir (Ed.), *Çağdaş dilbilim yazıları I: İMÜ Dilbilimi 15. yıl armağanı* (ss. 81-108). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X004>

ÖZ

Son yıllarda deneysel dilbilim ve psikodilbilim alanlarında çevrim içi veri toplama yöntemlerine olan ilgi belirgin biçimde artmıştır. Bu eğilim, katılımcı erişimini kolaylaştıran ve deneysel süreçleri daha esnek hale getiren web-tabanlı platformların yaygınlaşmasıyla yakından ilişkilidir. Bu çalışma, PCibex platformu ve PennController eklentisinin deneysel dilbilim araştırmalarında sunduğu olanakları kapsamlı biçimde tanıtmayı amaçlamaktadır. Öncelikle, PCibex kullanılarak yürütülebilen başlıca deney türleri alanda yapılan çalışmalarla örneklendirilmektedir. Bu genel çerçevenin ardından, platformun somut kullanımını göstermek amacıyla örnek bir öz-denetimli okuma deneyi ayrıntılı biçimde sunulmaktadır. Uygulama kapsamında, Türkçede içe yerleşik tümcelerde boş ve açık özne adılarının öncüllerine gönderiminin işlenmesinin çevrim içi ortamda nasıl test edilebileceği gösterilmiştir. Deney tasarımı süreci; bilgilendirme ve gönüllü katılım onamı, yönergeler, katılımcı bilgi formu, alıştırma aşaması, ana test ve veri kaydı adımlarıyla sistematik biçimde ele alınmıştır. Ayrıca PennController tarafından üretilen olay-temelli sonuç dosyasının yapısı açıklanmış ve verilerin R ortamına aktararak betimleyici istatistikler aracılığıyla nasıl çözümlenebileceği gösterilmiştir. Çalışma, PCibex'in çevrim içi deneysel dilbilim araştırmaları için erişilebilir, yeniden üretilebilir ve ölçeklenebilir bir altyapı sunduğunu ortaya koymaktadır.

Anahtar Sözcükler: PCibex, PennController, Deneysel dilbilim, Öz-denetimli okuma, Çevrim içi deneyler

ABSTRACT

In recent years, experimental linguistics and psycholinguistics have increasingly adopted online data collection methods. This trend is closely related to the widespread use of web-based platforms that facilitate participant recruitment and enable flexible experimental designs. This study aims to provide a comprehensive overview of the PCibex platform and the PennController extension in experimental linguistic research. First, the paper outlines and exemplifies the main types of studies that can be conducted

using PCIBex. Building on this overview, the study then presents a detailed example application to illustrate the practical use of the platform. Specifically, a self-paced reading experiment is introduced to demonstrate how the processing of the coindexation of null and overt subject pronouns in embedded clauses with their antecedents in Turkish can be investigated in an online environment. The experimental workflow is described step by step, covering informed consent, instructions, participant information forms, practice trials, the main test phase, and data storage. Special attention is given to the event-based structure of the results file generated by PennController and to the procedures for transferring and summarizing the data in the R environment using descriptive statistics. Overall, the study shows that PCIBex and PennController offer an accessible, reproducible, and scalable framework for conducting a wide range of online experimental studies in linguistics.

Keywords: PCIBex, PennController, Experimental linguistics, Self-paced reading, Online experimentation

1. Giriş

Son yıllarda dilbilim araştırmalarında deneysel yöntemlerin kullanımı önemli bir artış göstermektedir. Bu yöntemler, dilin nasıl işlendiğini, üretildiğini ve algılandığını anlamak amacıyla tasarlanan kontrollü çalışmalar aracılığıyla, dilsel süreçlerin bilişsel yönlerini inceleme olanağı sunmaktadır. Öz-denetimli okuma, kabul edilebilirlik yargı testleri, sözcüksel karar verme ve işitme temelli testler bu deneysel yaklaşımlara örnek olarak gösterilebilir. Söz konusu bu çalışmalar genellikle bilgisayar ortamında yürütülmekte, katılımcıların tepkileri özel yazılımlar aracılığıyla ölçülmekte ve analiz edilmektedir. Bu yaklaşım, araştırmacılara yüksek düzeyde kontrol imkanı sağlayarak güvenilir ve tekrarlanabilir sonuçlar üretme açısından önemli avantajlar sunmaktadır. Ancak bu avantajlara karşın, söz konusu çalışmalar katılımcı erişimi ve zaman yönetimi bakımından belirli sınırlılıklara sahiptir. Veri toplamak için kullanılan dilbilimsel deney araçları, genellikle aynı anda yalnızca tek bir katılımcı tarafından kullanılabilirliğinden, veri toplama süreci ardışık biçimde yürütülmektedir. Bu durum, çok sayıda katılımcıdan veri toplanması gerektiğinde sürecin oldukça zaman almasına yol açmaktadır. Ayrıca, bu çalışmaların belirli fiziksel koşullarda yürütülmesi gerekliliği, araştırmacıların daha geniş ve çeşitlendirilmiş katılımcı gruplarına erişimini de sınırlandırmaktadır.

Bunun dışında, COVID-19 pandemisi ile birlikte bu sınırlılıklar daha da belirgin hale gelmiş ve yüz yüze veri toplama süreçleri büyük ölçüde güçleşmiştir. Bu gelişme, araştırmacıları çevrim içi ortamda deney yürütmeye yönlendirmiştir. Bu gibi ihtiyaçlara yanıt olarak, web-tabanlı deney platformları, son yıllarda daha fazla kullanılmakta ve daha geniş bir katılımcı yelpazesiyle veri toplama olanağı sağlamaktadır. Bu platformlar arasında öne çıkan PCIBex (<https://farm.pcibex.net/>) tarayıcı üzerinden erişilebilen, çeşitli dilbilimsel çalışmaların yürütülebileceği, açık kaynaklı ve ücretsiz web-tabanlı bir deney platformudur. Bu açıdan, PCIBex, çeşitli

dilbilimsel deneylerin web-tabanlı yürütülebilmesi açısından araştırmacılar için önemli kolaylıklar sunmaktadır.

Bu çalışmada, PCIBex platformunun temel altyapısı, dilbilimsel deney tasarımı süreçleri ve dilbilim alanındaki uygulamaları incelenecektir. Ayrıca, platformun hangi tür dilbilimsel çalışmalar için kullanılabilceği, araştırmacılara sunduğu yenilikçi olanaklar tartışılacaktır. Bunun dışında çevrim içi deney yürütmenin getirdiği sınırlılıklara da değinilecektir.

Çalışmanın genel yapısı şu şekildedir. 2. bölümde, PCIBex platformunun tarihçesi, genel özellikleri ve getirdiği yenilikler ele alınacaktır. 3. bölümde, platform üzerinden yürütülebilecek dilbilimsel deney türleri, önceki çalışmalardan örnekler ve uygulama alanları tanıtılacaktır. 4. bölümde, PCIBex kullanılarak geliştirilen örnek bir deney tasarımı ayrıntılı biçimde sunulmakta; deneyin kurulum aşamaları, katılımcı yönetimi ve uygulama ve veri çözümleme süreci adım adım açıklanmaktadır. 5. ve son bölümde ise PCIBex'in deneysel dilbilim araştırmalarındaki güçlü ve zayıf yönleri, sunduğu avantajlar ve çevrim içi deneylere özgü sınırlılıklar bütüncül bir değerlendirme çerçevesinde ele alınacaktır.

2. PCIBex: PennController, Arayüz ve Kullanım Özellikleri

2.1 PennController

PCIBex (internet tabanlı deneyler için PennController)^{1 2}, geniş ölçekli davranışsal deneylerin çevrimiçi olarak yürütülmesi amacıyla tasarlanmış, ücretsiz bir web tabanlı platformdur. Ibex, ilk olarak 2007 yılında Alex Drummond tarafından geliştirilmiş; 2018 yılında ise Jeremy Zehr ve Florian Schwarz tarafından PennController adlı bir kütüphane üzerine inşa edilerek genişletilmiş ve böylece Ibex'in bir uzantısı hâline gelmiştir (Zehr & Schwarz, 2018).

Platform, 2010'dan 30 Eylül 2021'e kadar <https://spellout.net/ibexfarm> adresinde çalışmıştır. 31 Aralık 2020'den 30 Eylül 2021'e kadar ise yeni bir Ibex Farm sürümü <https://ibex.spellout.net> adresi üzerinden hizmet vermiştir. Güncel ve aktif web sitesi ise şu linktedir: <https://farm.pcibex.net/>

PennController, bir JavaScript tabanlı eklenti kütüphanesidir ve açık kaynak kodludur. Bu eklenti, Ibex'in sınırlı deney yapısını (özellikle metin tabanlı görevleri) genişleterek, görsel, işitsel ve etkileşimli dilsel deney

¹“PennController” sözcüğünde “Penn” Pennsylvania Üniversitesi'nin kısaltması olarak kullanılmaktadır. İlk olarak Pennsylvania Üniversitesi tarafından desteklendiği için bu adla kısaltılmıştır. Dolayısıyla “Penn Denetleyici” olarak da Türkçeleştirilebilir; ancak çalışmada özel isim olarak kabul edilip orijinal adı ile kullanılacaktır.

² “PCIBex” kavramı ise “PennController for Internet-Based Experiments” (internet tabanlı deneyler için PennController) sözcüğündeki sözcüklerin ilk harfleri kısaltılarak oluşturulmuş bir sözcüktür.

tasarımlarının kolayca tasarlanmasını sağlar. Bu sayede, araştırmacının basit sözdizimiyle karmaşık deneyleri oluşturmaya olanak sağlar.³

PCIBex arayüzünde, `PennController` komutlarıyla deney nesneleri (ör. `text`, `image`, `audio`, `button`, `keypress`) tanımlanmaktadır. Aşağıda *PennController* kütüphanesi kullanılarak oluşturulmuş basit bir deney yapısı gösterilmektedir:

```
PennController.ResetPrefix(null);

newTrial("talimat",
  newText("Bu bir okuma testidir. Boşluk tuşuyla ilerleyin.")
    .print()
  ,
  newKey(" ")
    .wait()
);

Template("deneyListesi.csv", row =>
  newTrial("okuma",
    newController("DashedSentence", { s: row.cumle })
      .print()
      .log() // her kelimenin gösterilme süresi kaydedilir
      .wait()
    )
  );
```

Şekil 1. Örnek bir deney yapısı

`PennController.ResetPrefix(null)` komutu, kodun daha sade biçimde yazılmasına olanak tanır. `newTrial()` her bir deney aşamasını (örneğin talimat, alıştırma veya uyarın gösterimi) tanımlar. `newText()` katılımcıya yazılı bir bilgilendirme sunarken, `newKey()` nesnesi katılımcının belirli bir tuşa (bu örnekte boşluk tuşu) basmasını tanımlar. Asıl deney kısmında ise benzer bir yapıyla `newController("DashedSentence", {...})` nesnesi kullanılarak, bu örnekte cümleler öz-denetimli okuma biçiminde, yani sözcük sözcük sunulur. Her sözcüğe geçişteki süre, katılımcının okuma tepki süresi (reaction time) olarak kaydedilir. `.log()` komutu, katılımcının yaptığı işlemlere ait verileri kaydetmek için kullanılır. Her deneme sonunda bu veriler PCIBex'in sonuç dosyasına (.csv) otomatik olarak yazılır.

Yukarıda da görüldüğü gibi, standard Javascript dilinin kullanan PCIBex dili `Penncontroller` komutlarıyla basittir. Bunlar dışında daha ileri seviye özelliklerle, iştirme testleri ve göz izleme (Papoutsaki et al., 2016) gibi uygulama programlama özelliklerine de sahiptir.

2.2 Arayüz ve kullanım özellikleri

Araştırmacılar, platformun “Script Editor” (Kod Düzenleyici) bölümünde hazır şablonları kullanabilir veya kendi deney kodlarını sıfırdan oluşturabilirler. Arayüzün ekran görüntüsü Şekil 2’de sunulmaktadır. Yapılan tüm değişiklikler sistem tarafından otomatik olarak kaydedilmektedir. Deneyin ön izlemesi, “Refresh” (Yenile) ile aynı sekmede veya “Open in new

³ <https://doc.pcibex.net/> bağlantısında deneylerin nasıl yürütüleceği, komutlarının nasıl kullanıldığı gibi temel kavramlar ve temel tutorial’a ulaşılabilir.

tab” (Yeni sekmede aç) seçeneğiyle yeni bir sekmede yapılabilir. “Resources” (Kaynaklar) sekmesinde araştırmacı, deneye ait dosyaları yükleyebilir (örn., uyarıların yer aldığı .csv tabloları, ek HTML sayfaları veya ses ve görsel dosyaları). “Scripts” (Kodlar) sekmesi yeni JavaScript dosyaları eklemeye, “Aesthetics” (Görsel Düzen) sekmesi ise deneyin tasarımını biçimlendirmeye olanak tanır. PennController kütüphanesi ise, “Modules” (Modüller) sekmesinde yer alır ve deneyin yürütülmesini sağlayan temel bileşendir.



Şekil 2. PCIBex arayüzü

Kod düzenleyicinin sağ kısmında yer alan “Actions” (Eylemler) paneli deney yönetimi için çeşitli işlevler içerir. “Results” (Sonuçlar) sekmesi, deneyin çıktı dosyasını ön izleme veya indirme olanağı sağlar. “Share” (Paylaş) sekmesi, deney için otomatik olarak bir paylaşım bağlantısı (link) oluşturur. Araştırmacılar, PennController sözdizimiyle oluşturdukları deneyleri platformun arayüzü üzerinden kaydeder ve doğrudan web tabanlı bir deney bağlantısı olarak yayımlayabilirler. Bu bağlantı, deneyin test edilmesi, hata ayıklanması veya yayımlanan bir çalışmanın tekrarlanabilirliğini sağlamak amacıyla başkalarıyla paylaşılabilir. Veri toplama sürecine geçildiğinde, katılımcılara veri toplama bağlantısı gönderilir ve bu bağlantı üzerinden gelen sonuçlar sistemde ayrı olarak kaydedilir. Katılımcılar deney bağlantısını açtıklarında, herhangi bir ek yazılım ya da eklenti yüklemeye gerek kalmadan, tarayıcı üzerinden doğrudan deneye katılabilirler.

“Download” (İndir) seçeneği, deney projesinin tamamını yerel bilgisayara indirme olanağı sunar. PCIBex, her katılımcının yanıtlarını otomatik olarak kaydeder. Tüm veriler sistemde oturum bazında depolanır ve araştırmacı tarafından daha sonra .csv biçiminde indirilebilir. Elde edilen bu ham veriler, istatistiksel analiz yazılımlarında veya programlama dillerinde (ör. R, Python) çözümlenebilir.

“Git Sync” seçeneği, mevcut bir Git deposunu projeye entegre etmeyi sağlar. “Settings” (Ayarlar) sekmesi ise deneye açıklama ekleme ve projenin kopyalanabilir olup olmadığını belirleme seçeneklerini içerir.

Sonuç olarak, PCIBex Farm araştırmacıya deneyin tüm aşamalarını tek bir çevrim içi arayüz üzerinden yönetme olanağı sunar. Bu yönüyle platform, çevrim içi dilbilimsel ve psikodilbilimsel deneylerde erişilebilir, yeniden üretilebilir ve paylaşıma açık bir araştırma ortamı sağlar.

3. Deneyisel Dilbilim Bağlamında PCIBex Kullanımı

Bu bölümde, platform üzerinden yürütülebilecek dilbilimsel deney türleri, önceki çalışmalardan örnekler ve uygulama alanları tanıtılmaktadır.

PCIBex, tarayıcı tabanlı yapısı sayesinde çevrim içi ortamda başta psikodilbilimsel olmak üzere çok sayıda deneyin yürütülmesine olanak tanır. Yazılı, görsel veya işitsel uyarıların gösterimi ve katılımcı tepkilerinin (tuş basışı, süre, seçim vb.) otomatik kaydını desteklediği için de farklı alanlarda deneyisel testlerin kolayca uygulanmasını mümkün kılar. Alanyazında en çok kullanılan deney türleri aşağıda alt başlıklar halinde anlatılmaktadır.

3.1 Öz-Denetimli okuma testleri

Katılımcı, deney sırasında sözcükleri *tek tek*, öbekleri ya da tümcecikleri *parça parça* ya da cümleleri *bir bütün olarak* okur ve her bir tuş basışı (genellikle boşluk tuşu olarak atanır) arasındaki süre otomatik olarak kaydedilir. Bu süreler, katılımcının okuma hızı ve işleme sürecine ilişkin temel ölçütleri oluşturur. Bu tür bir sunum tekniği, PennController’ın DashedSentence modülü kullanılarak uygulanır ve çevrim içi okuma deneylerinde yaygın şekilde tercih edilir.

Bunun yanı sıra, katılımcıların okudukları ifadelerle ilgili anlama düzeylerini veya karar verme süreçlerini ölçmek amacıyla farklı modüller de kullanılabilir. Örneğin, `newText()` ve `newScale()` modülleri aracılığıyla doğru–yanlış gibi iki seçenekli yargı soruları oluşturulabilir; katılımcıların F veya J gibi belirli tuşlarla yanıt vermesi gereken testler ise `newKey()` ya da `newController("Question")` modülüyle düzenlenir. Bu araçlar, araştırmacıların hem dil işleme süreçlerini hem de yargı hızlarını aynı deney içinde bütüncül olarak ölçmelerine olanak tanır.

Bu tür zamana duyarlı testler PCIBex’te en sık kullanılan deney araçlarından biridir ve bugüne kadar çok sayıda çalışma bu platform üzerinden yürütülmüştür. Özellikle son yıllarda çevrimiçi veri toplama yöntemlerinin yaygınlaşmasıyla birlikte PCIBex kullanılarak gerçekleştirilen araştırmaların sayısı önemli ölçüde artmıştır. Bu araştırmalara Kim & Kim (2025); D’Alesio vd. (2025); Kim & Hwang (2025); Cairncross vd. (2024); Park & Lee (2024); Britton vd. (2024); Lu (2024); Yamaguchi & Ohta (2023); Moulton vd. (2022); Çakır & Çınar (under review); Çınar (2024); Uygun (2022); Gračanin-Yuksek vd. (2020); Gračanin-Yuksek vd. (2017) gibi çalışmalar örnek olarak verilebilir.

3.2 Kabul edilebilirlik yargı testleri

PCIBex'te kullanılan bir diğer veri toplama yöntemi kabul edilebilirlik yargı testleridir. Kabul edilebilirlik yargısı testleri, laboratuvar ortamına ihtiyaç duyulmadan uygulanabilen; hatta basit bir kağıt-kalem formatında dahi yürütülebilen yöntemlerdir. Bununla birlikte PCIBex platformu, bu tür testlerin çevrimiçi, kontrollü ve katılımcı dostu bir biçimde uygulanmasını önemli ölçüde kolaylaştırmaktadır. Bu testlerde katılımcılardan, verilen bir cümleğin dilbilgisel ya da anlamsal açıdan ne kadar doğal olduğunu değerlendirmeleri beklenir. Değerlendirme genellikle iki tür yanıt formatından biriyle yapılır:

- (i) Likert ölçeği (örneğin 1–7: “1 = hiç kabul edilebilir değil” → “7 = tamamen kabul edilebilir”),
- (ii) İkili karar (binary judgement): “Kabul edilebilir / Kabul edilemez”.

PCIBex'te ikili karar temelli kabul edilebilirlik yargıları, farklı modüller aracılığıyla yapılandırılabilir. Bu tür görevler `newScale()` kullanılarak görsel seçenekler üzerinden sunulabileceği gibi, katılımcıların klavyedeki belirli tuşlara (ör. F = kabul edilebilir, J = kabul edilemez) basarak yanıt verdikleri biçimde de tasarlanabilmektedir. Bu ikinci yaklaşımda, `newKey()` ya da `newController("Question")` denetleyicisi kullanılarak yanıtlar otomatik olarak kaydedilmektedir.

Kabul edilebilirlik yargı testleri oluşturulurken `Template()` işlevi, cümleler, koşullar ve diğer deneysel bilgilerin bulunduğu .csv dosyasındaki her bir maddeyi okuyarak deneme setinin temelini oluşturur. Ardından `newTrial()`, her bir yargı denemesini tanımlar ve deney akışının yapısını belirler. Katılımcıya gösterilecek cümleler ile gerekiyorsa kısa yönergeler `newText()` komutu aracılığıyla ekrana yansıtılır; yargılamanın kendisi ise az önce bahsedildiği gibi `newScale()/newKey()` ya da `newController("Question")` denetleyicisi aracılığıyla gerçekleştirilir.

Alan yazın incelendiğinde, kabul edilebilirlik yargı testlerinin çoğu zaman öz-denetimli okuma gibi başka deney türleriyle birlikte kullanıldığı görülmektedir. Yukarıda değinilen öz-denetimli okuma çalışmaları arasından Lu (2024), Park & Lee (2024), D'Alesio vd. (2025), Moulton vd. (2022) ve Çınar (2024), bu yöntemin farklı deneysel paradigmalara bir arada uygulanmasına örnek teşkil etmektedir.

3.3 Dinleme/İşitsel Uyarı Temelli Testler

PCIBex platformunda işitsel testler, katılımcılara ses uyarıları—örneğin sözcükler, cümleler veya en küçük çiftler—sunarak onların karar, algı ya da yargı temelli tepkilerini toplamaya dayanır. Bu tür testler özellikle sesbilgisi, sesbilim, işitsel algı ve cümle işleme alanlarında yaygın biçimde kullanılmaktadır. Kabul edilebilirlik yargı testlerinde olduğu gibi, katılımcılardan dinledikleri ses uyarısına ilişkin doğru/yanlış ya da uygun/uygun değil türü kararlar vermeleri istenebilir.

İşitsel testlerin genel işleyişi birkaç temel aşamadan oluşur. Öncelikle, ses dosyaların olarak uyarılar sisteme yüklenir. Ardından, katılımcı

cıya sunulan ses uyarısı—örneğin tek bir sözcük, bir cümle veya daha uzun bir ifade—çalıştırılır ve katılımcı bu uyarıyı dinler. Sonrasında, katılımcının tepkisini kaydetmek amacıyla bir düğmeye tıklaması, bir tuşa basması ya da bir ölçek üzerinden seçim yapması istenir. Son aşamada ise sistem, verilen yanıt ve tepki süresini otomatik olarak kaydeder.

Bu yapıyı uygulamak için PennController’ın `newAudio()` modülü ses dosyalarının (.wav veya .mp3 formatında) çalınmasını sağlar. Katılımcı tepkilerini toplamak için `newButton()` veya `newKey()` kullanılabilir; dereceli ya da kategorik değerlendirmelerin gerektiği durumlarda ise `newScale()` modülü tercih edilir. Bu tür işitsel karar görevlerinin kullanımına Plaut-Forckosh & Meir (2025) ile Ruda & Huang (2025) çalışmalarında örnekler bulunmaktadır.

3.4 Sözcüksel Karar Verme Testleri

Sözcüksel karar verme testleri (lexical decision tasks), katılımcılardan dinledikleri ya da gördükleri bir birimin—örneğin bir harf dizisi ya da ses dizisi—gerçek bir sözcük olup olmadığını mümkün olan en kısa sürede belirlemelerini ister. Bu testlerin temel amacı, sözcük tanıma sürecinin hızını ve bilişsel yükünü ölçmektir. Böylece şu tür soruların yanıtlanmasına katkı sağlar: Katılımcılar gerçek sözcükleri ne kadar hızlı tanır? Sözcük sıklığı, uzunluğu ve biçimbilimsel yapısı tanıma süresini nasıl etkiler? Hangi tür sözcüksel özellikler tanımayı kolaylaştırır ya da zorlaştırır?

PCIBex platformunda sözcüksel karar verme görevleri, görsel uyarıların sunumu için `newText()`, işitsel uyarıların sunumu için `newAudio()` modüllerini kullanarak yapılandırılır. Katılımcıların verdikleri yanıtları toplamak amacıyla `newKey()` veya `newButton()` komutları kullanılır. Her bir uyarının ekranda ne kadar süreyle kalacağını belirlemek ve deney akışını zamanlamak için ise `newTimer()` işlevinden yararlanılır. Bu yapı sayesinde deneysel koşullar arasında tutarlı, ölçülebilir ve karşılaştırılabilir tepki süreleri elde edilir.

Sözcüksel karar verme paradigması, psikodilbilim alanında yaygın biçimde uygulanmakta olup farklı dillerde, yazı sistemlerinde ve bilişsel süreçlerde sözcük tanımanın doğasını inceleyen pek çok çalışma tarafından kullanılmıştır. Bu çalışmalara Lee & Kaiser (2021); Hatzidaki & Santesteban (2024); Zhang vd. (2023); Liu vd. (2025); Ariga & Hirose (2025) gibi araştırmalar örnek olarak verilebilir.

3.5 Göz İzleme Testleri

PCIBex, donanım tabanlı bir göz izleme (eye-tracking) sistemi sunmamakla birlikte, çevrim içi ortamlarda sıklıkla kullanılan göz-hareketi benzeri ölçütlerin elde edilmesine imkan veren çeşitli araçlar ve modüller sağlamaktadır. Bu yaklaşımların amacı, katılımcının görsel bir alanı nasıl taradığını, hangi bölgelere ne kadar süre odaklandığını veya belirli bir uyarıyı ne zaman hedeflediğini dolaylı yollarla ölçmektir. Böylece klasik göz izleme çalışmalarında kullanılan ilk bakış süresi, toplam bakış süresi veya yeniden

yönelme gibi ölçümlerin çevrimiçi koşullar altında yaklaşık olarak modellenmesi mümkün hale gelir.

PCIBex bu tür görevleri üç farklı teknik üzerinden desteklemektedir:

- (i) görsel bölge tabanlı seçim ve süre ölçümleri,
- (ii) fare (mouse) hareketine dayalı dikkat benzeri izleme,
- (iii) WebGazer.js aracılığıyla webcam-tabanlı bakış tahmini.

Görsel bölge tabanlı uygulamalarda, ekranda yer alacak bölgeler `newImage()`, `newCanvas()` ve `newSelector()` komutlarıyla yapılandırılır; bu sayede nesnelerin konumları hassas biçimde ayarlanabilir. Katılımcının hangi bölgeyi ne zaman seçtiğini belirlemek için `newSelector()` sıklıkla kullanılır; çünkü her seçim zaman damgasıyla birlikte otomatik olarak kaydedilir. Fare hareketinin dikkat yönelimini dolaylı biçimde yansıtabildiği durumlarda `newMouseTracker()` modülü kullanılarak katılımcının imleç hareketleri zaman içinde izlenebilir ve bakış benzeri bir veri elde edilebilir.

Zamanlama gerektiren deneylerde `newTimer()` modülü, uyarıların ekranda ne kadar süreyle kalacağını ve deney akışının hangi hızda ilerleyeceğini kontrol eder. Bu sayede bir bölgeye yönelme hızı, bölgeler arası geçiş süreleri veya toplam bakış benzeri süreler hesaplanabilir.

PCIBex ayrıca, çoklu görsel nesneler ve bir sesli uyarı içeren görsel dünya (visual-world) türü deneyleri de destekler. Bu tür bir kurulumda görseller `newCanvas()` ile ekrana yerleştirilir, uyarı `newAudio()` ile sunulur ve katılımcının hangi nesneye yöneldiği, hangi nesneyi seçtiği ya da fare imlecini ilk olarak hangi bölgeye götürdüğü kaydedilir. Böylece çevrim içi bir görsel dünya analizi yapılabilir.

Son olarak, PCIBex'in WebGazer.js ile uyumu sayesinde webcam tabanlı deneyler de yürütülebilmektedir. Bu yöntem, tarayıcı içi kalibrasyonla çalışır ve katılımcının kısa süreli bakış yönelimlerini izlemeye olanak sağlar. Bu yöntemlerin kullanımına Slim vd. (2023) ve Ovans (2025) çalışmalardan örnek gösterilebilir.

Bununla birlikte, PCIBex üzerinde webcam-tabanlı göz izleme kullanımının belirli sınırlılıkları bulunmaktadır. Her şeyden önce, WebGazer.js gibi tarayıcı tabanlı izleme teknolojileri, laboratuvar ortamında kullanılan Tobii veya EyeLink gibi yüksek örnekleme hızına sahip aygıtlara kıyasla hem zaman çözünürlüğü hem de konumsal hassasiyet açısından daha düşük performans göstermektedir. Özsoy vd. (2023), PCIBex/WebGazer ile yürütülen çevrim içi deneylerle karşılaştırıldığında laboratuvar ortamında kullanılan Tobii sisteminin daha güvenilir sonuçlar ürettiğini göstermiştir. Benzer biçimde Patterson vd. (2025), webcam tabanlı göz izleme kullanımının doğal büyüklükte metin okuma gibi göz hareketi analizi gerektiren deneylerde uygun olmadığını belirtmektedir. Ek olarak, webcam-tabanlı izlemenin başarılı olabilmesi katılımcının donanım özelliklerine, ekran çözünürlüğüne, aydınlatma koşullarına ve kalibrasyon doğruluğuna büyük ölçüde bağlıdır. Patterson vd. alandaki çalışmaların önemli bir kısmında minimum donanım gereksinimleri ile kalibrasyon prosedürlerinin yeterli ayrıntıyla rapor-

lanmadığını, bunun da deneylerin tekrarlanabilirliğini zayıflattığını vurgulamaktadır.

4. Örnek bir Deney Tasarımı ve Uygulama Süreci: Öz-Denetimli Okuma Testi

Bu bölümde, PennController eklentisi kullanılarak geliştirilen çevrim içi bir öz-denetimli okuma deneyinin tasarım ve uygulama süreci ayrıntılı biçimde sunulmaktadır. Deney, katılımcı etkileşimini aşamalı olarak yöneten ve belirli görevleri yerine getiren bir dizi modülden oluşmaktadır. Aşağıda, deneyin temel bileşenleri ve bu bileşenlerin işlevleri sırasıyla açıklanmaktadır. Deneyin çevrim içi uygulamasına ise şu adresten erişilebilir: <https://farm.pcibex.net/p/fMvTSf/>

Deneye ait .csv veri dosyaları, deney kodları ve ilgili materyaller ise Açık Bilim Çerçevesi (OSF) üzerinden şu adreste paylaşılmaktadır: https://osf.io/uyq82/overview?view_only=a41eca13771149e39be8260f8d7cbec24

4.1 Çalışma içeriği ve veri hazırlama

Örnek verilecek bu çalışma, Türkçede içe yerleşik tümcelerde boş ve açık özne adıllarının ana tümcedeki öncüllere (AÖ Özne, ne-sorusu özne, niceleyici özne) gönderiminin nasıl işlemlendiğini, bağlamli cümleler aracılığıyla incelemektedir. Deney kapsamında katılımcılara her denemede önce kısa bir bağlam sunulmakta, ardından bu bağlamla ilişkili hedef cümle ekrana getirilmektedir. Hedef cümle okunduktan sonra ise katılımcılar adılın gönderimine yönelik Evet ya da Hayır tuşlarından birini seçmektedir. Örnek bir deneme aşağıda verilmiştir:

- (1) **Bağlam:**
Tüm öğrenciler bilgi yarışmasına çok iyi hazırlandılar.
Hedef cümle:
Onlar birinci olacaklarını düşünüyorlar.
Evet (Özne kendisi hakkında konuşuyor)
Hayır (Özne başkası hakkında konuşuyor)

Katılımcılar, bağlamı okuduktan sonra hedef cümleyi sözcük sözcük ilerleyen bir öz-denetimli okuma düzeninde okumaktadır. Çalışmanın temel amacı, hedef cümlede yer alan kritik sözcüklerin (boş ya da açık adıldan sonraki ilk iki sözcük) işlemlenmesine ilişkin tepki sürelerinin bağlama bağlı olarak nasıl farklılaştığını incelemektir. Bu doğrultuda, katılımcıların kritik sözcükte ve izleyen bölgelerde gösterdikleri okuma süreleri karşılaştırmalı olarak çözümlenmektedir.

Bu aşamada, deneyde kullanılacak bağlam metinleri ve hedef cümlelerin, bunlara karşılık gelen deneysel tür etiketleriyle birlikte bir .csv dosyasında tanımlanması gerekmektedir. Çalışma kapsamında hazırlanan bu dosya, virgülle ayrılmış sütunlardan oluşmakta olup her bir satır bir deney maddesine karşılık gelmektedir. İlgili .csv dosyasında yer alan temel sütunlar sırasıyla group (deneysel grup), item (soru/uyaran numarası), context (bağlam

metni), sentence (hedef cümle) ve type (soru/koşul türü) biçiminde yapılandırılmıştır. Bu düzenleme, deney sırasında her bir uyarının sistematik biçimde sunulmasını ve sonuç dosyasında koşul-temelli çözümlemelerin tutarlı biçimde yürütülmesini sağlamaktadır.

4.2 Bilgilendirme ve gönüllü katılım onamı

Katılımcı içeren deneysel çalışmalarda etik kurul onayının alınması ve katılımcıların araştırmaya gönüllü olarak katıldıklarını açık biçimde beyan etmeleri zorunludur. Bu gereklilik, çevrim içi ortamda yürütülen deneyler için de geçerlidir. Bu nedenle, katılımcı onamının veri toplama süreci başlamadan önce, çalışmanın en başında alınması gerekmektedir.

Bu çalışma kapsamında deney, katılımcılara araştırmanın amacı, kapsamı ve etik ilkeleri hakkında bilgi veren bir gönüllü katılım onam sayfası ile başlatılmıştır. Söz konusu bilgilendirme aşamasında, çalışmanın akademik bir araştırma olduğu, katılımcılardan toplanan verilerin anonim biçimde saklanacağı ve bu verilerin yalnızca bilimsel amaçlarla kullanılacağı açıkça belirtilmiştir. Katılımcıların deneye devam edebilmeleri için, onamlarını açık ve bilinçli bir şekilde vermeleri zorunlu tutulmuştur.

Bilgilendirme ve onam sürecine ilişkin kullanılan örnek kod aşağıda sunulmaktadır. Verilen örnek kodlarda çalışmanın içeriği daha kısa verilmiş olup, ayrıntılı açıklamalar OSF üzerinde paylaşılan deney dosyalarında yer almaktadır.

```
newTrial("notice",
  newText("consentText", `
    <h2 style="text-align:center;">Gönüllü Katılım Onay Formu</h2>

    <p><b>"The Effect of Antecedents in Processing of Pronominal Binding in
    Turkish"</b>
    başlıklı çalışma, Necmettin Erbakan Üniversitesi ve İstanbul Medeniyet
    Üniversitesi akademik
    personellerinin yürüttüğü bir araştırmadır.</p>

    <ol>
    <li>Çalışmanın amacı ve süresi hakkında bilgilendirildim.</li>
    <li>Katılımımın gönüllü olduğunu ve dilediğim zaman çekilebileceğimi
    biliyorum.</li>
    </ol>
  `).print(),

  newButton("consentBtn", "ÇALIŞMAYA KATILMAYA ONAY VERİYORUM")
    .print()
    .wait()
);
```

Kullanılan modüllerin işlevleri aşağıda açıklanmaktadır:⁴

```
newTrial("notice", ...)
```

⁴ Deney kodunda yer alan ve modüler yapıların dışında kalan bazı bölümler, metnin görsel sunumuna ilişkin teknik ayarları içermektedir. Bu ayarlar; metinlerin punto büyüklüğü, yazı tipi, hizalama, kenar boşlukları ve benzeri biçimsel özellikleri düzenlemeye yöneliktir. Söz konusu düzenlemeler, deneyin kuramsal yapısı ya da işlevsel akışıyla doğrudan ilişkili değildir. Bu nedenle, çalışmanın yöntemsel açıklaması kapsamında yalnızca deneyin işleyişini ve bilişsel süreçleri etkileyen modüller ele alınmış; biçimsel sunuma ilişkin bu tür teknik ayrıntılar ayrıntılı olarak tartışılmamıştır. İlgili ayarlar, genel JavaScript kullanımında yaygın olarak bilinen ve PennController dokümantasyonunda standart biçimde sunulan uygulamalara dayanmaktadır.

Bu komut bir deneme/sayfa oluşturmaktadır. Söz konusu sayfa, deneyin ilk aşaması olarak yapılandırılmış olup yalnızca bilgilendirme ve gönüllü katılım onamı sürecine ayrılmıştır.

```
newText("consentText", ...)
```

Katılımcılara sunulan onam metninin ekrana yazdırılması amacıyla kullanılmıştır. Çalışmanın amacı, gizlilik ilkeleri, gönüllülük esasları ve katılımcı haklarına ilişkin tüm açıklamalar bu modül içerisinde tanımlanmıştır.

```
newButton("consentBtn", "...")
```

Katılımcının çalışmaya katılmayı kabul ettiğini belirtmesi için kullanılan onay butonunu ve ilgili metni oluşturmaktadır. Bu buton, katılımcının bilinçli onam verdiğini açık biçimde işaretlemesini sağlamaktadır.

```
.print()
```

Tanımlanan onay butonunun ekranda görünür hale getirilmesini sağlamaktadır.

```
.wait()
```

Katılımcı ilgili butona tıklayana kadar sayfayı bekleme durumunda tutmaktadır. Bu işlev sayesinde, katılımcının açık onam vermeden deneyin bir sonraki aşamasına geçişi mümkün olmamaktadır.

4.3 Yönergeler

Gönüllü katılım onam sayfasını takiben, deneyin nasıl yürütüleceğini ayrıntılı biçimde açıklayan bir yönerge ekranı katılımcılara sunulmuştur. Bu bölümde katılımcılara, deneyin yalnızca bilgisayar üzerinden gerçekleştirileceği; deney sırasında cümlelerin maskeleyme (moving-window technique; Just vd. 1982) yöntemiyle sunulacağı; ilerlemek için her aşamada boşluk (space) tuşunun kullanılacağı ve okuma tamamlandıktan sonra Evet (F tuşu ile)/Hayır (J tuşu ile) seçeneklerinden birinin seçilmesinin beklendiği açıkça belirtilmiştir. Yönergeler, örnek bağlamlar ve örnek cümleler üzerinden ayrıntılı biçimde açıklanarak katılımcıların deneyi doğru şekilde anlamaları sağlanmıştır. Yönerge sayfası oluşturmak için kullanılan örnek kod aşağıda sunulmuştur:

```
newTrial("intro",
  newText("instructions", `
    <h2 style="text-align:center;">Teste Başlamadan Önce Lütfen Okuyunuz</h2>

    <p>Bu çalışma yalnızca <b>bilgisayar</b> üzerinden yürütülmektedir.</p>

    <p>Cümleler <b>maskeleyme (moving-window)</b> yöntemiyle sunulacaktır.
    İlerlemek için her zaman <kbd>BOŞLUK</kbd> (Space) tuşuna basınız.</p>

    <p>Cümle tamamlandığında:</p>
    <ul>
      <li><b>Evet</b> için <kbd>F</kbd> tuşuna basınız.</li>
      <li><b>Hayır</b> için <kbd>J</kbd> tuşuna basınız.</li>
    </ul>

    <p>Lütfen cümleleri normal hızınızda okuyunuz ve testi ara vermeden
    tamamlayınız.</p>
  `).print(),

  newButton("continue", "Devam Et")
    .print())
```

```
.wait()
);
```

Yönerge ekranının oluşturulmasında kullanılan modüller ve işlevleri ise aşağıda sunulmaktadır:

```
newTrial("intro", ...)
```

Yönergeler için deney akışı içerisinde ayrı bir sayfa/deneme oluşturmaktadır. Bu deneme, onam sayfasının hemen ardından yer almakta olup herhangi bir deneysel veri toplamamaktadır.

```
newText("instructions", ...)
```

Yönergelerin tamamının ekrana yazdırılması amacıyla kullanılmaktadır. Görev tanımı, okuma biçimi ve tuşların kullanımıyla ilgili tüm açıklamalar bu modül içerisinde tanımlanmaktadır. Sunulan metin yalnızca ekranda görüntülenmekte olup sonuç dosyasına kaydedilmemektedir.

```
newButton("continue", ...)
```

Katılımcının yönergeleri okuduğunu ve deneye devam etmeye hazır olduğunu belirtmesini sağlayan butonu oluşturmaktadır.

```
.print()
```

Tanımlanan butonun ekranda görünür hale getirilmesini sağlamaktadır.

```
.wait()
```

Katılımcı butona tıklayana kadar sayfayı bekleme durumunda tutmaktadır. Bu işlev sayesinde katılımcıların bir sonraki aşamaya ancak yönergeleri onayladıktan sonra geçmeleri sağlanmaktadır.

4.4 Katılımcı bilgi formu

Yönergelerin ardından katılımcılardan kısa bir katılımcı bilgi formu doldurmaları istenmiştir. Bu form aracılığıyla katılımcıların ad-soyad, yaş ve ana dili demografik bilgileri toplanmıştır.

Girilen bilgiler, deney süresince her denemeye ait verilerle ilişkilendirilmek üzere kaydedilmiştir. Veri girişinin eksiksiz ve tutarlı biçimde sağlanabilmesi amacıyla form alanlarına belirli kısıtlamalar uygulanmıştır. Buna göre, katılımcıların ad ve soyad alanına en az iki sözcükten oluşan metin girmeleri; yaş alanına yalnızca sayısal değer girmeleri; ana dili alanına ise metin girdisi sağlamaları zorunlu tutulmuştur. Örnek kod aşağıda verilmiştir:

```
newTrial("form",
  newText("t1", "Adınız ve Soyadınız").print(),
  newTextInput("fullname")
    .length(64)
    .print(),

  newText("t2", "Yaşınız").print(),
  newTextInput("age")
    .length(4)
    .print(),

  newText("t3", "Ana Diliniz").print(),
  newTextInput("language")
    .length(64)
    .print(),
```

```

newButton("send", "Onayla ve İlerle")
    .print()
    .wait {
        getTextInput("fullname").test.text(/^\s*[a-zıİöüşçğ]+(\s+[a-
zıİöüşçğ]+)\s*$/i)
        .and(getTextInput("age").test.text(/^\s*\d+\s*$/))
        .and(getTextInput("language").test.text(/^\s*[a-zıİöüşçğ]+(\s+[a-
zıİöüşçğ]+)\s*$/i))
    },

newVar("fullname").global().set(getTextInput("fullname")),
newVar("age").global().set(getTextInput("age")),
newVar("language").global().set(getTextInput("language"))
);

```

Katılımcı bilgi formunun PennController ortamında çalıştırılmasında aşağıda açıklanan modüller kullanılmıştır:

```
newTrial("form", ...)
```

Demografik bilgi toplamak amacıyla deney akışında ayrı bir sayfa/deneme oluşturmaktadır. Bu deneme, yönergelerden sonra yer almakta ve test aşamasından önce sunulmaktadır.

```
newText (...)
```

Form alanlarına ait başlıkların (Ad-Soyad, Yaş, Ana Dil) ekrana yazdırılmasını sağlamaktadır. Bu metinler yalnızca görsel etiket işlevi görmektedir olup deney verisi olarak kaydedilmemektedir.

```
newTextInput("fullname" / "age" / "language")
```

Katılımcılardan metin girdisi almak için kullanılmaktadır. Her bir TextInput öğesi, farklı bir demografik değişkene karşılık gelmektedir.

```
.length (...)
```

Girdi alanlarına yazılabilecek maksimum karakter sayısını sınırlandırmaktadır.

```
.print()
```

Tanımlanan öğelerin ekranda görünür hale getirilmesini sağlamaktadır.

```
newButton("send", "Onayla ve İlerle")
```

Katılımcının formu tamamladığını ve bir sonraki aşamaya geçmek istediğini belirtmesini sağlayan “Onayla ve İlerle” butonunu oluşturmaktadır.

```
.wait (...)
```

Butonun yalnızca tüm alanlar geçerli biçimde doldurulduğunda etkinleşmesini sağlamaktadır. Bu amaçla düzenli ifadeler (*regular expressions, regex*) kullanılarak, ad-soyad alanında en az iki sözcük bulunması, yaş alanına sayısal bir değer girilmesi ve ana dili alanına metin girdisi sağlanması zorunlu tutulmuştur.

```
newVar (...).global().set (...)
```

Formdan elde edilen bilgileri global değişkenlere atamaktadır. Bu sayede katılımcı bilgileri, test denemeleri sırasında her satıra eklenerek sonuç dosyasında sistematik biçimde kaydedilmektedir.

4.5 Alıştırma aşaması

Ana test aşamasına geçilmeden önce, katılımcıların teste alışmalarını sağlamak amacıyla bir alıştırma bölümü uygulanmıştır. Alıştırma maddeleri deney maddeleriyle aynı yapıya sahip olmakla birlikte `.log()` kullanılmadığı için okuma süreleri ve yanıtlar kaydedilmemektedir. Bu aşama, katılımcıların sözcük sözcük okuma tekniğine ve yanıt verme biçimine uyum sağlamalarını amaçlamaktadır. Alıştırma aşamasını oluşturmak için kullanılan örnek kod aşağıda sunulmuştur:

```
Template("practice.csv", row =>

  newTrial("practice",
    newText("label", "Alıştırma").print(),

    newText("context", row.context).print(),
    newNexter("nextContext", getText("context").remove()),

    newController("DashedSentence", { s: row.sentence.trim() })
      .wait()
      .remove(),

    newController("Question", {
      as: [{"f", "Evet (F)"}, {"j", "Hayır (J)"}],
      randomOrder: false,
      presentHorizontally: true
    })
      .wait()
      .remove()
  )
);
```

Kullanılan modüller ve işlevleri ise aşağıda verilmiştir:

```
Template("practice.csv", row => ...)
```

Alıştırma uyarılarını `practice.csv` dosyasından satır satır okumakta ve her satır için bir deneme üretmektedir.

```
newTrial("practice", ...)
  practice etiketli bir alıştırma denemesi oluşturmaktadır.
```

```
newText("label", ...)
```

Denemenin alıştırma olduğunu belirten kısa bir etiket/metin sunmaktadır.

```
newText("context", row.context)
```

`.csv`'den gelen bağlam metnini ekranda görüntülemektedir.

```
newNexter(..., getText("context").remove())
```

Boşluk tuşu girdisini algılamaktadır. Tuşa basıldığında bağlam metnini kaldırarak bir sonraki aşamaya geçişi sağlamaktadır.

```
newController("DashedSentence", ...)
```

Öz-denetimli okuma sunumunu gerçekleştirmektedir. Bu modülle katılımcı boşluk tuşuyla sözcük sözcük ilerlemektedir. Cümle tamamlanmadan sonraki aşamaya geçişe izin verilmez.

```
newController("Question", ...)
```

Okuma tamamlandıktan sonra ikili karar sorusunu (Evet / Hayır) sunmaktadır. Bu noktada F ve J tuşları yanıt seçeneklerine atanmıştır. Yanıt verilmeden deneme sona ermez.

4.6 Test aşaması

Deneyin temel bölümünü oluşturan test aşamasında, deneysel uyaranlar katılımcılara rastgele sırayla sunulmuştur. Her test denemesi aşağıdaki bileşenlerden oluşmaktadır:

4.6.1 Bağlam sunumu

Her deneme, hedef cümleden önce sunulan kısa bir bağlam metni ile başlamaktadır. Katılımcı boşluk tuşuna bastığında bağlam metni ekrandan kaldırılır ve hedef cümleye geçilir.

4.6.2 Öz-denetimli okuma (sözcük sözcük sunum)

Hedef cümle, öz-denetimli okuma modülü kullanılarak sözcük sözcük sunulmuştur. Katılımcı her boşluk tuşuna bastığında yalnızca bir sözcük görünüp, önceki sözcük maskeleye yöntemiyle gizlenmektedir. Bu yöntem, katılımcıların cümleyi kendi okuma hızlarında işlemlemelerine olanak tanımaktadır. Okuma süreci boyunca her sözcük için tepki süreleri kaydedilmektedir.

4.6.3 Soru ve yanıt

Hedef cümlelerin tamamlanmasının ardından katılımcılara ikili bir karar sorusu yöneltilmektedir. Katılımcılar, cümlede yer alan öznenin kendisinden mi yoksa başka bir kişiden mi bahsettiğine karar vererek: Evet (F tuşu) ya da Hayır (J tuşu) seçeneklerinden birini işaretlemektedir. Verilen yanıtlar ve yanıt süreleri de otomatik olarak kaydedilmektedir. Denemelerde ayrıca, deneysel koşullara bağlı olarak doğru yanıt tanımlamaları yapılmış ve bu bilgiler analiz sürecinde kullanılmak üzere kaydedilmiştir. Test aşamasını oluşturmak için kullanılan örnek kod ise aşağıda sunulmuştur:

```
Template("test.csv", row =>
  newTrial("test",
    // 1) Bağlam
    newText("context", row.context).print(),
    newNexter("nextFromContext", getText("context").remove()),

    // 2) Moving-window öz-denetimli okuma (RT kaydı açık)
    newController("DashedSentence", { s: row.sentence.trim() })
    .log()
    .wait()
    .remove(),

    // 3) İkili karar sorusu (F/J) - yanıt + RT kaydı açık
    newController("Question", {
      as: [{"f", "Evet (F)"}, {"j", "Hayır (J)"}],
      hasCorrect: row.type.trim().endsWith("/null") ? 0 : 1,
      randomOrder: false,
      presentHorizontally: true
    })
    .log()
    .wait()
    .remove()
  )
  // 4) Deneme-etiketleri (koşul/uyaran bilgisi)
```

```
.log("group", row.group)
.log("item", row.item)
.log("type", row.type)

// 5) Katılımcı bilgileri (formdan gelen global değişkenler)
.log("fullname", getVar("fullname"))
.log("age", getVar("age"))
.log("language", getVar("language"))
);
```

Kullanılan modüllerin ayrıntılı açıklaması aşağıdadır:

```
Template("test.csv", row => ...)
```

Ana test uyarılarını test.csv dosyasından satır satır okumakta ve her satır için bir test denemesi üretmektedir.

```
newTrial("test", ...)
```

Bu modül ana test denemesini oluşturmaktadır. Denemeler tipik olarak `Sequence(..., randomize("test"), ...)` ile rastgele sırada sunulmaktadır.

```
newText("context", row.context).print()
```

Bağlam metnini katılımcılara ekranda sunmaktadır (.csv'den okur).

```
newNexter("...", getText("context").remove())
```

Katılımcının boşluk tuşuna basmasını bekler; basılınca bağlam metnini kaldırarak hedef cümle aşamasına geçişi sağlamaktadır.

```
newController("DashedSentence", ...).log().wait().remove()
```

Öz-denetimli okuma sunumunu boşluk tuşuyla sözcük sözcük gerçekleştirmektedir.

```
.log(): sözcük/tuş basışı düzeyinde okuma sürelerini kaydeder.
```

```
.wait(): cümle tamamlanmadan sonraki aşamaya geçişe izin vermez.
```

```
.remove(): tamamlandığında öğeyi ekrandan kaldırır.
```

```
newController("Question", ...).log().wait().remove()
```

Okuma sonrasında ikili karar sorusunu sunmaktadır (F=Evet, J=Hayır).

```
as: seçenekleri tuşlara bağlar.
```

```
presentHorizontally:true seçenekleri yatay; randomOrder:false sıralamayı sabit gösterir.
```

```
hasCorrect: doğru yanıt tanımlamaktadır (0=ilk seçenek/F, 1=ikinci seçenek/J).
```

```
.log(): yanıt ve yanıt süresini kaydeder.
```

```
.wait(): yanıt verilmeden denemeyi ilerletmez.
```

```
.remove(): yanıt sonrası soruyu kaldırır.
```

```
.log("group"/"item"/"type", ...)
```

Denemeyi koşul/uyaran etiketleriyle işaretlemektedir; sonuç dosyasında filtreleme ve analiz için kullanılmaktadır (bkzn., bölüm 4.7).

```
.log("fullname"/"age"/"language", getVar(...))
```

Katılımcı bilgi formundan gelen global değişkenleri her denemeye eklemektedir; böylece sonuç dosyasında deneme düzeyinde katılımcı-demografik bilgi eşleştirmesi yapılabilir.

4.6.4 Deneyin sonlandırılması ve veri kaydı

Tüm test denemeleri tamamlandıktan sonra deney bir kapanış ekranı ile sonlandırılmıştır. Deney süresince toplanan okuma süreleri, yanıtlar, ve katılımcı bilgileri sistem tarafından otomatik olarak kaydedilmiş ve analiz için hazır hale getirilmiştir. Buna ilişkin örnek kod aşağıda sunulmaktadır:

```
Sequence(
  "notice",
  "intro",
  "form",
  "practice",
  randomize("test"),
  SendResults(),
  "outro",
);
newTrial("outro",
  newText("outro", "Görüşürüz! See you! Arrivederci! Auf wiedersehen! じゃあね!"),
  newButton("wait4me", "").wait(),
);
```

Kullanılan modül ve işleri ise aşağıda verilmiştir:

```
SendResults()
```

Deney boyunca toplanan tüm verilerin (tepki süreleri, yanıtlar, ve katılımcı bilgileri) PCIBex sunucusuna gönderilmesini sağlamaktadır. Bu modül çalıştırılmadan deney sonlandırılırsa veriler kaydedilmez. Bu nedenle, deney akışında test aşamasından hemen sonra konumlandırılmıştır.

```
Sequence(...)
```

Deneyin genel akış sırasını belirlemektedir.

```
newTrial("outro", ...)
```

Deneyin son ekranını (kapanış sayfasını) oluşturmaktadır.

```
newText("outro", ...)
```

Kapanış mesajını katılımcılara ekranda sunmaktadır.

```
newButton("wait4me", "").wait()
```

Deneyin otomatik olarak kapanmasını engeller. Katılımcı bu ekranda kalıp tarayıcıyı kendi isteğiyle kapatabilmektedir. Bu yaklaşım, SendResults() işleminin tamamlanması için güvenli bir bitiş sağlamaktadır.

4.7 Sonuç Dosyasını Okuma

Bu çalışmada elde edilen deneysel veriler, PennController altyapısı tarafından üretilen olay-temelli (event-based) bir sonuç dosyasında kaydedilmiştir. Sonuç dosyası, deney sırasında gerçekleşen her etkileşimi ayrı bir satır olarak kaydeden uzun formatlı bir yapıya sahiptir. Aşağıda, sonuç dosyasında yer alan sütunların nasıl yorumlanacağı ve bu sütunların deney kodunda tanımlanan aşamalarla nasıl ilişkilendirileceği, tek bir denemeye ait örnek satırlar üzerinden ayrıntılı biçimde açıklanmaktadır.

Dosyanın en üst kısmında # işaretiyle başlayan satırlar yer almaktadır. Bu satırlar, deneyin çalıştırıldığı tarih, tarayıcı bilgisi ve teknik ayarlara ilişkin meta-bilgiler içermektedir. Bu bilgiler istatistiksel analizlerde kullanılmamış ve veri işleme aşamasında dışlanmıştır. Bu bilgiler aşağıda örneklendirilmiştir:

```
#
# Results on Tue, 17 Dec 2024 12:37:49 GMT
# USER AGENT: Mozilla/5.0 (Windows NT 10.0
# Design number was non-random = 0
```

Analiz edilen kısım ise, bu meta-bilgilerin ardından gelen ve gerçek deney olaylarını içeren satırlardan oluşmaktadır. Sonuç dosyasında yer alan her satır, deney sırasında gerçekleşen tek bir olayı temsil etmekte ve bu olaylara ilişkin bilgiler PennController tarafından önceden tanımlanmış sabit bir sütun sırası içerisinde kaydedilmektedir. Bu sütunlar sırasıyla⁵ Controller name, Order number of item, Inner element number, Label, Latin Square Group, PennElementType, PennElementName, Parameter, Value ve EventTime alanlarını içermektedir. Bu temel teknik sütunları, deney kodunda .log() komutları aracılığıyla eklenen group, item, type, fullname, age ve language sütunları izlemektedir.

Bu sütunların sonuç dosyasındaki görünümü aşağıda sunulmaktadır. Her bir katılımcıya ait bu bilgiler, sonuç dosyasında bir kez yer almaktadır:

```
#
# Columns below this comment are as follows:
# 1. Results reception time.
# 2. MD5 hash of participant's IP address.
# 3. Controller name.
# 4. Order number of item.
# 5. Inner element number.
# 6. Label.
# 7. Latin Square Group.
# 8. PennElementType.
# 9. PennElementName.
# 10. Parameter.
# 11. Value.
# 12. EventTime.
# 13. Comments.

1734439069,8a49962f9d12346acdf3b08b19f23e15,PennController,1,0,notice,NULL,Penn
Controller,0,_Trial_,Start,1734435806773,NULL
```

⁵ Sonuç dosyasında yer alan sütun adlarının Türkçe karşılıkları aşağıdadır:

Controller name → Denetleyici adı
Order number of item → Denemenin deney içindeki sıra numarası
Inner element number → Deneme içi öge numarası
Label → Deneme etiketi
Latin Square Group → Latin kare grubu
PennElementType → PennController öge türü
PennElementName → PennController öge adı
Parameter → Parametre / olay türü
Value → Değer (uyarıcı ya da yanıt içeriği)
EventTime → Olay zamanı (ms)
group → Deneyisel grup
item → Uyarıcı numarası
type → Deneyisel koşul
fullname → Katılımcı adı-soyadı
age → Yaş
language → Ana dili
Reading time → Okuma süresi
Whether or not answer was correct → Yanıtın doğru olup olmadığı
Time taken to answer → Yanıt verme süresi
DashedSentence → Öz-denetimli okuma (maskleme) sunumu

```
# 13. group.
# 14. item.
# 15. type.
# 16. fullname.
# 17. age.
# 18. language.
# 19. Comments.

1734439069,8a49962f9d12346acdf3b08b19f23e15,PennController,12,0,test,NULL,PennC
ontroller,11,_Trial_,Start,1734438850287,1,5,Q-universalNEG /null,AD-SOYAD,19,
Türkce,NULL
```

Öz-denetimli okuma ve karar aşamasına ilişkin ölçümler ise dosyanın devamında Reading time, Whether or not answer was correct ve Time taken to answer sütunlarında yer almaktadır. Her okunan cümlelerin başında bu sütunlar yer almaktadır. Deney kapsamında kaç cümle okunmuşsa, bu sütunlar da her cümle için sonuç dosyasında aynı biçimde yinelenmektedir. Sözcüklerin kaç milisaniye (ms) içinde okunduğuna ilişkin süre bilgileri, öz-denetimli okuma ölçümü kapsamında bu sütunlar aracılığıyla kaydedilmektedir:⁶

```
# 19. Reading time.
# 20. Newline?.
# 21. Sentence (or sentence MD5).
# 22. Comments.

1734439069,8a49962f9d12346acdf3b08b19f23e15,PennController,12,0,test,NULL,Contr
oller-DashedSentence,DashedSentence,1,Kimse,1734438868224,1,5,Q-universalNEG
/null, AD-SOYAD,19,Türkce,337,false,Kimse kopya cektigini söylemedi.,Any
additional parameters were appended as additional columns

1734439069,8a49962f9d12346acdf3b08b19f23e15,PennController,12,0,test,NULL,Contr
oller-DashedSentence,DashedSentence,2,kopya,1734438868224,1,5,Q-universalNEG
/null, AD-SOYAD,19,Türkce,377,false,Kimse kopya cektigini söylemedi.,Any
additional parameters were appended as additional columns

1734439069,8a49962f9d12346acdf3b08b19f23e15,PennController,12,0,test,NULL,Contr
oller-DashedSentence,DashedSentence,3,cektigini,1734438868224,1,5,Q-
universalNEG /null, AD-SOYAD,19,Türkce,401,false,Kimse kopya cektigini
söylemedi.,Any additional parameters were appended as additional columns

1734439069,8a49962f9d12346acdf3b08b19f23e15,PennController,12,0,test,NULL,Contr
oller-DashedSentence,DashedSentence,4,söylemedi.,1734438868224,1,5,Q-
universalNEG /null, AD-SOYAD,19,Türkce,2768,false, Kimse kopya cektigini
söylemedi.,Any additional parameters were appended as additional columns
```

Bu sütunların işleyişi, tek bir denemeye ait örnek satırlar üzerinden açıkça görülebilmektedir. Örneğin ana test aşamasında sunulan “Kimse kopya cektigini söylemedi.” cümlesi, sonuç dosyasında PennElementName = DashedSentence olarak etiketlenmiş dört ayrı satırla temsil edilmiştir. Bu satırlarda Controller name alanı olayın Controller-DashedSentence tarafından üretildiğini, Label alanı olayın test denemesine ait olduğunu göstermektedir. Parameter sütunu kelimenin cümle içindeki sırasını (1–4), Value sütunu ise ekranda sunulan kelimeyi (*Kimse*, *kopya*, *cektigini*, *söylemedi*.) içermektedir. Aynı satırlarda yer alan Reading time sütunu, her bir kelimenin ekranda kaldığı süreyi ms cinsinden doğrudan vermektedir.

⁶ Okunabilirliği artırmak amacıyla, örnekte okunan sözcükler ile bu sözcüklerin kaç ms içinde okunduğuna ilişkin süre bilgileri koyu vurgulanmıştır.

Yukarıda verilen örnekte okuma süreleri sırasıyla **Kimse** için 337 ms, **kopya** için 377 ms, **çektğini** için 401 ms ve **söylemedi.** için 2768 ms olarak kaydedilmiştir. Bu sözcük okuma satırlarının tamamında `group = 1`, `item = 5` ve `type = Q-universalNEG /null` (öncül: olumsuzluk taşıyan evrensel niceleyici/içe yerleşik tümce öznesi: boş) değerlerinin tekrar etmesi, bu olayların aynı deneysel denemeye ve koşula ait olduğunu açıkça göstermektedir. Aynı şekilde `age` ve `language` sütunlarında yer alan değerler, bu okuma olaylarının ilgili katılımcıyla doğrudan ilişkilendirildiğini ortaya koymaktadır.

Okuma aşamasını takiben sunulan ikili karar sorusuna verilen yanıt, aynı denemeye ait ayrı bir satırla kaydedilmiştir. Bu satırda `PennElementName = Question` olarak belirtilmekte; `Value` sütununda katılımcının verdiği yanıt **Evet (F)** biçiminde açıkça yer almaktadır. Bu yanıt, deney yönergelerinde belirtildiği üzere klavyedeki `F` tuşuna karşılık gelmektedir. Aynı satırda bulunan `Whether or not answer was correct` sütununda değer 1 olarak kaydedilmiş olup, bu durum verilen yanıtın doğru olduğunu göstermektedir. Bu doğruluk bilgisi, deney kodunda tanımlanan `hasCorrect` parametresiyle birebir uyumludur. `Time taken to answer` sütununda yer alan 2008 ms değeri ise katılımcının soruya yanıt vermek için harcadığı süreyi göstermektedir. Böylece bu deneme için hem sözcük düzeyindeki okuma süreleri hem de karar aşamasına ait yanıt ve yanıt süresi, aynı sonuç dosyası içerisinde açık biçimde izlenebilmektedir.

```
# 19. Whether or not answer was correct (NULL if N/A).
# 20. Time taken to answer..
# 21. Comments.

1734439069,8a49962f9d12346acdf3b08b19f23e15,PennController,12,0,test,NULL,Controller-Question,Question,NULL,Evet (F),1734438870233,1,5,Q-universalNEG /null,AD-SOYAD,19, Türkce,1,2008,Any additional parameters were appended as additional columns

# 19. Comments.
1734439069,8a49962f9d12346acdf3b08b19f23e15,PennController,12,0,test,NULL,PennController,11,_Trial_,End,1734438870234,1,5,Q-universalNEG /null,AD-SOYAD,19,Türkce,NULL
1734439069,8a49962f9d12346acdf3b08b19f23e15,PennController,18,0,test,NULL,PennController,17,_Trial_,Start,1734438870237,1,11,wh-distributive /null,AD-SOYAD,19,Türkce,NULL
```

4.8 Veri çözümleme

Bu bölümde, `PennController` tarafından üretilen sonuç dosyasının R ya da Python gibi programlama tabanlı istatistik ortamına aktarılması ve araştırma sorusuna uygun biçimde çözümlemesi için izlenen çözümleme akışı sunulmaktadır. Çalışmanın temel hedefi, deneysel koşullara (ör. öncül türü + boş/açık özne adlı türü bilgisi) göre sözcük düzeyinde okuma sürelerinin (ms) ve ikili karar bölümünde yanıt doğruluğu ile yanıt verme sürelerinin (ms) özetlenmesidir.

Analiz süreci üç ana adımdan oluşmaktadır. İlk adımda sonuç dosyası R ya da Python'a aktarılır; PCIBex dosyalarında `#` ile başlayan açıklama satırları veri olmadığı için içe aktarma sırasında dışarıda bırakılmaktadır. İkinci adımda yalnızca ana deney denemeleri seçilmektedir; bunun için `label` değişkeni

“test” olan satırlar filtrelendir. Üçüncü adımda ise veri iki ayrı ölçüm türüne ayrılmaktadır: öz-denetimli okuma verisi ve ikili karar verisi. Öz-denetimli okuma verisi, `penn_element_name` değeri “DashedSentence” olan satırlardan alınır; bu satırlarda her kelime için `Reading time (ms)` doğrudan mevcuttur. Karar verisi ise `penn_element_name` değeri “Question” olan satırlardan elde edilmektedir; burada yanıt (Evet–F / Hayır–J), yanıtın doğruluğu (1/0) ve yanıt verme süresi (ms) ayrı değişkenler olarak ele alınmaktadır.

Son aşamada, deneysel koşulları temsil eden `type` değişkenine göre grupta yapılarak her koşul için ortalama okuma süresi ve standart sapma (SS) hesaplanmıştır. İkili karar aşaması için de aynı şekilde koşul bazında doğruluk oranı ve ortalama yanıt verme süresi üretilmiştir. Böylece, araştırma sorusunun merkezinde yer alan “öncül türüne göre boş/açık özne adlı içeren cümlelerde sözcükler ortalama kaç ms’de okunuyor?” sorusu, koşul bazında özet istatistiklerle doğrudan yanıtlanabilir hale gelmiştir. Aşağıda bu deneyin verilerini çözümlemeye ilişkin örnek bir R kodu bulunmaktadır⁷.

```
library(readr)
library(dplyr)
library(stringr)
library(tidyr)

df <- read_csv(
  "results_prod (4).csv",
  comment = "#",
  col_names = FALSE,
  locale = locale(encoding = "UTF-8"),
  show_col_types = FALSE
)

colnames(df) <- c(
  "reception_time", "ip_md5",
  "controller_name", "item_order", "inner_element",
  "label", "latin_square",
  "penn_element_type", "penn_element_name",
  "parameter", "value", "event_time",
  "group", "item", "type", "fullname", "age", "language",
  "col19", "col20", "sentence", "comments"
)

test <- df %>%
  filter(label == "test")

spr <- test %>%
  filter(penn_element_name == "DashedSentence") %>%
  mutate(
    word_pos = as.integer(parameter),
    word = as.character(value),
    reading_time = as.numeric(col19),
    group = as.integer(group),
    item = as.integer(item),
    age = suppressWarnings(as.integer(age))
  )
```

⁷ Sonuçlar Python ortamına da aktarılabilmeyle birlikte, bu çalışmada veri çözümlemesi R ortamı üzerinden anlatılmaktadır.

```

qa <- test %>%
  filter(penn_element_name == "Question") %>%
  transmute(
    item_order = as.integer(item_order),
    group = as.integer(group),
    item = as.integer(item),
    type = as.character(type),
    age = suppressWarnings(as.integer(age)),
    language = as.character(language),
    answer = as.character(value),
    correct = as.numeric(col19),
    answer_time = as.numeric(col20)
  )

spr_clean <- spr %>%
  filter(!is.na(reading_time)) %>%
  filter(reading_time >= 100, reading_time <= 5000)

qa_clean <- qa %>%
  filter(!is.na(answer_time)) %>%
  filter(answer_time >= 200, answer_time <= 10000)

spr_summary <- spr_clean %>%
  group_by(type) %>%
  summarise(
    mean_rt = mean(reading_time),
    sd_rt = sd(reading_time),
    n = n(),
    .groups = "drop"
  )

qa_summary <- qa_clean %>%
  group_by(type) %>%
  summarise(
    acc = mean(correct, na.rm = TRUE),
    mean_answer_time = mean(answer_time),
    sd_answer_time = sd(answer_time),
    n = n(),
    .groups = "drop"
  )

spr_summary
qa_summary

```

Bu kod R programında çalıştırıldığında ortalama okuma süreleri aşağıdaki gibi bir sonuç üretmektedir:

Tablo 1. Sözcük düzeyi ortalama okuma süreleri

type ⁸	Ortalama okuma süresi (ms)	SS
DP/null	792.43	1315.22
DP/overt	1013.17	2448.52

⁸ Bu kısımda öncül türü ve adıl türleri verilmektedir, Örneğin DP/null etiketinde DP, ana tümcedeki öncülün (öznenin) türünü; null ise içe yerleşik tümcede öznenin boş olarak gerçekleştiğini göstermektedir.

type ⁸	Ortalama okuma süresi (ms)	SS
Q-distributive /null	615.13	560.63
Q-distributive /overt	757.29	1376.08
Q-distributiveNEG /null	809.92	1386.99
Q-distributiveNEG /overt	838.85	1537.44
Q-universal /null	694.69	708.47
Q-universal /overt	807.43	1206.74
Q-universalNEG /null	980.29	1479.19
Q-universalNEG /overt	801.77	1271.14
wh /null	985.55	1490.49
wh /overt	858.47	1460.43
wh-distributive /null	888.08	1904.94
wh-distributive /overt	924.65	1576.15

R'dan elde edilen ikili karar sorularının doğruluk oranları ve ortalama yanıt süreleri ise aşağıdaki tabloda sunulmaktadır:

Tablo 2. Doğruluk ve yanıt süreleri

type	Doğruluk (oran) ⁹	Ortalama yanıt süresi (ms)	SS
DP/null	0.9211	1322.95	868.78
DP/overt	0.3421	2376.45	3579.29
Q-distributive /null	0.8421	2307.08	2650.14
Q-distributive /overt	0.7368	2953.84	4265.46
Q-distributiveNEG /null	0.7895	2236.47	2640.54
Q-distributiveNEG /overt	0.6579	1921.79	1913.04
Q-universal /null	0.7895	2000.71	2197.96
Q-universal /overt	0.6842	2265.45	2795.61
Q-universalNEG /null	0.7632	3515.71	5687.32
Q-universalNEG /overt	0.7368	3168.97	4963.12
wh /null	0.5789	2136.79	2010.94
wh /overt	0.7632	2392.16	2840.21
wh-distributive /null	0.6053	2670.58	3241.79
wh-distributive /overt	0.7105	2462.13	1936.01

Bu noktada vurgulanması gereken önemli bir husus, çalışmada kullanılan okuma sürelerinin (ms) ve yanıt sürelerinin (ms) ölçüm olarak zaten PennController tarafından sonuç dosyasına kaydedildiğidir. Dolayısıyla R

⁹ 1=doğru, 0=ise yanlış olarak alınmıştır.

ortamında yapılan işlem, yeni bir tepki süresi üretmek değil, sonuç dosyasında mevcut olan bu süreleri deneysel koşullara göre filtreleyerek basit ve sade biçimde özetlemek (ör. koşul bazında ortalama, standart sapma) şeklinde gerçekleşmiştir. Bu bağlamda, R programı burada temel olarak veriyi düzenleme ve tanımlayıcı istatistikleri raporlama işlevi görmektedir.

Elde edilen koşul bazlı ortalama okuma süreleri ve karar verileri, çözümleme sürecinin ilk ve betimleyici aşamasını oluşturmaktadır. Bu özet çıktılar alındıktan sonra, deneysel koşullar arasındaki farkların çıkarımsal olarak değerlendirilmesi amacıyla ileri düzey istatistiksel analizlere geçilecektir. Araştırma sorusunun yapısına bağlı olarak, iki koşulun karşılaştırılması gereken durumlarda t-testi gibi parametrik testler uygulanabilir; ancak öz-denetimli okuma verilerinin katılımcı ve madde düzeyinde tekrarlı ölçümler içermesi nedeniyle, katılımcı ve madde değişkenlerini aynı anda modelleyebilen doğrusal karma etkili modeller (linear mixed-effects models) çoğu durumda daha uygun bir çözümleme yaklaşımı sunmaktadır. Bu nedenle çalışmanın sonraki aşamalarında, okuma süreleri ve karar verileri üzerinde uygun varsayımlar gözetilerek hem temel karşılaştırmaların hem de karma etkili model temelli çözümlemelerin yürütülmesi planlanabilir.

6. Sonuç ve Değerlendirme

Bu çalışmada, PCIBex platformu ve PennController eklentisinin deneysel dilbilim araştırmalarında sunduğu olanaklar, yürütülebilecek uygulamalar ve örnek bir öz-denetimli okuma testi tasarımı hem kuramsal hem de uygulamalı düzeyde tartışılmıştır. Son yıllarda psikodilbilimsel yöntemlerin yaygınlaşması ve çevrim içi veri toplama pratiklerinin araştırmalar için son derece yaygın olması tarayıcı tabanlı deney altyapılarını yalnızca bir seçenek olmaktan çıkarıp giderek daha merkezi bir konuma taşımıştır. Bu bağlamda PCIBex erişilebilir bir ortamda, deney akışını baştan sona yönetebilmesi bakımından araştırmacılar için pratik bir çözüm sunmaktadır.

Özellikle PennController'ın JavaScript tabanlı, ancak araştırmacının hızlıca uygulamaya geçebilmesine olanak sağlayan görece sade söz dizimi, psikodilbilimsel paradigmaların web ortamına aktarılmasında önemli bir avantajdır. Örnek uygulamada da görüldüğü üzere, okuma süresi gibi zamana duyarlı ölçümler ile karar yanıtları aynı denemede birleştirilebilmekte; deneysel koşullara ilişkin etiketler (group, item, type) ile katılımcı bilgileri (age, language vb.) `.log()` komutları aracılığıyla her olay satırına sistematik biçimde ilişkilendirilebilmektedir.

Bununla birlikte, çevrim içi deneylerin doğası gereği bazı yöntembilimsel sınırlılıklar kaçınılmazdır ve sonuçların yorumlanmasında bu sınırlılıkların açıkça dikkate alınması gerekir. Katılımcıların ekran boyutu, klavye düzeni, tarayıcı performansı, dikkat düzeyi ve ortam gürültüsü gibi araştırmacının doğrudan kontrol edemediği değişkenler, özellikle tepki süresi ölçümlerinde güçlükler yaratabilir. Bu nedenle PCIBex üzerinden yürütülen zamana duyarlı çalışmalarda, veri toplama öncesinde minimum teknik gereksinimlerin (ör. yalnızca bilgisayar, belirli tarayıcılar, tam ekran tercihleri)

yönergelerde net biçimde belirtilmesi; veri sonrası aşamada ise dışlama ve temizleme adımlarının (aşırı kısa/uzun süreler, eksik yanıtlar) raporlanması önemlidir. Diğer bir deyişle, PCIBex çevrim içi çalışmayı kolaylaştırırken, araştırmacının ölçüm kalitesini korumak için tasarım ve çözümleme düzeyinde daha sistemli önlemler almasını da gerektirebilir.

Sonuç olarak bu çalışma PCIBex'in, deney tasarımı ve veri yönetimi açısından araştırmacıya sağladığı pratik avantajları örnek bir öz-denetimli okuma senaryosu üzerinden görünür kılmış; aynı zamanda çevrim içi veri toplamanın taşıdığı sınırlılıkları da ortaya koymuştur. Bu çerçevede PCIBex, hem psikodilbilimsel araştırmaların ölçeğini büyütme hem de farklı katılımcı profillerine erişimi artırma potansiyeli taşımaktadır.

Teşekkür

Bu çalışmada kullanılan deney kodlarının hazırlanması ve uygulanmasındaki katkılarından dolayı İstanbul Medeniyet Üniversitesi Dilbilimi Lisans öğrencisi Deniz Engin'e teşekkür ederiz.

Kaynakça

- Ariga, T., & Hirose, Y. (2025). Recognition of spoken words with mispronounced lexical prosody in Japanese. *The Journal of the Acoustical Society of America*, 157(6), 4102–4118.
- Britton, J., Cong, Y., Hsu, Y. Y., Chersoni, E., & Blache, P. (2024). On the influence of discourse connectives on the predictions of humans and language models. *Frontiers in Human Neuroscience*, 18, 1363120.
- Cairncross, A., Vogelzang, M., & Tsimpli, I. (2024). Evaluating the pseudorelative-first hypothesis: Evidence from self-paced reading and persistence effects. *Glossa Psycholinguistics*, 3(1).
- Çakır, S., & Çınar, O. (under review). The contextual influences on the pronominal binding preferences in Turkish: A processing-based study.
- Çınar, O. (2024). Revisiting native grammar through L2 theories: Knowledge and processing of null and overt subject pronouns in Turkish. *Nesne Psikoloji Dergisi*, 12(33), 331–350.
- D'Alesio, V., Porrini, A. T., Greco, M., & Moro, A. (2025). An investigation of nominal copular sentences in three reading paradigms: Acceptability judgments, self-paced reading, and eye-tracking. *Lingua*, 326, 104019.
- Drummond, A. (2013). Ibox Farm. Available at: <https://farm.pcibex.net/>
- Gračanin-Yukse, M., Lago, S., Şafak, D. F., Demir, O., & Kırkıcı, B. (2017). The interaction of contextual and syntactic information in the processing of Turkish anaphors. *Journal of Psycholinguistic Research*, 46(6), 1397–1425.
- Gračanin-Yukse, M., Lago, S., Şafak, D. F., Demir, O., & Kırkıcı, B. (2020). The interpretation of syntactically unconstrained anaphors in Turkish heritage speakers. *Second Language Research*, 36(4), 475–501.
- Hatzidaki, A., & Santesteban, M. (2024). Emotion effects survive non-standard orthographic representations. *Cognition and Emotion*, 1–8. <https://doi.org/10.1080/02699931.2024.2362381>

- Just, M. A., Carpenter, P. A., & Woolley, J. D. (1982). Paradigms and processes in reading comprehension. *Journal of experimental psychology: General*, 111(2), 228.
- Kim, D., & Kim, H. (2025). Effects of L1 knowledge and task type on bilingual acquisition and processing of English count and mass nouns. *International Journal of Bilingualism*.
<https://doi.org/10.1177/13670069251360171>
- Kim, H., & Hwang, H. (2025). Learning contexts and proficiency matter: L2 real-time sensitivity to conventional and unconventional dative pattern. *Language and Cognition*, 17, e21.
<https://doi.org/10.1017/langcog.2024.65>
- Lee, S. H., & Kaiser, E. (2021). Does hitting the window break it? Investigating effects of discourse-level and verb-level information in guiding object state representations. *Language, Cognition and Neuroscience*, 36(8), 921–940.
<https://doi.org/10.1080/23273798.2021.1896013>
- Liu, J., Fan, L., Zhang, X., & Xu, Z. (2025). Affective and psycholinguistic norms for 1200 two-character Chinese words. *International Journal of Applied Linguistics*.
- Lu, X. (2024). Enhancing the processing advantage: Two psycholinguistic investigations of formulaic expressions in Chinese as a second language. *International Review of Applied Linguistics in Language Teaching*. <https://doi.org/10.1515/iral-2023-0262>
- Moulton, K., Block, T., Gendron, H., Storoshenko, D., Weir, J., Williamson, S., & Han, C.-H. (2022). Bound variable singular *they* is underspecified: The case of *all* vs. *every*. *Frontiers in Psychology*, 13, 880687.
- Ovans, Z., Ayala, M. R., Asmah, R., Hu, A., Montoute, M., Van Horne, A. O., Qi, Z., Morini, G., & Huang, Y. T. (2025). The feasibility of remote visual-world eye-tracking with young children. *Open Mind*, 9, 992–1019. <https://doi.org/10.1162/opmi.a.16>
- Özsoy, O., Çiçek, B., Özal, Z., Gagarina, N., & Sekerina, I. A. (2023). Turkish-German heritage speakers' predictive use of case: webcam-based vs. in-lab eye-tracking. *Frontiers in Psychology*, 14, 1155585.
- Papoutsaki, A., Sangkloy, P., Laskey, J., Daskalova, N., Huang, J., & Hays, J. (2016). *WebGazer: Scalable webcam eye tracking using user interactions*. In Proceedings of the 25th International Joint Conference on Artificial Intelligence (IJCAI-16) (pp. 3839–3845). AAAI Press.
<https://dl.acm.org/doi/10.5555/3061053.3061156>
- Park, S.-H., & Lee, E.-K. (2024). Definiteness and context in double-accusative ditransitives in Korean: An experimental approach. *Language Research*, 41(3), 345–366.
- Patterson, A. S., Nicklin, C., & Vitta, J. P. (2025). Methodological recommendations for webcam-based eye tracking: A scoping review. *Research Methods in Applied Linguistics*, 4(3), 100244.

- Plaut-Forckosh, D., & Meir, N. (2025). Insights into linguistic performance: Investigating pragmatic and syntactic violations of Hebrew definiteness through acceptability judgment and self-paced listening and reading tasks. *Applied Psycholinguistics*, 46, e5.
<https://doi.org/10.1017/S0142716425000037>
- Ruda, M., & Huang, N. (2025). Overt pronouns in null subject languages: Experimental investigation of Kashubian, Polish, and Silesian. *Linguistics*.
- Slim, M. S., & Hartsuiker, R. J. (2023). Moving visual world experiments online? A web-based replication of Dijkgraaf, Hartsuiker, and Duyck (2017) using PCIBex and WebGazer.js. *Behavior Research Methods*, 55(7), 3786–3804. <https://doi.org/10.3758/s13428-022-01989-z>
- Uygun, S. (2022). Processing pro-drop features in heritage Turkish. *Frontiers in Psychology*, 13, 988550.
- Yamaguchi, K., & Ohta, S. (2023). Dissociating the processing of empty categories in raising and control sentences: A self-paced reading study in Japanese. *Frontiers in Language Sciences*, 2, 1138749.
- Zehr, J., & Schwarz, F. (2018). PennController for Internet Based Experiments (IBEX). <https://doi.org/10.17605/OSF.IO/MD832>
- Zhang, N., Ren, J., Wang, M., & Jiang, N. (2023). Automatic phonological access among bilinguals with cross-script languages. *Quarterly Journal of Experimental Psychology*, 77(7), 1399–1417.
<https://doi.org/10.1177/17470218231198500>

DISTINCTIVE DIACRITICS IN THE ATIL-TURKIC EPITAPHS AND IN THE YARKAND DOCUMENTS

Marcel ERDAL

Prof., Goethe University Frankfurt a. M., Department of Turcology, merdal4@gmail.com, ORCID: 0000-0002-1604-4193

Erdal, M. (2025). Distinctive diacritics in the Atil-Turkic epitaphs and in the Yarkand documents. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 109–116). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X005>

ABSTRACT

Scribes using the Arabic alphabet for writing Arabic, Persian or early varieties of Turkic used non-obligatory ‘distinctive’ diacritics beside the standard diacritics of the writing systems of their languages. After quoting Onur (2024) on the presence of such diacritics in Arabic, Persian, Khāqānī Turkic and Early Middle Turkic, the author documents their appearance in 13th-14th century epitaphs in Arabic and Atil Turkic, a language closely related to Chuvash, spoken in the Volga-Kama area of Russia, and in 11th-12th century legal documents in Uyghur script discovered in Yarkand, East Turkestan. The distinctive diacritics consist either of dots or of small Arabic letters placed either below or above the line of writing.

Keywords: The Arabic writing system, Atil Turkic, Early middle Turkic texts, Yarkand documents

ÖZ

Arap alfabesini Arapça, Farsça ya da Türkçenin erken dönem değışkelerini yazmak için kullanan kâtipler, dillerin yazı sistemlerinde yer alan ölçünlü diyakritik işaretlerin yanı sıra, zorunlu olmayan “ayırt edici” diyakritik işaretler de kullanmışlardır. Yazar, Onur (2024)’te Arapça, Farsça, Hākānī Türkçesi ve Erken Orta Türkçede bu tür diyakritiklerin varlığına yapılan atıftan sonra, söz konusu işaretlerin 13.–14. yüzyıllara ait Arapça ve Rusya’nın İdil–Kama bölgesinde konuşulmuş ve Çuvaşça ile yakından ilişkili bir dil olan Atil Türkçesi mezar kitabelerinde ve Doğu Türkistan’da, Yarkent’te bulunmuş 11.–12. yüzyıla ait Uygur yazılı hukuk belgelerinde de görüldüğünü belgelemektedir. Ayırt edici diyakritik işaretler, ya noktalardan ya da yazı satırının altına ya da üstüne yerleştirilen küçük Arap harflerinden oluşmaktadır.

Anahtar Sözcükler: Arap yazı sistemi, Atil Türkçesi, Erken Orta Türkçe kaynaklar, Yarkent belgeleri

1. Introduction

The Arabic alphabet consists of 28 letters, representing consonants. Some of these letters differ from each other in that their simple (*rasm*) forms get a dot or two dots placed above or below them, or three dots placed above them: Thus, *b* differs from *t*, from *n* and from *θ* just by having a dot under it while *t* has two dots above it, *n* (in the beginning and middle of a word) one dot above it and *θ* has three dots above it. *y*, *δ*, *d*, *z* and *z* differ from ^c (*‘ayn*), *d*, *ṣ*, *r* and *ṭ* respectively just by having a single dot above them. *j* [dʒ] differs from *h* and from *x* just by having a single dot under it while *x* has one dot above it. *š* differs from *s* just by having three dots above it.¹ This consonant pointing, called *i‘jām*, is obligatory; however, some early Qur’ān manuscripts still only have the *rasm* forms of the letters without the *i‘jām* pointing.

The standard Arabic writing system also has a system of non-obligatory diacritic signs, called *taškil*; these include the vowel marks called *ḥarakāt*, the *sukūn* to show that no vowel is to be pronounced, the *šadda* above consonants which are to be pronounced doubly, the *tanwīn*, nunation, for the marking of indefinites and for adverb formation, the *madda* and the *waṣla* above the letter *alif* (the first to indicate a long /a:/, the second to mark positions where the *hamza* – the glottal stop – is deleted).

This paper is about a different set of non-obligatory diacritics hardly mentioned in descriptions of the writing system: the distinctive signs, called ‘Differenten’ in German, which were in use till the 14th century. The job of the distinctive signs was to make sure that *rasm* letters differing by *i‘jām* pointing were read and understood correctly, whether the *i‘jām* dots were actually present or were lacking: The central idea would have been that the *i‘jām* markings could have been effaced on the stone or in the manuscript, leading to a wrong reading. The distinctive diacritics consist of single dots under the letters to be read as ^c, *d*, *ṣ*, *r* or *ṭ*, and mostly of three dots under the letter to be read as *s*; there are some other distinctive signs as well, to which we will come when discussing the Turkic Yarkand manuscripts.

2. The Arabic Evidence

The oldest instance mentioned by Grohmann (1971: 42) for such signs consists of a dot placed under the letter *ṭ*, which would be read as *z* if there was a dot above it, in l. 2 of an epitaph from the year 883. In an 1163/64 epitaph from Baṣra, *d* is marked with a dot under it; if the letter had a dot above it, it would be read as *δ*. In a Baṣra epitaph from the year 1159 mentioned by Grohmann, the *s* is distinguished from *š* with a single dot above it, but in widespread use elsewhere, *s* is made to differ from *š*

¹ I am referring to the standard system used for Classical Arabic. I will not be going into phonetic or graphemic detail here, as the facts are well-known. Nor is knowledge of the terms of the next paragraph, part of Arabic grammar, necessary for understanding the present paper.

with three dots under it, whereas *š* has three dots above it: Grohmann thus mentions a 1226 inscription in Syria in the word شمس (*šams*, ‘the sun’), and in the word بسم (*bismi* ‘in the name of ...’) in the *basmala* on the sarcophagus of a certain Asʿad ibn ʿAli baq in Ġazna (Afghanistan, late 11th century) where the *s* is spelled as پ, although the context would have made the reading clear enough. These examples also show the geographical spread of the phenomenon in Arabic. Onur (2024: 850), footn. 47 mentions an Arabic manuscript with this feature, from the year 1356.

The early 12th century Arabic legal manuscripts edited by Gronke (1986) also show distinctive diacritics. She discusses the phenomenon on p. 458 with the phrase *aḥad ḥudūdihī* ‘one of its borders’ with dots under its first and second *dāls* to show that they are not *ḍāls*. No dot appears under the third *dāl*; it very often happens in all our sources that – without explanation – some of the relevant letters are marked with distinctive diacritics while others are not. The term *waṣiyy* ‘guardian’ has a dot under its *šād* to mark it as *šād*. The names *Yūsuf* and *Ilyās* and the title *sultān* are, e.g., spelled with three dots under their *sīn*, as is the Turkic title *sübaşı*: The documents were commissioned by the same group of Turkic persons who were responsible for the ones edited by Erdal (1984). Gronke also mentions distinctive diacritics in the 13th-14th century documents from Ardabil (Persian Azerbaijan) which she edited, and in Arabic papyri discussed in another work by Grohmann.

3. Evidence in Other Languages

The Arabic script was used also for writing Persian and Turkic varieties and, as documented by Onur (2024: 849-852), the distinctive diacritics were put to use in texts written in those languages as well. For Persian, Onur mentions a pharmacological tractate from 1055/56, whose distinctive signs are already discussed by Orsatti (2019: 53). Further, an Andarz-name (a text with advice and injunctions for proper behaviour) from 1090, a 12th/13th century manuscript from Bukhara, and a document from Bāmiyān (Afghanistan) from the year 1211; full bibliographical details for all of these are given in the publication.

4. Distinctive Diacritics in Turkic

Onur discovered distinctive diacritics also in four Turkic manuscripts, examples of which are presented in facsimiles. The topic was brought up by the author in connection with the dating of the Tahrān manuscript of the Khwarazm-Turkic *Qışaṣu ʿl-Anbiyā*, completed by Rabḡūzī in 1311: With a number of orthographical, grammatical and lexical arguments, he convincingly shows this manuscript to be much closer to the original than the British Museum manuscript, probably written down in the 15th century but hitherto thought to be the oldest and best source. One of the arguments which Onur gives for his early dating of the Tahrān manuscript is based on the recourse to the diacritics discussed in the present

paper.² Another source with distinctive diacritics which Onur gives is the Ferghana copy of the Kutadgu Bilig, clearly the earliest of three manuscripts of this Khākānī text: This copy is taken to have been written down in the first half of the 14th century. A third manuscript with this orthographic feature is the Middle Turkic-Persian interlinear translation of the Qur'an, kept in the Abu Rayhan Biruni Institute of Oriental Studies of the Academy of Sciences of Uzbekistan with the archive number 2008; the editor of this text dated it to the 13th century. Fourthly, Onur found such diacritics in the first few pages of the 1365 manuscript of the Nahju 'l-Farādīs, another Khwarazm-Turkic text.

5. The Atil-Turkic Texts

We will here discuss two other Turkic corpuses showing distinctive diacritics; one is the corpus of Atil-Turkic epitaphs dealt with in Erdal 1993, the other corpus the Muslim manuscripts discovered in Yarkand in East Turkestan and published in Erdal (1984). The Atil-Turkic epitaphs, mostly discovered in present-day Tatarstan, were hitherto called 'Volga Bulgarian', but the name proposed here is more appropriate: The Kama-Volga river system is in Greek, Arabic, Persian and Chaghatay sources called *Atil*.³ The name Bulgarian/Bolgarian is confusing for readers, as modern Bulgarian is a Slavic language spoken in Bulgaria. Atil Turkic is an important branch of the Turkic languages, which is abundantly represented in Hungarian loans and, prehistorically, had an important connection to Mongolic. In their editions of the Atil-Turkic epitaphs, Róna-Tas & Fodor (1973), Khakimzyanov (1978), Khakimzyanov (1987) and Tekin (1988) mark distinctive diacritics both in their transliterations and their transcriptions, and Róna-Tas and Fodor mention them in their notes as well.⁴ This was clearly a general practice, in view of the fact that a number of the marks must, with time, have become unclear on the stone and could not be seen by the scholars.

The epitaphs are bilingual: Beside the Atil-Turkic texts formulated by the relatives of the deceased or by local professionals employed by

² Edited in Çelebi Çam (2025: 203-663). Note that the Tahrān ms. also has numerous errors.

³ Concerning Greek sources, see the entry Ἀτὶλ in Moravcsik 1983. Aḥmad Ibn Faḡlān, who participated in the Abbasid expedition to the Atil Turks in 921-922, calls the large river flowing through the territories of the Bulgarian and Khazar Turks *Atil*, with *madda* above the onset *alif* in all instances, in his travelogue edited by Togan (1939). According to Togan, the name is spelled like this also by al-Bīrūnī in the 11th/12th century and by al-Qazwīnī in the 13th/14th century. In his *Shejere-yi Terākime* written in 1659/60, Abulgazi Bahadır Khan also mentions the *Atil suyu*, the Atil river: According to the note in Ölmez (2020): 266, the name appears with (non-obligatory) *madda* over the *alif* in 5 among the 7 instances in the T manuscript.

⁴ In each of these publications, only parts of the corpus were edited; Tekin (1988), e.g., only has 90 among the 139 epitaphs discussed in Erdal (1993). I have not tried to decipher the facsimiles published in these and in earlier editions (referred to in the mentioned editions), as many of them are not clear and some appear to have been redrawn; some of the distinctive diacritics may have been missed by the editors, but there is no need for the present paper to be exhaustive. There is no reason to think that the editors added distinctive diacritics which are not there, as these marks were not topics of scientific discussions (e.g. whether the language is closer to Chuvash or to Tatar).

them, there are purely Arabic sentences traditionally found in Muslim epitaphs, showing a limited variation. The language used by the community which left us their epitaphs also had Arabic and Persian loans, partly subject to Atil-Turkic sound laws. The distinctive diacritics found in our epitaphs appear to concern only graphemes, whether the words in which we find them were part of the international Arabic epitaph corpus not necessarily meant to be understood, or of the Atil-Turkic corpus (including its loans).

In the Arabic sentences we find three dots under *sīn* to show that it is not *šīn* e.g. with the noun *nās* ‘people’ or the verb *sayamūtu* ‘they will die’; single dots under *dāl* to show that it is not *ḏāl* e.g. with the verb *dāhiluhu* ‘they will enter it’, under *rā* to show that it is not *zāy* e.g. with the adjective *kabīr* ‘great’, under ‘*ayn* to show that it is not *ḡayn* e.g. with *al-‘aliyy* ‘the superior’.

The Turkic text obviously includes many proper names, e.g. *Ismā‘īl* with dots under the *sīn* and the ‘*ayn*, *Sāra* with dots under the *sīn* and the *rā* or *Xar(i)mās* with a dot under the *rā*. There are Arabic loans, e.g. the dative / accusative forms *āxirete* ‘to the thereafter’ with a dot under the *rā* and *mesjidsemne* ‘mosques’ with a dot under *dāl* and three dots under *sīn*. The names of the months are Arabic as well, e.g. *ṣafar* with a dot under the *ṣād* or *rabi‘ul āxir*⁵ with a dot under the ‘*ayn* and a second one under the second *rā*. The numbers are well represented: Note, e.g., the spellings of *toxur* ‘nine’ with dots under the *tā* and *rā*, or *sekir* ‘eight’ with three dots under the *sīn* and one under the *rā*. Even better represented are the ordinal numbers, e.g. *altıši* ‘sixth’ with a dot under the *tā* and *bīrim* with a dot under the *rā*. *hīr* ‘daughter’ is often spelled with a dot under the *rā* and *erle-sa‘at* ‘early hour’ has dots under the *rā*, *sīn* and ‘*ayn*. Finally there are infinite verb forms, among them *batuwi* ‘that (s)he went’, *wafāt(i) baltuwi* ‘that (s)he died’, both with dots under their *tās*, and *tanruwi* ‘that he did’; the modal form *baltur* ‘let it be!’ has distinctive dots under its *tā* and its *rā*.

Differentiae are noted in 49 of the Atil-Turkic epitaphs, dated between 1291 and 1358.

6. The Yarkand Documents⁶

The Turkic Yarkand documents, linked to the Arabic ones of Gronke (1986) both by content and by persons involved, were edited in Erdal (1984); Salur (2024) is an actualized Turkish translation of this paper, supervised by the present author. The graphemic system of these texts is Uyghur: no letters are taken over from the Arabic script, as Şinasi Tekin, a previous editor, thought. Uyghur /n/, /’/ and /r/ look alike, but the first often has one dot above it (as in Uyghur script and not just in Arabic),

⁵ Spelled like this.

⁶ Clark (2010: 93-104) already discusses the features of the script described here, connecting them with Kāšgari’s statements about the writing systems.

the third two dots; I will not be mentioning these normal features of the Uyghur writing system. However, various Arabic signs are used diacritically to distinguish cases in which the same Uyghur letter, or two or more Uyghur letters looking identical in certain positions, represent more than one phoneme or allophone. /z/ is often distinguished from final /n/ by two dots placed under it, while in the word *Težik* in I,32 there are three dots under an R (with the regular two dots above it): This must be reminiscent of the Persian spelling of the specifically Persian sound *ž* with three dots over R.

In the Arabic tradition, three dots under the letter S show that it is not to be read as /š/. This is opposite to Uyghur practice, where /š/ is occasionally marked against /s/ with two dots underneath.⁷ Our texts also often mark the /s/ with dots, and not the /š/: *qanmasa* ‘if he is not satisfied’ in I,17, the proper name *Arslan* in II,b3, the proper name *Narsi* in III,1 and *Ishaq jal(l)ab-qa* ‘to Mr. *Ishaq*’ as well as *keseḳ* ‘piece’ of II,3c3 have three dots each under their S's. In IV,2, the title to be read as *sübaşı* has a single dot under its fifth letter. *igsiz* ‘healthy’ in IV,3 has three dots under its S but there is a single dot under the S of the proper name *Sökmen* in IV,4. The S's of *keseḳ*, of *sat(t)ım* ‘I sold’ in IV,5 and IV,6 and of a verb derived from a Persian loan, *usparladım* in IV,8, have double dots underneath. *Ishaq* in IV,5 has a single light dot under its S. *Sulman* in V,a1 and *siliḥdar* ‘man-at-arms’ in V,a5 both have three dots under their S. The title *teksin* in V,b7 is written with two or three dots under S. In V,b8 *sarıḡ* ‘blond’, two dots are visible under the S. *inisi* ‘his younger brother’ in V,b9 has three dots under the S, while the same word in IV,3 has two joined dots above its S. The last instance being an exception, all /s/-dots – mostly but not always three – are *under* the letter.

7. Diacritic Letters

A single Uyghur letter, Q, stands for [q], [ɣ], /^c/, and /h/. [q] can be marked by a superscribed single dot, in *Ishaq* in II,3c3, in *tanuqlar* ‘witnesses’ in IV,1 and in *qılsar* ‘if he does’ in IV,10. All other distinctive markings come from Arabic letters: *ḡayn* or ‘*ayn* for [ɣ] and /^c/, a small *ḡā* for /h/ and /x/ and small *ḡā* for /h/. In I,15, an Arabic *ḡayn* is visible under the Q of *oylumız* ‘our son’, in I,23 small *ḡayns* under the Q of *siċġan* ‘mouse’ and under the ‘*ayn* of *rabi^cal aḡir* (spelled like this, and with Q as fifth letter). *ḡallāj* ‘cotton carder’ in I,30 has a *ḡā*-like sign under its first letter; in V,b5, the diacritic *ḡā* is written *over* its first letter. The instances of *ḡaddi* ‘its border’ in II,c1, c5 and c7 have small *ḡās* under their first letter, but *ḡaddi* in IV,6 has something like an Arabic *ḡā* *above* it. In *Ishaq* of II,3c3 a small Arabic *ḡā* appears under the third letter. The Persian loan *baha* ‘price’ in II,c6 and twice in IV,7 has the Arabic letter *ḡā* above it.

⁷ It is also contrary to what we find in p. 6 of the *Dīvān Luġātī* ‘t-Turk: In a list of the “Turkic” alphabet which Maḡmūd offers with Arabic correspondences, the Uyghur letter corresponding to Arabic *šīn* has two dots under it, while the one corresponding to *šīn* has nothing.

taqaġu ‘chicken’ in II,b1 has an ‘*ayn*’ below its second Q. The names to be read as *‘Uθmān* in V,b2 and *‘Alī* in V,b4 have little ‘*ayns*’ under their first letters.

The early practice for writing Arabic texts is not limited to dots either: The *hamza* also, after all, comes from a small ‘*ayn*’. Grohman (1971: 43-46) mentions other signs, mostly simplified letters, which were placed under or over the *rasm* letters to ease their identification. Gronke (1986: 458), e.g., mentions a letter *fā* added above the *fā* of the text.

8. Conclusion

What is unique in the Turkic Yarkand documents is that the distinctive diacritics appear in texts which are not only in a different language from Arabic, but also are imposed on a different writing system – albeit used in situations where the Arabic language had a prominent position. The paper gives evidence for a widespread graphic practice in early Muslim epigraphic and manuscript texts for markings beyond the standard system, originally intended to prevent misreadings, but often also used when misreadings were unlikely.

References

- Clark, L. (2010). The Turkic script and the *Kutadgu Bilig*. In H. Boeschoten & J. Rentzsch (eds.) *Turcology in Mainz*. Wiesbaden.
- Çelebi Çam, S. (2025). *Kışaşü'l-Enbiyā (Rabġūzī) Tahran Nūshası (transkripsiyonlu metin). Tahran ve Londra nüshaları arası eş anlamlılık*. Ankara.
- Erdal, M. (1984). The Turkish Yarkand documents. *Bulletin of the School of Oriental and African Studies*, 47, 260–301.
- Erdal, M. (1993). *Die Sprache der wolgalbolgarischen Inschriften*. Harrassowitz. Wiesbaden.
- Grohmann, A. (1971). *Arabische Paläographie. 2. Teil: Das Schriftwesen. Die Lapidarschrift*. Böhlau Nachf.
- Gronke, M. (1986). The Arabic Yarkand documents. *Bulletin of the School of Oriental and African Studies*, 49(3), 454–507.
- Khakimzjanov, F. S. (1978). *Jazyk epitaŋii volġskix Bulgar*. Moscow.
- Khakimzjanov, F. S. (1987). *Epigrafičeskie pamjatniki volġskoj Bulgarii i ix jazyk*. Moscow.
- Moravcsik, G. (1983). *Byzantinoturcica II: Sprachreste der Turkvölker in den byzantinischen Quellen*. Brill.
- Onur, S. (2024). Which is the oldest manuscript of Rabġhūzī’s *Qisas al-Anbiyā*? *Türkiyat Mecmuası*, 34(2), 839–875.
- Orsatti, P. (2019). Persian language in Arabic script: The formation of the orthographic standard and the different graphic traditions of Iran in the first centuries of the Islamic era. In D. Bondarev, A. Gori, & L. Souag (eds.), *Creating standards: Interactions with Arabic script in 12 manuscript cultures* (pp. 39–72). De Gruyter.
- Ölmez, Z. (ed.). (2020). *Ebulgazi Bahadır Han. Şecere-yi Terākime*. TDK.

- Róna-Tas, A., & Fodor, S. (1973). *Epigraphica Bulgarica. A volgai bolgár-török feliratok* (Studia Uralo-Altaica, 1). Szeged.
- Salur, G. (2024). Türkçe Yarkent belgeleri. *Hacettepe Üniversitesi Türkiyat Araştırmaları Dergisi*, 41, 261–327.
- Tekin, T. (1988). *Volga Bulgar Kitabeleri ve Volga Bulgarcası*. TDK.
- Togan, A. Z. V. (1939). *Ibn Faḍlān's Reisebericht* (Abhandlungen für die Kunde des Morgenlandes XXIV, 3). Leipzig.

REVISITING VERBAL REFLEXIVITY: ON THE MORPHOSYNTAX AND ARGUMENT STRUCTURE OF IN DERIVED VERBS IN TURKISH

Zeynep ERDEMİR

Res. Asst. Istanbul Medeniyet University, Department of Linguistics,
Graduate Student, Boğaziçi University, Department of Linguistics, zeynepferdemir@gmail.com, ORCID:
0009-0000-5704-9444

Ümit ATLAMAZ

Asst. Prof., Boğaziçi University, Department of Linguistics, umit.atlamaz@bogazici.edu.tr, ORCID: 0000-
0003-1657-9654

Ömer DEMİROK

Asst. Prof., Boğaziçi University, Department of Linguistics omerfaruk.demirok@bogazici.edu.tr, ORCID:
0000-0002-2536-5247

Erdemir, Z., Atlamaz, Ü. & Demirok, Ö. (2025). Revisiting verbal reflexivity: On the morphosyntax and argument structure of IN derived verbs in Turkish. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 117–146). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X006>

ABSTRACT

In Turkish, verbal reflexives are formed with two different suffixes, *-Il* and *-In*, as observed in the literature. It has been argued that of these verbs, those derived with the suffix *-Il* express only figure reflexivity; while reflexives derived with the suffix *-In* express ground reflexivity (Key, 2021, 2025). We observe that the *-In* suffix appears in two additional contexts: (i) in building figure reflexivity and (ii) in some unergative verbs that do not express any reflexivity. Based on these observations, we propose that the *-In* suffix is essentially a realization of a verbal head that does not introduce arguments. These non-argument-introducing verbal heads emerge within the context of a head, which we call Ref and propose to be subject to contextual allosemy.

Keywords: Verbal reflexivity, Voice, Turkish

ÖZ

Türkçe’de dönüşlü fiillerin *-Il* ve *-In* olmak üzere iki farklı ek ile türetildiği alanyazında gözlemlenmiştir. Bu fiillerden *-Il* ekiye türetilenlerin yalnızca figür dönüşlülüğü ifade ettiği, *-In* ekiyle türetilen dönüşlülerin ise zemin dönüşlülüğü ifade ettiği savunulmuştur (Key, 2021, 2025). Biz bu çalışmada *-In* ekinin iki ayrı durumda daha görüldüğünü gösteriyoruz: (i) figür dönüşlülüğünün üretiminde ve (ii) dönüşlülük ifade etmeyen bazı özneli geçişsiz eylemlerde. Bu gözlemlerimize dayanarak, *-In* ekinin özünde, argüman getirmeyen fiil başların sesletimi olduğunu ileri sürüyoruz. Bu argüman getirmeyen fiil başlarının, Ref adını verdiğimiz ve bağlamdan anlam koşullanmasına tabi olan başın bulunduğu yapılarda ortaya çıktığını savunuyoruz.

Anahtar Sözcükler: Dönüşlü eylemler, Eylem çatısı, Türkçe

Introduction

This paper examines the verbal reflexive structures in Turkish. There are two primary ways to express reflexive meaning in Turkish. The first one is pronominal reflexives, which involve the Turkish self-anaphor *kendi* as exemplified in (1a)-(2a). The second is verbal reflexives, which derive the relevant argument structure within the verbal domain. Verbal reflexives in Turkish feature the suffixes *-In* or *-Il*¹.

Although the pronominal and verbal reflexives are often thought to be truth-conditionally equivalent and can indeed be used interchangeably, they have been shown to have different syntactic structures and semantic interpretations (e.g. Akkuş & Paparounas, 2024). The present study sets aside pronominal reflexives and focuses on the syntactic and semantic properties of verbal reflexives in Turkish, as illustrated in (1b) and (2b).

- (1) a. Ayça kendin-i çok öv-dü.
Ayça self-ACC much praise-PST.3SG
‘Ayça praised herself too much.’
b. Ayça çok öv-**ün**-dü.
Ayça much praise-**In**-PST.3SG
‘Ayça praised (herself) too much.’
- (2) a. Ali kendin-i yol-dan çekti.
Ali self-ACC road-ABL pull.back-PST.3SG
‘Ali pulled himself back from the road.’
b. Ali yol-dan çek-**il**-di.
Ali road-ABL pull.back-**Il**-PST.3SG
‘Ali pulled himself back from the road.’

¹ In this paper, we work within the framework of Distributed Morphology, where affixes are the surface realizations of abstract morphemes, which are syntactic terminals. Throughout, we use *-Il* and *-In* to represent the suffixes and IL and IN to represent the abstract morphemes.

Previous studies on Turkish verbal reflexives observed the intransitive nature of reflexive verbs (Akkuş & Paparounas, 2024) and even argued that reflexive verbs can be distinguished by the suffix that marks them, with *-In* and *-Il* reflexives corresponding to different argument structures (Key, 2021).

The *-Il* suffix, which appears in (2b), has been characterized as the realization of the non-active Voice morpheme IL because it is able to derive reflexive verbs within the verbal domain alongside passive, impersonal, and anti-causative verbs (Taneri, 1993; Göksel, 1993; Nakipoğlu, 1998; Zeyrek, 2006; Gündoğdu, 2017; Key, 2021, 2025). This type of distribution is commonly described as a syncretism, referred to as *u-syncretism* (Embick, 2004), which arises from similar structural properties of the non-active Voice head. This account, also supported by cross-linguistic observations (Spathas et al., 2015; Alexiadou & Schäfer, 2013), offers an explanation for the ambiguities observed among reflexive-passive (3), passive-anticausative (4), and even reflexive-anticausative (5) constructions marked by the morpheme in question.

- (3) Sporcu yarışma-dan çek-**il**-di.
Athlete competition-ABL pull-**II**-PST.3SG
Reflexive: ‘The athlete withdrew(oneself) from the competition.’
Passive: ‘The athlete was withdrawn from the competition.’
- (4) Kapı aç-**ıl**-dı.
Door open-**II**-PST.3SG
Passive: ‘The door was opened (by someone).’
Anti-causative: ‘The door opened (by itself).’
(From Göksel, 1993)
- (5) Adam koyunlar-a tak-**ıl**-dı.
Man sheep-DAT attach-**II**-PST.3SG
Reflexive: ‘The man attached himself after a flock of sheep.’
Anti-causative: ‘The man stumbled over a flock of sheep.’
- (6) Deniz hızlıca giy-**in**-di.
Deniz quickly wear-In-PST.3SG
Reflexive: ‘Deniz dressed herself up quickly.’

Although, the analyses in accord with the *u-syncretism* account considers the suffix *-In* an allomorph of the non-active Voice morpheme IL, it has also been observed that verbs marked with the suffix *-In* show no ambiguity, as illustrated in (6), which only has the reflexive interpretation. Thus, most studies consider it as the realization of an independent morpheme responsible for deriving reflexive verbs (Lewis, 2000; Göksel, 1993; Kornfilt, 2013). Lastly, the most recent studies have proposed that the morpheme IN functions as an Applicative-head deriving a specific type of reflexivity called *ground-reflexivity* (Key, 2021, 2025).

The main objective of this study is to present a new set of data that previous proposals for IN derived reflexive verbs fail to capture and show that these data constitute counterevidence to the aforementioned analyses. We then propose an analysis for IN derived verbs that aims to account for the complete set of data in an exhaustive and consistent manner. It is worth noting, though, that IL derived reflexives will remain outside the scope of this work.

More concretely, we make two proposals regarding verbal reflexives, particularly the derivation of IN derived verbs in Turkish. First, contrary to the previous proposals which associated the meaning of reflexivity with a certain morpheme, we argue that reflexivity in verbal structure is derived through a separate reflexivizing head, which we will call Ref-head in this study. Second, the suffix *-In* is in fact the realization of these verbal heads, introducing theta roles within the verbal domain. We argue that while theta-role introducing verbal heads in Turkish are generally null, they are realized by the suffix *-In* when they do not project a specifier, i.e., when they do not introduce any argument in the structure.

This paper is structured as follows: Section 2 introduces how Turkish verbal reflexives are analyzed in the literature. Section 3 presents our observations regarding the IN derived reflexive verbs and discusses their argument structure, guiding us towards an analysis. Section 4 fleshes out the analysis while Section 5 discusses further data, along with the predictions that the proposed analysis makes. Section 6 is the conclusion.

2. Previous Approaches to Turkish Verbal Reflexives

2.1 Argument structure to Turkish verbal reflexives

It is possible to express reflexivity through various methods, such as accidental coreference (7), pronominal reflexive structures (8) and verbal reflexives (9). While the first two methods observably constitute a transitive structure, verbal reflexives in Turkish have been shown to be syntactically intransitive and semantically monadic (Akkuş & Paparounas, 2024).

- (7) Aristoteles Büyük İskender-in hoca-sın-ı çok
Aristotle Great Alexander-GEN tutor-POSS.3SG-ACC much
öv-dü.
praise-PST.3SG
‘Aristotle praised Alexander the Great’s tutor a lot.’

Aristotle = Alexander the Great’s tutor

- (8) Aristoteles kendin-i çok öv-dü.
Aristotle self-ACC much praise-PST.3SG
‘Aristotle praised himself a lot.’ Aristotle = himself

- (9) Aristoteles çok öv-ün-dü.
Aristotle much praise-REFL-PST.3.SG
‘Aristotle praised himself a lot.’ Aristotle = AGENT = THEME

This observation is significant because it tells us both types of analysis are possible for verbal reflexive structures. Although these verbs are intuitively perceived as intransitive, crosslinguistic research on verbal reflexive verbs has shown that they can also be analyzed as syntactically transitive and semantically dyadic. In this approach, the derivation usually involves DP-movement and binding between the DP and a reflexive pronoun (Heim & Kratzer, 1998).

Akkuş and Paparounas (2024) demonstrate the intransitive nature of Turkish verbal reflexives through several diagnostics, including proxy readings, ellipsis ambiguities, focus-bound readings, *de dicto* interpretations, and indirect causative structures. Their analysis of intransitive verbs puts forward no binding relation and treats the internal argument as the sole argument of the verb. For transitive reflexives, the DP argument and the reflexive pronoun are distinct entities, and it is possible to manipulate one of the arguments, leaving the other one unaffected.

For reasons of space, we will illustrate the differences between overt anaphors and verbal reflexives, drawing from the examples discussed in Akkuş and Paparounas (2024). Consider the sentences (10a) and (10b) in the given context. In (10a) the embedded clause has a pronominal reflexivity while (10b) features a reflexive verb, lacking an overt pronominal object. The observation is that while (10b) is not judged as true in the given context, (10a) is judged true. This shows that reflexive verbs only permit *de se* construal while pronominal reflexives allow *de re* readings as well.

(10) *Context*: Ali, the leader of a religious cult, must once a year ceremonially wash the oldest member of the community. He hasn't realized that, as of this year, he himself is the oldest member. On the day, he announces: 'I must now wash the oldest member of the community!'

- a. Ali kendini yıka-mak isti-yor.
Ali self.ACC wash-INF want-PROG
'Ali wants to wash himself.'
- b. Ali yıka-n-mak isti-yor.
Ali wash-REF-INF want-PROG
'Ali wants to wash himself.'

(From Akkuş & Paparounas, 2024)

In the analysis proposed in Akkuş and Paparounas (2024), reflexive verbs are assumed to be unaccusative in their base, underived form, which means that they have only an internal argument as their sole argument. They propose a reflexive variant of the Voice head which they call Voice_{REFL}. This head does not introduce an external argument to the structure, but it does assign an AGENT role. The sole DP of the reflexive verb is interpreted in its base position as the internal argument, which naturally carries the THEME role.

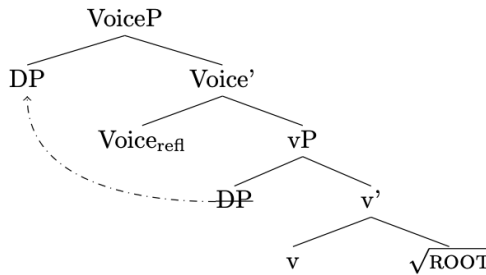
This assumption accounts for the unaccusative-like behavior and intransitivity of Turkish verbal reflexives.

Reflexive meaning arises through the movement of the internal argument to the external argument position. In this way, the single argument of the verb comes to bear both the THEME and the AGENT roles. This movement analysis, although it remains unclear how to interpret the resulting logical form as desired, intends to capture two main points:

- (i) how reflexive meaning is derived compositionally, and
- (ii) how the Turkish verbal reflexives display both unaccusative and unergative behavior under different syntactic diagnostics.

The details of the proposal can be observed more clearly in the syntactic tree illustrated in (11).

(11)



(From Akkuş & Paparounas 2024)

It is also important to note that Akkuş and Paparounas (2024) do not distinguish between IL and IN derived reflexives in terms of the argument structures they are licensed in, applying all their tests uniformly to both types of verbs. The present study, however, will focus on the distinction between IL and IN morphemes, focusing on IN derived verbs.

2.2 Reflexive morphology

The reflexive morphemes IL and IN in Turkish have either been analyzed as uniformly with the passive and anti-causative morphemes under the *u-syncretism* account (Embick, 2004) or they have been taken to be distinct morphemes. Key (2021) was the first to argue that Turkish uses these morphemes to build distinct types of reflexivity. He observed that the two suffixes are in fact not entirely in complementary distribution (with respect to the phonological contexts they appear in), as they reveal their underlying forms in post-voiceless context (Table 1). Rather, it appears that there are indeed two distinct suffixes *-Il* and *-In*, whose phonological contrast is neutralized along with {-l}-final and V-final stems.

Table 1. Phonological environment of -Il and -In suffixes

Post Vowel		Other	
<i>tıka-n</i> (unaccusative)	‘to get clogged’	<i>çek-in</i> (reflexive)	‘to hold oneself back’
<i>koru-n</i> (reflexive)	‘to protect oneself’	<i>çek-il</i> (reflexive)	‘to pull oneself back’
		<i>tık-in</i> (reflexive)	‘to get oneself stuffed’
		<i>tık-ıl</i> (reflexive)	‘to get oneself tucked in’
Post /-l/ Phoneme			
<i>del-in</i> (unaccusative)	‘to get pierced’		
<i>sal-in</i> (reflexive)	‘to oscillate’ (lit: to release oneself)		

Building on this observation, Key further identified different argument structures that seem to be dependent on the morpheme that we observe on the verb. Even when the verb stems are identical, the two morphemes can derive reflexive verbs with distinct argument structures. Some of the minimal pairs Key reported are illustrated in (12)-(13).

- (12)a. Ahmet ağac-a sar-**ıl**-dı.
 Ahmet tree-DAT wrap-**REF**-PST.3SG
 ‘Ahmet hugged (‘wrapped himself around’) a tree.’

- b. Ahmet havlu-yu sar-**ın**-dı
 Ahmet towel-ACC wrap-**REF**-PST.3SG
 ‘Ahmet wrapped the towel wound himself.’

- (13)a. Küçük yatağ-a tık-**ıl**-dım.
 Small-DIM bed-DAT stuff-**REF**-PST.1SG
 ‘I stuffed myself into the tiny bed.’

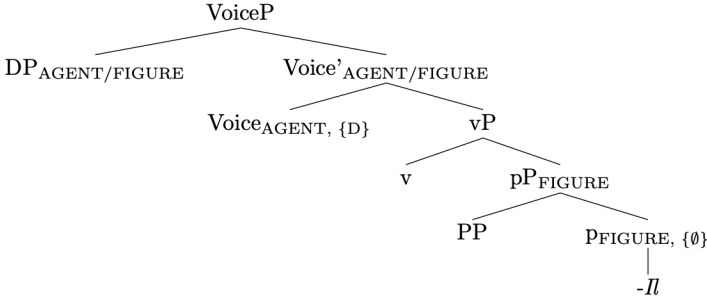
- b. Bir şey-ler tık-**ın**-dım.
 A think-PL stuff-**REF**-PST.1SG
 ‘I stuffed myself (eating) this-and-that’

(adapted from Key, 2021)

Key’s (2021) analysis draws on the figure-ground distinction (Talmy, 1978). He argues that the IL derived reflexives encode the co-reference between the agent and the figure within the argument structure of the verb. Consequently, the sole argument of such verbs bears both AGENT and THEME roles. In his examples, the reflexive meaning is established between these two roles referring to the same entity *Ahmet* in (12a) and first-person singular in (13a). Following previous works, he referred to this class as *figure reflexives* (Wood, 2023). These verbs are derived from transitive bases and

become de-transitivized once the IL morpheme is attached. The derivation of the figure reflexive verbs can be seen in (14).

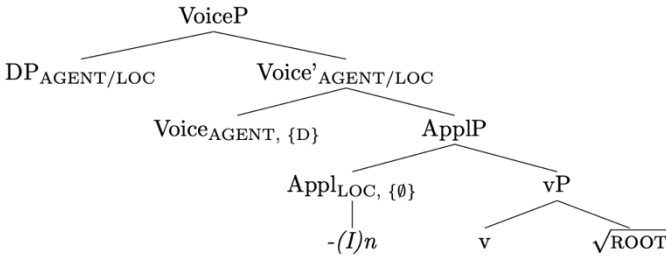
(14) Delayed Gratification for Turkish Figure Reflexives



(From Key, 2021)

In contrast, IN derived reflexives express a co-reference between the agent and the ground on which the motion of the figure applies. The presence of the IN morpheme indicates that an applied argument is introduced to the structure by a specifier-less Appl_[-D, +GROUND] head bearing a GROUND² role (Key, 2025). In this way the GROUND role is assigned to the external argument through a process called Delayed Gratification (15). Thus, the agent and the applied argument end up referring to the same entity. Since reflexivity does not target the THEME in ground reflexives, some IN derived reflexive verbs may also appear as transitive, licensing objects with the THEME role.

(15) Delayed Gratification for Turkish Ground Reflexives/Applicatives



(From Key 2021, 2025)

The Delayed Gratification analysis by Key (2021) successfully captures the distinct behavior of the IN derived reflexives in Turkish, which

² Here, the term *ground* serves as an umbrella term for the thematic roles BENEFACTIVE, MALEFACTIVE, and LOCATION (Key, 2025).

the previous accounts could not capture. However, this analysis makes two key predictions regarding IN derived verbs:

- (i) that the Applicative-head in Turkish is realized by an overt morpheme, implying that all IN derived verbs in Turkish should express applied meanings.
- (ii) that the IN reflexives only express ground reflexives but never figure reflexives.

In the following section, we present counterexamples to both predictions and motivate a new analysis to be fleshed out in section 4.

3. Against the Analysis of IN as the Exclusive Marker of Ground Reflexivity

Both of the two recent studies on Turkish reflexives reviewed in the previous section agree that reflexive verbs derived with the morpheme IL express *figure reflexivity*, while those derived with the morpheme IN express *ground reflexivity*. While Akkuş and Paparounas (2024) propose to account for this distribution through DP-movement, Key (2021, 2025) introduces a separate Applicative-head to capture the GROUND role expressed in IN reflexives.

In what follows we show that once we look beyond the data that has already been reported, we find patterns that existing analyses fail to account for.

3.1 IN can build figure reflexives

We show that IN derived reflexive verbs do not always express ground reflexivity. In some cases, they can display figure reflexivity (*pace* Key 2021, 2025). In other words, IN derived reflexive verbs in general can encode both agent-ground (16) and agent-theme coreference (17). This can be illustrated by the verb *sürün-* ‘to apply on’ where a single form is ambiguous between a figure and a ground reflexive.

- (16) a. Misafir kendin-e kolonya sür-dü.
 Guest self-DAT cologne apply-PST.3SG
 ‘The guest applied cologne on herself.’

- b. Misafir kolonya sür-**ün**-dü.
 Guest cologne apply-**REF**-PST.3SG
 ‘The guest applied cologne on herself.’

The guest = AGENT = GROUND

- (17) a. Çocuk kendin-i yer-e sür-dü.
 Child self-ACC ground-DAT apply-PST.3SG
 ‘The child dragged herself on the ground.’

- b. Çocuk yer-de sür-ün-dü.
 Child ground-LOC apply-REF-PST.3SG
 ‘The child dragged herself on the ground.’

The child = AGENT = THEME

Since Key (2025) adopts the u-syncretism account in his analysis of the IL morpheme, the agent-theme relationship of the IL derived verbal structures is consistently accounted for under his approach. However, the Delayed Gratification account, adopted for deriving ground-reflexivity of IN derived verbs, fails to capture the fact that IN derived verbs can also express theme-agent relations. This is so because Key treats the morpheme IN as a realization of an Applicative-head introducing the GROUND role. However, except for a few minimal pairs such as *sürün-* ‘to apply on’, none of the following verbs express an agent-ground relation. A list of the IN derived figure reflexive verbs that we have observed is provided in Table 2³.

Table 2. A list of IN derived figure reflexive verbs in Turkish

Verbs Stem		Figure Reflexive	
<i>dit-</i>	‘to pick apart’	<i>did-in</i>	‘to wear oneself out’
<i>sürt-</i>	‘to rub’	<i>sürt-ün</i>	‘to rub oneself against’
<i>silk-</i>	‘to shake off’	<i>silk-in</i>	‘to shake oneself off’
<i>ölç-</i>	‘to measure’	<i>ölç-ün</i>	‘to measure one’s height’
<i>sık-</i>	‘to squeeze’	<i>sık-in</i>	‘to squeeze oneself’
<i>yırt-</i>	‘to tear apart’	<i>yırt-in</i>	‘to strive’ (lit: to tear oneself apart)
<i>sak-</i> (bound)	‘hidden’	<i>sak-in</i>	‘to refrain oneself’
<i>çek-</i>	‘to pull’	<i>çek-in</i>	‘to hesitate’ (lit: to hold oneself back)
<i>kas-</i>	‘to tighten	<i>kas-in</i>	‘to swagger’ (lit: to straighten one’s posture)
<i>kaç-</i>	‘to run away’	<i>kaç-in</i>	‘to avoid’
<i>çirp-</i>	‘to whisk’	<i>çirp-in</i>	‘to flounder’ (lit: to whisk oneself)
<i>döv-</i>	‘to beat up’	<i>döv-ün</i>	‘to beat one’s chest’
<i>yak-</i>	‘to burn’	<i>yak-in</i>	‘to whimper’ (lit: to burn oneself)
<i>ger-</i>	‘to stretch’	<i>ger-in</i>	‘to stretch oneself out’
<i>ov-</i>	‘to scrub’	<i>ov-un</i>	‘to scrub oneself’
<i>sil-</i>	‘to wipe’	<i>sil-in</i> ⁴	‘to wipe one’s body’
<i>sür-</i>	‘to apply on’	<i>sür-ün</i>	‘to crawl’ (lit: to apply oneself to the ground)
<i>vur-</i>	‘to hit’	<i>vur-un</i>	‘to hit oneself’

³ For the verbs list on Table 2, the following studies on Turkish verbal reflexives literature are consulted such as Arslan, 2022; Baturay-Meral & Meral, 2018; Koşaner, 2005; Gülsevin, 2022; Üstünova, 2016 in addition to our own observations.

⁴ Since post /-l/ phoneme is an undistinctive environment the underlying form of the suffix could be interpreted as *-ll* which can also result in the passive meaning ‘to get erased’. The reflexive verb *silin-* here is specifically used for ‘getting a shower through wet wipes’ in daily language.

<i>tep-</i>	‘to kick’	<i>tep-in</i>	‘to stamp oneself on the ground’
<i>gör-</i>	‘to see’	<i>gör-ün</i>	‘to make oneself visible’
<i>öv-</i>	‘to praise’	<i>öv-ün</i>	‘to praise oneself’
<i>kır-</i>	‘to break’	<i>kır-in</i>	‘to perform a dance’ (lit: to break oneself)

Notably, most of these verbs do not even contain any ground arguments, which challenges Key’s claim that IN is an Applicative-head that introduces a ground argument to the structure. Putting that aside, the DP-movement analysis could in principle account for these verbs without any ground arguments. Nevertheless, the DP-movement cannot derive figure reflexives containing a ground argument, as the argument introduced by Applicative-head would cause an intervention in the structure between the agent and the theme arguments. These are verbs like *sürtün-* ‘to rub oneself against’ and *tepin-* ‘to stamp oneself on the ground’ in which a ground argument is licensed. Examples such as (18)-(19) have ground denoting DPs like *duvar* ‘a wall’ or *masa* ‘a desk’; however, crucially they do not encode any agent-ground relation.

- (18) Kedi duvar-a sürt-**ün**-dü.
 Cat wall-DAT rub-**REF**-PST.3SG
 ‘The cat rubbed itself against the wall’

The cat = AGENT = THEME

- (19) Öğrenci masa-da tep-**in**-di.
 Student desk-LOC kick-**REF**-PST.3SG
 ‘The student stamped on the desk.’

The student = AGENT = THEME

Given that the existing analyses fail to account for the fact that IN derived reflexives can express both figure and ground reflexivity, a comprehensive analysis capturing the more complex distribution indicated by the complete set of data is lacking.

3.2 IN does not always encode reflexivity

The IN morpheme does not necessarily occur solely within the structures which we conventionally perceive as reflexives (see section 4 for the complete set of data). To illustrate, the example in (20) does not convey the reflexive meaning (neither figure nor ground) even though it appears to be marked with the *-In* suffix on the surface. Since direct objects of verbs are marked with an Accusative Case in Turkish (Kornfilt, 2013; Lewis, 2000), the transitive verb *gez-* ‘to visit’ in (20) expectedly has an accusative-marked internal argument. However, *gez* ‘visit’ → *gez-in* ‘to wander’ clearly de-transitivizes the verb, as it is ungrammatical to express direct objects within the structure. Here, the sentence (21) does not refer to a specific ground on

which the students wandered upon. The data further indicates that the presence of IN correlates with another purpose rather than a reflexivizing function, which needs to be discussed more thoroughly.

- (20) Öğrenci bahçe-yi gez-di.
 Student garden-ACC wander-PST.3SG
 ‘The student visited the garden.’

- (21) Öğrenci (*bahçe-yi) gez-in-di.
 Student garden-ACC wander-In-PST.3SG
 ‘The student wandered (*intended*: in the garden).’

Two inferences can be drawn from this. First, the morpheme IN seems to be altering the argument structure of the verb, implying that internal argument is no longer syntactically expressed. It is also important to notice that this is a common property it shares with the IN forms that are genuinely reflexive. Second, the data provides strong evidence that the function of the morpheme IN and the reflexivizing function must be handled separately from one another. In section 4, we present an analysis that accommodates both inferences.

4. Proposal: Deriving IN Verbal Structures in Turkish

In this section we provide a novel analysis that accounts for the distribution of the IN morpheme as well as its functions that previous analyses fail to capture. This section is structured as follows. First, we provide the structure we assume for canonical transitive verbs in Turkish. Second, we show the derivation of the IN derived reflexives. Next, we provide an analysis of non-reflexive IN verbs in Turkish, which we will call IN unergatives. Some predictions of the proposal are discussed in Section 5.

4.1 IN figure reflexives

For our analysis, adopting the Neo-Davidsonian framework (Davidson, 1967), we assume all verbs to denote predicates of events regardless of how they are classified conventionally in terms of their transitivity (22).

- (22) a. $\llbracket \text{ÖV} \rrbracket = \lambda e. \text{praise}(e)$
 b. $\llbracket \text{GİY} \rrbracket = \lambda e. \text{wear}(e)$
 c. $\llbracket \text{SÜRT} \rrbracket = \lambda e. \text{rub}(e)$

In addition, we assume that all arguments form thematic relations with the event via verbal heads introduced to the structure (Borer, 1994). The AGENT role is introduced by the Voice-head and thematic roles such as THEME or GROUND are by v_{theme} and v_{appl} heads respectively (23). The denotations of these heads are provided below.

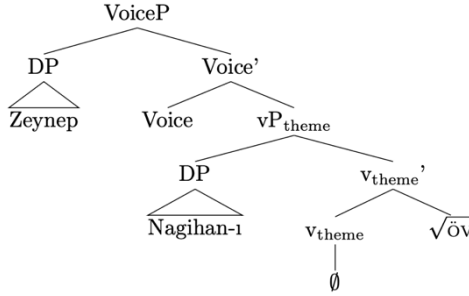
- (23) a. $\llbracket \text{Voice} \rrbracket = \lambda P_{\langle v, t \rangle}. \lambda x. \lambda e. P(x)(e) \wedge \text{agent}(e) = x$

- b. $\llbracket v_{\text{theme}} \rrbracket = \lambda P_{\langle v, t \rangle}. \lambda x. \lambda e. P(x)(e) \wedge \text{theme}(e) = x$
c. $\llbracket v_{\text{appl}} \rrbracket = \lambda P_{\langle v, t \rangle}. \lambda x. \lambda e. P(x)(e) \wedge \text{ground}(e) = x$

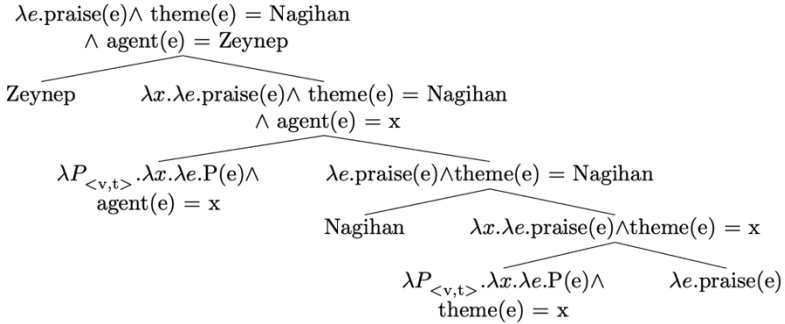
With all these assumptions in place, we go over the syntactic representation (25) and the semantic derivation (26) of the sentence in (24) that features a transitive verb.

- (24) Zeynep Nagihan-ı öv-dü.
Zeynep Nagihan-ACC praise-PST.3SG
Zeynep praised Nagihan.'

(25)



(26)



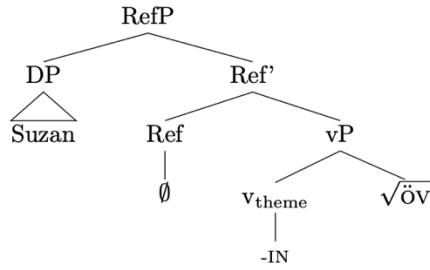
The verbal root *öv-* ‘praise’ combines with the v_{theme} head that both equips the verb with the THEME role in semantic derivation and opens a specifier position in syntax for an argument that will acquire that role. This way *Nagihan* is introduced to the structure as an internal argument. The same is true for the external argument *Zeynep*, which saturates the AGENT role that the Voice head introduces. It is important to notice that v_{theme} head is null here, thus the verb is not marked with any suffix indicating a change in valence on the surface (27).

- (27) a. $\llbracket v_{\text{theme}} \rrbracket = \lambda P_{\langle v, t \rangle}. \lambda x. \lambda e. P(x)(e) \wedge \text{theme}(e) = x$
b. $\llbracket v_{\text{theme}} \rrbracket \leftrightarrow \emptyset$

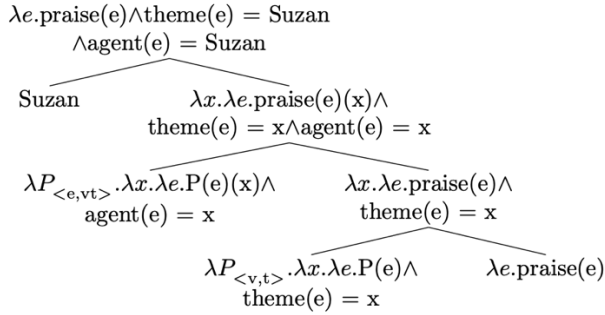
As for IN derived verbal reflexives, we assume their syntactic intransitivity and semantic monadicity, following Akkuş and Paparounas (2024). A sentence containing one of the IN derived figure reflexive verbs we have shown in the previous section is given in (28). Here, the verb *övün-* expresses the event of praising oneself. The syntactic representation (29) and the semantic composition (30) of the sentence are also provided below.

- (28) Suzan öv-ün-dü.
 Suzan praise-IN-PST.3SG
 ‘Suzan praised herself.’

(29)



(30)



As in the previous structure, the verb is combined with the head v_{theme} and gains the THEME role thanks to it. However, unlike what we saw with the transitive structure, no DP is introduced into Spec- vP_{theme} . We propose that the key piece that establishes the reflexive meaning is what we call the Ref-head, which combines with vP_{theme} .

This analysis of IN derived verbs aims to account for the occurrence of the IN morpheme indirectly. In particular, we associate the occurrence of IN morpheme with the presence of a null reflexivizer head in the structure

which we call Ref⁵. The denotation of the Ref-head is given in (31). Notably, the Ref-head is always null.

- (31) a. $\llbracket \text{Ref} \rrbracket = \lambda P_{\langle e, vt \rangle}. \lambda x. \lambda e. P(x)(e) \wedge \text{agent}(e) = x$
 b. $\llbracket \text{Ref} \rrbracket \leftrightarrow \emptyset$

As the Ref-head combines with an unsaturated vP_{theme} , its complement remains specifier-less. Consequently, the THEME role is not, at this point, identified with any argument. This pattern, we argue, results in a contextual allomorphy among v -heads. When they do not express their argument overtly in the syntax, they are spelled out as IN and when they do, they are realized by a null exponent (32). The separation of the reflexivizing function and the morpheme IN was one of the main points motivating our analysis. The Ref head makes this possible, not only mediating the contextual allomorphy in how the v -head is realized but also establishing identity between an unsaturated thematic role and an agent of the event.

- (32) a. $\llbracket v_{\text{theme}} \rrbracket = \lambda P_{\langle v, t \rangle}. \lambda x. \lambda e. P(x)(e) \wedge \text{theme}(e) = x$
 b. $\llbracket v_{\text{theme}} \rrbracket \leftrightarrow \text{-IN/} ___ \llbracket \text{Ref} \rrbracket$
 c. $\llbracket v_{\text{theme}} \rrbracket \leftrightarrow \emptyset$ (elsewhere)

A key property of the Ref-head we propose is that it introduces the AGENT role. In particular, the Ref-head takes a function from individuals to predicates of events and asserts that the argument position to be filled by the sole argument of the verb is also the AGENT of that event. Thus, it makes the external argument to fill a position on which it is associated with both AGENT and THEME roles at the same time. For example, the [Spec, RefP] is filled by *Suzan* in (30) and the whole structure is interpreted as follows: *there is a praising event whose theme is Suzan and whose agent is Suzan*. These are the correct truth-conditions, capturing the agent-figure reflexivity. It is also important to notice that our analysis, incorporating the Ref head, assumes the viability of agentive structures without a VoiceP layer on top.

In the next part of this section, we discuss the properties of the ground reflexives which are also marked through the same IN morpheme.

4.2 Turkish ground reflexives

We have shown how Turkish can derive IN derived figure reflexives in the previous section. However, the ground reflexivity is what previous studies have observed when it comes to the IN derived verbs that allow reflexive reading. Recall that Key (2021, 2025) argues that IN realizes an Applicative-head. While acknowledging his observation regarding the relation between IN morpheme and ground reflexivity, we argue that our observations necessitate a more comprehensive approach to the verbal heads in Turkish.

⁵ We will later extend the Ref feature to the analysis of the IN unergatives, arguing for a variation in its meaning.

We argue the IN morpheme found on ground reflexive verbs to be projection of another v -head, the v_{appl} head to be specific (as suggested by Key), that renders the expression of the applied argument in the structure possible. The denotation for the v_{appl} is given in (33a). This head, too, is spelled out as IN in the context of the Ref-head. Recall that the context in which the Ref is licensed are also the contexts where we have spec-less v Ps. This is simply another environment where this situation arises.

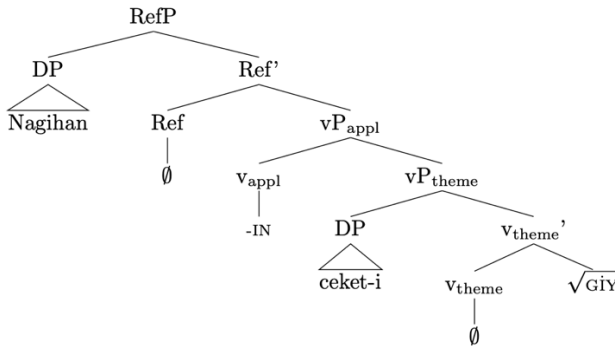
- (33) a. $\llbracket v_{\text{appl}} \rrbracket = \lambda P_{\langle v, t \rangle}. \lambda x. \lambda e. P(x)(e) \wedge \text{ground}(e) = x$
 b. $[v_{\text{appl}}] \leftrightarrow \text{-IN/} ___ [\text{Ref}]$

The ground reflexive meaning expressed in sentence (34) is derived in a similar manner. The verb *giyin-* ‘to dress oneself’ expresses ground reflexivity, which means the sole argument in this structure bears the AGENT and GROUND roles. This sentence additionally contains a THEME argument. As it can be observed from the semantic derivation in (36), the verb *giy-* ‘to dress up’ respectively gains the THEME role introduced by the v_{theme} -head which the internal argument *ceket* ‘the jacket’ saturates.

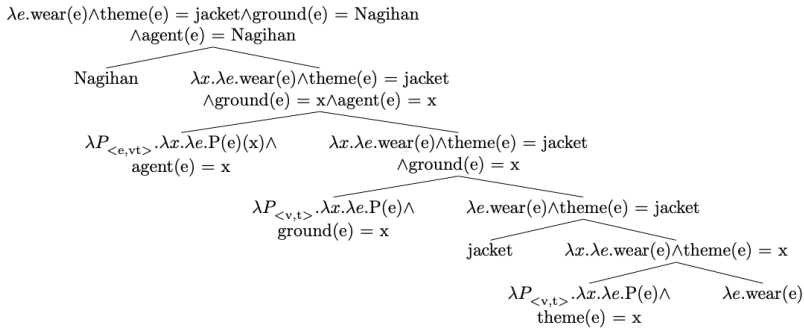
- (34) Nagihan ceket-in-i giy-in-di.
 Nagihan jacket-POSS-ACC wear-IN-PST.3SG
 ‘Nagihan dressed herself with her jacket.’

Next, the vP_{theme} is attached to v_{appl} introducing the GROUND role. Since the Ref-head is present, the specifier-less v_{appl} head is realized as IN here, as well. Lastly, the Ref-head attaches to the structure, and it re-defines the argument position to bear the roles AGENT and GROUND. At the end of the derivation, ground reflexivity is interpretable in the structure (36) as follows: *There is an event of wearing whose theme is the jacket and whose agent and ground is Nagihan.*

(35)



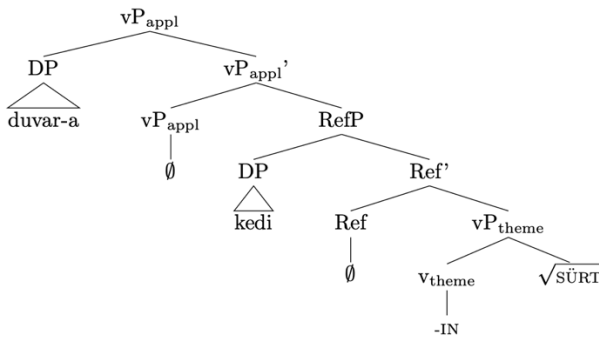
(36)



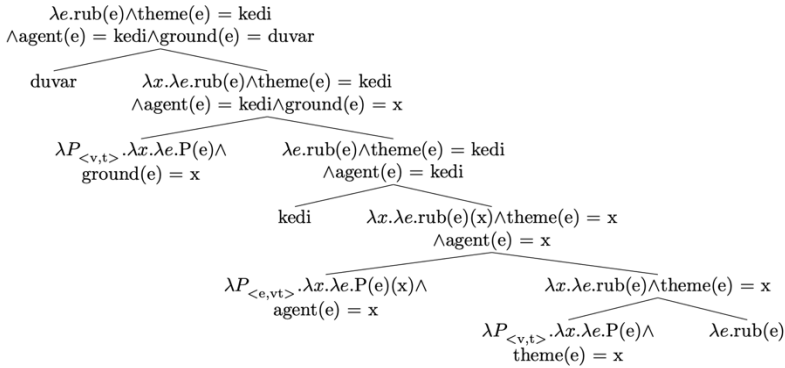
Verbs like *ıepin-* 'to stamp oneself on the ground' and *sürtün-* 'to rub oneself against', which contain arguments bearing the GROUND role, but still express figure reflexivity also follow from this analysis. The syntactic structure (38) and the semantic derivation (39) of the relevant type of figure-reflexive is given below. Given that these structures have ground arguments, it would be problematic to account for the agent-theme co-reference via DP-movement. The ground argument would be an intervenor for movement to theta position.

(37) Kedi duvar-a sürt-ün-dü.
Cat wall-DAT rub-In-PST.3SG
'The rubbed itself against the wall.'

(38)



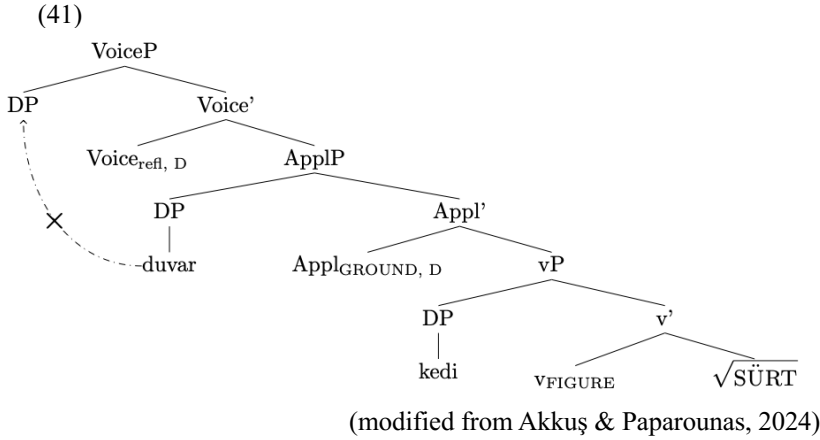
(39)



In (39), we can better observe that even though the ground argument is higher than the theme argument, the reflexivity is formed through an AGENT-THEME relation. This data point is best explained by a separate RefP that establishes reflexivity in the position it is introduced to the structure. The figure reflexivity is derived the same way as in (30). At this point it is important to notice that, while the v_{theme} head in (39) is spelled out as IN in the environment of Ref-head, the other verbal head v_{appl} which is right above the RefP is spelled out as null. Here, we observe the same allomorph for v_{appl} that we observed for v_{theme} (40).

- (40) a. $\llbracket v_{\text{appl}} \rrbracket = \lambda P_{\langle v, t \rangle}. \lambda x. \lambda e. P(x)(e) \wedge \text{ground}(e) = x$
 b. $[v_{\text{appl}}] \leftrightarrow \emptyset$
 c. $[v_{\text{appl}}] \leftrightarrow \text{-IN/___ [Ref]}$

That the morpheme IN can also derive figure reflexives in structures that contain a ground argument poses a problem for the DP-movement analysis as briefly mentioned previously in Section 3 (Akkuş & Paparounas, 2024) for the following reason: When a ground argument is overtly present in Spec- $v_{app}P$ position, it would occupy a higher position than the THEME argument. Here, contrary to what the data has shown, the DP movement analysis incorrectly assigns the GROUND role to the external argument, rather than the THEME role under the assumption that this is an A-movement (41). Consequently, this morpheme appearing in two different verbal reflexive structures makes it difficult to analyze IN as the head of a certain functional projection, since it seems capable of occurring in more than one position in the structure.



4.3. The IN unergatives

The two previous parts of this section discussed the syntactic structure and derivation of reflexive verbs marked with IN morpheme. The present study aims to include an analysis of IN derived verbs that are not reflexive that we reviewed in Section 3. Recall that there are a certain set of verbs marked with IN in Turkish, that are not expressing any reflexive meaning in the sense we described above.

- (42) a. Suzan bahçe-yi gez-di.
 Suzan garden-ACC visit-PST.3SG
 ‘Suzan student visited the garden.’
- b. Suzan gez-in-di.
 Suzan visit-IN-PST.3SG
 ‘Suzan wandered.’

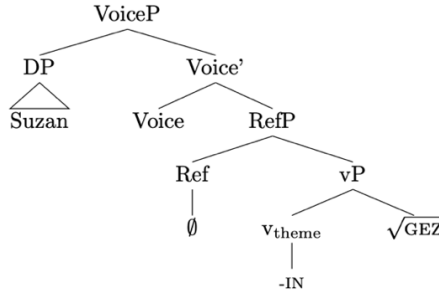
Our analysis for these structures will extend the function of the Ref-head and argue for root-sensitive contextual allosemy in the interpretation of the Ref-head. A complete list of these verbs is given in Table 3 below. In these structures, the Ref does not hold the reflexivizer function but an (anti)passive-like existential semantics. Hence, the Ref-head will have an alternative semantic denotation in the context of a certain class of roots, as shown in (43).

- (43) $\llbracket \text{Ref} \rrbracket = \lambda P_{\langle e, vt \rangle}. \exists x. \lambda e. P(x)(e) / _ \{ \text{BAK, GEZ, TAP, DEĞ, EŞ, SÖYLE, BAĞIR} \}$

We previously argued for the emergence of IN morpheme in the environment of the Ref-head and these examples do not constitute counterexamples. The v_{theme} head does not project the argument within the structures that contain the Ref head and is therefore spelled out as IN, as it can be more clearly observed in the syntactic representation of the sentence (42b).

Like its role in reflexives, the Ref-head introduces an argument into the structure when deriving the IN unergatives in (45). However, instead of assigning the AGENT role, Ref-head existentially binds the argument, preventing it from projecting as a full DP. This mechanism resembles what we observe in passive and anti-passive structures, in which an argument role is semantically present but does not correspond to a DP that saturates that role. Consequently, speakers perceive the internal argument to be more generalized or abstract, which existential saturation captures.

(44)



(45)

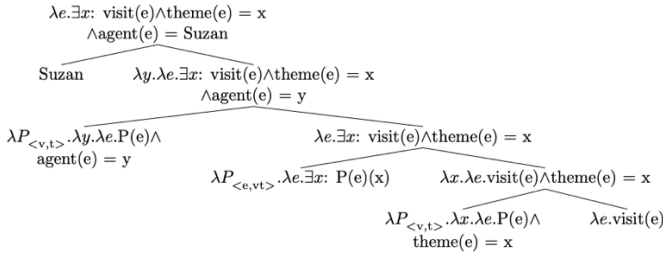


Table 3. A list of IN unergative verbs

Verbs Stem		IN Unergatives	
<i>bak-</i>	‘to look at’	<i>bak-in</i>	‘to seek’
<i>tap-</i>	‘to adore’	<i>tap-in</i>	‘to worship a deity’
<i>değ⁶-</i>	‘to touch’	<i>değ-in</i>	‘to mention’
<i>bağır-</i>	‘to shout at’	<i>bağır-in</i>	‘to scream’
<i>eş-</i>	‘to dig up’	<i>eş-in</i>	‘to scratch’
<i>gez-</i>	‘to travel’	<i>gez-in</i>	‘to wander’
<i>söyle-</i>	‘to tell’	<i>söyle-n</i>	‘to complain’ (lit: to tell something without an audience)

⁶ The verb *değ-* ‘to touch upon something’ could be both used as physically or abstractly in terms of its meaning.

To summarize, in this section, we provided an analysis of the syntactic structures and semantic derivations of IN derived verbs in Turkish. We introduced a head which we called Ref, to exhaustively generate reflexive meanings in these verbs. We also associated the overt spelling of the IN morpheme not directly with the Ref but with its presence in the structure. More generally, our analysis therefore takes IN to be the realization of spec-less v -heads (which can arise in the context of the proposed Ref-head). We also extended our analysis to what we called IN unergatives, arguing that they arise from a contextually available alternative existential/non-reflexive interpretation of the Ref-head.

Our analysis, while comprehensively accounting for our current data set, also makes some predictions, which we discuss in some detail in the next section.

5. Predictions & Discussion

Section 4 presented our analysis of IN derived verbs. The present section discusses the empirical predictions following from our analysis and examines how the new set of data is accounted for within the proposed framework. Based on the data, we have sufficient evidence to suggest that reflexive verbs in Turkish exhibit distinct syntactic structures depending on the morpheme with which they are derived. In this section we discuss our current data set more closely and illustrate their properties accordingly with the predictions made by our analysis.

5.1 The IN figure reflexives

In our analysis the morpheme IN realizes ‘expletive’ verbalizer heads, which do not introduce an argument into the structure⁷. We have also shown that this same $v_{[-D]}$ -head may exhibit an allosemantic variation, introducing either a THEME or a GROUND role depending on its context. In the context of the verb roots that are of the un-reflexivizable type, the Ref head functions to suppress their internal argument by existentially saturating it. In the presence of this version of the Ref-head which does not introduce the AGENT role, we expect the derivation to include an active VoiceP layer that introduces an external argument bearing the AGENT role. If a verb root is of the reflexivizable type, the Ref-head in its interpretation introduces the AGENT role to the relevant DP argument it introduces to the structure in its spec. In both cases, the resulting IN derived verbs are predicted to behave as intransitive verbs and to exhibit unergative properties, as no internal argument introduced by the v -head is in the structure.

So far, we have assumed that Turkish verbal reflexives are single-argument verbs, as clearly demonstrated by Akkuş and Paparounas (2024) and this is not a property that depends on the morpheme that derived them. What

⁷ The term ‘expletive’ here is defined as a syntactic head that does not project an argument to its Spec position. This does not apply to the semantic properties of the verbal heads mentioned here as both the v_{theme} and v_{appl} opens an argument slot in the semantic derivation.

remains to be determined is the nature of the single argument of the IN derived verbs and the position of this sole argument within argument structure. Following the methods adopted in previous literature (Sezer, 1991; Nakipoğlu 1998, 2001; Acartürk & Zeyrek 2010; Özkaragöz, 2011; Baturay-Meral & Meral, 2018; Kornfilt, 2009; Dikmen et al., 2022), we can apply several tests to examine whether these verbs display unergative or unaccusative behavior.

Examining an asymmetry in the distribution of genitive subjects in nominalizations provides a way to test the presence of an external argument in Turkish verbal structure (Kornfilt, 2009; Dikmen et al., 2022). Under this test, nonspecific subjects of embedded (nominalized) unaccusative verbs can remain caseless (or get genitive marking). However, with unergative verbs, subjects are ungrammatical in their caseless forms and must bear the genitive case as in (46).

Sentences formed with an unaccusative verb such as *boğul-* ‘to drown’ allows both bare noun subjects and genitive marked subjects in the embedded clause. In contrast, the IN-figure reflexives display another pattern: while the genitive marked embedded subjects are grammatical (46c), bare noun subjects render the sentence ungrammatical (46d). This pattern provides evidence that IN derived verbs all need to include an argument that bears the AGENT role, namely either an active VoiceP layer, or a RefP layer that maps an argument to the AGENT role, as also proposed in the current analysis.

(46) Geçen ay bu havuz-da,
Last month this pool-LOC
‘Last month in this pool,’

a. bir yüzücü boğul-duğ-un-u duy-dum.
a swimmer drown-NMLZ-POSS-ACC hear-PST.3SG
‘I heard a swimmer drowned.’

b. bir yüzücü-nün boğul-duğ-u-nu duy-dum.
a swimmer-GEN drown-NMLZ-POSS-ACC hear-PST.3SG
‘I heard a swimmer drowned.’

c. *bir yüzücü çırp-ın-dığ-ı-nı duy-dum.
a swimmer whisk-IN-NMLZ-POSS-ACC hear-PST.3SG
‘I heard a swimmer floundering.’

d. bir yüzücü-nün çırp-ın-dığ-ı-nı duy-dum.
a swimmer-GEN whisk-IN-NMLZ-POSS-ACC hear-PST.3SG
‘I heard a swimmer floundering.’

When talking about reflexivization more generally, we observe that all verbal reflexives in Turkish are encoding agency semantics, regardless of the morpheme they are derived with. Consequently, both types can successfully pass the diagnostics used for agency semantics. However, only IN derived reflexives seem to have an external argument in their syntactic

structure. Reflexive verbs derived with the non-active Voice morpheme IL, by contrast, lack such a projection.

One way to demonstrate this difference is through observing their behavior under causativization, as verbs with a non-active Voice layer cannot be further causativized. This restriction is consistent with the fact that the non-active Voice morpheme, also found with passives, impersonal passives and anticausatives, also projecting a non-active Voice layer, blocks additional causativization in the structure. The IL derived reflexives are no exception to this pattern (Key, 2025). As predicted, they cannot undergo causativization because the external argument is not present in the argument structure and thus cannot further function as a causee (47)–(48).

- (47) a. Pırnç toprak-tan (Zeynep tarafından) ayr-ıl-dı.
 Rice ground-ABL Zeynep by separate-II-PST.3SG
 ‘The rice was separated (by Zeynep) from the ground.’
 [passive]

- b. *Asemin pırnç-i toprak-tan (Zeynep tarafından) ayr-ıl-dır-dı.
 Asemin rice-ACC ground-from Zeynep by separate-II-CAUS-PST.3SG
 ‘Asemin made the rice be separated from the ground by Zeynep.’

[causativized passive]

- (48) a. Turistler otel-den ayr-ıl-dı.
 Tourists hotel-ABL separate-II-PST.3SG
 ‘The tourists left (separated themselves) from the hotel.’
 [reflexive]

- b. *Otel sahibi turistler-i otel-den ayr-ıl-dır-dı.
 Hotel owner tourists-ACC hotel-ABL separate-II-CAUS-PST.3SG
 Intended: ‘The hotel owner made the tourists leave the hotel’

[causativized reflexive]

In contrast, we observed the IN derived figure and ground reflexives lack an overt or a specific internal argument in their structure. These verbs may appear de-transitivized on the surface. However, this does not indicate the presence of a non-active Voice head or a demoted external argument. In fact, IN verbs retain their external argument and associated projections and therefore can undergo causativization without restriction. Their apparent intransitivity results from *v*-head layers within the verbal domain, which prevents the introduction of a *v*P-internal argument. Accordingly, it is entirely expected that these verbs surface as intransitive while remaining causativizable, as illustrated in (49).

- (49) Satıcı müşteri-yi alışveriş-ten çek-in-dir-di.
 Seller customer-ACC shopping-ABL pull-IN-CAUS-PST.3SG
 ‘The seller made the customer hesitate (hold oneself back) from shopping’

In the next section, we turn to non-reflexive IN derived verbs in Turkish and argue that they pattern with unergatives, as predicted under our analysis.

5.2 The IN unergatives

The verbs in the examples (50) and (51) pointed to another type of IN derived verb in which reflexive meaning is not conveyed. These structures, just like the IN figure and ground reflexives, have a *v*-head that does not project a DP argument. However, we hypothesized that the Ref-head has a distinct meaning in this case.

- (50) Öğrenci bahçe-yi gez-di.
 Student garden-ACC wander-PST.3SG
 ‘The student visited the garden.’
- (51) Öğrenci (*bahçe-yi) gez-in-di.
 Student garden-ACC wander-IN-PST.3SG
 ‘The student wandered in the garden.’

In sentence (50), the verb *gez-* ‘to visit’ expresses an event performed on a specific ground. However, its derived intransitive counterpart *gezin-* ‘to wander’ is ungrammatical with the accusative case marked argument, that is the direct object. This indicates that the ground on which the action is performed cannot be specified within the argument structure. Like other figure and ground reflexives derived with IN, these verbs have intransitive syntactic structures and exhibit agentive semantics when tested.

- (52) *Dün bu park-ta bir öğrenci gez-in-diğ-i-ni
 Yesterday this park-LOC a student visit-IN-NMLZ-POSS-ACC
 duy-dum.
 hear-PST.3SG
 ‘I heard that a student was wandering at this park yesterday.’
- (53) Öğretmen öğrenciler-i bahçe-de gez-in-dir-di.
 Teacher students-ACC garden-LOC visit-IN-CAUS-PST.3SG
 ‘The teacher made the students wander in the garden.’

As mentioned earlier, the allosemantic variant of the Ref-head does not introduce an AGENT role or an additional argument into the verbal structures; it simply existentially binds the internal argument. Thus, it is convenient to argue for a VoiceP layer within this structure, unlike reflexives. The results from the embedded nominalization test support this prediction (52): like IN reflexives, IN derived unergatives require their subjects to bear the Genitive

Case in the embedded environments. Furthermore, these verbs can be causativized as shown in (53), just like other figure and ground reflexives and unlike the IL derived verbs. This finding reinforces the claim that IN verbs have agentive semantics and patterns both syntactically and semantically with unergatives.

5.3 The base root transitivity

Our current analysis on IN morpheme makes another key prediction. Regarding the IN derived figure reflexives, all the verb roots discussed so far have been conventionally considered as transitive. However, strictly speaking, IN morpheme should not function as a de-transitivizer, as we have previously stated. Within the theoretical framework adopted in this study, we assume that verb roots are not lexically encoded as requiring a specific argument structure. Therefore, as shown in Table 4 below, we predict that it is possible to derive IN figure reflexives from verb roots that would normally be considered intransitives, as predicted under our current analysis.

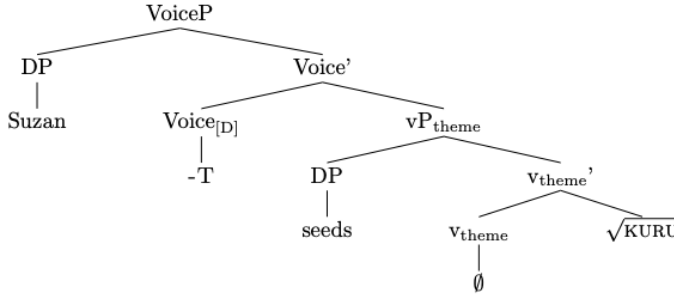
Table 4. A list of figure reflexive verbs derived from unaccusative verb stems

Verb Stem		Figure Reflexive	
<i>kuru-</i>	‘to dry’	<i>kuru-n</i>	‘to dry oneself’
<i>şiş-</i>	‘to swell out’	<i>şiş-in</i>	‘to boast oneself’ (lit :to make oneself swell out)
<i>uza-</i>	‘to extend’	<i>uza-n</i>	‘to reach forth’
<i>sığ-</i>	‘to fit in’	<i>sığ-in</i>	‘to take shelter in’ (to make oneself fit in)
<i>çöz-</i>	‘to solve’	<i>çöz-ün</i>	‘to dissolve’
<i>boz-</i>	‘to ruin’	<i>boz-un</i>	‘to decay’
<i>geç-</i>	‘to pass by’	<i>geç-in</i>	‘to make living’ (lit: to make oneself pass from one day to the other)
<i>kalk-</i>	‘to get up’	<i>kalk-in</i>	‘to develop, to progress’ (lit: to get oneself up)

To discuss some concrete examples, the verb roots *kuru-* ‘to dry’, *şiş-* ‘to swell out’ and *uza-* ‘to extend’ are unaccusatives in Turkish, all denoting a change-of-state and having their sole argument bearing the THEME role. Therefore, in transitive versions of these verbs, such as example (54), the causative morpheme introduced by the Voice head can be observed in the structure (55).

- (54) Suzan çekirdekleri *kuru-t*-tu.
 Suzan seeds.ACC dry-CAUS-PST.3SG
 ‘Suzan dried the seeds.’

(55)

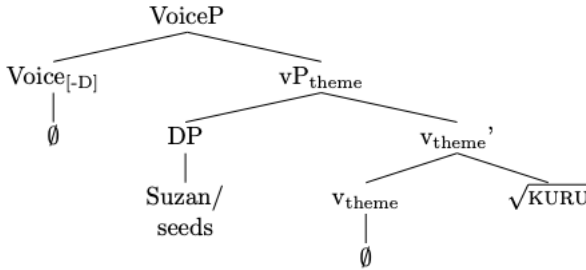


The connection between animacy and agency has been previously examined in the literature (Perlmutter, 1978; Burzio, 1981), and several studies have observed that unergative verbal constructions in Turkish are observant to the animacy of the external argument (Sezer 1979; Özsoy, 2010), with certain exceptions. At this point, we do not observe a constraint on animacy in structures containing unaccusative verbs, as both *Suzan* and the *seeds* can be the arguments of the event drying in (56)-(57).

(56) a. Suzan güneş-te kuru-du.
Suzan sun-DAT dry-PST.3SG
'Suzan dried under the sun.'

b. Çekirdekler güneş-te kuru-du.
Seeds sun-DAT dry-In-PST.3SG
'The seeds dried under the sun.'

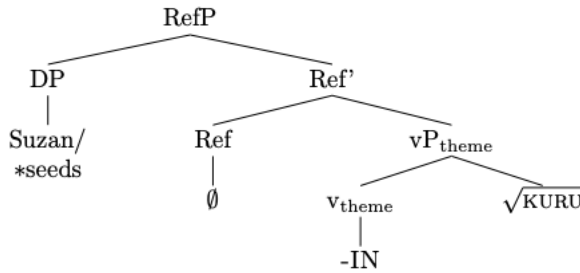
(57)



But when we observe the -IN derived versions of these verbs, they result in their sole argument bearing the AGENT and THEME roles at the same time, hence expressing figure reflexivity. As a result of this derivation, we see that verbal structures no longer accept inanimate entities as arguments (58)-(59). This result indicates our analysis for reflexive verbs being on the right track. Here, *Suzan* can be the argument of *kurun-* 'drying oneself' while the

seeds cannot, since the Ref-head associates the sole argument with the AGENT role, as well.

- (58) a. Suzan güneş-te kuru-n-du.
 Suzan sun-DAT dry-In-PST.3SG
 ‘Suzan dried herself under the sun.’
 b. *Çekirdekler güneş-te kuru-n-du.
 Seeds sun-DAT dry-PST.3S
 ‘The seeds dried themselves under the sun.’
 (59)



The discussion here shows that there is no requirement for the base root to be transitive to build a figure reflexivity. As a matter of fact, this type of requirement would have been incompatible with the framework adopted in this study. Most importantly, IN morpheme which indirectly signals the presence of the Ref head in the structure directly precludes the presence of the regular Voice in the structure, as confirmed by the lack of the causative suffix {-t} in the figure reflexive form.

6. Conclusion

There are two ways to express reflexive meaning in Turkish: pronominal reflexives and verbal reflexivity. Verbal reflexives have been observed to be derived from two different morphemes such as IL and IN. Previous studies have observed that IL reflexives, which produce figure reflexives, tend to view them as unified with other verb types derived from non-active Voice. On the other hand, the analyses of IN morphemes have seen them limited to their ability to derive ground reflexivity and considered this morpheme to be realization of an Applicative-head. However, these analyses were insufficient to analyze other verbal structures where the IN morpheme can appear.

The contribution of this study is to provide counter-example data indicating that IN derived verbs can be IN figure reflexives or non-reflexive unergatives. We also proposed an analysis that can derive all observed IN derived verbs. Our proposal has two main assumptions. First, we propose that the IN morpheme, rather than being a reflexivizer or a de-transitivizer head, is the realization of a spec-less *v*-head. Second, we have shown that such spec-

less *v*-heads are licensed in the context of a Ref-head that we claim to derive agent-ground and agent-figure reflexivity as well as non-reflexive existential saturation (limited to a certain class of roots).

Finally, all these things being said, our analysis has not presented any proposal regarding IL derived verbs, more specifically IL derived reflexives. However, we observe that there are almost no reflexives that are derived with both morphemes, with the exception of a few minimal pairs. This result strongly indicates that the two morphemes are almost in a complementary distribution and that their productivity is limited in a way that suggests root-sensitive distribution. At this point, it seems reasonable that the system derives reflexive meanings via the IL morpheme in contexts where a reflexive derivation via the IN morpheme is not available. We believe that a comprehensive analysis along these lines can be developed for the derivation of Turkish verbal reflexives in a way that is fully consistent with our current proposal for IN derived verbs. However, fleshing out this idea is left for a future study.

References

- Acartürk, C., & Zeyrek, D. (2010). Unaccusative/unergative distinction in Turkish: A connectionist approach. In *Proceedings of the Eighth Workshop on Asian Language Resources* (pp. 111–119).
- Alexiadou, A., & Schäfer, F. (2013). Non-canonical passives. In A. Alexiadou & F. Schäfer (Eds.), *Non-canonical passives* (pp. 1–20). John Benjamins Publishing Company.
- Akkuş, F., & Paparounas, L. (2024). *The argument structure of Turkish verbal reflexives* (Unpublished manuscript).
- Arslan, A. (2022). Türkiye Türkçesinde dönüşlü çatının oluşum yöntemleri. *Uluslararası Beşeri Bilimler ve Eğitim Dergisi*, 8(17), 226–240.
- Baturay-Meral, S., & Meral, H. M. (2018). On single argument verbs in Turkish. *Bilig*, (86), 115–136.
- Borer, H. (1994). *The projection of arguments*. University of Massachusetts Occasional Papers in Linguistics (Vol. 20).
- Burzio, L. (1981). *Intransitive verbs and Italian auxiliaries* (Doctoral dissertation). Massachusetts Institute of Technology.
- Davidson, D. (1967). Truth and meaning. In J. Hintikka & P. Suppes (Eds.), *Philosophy, language, and artificial intelligence: Resources for processing natural language* (pp. 93–111). Springer.
- Dikmen, F., Demirok, Ö., & Öztürk, B. (2022). How can a language have double passives but lack antipassives? *Glossa: A Journal of General Linguistics*, 7(1).
- Embick, D. (2004). Unaccusative syntax and verbal alternations. In A. Alexiadou, E. Anagnostopoulou, & M. Everaert (Eds.), *The unaccusativity puzzle: Explorations of the syntax–lexicon interface* (pp. 137–158). Oxford University Press.

- Göksel, A. (1993). *Levels of representation and argument structure in Turkish*. University of London, School of Oriental and African Studies (United Kingdom)
- Gülsevin, S. (2022). Üst üste kullanılan *-In-Il-* ekleri üzerine. *Turkish Studies – Language*, 17(3), 857–867.
- Gündoğdu, S. (2017). *(I)l / (I)n morphology in Turkish: Implications for u-syncretism*. In *Proceedings of the 4th Patras International Conference for Graduate Students in Linguistics (PICGL-4)* (pp. 85–105).
- Heim, I., & Kratzer, A. (1998). *Semantics in generative grammar*. Blackwell.
- Key, G. (2021). Figure and ground reflexives in Turkish. In *Proceedings of the Workshop on Turkic and Languages in Contact with Turkic* (Vol. 6, pp. 50–64).
- Key, G. (2025). Voice in Turkish: Re-thinking u-syncretism. *Natural Language & Linguistic Theory*, 43(1), 331–384.
- Kornfilt, J. (2013). *Turkish*. Routledge.
- Kornfilt, J., & von Stechow, K. (2009). Specificity and partitivity in some Altaic languages. *MIT Working Papers in Linguistics*, 58, 19–40.
- Koşaner, Ö. (2005). *Türkçede dönüşlü yapıların biçim-sözdizimsel özellikleri* (Master's thesis). Dokuz Eylül Üniversitesi.
- Lewis, G. (2000). *Turkish grammar*. Oxford University Press.
- Nakipoğlu, M. H. (1998). *Split intransitivity and the syntax–semantics interface in Turkish* (Doctoral dissertation). University of Minnesota.
- Özkaragöz, İ. (2011). Monoclausal double passives in Turkish. In S. Ay, Ö. Aydın, İ. Ergenç, S. Gökmen, S. İşsever, & D. Peçenek (Eds.), *Studies in Turkish linguistics* (pp. 77–92). John Benjamins Publishing Company.
- Özsoy, A. S. (2010). Argument structure, animacy, syntax and semantics of passivization in Turkish: A corpus-based approach. In A. Okrent & J. Boyd (Eds.), *Corpus analysis and variation in linguistics* (pp. 259–279). John Benjamins Publishing Company.
- Perlmutter, D. M. (1978). Impersonal passives and the unaccusative hypothesis. In *Proceedings of the Annual Meeting of the Berkeley Linguistics Society* (pp. 157–190).
- Sezer, E. (1979). Eylemlerin Çoğul Özne Uyumunda Anlamsal Özelliklerin Rolü. *Dilbilim Seçmeleri*. Ankara Üniversitesi Yayınları.
- Sezer, F. E. (1991). *Issues in Turkish syntax* (Doctoral dissertation). Harvard University.
- Spathas, G., Alexiadou, A., & Schäfer, F. (2015). Middle voice and reflexive interpretations: Afto-prefixation in Greek. *Natural Language & Linguistic Theory*, 33, 1293–1350.
- Talmy, L. (1978). Figure and ground in complex sentences. In J. H. Greenberg (Ed.), *Universals of human language* (Vol. 4, pp. 625–649). Stanford University Press.
- Taneri, M. (1993). *The morpheme -Il/(I)n: The syntax of personal passives, impersonal passives and middles in Turkish* (Doctoral dissertation). University of Kansas.

- Wood, J. (2023). *Icelandic nominalizations and allosemy* (Vol. 84). Oxford University Press.
- Zeyrek, D. (2006). Anticausatives in Turkish: The role of the suffix *(I)l/(I)n*. In É. Á. Csató, B. Karakoç, & A. Menz (Eds.), *The Uppsala meeting: 13th International Conference on Turkish Linguistics*. Harrassowitz Verlag.

DOES CHATGPT ARGUE THROUGH STRUCTURE OR STANCE? A METADISOURSE ANALYSIS OF ENGLISH ARGUMENTATIVE ESSAYS ACROSS TOPICS

Ruhan GÜÇLÜ

Asst. Prof. Dr., Gaziantep University, Department of English Language and Literature,
gucluruhan@gmail.com, ORCID: 0000-0002-2748-8363

Güçlü, R. (2025). Does ChatGPT argue through structure or stance? A metadiscourse analysis of English argumentative essays across topics. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 147–169). Artsürem. <https://doi.org/10.7816/imuling-15-25-01X007>

ABSTRACT

This study examines the use of interactive and interactional metadiscourse markers in English argumentative essays generated by ChatGPT across a range of social science topics. Drawing on Hyland's (2005) Interpersonal Model of Metadiscourse, a corpus of seventy-five argumentative essays, comprising fifteen topics with five essays each, was analysed using both qualitative and quantitative methods. The log-likelihood analysis demonstrates a clear preference for interactional markers, particularly hedges, while transitions emerge as the most frequent interactive devices in the overall corpus. These patterns indicate that ChatGPT-generated essays foreground stance-taking and reader engagement rather than explicit structural guidance. At the topic level, interactional markers are especially prominent in essays on sports, gender studies, psychology, and economics, whereas topics such as food and nutrition, and history rely more heavily on interactive resources. Overall, these findings expand the understanding of AI discourse by highlighting both the strengths and limitations of AI-generated writing.

Keywords: Interactive markers, Interactional markers, Metadiscourse, Topic, ChatGPT-produced argumentative essays

ÖZ

Bu çalışma, ChatGPT tarafından üretilen İngilizce tartışmacı denemelerde bilgi odaklı etkileşimli ve alıcı odaklı etkileşimli üstsöylem belirleyicilerinin, farklı sosyal bilimler konuları bağlamındaki kullanımını incelemektedir. Hyland'ın (2005) Kişilerarası Üstsöylem Modeli temel alınarak, her biri beş denemeden oluşan on beş farklı konuya ait toplam yetmiş beş tartışmacı denemeden oluşan bir derlem, nitel ve nicel yöntemler kullanılarak çözümlenmiştir. Log-olabilirlik çözümlemesi, özellikle kaçınılmaların öne çıktığı alıcı odaklı etkileşimli üstsöylem belirleyicilerine yönelik belirgin bir tercihi

ortaya koyarken, bağlayıcıların ise derlemin genelinde en sık kullanılan bilgi odaklı etkileşimli araçlar olduğunu göstermektedir. Bu örüntüler, ChatGPT tarafından üretilen denemelerin açık yapısal yönlendirmeden ziyade duruş kurma ve okurla etkileşimi ön plana çıkardığını göstermektedir. Konu düzeyinde bakıldığında, spor, toplumsal cinsiyet çalışmaları, psikoloji ve ekonomi alanlarındaki denemelerde alıcı odaklı etkileşimli belirleyicilerin özellikle baskın olduğu; buna karşılık beslenme ve gıda ile tarih gibi konularda bilgi odaklı etkileşimli kaynakların daha yoğun biçimde kullanıldığı görülmektedir. Genel olarak bu bulgular, yapay zekâ tarafından üretilen yazının hem güçlü hem de sınırlı yönlerini ortaya koyarak yapay zekâ söylemine ilişkin anlayışımızı genişletmektedir.

Anahtar sözcükler: Bilgi odaklı etkileşimli belirleyiciler, Alıcı odaklı etkileşimsel belirleyiciler, Üstsöylem, Konu, ChatGPT tarafından üretilen tartışmacı denemeler

1. Introduction

Artificial Intelligence (AI) has, in a remarkably short span, become part of education and academic writing, offering students, teachers, and researchers tools that deliver structure and clarity with remarkable ease. As an everyday reality in classrooms and research settings, students frequently turn to these tools for guidance, teachers use them to support lesson design, and researchers incorporate them into their academic routines as part of their experiments. This sudden shift has brought both excitement, due to the speed and accessibility AI provides, and hesitation, as questions of originality, authorship, and academic value remain very much open.

Among these tools, ChatGPT, created by OpenAI, has received particular attention. Part of its appeal lies in its versatility: it can generate essays, summaries, and arguments on a broad range of topics. Besides being attractive, it also raises questions about the quality of the academic prose it produces: What kind of writing does ChatGPT actually produce? And how closely do its texts resemble the rhetorical and stylistic qualities that characterize academic prose? One effective way to address such questions is to examine metadiscourse markers.

Metadiscourse markers (MDMs) are the rhetorical devices writers use to organize their arguments and to establish a relationship with their readers in their academic writing. According to Hyland (2005), interactive metadiscourse markers such as transitions, frame markers, endophoric markers, evidentials, code-glosses help the writers organize the text and guide the reader through the text by signaling structure and connections while interactional metadiscourse markers such as hedges, boosters, attitude markers, self-mentions, engagement markers allow writers to show their stance and involve the reader more directly in the discussion. The effective deployment of both categories is central to argumentative writing, where persuasiveness depends not only on logical reasoning but also on how writers signal credibility, involvement, and consideration of the reader's perspective. In addition, a large body of research confirms that

proficient writers rely on a balance of both interactive and interactional categories to produce texts that are coherent, dialogic, and persuasive (Ädel, 2006; Hyland, 2019; Hyland & Tse, 2004).

The growing presence of AI-generated writing has increased interest in how these systems employ metadiscourse. Scholars have begun to compare AI-generated texts with those written by humans, and the results so far reveal both overlaps and distinct gaps. Amirjalili et al. (2024), for example, noted that ChatGPT often falls back on formulaic coherence devices, while Shalevska (2024) emphasized its limited ability to capture more subtle rhetorical strategies. Similarly, Alshalan and Alyousef (2024) reported that AI-generated essays often lack a genuine authorial voice, particularly in their use of self-mentions and boosters. Dynel (2023) showed that ChatGPT prompts metadiscursive awareness in playful human–AI interactions, but fails to replicate authentic rhetorical presence. Kobzová (2023) further noted that ChatGPT’s hedging repertoire is narrower and less varied than in published academic texts. Similarly, Yao and Liu (2025) observed that while GPT-generated book reviews contained more interactional markers than human-authored reviews, the overuse of attitude markers and the lack of hedges led to an unbalanced stance.

Other studies focused on AI writing in various academic genres. Zhang and Zhang (2025a, 2025b) examined research article abstracts, demonstrating that ChatGPT tends to overuse reflexive and text-oriented markers but lacks the disciplinary nuance characteristic of human writing. Hongyan and Saeed (2025) also found that, although AI-generated abstracts heavily utilized stance and engagement markers, they were formulaic in comparison to the subtle variation found in human-authored texts. These concerns align with Flowerdew and Wang’s (2023) argument that academic writing in the age of AI faces both opportunities and challenges. In addition to these studies, recent research has also compared student-written essays with those produced by ChatGPT. Jiang and Hyland (2025a, 2025b, 2025c, 2025d) demonstrated that although AI-produced texts are generally coherent and frequently employ interactive markers and lexical bundles, they fall short in utilizing engagement markers, hedging devices, and evaluative expressions. These elements contribute to the persuasiveness of human writing. In a similar vein, Shalevska (2024) found that while ChatGPT regularly used cautious expressions like *may*, it rarely employed emphatic terms such as *clearly* or *definitely*. Yao and Liu’s (2025) book review analysis and Amirjalili, Neysani, and Nikbakht’s (2024) investigation of authorship and voice further confirmed that AI-generated texts, while contextually relevant, often lack depth, specificity, and authentic authorial presence. These studies consistently show that ChatGPT handles coherence well by relying on interactive markers; however, it has difficulty with stance-taking strategies, such as hedges, boosters, self-mentions, and engagement, which are essential for making academic writing more persuasive.

More recently, however, attention has turned to how metadiscourse patterns differ across disciplines. It has also been shown that topic-shaping metadiscourse practices are present in argumentative writing. For example, EFL students demonstrate distinct preferences, often employing fewer hedges but more reader pronouns than native writers (Yoon, 2021). This topic-based perspective reveals that it is not enough to ask whether AI can produce metadiscourse; we must also consider how its rhetorical choices shift in relation to the argument's topic. Nevertheless, existing studies on AI-generated writing and metadiscourse have mainly focused on broad comparisons between human- and machine-produced texts, with limited attention to how these resources vary across academic topics. What remains unexplored, and what this study seeks to address, is whether ChatGPT's argumentative essays shift their metadiscourse practices across topics, and whether stance or structure emerges as their dominant rhetorical strategy, in addition to focusing exclusively on the distribution and use of interactive and interactional metadiscourse markers in AI texts. Accordingly, this research is guided by the following research questions:

1. What are the metadiscourse categories and their frequencies used in ChatGPT-generated argumentative essays?
2. Is there any significant difference between the overall use of interactive metadiscourse markers and interactional metadiscourse markers in ChatGPT-generated argumentative essays?
3. Does the use of interactive and interactional metadiscourse markers vary depending on the topic (e.g., education, technology, health, environment, politics and law, culture, economics, gender studies, psychology, sports, religion, history, arts and aesthetics, food and nutrition, travel and tourism)?

By combining corpus-level analysis with topic-based comparison, this study aims to reveal both the general rhetorical orientation of AI writing and its adaptation to topic-specific contexts. This knowledge is not only theoretically valuable but also practically useful, as it can help educators and students determine the best approach to working with AI-generated texts in settings where developing argumentative skills is crucial.

The paper is structured as follows. Section 2 presents the methodology, including corpus design, data collection, and data analysis procedures with their analytical framework. Section 3 reports and discusses the results of the quantitative and qualitative analyses, pointing to the implications of these findings in relation to existing research. Finally, Section 4 presents the conclusion, providing a summary and outlining directions for future research.

2. Methododology

2.1 Corpus and data collection procedure

The dataset for this study was assembled through a systematic two-step process, with the goal of creating a balanced corpus of argumentative essays that captures a wide range of globally relevant topics while maintaining both topical diversity and textual variation.

In the first step, the following prompt was submitted: “*Identify the most debated and up-to-date fifteen topics in the world.*” This request generated fifteen topics that are widely discussed in contemporary global discourse, such as education, technology, health, politics and law, culture, economics, gender studies, psychology, sports, religion, history, arts and aesthetics, food and nutrition, travel and tourism, and environment.

In the second stage, each topic was used as the basis for essay generation with the instruction: “*Write five argumentative essays on ... issue, focusing on its most current and debated aspects.*” A total of seventy-five essays were collected, with 15 topics, each comprising five essays in their original form without modification.

Each prompt was submitted independently to capture lexical and rhetorical variation within topics, and the outputs were collected in their original form. The essays were systematically labelled (e.g., Education_1, Education_2, ..., Education_5) to enable both quantitative and qualitative analyses. A complete list of the generated essays, including their titles and associated topics, is provided in the Appendix.

The corpus compiled through this collection amounts to 33,801 words. Table 1 provides the word count distribution across the fifteen topics:

Table 1. Distribution of word counts across topics

Topic	Word Count	Topic	Word Count	Topic	Word Count
Education	2,988	Economics	2,790	Arts and Aesthetics	2,025
Technology	2,992	Gender Studies	2,044	Food and Nutrition	2,025
Health	2,369	Psychology	2,000	Travel and Tourism	1,931
Politics and Law	2,239	Sports	2,125	Environment	2,296
Culture	2,096	Religion	1,798	Arts and Aesthetics	2,025

This two-step data collection strategy ensured that the topics were drawn from prompts focusing on widely debated and current global issues, rather than being selected in advance by the researcher. In addition, generating multiple essays for each topic introduced internal variation, which helped strengthen the reliability of the dataset.

2.2 Analytical framework and data analysis procedure

The analysis followed Hyland’s (2005) taxonomy of metadiscourse, dividing markers into interactive (e.g., transitions, frame markers, evidentials, code glosses, endophoric markers) and interactional categories (e.g., hedges, boosters, attitude markers, engagement markers, self-mentions). Table 2 illustrates functions and examples of these categories.

Table 2. Hyland’s (2005) interpersonal model of metadiscourse

Category	Function	Examples
Interactive	Help to guide the reader through the text	Resources
Transitions	Express semantic relation between main clauses	And, in addition, but, consequently
Frame markers	Refer to discourse acts, sequences, or text stages	Finally, to conclude, my purpose is
Endophoric markers	Refer to information in other parts of the text	Noted above, see Fig.,in Section 2
Evidentials	Refer to source of information from other texts	According to X, (Y, 1990), Z states
Code-glosses	Help readers grasp the meanings of ideational material	Namely, e.g., such as, in other words
Interactional	Involve the reader in the text	Resources
Hedges	Withhold the writer’s full commitment to the proposition	Might, perhaps, possible, about
Boosters	Emphasize force or writer’s certainty in proposition in fact / definitely / it is clear that	In fact, definitely, it is clear that
Attitude Markers	Express writer’s attitude to proposition	Unfortunately, I agree, surprisingly
Engagement Markers	Explicitly refer to or build a relationship with the reader	Consider, note that, you can see that
Self-Mentions	Explicit reference to author(s)	I, we, my, our

Interactive resources help to organize discourse in ways that guide readers through the text: Transitions (e.g., however, therefore) signal relations between clauses. Frame markers (e.g., in conclusion, first)

indicate stages in discourse organization. Endophoric markers (e.g., as noted above) direct readers to other parts of the text, whereas evidentials (e.g., according to X, research shows) attribute claims to external sources. Code glosses (e.g., namely, in other words) are used to provide clarification or elaboration.

Interactional resources, by contrast, reveal the writer's stance and regulate how the audience is addressed. Hedges (e.g., may, might) introduce tentativeness, whereas boosters (e.g., clearly, in fact) stress conviction. Boosters (e.g., clearly, in fact, it is obvious that) do the opposite, strengthening claims by signalling certainty and authority. Attitude markers (e.g., unfortunately, importantly) express how the writer evaluates or positions an idea or argument. Self-mentions (e.g., I argue, we suggest) foreground the writer's role, and engagement markers (e.g., you can see, note that) directly address the reader.

This framework was adopted because this model is specifically designed for academic writing (Zarei & Mansoori, 2011) and provides a comprehensive model that incorporates previous approaches while addressing their gaps and overlaps (Hyland, 2005). Given its extensive use in metadiscourse research (Hyland & Tse, 2004; Triki, 2019), it provides an effective tool for exploring topic-based variation.

The data analysis process involved three steps, combining quantitative and qualitative approaches to examine both the distribution and the contextual functions of metadiscourse markers.

Initially, the corpus of seventy-five essays was processed within a text analysis environment. Each text was scrutinized to locate metadiscourse markers, which were then classified and coded according to the analytical framework. Ambiguous cases were resolved by checking the operational definitions and examples provided in the literature. Following identification and coding, quantitative analyses were conducted. The raw frequencies of each metadiscourse category were calculated, and figures were normalized to 1,000 words to facilitate cross-topic comparison. The distribution of metadiscourse markers was compared across the fifteen topics, allowing for the identification of both overlapping tendencies and topic-specific distinctions. The analysis was then complemented by a qualitative stage, in which markers were examined within their textual environments, and illustrative excerpts were provided to demonstrate their rhetorical functions. Through this twofold procedure, the research was able to account for both the general distribution of markers and their rhetorical realizations in context.

3. Results and Discussion

This section presents the findings and discusses them in relation to the three research questions. It first outlines the overall distribution of metadiscourse categories in the dataset, then compares interactive and interactional markers across the corpus, and finally examines how these

markers vary across topics. Each stage of the analysis is supported by quantitative evidence and illustrated with tables and figures.

3.1 Distribution of metadiscourse categories in the corpus

The first research question aimed to determine the general distribution of metadiscourse categories across the entire corpus. As shown in Figure 1, all of the interactive categories, such as transitions, frame markers, endophoric markers, evidentials and code-glosses and interactional categories, such as hedges, boosters, attitude markers, self-mentions and engagement markers, were found in English argumentative essays produced by ChatGPT.

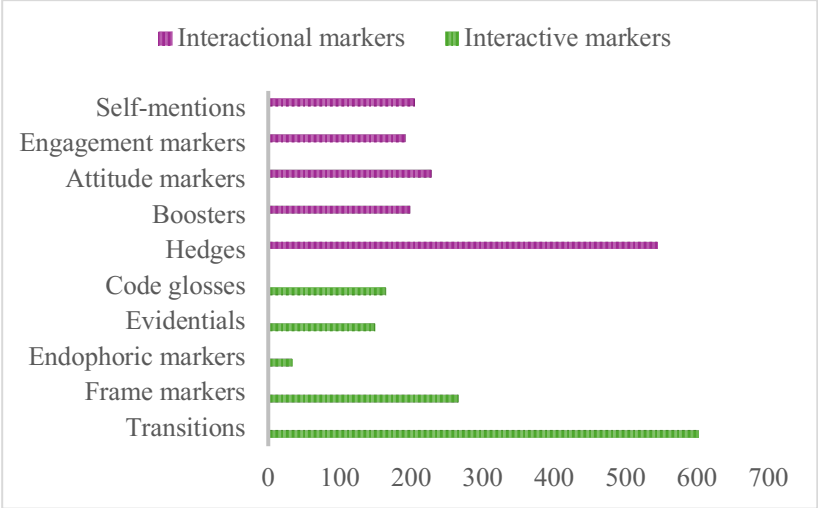


Figure 1. Distribution of metadiscourse categories in the ChatGPT’s argumentative essays

Figure 1 demonstrates that both interactive and interactional markers appear frequently in the analyzed ChatGPT’s argumentative essays, although their subcategories show marked variation in distribution. Among interactive devices, transitions are the most commonly used, followed by frame markers, code glosses, and evidentials. Endophoric markers are scarcely used in the corpus. Within the interactional group, hedges represent nearly half of all instances, whereas boosters, attitude markers, self-mentions, and engagement markers occur in more even proportions. In a similar vein, Abdi (2002) observed that non-native students make extensive use of hedges to soften claims, a tendency also evident in the present data. The rarity of endophoric references corresponds with Hyland and Tse’s (2004) findings, which suggest that explicit cross-referencing is more typical of professional academic prose than of student essays. Below, the frequency of use for each category will be discussed, along with examples extracted from the corpus.

The high frequency of transitions (1.78 per 100 words) indicates that cohesion and the logical sequencing of ideas played a central role in the essays. This is exemplified below:

(1) “**Moreover**, compulsory voting can promote political equality by ensuring that all citizens participate in elections.” Here, the metadiscourse marker “moreover” creates a cumulative effect that strengthens the flow of reasoning. (2) “**However**, globalization can also result in economic inequality despite its many benefits.” The transition marker “however” establishes contrast and balance. (3) “**Additionally**, the relevance of beauty pageants has been questioned in recent years due to shifting cultural norms.” As an additive transition marker, “additionally” signals the extension of argument. These instances confirm that the discourse is heavily reliant on transitions to achieve textual flow, a pattern consistent with Triki’s (2019) finding that student essays prioritise connectors to guide readers. In addition to the transitional devices given in (1), (2) and (3), the dataset also contains a range of transition markers such as: *and, but, while, furthermore, although, as, by, in addition, conversely, when, if, because, as a result, especially, instead, therefore, since, on the other hand, in contrast*.

Within the interactive categories, frame markers rank second after transitions, with a normalized frequency of 0.79 per 100 words. The essays employed these devices to signal sequencing and conclusions:

(4) “**In conclusion**, therapy can be highly effective in reducing stress among patients.” The marker “in conclusion” closes the argument explicitly. (5) “**First**, it is necessary to acknowledge the economic consequences of raising the minimum wage.” The marker “first” orients the reader to a staged argument. (6) “**Overall**, feminism continues to shape contemporary discussions of equality.” Here, “overall” generalizes the argument. Such uses support Hyland and Tse’s (2004) claim that frame markers help structure complex reasoning, though their relatively modest presence suggests that the essays often relied on implicit structuring rather than explicit staging. Other frame markers occurring frequently in the corpus are listed as follows: *in conclusion, one of the primary reasons, one of the strongest/primary concerns, first, second, furthermore, overall, the question of whether, the issue of, the debate surrounding, first and foremost, another issue/benefit, another criticism, ultimately arguing, in recent years*.

When compared to the transitions and frame markers, code glosses were moderately used (0.49 per 100 words), mainly to clarify or exemplify claims.

(7) “Globalization creates cultural homogenization, **for example**, the spread of fast-food chains across the globe.” (8) “Beauty pageants promote unhealthy ideals, **such as** unrealistic body expectations.” (9) “Voting reforms can ensure fairness, **that is**, they can prevent manipulation and fraud.” The markers “for example”, “such as”, “that is” show that clarification was strategically used to make arguments accessible. This confirms Abdi’s (2002) suggestion that code glosses help writers anticipate reader needs by offering additional explanation. Several different code-glosses appeared with notable

frequency, for example: *for example, for instance, such as, that is, namely, including, particularly, in other words.*

The fourth most frequently used interactive category is evidentials, with the frequency of 0.44 per 100 words, showing limited engagement with external sources in ChatGPT's argumentative essays. The few instances illustrate general references rather than disciplinary grounding:

(10) "**Studies have shown that** therapy reduces symptoms of depression." The marker "studies have shown that" invokes authority but in a broad, unsourced manner. (11) "**Evidence from** countries with compulsory voting indicates higher turnout rates." The marker "evidence from" suggests research, yet it is generic. (12) "**Research from** the Alaska Permanent Fund shows that basic income schemes can work in practice." The marker "research from... shows that" again signals authority but without detailed citation. Compared with research writing in hard sciences, where evidentials dominate (Hyland & Jiang, 2017), this underuse demonstrates that the essays lack intertextual depth and remain closer to opinion pieces. Further examples of evidentials that appeared frequently in the dataset are as: *studies have shown, research suggests, research indicates, evidence suggests, evidence from, history has shown, history shows, according to, critics argue, proponents contend, opponents may argue, many argue, often cite, historians risk.*

Endophoric markers were relatively rare in the dataset, appearing only in a few instances (0.10 per 100 words), where the essays referred to other parts of the text. For example, (13) "**As noted above**, the evidence strongly supports this claim" employs as noted above to direct the reader backward in the text. In (14) "**The following section** will further illustrate this point," the phrase the following section signals a forward reference, preparing the reader for upcoming discussion. Similarly, (15) "See the argument **mentioned earlier** in this essay" uses mentioned earlier to maintain internal cohesion. The overall low frequency of such devices indicates that students relied more on linear progression rather than explicit cross-referencing, a finding consistent with Hyland and Tse's (2004) observation that endophoric markers are more typical of advanced academic texts than of short student essays. Some of the endophoric markers detected in the dataset are as follows: *as mentioned above, as noted earlier, as outlined next, as will be demonstrated.*

Among the interactional metadiscourse categories, hedges were by far the most frequent (1.61 per 100 words), underlining a strong tendency toward caution. This indicates that claims are typically framed as tentative rather than categorical, leaving space for alternative viewpoints and reader interpretation. Such cautious positioning is characteristic of argumentative writing in soft disciplines, where persuasion often relies on moderation rather than assertive certainty. At the same time, the heavy reliance on hedging suggests that epistemic caution sometimes compensates for limited evidential support in the essays. The essays frequently included modal verbs and tentative expressions to soften their claims. For instance, (16) "**Mandatory voting may** enhance democratic participation but could also reduce individual freedom." The hedge "may" acknowledges possibility. (17) "**Feminism might** still be relevant

in today's society despite significant progress." The hedge "might" limits the force of the claim. (18) "*Globalization **could** create new opportunities, but it also threatens local cultures.*" The hedge "could" positions the statement as probable but not absolute. These examples align with Hyland's (1996) observation that hedging creates a space of negotiation, reflecting an awareness of alternative viewpoints. A variety of hedging devices were frequently used in the dataset, including: *may, might, can, could, often, sometimes, generally, typically, usually, frequently, largely, potentially, arguably, perhaps, tends to, appears to, seems to, is likely to, it is possible that, may not necessarily, in many cases, could arguably, often argued that.*

Attitude markers were also used frequently (0.68 per 100 words), showing that evaluation is an integral part of how the AI-generated essays frame their arguments. For example, (19) "*It is **important** to recognize the role of education in fostering equality.*" The marker "important" conveys evaluative emphasis. (20) "*The relevance of beauty pageants is **questionable** in modern societies.*" The marker "questionable" signals critical stance. (21) "*Mandatory voting remains a **vital** tool for increasing political participation.*" The marker "vital" emphasizes necessity. These choices give the essays a more evaluative and engaged tone, indicating that the AI-generated texts move beyond a purely detached or impersonal style, which echoes Hyland's (2005) observation that attitude markers function to position claims interpersonally. Hyland's (2005) claim that attitude markers shape interpersonal positioning. The analysis revealed frequent use of various attitude markers, including the following: *important, significant, essential, crucial, vital, meaningful, responsible, effective, preferable, desirable, unfortunate, alarming, problematic, promising, remarkable, inevitable, questionable, troubling, inclusive, beneficial, harmful, valuable, worthwhile, pressing, healthier.*

Self-mentions were moderately present (0.61 per 100 words), revealing some authorial visibility. Illustrations include: (22) "***I argue** that compulsory voting strengthens democracy.*" The phrase "I argue" highlights personal stance. (23) "***We suggest** that parental control should be balanced with children's autonomy.*" The phrase "we suggest" signals collective authorial presence. (24) "***In this essay, we contend** that globalization requires nuanced evaluation.*" This explicit self-mention frames the argument as a discursive claim rather than a purely impersonal statement. Their frequency suggests that the essays move between impersonal formulation and explicit authorial presence, producing a rhetorical stance in which self-mention is used selectively rather than extensively, a pattern widely observed in academic discourse (Hyland & Jiang, 2017). The corpus featured several self-mentions occurring regularly, such as: *I, we, my, our, us.*

Boosters, while less frequent (0.59 per 100 words), were nonetheless present, used to reinforce conviction. Illustrations include: (25) "*Research **clearly** demonstrates that education reduces inequality.*" The booster "clearly" heightens certainty. (26) "*Movements such as #MeToo have **undoubtedly** demonstrated the persistence of gender inequality.*" The booster "undoubtedly" asserts a strong claim. (27) "*Universal basic income is*

certainly one of the most debated economic policies of our time.” The booster “certainly” projects firm stance. However, the lower frequency of these markers in comparison with hedges points to an orientation toward caution rather than strong assertion, reflecting rhetorical practices typically observed in the soft disciplines (Hyland & Tse, 2004). A diverse set of boosters was observed in the corpus, as in the following: *clearly, in fact, indeed, undoubtedly, of course, always, definitely, certainly, must (emphatic use), will (predictive/strong certainty), it is essential, it is crucial, it is obvious that, there is no doubt that, surely, truly, without doubt, research has shown, overwhelmingly.*

Engagement markers occurred at moderate levels (0.57 per 100 words), inviting the reader into the discussion. For example, (28), “*As we can see, patients respond positively to therapy.*” The phrase, “as we can see”, works to create a sense of shared observation. In (29), “*Policymakers must recognize the implications of globalization,*” the modal “must” is a direct call to the reader’s sense of duty. Example (30), “*Let us consider the potential drawbacks of universal basic income,*” uses “let us consider” to bring readers into the reasoning process. Such uses illustrate Hyland’s (2001) observation that engagement markers foster a dialogic relationship by shaping how readers respond to the text. The analysis further identified the following engagement markers as commonly used: *you, we, our, us, should, must, let us consider, note that, as we will see, remember that, it is worth noting, one might argue, governments should, citizens must, we need to, we contend, we propose, let’s, consider, it is essential to recognize, need to, opponents may argue.*

Overall, the findings suggest that argumentative essays rest on two key features: cohesion and caution. Cohesion comes mainly from transition markers, while caution is expressed through hedging. The two-work hand in hand to create a persuasive style of writing, one that opens space for negotiation rather than pressing claims with absolute certainty. There is also a weakness worth noting. Evidentials are seldom used, and endophoric connections are almost absent. They reduce the essays’ connection to broader scholarship and weaken their authority. The writing therefore risks sounding like academic prose in form but not in substance, a pattern reminiscent of disciplines that lean more on stance than on evidence. As Hyland (2005) and Abdi (2002) remind us, this reflects the traditions of fields in which knowledge is treated not as fixed truth but as something open to debate, contestation, and reinterpretation.

When compared with earlier research on AI-generated texts, the present findings reveal both points of overlap and areas of difference. One striking similarity is the dominance of hedges, a pattern also noted in previous studies (Liu & Deng, 2023; Yuan, 2024). This pattern is also genre-sensitive: in academic book reviews, ChatGPT has been shown to overuse interactional resources overall, especially attitude markers, while underusing hedges and self-mentions compared with human reviewers (Yao & Liu, 2025). The frequent use of expressions such as *may, might, could, and seems* suggests that ChatGPT prefers to frame its claims cautiously rather than with certainty. In

the same way, the relative absence of evidentials confirms what others have observed; rather than anchoring arguments in external sources, the texts often rely on general reasoning. Accordingly, these tendencies place AI writing closer to the rhetorical style of soft disciplines, where persuasion rests on positioning and interpretation rather than on evidential authority (Hyland, 2005).

At the same time, the analysis also uncovers some differences. While several scholars have argued that ChatGPT frequently produces categorical or overconfident statements (Gao, 2023), the essays in this corpus suggest otherwise. In this context, hedging is far more common than boosting, creating an overall tone of caution rather than certainty. Similarly, although AI writing is sometimes portrayed as impersonal or detached (Li, 2025), the presence of engagement markers such as *we*, *you*, and *let us consider* indicates that the model does make an effort to involve the reader.

Overall, the results suggest that AI-generated essays cannot be neatly described in one way. They support earlier claims about heavy reliance on hedging, limited use of evidentials, and the preference for interactional markers, but they also complicate this view. The essays show a mixture of caution and reader involvement, shaped not only by the model’s probabilistic nature but also by the conventions of argumentative writing. According to Hyland (2005), strong academic argumentation requires writers to combine interactive and interactional resources to both guide readers and establish authority. However, many studies show that L2 students struggle with this balance. According to Abdi (2002), Iranian EFL learners tended to rely heavily on interactional resources, especially attitude markers, while making limited use of evidentials and transitions. Likewise, Hyland and Tse (2004) pointed out that undergraduate essays often showed an “evaluative excess,” with stance dominating over textual organization. The findings here suggest that ChatGPT mirrors these same patterns and, in certain instances, pushes them even further.

3.2 Comparison of interactive and interactional categories

This section addresses the second research question and presents the findings concerning the overall use of interactive and interactional metadiscourse markers in ChatGPT-generated argumentative essays. The normalized figures reveal a clear imbalance between the two categories: interactional markers occur at an average rate of approximately 61 per 1,000 words, whereas interactive markers reach only 38 per 1,000 words. A log-likelihood test confirms that this difference is statistically significant (LL = 178.64, $p < .0001$), as reported in Table 3.

Table 3. Log-likelihood results of the distribution of interactive and interactional markers

	Interactive		Interactional		LL Ratio
	(f)	(%)	(f)	(%)	
Metadiscourse Categories	1969	5.83	1218	3.60	+178.64****

Table 3 demonstrates that ChatGPT-produced argumentative essays show a marked preference for interactional resources associated with stance-taking and reader engagement, rather than for interactive resources that primarily serve organizational functions. This pattern is consistent with earlier findings on student argumentative writing. Hyland and Tse (2004) observed that student writers tend to privilege interactional resources in order to enhance persuasive force, while organizational guidance remains comparatively underdeveloped. Similarly, Hyland (1998) demonstrates that academic argumentative writing extensively uses boosters as stance-related resources to negotiate certainty and strengthen claims, particularly in contexts where persuasion is foregrounded. In this respect, the AI-generated corpus closely mirrors patterns documented in human student writing. However, the imbalance observed here appears more pronounced than that typically reported for expert academic prose. Dahl (2004) noted that in professional writing, interactional and interactive resources are more evenly distributed, with only a slight predominance of interactional markers. Overall, these findings suggest that ChatGPT-generated essays resemble student writing more closely than expert prose, foregrounding a persuasive stance while placing less emphasis on explicit structural scaffolding.

3.3 Topic-based variation in metadiscourse use

Figure 2 illustrates the topic-based distribution of interactive and interactional markers in ChatGPT-generated argumentative essays.

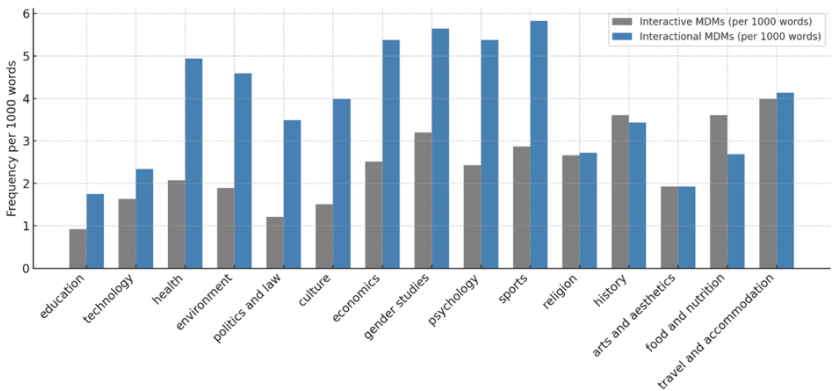


Figure 2. Topic-based distribution of interactive and interactional categories in the corpus of ChatGPT argumentative essays

The largest group of topics is characterized by a clear predominance of interactional resources, indicating a strong emphasis on stance-taking, evaluation, and reader engagement. This group includes education, technology, health, environment, politics and law, culture, economics, gender studies, psychology, and sports. Across these domains, the essays display a marked preference for interpersonal positioning over explicit textual guidance. Rather than foregrounding structural signposting, they rely more on evaluative

expressions, cautious claim formulation, and the consideration of alternative viewpoints. Consequently, the emphasis shifts from formal organizational cues to the way arguments are framed for the reader and how degrees of commitment, distance, or caution are conveyed within the discourse. This tendency is consistent with long-standing findings in metadiscourse research, which indicate that topics involving evaluation, judgement, or ideological positioning typically prompt greater use of interactional resources, as persuasion in such contexts entails heightened interpersonal risk (Hyland, 2005). More recent empirical research on argumentative writing further supports this view, showing that topic exerts a decisive influence on interactional choices, with evaluative prompts eliciting denser use of stance and engagement features irrespective of writers' proficiency level (Yoon, 2021).

Within this group, topics such as politics and law, economics, gender studies, psychology, and sports stand out as areas where disagreement and value-based reasoning are hard to avoid. Arguments in these domains rarely unfold in neutral terms; instead, they tend to draw on competing perspectives, moral judgements, and forms of social positioning. This, in turn, creates a stronger need to signal stance and to engage the reader more directly. Health and environment show a comparable tendency, not because they are inherently argumentative, but because they are widely perceived as matters of public concern, where information is often accompanied by evaluation, caution, and a sense of responsibility. Technology and education, by contrast, follow a slightly different path. Although interactional resources remain prominent, engagement here is generally more measured, suggesting that these topics are approached as familiar or semi-technical rather than openly confrontational. Similar patterns have been reported in recent corpus-based studies comparing AI-generated and human-written argumentative essays, which note a pronounced reliance on evaluative language when responding to socially sensitive prompts, alongside more limited forms of dialogic engagement (Jiang & Hyland, 2025; Alghazo et al., 2024).

Similar patterns have been noted in recent work on AI-generated academic and argumentative writing. Studies focusing on ChatGPT and comparable large language models suggest that, when dealing with controversial, opinion-oriented, or socially sensitive prompts, these texts tend to rely more heavily on stance-taking, evaluation, and reader engagement than on neutral or purely descriptive exposition (Cao et al., 2025; Gao et al., 2025). It could be suggested that this preference reflects an attempt to respond to task demands and prevailing discourse expectations in academic contexts. This shows that the dominance of interactional resources in such domains is unlikely to be incidental, but instead points to a systematic sensitivity to topic and rhetorical context in AI-produced discourse. Evidence from genre-focused analyses further indicates that stance serves not only an evaluative function but also contributes to rhetorical visibility and persuasive force, particularly in high-stakes academic genres where positioning and credibility are closely intertwined (Zou & Fan, 2025).

By contrast, food and nutrition and history stand out as the domains in which interactive resources exceed interactional ones. In this regard, the emphasis shifts toward textual organization, coherence, and reader guidance, with essays favouring explanation, sequencing, and informational clarity over overt stance-taking. Rather than persuading readers through evaluation, the texts tend to present information in a structured and accessible way, a pattern that aligns well with the instructional or advisory character of food- and history-related discourses. Previous research has shown that explanatory and informational writing typically privileges interactive metadiscourse resources, since its primary aim is to guide readers through content and maintain coherence rather than to foreground interpersonal positioning (Ådel, 2006). Comparable distributions have been reported in genre-based studies of academic writing more generally, where informational and procedural topics consistently place greater weight on organizational guidance than on stance (Biber et al., 2011). The present findings suggest that AI-generated essays reproduce this familiar functional distinction, adjusting their rhetorical organization in line with genre expectations rather than applying a uniform strategy across topics.

A third pattern emerges in religion, arts and aesthetics, and travel and accommodation, where interactive and interactional resources appear at broadly comparable levels. This relative balance points to a rhetorical orientation in which engagement with the reader is tempered by careful textual organization. In the case of religion, the near parity is associated with a largely expository mode, in which interpersonal positioning remains restrained, arguably as a way of avoiding overt ideological confrontation. Arts and aesthetics tend to adopt an interpretive stance in which evaluation is present but not strongly foregrounded, allowing space for commentary without sustained stance marking. Travel and accommodation occupy a different position again, combining organizational guidance with a moderate level of engagement that reflects the experiential and descriptive nature of the topic, where texts both guide readers and sustain an involved tone. Comparable balanced distributions have also been noted in studies comparing human- and AI-generated texts, particularly in genres which prioritize description and interpretation over overt argumentation (Alghazo et al., 2024).

Overall, Figure 2 highlights that AI-generated essays do not follow a single, uniform metadiscursive pattern across topics. Instead, they exhibit systematic variation in the use of interactional and interactive resources across domains and communicative purposes. It could be argued that this variability mirrors patterns long observed in corpus-based studies of student and academic writing, where stance and organizational choices shift in relation to register, discipline, and rhetorical goals (Biber et al., 2011). This shows that AI-generated texts cannot be reduced to mechanically reproduced templates; rather, they reflect topic-sensitive rhetorical behaviour that aligns with emerging accounts of AI discourse as contextually situated and responsive to discourse expectations (Jiang & Hyland, 2025a; Majdik & Graham, 2024).

4. Conclusion

This study set out to examine how ChatGPT employs interactive and interactional metadiscourse markers in argumentative essays across a range of topics, drawing on Hyland's (2005) Interpersonal Model of Metadiscourse. By combining quantitative and qualitative analyses, it explores the distribution of metadiscourse categories and their use across topics. Thereby, the study offers a clearer picture of the extent to which AI-generated essays reflect the rhetorical practices of academic argumentation.

The analysis shows that ChatGPT-generated argumentative essays draw on both interactive and interactional metadiscourse markers, shaping arguments through textual organization and the positioning of claims for readers. Hedging is particularly noticeable among interactional resources, pointing to a tendency to frame arguments cautiously and leave room for alternative positions rather than making categorical claims. On the interactive side, transitions are especially prominent, highlighting their role in maintaining cohesion and guiding readers through the line of argument. This pattern suggests that persuasion in these essays is achieved largely through careful positioning and textual flow.

The comparison across the corpus pointed to a notable imbalance: interactional markers appeared almost twice as often as interactive ones. Higher use of interactional markers than interactive ones shows that the essays are primarily oriented toward stance construction and interpersonal negotiation rather than toward explicit textual organization. This pattern reflects the rhetorical priorities of argumentative writing, where persuasion depends on evaluating positions, qualifying claims, and anticipating alternative viewpoints. From this perspective, interactional resources play a central role in shaping how arguments are presented to the reader. By contrast, interactive resources, which are more closely associated with structural guidance and source-based organization, appear less salient in a genre that privileges positioning and evaluation over exposition. The observed imbalance therefore points to a genre-driven rhetorical orientation in AI-generated argumentative essays, rather than to limitations related to writing competence or academic expertise.

Importantly, the topic-based analysis indicates that the relative prominence of these resources varies systematically across topics. Evaluative and socially sensitive domains such as sports, gender studies, psychology, economics, and politics and law showed a strong preference for interactional resources, reflecting the need to manage disagreement, moral judgement, and interpersonal risk. By contrast, food and nutrition and history emerged as the only domains in which interactive resources clearly predominated, indicating an orientation toward explanation, sequencing, and clarity of information. Other topics, including, arts and aesthetics, travel and accommodation, and religion displayed a more balanced distribution, combining measured engagement with careful textual organization. These patterns report that AI-generated

argumentation is responsive to topic-specific rhetorical demands rather than being mechanically uniform.

Overall, the findings suggest that ChatGPT produces argumentative essays that look convincingly academic on the surface and adopt many of the rhetorical features associated with human writing. At the same time, this resemblance appears somewhat intensified, as the texts rely heavily on stance and evaluation. Although the model shows clear awareness of genre conventions and topic expectations, it consistently gives priority to interpersonal positioning. As a result, AI-generated essays often come across as persuasive and dialogic; however, they tend to fall short in structural balance and evidentiary grounding.

From a pedagogical perspective, these findings point to both promise and caution. AI-generated essays can be useful teaching resources for demonstrating how stance, hedging, and reader engagement function in academic argumentation, helping students notice the interpersonal work that writing performs. However, they should not be treated as fully developed models of academic prose. Used uncritically, such texts may encourage an overemphasis on evaluative language at the expense of evidence, structure, and disciplinary grounding. Pedagogical engagement with AI writing therefore needs to be deliberate and reflective, foregrounding balance and critical awareness rather than simple imitation.

At a theoretical level, the study contributes to ongoing debates on metadiscourse, genre, and AI-mediated writing by showing that AI-generated argumentation closely resembles patterns commonly associated with novice writing and soft-disciplinary discourse. By bringing both the strengths and the limitations of AI writing into focus, the study reinforces the importance of genre- and topic-sensitive approaches to academic discourse and highlights the need for continued research into how emerging AI technologies are shaping, and potentially reshaping, academic communication.

References

- Abdi, R. (2002). Interpersonal metadiscourse: An indicator of interaction and identity. *Discourse Studies*, 4(2), 139–145. <https://doi.org/10.1177/14614456020040020101>
- Ädel, A. (2006). *Metadiscourse in L1 and L2 English*. John Benjamins.
- Alghazo, S., Rababah, G., El-Dakhs, D. A. S., & Mustafa, A. (2025). Engagement strategies in human-written and AI-generated academic essays: A corpus-based study. *Ampersand*, 100237.
- Alshalan, K., & Alyousef, H. (2024). A corpus-based study of the experiential meaning in business students' handwritten and ChatGPT-generated argumentative essays. *Middle East Research Journal of Linguistics and Literature*, 4(5), 93–103.
- Amirjalili, F., Neysani, M., & Nikbakht, A. (2024). Exploring the boundaries of authorship: A comparative analysis of AI-generated text and human

- academic writing in English literature. *Frontiers in Education*, 9, 1347421. <https://doi.org/10.3389/feduc.2024.1347421>
- Cao, X., Lin, Y. J., Zhang, J. H., Tang, Y. P., Zhang, M. P., & Gao, H. Y. (2025). Students' perceptions about the opportunities and challenges of ChatGPT in higher education: a cross-sectional survey based in China. *Education and Information Technologies*, 1-20.
- Dahl, T. (2004). Textual metadiscourse in research articles: A marker of national culture or of academic discipline? *Journal of Pragmatics*, 36(10), 1807–1825. <https://doi.org/10.1016/j.pragma.2004.05.004>
- Dynel, M. (2023). Lessons in linguistics with ChatGPT: Metapragmatics, metacommunication, metadiscourse and metalanguage in human–AI interactions. *Language & Communication*, 93, 107–124. <https://doi.org/10.1016/j.langcom.2023.09.002>
- Flowerdew, J., & Wang, S. H. (2023). Academic writing in the age of AI: Opportunities and challenges. *Journal of English for Academic Purposes*, 61, 101230.
- Gao, Y. (2023). Exploring stance in AI-generated academic writing: A metadiscourse perspective. *Journal of English for Academic Purposes*, 62, 101188. <https://doi.org/10.1016/j.jeap.2023.101188>
- Gao, R., Yu, D., Gao, B., Hua, H., Hui, Z., Gao, J., & Yin, C. (2025). Legal regulation of AI-assisted academic writing: challenges, frameworks, and pathways. *Frontiers in Artificial Intelligence*, 8, 1546064.
- Hongyan, Wang and Saeed, Murad Abdu, A Comparative Study of Stance and Engagement in Ai-Generated vs. Human-Written Article Abstracts Across Multidisciplinary Research. Available at SSRN: <https://ssrn.com/abstract=5343337> or <http://dx.doi.org/10.2139/ssrn.5343337>
- Hyland, K. (1996). Writing without conviction? Hedging in science research articles. *Applied Linguistics*, 17(4), 433–454. <https://doi.org/10.1093/applin/17.4.433>
- Hyland, K. (1998). Boosting, hedging and the negotiation of academic knowledge. *Text*, 18(3), 349–382.
- Hyland, K. (2001). Bringing in the reader: Addressee features in academic writing. *Written Communication*, 18(4), 549–574. <https://doi.org/10.1177/0741088301018004005>
- Hyland, K. (2005). *Metadiscourse: Exploring interaction in writing*. Continuum.
- Hyland, K. (2019). *Metadiscourse: Exploring interaction in writing* (2nd ed.). Bloomsbury.
- Hyland, K., & Jiang, F. (2017). Is academic writing becoming more informal? *English for Specific Purposes*, 45, 40–51. <https://doi.org/10.1016/j.esp.2016.09.001>
- Hyland, K., & Tse, P. (2004). Metadiscourse in academic writing: A reappraisal. *Applied Linguistics*, 25(2), 156–177.

- Jiang, F. K., & Hyland, K. (2025a). Rhetorical distinctions: Comparing metadiscourse in essays by ChatGPT and students. *English for Specific Purposes*, 79, 17–29. <https://doi.org/10.1016/j.esp.2025.03.001>
- Jiang, F. K., & Hyland, K. (2025b). Metadiscursive nouns in academic argument: ChatGPT vs student practices. *Journal of English for Academic Purposes*, 75, 101514. <https://doi.org/10.1016/j.jeap.2025.101514>
- Jiang, F. K., & Hyland, K. (2025c). Does ChatGPT write like a student? Engagement markers in argumentative essays. *Written Communication*, 42(3), 463–492.
- Jiang, F. K., & Hyland, K. (2025d). Does ChatGPT argue like students? Bundles in argumentative essays. *Applied Linguistics*, 46(3), 375–391.
- Kobzová, D. (2023). Content-oriented hedging in academic English generated by ChatGPT (Unpublished master's thesis). Charles University.
- Li, S. (2025). Generative AI and Second Language Writing. *Digital Studies in Language and Literature*, (0).
- Liu, Z., & Deng, J. (2023). Hedging and stance in AI-assisted student writing. *Lingua*, 293, 103457. <https://doi.org/10.1016/j.lingua.2023.103457>
- Majdik, Z. P., & Graham, S. S. (2024). Rhetoric of/with AI: An introduction. *Rhetoric Society Quarterly*, 54(3), 222–231
- Shalevska, E. (2024). Hedges and boosters in AI and human writing: A comparative analysis. *Knowledge-International Journal*, 65(5), 505–510.
- Triki, A. (2014). Metadiscourse in argumentative writing: A cross-sectional study of Tunisian EFL learners. *International Journal of Humanities and Social Science*, 4(12), 108–119.
- Triki, A. (2019). Metadiscourse in EFL university students' writing: A corpus-based study of interactive and interactional markers. *International Journal of English Linguistics*, 9(2), 69–81. <https://doi.org/10.5539/ijel.v9n2p69>
- Yao, G., & Liu, Z. (2025). Can AI simulate or emulate human stance? Using metadiscourse to compare GPT-generated and human-authored academic book reviews. *Journal of Pragmatics*, 247, 103–115. <https://doi.org/10.1016/j.pragma.2025.07.018>
- Yoon, H.-J. (2021). Interactions in EFL argumentative writing: Effects of topic, L1 background, and L2 proficiency on interactional metadiscourse. *Reading and Writing*, 34, 705–725.
- Yuan, X. (2024). Uncertainty in machine writing: The role of hedges in ChatGPT texts. *Journal of Pragmatics*, 210, 47–61. <https://doi.org/10.1016/j.pragma.2023.10.013>
- Zarei, G. R., & Mansoori, S. (2011). A contrastive study on metadiscourse elements used in humanities vs. non humanities across Persian and English. *English Language Teaching*, 4(1), 42–50.
- Zhang, M., & Zhang, J. (2025a). Reflexivity in human-written and ChatGPT-generated English research article abstracts: A comparison of metadiscourse. *Applied Linguistics*.

<https://doi.org/10.1093/applin/amaf032>

Zhang, M., & Zhang, J. (2025b). Disciplinary variation of metadiscourse: A comparison of human-written and ChatGPT-generated English research article abstracts. *Journal of English for Academic Purposes*, 76, 101540. <https://doi.org/10.1016/j.jeap.2025.101540>

Zou, H. (J.), & Fan, Y. (2025). Promotion of medical research: Stance in AI-generated and scholar-written highlights. *Journal of English for Academic Purposes*, 78, 101593. <https://doi.org/10.1177/14614456020040020101>

Appendix

Education

1. Should college education be free for everyone?
2. Are standardized tests an accurate measure of student intelligence?
3. Should schools replace textbooks with digital devices?
4. Is homework beneficial to students' learning experience?
5. Should students be allowed to grade their teachers?

Technology

1. Does technology isolate people or bring them closer together?
2. Should artificial intelligence be regulated by governments?
3. Is privacy more important than national security in the age of digital surveillance?
4. Should there be age restrictions on smartphone use for children?
5. Are video games to blame for violent behavior in children?

Health

1. Should vaccinations be mandatory for all citizens?
2. Is universal healthcare a right or a privilege?
3. Should the use of performance-enhancing drugs in sports be allowed?
4. Is the legalization of marijuana beneficial to society?
5. Should fast food companies be held accountable for the obesity epidemic?

Environment

1. Should developed countries be required to do more to combat climate change?
2. Is nuclear energy a safe alternative to fossil fuels?
3. Should animal testing be banned for scientific research?
4. Can renewable energy fully replace fossil fuels?
5. Is overpopulation the greatest threat to the environment?

Politics and Law

1. Should voting be mandatory in democratic countries?
2. Is censorship ever justified in a free society?

3. Should countries have open borders for immigrants?
4. Should lobbying be banned in politics?
5. Is socialism a better system than capitalism?

Culture

1. Is modern pop culture contributing to the decline of morality?
2. Should cultural appropriation be considered a crime?
3. Is it necessary to preserve indigenous languages?
4. Are beauty pageants outdated in modern society?
5. Should parents control what their children watch and read?

Economics

1. Should universal basic income be implemented to address poverty?
2. Is income inequality a serious problem that needs immediate attention?
3. Should the minimum wage be raised to a living wage?
4. Is globalization beneficial or detrimental to local economies?
5. Should governments provide financial assistance to struggling industries?

Gender Studies

1. Is feminism still relevant in today's society?
2. Should gender-neutral bathrooms be mandated in public spaces?
3. Are traditional gender roles harmful to society?
4. Is the representation of women in media improving?
5. Should men face consequences for sexist behavior?

Psychology

1. Is mental illness adequately addressed in today's society?
2. Should therapy be a mandatory part of education for young people?
3. Is social media addiction a real psychological condition?
4. Should parents be held accountable for their children's mental health?
5. Is it ethical to use psychological tactics in marketing?

Sports

1. Should college athletes be paid for their performance?
2. Is the role of sports in society primarily positive?
3. Should performance-enhancing drugs be allowed in professional sports?
4. Are sports fandoms detrimental to mental health?
5. Should violent sports, such as boxing, be banned?

Religion

1. Should religion play a role in government policy?
2. Is religious tolerance essential for global peace?

3. Are atheists discriminated against in modern society?
4. Should religious organizations be taxed?
5. Is faith-based education beneficial or harmful?

History

1. Should controversial historical monuments be removed?
2. Is history best understood through the lens of the victors?
3. Did colonialism have a net positive effect on the world?
4. Should reparations be paid for historical injustices?
5. Is the teaching of history in schools biased?

Art and Aesthetics

1. Is art essential for a well-rounded education?
2. Should public funding be used to support the arts?
3. Is modern art a legitimate form of artistic expression?
4. Should controversial art be censored?
5. Is street art a valid form of artistic expression or vandalism?

Food and Nutrition

1. Should governments regulate fast food companies?
2. Is a vegan diet healthier than a meat-based diet?
3. Should schools provide free meals to all students?
4. Is organic food worth the extra cost?
5. Should the consumption of sugary drinks be taxed?

Travel and Tourism

1. Is tourism harmful to local cultures?
2. Should countries limit the number of tourists they receive?
3. Is eco-tourism a sustainable alternative to traditional tourism?
4. Should travel agencies be required to disclose the carbon footprint of trips?
5. Is it ethical to visit destinations affected by disasters?

TÜRKÇEDE ÜNLÜ FORMANTLARININ META-ANALİTİK İNCELENMESİ: YÖNTEMBİLİMSEL FARKLILIKLAR VE ETKİLERİ¹

Ali Çağan KAYA

Lisans Öğr., Hacettepe Üniversitesi, İngiliz Dilbilimi Bölümü, ali-kaya@hacettepe.edu.tr, ORCID: 0009-0007-9386-9766

Emre YAĞLI

Doç. Dr., Hacettepe Üniversitesi, İngiliz Dilbilimi Bölümü, yaqli@hacettepe.edu.tr, ORCID: 0000-0002-1044-9018

Kaya, A. Ç. & Yağlı, E. (2025). Türkçede ünlü formantlarının meta-analitik incelenmesi: Yöntembilimsel farklılıklar ve etkileri. İçinde O. Çınar, F. Başbuğ & H. Aydemir (Ed.), *Çağdaş dilbilim yazıları I: İMÜ Dilbilimi 15. yıl armağanı* (ss. 170-192). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X008>

ÖZ

Türkçede ünlü formantlarına odaklanan araştırmalar niceliksel olarak zengin bir birikim oluşturmaya karşın, bu çalışmalar arasında belirgin bir yöntem birliği bulunmamaktadır. Bu çalışma, ölçünlü Türkçenin ünlü formantlarını konu alan araştırmalardaki yöntembilimsel farklılıkları ve bu farklılıkların formant değerleri üzerindeki etkisini inceleyen bir meta-analiz sunmaktadır. Çalışma kapsamında, ölçünlü Türkçe üzerine yürütülen 16 araştırma taranmış ve bu araştırmaların yöntembilimsel ölçütleri bir veri tablosuna kodlanarak çalışmanın verisi oluşturulmuştur. Verilerin çözümlenmesinde R programlama dili kullanılarak her bir ünlü ses için ayrı ayrı karma etkiler modelleri oluşturulmuştur. İncelme sonuçları, modellerdeki rastgele etkilerin (çalışma ve çalışma-ünlü etkileşimi), sabit etkilere (cinsiyet, materyal türü, konuşmacı sayısı, kayıt ortamı, kayıt aleti, örnekleme hızı, normalizasyon yöntemi) kıyasla değişkenliği açıklamada istatistiksel olarak daha belirleyici olduğunu göstermiştir. Bu bulgu, formant değerlerindeki değişkenliğin büyük ölçüde çalışmalarda açıkça raporlanmayan veya standartlaştırılmayan örtük yöntemsel tercihlerden kaynaklandığını ortaya koymaktadır. Çalışma, Türkçedeki sesbilgisel betimlemelerin güvenilirliği ve karşılaştırılabilirliği için yöntemsel şeffaflığın, açık veri paylaşımının ve standart ölçüm protokollerinin gerekliliğini vurgulamaktadır.

¹ Bu çalışma TÜBİTAK 2209-A Üniversite Öğrencileri Araştırma Projeleri kapsamında desteklenen “Ölçünlü Türkçenin ünlü formant frekansları” başlıklı proje kapsamında üretilmiştir. Bununla birlikte çalışma, 22-23 Mayıs 2025 tarihlerinde Erzurum Atatürk Üniversitesi’nde düzenlenen 38. Ulusal Dilbilim Kurultayı’nda sunulan aynı başlıklı bildirinin genişletilmiş sürümüdür.

Anahtar Sözcükler: Ünlü formantları, Sesbilgisel betimleme, Meta-analiz, Yöntembilim

ABSTRACT

Although research on vowel formants in Turkish constitutes a substantial body of work, there is a lack of methodological consensus among these studies. This study presents a meta-analysis examining methodological differences in research on vowel formants in Standard Turkish and the effects of these differences on formant values. Within the scope of this study, 16 studies on Standard Turkish were adopted, and the methodological criteria of these studies were coded in a data table to create the dataset. In the data analysis, separate mixed-effects models were fitted for each vowel in the R environment. The results indicated that random effects (e.g., study and study-vowel interaction) in the models were statistically more determinant in explaining variability compared to fixed effects (e.g., gender, material type, number of speakers, recording environment, recording equipment, sampling rate, and normalisation method). This finding indicates that variability in formant values largely stems from implicit methodological choices that are neither clearly reported nor standardised across studies. The study emphasises the necessity of methodological transparency, open data sharing, and standardised measurement protocols to ensure the reliability and comparability of phonetic descriptions in Turkish.

Keywords: Vowel formants, Phonetic description, Meta-analysis, Methodology

1. Giriş

Ünlü formant frekansları, konuşmaya ilişkin en sık betimlenen akustik ölçümlerden biridir. Bu formantlar, bir yandan seslerin üretiminde görev alan eklemleyicilerin fiziksel özelliklerini yansıtarak sesin üretimsel doğasının anlaşılmasını sağlarken, öte yandan konuşmanın algısal ve işitsel bileşenlerinin çözülmesine olanak tanır. Bu çok yönlü yapısı sayesinde formant frekansları; dil değişkenliğinin incelenmesinden konuşmacı tanımlama çalışmalarına, konuşma bozukluklarının tanı ve takibinden çocuklarda dil gelişiminin izlenmesine ve ikinci dil öğrencilerinin hedef seslere ne ölçüde yaklaştığının değerlendirilmesine kadar pek çok alanda yaygın biçimde kullanılmaktadır. Ayrıca, konuşma teknolojileri alanında, otomatik konuşma tanıma ve sentez sistemlerinin geliştirilmesinde de formant bilgisi temel alınan değişkenlerdendir. Bu çok yönlü kullanımlarına rağmen, Türkçenin ünlü formant verilerine ilişkin mevcut betimlemelerin yöntemsel karşılaştırmalardan yoksun olduğu görülmektedir. Bu çalışma, Türkçede ünlü formantlarının ölçümüne ilişkin betimlemeleri karşılaştırmalı biçimde ele alarak, alandaki yöntembilimsel çeşitliliği ve bunun formant değerleri üzerindeki etkisini değerlendiren bir meta-analiz niteliğindedir.

Ünlü formantları, dilbilimsel betimlemede merkezi bir rol oynamaktadır. 1948’de yapılan öncü çalışmalarla temelleri atılan bu alan, günümüze kadar uzanan süreçte yoğun biçimde ele alınmış ve dikkate

değer bir birikim oluşturmıştır. Ancak Maurer (2024, s. 84), bugüne dek gerçekleştirilen ünlü ses betimlemelerinin sayıca fazla olmasına karşın ciddi yöntembilimsel sorunlar barındırdığına dikkat çekmektedir. Evrensel alanyazında gözlemlenen bu durum, Türkçe için de geçerlidir. Örneğin Aydın ve Uzun (2020, s. 297), Türkçedeki ünlü seslere yönelik çalışmalarda, karşılaştırmalı inceleme açısından önemli bir adım olan normalizasyon işleminin yalnızca sınırlı sayıda araştırmada yer aldığını belirtmektedir. Aydın ve Uzun'un (2020) bu gözlemi, Türkçede ünlü formant değerlerinin güvenilir ve karşılaştırılabilir biçimde sunulmasındaki önemli bir eksikliğe işaret etmektedir.

Bununla birlikte, formant ölçümlerinin karşılaştırılabilirliğini etkileyen tek etmen normalizasyon değildir. Betimleyici çalışmaların yöntemsel tasarımı; kullanılan söyletim materyalinin türü, konuşmacı profili, kayıt ortamı, örnekleme hızı, ölçüm tekniği ve veri işleme süreçleri gibi çok sayıda etmenden etkilenmektedir. Formant frekanslarının bu denli çok boyutlu bir üretimsel doğaya sahip olması, aynı ünlünün farklı bireylerce üretildiğinde önemli ölçüde farklı formant değerleriyle karşılaşılmaya neden olmaktadır (Peterson ve Barney, 1952). Nitekim Hillenbrand vd. (1995) ve Clopper, Pisoni ve De Jong (2005), yaş, cinsiyet, konuşma biçimi ve üretim bağlamı gibi etkenlerin formant değerleri üzerinde sistematik etkiler oluşturduğunu ortaya koymuştur. Bu durum, yalnızca üretimsel değil, aynı zamanda algısal düzlemde de formantların değerlendirilmesini güçleştirmekte; elde edilen verilerin genellenebilirliğini sınırlamaktadır.

Bu bağlamda, yalnızca formant değerlerinin sayısal karşılaştırmalarla değil, bu değerlerin üretildiği yöntembilimsel koşulların bütüncül biçimde değerlendirilmesi gerekmektedir. Araştırmaların veri toplama, inceleme ve raporlama süreçlerinde karşılaşılan şeffaflık eksikliği, alanyazında birikimin bütünleştirilmesini zorlaştırmakta ve yeni çalışmaların önceki bulgularla tutarlı biçimde ilerlemesini engellemektedir (Winter ve Grice, 2021).

Bu çalışma, söz konusu yöntemsel etmenlerin formant değerlerine olan etkisini ele almanın yanı sıra, Türkçeye odaklanan çalışmaları yöntemsel açıdan karşılaştırmalı bir bakışla değerlendirmektedir. Böylelikle, Türkçedeki ünlü formantlarına ilişkin daha bütüncül, şeffaf ve tutarlı bir veri zemininin oluşturulmasına katkı sunmayı amaçlamaktadır. Ayrıca, yöntembilimsel çeşitliliğin sesbilgisel veriler üzerindeki etkisini nicel olarak değerlendiren bu çalışma, dilsel veri yorumunda açıklık ve yeniden üretilebilirlik ilkelerinin önemini vurgulayan çağdaş eğilimlere paralel bir yaklaşım benimsemektedir.

2. Ünlü Formantlarının Dilbilimsel Betimlemedeki Önemi

Konuşma, doğası gereği sürekli değişim halinde olan sesletimsel yapılardan oluşur. Bu dinamik yapı, akustik çözümleme açısından önemli güçlükler doğurmaktadır. Akustik betimlemeye odaklanan çalışmalar, bu güçlüklerle başa çıkabilmek amacıyla, ses üretiminin en sabit olduğu alanlara

yönelir. Bu sabit alanlar, ünlü sesler açısından bakıldığında birinci (F1) ve ikinci (F2) formant değerleridir.

Formantlardaki temel değişkenliğin kaynağı, ses üretimi sırasında başta dil ve dudak olmak üzere ses yolunda gerçekleşen biçimsel farklılıklardır. Ünlülerin üretimi sırasında ses yolu boyunca gerçekleşen bu farklılıklar, her bir ünlü için karakteristik bir formant değerinin ortaya çıkmasına neden olur. Alanyazında bu biçimlenimlerle formant frekansları arasındaki ilişkinin oldukça karmaşık olduğu vurgulanmaktadır (Fant, 1960; Stevens, 1998). Bu karmaşayı en aza indirmek amacıyla, dilsel çözümlemelerde genellikle ilk iki formant değerine odaklanılır. Bu iki formant, ünlülerin ağız boşluğundaki en yüksek noktalarının ön-arka ve yüksek-alçak eksenlerinde konumlandırılmasına olanak tanır.

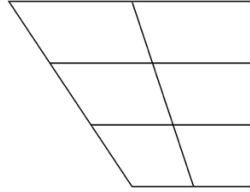
Bu bağlamda, birinci formant olarak adlandırılan F1 değeri ünlülerin dikey düzlemdeki konumu ile ilgilidir. Açık ünlüler genellikle yüksek F1 değerleriyle, kapalı ünlüler ise düşük F1 değerleriyle gözlemlenir. İkinci formant olarak nitelenen F2 değeri ise ünlünün yatay düzlemdeki konumunu, yani ön ya da arka ünlü olma niteliğini yansıtır. Ön ünlüler daha yüksek F2 değerine sahipken, arka ünlüler daha düşük F2 değerleriyle tanımlanır. Bu tanımlamalar doğrultusunda, ünlülerin F1 ve F2 değerleriyle temsil edilmesi, sesbilgisel betimlemelerde yaygın biçimde kullanılan bir yaklaşımdır. Söz konusu formant değerleri, genellikle ters çevrilmiş yatay ekseninde F2 ve dikey ekseninde F1 olacak şekilde (Bkz. Şekil 1), Uluslararası Sesbilgisi Alfabesi (İng. International Phonetic Alphabet, IPA) ünlü dörtgeni ile hizalı dağılım grafikleri aracılığıyla gösterilir (Bkz. Şekil 2).^{1 2}

¹ Ünlü dörtgeni, dünya dillerinde yaygın biçimde gözlemlenen temel ünlü seslerin sesletimsel ve akustik özelliklerini görselleştirmek amacıyla kullanılan kavramsal bir çerçevedir. Bu gösterim, ilk olarak Jones (1917) tarafından geliştirilen *asal ünlü sistemi* (İng. cardinal vowel system) temel alınarak oluşturulmuş ve Uluslararası Sesbilgisi Derneği (İng. International Phonetic Association) tarafından zamanla güncellenmiştir. Ünlüler, ağız boşluğundaki eklemleyici konumlarına göre yatay ekseninde dilin ön-arka pozisyonu, dikey ekseninde ise dilin yüksekliği dikkate alınarak dörtgen biçiminde yerleştirilir (Pfitzinger ve Niebuhr, 2011). Bu görsel yapı hem sesbilgisi öğretiminde hem de dilsel verilerin betimlenmesinde temel bir araç işlevi görür.

² Şekil 1 ve Şekil 2 arasındaki temel fark, ilkinin formant verilerine dayalı bir ünlü görselleştirmesi, ikincisinin ise idealize edilmiş bir ünlü dörtgeni sunmasıdır. Ünlü dörtgeni, ünlülerin yaklaşık sesletimsel yerleşimlerini temsil eden kavramsal bir şemadır; bu yerleşimler, ölçüme değil, dilin ön-arka ve yüksek-alçak eksenlerine göre belirlenen ideal konumlara dayanır. Buna karşılık, ünlü görselleştirmesi her bir ünlünün gerçek F1 ve F2 formant frekansları esas alınarak oluşturulur ve bu yönüyle daha ayrıntılı, ölçüme dayalı bir betimleme sağlar. Günümüzde deneysel sesbilgisi ve toplumdilbilim alanındaki çalışmalar, bu tür formant tabanlı görselleştirmeleri yaygın biçimde kullanmaktadır.



Şekil 1. F1 ve F2 değerlerinin gösterimi



Şekil 2. Ünlü dörtgeni

Ünlü formantları bireyden bireye anlamlı ölçüde değişkenlik gösterebilir. Nitekim aynı ünlünün farklı konuşurlar tarafından üretimi sırasında ölçülen formant frekanslarının önemli ölçüde farklılık gösterdiği uzun zamandır bilinmektedir (Peterson ve Barney, 1952). Hatta, neredeyse aynı fizyolojik yapıya sahip olduğu düşünülen tek yumurta ikizleri üzerinde yapılan çalışmalar bile bu tür farklılıkların izlenebildiğini ortaya koymuştur (Loakes, 2006; Nolan ve Oh, 1996). Bu durum, bireyler arasındaki farklılıkların yalnızca fizyolojik temellere dayanmadığını; aynı zamanda öğrenilmiş bireysel davranışlarla da şekillendiğini göstermektedir.

Formant frekanslarını etkileyen etmenler oldukça çeşitlidir ve hem konuşucu temelli hem de yöntemsel değişkenleri kapsamaktadır. Öncelikle, konuşucunun yaşı ve cinsiyeti, formant frekansları üzerinde belirleyici bir etkiye sahiptir. Ses kıvrımı uzunluğu ve biçimi gibi anatomik özellikler, akustik çıktının karakterini doğrudan etkiler (Chung vd., 2012; Hillenbrand vd., 1995; Lee, Potamianos ve Narayanan, 1999; Peterson ve Barney, 1952; Bickley, 1989). Bunun yanı sıra, konuşma biçimi ve kesiti de önemli bir değişkendir. Özenli ve dikkatli konuşma biçimi, çoğu zaman daha uç değerlerde formant üretimiyle sonuçlanır (Hillenbrand, 2013; Lindbloom, 1963; Nearey, 1989; Strange, 1989). Konuşma hızı ve tempo da formant yapısını etkileyen bir başka etmendirdir; hızlı konuşma, ünlülerin üretim süresini kısaltarak formant alanını daraltmakta ve değerlerin merkezileşmesine neden olmaktadır (Lindbloom, 1963; Turner, Tjaden ve Weismer, 1995). Bunun yanı sıra, söyleyişin duygusal ya da vurgusal tonunu ifade eden bürünsel özellikler de formant frekanslarında değişikliklere yol açabilir (Fletcher vd., 2015). Sesbilgisel bağlam da göz ardı edilemeyecek bir etkindir; ünlülerin önünde ya da sonrasında yer alan sesler, özellikle eşsesletim (İng. coarticulation) süreçleri aracılığıyla formant yapılarını değiştirebilir

(Lindbloom, 1963; Nearey ve Assmann, 1986; Strange, 1989).³ Ayrıca, konuşucunun toplumsal kimliği, toplumsal konumu ya da bölgesel aidiyeti gibi toplumdilbilimsel etmenler, ses üretiminde sistemli farklılıklara neden olabilmektedir (Clopper, Pisoni ve De Jong, 2005; Labov, Ash ve Boberg, 2006).

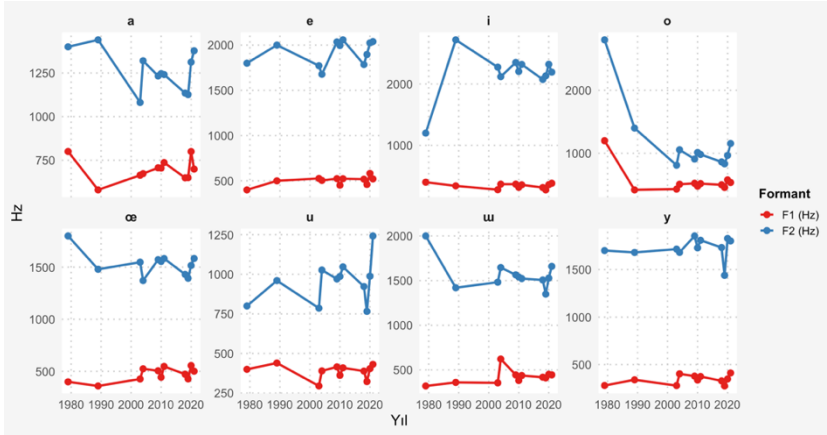
Yöntembilimsel açıdan değerlendirildiğinde, kayıt koşulları ve ölçüm araçları da formant frekanslarının güvenilirliğini doğrudan etkiler. Kullanılan mikrofona türü, örnekleme oranı ve analiz yazılımı gibi unsurlar, ölçüm doğruluğunu belirleyen temel faktörler arasında yer alır (Hillenbrand vd., 1995). Ölçüm işlemlerinin ayrıntıları da bu doğrultuda önem kazanır; formant izleme yöntemleri, analiz penceresinin genişliği ve kullanılan otomatik çıkarım algoritmaları gibi değişkenler, elde edilen sonuçların farklılaşmasına neden olabilir (Bickley, 1989; Fox ve Jacewicz, 2009; Jin ve Liu, 2013; Ladefoged, 1967; Morrison ve Nearey, 2007). Ayrıca, örneklem büyüklüğü ve konuşmacı seçimi de veri yapısını etkileyen önemli etmenlerdendir. Aynı konuşmacının birden fazla üretiminin ortalamasını almak değişkenliği azaltabilirken, farklı konuşucuların verilerini doğrudan bir araya getirmek karışıklıklara yol açabilir (Hillenbrand, 2013; Peterson ve Barney, 1952). Son olarak, normalizasyon yöntemleri, konuşucular arası anatomik farklılıkları denetim altına almak amacıyla geliştirilmiş analitik araçlar arasında yer alır. Lobanov (1971) ve Nearey (1978) gibi yöntemler, bu amaçla yaygın biçimde kullanılmakta; ancak her bir normalizasyon yöntemi, farklı varyans kaynaklarını azaltma konusunda değişen derecelerde etkilidir.⁴ Bu nedenle, yöntem seçimi, elde edilen verilerin karşılaştırılabilirliği açısından kritik öneme sahiptir.

Tüm bu etmenler dikkate alındığında, ünlü formantlarının dilbilimsel betimlemesi yalnızca bir ölçüm süreci değil; çok katmanlı, disiplinlerarası ve yöntembilimsel olarak hassas bir değerlendirme sürecidir. Bu karmaşık yapı, formant verilerinin dikkatli yorumlanmasını ve karşılaştırmalı çalışmalarda yöntemsel çeşitliliğin açık biçimde raporlanmasını gerekli kılmaktadır.

³ Eşsesletim, bir sesin üretiminin, öncesindeki veya sonrasındaki seslerin eklemleyici özelliklerinden etkilenmesi durumudur. Konuşma akışı içindeki bu etkileşim, seslerin birbirine yaklaşmasına, özelliklerinin kaynaşmasına ya da değişmesine yol açabilir. Örneğin, bir ünlünün formant değerleri, onu izleyen ya da önceleyen ünsüzün yeri ve biçimine bağlı olarak değişiklik gösterebilir. Türkçeye yönelik bu konuda son dönemde yürütülen bir çalışma için bakınız Arslan ve Uzun (2025).

⁴ Ünlü formant frekanslarında gözlemlenen toplumsal ve anatomik farklılıkları sabitlemek amacıyla literatürde çeşitli normalizasyon yöntemleri kullanılmaktadır (Studdert-Kennedy, 1964). Normalizasyon, geçerli bilgisayarlı işlemler aracılığıyla ünlü seslerin farklı konuşur grupları ve dil değişimleri bağlamında karşılaştırılabilir hale getirilmesini sağlar (Clopper, 2009). Örneğin, dil değişimi temelli bir çalışmada cinsiyet ve yaş gibi etkilerin giderilmesi hedeflenirken; cinsiyet temelli bir çalışmada etnisite, bölgesel kimlik ve yaş gibi değişkenler kontrol altına alınır. En yaygın kullanılan iki yöntem Lobanov (Lobanov, 1971) ve bark (Traunmüller, 1990) normalizasyonudur. Lobanov yöntemi, fizyolojik (Örneğin, yaş, cinsiyet) ve toplumdilbilimsel (Örneğin, dil değişimi kullanımı) farklılıkları en aza indirerek özellikle ölçünlü dil betimlemelerinde tercih edilir. Bark normalizasyonu ise işitsel algıya odaklanır ve genellikle algı araştırmalarında kullanılır.

Şekil 3 takip edildiğinde, ölçünlü Türkçeye yönelik çalışmalarda betimlenen formant değerlerinin kendi içinde değişkenlik sergilediği görülmektedir. Bu duruma yönelik bir fikir vermesi amacıyla sunulan Şekil 4 ise, F1 ve F2 değerlerinin yıllar içerisindeki değişimini görselleştirmektedir.



Şekil 4. Ölçünlü Türkçeye odaklanan çalışmalarda sunulan F1 ve F2 değerlerinin yıllara göre değişimi

Şekil 4'te izlenen eğilim, ölçünlü Türkçeye yönelik formant betimlemesinin zaman içinde farklılaştığını ve yöntemsel çeşitliliğin yalnızca bireysel çalışmalarla değil, dönemsel olarak da izlenebilir olduğunu göstermektedir. Bu açıdan Şekil 3 ve 4, alanyazındaki ölçünlü Türkçe verilerinin yöntemsel çoğulluğunu ilk bakışta görünür kılmak açısından önem taşımaktadır. Ancak, formant verilerindeki bu değişkenliğin kaynaklarına ve olası sonuçlarına ilişkin ayrıntılı değerlendirme, bu çalışmanın izleyen bölümlerinde ele alınmıştır.

4. Veri ve Yöntem

Bu bölümde, çalışmanın temelini oluşturan veri kümesi ile bu verilerin incelenmesinde izlenen yöntembilimsel yaklaşım ayrıntılı biçimde sunulmaktadır. İlk olarak, çalışmada incelenen araştırmaların seçimi, kapsamı ve özellikleri açıklanmakta; ardından, elde edilen verilerin çözülmesinde benimsenen analitik araçlar tanıtılmaktadır.

4.1 Çalışmanın verisi

Bu çalışmanın verisi, ölçünlü Türkçenin ünlü formant frekanslarını betimleyen araştırmalardan meta-analiz amacı doğrultusunda oluşturmaktadır. Çalışmanın verisi üç aşamada elde edilmiştir: (1) çalışmalara ilişkin kapsamın belirlenmesi, (2) seçilen çalışmaların yöntembilimsel özelliklerinin ayrıntılı olarak incelenmesi ve (3) elde edilen bilgilerin sistematik bir veri tablosuna dönüştürülmesi.

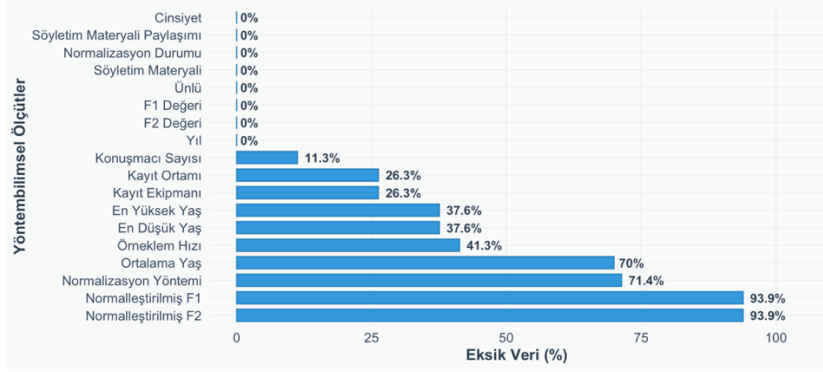
İlk aşamada, Türkçeye yönelik ünlü formant ölçümlerini içeren 22 çalışma belirlenmiştir (Selen (1979), Ergenç (1989), Kılıç (2003), Kılıç & Ögüt (2004), Türk vd. (2004), Çiyiltepe & Orberk (2008), Malkoç (2009, 2011), Çiyiltepe, Bekar & Ergenç (2009), Yılmaz Davutoğlu (2010), Kopkallı-Yavuz (2010), Güven (2015), Korkmaz & Boyacı (2018), Esen Aydın, Kulak Kayıkçı & Suslu (2019), Korkmaz, Boyacı & Tuncer (2019), Aydın & Uzun (2020), Erdem Nas (2020), Koçak (2020), Börtlü (2020), Akbulut & Erdem Nas (2021), Erdem Nas & Yalçınkaya (2021), Korkmaz & Boyacı (2022)). Belirlenen bu çalışmalardan altısı ağırlıklı olarak bölgesel değişkenlik, ikidillilik ve yabancı dil öğretimi bağlamlarında gerçekleştirildiği için kapsam dışı bırakılmıştır. Bunun sonucunda, yalnızca ölçünlü Türkçede, anadili Türkçe olan konuşurlarla gerçekleştirilen ve üretim verisine dayalı formant ölçümlerini içeren 16 çalışma veri havuzuna dahil edilmiştir. Bu kapsamda çalışmanın çözümlemesine dahil edilen çalışmalar şunlardır: Selen (1979), Ergenç (1989), Kılıç (2003), Kılıç & Ögüt (2004), Türk vd. (2004), Malkoç (2009, 2011), Yılmaz Davutoğlu (2010), Kopkallı-Yavuz (2010), Korkmaz & Boyacı (2018), Esen Aydın, Kulak Kayıkçı & Suslu (2019), Aydın & Uzun (2020), Erdem Nas (2020), Börtlü (2020), Akbulut & Erdem Nas (2021), Erdem Nas & Yalçınkaya (2021).

İkinci aşamada, bu 16 çalışmaya yönelik ayrıntılı bir yöntembilimsel okuma gerçekleştirilmiştir. Bu okuma sırasında, her bir çalışmada raporlanan konuşucu sayısı, cinsiyet, yaş, söyletim materyali, söyletim materyalinin paylaşımı, kayıt ortamı, kayıt ekipmanı, kayıt örnekleme hızı, formant değerlerinin normalleştirilip normalleştirilmediği, normalizasyon uygulanmışsa kullanılan yöntem, ünlü ses ve hem normalleştirilen hem de normalleştirilmeyen F1 ve F2 değerleri gibi yöntembilimsel sistematik biçimde değerlendirilmiştir. Bu ayrıntılı okuma sonucunda, çalışmaların önemli bir kısmında yöntemsel ölçütlerin eksik veya farklı biçimlerde sunulduğu görülmüştür.

Üçüncü aşamada, yöntembilimsel okuma sonucunda elde edilen ölçütlerin yer aldığı bir veri tablosu oluşturulmuştur. Ancak bu süreçte bazı güçlüklerle karşılaşmıştır. Öncelikle, birçok çalışmada konuşucu yaşına ilişkin bilginin farklı biçimlerde raporlandığı görülmüştür. Örneğin bazı çalışmalar yalnızca yaş ortalamasını verirken, bazıları en düşük ve en yüksek yaş aralıklarını belirtmiştir (Örneğin, 21-25 yaş arasındaki katılımcılar). Bu nedenle, yaş bilgisi için veri tablosunda üç ayrı sütun açılmış ve bu sütunlar aracılığıyla her çalışmanın sunduğu biçime göre veri girilmiştir. Benzer şekilde, kayıt ortamı ve kullanılan ekipmana ilişkin bilgiler, çoğu çalışmada hiç yer almamış; yer alanlarda ise oldukça genel ifadeler kullanılmıştır. Ayrıca, Aydın ve Uzun'un (2020) da belirttiği gibi, formant ölçümlerine yönelik normalizasyon işlemi yalnızca sınırlı sayıda çalışmada yer almaktadır. Bu durum, normalleştirilmiş F1 ve F2 değerlerinin büyük bir kısmında eksik veri oluşmasına neden olmuştur.

Söz konusu eksiklikleri daha şeffaf biçimde ortaya koymak ve çözüm üretmek amacıyla veri yapısı esnek biçimde tasarlanmış ve veri kaybını en aza indirmek için bazı stratejiler uygulanmıştır. Örneğin, yaş bilgisi için

oluşturulan üçlü sütun yapısı, farklı raporlama biçimlerine uyum sağlama amacı gütmekle kalmamış; çözümlemenin kapsayıcılığını da artırmıştır. Bununla birlikte, örnekleme hızı, kayıt ortamı, kayıt ekipmanı, normalizasyon yöntemi ve normalleştirilmiş formant değerlerine ilişkin sütunlarda önemli ölçüde veri kaybı gözlemlenmiştir. Bu eksik veri durumu, aşağıda sunulan Şekil 5'te görselleştirilmiştir.



Şekil 5. Veri tablosundaki eksik veri

Yukarıda açıklanan süreçlerin sonucunda toplam 219 satırdan oluşan bir veri tablosu elde edilmiştir. Veri tablosundaki her satır, bir çalışmada raporlanan belirli bir ünlüye ilişkin formant ölçümünü temsil etmektedir. Bu veri yapısının sütun mimarisi Tablo 1'de sunulmuştur.

Tablo 1. Veri yapısı

Sütun adı	Tanımı
id	Çalışmanın kimliği
year	Çalışmanın yayımlandığı yıl
material	Kullanılan veri söyletim materyalinin türü
mat_share	Veri söyletim materyalinin paylaşım durumu
spcount	Konuşucu sayısı
spgender	Konuşucuların cinsiyet bilgisi
age_min	En düşük yaş
age_max	En yüksek yaş
age_mean	Ortalama yaş
env	Kayıt ortamı
equipment	Kayıt sırasında kullanılan ekipman
sample_khz	kHz cinsinden örnekleme hızı
norm	Normalizasyon işleminin uygulanıp uygulanmadığı
norm_method	Normalizasyon yapıldıysa kullanılan yöntem
f1nonr	Normalleştirilmemiş F1 değeri
f2nonr	Normalleştirilmemiş F2 değeri
f1nor	Normalleştirilmiş F1 değeri
f2nor	Normalleştirilmiş F2 değeri

Elde edilen bu veri kümesi, araştırmacılarla şeffaf ve yeniden üretilebilir biçimde paylaşılacak üzere OSF (Open Science Framework, osf.io) platformu üzerinden açık erişime sunulmuştur (Kaya ve Yağlı, 2025).⁵

4.2 Veri inceleme yöntemi

Bu çalışmada, meta-analiz kapsamında oluşturulan veri tablosunun istatistiksel çözümlemesi R programlama dili (R Core Team, 2021) kullanılarak gerçekleştirilmiştir. Veri çözümlemesinde karma etkiler modeli (İng. mixed-effects model) benimsenmiş ve bu modellerin uygulanmasında *lme4* (Bates vd., 2015) ve *lmerTest* (Kuznetsova vd., 2017) paketlerinden yararlanılmıştır.

Karma etkiler modelinin benimsenmesinin arka planında, veri yapısına ilişkin bazı temel gözlemler bulunmaktadır. Bu gözlemlerden ilki, ünlü verisi tablosunun karmaşık ve hiyerarşik bir yapıya sahip olmasıdır. İncelenen çalışmalar, farklı örneklem büyüklükleriyle ve değişen yöntembilimsel yaklaşımlarıyla veri üretmiş; buna bağlı olarak her bir çalışmada raporlanan ünlü formant değerleri hem çalışmalar arasında hem de çalışmaların kendi içindeki ünlü türleri arasında sistematik farklılıklar gösterebilecek potansiyele sahiptir. Ayrıca, aynı çalışmaya ait birden fazla gözlem satırı bulunduğundan, gözlemler arasında olası bağımlılık ilişkileri söz konusudur. Bu nedenle, veriye uygulanacak modelin bu çok düzeyli yapıyı dikkate alması gerekliliği ortaya çıkmıştır.

Bu doğrultuda geliştirilen karma etkiler modelinde iki düzeyde rastgele etki tanımlanmıştır: (1 | Çalışma) ve (1 | Çalışma:Ünlü). İlk olarak, her çalışmanın kendi örnekleme dayalı olarak veri sunduğu dikkate alınarak, çalışmalar rastgele etki olarak modele dahil edilmiştir. İkinci olarak ise, çalışmaların kendi içlerinde ünlüleri raporlama biçimlerinde farklılık gösterebileceği öngörüsüne dayanarak, çalışma ve ünlü etkileşimi (Örneğin, bir çalışmanın bazı ünlüleri raporlayıp diğerlerini dışarıda bırakması ya da olağanın dışında sapma ile raporlaması olasılığı) ayrı bir rastgele etki olarak modele eklenmiştir.

Modeldeki sabit etkiler ise çalışmalarda raporlanan temel yöntembilimsel değişkenler üzerinden yapılandırılmıştır. Bu değişkenler arasında çalışmanın yılı, konuşucu cinsiyeti, kullanılan söyletim materyali, kayıt ortamı, kayıt ekipmanı, ortalama yaş, en düşük yaş, en yüksek yaş ve örnekleme hızı yer almaktadır.

Analiz sürecinde, geliştirilen karma etkiler modeli her bir ünlü ses için ayrı ayrı uygulanmıştır. Bu tercihin arkasında, ünlülerin üretimsel özelliklerine ilişkin bazı gözlemler yer almaktadır. Her ünlü ses, farklı

⁵ OSF, bilimsel araştırma süreçlerinde şeffaflığı, açıklığı ve yeniden üretilebilirliği desteklemek amacıyla geliştirilen çevrim içi bir açık bilim platformudur. Araştırmacılar bu platform aracılığıyla veri kümelerini, analiz kodlarını, görselleri ve diğer araştırma çıktılarını açık erişimli biçimde paylaşabilir, böylece bilimsel işbirliğini ve yeniden kullanım olanaklarını artırabilirler. Türkiye bağlamında yürütülen dilbilim araştırmalarında açık bilime yönelik kavramsal çerçeveyi tartışan ve ön kayıt, veri paylaşımı ve yeniden üretilebilir iş akışları gibi temel uygulamalara dikkat çeken bir çalışma için bakınız Ataman, Çağlar ve Kırkıcı (2021).

eklemleyici pozisyonları, hava akımı düzenlemeleri ve rezonans özelliklerine sahiptir. Bu fizyolojik ve akustik farklılıklar, tüm ünlülerin tek bir modelle birlikte analiz edilmesini istatistiksel olarak sorunlu hale getirebilmekte ve modelin açıklayıcılığını sınırlayabilmektedir. Bu nedenle, daha düşük hata oranı ve daha yüksek açıklayıcılık düzeyi elde edebilmek adına, modelleme süreci ses türü düzeyinde gerçekleştirilmiştir. Bu bağlamda kullanılan model mimarisi (1)'de sunulmuştur.

(1) Analizde kullanılan modeller⁶

- (a) $F1_{Nonr} \sim Yıl + Cinsiyet + Materyal + Ortam + Araç + Ortalama\ yaş + En\ düşük\ yaş + En\ yüksek\ yaş + Örneklem\ hızı (1 | Çalışma) + (1 | Çalışma:Ünlü)$
- (b) $F2_{Nonr} \sim Yıl + Cinsiyet + Materyal + Ortam + Araç + Ortalama\ yaş + En\ düşük\ yaş + En\ yüksek\ yaş + Örneklem\ hızı (1 | Çalışma) + (1 | Çalışma:Ünlü)$

(1)'de sunulan model yapısı, çalışmalarda gözlemlenen yöntemsel çeşitliliğin formant değerleri üzerindeki etkilerini değerlendirmeye olanak tanımaktadır. Bu değerlendirme, ölçünlü Türkçeye ilişkin formant verilerini yöntemsel bağlamlarıyla birlikte ele alan sistematik bir meta-analiz yaklaşımıyla gerçekleştirilmiştir. Böylece, alanyazındaki bulguların karşılaştırmalı biçimde yorumlanması mümkün hale gelmiştir.

5. Bulgular ve Tartışma

Bu bölümde, ölçünlü Türkçenin ünlü formantlarına yönelik gerçekleştirilen çalışmalara dayalı olarak yürütülen meta-analizin bulguları sunulmakta ve bu bulgular çerçevesinde ünlü formant araştırmalarının yöntembilimsel nitelikleri tartışılmaktadır. Meta-analiz kapsamında her bir ünlü ses için ayrı ayrı geliştirilen modeller elde edilmiştir. Bu modellere ilişkin en temel gözlem, modellerdeki rastgele etkilerin sabit etkilere kıyasla daha baskın bir açıklayıcılığa sahip olmasıdır. Bu durum, formant verilerine ilişkin çalışmalarda araştırmacılar tarafından açıkça raporlanmayan ancak modele yansıyan örtük yöntembilimsel farklılıkların varlığına işaret etmektedir.

Özellikle her çalışmanın farklı veri toplama süreçleri, konuşucu profilleri ve kayıt koşulları gibi yöntemsel ayrımlar içermesi, modellerde rastgele etkilerin yüksek varyans açıklayıcılığına sahip olmasına neden olmuştur. Ayrıca, verinin her bir çalışmada belirli bir ünlüye ilişkin olarak elde edilmesi sırasında kullanılan söyletim materyali, kayıt ortamı ve konuşucu yaşı gibi değişkenlerin bu etkiyi artırdığı düşünülmektedir. Bu gözlemleri sınamak

⁶ Modelin R ortamı ile uyumlu mimarisi şu şekilde sunulabilir:

$F1_{nonr} \sim year + spgendor + material + env + equipment + age_mean + age_min + age_max + samplekhz + (1 | study) + (1 | study:vowel)$

$F2_{nonr} \sim year + spgendor + material + env + equipment + age_mean + age_min + age_max + samplekhz + (1 | study) + (1 | study:vowel)$

amacıyla, rastgele etkilerin modele anlamlı katkı yapıp yapmadığını belirlemek için *olabilirlik oranı testi* (İng. likelihood ratio test, LRT) uygulanmıştır. Bu test, hem çalışmalar arası değişkenliğin (1 | Çalışma) hem de çalışmalar içindeki farklı ünlülerin (1 | Çalışma:Ünlü) ayrı ayrı değerlendirilebilmesini sağlamıştır.

Aşağıda sunulan Tablo 2 ve 3, [e] ünlüsü için geliştirilen modelin çıktıları vermektedir.

Tablo 2. [e] ünlüsüne ilişkin geliştirilen modelin çıktıları: Rastgele etkiler

Rastgele Etkiler	Varyans (F1)	Standart Sapma (F1)	Varyans (F2)	Standart Sapma (F2)	Toplam Varyansa Katkı (F1)	Toplam Varyansa Katkı (F2)
Çalışma	4.32	2.08	9.42	3.07	%23	%23
Çalışma: Ünlü	6.65	2.58	17.58	4.19	%35	%43
Artık	7.99	2.83	14.02	3.74	%42	%34

Tablo 3. [e] ünlüsüne ilişkin geliştirilen modelin çıktıları: Sabit etkiler

Sabit Etkiler	Tahmin (F1)	Standart Hata (F1)	Serbestlik Derecesi (F1)	t Değeri (F1)	p Değeri (F1)	Tahmin (F2)	Standart Hata (F2)	Serbestlik Derecesi (F2)	t Değeri (F2)	p Değeri (F2)
Intercept	477.8	162800.0	0.002	0.003	0.99987	1938	224100.0	0.0175	0.009	0.999
Yıl	-0.000000019	79.73	0.002	0.000	1.00000	-0.0000003879	109.7	0.0175	0.000	1.000
Cinsiyet (Karma)	64.69	2553.0	0.002	0.025	0.99891	175.8	3513.0	0.0175	0.050	0.994
Cinsiyet (Kadın)	74.38	14.78	7.0	5.033	0.00151	334.6	19.51	7.0	17.152	0.001
Materyal (Paragraf)	-114.5	2277.0	0.002	-0.050	0.99805	141.2	3134.0	0.0175	0.045	0.994
Materyal (Ses)	4.70	27.65	7.0	0.170	0.86984	-1.85	36.50	7.0	-0.051	0.961
Materyal (Sözcük)	-126.0	3789.0	0.002	-0.033	0.99862	-53.66	5214.0	0.0175	-0.010	0.999
Kayıt ortamı (Stüdyo)	180.5	5465.0	0.002	0.033	0.99862	34.91	7522.0	0.0175	0.005	0.999
Kayıt ortamı (Oda)	97.30	2474.0	0.002	0.039	0.99840	92.60	3405.0	0.0175	0.027	0.996
Ekipman (Mikrofon)	102.3	1758.0	0.002	0.058	0.99783	65.6	2419.0	0.0175	0.027	0.996

Tablo 3 ve 4 takip edildiğinde, toplam varyansın büyük bir bölümünün rastgele etkiler tarafından açıklandığı görülmektedir. Örneğin, [e] ünlüsüne ilişkin F1 değerlerindeki değişkenliğin %35'i çalışma-ünlü etkileşiminden, %23'ü ise farklı çalışmalar arasındaki yöntembilimsel farklardan kaynaklanmaktadır ($p < .05$). Buna karşın sabit etkilerin varyans üzerindeki etkisi daha sınırlıdır; örneğin [e] için yalnızca cinsiyet değişkeninin anlamlı bir sabit etki gösterdiği görülmektedir ($p < .05$). Bu durum, formant değerleriyle ilgili verilerin yalnızca sabit değişkenlerle değil, aynı zamanda çalışmalara özgü daha derin yapısal farklılıklarla da şekillendiğini ortaya koymaktadır. Bu çalışmada diğer ünlüler için geliştirilen modellerin rastgele etki çıktıları da hemen hemen benzer eğilimler göstermektedir.⁷

Bu bulgular, bir önceki bölümde de belirtildiği üzere, veri tablosunun oluşturulması sürecinde yapılan ayrıntılı yöntembilimsel okuma sırasında da gözlemlenen duruma paraleldir. Çalışmalarda belirli bir yöntembilimsel ölçütün izlenmemesi, verilerin ortak bir zeminde değerlendirilmesini büyük ölçüde güçleştirmiştir. Bu yöntembilimsel dağınıklık, veri tablosunda eksik verilerin oluşmasına yol açmış ve formant ölçümlerinin karşılaştırılabilirliğini sınırlamıştır.

Aslında yalnızca ölçünlü Türkçeye odaklanan 16 çalışmanın yöntembilimsel ölçütlerle oluşturulduğu genel tablo değil, tek bir çalışmada ele alınan bir ünlü sesin verisi dahi kendi içinde oldukça hiyerarşik ve karmaşık bir yapıya sahiptir. Çalışmalar çoğu zaman yaş, cinsiyet, söyletim materyali ve kayıt ortamı gibi temel bilgileri sunsa da bu bilgilerin ardında yatan daha derin etkileyen faktörlerin çoğu genellikle raporlanmamıştır. Örneğin, konuşucunun kimliği, dilsel tutumu ve kullandığı değişke gibi toplumsal arka plan özellikleri ya da ünlünün öncesindeki ve sonrasındaki sesler gibi sesbilgisel bağlamı, formant değerlerini ciddi biçimde etkileyebilecek unsurlardır. Bir çalışmada söyletim materyali olarak yalnızca sözcüklerin belirlenip şeffaf biçimde sunulması, yöntembilimsel raporlama açısından olumlu bir yaklaşım olsa da bu sözcüklerdeki ünlülerin eşsesletimden nasıl etkilendiği gibi kritik etmenlerin göz ardı edilmesi riskini ortadan kaldırmaz.

Bu noktada tartışmaya açık olan temel soru şudur: Türkçenin ünlü formantlarını betimleyen çalışmalarda hangi yöntembilimsel ölçütler raporlanmamış olabilir? Bu soruya verilecek yanıt, yalnızca eksik bilgilerin saptanması açısından değil, aynı zamanda alandaki betimlemelerin karşılaştırılabilirliğini etkileyen yapısal sorunların anlaşılması açısından da önemlidir. Bu bölümün takip eden satırlarında, Türkçenin ünlü formantları alanyazınında sıklıkla göz ardı edilen ya da yetersiz biçimde raporlanan bazı temel yöntembilimsel etkenler sunulmaktadır.

Formant verilerindeki farklılaşmanın önemli bir bölümü, çoğu çalışmada açıkça belirtilmeyen fakat ölçümleri doğrudan etkileyen yöntemsel değişkenlerle ilişkilendirilebilir. Bunların başında kayıt koşulları gelmektedir. Stüdyo mikrofonları, bilgisayar mikrofonları, taşınabilir ses kayıt cihazları ya

⁷ Çalışmada ele alınan diğer ünlülere yönelik rastgele etkilere ilişkin tablolar çalışma ekinde sunulmuştur.

da cep telefonu mikrofonları gibi kayıt sürecinde kullanılan ekipman türleri arasında formant değerlerinde anlamlı farklar oluşabilmektedir (Fahed vd., 2022; Zhang, Jepson ve Chuang, 2024). Ölçümün gerçekleştirildiği donanımın teknik kapasitesi, ses sinyalinin kalitesini doğrudan etkileyerek akustik verilerin güvenilirliğini belirlemektedir (Jannets vd., 2019). Bu bağlamda, laboratuvar ortamında profesyonel ekipmanlarla elde edilen verileri ile uzaktan veri toplama araçları kullanılarak toplanan veriler arasında formant düzeyinde farklılıkların gözlemlenmesi şaşırtıcı değildir (Chai ve Garellek, 2022; Conklin, 2023).

Kayıt ekipmanının yanı sıra, konuşucu seçimi de formant verilerinin doğasını etkileyen önemli bir değişkendir. Konuşucuların yaş ve konuşma biçimi gibi özellikler, ünlü üretimi üzerinde sistematik etkiler yaratmaktadır. Özellikle yaş değişkeni, bölgesel değişkenlik çalışmalarında formant yapısını doğrudan etkilemektedir (Clopper, Pisoni ve De Jong, 2005; Jacewicz, Fox ve Salmons, 2011). Peterson ve Barney'in (1952) ölçünlü Amerikan İngilizcesi üzerine gerçekleştirdiği klasik çalışma bile, sonradan yapılan araştırmalarla aslında ciddi düzeyde bölgesel değişkenlik içerdiği gösterilen bir örneğe dönüşmüştür. Buna ek olarak, resmi ve gündelik konuşma gibi konuşma biçimindeki farklılıklar, ünlü üretim süresi ve formant değerleri üzerinde anlamlı etkilere sahiptir (Clopper, Pisoni ve De Jong, 2005; Nolan vd., 2009).

Söyletim yöntemi de ünlü formantlarının ölçümünü etkileyen kritik bir unsurdur. Sözcük listeleri, taşıyıcı tümceler, yalın sesler veya sürekli konuşmadan oluşan veri toplama bağlamları, elde edilen formant frekanslarını doğrudan biçimlendirmektedir (Kent ve Vorperian, 2018).⁸ Söyletim materyalinin içeriği, konuşucunun deney ortamına tepkisi ve eşsesletim gibi etmenler, verinin geçerliliğini etkileyen önemli parametrelerdir (Clopper, Pisoni ve De Jong, 2005). Bu nedenle, yalnızca sözcüğün yüzeysel biçimde sunulması, o sözcükteki ünlünün sesbilgisel bağlamının göz ardı edilmesine yol açabilmektedir.

Son olarak, verilerin incelenmesi ve işlenmesinde kullanılan yöntemler ve yazılımlar da formant değerlerinde farklılık yaratabilmektedir. Aynı akustik veri, kullanılan analiz aracına bağlı olarak farklı biçimlerde ölçülebilmektedir (Burris vd., 2014). Ayrıca, formant analizinde benimsenen çözümleme parametreleri, elde edilen sonuçların güvenilirliğini ve karşılaştırılabilirliğini doğrudan etkilemektedir (Kendall ve Vaughn, 2015, 2020). Bu bağlamda, hem veri üretim süreçlerinde hem de analiz aşamalarında benimsenen standartların belirginleştirilememesi, Türkçeye ilişkin formant çalışmalarında tutarlı ve karşılaştırılabilir sonuçlara ulaşılmasını güçleştirmektedir.

⁸ Taşıyıcı tümce, konuşma verisinin sabit bir sözdizimsel bağlam içinde elde edilmesini sağlamak amacıyla kullanılan sabit yapı ifadeleridir. Bu veri söyletim yönteminde genellikle hedef sözcük ya da ses bir boşluğa yerleştirilir (Örneğin, Ali ____ dedi.) Bu yöntem, seslerin üretimini etkileyebilecek bağlamsal (önceki/sonraki) ses dizimlerini kontrol altına alarak ölçümlerdeki değişkenliği azaltmayı amaçlar.

Sonuç olarak, bu bölümdeki inceleme aracılığıyla elde edilen bulgular, Türkçenin ünlü formantlarını betimleyen çalışmaların yöntembilimsel çeşitliliğinin yalnızca formant ölçümleri üzerindeki etkilerine değil, aynı zamanda alanyazında karşılaştırılabilirliğin sınırlarına da vurgu yapmaktadır. Rastgele etkilerin baskınlığı, araştırmacılar tarafından çoğu zaman açık biçimde raporlanmayan ancak veri üretim ve analiz süreçlerine için olan yöntemsel farklılıkların güçlü bir göstergesi niteliğindedir.

6. Sonuç

Bu meta-analiz, Türkçenin ünlü formantlarına ilişkin yürütülen çalışmaların sistematik bir çerçevede değerlendirilmesine olanak tanımış; böylece alanyazındaki yöntembilimsel çeşitliliğin ve dağınıklığın ölçülebilir etkilerini görünür kılmıştır. İncelenen çalışmalarda gözlemlenen örtük yöntemsel farklılıkların, geliştirilen modellerde sabit etkilerden çok rastgele etkiler aracılığıyla açıklanması, bu çeşitliliğin istatistiksel olarak anlamlı bir düzeyde modellendiğini ortaya koymuştur. Bu durum, yalnızca geçmiş çalışmalardaki örtük varsayımların açığa çıkarılmasına değil, aynı zamanda gelecekteki araştırmalar için yöntemsel bir uyarı işlevi görmesine de olanak sağlamaktadır.

Elde edilen bulgular, yöntembilimsel şeffaflığın sağlanmasının dilbilimsel verilerin karşılaştırılabilirliği açısından ne denli kritik olduğunu açıkça ortaya koymaktadır. Özellikle modellerdeki varyansın büyük ölçüde çalışmalara özgü rastgele etkilerle açıklanması, veri üretim sürecinde açık biçimde raporlanmayan ya da sistematikleştirilmeyen değişkenlerin önemli rol oynadığını göstermektedir. Nitekim çalışmalarda veri söyletim yöntemi, kayıt ortamı, kullanılan ekipman ve konuşucu özellikleri gibi temel değişkenler çoğu zaman rapor edilmekle birlikte, ölçüm yöntemleri, sesbilgisel bağlam ve konuşucuların toplumsal arka planı gibi daha derin düzeydeki etmenler sıklıkla göz ardı edilmektedir. Bu eksiklik, yalnızca tekil çalışmaların iç tutarlılığını değil, aynı zamanda çalışmalar arası karşılaştırmaların geçerliliğini de etkilemektedir. Böyle bir yöntembilimsel belirsizlik, alandaki bilgi birikiminin birikimli biçimde gelişmesini sınırlandırmakta ve farklı çalışmaların birbiri ile iletişim kurmasını da zorlaştırmaktadır.

Winter ve Grice'in (2021) de belirttiği gibi, dilbilim araştırmalarının yalnızca istatistiksel analiz tekniklerine odaklanması yeterli değildir; aynı zamanda tutarsızlıkların nasıl ele alınacağına ilişkin açık bir yöntem tasarımı da benimsenmelidir. Katılımcı demografisi, kayıt koşulları ve veri söyletim yöntemleri gibi temel değişkenlerde tutarsızlıklar barındıran çalışmaların bulguları, yalnızca o çalışmanın bağlamına özgü kalır ve daha geniş bir dilsel olguyu temsil etme iddiasından uzaklaşır. Bu çalışma, yalnızca formant ölçüm yöntemlerine ilişkin bir değerlendirme sunmakla kalmamakta, aynı zamanda Türkçede deneysel sesbilgisi ve sesbilim araştırmalarında ve dil değişkenliği çalışmaları için yeniden üretilebilir ve karşılaştırılabilir veri üretiminin önemini ortaya koymaktadır. Bu bağlamda, Türkçenin ünlü formantları gibi sesbilgisel ve sesbilimsel betimlemelerine ilişkin yapılacak gelecekteki çalışmalarda, yöntemsel şeffaflık, standartlaştırma ve yeniden üretilebilirlik

ilkelerinin öncelikli hedefler arasında yer alması gerektiği açıktır. Özellikle ortak bir ölçüm protokolü ve raporlama ölçütünün belirlenmesi hem veri niteliğini artıracak hem de farklı araştırmalar arasında doğrudan karşılaştırma yapılabilmesini olanaklı kılacaktır. Bunun yanı sıra, açık veri paylaşımı kültürünün güçlendirilmesi, alandaki bulguların doğrulanabilirliğini artıracak ve Türkçenin sesbilgisel özelliklerine ilişkin daha bütüncül ve erişilebilir bir bilgi zemini oluşturacaktır.

Kaynaklar

- Akbulut, M., & Erdem Nas, G. (2021). Türkçe öğrenen Boşnak konuşurlarının Türkçe ünlü üretimi. *Karadeniz Araştırmaları*, 71, 735-762.
- Arslan, M. N., & Uzun, İ. P. (2025). An acoustic analysis of surrounding vowel effects on intervocalic /h/ in Turkish. *Hacettepe Üniversitesi Edebiyat Fakültesi Dergisi*, 42(1), 248-261.
<https://doi.org/10.32600/huefd.1534383>
- Ataman, E., Çağlar, O. C., & Kırkıcı, B. (2021). Dilbilim araştırmalarında açık bilim. *Dilbilim Araştırmaları Dergisi*, 32(2), 149-175.
<https://doi.org/10.18492/dad.936072>
- Aydın, Ö., & Uzun, İ. P. (2020). Ünlü formlant normalizasyonu: R programlama dilinde bir uygulama. İ. P. Uzun (Yay. haz.), *Kuramsal ve uygulamalı sesbilim* içinde (ss. 297-322). Seçkin Yayıncılık.
- Bates, D., vd. (2015). Fitting Linear Mixed-Effects Models using lme4. *Journal of Statistical Software*, 67(1), 1-48.
<https://doi.org/10.18637/jss.v067.i01>
- Bickley, C. A. (1989). *Acoustic evidence for the development of speech*. Research Laboratory of Ethics (RLE) Technical Report No 548. MIT.
- Börtlü, G. (2020). *The vowel triangle of Turkish and phonological processes of laxing and fronting in Turkish* [Yüksek Lisans Tezi]. Hacettepe Üniversitesi, Sosyal Bilimler Enstitüsü, İngiliz Dilbilimi Anabilim Dalı. <http://hdl.handle.net/11655/22583>
- Burris, C., vd. (2014). Quantitative and descriptive comparison of four acoustic analysis systems: Vowel measurement. *Journal of Speech, Language, and Hearing Research*, 57(1), 26-45.
[https://doi.org/10.1044/1092-4388\(2013\)12-0103](https://doi.org/10.1044/1092-4388(2013)12-0103)
- Chai, Y., & Garellek, M. (2022). On H1-H2 as an acoustic measure of linguistic phonation type. *The Journal of the Acoustical Society of America*, 152, 1856-1870. <https://doi.org/10.1121/10.0014175>
- Chung, H., vd. (2012). Cross-linguistic studies of children's and adults' vowel spaces. *The Journal of the Acoustical Society of America*, 131, 442-454.
<https://doi.org/10.1121/1.3651823>
- Clopper, C. G. (2009). Computational methods for normalizing acoustic vowel data for talker differences. *Language and Linguistics Compass*, 3(6), 1430-1442.
<https://doi.org/https://doi.org/10.1111/j.1749818X.2009.00165.x>
- Clopper, C. G., Pisoni, D. B., & De Jong, K. (2005). Acoustic characteristics of the vowel systems of six regional varieties of American English. *The*

- Journal of the Acoustical Society of America*, 118, 1661-1676.
<https://doi.org/10.1121/1.2000774>
- Conklin, J. (2023). Examining recording quality from two methods of remote data collection in a study of vowel reduction. *Laboratory Phonology* 14(1). doi: <https://doi.org/10.16995/labphon.10544>
- Çiyiltepe, M., & Orberk, M. S. (2008). Konuşmacı tanıma/tanımlamada Türkiye Türkçesindeki ünlülerin formant analizi. M. Sarıca, N. Sarıca, & A. Karaca (Yay. haz.), *XXII. Dilbilim Kurultayı bildirileri* içinde (ss. 293-298). Cantekin Matbaası.
- Çiyiltepe, M., Bekar, İ. P., & Ergenç, İ. (2009). Changing formant values: Synthesis of four regions of Turkey. S. Ay vd. (Yay. haz.), *Essays on Turkish linguistics* içinde (ss. 3-9). Harrassowitz Verlag.
- Erdem Nas, G. (2020). Ünlüler ile ilgili kuramlar çerçevesinde Türkçedeki ünlüler. *RumeliDE Dil ve Edebiyat Araştırmaları Dergisi*, 7, 237-252.
<https://doi.org/10.29000/rumelide.813412>
- Erdem Nas, G., & Yalçınkaya, S. (2021). Endonezya dili konuşurlarının Türkçe ünlü üretimi (Bir R dili uygulaması). *Disiplinler Arası Dil Araştırmaları*, 2(2), 1-24.
- Ergenç, İ. (1989). *Türkiye Türkçesinin görevsel sesbilimi: Sesbirimlere genel bir bakış*. Engin Yayıncılık.
- Esen Aydın, F., Kulak Kayıkçı, M. E., & Suslu, N. (2019). Temporal and frequency characteristics of Turkish vowels in laryngectomized speakers: Preliminary study. *Medeniyet Medical Journal*, 34, 149-159.
<https://doi.org/10.5222/MMJ.2019.42744>
- Fahed, V. S., Doheny, E. P., Busse, M., Hoblyn, J., & Lowery, M. M. (2022). Comparison of acoustic voice features derived from mobile devices and studio microphone recordings. *Journal of Voice*. <http://doi.org/10.1016/j.jvoice.2022.10.006>
- Fant, G. (1960). *Acoustic theory of speech production*. Mouton.
- Fletcher, A. R. vd. (2015). The relationship between speech segment duration and vowel centralization in a group of older speakers. *The Journal of the Acoustical Society of America*, 138, 2132-2139.
<https://doi.org/10.1121/1.4930563>
- Fox, R. A., & Jacewicz, E. (2009). Cross-dialectal variation in formant dynamics of American English vowels. *The Journal of the Acoustical Society of America*, 126, 2603-2618. <https://doi.org/10.1121/1.3212921>
- Güven, M. (2015). Türkçedeki ünlülerin formant frekans değerleri ve enkliz “ses” olayı. *Turkish Studies*, 10(16), 689-712.
<https://doi.org/10.7827/TurkishStudies.8941>
- Hillenbrand, J. (2013). Static and dynamic approaches to vowel perception. G. S. Morrison & P. F. Assmann (Yay. haz.), *Vowel inherent spectral change* içinde (ss. 9-30). Springer.
- Hillenbrand, J., vd. (1995). Acoustic characteristics of American English vowels. *The Journal of the Acoustical Society of America*, 97, 3099-3111. <https://doi.org/10.1121/1.411872>

- Jacewicz, E., Fox, R. A., & Salmons, J. (2011). Vowel change across three age groups of speakers in three regional varieties of American English. *Journal of Phonetics*, 39(4), 683-693. <https://doi.org/10.1016/j.wocn.2011.07.003>
- Jannetts, S., vd. (2019). Assessing voice health using smartphones: Bias and random error of acoustic voice parameters captured by different smartphone types. *International Journal of Language & Communication Disorders*, 54, 292-305. <https://doi.org/10.1111/1460-6984.12457>
- Jin, S.-H., & Liu, C. (2013). The vowel inherent spectral change of English vowels spoken by native and non-native speakers. *The Journal of the Acoustical Society of America*, 133, EL363-EL369. <https://doi.org/10.1121/1.4798620>
- Jones, D. (1917). *Speech sounds: Cardinal vowels*. The Gramophone.
- Kaya, A. Ç., & Yağlı, E. (2025, 17 Nisan). *Türkçe ünlülerin meta-analizi* [Veri]. <https://doi.org/10.17605/OSF.IO/UA465>
- Kendall, T., & Vaughn, C. (2015). Measurement variability in vowel formant estimation: A simulation experiment. M. Wolters vd. (Yay. haz.), *Proceedings of the International Congress on Phonetics (ICPhS) 2015*. International Phonetic Association.
- Kendall, T. & Vaughn, C. (2020). Exploring vowel formant estimation through simulation-based techniques. *Linguistics Vanguard*, 6(s1), 20180060. <https://doi.org/10.1515/lingvan-2018-0060>
- Kent, R. D., & Vorperian, H. K. (2018). Static measurements of vowel formant frequencies and bandwidths: A review. *Journal of Communication Disorders*, 74, 74-97. <https://doi.org/10.1016/j.jcomdis.2018.05.004>
- Kılıç, M. A. (2003). Türkiye Türkçesindeki ünlülerin sesbilgisel özellikleri. D. Akar & S. Özsoy (Yay. haz.), *Proceedings of the 10th International Conference on Turkish Linguistics* içinde (ss. 3-18). Boğaziçi Üniversitesi Yayınları.
- Kılıç, M. A., & Ögüt, F. (2004). A high unrounded vowel in Turkish: Is it central or back vowel? *Speech Communication*, 43(1-2), 143-154. <https://doi.org/10.1016/j.specom.2004.03.001>
- Koçak, A. (2020). *Yabancı dil olarak Türkçe öğretiminde telaffuz: Bir karma yöntem araştırması ve ünlülerin telaffuz öğretimine yönelik yeni bir teknik* [Doktora Tezi]. Necmettin Erbakan Üniversitesi, Sosyal Bilimler Enstitüsü, Türk Dili ve Edebiyatı Anabilim Dalı.
- Kopkallı-Yavuz, H. (2010). The sound inventory of Turkish: Consonants and vowels. S. Topbaş & M. Yavaş (Yay. haz.), *Communication disorders in Turkish* içinde (ss. 27-47). Multilingual Matters.
- Korkmaz, Y., & Boyacı, A. (2018). Examining vowels' formant frequency shifts caused by preceding consonants for Turkish language. *Journal of Engineering and Technology*, 2(2), 38-47.
- Korkmaz, Y., & Boyacı, A. (2022). A comprehensive Turkish accent/dialect recognition system using acoustic perceptual formants. *Applied Acoustics*, 193, 108761.

- <https://doi.org/10.1016/j.apacoust.2022.108761>
- Korkmaz, Y., Boyacı, A., & Tuncer, T. (2019). Turkish vowel classification based on acoustical and decompositional features optimized by genetic algorithm. *Applied Acoustics*, 154, 28-35.
<https://doi.org/10.1016/j.apacoust.2019.04.027>
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2017). lmerTest package: Tests in Linear Mixed Effects Models. *Journal of Statistical Software*, 82(13), 1-26. <https://doi.org/10.18637/jss.v082.i13>
- Labov, W., Ash, S., & Boberg, C. (2006). *The atlas of North American English: Phonetics, phonology, and sound change: A multimedia reference tool*. Mouton de Gruyter.
- Ladefoged, P. (1967). *Three areas of experimental phonetics*. Cambridge University Press.
- Lee, S., Potamianos, A., & Narayanan, S. (1999). Acoustics of children's speech: Developmental changes of temporal and spectral parameters. *The Journal of the Acoustical Society of America*, 105, 1455-1468.
<https://doi.org/10.1121/1.426686>
- Lindblom, B. (1963). Spectrographic study of vowel reduction. *The Journal of the Acoustical Society of America*, 35, 1773-1781.
<https://doi.org/10.1121/1.1918816>
- Loakes, D. (2006). *A forensic phonetic investigation into the speech patterns of identical and non-identical twins* [Doktora Tezi]. School of Languages and Literatures, The University of Melbourne.
- Lobanov, B. M. (1971). Classification of Russian vowels spoken by different speakers. *The Journal of the Acoustical Society of America*, 49, 606-608. <https://doi.org/10.1121/1.1912396>
- Malkoç, E. (2009). Türkçe ünlü formant frekans değerleri ve bu değerlere dayalı ünlü dördgeni. *Dil Dergisi*, 146, 71-85.
https://doi.org/10.1501/Dilder_00000000122
- Malkoç, E. (2011). *Konuşmanın kişiye özgü değişmezleri: Ünlüler üzerine sesbilgisel bir çalışma* [Doktora Tezi]. Ankara Üniversitesi, Sosyal Bilimler Enstitüsü.
- Maurer, D. (2024). *Acoustics of the vowel: Indices*. Peter Lang.
- Morrison, G. S., & Nearey, T. M. (2007). Testing theories of vowel inherent spectral change. *The Journal of the Acoustical Society of America*, 122, EL15-EL22. <https://doi.org/10.1121/1.2739111>
- Nearey, T. M. (1978). *Phonetic feature systems for vowels*. Indiana University Linguistics Club.
- Nearey, T. M. (1989). Static, dynamic, and relational properties in vowel perception. *The Journal of the Acoustical Society of America*, 85, 2088-2113. <https://doi.org/10.1121/1.397861>
- Nearey, T. M., & Assmann, P. F. (1986). Modeling the role of inherent spectral change in vowel identification. *The Journal of the Acoustical Society of America*, 80, 1297-1308. <https://doi.org/10.1121/1.394433>

- Nolan, F., & Oh, T. (1996). Identical twins, different voices. *The International Journal of Speech, Language and the Law*, 3(1), 39-49. <https://doi.org/10.1558/ijssl.v3i1.39>
- Nolan, F., vd. (2009). The DyViS database: Style-controlled recordings of 100 homogenous speakers for forensic phonetic research. *The International Journal of Speech, Language and the Law*, 16(1), 31-57. <https://doi.org/10.1558/ijssl.v16i1.31>
- Peterson, G. E., & Barney, H. L. (1952). Control methods in a study of the vowels. *The Journal of the Acoustical Society of America*, 24, 175-184. <https://doi.org/10.1121/1.1906875>
- Pfützing, H. R., & Niebuhr, O. (2011). Historical development of phonetic vowel systems: The last 400 years. *ICPhS XVII*, 160-163.
- R Core Team. (2021). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing.
- Selen, N. (1979). *Söyleyiş sesbilimi, akustik sesbilim ve Türkiye Türkçesi*. TDK Yayınları.
- Stevens, K. N. (1998). *Acoustic phonetics*. MIT Press.
- Strange, W. (1989). Evolving theories of vowel perception. *The Journal of the Acoustical Society of America*, 85, 2081-2087. <https://doi.org/10.1121/1.397860>
- Studdert-Kennedy, M. (1964). The perception of speech. In T. A. Sebeok (Yay. haz.), *Current trends in linguistic* içinde (pp. 2349-2385). Mouton.
- Trautmann, H. (1990). Analytical expressions for the tonotopic sensory scale. *Journal of the Acoustic Society of America*, 88(1), 97-100. <https://doi.org/10.1121/1.399849>
- Turner, G. S., Tjaden, K., & Weismer, G. (1995). The influence of speaking rate on vowel space and speech intelligibility for individuals with amyotrophic lateral sclerosis. *Journal of Speech, Language, and Hearing Research*, 38, 1001-1013. <https://doi.org/10.1044/jshr.3805.1001>
- Türk, O., vd. (2004). Türkçede ünlülerin formant incelemesi. İ. Ergenç, vd. (Yay. haz.), *Dilbilim incelemeleri* içinde (ss. 67-76). Doğan Yayıncılık.
- Wickham, H. (2016). *ggplot2: Elegant graphics for data analysis* [R Paketi]. Springer-Verlag New York. <https://ggplot2.tidyverse.org>
- Winter, B., & Grice, M. (2021). Independence and generalizability in linguistics. *Linguistics*, 59(5), 1251-1277. <https://doi.org/10.1515/ling-2019-0049>
- Yılmaz Davutoğlu, A. (2010). *Standart Türkçedeki ünlülerin akustik analizi ve fonetik altyapılar*. [Sanatta Yeterlilik Tezi]. İstanbul Üniversitesi Sosyal Bilimler Enstitüsü, Tiyatro Sanat Dalı.
- Zhang, C., Jepson, K. & Chuang, Y., (2024) Investigating differences in lab-quality and remote recording methods with dynamic acoustic measures. *Laboratory Phonology* 15(1). <https://doi.org/10.16995/labphon.10492>

Ek

Bu çalışma kapsamında geliştirilen modellerde rastgele etkilerin model geneline anlamlı katkı yapıp yapmadığına yönelik uygulanan olabilirlik oranı testini içeren tablolar aşağıda sunulmuştur.

Tablo 1. [i] ünlüsüne ilişkin geliştirilen modelin çıktıları: Rastgele etkiler

Rastgele etkiler	Varyans (F1)	Standart Sapma (F1)	Varyans (F2)	Standart Sapma (F2)	Toplam Varyansa Katkı (F1)	Toplam Varyansa Katkı (F2)	p LRT
Çalışma	16.20	4.02	21.80	4.67	%16	%22	< .05
Çalışma:Ünlü	52.04	7.21	41.97	6.48	%52	%42	< .05
Artık	31.76	5.64	36.23	6.02	%32	%36	-

Tablo 2. [œ] ünlüsüne ilişkin geliştirilen modelin çıktıları: Rastgele etkiler

Rastgele etkiler	Varyans (F1)	Standart Sapma (F1)	Varyans (F2)	Standart Sapma (F2)	Toplam Varyansa Katkı (F1)	Toplam Varyansa Katkı (F2)	p LRT
Çalışma	0.00	0.00	14.00	3.74	%9	%39	> .05
Çalışma:Ünlü	12.00	3.46	21.20	4.61	%61	%59	< .05
Artık	6.10	2.47	0.72	0.85	%30	%2	-

Tablo 3. [y] ünlüsüne ilişkin geliştirilen modelin çıktıları: Rastgele etkiler

Rastgele etkiler	Varyans (F1)	Standart Sapma (F1)	Varyans (F2)	Standart Sapma (F2)	Toplam Varyansa Katkı (F1)	Toplam Varyansa Katkı (F2)	p LRT
Çalışma	15.00	3.87	10.00	3.16	%79	%24	< .05
Çalışma:Ünlü	2.50	1.58	16.50	4.06	%13	%40	< .05
Artık	1.50	1.22	14.50	3.81	%8	%36	-

Tablo 4. [a] ünlüsüne ilişkin geliştirilen modelin çıktıları: Rastgele etkiler

Rastgele etkiler	Varyans (F1)	Standart Sapma (F1)	Varyans (F2)	Standart Sapma (F2)	Toplam Varyansa Katkı (F1)	Toplam Varyansa Katkı (F2)	p LRT
Çalışma	4.00	2.00	8.00	2.83	%20	%25	< .05
Çalışma:Ünlü	11.20	3.35	11.90	3.45	%56	%37	< .05
Artık	4.80	2.19	12.10	3.48	%24	%38	-

Tablo 5. [u] ünlüsüne ilişkin geliştirilen modelin çıktıları: Rastgele etkiler

Rastgele etkiler	Varyans (F1)	Standart Sapma (F1)	Varyans (F2)	Standart Sapma (F2)	Toplam Varyansa Katkı (F1)	Toplam Varyansa Katkı (F2)	<i>p</i> LRT
Çalışma	2.00	1.41	8.00	2.83	%10	%25	< .05
Çalışma:Ünlü	15.00	3.87	11.40	3.38	%75	%36	< .05
Artık	3.00	1.73	12.60	3.55	%15	%39	-

Tablo 6. [o] ünlüsüne ilişkin geliştirilen modelin çıktıları: Rastgele etkiler

Rastgele etkiler	Varyans (F1)	Standart Sapma (F1)	Varyans (F2)	Standart Sapma (F2)	Toplam Varyansa Katkı (F1)	Toplam Varyansa Katkı (F2)	<i>p</i> LRT
Çalışma	0.50	0.71	12.00	3.46	%1	%24	> .05
Çalışma:Ünlü	28.00	5.29	19.00	4.36	%56	%38	< .05
Artık	21.50	4.64	19.00	4.36	%43	%38	-

Tablo 7. [u] ünlüsüne ilişkin geliştirilen modelin çıktıları: Rastgele etkiler

Rastgele etkiler	Varyans (F1)	Standart Sapma (F1)	Varyans (F2)	Standart Sapma (F2)	Toplam Varyansa Katkı (F1)	Toplam Varyansa Katkı (F2)	<i>p</i> LRT
Çalışma	13.00	3.61	12.00	3.46	%26	%24	< .05
Çalışma:Ünlü	36.00	6.00	17.00	4.21	%72	%38	< .05
Artık	1.00	1.00	19.00	4.36	%2	%38	-

FONETİK–FONOLOJİK SÜREKLİLİK VE ÜNLÜ DÖRTGENİ BAĞLAMINDA TÜRKÇEDEKİ /a/ ÜNLÜSÜNÜN KAVRAMSAL KONUMU¹

Mehmet Akif KILIÇ

Prof., Sağlık Bilimleri Üniversitesi, makilic@yahoo.com, ORCID: 0000-0003-4482-3490

Kılıç, M.A. (2025). Fonetik–fonolojik süreklilik ve ünlü dörtgeni bağlamında Türkçedeki /a/ ünlüsünün kavramsal konumu. İçinde O. Çınar, F. Başbuğ & H. Aydemir (Ed.), *Çağdaş dilbilim yazıları I: İMÜ Dilbilimi 15. yıl armağanı* (ss. 193-199) Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X009>

ÖZ

Bu çalışma, Türkçedeki /a/ ünlüsünün ünlü dörtgeni üzerindeki konumunu belirlemeyi ve sesbirim düzeyinde hangi Uluslararası Fonetik Alfabe (UFA) karakteriyle gösterilmesinin en uygun olacağını ortaya koymayı amaçlamaktadır. Ayrıca, /a/ ünlüsünün Türkçedeki alt sesbirim varyasyonları örneklenerek, fonetik ve fonolojik düzlemler arasındaki sınırın kuramsal olarak nasıl daha tutarlı biçimde tanımlanabileceği tartışılmaktadır. Çalışmada, artikülatuvar temelli üçgen modelden akustik dörtgen modele geçişin tarihsel arka planı özetlenmekte; fonetik temsillerin fonolojik açıklamalara yanlış biçimde aktarılmasının doğurduğu kavramsal karışıklık incelenmektedir. Fonetik–fonolojik süreklilik (*continuum*) kavramı, bu iki düzlemin sıklıkla birbirine karıştırılmasına yol açan temel etmenlerden biri olarak değerlendirilmektedir. Sesbirimlerin fonolojik özelliklerinin fonetik veriler üzerinden yorumlanması, bu kavramsal karışıklığın en belirgin sonucudur. Oysa fonolojik özellikler, fiziksel ölçümlerle değil, soyutlama yoluyla belirlenebilir. Bu bağlamda, algısal fonetik araştırmalar ve fonolojideki çağdaş bir yaklaşım olan Element Teorisi, sesbirimlerin fonolojik kategorilerinin daha tutarlı biçimde tanımlanmasına katkı sağlayabilir. Sonuç olarak, Türkçedeki /a/ ünlüsünün temel fonolojik özelliği açıklık (genişlik) olup, arka–orta–ön konumu fonetik düzleme aittir; bu nedenle sesbirim düzeyinde /a/ biçimiyle temsil edilmesi gerekir.

Anahtar Sözcükler: Fonetik, Fonoloji, Ünlü dörtgeni, Türkçe

¹ Okumayı kolaylaştırmak için metin içinde Uluslararası Fonetik Alfabe (UFA) kodları kullanılmamıştır. Dolayısıyla, “Türkçe /a/ ünlüsü”, “/a/ sesbirimi”, “/a/ ünlüsünün ince alt sesbirimi”, “/a/ ünlüsünün kalın alt sesbirimi” gibi ifadelerde geçen a karakteri UFA kodu değildir. Gerekli olduğu durumlarda kodlar metin akışından ayrılarak ayrı bir satırda gösterilmiştir.

ABSTRACT

This study aims to determine the position of the Turkish vowel /a/ within the vowel quadrilateral and to identify the most appropriate International Phonetic Alphabet (IPA) symbol for representing it at the phonemic level. In addition, by exemplifying the sub-phonemic variations of /a/ in Turkish, the paper discusses how the boundary between the phonetic and phonological levels can be more accurately defined from a theoretical perspective. The historical background of the transition from the articulatory triangle model to the acoustic quadrilateral model is summarized, and the conceptual confusion arising from the misinterpretation of phonetic representations as phonological explanations is examined. The concept of the phonetic–phonological continuum is addressed as a key factor underlying the frequent conflation of these two domains. Interpreting the phonological properties of phonemes solely on the basis of phonetic data represents a major consequence of this conceptual ambiguity. However, phonological properties should be established not through physical measurement but through abstraction. In this regard, perceptual phonetic studies and the contemporary phonological framework known as Element Theory may offer a more consistent means of defining phonological categories. In conclusion, the essential phonological property of the Turkish /a/ vowel is its openness (vowel height), while its backness or frontness belongs to the phonetic plane; therefore, it should be represented at the phonemic level with the symbol /a/.

Keywords: Phonetics, Phonology, Vowel quadrilateral, Turkish

1. Giriş

Fonetik ve fonoloji birbirinden ayrı düzlemler olarak tanımlansa da fonetik–fonolojik süreklilik (*continuum*) kavramı bu iki alan arasındaki sınırın sanıldığı kadar keskin değil, bulanık olduğunu ima eder. Bu bulanıklık, iki alanın uygulamada sıkça birbirine karıştırılmasına neden olur.

Bir kavramın sözlükteki tanımını bilmek, onun temsil ettiği düşünsel alanı kavramaya yetmez; uygulamadaki tutum, kavramı içselleştiremeyen kişiyi hemen ele verir.

Bu yazı, temeli fonetik–fonoloji ilişkisine dair zihinsel bulanıklıktan kaynaklanan bir sorunu, Türkçedeki /a/ ünlüsünün ünlü dörtgeni üzerinde grafiksel temsili kavramsal açıdan tartışmayı amaçlamaktadır.

Konunun detayına girmeden farklı dilbilimsel ifade yöntemlerini ve bu amaçla kullanılan kod sistemini hatırlatmak yararlı olacaktır.

Ortografik gösterim: Bir dilsel birimin yazıdaki biçimi, her dil için geçerli olmasa da Türkçede harf olarak düşünülebilir, iki açılı parantez arasında gösterilir: {a}

Sesbirim (fonem): Dilde anlam ayırt eden en küçük birim, iki eğik çizgi arasında gösterilir. Kılıç (2017), yine iki eğik çizgi arasında gösterilen fonolojik alt sesbirimden ayırmak için üst karakter olarak “S” (sesbirim kelimesinin ilk karakteri) eklemeyi önermektedir: ² /a/ veya /a/^S

² Burada “a” karakteri işaretli (*unmarked*) olduğu yani rutin olarak kullanılan şekil olduğu için tercih edilmiştir.

Fonolojik alt sesbirim: Bir sesbirimin belirli bir dildeki varyantları; fonolojik alt sesbirimler konuşma şekliyle değil dilin kurallarıyla belirlenir. Sesbirimler gibi iki eğik çizgi arasında gösterilir. Sesbirimlerden ayırmak gerektiğinde üst karakter olarak “DA” eklenebilir: /a/, /a/ veya /a/^{DA}, /a/^{DA}

Fonetik alt sesbirim: Konuşmacıdan, komşu konuşma seslerinden (eşsöyleyiş/koartikülasyon) veya hız gibi konuşmanın farklı özelliklerden kaynaklanan, dilbilimsel açıdan önemi olmayan varyantlar, iki köşeli parantez arasında gösterilir; gerektiğinde değiştirme işaretleri de kullanılabilir: [a], [a], [ʌ], [ä], [ɑ:]

2. Ünlü Dörtgeninin Ortaya Çıkışı ve Kullanım Amacı

Ünlü dörtgeninin kökeni 18. yüzyıla kadar uzanır. Alman dilbilimci Johann Christian Hellwag (1781), ünlüleri dilin ön–arka ve yüksek–alçak konumlarına göre artikülatuvar bir üçgen biçiminde düzenlemiştir. 20. yüzyılın başlarında akustik fonetiğin gelişmesiyle birlikte, ünlülerin akustik özelliklerini daha iyi yansıtmak amacıyla bu model dörtgen ya da trapez biçiminde yeniden yorumlanmıştır (Vilain, Berthommier & Boë, 1999).

Artikülatuvar temelde gelişen, ardından akustik fonetik düşünceyle biçimlenen bu temsil, zamanla fonolojik bir yapıymış gibi yanlış yorumlanmaya başlanmıştır.

“Ünlü” kavramı fonolojik yani zihinsel, “üçgen” ve “dörtgen” ise fonetik diğer bir ifadeyle fiziksel düzleme ait terimlerdir; bu durum, kavramsal karışıklığı derinleştiren önemli bir etkidir.

3. Ünlüler Ünlü Dörtgeni Üzerinde Nasıl Gösterilir?

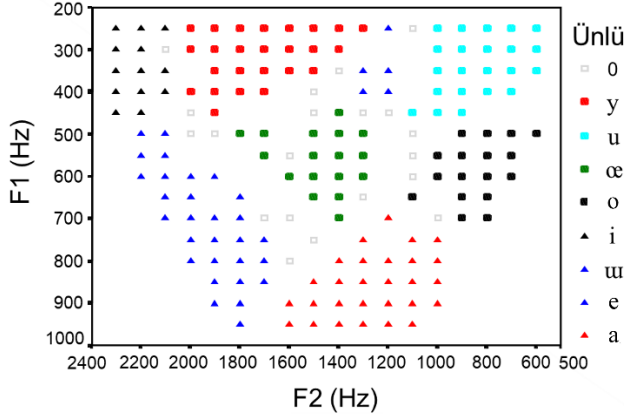
Ünlüler, ünlü dörtgeni üzerinde farklı akustik özelliklerine göre konumlandırılır. Bu gösterimde temel olarak formant³ frekansları (F1, F2, bazen F3) ve temel frekans (F0) değerleri kullanılır. Ölçüm birimi genellikle Hertz (Hz) olmakla birlikte, bazı çalışmalarda Bark ya da Mel gibi algısal frekans ölçekleri de kullanılmıştır.

F1, ünlünün kapalı–açık (dar–geniş) olma derecesini; F2 ise arka–ön (kalın–ince) olma özelliğini gösterir. F1 değerinin artması ağız açıklığının arttığını; F2 değerinin artması ise dil öne geldiğini, ünlünün incelendiğini gösterir. Bu amaçla, erişkin erkek veya kadın konuşuculara ait ortalama formant değerleri nokta olarak; değişkenliği göstermek içinse standart sapma elipsleri biçiminde görselleştirilebilir.

Daha seyrek olmakla birlikte, ünlüler algısal (perseptüel) özelliklerine göre de ünlü dörtgeni üzerinde gösterilebilir. Kılıç ve ark. (2003), 220 sentetik ünlü kullanarak gerçekleştirdikleri işitsel fonetik

³ Ünlülerin ve tınlı ünsüzlerin akustik analizi sırasında görülen enerji yoğunlaşmaları. 0-5000 Hz arasında beş formant olduğu kabul edilir ve F1, F2... şeklinde gösterilir.

çalışmasında, ünlü dörtgeninin işitsel algıya göre nasıl bölümlendiğini ortaya koymuştur (Şekil 1).



Şekil 1. Türkçe ünlülere ait algısal ünlü dörtgeni. Bu çalışma, çocuklarda işitme kaybının ünlü algısı üzerindeki etkilerini inceleyen geniş kapsamlı bir araştırmanın parçasıdır (Kılıç, 2003).

Bir dildeki ünlüleri fonolojik açıdan göstermek istediğimizde, bu tür fiziksel temsillere gerek yoktur; sesbirim envanteri basit bir tablo biçiminde sunulabilir.

Bazen küçük bir akustik fark, fonolojik kategoriye değiştirebilirken; bazen de büyük akustik farklar herhangi bir fonolojik ayırım yaratmayabilir. Fonetik özelliklerle fonolojik özellikler arasındaki ilişki her zaman doğrusal (lineer) bir biçimde ilerlemez.

Ünlüler fonolojik olarak genellikle ikili karşıtlıklar üzerinden kategorize edilir; ancak bazı dillerde üçlü ya da daha çoklu kategorizasyonlar da mümkündür. Türkçede ise sistem, esas olarak ikili karşıtlıklar üzerine kuruludur:

- i. Düz – Yuvarlak
- ii. Ön – Arka (ince – kalın)
- iii. Açık – Kapalı (geniş – dar)
- iv. Uzun – Kısa: Türkçede fonolojik değil, alofonik bir özelliktir.

Alt sesbirimlerin farklı soyutluk düzeyleri ve Türkçedeki ünlülerin alt sesbirimleri konusunda ayrıntılı bilgi için bkz. Kılıç (2017).

4. Fonetik ve Fonolojik Düzlemlerin Karışması

Soyutluk–somutluk ekseninde fonoloji–fonetik ilişkisi:

Bir konuşma sesi için “ünlü” ifadesi, fiziksel özelliğini değil, o dildeki işlevini anlatır.⁴

(boy) kelimesi bu duruma küçük ama açıklayıcı bir örnek oluşturur: Türkçede bir varlığın dikey doğrultudaki uzunluğunu, İngilizcede ise erkek çocuğu ifade eder. Yazılışları aynı, telaffuzları birbirine oldukça benzerdir; ancak her iki dilde farklı fonolojik birimlere karşılık gelir.

Türkçedeki biçim iki ünsüz ve bir ünlüden oluşurken, İngilizcedeki biçim bir ünsüz ve bir ünlüden (diftong) oluşur:

/bɔɪ/ (İngilizce)⁵

/boj/ (Türkçe)

İşitsel algısı İngilizceye göre biçimlenmiş bir dinleyicinin Türkçede ⟨ğ⟩ harfiyle gösterilen yumuşak damak ünsüzünü bağımsız bir birim olarak ayırt edememesi de benzer bir örnektir; aynı fiziksel sinyal, farklı dilsel sistemlerde farklı biçimlerde çözülür.

Fonetik fonolojiyi şekillendirmez; sadece veri sağlar, asıl belirleyici olan fonolojidir

5. Türkçede /a/ Ünlüsünün Durumu

Türkçedeki ⟨a⟩ ünlüsü fonolojik açıdan kalın, geniş ve düz ünlüdür. Ünlü dörtgeni üzerinde ince geniş ünlü alanını da kapsar. /a/ ünlüsünün kalın, ince, uzun, kısa fonolojik alt sesbirimleri vardır. Fonolojik alt sesbirim ifadesi dilin yapısında olduğunu gösterir, bu sesbirimler soyutluk-somutluk çizgisi üzerinde alt sesbirim ile fonetik alt sesbirim arasında yer alır.

/a/ ünlüsünün ince alt sesbiriminin varlığı sadece komşu seslerin etkisine bağlanamaz. ⟨g⟩, ⟨l⟩ harfleriyle gösterilen ünsüzlerin ince alt sesbirimleri eşşöyleyiş (koartikülasyon) nedeniyle komşu ünlüleri inceltir. Bu inceltme ⟨kar⟩ - ⟨kâr⟩ örneğinde olduğu gibi bazen ince ⟨a⟩ ünlüsü yapmaya yetmez, kelime ek aldığı kalın ek gelir (⟨karlı⟩ - ⟨kârlı⟩).

Her iki durumda da /a/ kalın (arka) ünlüdür; bu, gelen ek biçiminden anlaşılmaktadır. Peki, fark nedir? Aralarındaki fark, kelimenin başındaki ünsüzdendir. Kâr kelimesinde /k/ etkisiyle /a/ kısmen inceliyor ama ince ünlüye dönüşmüyor.

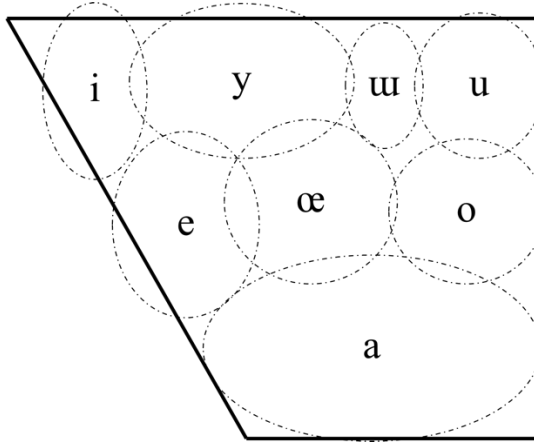
⟨harf⟩, ⟨harp⟩, ⟨sürat⟩ örneklerinde olduğu gibi palatal ünsüzlerden bağımsız olarak da ince ⟨a⟩ ünlüsü mevcuttur. Bu örneklerin başka dillerden Türkçeye geçmiş olması ince ⟨a⟩ ünlüsünü yok saymamız için gerekçe olamaz, bu ses sesbirim olarak olmasa da fonolojik alt sesbirim olarak Türkçede mevcuttur.

Yönetimsel Fonoloji (Government Phonology) ve onun alt bileşeni olan Element Teorisi, fonolojide soyut temsile dayalı yaklaşımın

⁴ Ünlü ve ünsüz kelimeleri fonolojik düzlemde yer alan kavramları işaret eder. Fonetik düzeyde ünlü ve ünsüzü belirtmek için literatürde *vocoid* (ünlümsü) ve *contoid* (ünsüzümsü) ifadeleri kullanılmaktadır.

⁵ /ɔɪ/ diftongu fonolojik olarak tek bir hece çekirdeği sayılır.

en uç biçimdir. Bu teoride sesler fiziksel niteliklerle değil, A (açıklık), I (incelik) ve U (yuvarlaklık) gibi temel elementlerin oluşturduğu yapısal ilişkilerle tanımlanır. Türkçedeki /a/ ünlüsünün durumunu Element Teorisi çerçevesinde yorumlayan Pöchtrager’e (2006) göre, Türkçedeki /a/ ünlüsünü tanımlayan tek element “açıklık (*openness*)”tır. Başka bir ifadeyle, Türkçedeki /a/ ünlüsünün temel fonolojik özelliği açıklıktır; bu ünlünün ön-orta-arka özelliği (hatta düzlük-yuvarlaklık özelliği) yoktur. Kılıç ve ark.’nın (2003) bulgularını Element Teorisine göre yorumladığımızda Türkçe ünlü dörtgenini Şekil 2’deki gibi gösterebiliriz.



Şekil 2. Ünlü dörtgeni üzerinde Türkçe ünlülerin kapladığı alanlar. /a/, /e/, /o/, /u/, /i/ sesbirimlerini göstermek için alfabede mevcut olan (işaretsiz) karakterler tercih edilmiştir. Alfabede (ü), (ı), (ö) harfleriyle gösterilen sesbirimler ise sırasıyla /y/, /u/ ve /œ/ şeklinde gösterilmektedir.

6. Sonuç

Türkçede sesbirim olarak mevcut olan /a/ ünlüsü düz, açık (geniş) ve kalın (arka) ünlüdür. Bu sesbirimin fonolojik olarak iki varyantı vardır: kalın ve ince. Konuşma sırasında ortaya çıkan akustik farklılıklar, dilin fonolojik düzlemini belirleyen öğeler olarak değerlendirilmemelidir. Bu tür akustik parametrelerden yola çıkarak Türkçenin fonolojik yapısına ilişkin yargılarda bulunmak, yöntemsel açıdan hatalı ve kuramsal dayanaklardan yoksun bir tutumdur.

Fonetik ölçütler dikkate alındığında, Türkçedeki /a/ ünlüsünün gösterimi için [ɐ], [ʌ], [a] ve [ɑ] gibi farklı UFA sembolleri kullanılabilmektedir. Ancak, işaretlilik teorisi açısından bu ünlü sistemin işaretsiz

(*unmarked*) birimi olarak değerlendirilmeli ve sesbirimsel düzeyde /a/ biçimiyle gösterilmelidir.⁶

Türkçedeki /a/ ünlüsü için arka-orta-ön olma durumu fonetik özelliktir, sesbirim olarak bu ünlüden bahsedilirken /a/ işareti kullanılmalıdır.

Kaynaklar

- Kılıç, M. A., Okur, E., Yıldırım, İ., Deniz, M., & Kara, E. (2003). *Normal ve işitme engelli çocuklarda vokallerin algılanması*. XXVII. Türk Ulusal Otorinolarenoloji ve Baş Boyun Cerrahisi Kongresi'nde sözlü bildiri olarak sunulmuştur, 4–9 Ekim 2003, Antalya, Türkiye.
- Kılıç, M. A. (2017). Türkiye Türkçesindeki sesbirimlerin dışsal ve içsel alt sesbirimleri: Katmanlı çevriyazı modeli. *Turkish Studies*, 12(7), 195–214. <https://doi.org/10.7827/TurkishStudies.11354>
- Pöchtrager, M. A. (2006). Is there phonological vowel reduction in Turkish? İçinde A. Asan, A. Genç & E. Erguvanlı-Taylan (Ed.), *Essays in honour of Eser Erguvanlı-Taylan* (ss. 25–40). İstanbul, Türkiye: Boğaziçi University Press.
- Vilain, C. E., Berthommier, F., & Boë, L.-J. (2015, September). *A brief history of articulatory-acoustic vowel representation*. In HSCR 2015 – 1st International Workshop on the History of Speech Communication Research (pp. 1–6). Dresden, Germany. <https://hal.science/hal-01197460>

⁶ İşaretlilik teorisi (markedness theory), dildeki kategorilerin veya özelliklerin “işaretsiz” (unmarked) ya da “işaretlili” (marked) olarak ikiye ayrıldığını öne süren bir teoridir. İşaretsiz olan biçim, sistemin varsayılan veya doğal halini; işaretlili olan biçim ise özel, belirli, sınırlı bir durumu temsil eder.

ALTAIC STUDIES – WHAT HAS NOT BEEN COMPARED SO FAR

Michael KNÜPPEL

Prof., Liaocheng University, e-mail: michaelknueppel@gmx.net, ORCID: 0000-0002-6348-5100

Knüppel, M. (2025). Altaic studies – what has not been compared so far. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 200–211). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X010>

ABSTRACT

This short article presents some of the author’s reflections on what has not yet been compared in Altaic studies, a field in which quite diverse approaches to language comparison have been employed over the past two centuries. While it is generally clear what cannot be compared, or should not be used for comparison, to clarify the linguistic relationships between the Turkic, Mongolian, and Tungus languages (onomatopoeia, children’s language, taboo language forms, etc.), it seems that hardly anyone has given much thought in recent times to what else could or should be compared. The author offers three examples: commands to call or shoo away animals, kinship terminology (or more precisely: systems of kinship relationships), and verbs relating to insect behavior.

Keywords: Commands to call or shoo away animals, Systems of kinship terminology, Verbs relating to insect behavior

ÖZ

Bu kısa makale, son iki yüzyıl boyunca dil karşılaştırmasına ilişkin oldukça farklı yaklaşımların benimsendiği Altayistik çalışmalarında, henüz karşılaştırılmamış olan alanlara dair yazarın bazı değerlendirmelerini sunmaktadır. Türk dilleri, Moğol dilleri ve Tunguz dilleri arasındaki dilsel ilişkileri açıklığa kavuşturmak amacıyla hangi unsurların karşılaştırılamayacağı ya da karşılaştırmada kullanılmaması gerektiği genel olarak açıktır (yansıma sözcükler, çocuk dili, tabu sözcükler vb.). Buna karşın, son dönemde başka hangi unsurların karşılaştırılabileceği ya da karşılaştırılması gerektiği üzerine neredeyse hiç düşünülmendiği görülmektedir. Yazar bu bağlamda üç örnek sunmaktadır: hayvanları çağırma ya da uzaklaştırma komutları, akrabalık terminolojisi (daha doğrusu akrabalık ilişkileri sistemleri) ve böcek davranışlarıyla ilişkili eylemler.

Anahtar Sözcükler: Hayvanları çağırma ya da uzaklaştırma komutları, akrabalık terminolojisi sistemleri, böcek davranışlarıyla ilişkili eylemler

1. Introduction

Since linguists do not receive Nobel Prizes – at least not for achievements in their specific field of work – and do not develop ground-breaking cures for ailments and diseases or make any other discoveries or inventions that change the world or even lead to a significant improvement in people's everyday lives, the most prestigious thing they can achieve seems to be the discovery of previously unknown or undescribed languages or the identification of linguistic relationships. Since Sir William Jones' (re)discovery of the Indo-European language family at the end of the 18th century,¹ the latter has fired the imagination, and over the past two and a half centuries, all possible (and impossible) languages, language groups and language families have been compared with each other for the purpose of proving some kind of relationship. The idea of an Altaic or even Ural-Altaic linguistic relationship is therefore by no means unusual, but rather part of a seemingly endless series of attempts to reap the laurels of v. Schlegel, Bopp or Pott, who proved the previously assumed relationship. What has not been compared in order to prove the relationship between the Turkic, Mongolian and Tungus languages (and perhaps other 'candidates' such as Japanese, Korean or the language of the Ryūkyū Islands)? The history of these endeavours alone could fill several volumes.

2. Phases of Altaic Language Comparisons – Roughly Summarized

If we consider the history and the essential developments of (Ural-) Altaic language comparisons in the broadest terms and with almost criminal neglect of any accuracy, we can distinguish between different 'phases':

1. Initially, the lexicon was compared on a rather narrow basis (many languages were not yet documented at all or only inadequately – apart from the fact that for the Turkic languages, most of the material was taken from Ottoman Turkish, for the Mongolian languages from classical Mongolian texts or Kalmyk, and for the Tungus languages from Manžū, of all places) – pure word similarity with somehow similar meanings dominated the endeavours, as undertaken, for example, by A. Boller or W. Schott.
2. With F. M. Müller, typology (already established by W. v. Humboldt) came to the fore when he popularised his 'Turanian languages', and
3. in the second half of the 19th century, it was primarily morphological comparisons, including so-called 'root etymology' ['Wurzel-etymologie'] (e.g. H. Winkler or W. Bang), as is still sometimes practised today in Uralic studies.

¹ Strictly speaking, Sir William Jones (28 September 1746 – 27 April 1794), who discovered the relationship of the Indo-European languages in 1786 (Jones [1786]), already had a whole series of predecessors, but there is no need to deal with them here.

4. In the 20th century, comparative syntax (e.g. V. Pröhle or D. R. Fokos-Fuchs) initially took centre stage or was added to previous endeavours,
5. Later, the focus shifted to the postulation of larger units – for example, in the context of the Nostratic language family or other so-called ‘super-phyla’, which literally incorporated the previously claimed ‘macro-families’, and ultimately even ‘global etymologies’, etc.

In addition, the approaches from previous ‘phases’ of comparative linguistic (Ural-)Altaic research were naturally continued, while others, such as comparative phonological considerations, were naturally employed in comparative studies at all times – we find them already in the early studies of the ‘Tatar languages’ [‘Tatarische Sprachen’] or the ‘Altai or Finnish-Tatar language family’ [‘Altai’sches oder Finnisch-Tatarisches Sprachengeschlecht’] by W. Schott.

3. What Cannot and Should not be Compared

While the ‘methods’ and approaches – or at least the main areas of investigation – changed over time, the question naturally arose as to what could actually be compared and what could not. In his articles, G. Doerfer occasionally summarised which lexical materials are not suitable for comparative linguistic analysis and in which cases comparisons to determine relationships are downright inadmissible. Roughly speaking, he excluded the following groups of word material:

1. Cultural words [Kulturwörter];
2. onomatopoeia;
3. children’s language forms [kindersprachliche Formen; Babysprache];²
4. taboo forms ~ language taboos;
5. loanwords (or forms that can be clearly identified as such);
6. ‘language tracts’ [‘Sprachtrakte’], i.e. pre-literary language contacts;³
7. expressive and impressive words;
8. ‘pre-relational’ forms [‘vorverwandtschaftliche’ Formen], e.g. pronouns,⁴ i.e. forms that originate from an early phase of human language development, or more precisely: a phase before the differentiation of individual languages.

Other contemporaries, such as Doerfer’s teacher K. H. Menges, who was much more ‘open’ to the study of ‘more distant linguistic relationships’ [‘entferntere Sprachverwandtschaften’] and was also convinced of the genetical relationship of the Altaic languages, made a very clear

² Doerfer (2001), e. g. p. 193.

³ Doerfer (2001), pp. 182 and 222-223.

⁴ Doerfer (2001), p. 182, for more details, see also pp. 223 ff.

distinction between words of primordial origin and genetically related languages. Some who were (and still are) enthusiastic about comparing languages across the boundaries of ‘established language families’ more or less adhered to this distinction (or were aware of the problems involved), while several of the “enthusiasts” ignored such objections and continued their endeavours in the ‘tried and tested manner’.

More interesting than the question of what has been compared, what can be compared and what cannot be compared is the question already raised above – albeit in a rather rhetorical or polemical form – of what has not been compared so far.

4. Calling and Shooing Commands

First and foremost, of course, is the complex of calling and shooing commands for animals, which has also been dealt with by the author of these lines – albeit only for individual languages and not in an overall Altaic context. Such commands have already been studied by A. v. Le Coq for the Turki language,⁵ by the author of this article for the question of differences between formerly nomadic and sedentary peoples on the basis of Kurmancî,⁶ and also by the same author for Tungus languages,⁷ on the basis of materials by S. M. Shirokogorov⁸ and B. Piłsudski.⁹

In lexicographical work, such commands to call and shoo away animals have been and continue to be omitted for various reasons, as the author has occasionally emphasised,¹⁰ even though the Brothers Grimm had already recognised the value of these forms.¹¹ This is mainly due to the fact that we are dealing with particularly archaic forms of calling and shooing commands, which are not easily borrowed (apart from their invention in the course of the introduction of new domesticated animals) or suppressed, they naturally appear particularly suitable for comparative linguistic studies, as the author of these lines has also pointed out in the past:

This can be explained by the fact that such exclamations demonstrate a particular ‘persistence’ – they are not subject to any (or virtually no) sound changes due to linguistic history, are rarely borrowed (if at all, then only through the appearance of new domesticated or domesticable animal species – whether as a result of changes in the settlement area or through

⁵ v. Le Coq (1919).

⁶ Knüppel (2017).

⁷ Knüppel (2013).

⁸ Here his ‘Tungus dictionary’ (Shirokogorov [1944/1953]) and its evaluation by G. Doerfer in the form of the ‘Etymologisch-Ethnologisches Wörterbuch’ (Doerfer [2004]).

⁹ Above all, the edition of his Tungusological notes by A. F. Majewicz (2011).

¹⁰ Knüppel (2013), p. 270.

¹¹ For example, in the case of the word ‘kitz!’ in: Grimm / Grimm (1984), col. 868; as the Brothers Grimm recorded commands to call and shoo away animals in general in their dictionary.

the introduction of such animals) and do not simply disappear in the course of changes caused by the transformation of material culture.¹²

5. Systems of Kinship Terminology

The systems of kinship terminology represent a completely different area for conceivable, but as yet unrealised comparisons in the field of Altaic studies. These are also of interest to us, of course, insofar as the terms can often be counted among the basic vocabulary of a language¹³ (apart from child language forms) and thus seem to allow conclusions of any kind to be drawn.¹⁴ Although we have V. I. Cincius' groundbreaking work on Altaic kinship terminology,¹⁵ this work focuses on the lexic itself and not on the comparison of the various systems of kinship terminology. The author of this article also addressed this issue several years ago.¹⁶

But first, a few brief fundamental remarks on the kinship terminology systems that have already been mentioned several times. It was recognised early on that the possible combinations of kinship terminologies are comparatively limited and that, precisely because of these limitations, a number of types can be identified that we encounter again and again. Attempts to classify these types under specific systems have been made since the late 19th century, apart from early endeavours in this direction, such as those of J. F. Lafiteau.¹⁷ Of course, L. H. Morgan should be mentioned here first, who divided the systems into 'descriptive' and 'classificatory'.¹⁸

Although these endeavours were, of course, 'children of evolutionism', the approach found its continuation in later ethnology, such as the various models of classification found in R. Lowie¹⁹ and G. P. Murdock, which ultimately formed the basis for the classifications based on ideal-typical systems still used today for description. In Murdock, we finally find the well-known classification comprising six basic types. In the following lines, the author of this article follows the explanations in his own article from 2014:

Strictly speaking, Murdock's classification was based on P. Kirchhoff's 'core group' ['Kerngruppe'], i.e. a classification comprising two

¹² Knüppel (2013), p. 273: "Dies erklärt sich daraus, daß solche Rufe ein besonderes 'Beharrungsvermögen' zeigen – sie unterliegen keinem (oder so gut wie keinem) resp. kaum nachweisbarem sprachgeschichtlich bedingten Lautwandel, werden höchst selten entlehnt (allenfalls durch das Auftreten neuer domestizierter oder domestizierbarer Tierarten – sei es infolge von Änderungen des Siedlungsraumes oder durch Einführung solcher Tiere) und verschwinden nicht einfach im Zuge von Änderungen, die durch den Wandel der materiellen Kultur bedingt sind".

¹³ See, for example, Gulya (1983).

¹⁴ The fact that kinship terms are also borrowed is usually not taken into account; cf. for example the remarks made by N. Poppe and, earlier, by E. Sapir (Poppe [1974], p. 124; Sapir [1920], p. 269 and [1921]).

¹⁵ Cincius (1972).

¹⁶ Knüppel (2014).

¹⁷ Here his remarks on the kinship system of the Iroquois compared to those in Europe (Lafiteau [1724]).

¹⁸ Morgan's work on the Iroquois marked the beginning of various descriptions of kinship systems and the examination of the terms used to describe them (Morgan [1851]).

¹⁹ Lowie (1928).

generations (1. parent generation and 2. ego generation [see below]).²⁰ It also included Lowie's system of naming schemes for the parent generation, which distinguished between four schemes: 1. linear system, 2. generational system, 3. bifurcated fusion system, and 4. bifurcated collateral system. All four types have in common that the relatives of the different generations are distinguished terminologically. Murdock took up Lowie's model, but further differentiated the bifurcated fusion system and named the different systems after societies in which they are found or in whose languages they are found. Murdock's designations are still used today for the basic types of kinship typology. Murdock's models are summarised below:²¹

5.1 Hawai'i system (simplified; male ego)

This is based on Lowie's generational scheme, i.e. all relatives of a generation – distinguished only by gender – are referred to by the same term. The emphasis here is on differentiation within the extended family, as no distinction is made between linearity and collaterality. The Hawai'i system is particularly common among Polynesian ethnic groups and in their languages (but also in Annamitic, for example).

Hawaii-System (vereinfacht, Ego männl.):

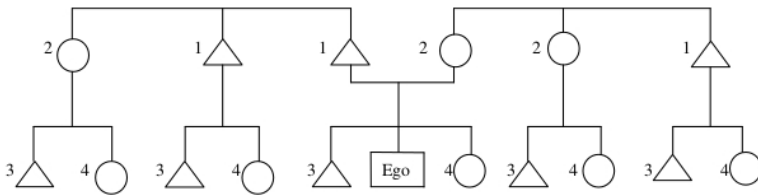


Figure 1. Hawai'i system (simplified; male ego)

5.2 Sudanese system

Here, both the gender principle and the bifurcated collateral system are reflected, i.e. both the members of the core group and the parallel and cross relatives are designated differently on the patrilineal and matrilineal sides. This system was first described on the basis of South Sudanese findings. In contrast to the Hawai'i system, the Sudanese system shows the most far-reaching terminological differentiation.

²⁰ Kirchhoff (1932).

²¹ See Murdock (1949) and (1970) for details.

Sudan-System:

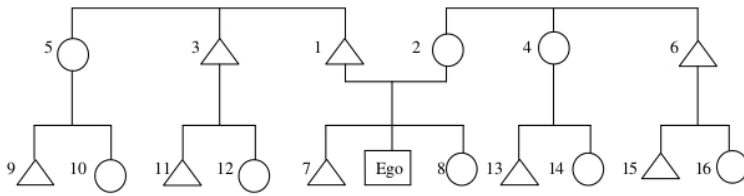


Figure 2. Sudan system

5.3 Eskimo system

In this system, a distinction is made primarily in terms of terminology between the core family and the extended family, i.e. the principle of kinship lines applies, whereby no distinction is made in terminology between parallel and cross kinship. This system is also found in Basque, for example.

Eskimo-System:

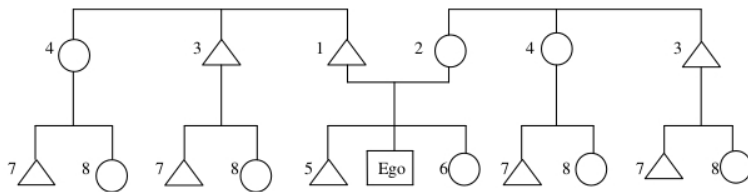


Figure 3. Eskimo system

5.4 Iroquois system

In this system, the only distinction made in terms of terminology is between gender and generation. The terms for the core group merge with those for parallel relatives, i.e. it corresponds to Lowie's scheme 3 (furcat-ed fusion system).

Irokesen-System:

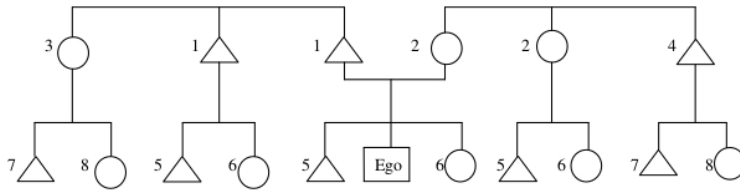


Figure 4. Iroquois system

5.5 Crow system (simplified; male ego)

In this system, which also corresponds to Lowie's scheme 3, descent takes precedence over the generational principle. Here, the terms used for parallel cousins correspond terminologically to those used for siblings, but a distinction is made between patrilineal and matrilineal cross cousins.

Crow-System (vereinfacht, Ego männl.):

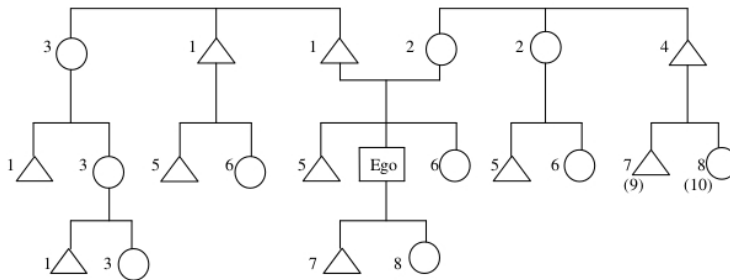
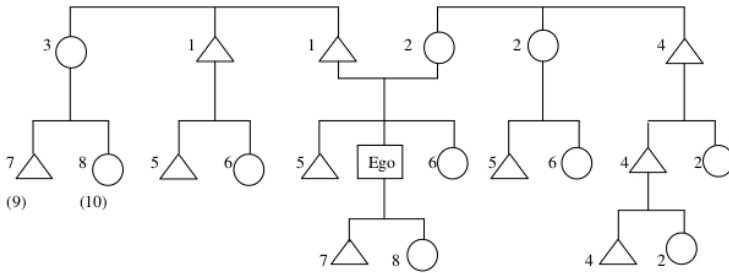


Figure 4. Crow system (simplified; male ego)

5.6 Omaha system (simplified; ego female)

As with the Crow system, the Omaha system is also a variant of Lowie's scheme 3, i.e. here too, descent takes precedence over the generational principle. The Omaha system is, so to speak, the patrilineal counterpart to the Crow system.

Omaha-System (vereinfacht; Ego weiblich):

**Figure 5.** Omaha system (simplified; ego female)

Of course, it quickly becomes clear that the conditions in the Turkic, Mongolian or Tungus linguistic world do not necessarily follow any of these ideal models, and the usability of the expected data for answering the so-called ‘Altaic question’ is limited insofar as the affiliation to certain types and systems within classifications and typologies can change over the course of centuries or millennia (just think of the possibility of a change in the typological status of a language, e.g. Northern Tažik, which was already in the process of transitioning from an inflecting language to a suffix-agglutinating one [a process that was only reversed by language policy in the Soviet Union]). But even knowledge of such possible changes was already a gain. Of course, it must also be taken into account that not all languages in a group of languages or a language family necessarily belong to the same typological system (for example, the typological affiliation of West Greenlandic differs from that of the other Eskimo-Aleut languages, which belong to the Eskimo system [see above]). Nevertheless, it should be clear that while the borrowing of individual kinship terms is conceivable, the borrowing of a complete kinship terminology system is not – and here it is important to compare these systems with each other (if necessary, also on the basis of reconstructions of earlier states), not the less meaningful terminology.

6. Verbs Relation to the Behaviour of Insects

A third complex of previously uncomparable terms is the use of verbs from the field of zoology. This does not refer to the obvious terms found in every dictionary and known to every child – the horse *neighs* and *trots* or *gallops*, the cow *moos*, the dog *barks* and *wags* its tail ... – but rather those that do not appear to be listed anywhere and are not related to domesticated animals (which, in case of doubt, are also ‘cultural words’), but rather to animals that are predominantly native to almost all of Eurasia and have always been part of the natural environment of humans in the

region in question. Insects in particular seem to be particularly suitable for this purpose – for example, the *chirping* or *singing* of crickets and cicadas, the *buzzing* of mosquitoes or even the *tapping* of certain beetle larvae.

The author of these lines has occasionally dealt with such terms relating to insects from various Turkic languages in a series of articles.²² Of course, corresponding studies can also be conducted beyond the boundaries of this language family on the basis of Mongolian and Tungus findings. However, when it comes to possible Altaic relations, one must not make the mistake of focusing on any similarities between words. Rather, it would be necessary to determine whether, for example, all Altaic languages originally distinguished between the *biting* and *stinging* of such insects, which have corresponding ‘mouthparts’ or stingers, or whether the distinction is made for larger animals but not for smaller ones, or whether corresponding ‘patterns’ can be found consistently in the languages of the three language families or not.

This is not about etymological playings (as we see in their most extreme form among various ‘long rangers’ – Matisoff spoke in this context of ‘megalocomparison’,²³ Doerfer of ‘omnicomparatism’²⁴), but about comparing structures that need to be uncovered and defined more precisely. In essence, we are of course dealing here with something completely different from what Ernst Lewy once mockingly remarked with regard to attempts to determine ‘more distant linguistic relationships’: ‘Give me two dictionaries and I will prove to you that the languages are related!’ [‘Geben Sie mir zwei Wörterbücher und ich beweise Ihnen, daß die Sprachen miteinander verwandt sind!’]²⁵ As with the comparative study of kinship systems, it is a matter of comparing structures, not words!

7. Conclusions

With a little creativity and serious effort, there are certainly many more areas that could be brought into play for one reason or another. What has not been compared is perhaps even more interesting than what has already been compared. On the one hand, because it is tiring to keep doing the same exercises over and over again – especially when nothing fundamentally new emerges (except perhaps the addition of ever new ‘empirical findings’ for or against the assumption of an Altaic language relationship) – and on the other hand, because this also raises the question of why completely different areas were not included in the considerations. Now, the previous proponents of conflicting views should not be accused of a lack of creativity at this point – some of them may even have been ‘blessed’

²² Knüppel (2025), (–a) and (–b).

²³ Matisoff (1990).

²⁴ Doerfer (1973).

²⁵ Personal communication from G. Doerfer (Göttingen); Lewy had explained this while speaking to an audience (including G. Doerfer) in one of the offices in Berlin, standing in front of a bookshelf. In somewhat simplified form, the quote can also be found in Doerfer, e.g. in Doerfer (2001), p. 183: “Gib mir zwei Sprachen, und ich beweise dir, daß sie verwandt sind”.

with a little too much imagination – but dogmatism has never led to truly new insights in science, as is well known. In other words, it is high time to change our perspective again or – to stick with the image given above – to move on to the next ‘phase’ of Altaic research, to open a new chapter.

Given the space available here, it is not possible to explore any of the proposed areas of investigation, let alone exhaustively. That would require a completely different scope. But perhaps one or two colleagues will be inspired by these lines, which are purely programmatic in nature (some may simply regard them as polemical nonsense), to undertake corresponding projects – and who knows what the possible outcome might be? One should always bear in mind that no one really gains anything if it is proven that the Altaic languages are related to each other in the sense of a genetic relationship, and no one loses anything if the opposite is irrefutably established. As I said at the beginning, no one will receive the Nobel Prize for this, be canonised by the church, or be remembered by humanity as a new Alexander Flemming, Louis Pasteur, Paul Ehrlich or Robert Koch!

References

- Cincius, V. I. (1972). K étimologii altaïskikh terminov rodstva. In *Ocherki sravnitel'noi leksikologii altaïskikh iazykov* (pp. 15–70). Leningrad.
- Doerfer, G. (1973). *Lautgesetz und Zufall: Betrachtungen zum Omnicomparatismus* (Innsbrucker Beiträge zur Sprachwissenschaft, 10). Innsbruck.
- Doerfer, G. (2001). Altaica gottingensia. *Central Asiatic Journal*, 45(2), 181–229.
- Doerfer, G. (2004). *Etymologisch-ethnologisches Wörterbuch tungusischer Dialekte (vornehmlich der Mandschurei)*. Hildesheim; New York.
- Grimm, J., & Grimm, W. (1984). *Deutsches Wörterbuch* (Vol. 11). München.
- Gulya, J. (1983). Survival-Erscheinungen in den Verwandtschaftstermini der uralischen Völker. *Ural-Altaische Jahrbücher, Neue Folge*, 3, 1–8.
- Jones, W. (1786/1799). The third anniversary discourse, on the Hindus, delivered 2nd of February, 1786. In *The works of Sir William Jones* (Vol. 1, pp. 19–34). London.
- Kirchhoff, P. (1932). Verwandtschaftsbezeichnungen und Verwandtenheirat. *Zeitschrift für Ethnologie*, 64, 41–71.
- Knüppel, M. (2013). Überlegungen zu Lock- und Scheuchrufen für Tiere in tungusischen Sprachen und Dialekten. *Journal of Oriental and African Studies*, 22, 270–273.
- Knüppel, M. (2014). Verwandtschaftstermini des Armanischen. *Journal of Oriental and African Studies*, 23, 7–19.
- Knüppel, M. (2017). Überlegungen zu Lock- und Scheuchrufen für Tiere im Kurmancî. *Journal of Oriental and African Studies*, 26, 386–388.
- Knüppel, M. (2025). Entomologia Turcica – Überlegungen zum Sprachgebrauch Insekten betreffend [I] (Nachtrag). *Journal of Oriental and African Studies*, 34, 181–183.
- Knüppel, M. (in press-a). Entomologia Turcica – Überlegungen zum Sprachgebrauch Insekten betreffend. *Bitig*.

- Knüppel, M. (in press-b). Entomologia Turcica – Überlegungen zum Sprachgebrauch Insekten betreffend (II). *Journal of Oriental and African Studies*.
- Lafiteau, J.-F. (1724). *Mœurs des sauvages américains, comparées aux mœurs des premiers temps* (Vols. 1–2). Paris.
- Le Coq, A. von. (1919). Osttürkische Lock- und Scheuchrufe für Tiere. *Mitteilungen des Seminars für Orientalische Sprachen an der Friedrich-Wilhelms-Universität zu Berlin, Westasiatische Studien*, 22, 110–111.
- Lowie, R. H. (1928). A note on relationship terminologies. *American Anthropologist*, 30, 263–267.
- Majewicz, A. F. (Ed.). (2011). *The collected works of Bronislaw Pilsudski* (Vol. 4: Materials for the study of Tungusic languages and folklore). De Gruyter.
- Matisoff, J. A. (1990). On megalocomparison. *Language*, 66(1), 106–120.
- Morgan, L. H. (1851). *League of the Ho-dé-no-sau-nee, or Iroquois*. Rochester, NY; Boston, MA.
- Murdock, G. P. (1949). *Social structure*. New York, NY.
- Murdock, G. P. (1970). Kin term patterns and their distribution. *Ethnology*, 9(2), 165–208.
- Poppe, N. (1974). Remarks on comparative study of the vocabulary of the Altaic languages. *Ural-Altaische Jahrbücher*, 46, 120–134.
- Sapir, E. (1920). Nass River terms of relationship. *American Anthropologist*, 22, 261–271.
- Sapir, E. (1921). A Haida kinship term among the Tsimshian. *American Anthropologist*, 23, 233–234.
- Shirokogorov, S. M. (1944/1953). *A Tungus dictionary: Tungus-Russian and Russian-Tungus* (S. Iwamura, Ed.). Tokyo.

YAPAY ZEKÂ ÇAĞINDA ADLİ DİLBİLİMDE YAZAR ANALİZİ

Hülya KOCAGÜL

Dr. Öğr. Gör., İstanbul Medeniyet Üniversitesi, Dilbilimi Bölümü, hulya.kocagul@medeniyet.edu.tr,
ORCID: 0009-0001-3188-0907

Kocagül, H. (2025). Yapay zekâ çağında adli dilbilimde yazar analizi. İçinde O. Çınar, F. Başbuğ & H. Aydemir (Ed.), *Çağdaş dilbilim yazıları I: İMÜ Dilbilimi 15. yıl armağanı* (ss. 212-243) Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X011>

ÖZ

Bu çalışma, büyük dil modellerinin (LLM) metin üretim kapasitesinin adli dilbilimsel yazarlık analizinin dönüşümünü incelemek üzere bir çerçeve sunmayı amaçlamaktadır. Buna göre, yazarlığın kavramsal temellerini sorguluyor, metodolojik dönüşümü ve hukuki boyutları tartışılmaktadır. Barthes'ın (1967) yazarın otoritesini sorgulayan eleştirisinden Foucault'nun (1969) yazar işlevi kavramına, Love'ın (2002) işlevsel yazarlık modeline uzanan teorik zemin, yapay zekâ (YZ) çağında yeni bir gerilim alanına dönüşmüştür. Sousa-Silva'nın (2024) vurguladığı gibi, yapay zeka tarafından üretilen metinler bireysel yazarlık özelliklerinden yoksun olduğundan, adli dilbilimsel analiz yöntemlerin stilistik unsurlardan derin dilbilimsel özelliklere kayması gerekmektedir (ss.164-165). Çalışmanın temel argümanı, adli dilbilimin geleneksel varsayımları olan yazar-içi tutarlılık (aynı yazarın farklı metinlerde benzer dilsel özellikler sergilemesi) ve yazarlar-arası ayırt edicilik (farklı yazarların birbirinden ayrılan örüntüler göstermesi) ilkelerinin (Grant, 2022), yapay zekâ aracılı metin üretimi bağlamında koşulsuz geçerliliklerini yitirme olasılıklarıdır. Bu durum, kavramsal çerçevelerin yeniden değerlendirilmesini zorunlu kılar. Çalışma, Love'ın (2002) işlevsel yazarlık modelinin yapay zekâ bağlamında yeniden yorumlanmasını, yapay bireydil kavramının insan bireydilinden (Coulthard, 2004) farklılığını ortaya koyarak katkılar olarak sunar. Metodolojik açıdan, Nini'nin (2023) Dilsel Bireysellik Teorisi'nin önerdiği çok katmanlı analizin önemi vurgulanmakta ve hukuki boyutta ise Daubert kriterlerinin yapay zekâ çağındaki uygulanabilirliği sorgulanmaktadır.

Anahtar Sözcükler: Adli dilbilim, Yazar analizi, Bireydil, Yapay zekâ yazar analizi, Yapay bireydil

ABSTRACT

This study aims to propose a framework for examining how the text-generation capacities of large language models (LLMs) are transforming forensic linguistic authorship analysis. In this context, the conceptual foundations of authorship are interrogated, while its methodological transformation and legal implications are discussed. The theoretical trajectory extending from Barthes's (1967) critique of

authorial authority to Foucault's (1969) concept of the author function, and further to Love's (2002) functional model of authorship, has evolved into a new field of tension in the age of artificial intelligence (AI). As emphasized by Sousa-Silva (2024), since AI-generated texts lack individual authorial characteristics, forensic linguistic methods need to shift from surface-level stylistic features to deeper linguistic properties (pp. 164–165). The central argument of this study is that the traditional assumptions of forensic linguistics—namely intra-author consistency (the expectation that the same author exhibits similar linguistic features across different texts) and inter-author distinctiveness (the assumption that different authors display distinguishable patterns) (Grant, 2022)—may lose their unconditional validity in the context of AI-mediated text production. This situation necessitates a re-evaluation of existing conceptual frameworks. As a contribution, the study proposes a reinterpretation of Love's (2002) functional model of authorship in the context of AI, highlighting the distinction between artificial idiolect and human idiolect (Coulthard, 2004). From a methodological perspective, the importance of the multi-layered analysis advocated by Nini's (2023) Theory of Linguistic Individuality is underscored, while, in legal terms, the applicability of the Daubert criteria in the age of AI is critically examined.

Keywords: Forensic linguistics, Authorship analysis, Idiolect, AI Authorship detection, Artificial idiolect

1. Giriş

Yazarlık kavramı, “Bu metnin sahibi kimdir?” sorusu ekseninde felsefe, edebiyat, hukuk gibi alanları kesiştiren, insanlık tarihinin merkezi entelektüel meselelerinden biridir. Metnin özgür ve yaratıcı kaynağı olarak görülen yazarlık Love’ın (2002) belirttiği gibi değişen toplumsal ve hukuki normların bir ürünü olarak sürekli dönüşüm geçirmektedir. Bu dönüşümün tarihsel izlerine bakıldığında yazarlık üzerine sorular yeni değildir. Roland Barthes 1967 yılında ‘Yazarın Ölümü’ adlı eserinde “okuyucunun doğuşunun yazarın ölümü pahasına gerçekleşmesi gerektiğini” (s.148) savunarak yazarın mutlak otoritesini sorgulamıştır. Daha sonra Michel Foucault (1969) ise bu düşüncüyü bir adım öteye taşıyarak “Kimin konuştuğunun ne önemi var ki?” (s.314) sorusunu sorar ve yazarı bir kişiden daha çok bir işlev gibi değerlendirir.

Yazarın önemini sorgulayan bu felsefi tartışmalara rağmen Love (2002), Barthes ve Foucault’un reddettiği romantik yazar mitini canlandırmaya çalışmadan edebi üretimin işbirlikçi doğasını kabul eder ancak onlardan farklı olarak her insanın benzersizliğinin nasıl ortaya çıktığına ve buradaki entelektüel işçiliğe odaklanarak bu tartışmaya farklı bir boyut getirir. Felsefi düzlemde yazarın kimliği önemsizleşmiş görünse de günümüzde üretken yapay zekâların milyonlarca sözcük ürettiği bir çağda Foucault’nun retorik sorusu (“Kimin konuştuğunun ne önemi var?”) yeniden somut bir nitelik kazanmıştır.

Büyük dil modelleri (LLM) çağında metnin anlamı için kimin konuştuğunun önemi ikincil öneme sahip olabilir ancak hukuki sorumluluk açısından kimliğin saptanması zorunluluğunu ortadan kaldırmaz. Böylece,

tartışma bugün teorik bir mesele olmanın ötesine geçerek uygulamada bir önem kazanır çünkü GPT-4 gibi modeller insan yazımına yakın taklit edebilirken (Kumarage vd., 2024) bilinçli bir niyet, eylem ve ahlaki sorumluluk taşımazlar. Bir yandan modellerin taklitçilik kapasitesi tespit sürecini zorlaştırırken, öte yandan üretim mekanizmalarının doğası gereği bıraktıkları izler insan yazarlığından ayrışır.

Sousa-Silva (2024, s. 164), büyük dil modellerinin olasılıksal yapısı nedeniyle ürettikleri her metnin benzersiz ve bireysel yazarlık özelliklerinden yoksun olduğunu, bu durumun bazı senaryolarda siber suçluların tespit edilmesini imkânsız hale getirebileceğini vurgular. Acemoğlu'nun (2021) değindiği gibi bu teknolojilerin düzenlenmemiş kullanımı demokrasiye, gizliliğe ve sosyal eşitliğe zarar verebilir. Sousa-Silva (2024) bu zararları "siber bağımlı" (hackleme vb.) ve "siber-olanaklı" (siber zorbalık, şantaj vs.) suçlar olarak ayırır. Özellikle siber-olanaklı suçlarda failer yapay zekâyı kullanarak geniş ölçekli ve yüksek inandırıcılığa sahip metinler üretebilmektedir. Bu kullanım senaryolarında yazarlık problemi bu zararların adli ve hukuki boyutunu oluşturan önemli bir mesele haline gelir.

Ceza hukuku sorumluluğu soyut bir yazar işlevine değil yaşayan bir faile bağlamak zorundadır. Bu zorunluluk, adli dilbilim uygulamalarını ve bireydir tabanlı yazar tanıma süreçlerini merkezi hale getirir. Somut olarak, Chaski'nin (2012) de belirttiği gibi yazar belirleme adli bağlamda pek çok farklı suç türünün soruşturmasında rol oynayabilir. Örneğin, anonim tehdit mektupları, fidye notları ve şüpheli belgeler gibi hukuki öneme sahip metinlerin yanı sıra; günlükler, e-postalar, iş yazışmaları ve blog paylaşımları gibi görünüşte daha az belirgin belgeler de cezai, hukuki ve güvenlik bağlamında bir kişinin olaya karışıp karışmadığını belirlemeye doğrudan katkıda bulunabilir.

Çalışmanın temel argümanı, tarihsel ve teknolojik gelişmeler ışığında, adli dilbilimsel yazar analizinin dönüşümünü incelemek üzere bir çerçeve sunmayı amaçlamaktadır. Günümüzde adli dilbilimci sadece "Bu metni kim yazdı?" sorusuyla değil: "Bu metin bir insan mı yoksa bir algoritma mı?" veya "Algoritma yardımıyla stilini gizleyen bir insan mı?" gibi çok daha karmaşık sorularla karşı karşıyadır. Love'ın (2002) yazarlığı tek bir öz değil, bir dizi bağlantılı teknik ve uygulama olarak gören yaklaşımı, bu yeni karmaşayı çözümlmek için uygun bir kuramsal zemin sunar. Bu çalışmada önerilen çerçeve Love'ın modelini genişleterek, metin üzerinde birden fazla failin (insan ve yapay zekâ) bulunabileceğini kabul eder ancak Grant ve MacLeod'un (2018) belirttiği gibi dilsel kimlik, bireydir teorisine bağlı kalarak bireysel sorumluluğu belirleme zorunluluğunu da korumaktadır.

1.1 Yazarlık kavramının dönüşümü

Yazarlık, niyet, eylem, sorumluluk üçgeninde kurulan karmaşık bir yapıdır. Bu yapının temelleri Barthes'ın (1967) yazarı metnin tanrısal yaratıcısı olarak gören bakış açısını sorgulamasıyla sarsılmıştır; Barthes, metni yazarın

niyetinin tek yönlü bir aktarımı değil, kültürden ve önceki metinlerden alınan alıntılarının bir dokusu olarak değerlendirmiştir. Barthes'ın yazarın otoritesini çözen yaklaşımı Foucault'nun sosyolojik bakış açısıyla birlikte ele alındığında yazarlığın sabit bir kimlikten ziyade değişken bir konum ve işlev olduğunu düşünülebilir. Günümüzde yapay zekâ milyonlarca veri girdisinden yaptığı derlemelerle ürettiği yazarı olmayan metin biçimlerini anlamlandırmak için teorik bir zemin sunar.

Foucault (1969) yazarlık sorununu sosyolojik bir boyuta taşır ona göre yazar işlevi bir söylemin gerçek bir bireye atfedilmesiyle kendiliğinden oluşmaz, tam tersine söylem alanını belirleyen sistemlere bağlı karmaşık işlemler içinde üretilir. Yazarlık işlevinin sabit veya evrensel olmadığını belirten Foucault (1969) buna örnek olarak Orta Çağ'daki bilimsel metinlerin yazar adını zorunlu kılarken, edebi metinlerin anonim olarak kabul edilmesini gösterir (s.308).

Barthes ve Foucault'un eleştirileri, yazarın sabit ve değişmez bir özne olduğunu varsayan geleneksel bakış açısını sarsmaları bakımından önemli ortaklıklar taşır fakat iki düşünürün odak noktaları farklılaşır. Barthes, yazarın tamamen ortadan kalkmasını ve okurun ayrıcalıklı konuma yükselmesini savunurken; Foucault, yazarın toplumsal ve söylemsel işlevi ile ilgilenir. Foucault'ya göre "yazarın kaybolduğu yönündeki boş onaylamayı yinelemek yeterli değildir" (1969, s.303) çünkü yazar işlevi, basitçe reddedilemeyecek kadar derin ve köklü bir yapıdır.

Yazarın ne anlama geldiğinin teorik olarak sorgulanması, pratik düzlemde yazara ait sorumluluk olgusunun nasıl çözüleceği sorusunu doğurabilir. Yazarın kimliğinin çözülmesine karşılık eylemin ve sorumluluğun nasıl tanımlanacağı ihtiyacı ortaya çıkar ve bu noktada J.L. Austin'in (1962) söz eylem teorisi önemli bir çerçeve sunar. Dilin pasif bir şekilde dünyayı betimlemekten çok eylemi gerçekleştiren bir araç olduğunu ileri süren bu teoriye göre; "söylemek" aynı zamanda "yapmak" sorumluluğunu da üstlenmektedir. Bu teoriye, konuşmak bir eylemi gerçekleştirmektir: 'Seni tebrik ederim' demek tebrik etme eylemidir, 'Özür dilerim' demek özür dileme eylemini gerçekleştirir. Bu bakış açısı, dilsel üretimi, gerçekleşmiş bir edim olarak konumlandırır konuşmacı veya yazar belirli bir iletişimsel etkiyi hedefleyerek sözcükleri bilinçli olarak seçer düzenler ve sunar. Dolayısıyla dil sadece bir bilgiyi aktarmaz aynı zamanda bir eylemi başlatma veya bir durumu değiştirme gücüne de sahiptir.

Bu çerçevede, bir tehdit mektubundaki "Seni bulacağım" ifadesi yalnızca bir cümle değil, aynı zamanda tehdit etme eylemidir ve hukuk bu eylemi gerçekleştiren faili araştırır. Böylece yazarlığı sadece peşpeşe dizilmiş sözcükler olmaktan çıkarır ve onu niyetli, sonuç doğuran bir iletişimsel performans olarak konumlandırır.

Love (2002), yazarlık sorumluluğunu dört işlev altında somutlaştırarak bu teorik tartışmaları pratik bir zemine taşır: (1) Metnin önceki kaynaklara bağımlılığını ifade eden öncül yazarlık, (2) sözcükleri bir araya getirme eylemi olan yürütücü yazarlık, (3) metnin düzenlenmesini içeren

düzeltilici yazarlık ve son olarak (4) eserin yaratıcısı olarak kamusal alanda görünüp sorumluluğu üstlenen beyan edici yazarlıktır.

Love'ın (2002) işlevsel ayrımı, yazarlığı tek bir eylem yerine katmanlı bir süreç olarak ele alır. Bu modelin sınırları, yapay zekâ tabanlı metin üretimi bağlamında yeniden düşünüldüğünde insan ve algoritma arasında paylaşılan farklı bir hibrit yazarlık yapısı ortaya çıkmıştır. Love'ın işlevsel modelinde insana özgü görülen “yürütücü yazarlık”, günümüzde yapay zekâ tarafından da gerçekleştirilebilen bir işlem hâline gelmiş böylece modelin bazı katmanları insan ve algoritma arasında bölüşülen yapıya dönüşmüştür. Love'ın modelindeki yürütücü işlevi yapay zekâyâ devredip, düzeltilici ve beyan edici işlevler de insan tarafından üstlenildiğinde failin tespiti daha zorlaşır.

Bir tehdit mektubunun taslağını yapay zekâyâ yazdırıp üzerinde bazı değişiklikler yapan fail, dil özelliklerini gizleyebilir. Bu soru adli dilbilimsel delil bağlamında elimize geçen metnin son hali ile önceden geçirdiği süreçlere doğru genişlemektedir. Elimizdeki metin, tek bir yazar tarafından yazılmış bir metin değil içinde hem insan hem de algoritma izleri taşıyan yani insanın yürütücü yazarlık fonksiyonunu üstlendiği yapay zekânın da düzeltilici yazarlık görevini taşıdığı çok katmanlı bir yapı olarak ele alınabilir.

Buna karşın, algoritmaların ürettiği söz dizimsel olarak iyi durumdaki metinler ne Austin'in söz eylem teorisindeki yapma sorumluluğunu ne de Love'ın beyan edici yazarlığını üstlenebilecek bilinçe sahiptir. Dolayısıyla, bir tehdit ve intihar mektubunun faili algoritma olamaz, hukuki bağlamda sözü söyleyen, o sözün arkasındaki toplumsal ve cezai yükü taşıyabilen irade ve niyete (*mens rea*) sahip özne olmalıdır. Yapay zekâ, bu temel hukuki unsurlara sahip olmadığından, Lima'nın (2018, s. 693) da belirttiği gibi ceza hukuku açısından sorumlu bir fail olarak değerlendirilemez ve sorumluluk zinciri mecburen kullanıcıya veya geliştiriciye yönelir (ss. 690–691).

Sonuç olarak, tarihsel süreç içinde ‘tanrısal yaratıcı’dan ‘söylemsel işlev’e oradan da ‘sorumluluk sahibi’ faile dönüşen yazarlık kavramı sabit değil hukuki ve toplumsal ihtiyaçlara göre sürekli yeniden tanımlanan dinamik ve sosyal bir yapıdır. Adli dilbilimin görevi ise bu yeni ve dinamik yapıdaki metnin doğasını anlamak, yazarlık teorilerini değerlendirmek, hukuki sorumluları işaret edebilecek dilsel kanıtları bu karmaşık yapıdan analiz etmektir.

1.2 Adli dilbilimde yazarlık teorileri ve bireydil

Yazarlık kavramının teorik dönüşümü, failin sorumluluğunu belirlemek için teorik bir zemin sunsa da adli dilbilimci için asıl soru bu sorumluluğun metinler üzerinden nasıl kanıtlanacağıdır. Hukuk sistemi metnin arkasındaki yazar işlevini değil biyolojik ve fiziksel bir varlığa sahip faili talep eder. Bu noktada adli dilbilim yazarı soyut bir kavramdan çıkartıp ölçülebilir, değerlendirilebilir dilsel özelliklere sahip bir birey olarak tanımlayan bireydil kavramı üzerine inşa edilir.

Bireydil kavramı ilk olarak Bloch (1948) tarafından kullanılmış olmakla birlikte kavramın kökleri Coleridge'in 19.yüzyıl başlarında yazdığı

Biographia Literaria eserine kadar uzanmaktadır (McMenamin, 2002). Coulthard ve arkadaşlarına (2017, s.15) göre “dilbilimci, tartışmalı yazarlık sorununa her anadil konuşucusunun konuştuğu ve yazdığı dilin kendine özgü ve bireysel bir versiyonuna [kendi bireydiline] sahip olduğu teorik konumundan yaklaşıp.” Coulthard (2004, s.431-432), bu benzersizliği bireyin sözcük seçimleri, eş dizimler, noktalama alışkanlıkları ve dil bilgisel stil gibi dilsel kaynakları kullanma biçimindeki farklılıklarla açıklamaktadır.

Adli dilbilim çalışmalarında sıklıkla kullanılan ‘dilsel parmak izi’ metaforu, her bireyin dilsel üretiminde ayırt edici özellikler bıraktığı ön kabulüne dayanır. Öbür yandan, Coulthard (2004, s.432) bu metaforun yanıltıcı olabileceğini vurgular zira fiziksel parmak izinden farklı olarak herhangi bir dilsel örnek yazarın bireydili hakkında kısmi bilgi sunar. Olsson (2008, s.56-58), bu metodolojik sınırı daha da netleştirerek biyolojik parmak izinin tam bir kesinlikle bireyi işaret etmesine rağmen dilsel parmak izi metaforunun yalnızca olasılıksal bir değerlendirme sunduğunu belirtir. Bu husus, hukukçulara ve jüriye yanlış bir kesinlik algısı verebilir oysa adli dilbilimci ancak “makul şüphenin ötesinde” değil “dengeli olasılıklar” (Olsson, 2008, s.58) çerçevesinde görüş bildirebilir.

Labov’un (1972) sosyodilbilim alanındaki çalışmaları bireylerin dil kullanımında sosyal sınıf, yaş ve bağlama göre sistematik çeşitlilik gösterdiğini ortaya koyulmuş; eğitim, medya, sosyal etkileşim gibi faktörlerin bireylerin dilsel repertuarını sürekli olarak şekillendirdiğini ileri sürülmüştür. Labov’un sosyal varyasyon bulguları bireysel farklılıkları ortadan kaldırmaz sadece bireydilin sabit ve değişmez olmadığını kanıtlar. Bununla birlikte, Coulthard (2004, s.432-433) yeterli uzunluktaki metinlerde özellikle sözcük dizileri ve ifade kalıplarında bireysel benzersizliğin tespit edilebileceğini öne sürer. Adli dilbilim uygulamalarında bu tez, sosyal varyasyonun varlığına rağmen bireydil tabanlı analizin ayırt edici güce sahip olduğunu gösteren somut vakalarla desteklenmiştir.

Bu bağlamdaki en önemli örneklerden biri Amerika Birleşik Devletleri’nde 1978-1995 yılları arasında bir dizi bombalı eylem gerçekleştiren Ted Kaczynski yani Unabomber vakasıdır. 1995 yılında, Unabomber olduğunu iddia eden kişi, “Endüstriyel Toplum ve Geleceği” başlıklı 35000 sözcüklük manifestosunu ulusal yayın kuruluşlarına göndermiştir. Ted Kaczynski’nin kardeşi David Kaczynski yayınlanan metinlerde yer alan “*cool-headed logicians*” (soğukkanlı mantıkçılar) gibi belirli ifadelerdeki kardeşine özgü bireydil özelliklerini tanımış ve yetkililere haber vermiştir. FBI Ajanı James R. Fitzgerald tarafından yapılan analizde on iki dilsel ögenin (*at any rate, clearly, presumably, thereabouts vb.*) sadece Ted Kaczynski’nin metinlerinde bir arada bulunması sosyal varyasyona rağmen bireydil tabanlı eş seçim örüntülerinin adli analiz için yeterli ayırt edicilik sunabileceğini göstermiştir (Coulthard, Johnson & Wright, 2017).

Bir diğer vaka, polis ifadelerinin manipülasyonunu konu alan Derek Bentley davasıdır. Kredens ve Coulthard (2012), Bentley’nin itirafının polis memurları tarafından kısmen yazılıp yazılmadığını araştırmak için derlem dilbilimi yöntemlerini kullanmıştır. Analizde, ifadedeki “*then*” (o zaman)

sözcüğünün kullanım sıklığı ve konumu incelenmiştir. Sıradan tanık ifadelerinde zamansal "*then*" yaklaşık her 500 sözcükte bir görülürken, polis memurlarının ifadelerinde ortalama her 78 sözcükte bir kullanılmaktadır. Bentley'nin ifadesindeki kullanım (her 53 sözcükte bir) açıkça polis ifadeleriyle örtüşmektedir. Dahası, "*I then*" yapısının konumu da önemli bir gösterge olarak bulunmuştur. COBUILD derleminde "*then I*" yapısı "*I then*" yapısından on kat daha sık görülürken, Bentley'nin kısa ifadesinde "*I then*" yapısı üç kez kullanılmıştır. Bu sıklık, 1.5 milyon sözcüklük COBUILD konuşma derlemindeki genel kullanım sıklığının yaklaşık bin katıdır.

Bu vakalar, birey dil tespitinin mümkün olduğunu gösterse de, metodolojik sınırlılıkları da gündeme getirmiştir. Turell (2010, s.238-240) birey dil tartışmalarına çok dilli bir bakış açısı kazandırarak, İspanyolca-Katalanca konuşucularının her iki dilde de benzer sözdizimsel alışkanlıklar sergilediklerini gösteren çalışmasıyla birey dilin sadece tek bir içinde değil diller-arası özellikler de taşıyabileceğini göstermiştir. Wright (2017), birey dil özelliklerinin metin türüne ve bağlamına göre önemli ölçüde değişiklik gösterdiğini, dolayısıyla sadece tek bir metin türünden hareketle yazar profilini çıkartılmasının ciddi riskler taşıdığını vurgular. Sousa-Silva (2024, s. 169) bu metodolojik kaygıyı deneysel verilerle destekleyerek yapay zekâ tarafından üretilen metinlerin analizinde, insan yazarların farklı metin türlerinde bile bireysel izlerini koruduğunu fakat yapay zekâ üretimlerinin bu varyasyondan yoksun olduğunu ortaya koymuştur.

Birey dil tartışmaları, adli dilbilim araştırmalarında özcü ve etkileşimci olmak üzere iki teorik yaklaşım arasındaki gerilimi de barındırır (Grant & MacLeod, 2018, s.81). Özcü yaklaşım, bireyin dilini geçmiş deneyimlerinin ve sosyal aidiyetinin bir kalıntısı olarak görürken, etkileşimci yaklaşım da dili belirli bir kimliğin o anki performansı için kullanılan dinamik bir kaynak olarak değerlendirir (Grant & MacLeod, 2018).

Bucholtz ve Hall (2005, s.588) kimliğin önceden var olan bir kaynak değil etkileşimde ortaya çıkan bir ürün olduğunu vurgular; buna göre kimlik dilsel ve göstergesel pratiklerin kaynağı değil sonucudur ve temelde sosyal ve kültürel bir olgudur. Johnstone'un (1996, s.3-4) "bireysel ses" in sosyal kategorilerin basit bir yansıması olmadığına dair tespiti de etkileşimci bakışını güçlendirir. Bu bakış, adli dilbilimci için şu soruyu gündeme getirir: Eğer dil bağlama göre değişen bir etkileşimci bir performans ise yazarın kimliği metinlerde nasıl tutarlı bir şekilde tespit edilebilir?

Bu teorik durumu aşmak için Grant ve MacLeod (2018, s. 82-83), kimliği "Kaynaklar ve Kısıtlamalar" ekseninde birleştiren bir model önermektedir. Buna göre bireyler, farklı bağlamlarda (kaynaklar) farklı kimlikler sergileyebilirler ancak bilişsel kapasiteleri ve alışkanlık haline gelmiş dilsel rutinleri (kısıtlamalar) bu performansı sınırlamaktadır. Özellikle kişinin kimliğini gizlemeye çalıştığı durumlarda gerçek dilsel alışkanlıklarının sahte kimliğinin altından sızması bu kısıtlamaların somut bir göstergesi olarak görülebilir (Grant & MacLeod, 2018, s. 93-94). Dolayısıyla birey dil, sabit bir öz değil, bireyin dilsel repertuarındaki kısıtlamaların yarattığı tutarlı örüntüler bütünü olarak yeniden kavramsallaştırılabilir.

Wright (2017), bu kavramsallaştırmayı deneysel verilerle destekler: aynı yazarın resmi ve gayri resmi bağlamlarda ürettiği metinler arasında sözcüksel yoğunluk ve cümle karmaşıklığı gibi özelliklerinde belirgin farklar gözlemlenirken, işlev sözcüklerinin kullanım örüntüleri ve belirli birlikte kullanım tercihlerinin bu stilistik varyasyonlar arasında bile tutarlı kaldığını gösterir. Wright'a göre, bu birey dil özellikleri, yazarın bilinçli stilistik manipülasyonuna karşı daha dirençlidir.

Farklı araştırmalar da bu sentezi verilerle desteklemektedir. Kredens ve Coulthard (2012), derlem dilbilimi araçlarını yazarlık analizine uygulayan sistematik bir çerçeve sunar. Bu yaklaşımın gücü, "normalde görünmez olan dil kullanım örüntülerinin, yeterli veriye erişildiğinde açıkça ortaya çıkması"dır (s. 504). Nini (2023, s. 15-16), kullanım-temelli dilbilim yaklaşımıyla geliştirdiği "Dilsel Bireysellik Teorisi"nde, her bireyin deneyimlerinin benzersizliği nedeniyle kendine özgü bir sözcüksel-dilbilgisel sisteme sahip olduğunu öne sürmektedir. Nini, dilsel bireyselliğin sözdizimsel yapılar, söylem özellikleri ve pragmatik örüntüler düzeyinde de aranması gerektiğini öne sürer (Nini, 2023, s. 18-20). Buna benzer bir şekilde Sousa-Silva (2024, s. 169) da yapay zekâ üretimi metinlerin tespiti için ymorfolojik, sözcüksel, sözdizimsel ve söylem düzeylerinde sistematik bir dilbilimsel analiz zorunlu olduğunu belirtmiştir.

Sonuç olarak, Foucault'nun tarihselleştirdiği yazar işlevi, adli bağlamda cezai sorumluluğun belirlenmesi pratikleriyle yeniden üretilmektedir. Adli dilbilimci, Barthes'ın işaret ettiği metnin çok anlamlılığının ve yorumsal risklerin farkında olarak; bireydili sabit bir parmak izi olarak değil, sosyo-bilişsel bir sürecin metindeki izleri olarak değerlendirmelidir. Bu teorik zemin, çalışmanın ilerleyen bölümlerinde ele alınacak olan yapay zekâ destekli metin üretimlerinde yazarlık ve sorumluluk kavramlarının nasıl dönüştüğünü anlamak için de bir başlangıç noktası sunar.

2. Adli Dilbilimde Yazarlık Analizi: Metodolojik Yöntemler ve Kanıt Standartları

Adli dilbilimde yazarlık analizi, belirli bir metnin kimin tarafından üretildiğini belirlemeye yönelik bir inceleme sürecidir. Bu süreç iki temel varsayıma dayanır: yazar-içi tutarlılık ve yazarlar-arası ayırt edicilik (Grant, 2022). İlk varsayım, aynı yazarın farklı metinlerde belirli dilsel özelliklerin tekrar edeceğini öngörür, ikincisi ise farklı bireylerin birbirinden ayrılan dilsel örüntüler sergileyebileceğini ileri sürer. Grant (2022) bu ikili varsayımın birey dil kavramının zeminini oluşturduğunu belirtir fakat her iki varsayımın da koşulsuz ve bağlamdan bağımsız şekilde geçerliliği tartışmaya açıktır.

Bu kavramsal zeminin pratiğe dönüştürülmesi farklı görevler gerektirir. Grant ve MacLeod (2020), bu görevleri üç temel kategoride ele alır: yazar analizi (*authorship analysis*), yazar doğrulaması (*authorship verification*) ve yazar profillemesi (*authorship profiling*). Her kategori farklı varsayımlar ve metodolojik stratejiler içerir dolayısıyla adli dilbiliminin hangi soruyu yanıtlamaya çalıştığını netleştirmesi analiz sürecinin ilk adımındır.

Yazar analizi, tartışmalı bir metnin belirlenmiş aday yazarlar havuzundan hangisine ait olduğunu bulmayı hedefler ve buradaki çalışma kapalı-küme ve açık-küme senaryolarına göre farklılaşır. Kapalı-küme analizinde gerçek yazarın aday listesinde kesinlikle bulunduğu varsayılır; adli dilbilimci adaylar arasından en olası eşleşmeyi seçer. Kopper ve arkadaşları (2009), bu senaryonun metodolojik açıdan görece basit olduğunu fakat adli bağlamda nadiren karşılaşılan bir durum oluşturduğunu belirtir. Açık-küme senaryosunda ise gerçek yazarın aday havuzunda bulunmama olasılığı göz önünde tutulmalıdır böylece hiçbir seçeneğinin de değerlendirilmesini zorunlu kılar. Chaski (2001), gerçek adli vakaların çoğunlukla açık-küme mantığı gerektirdiğini vurgular çünkü soruşturma sürecinde henüz tanımlanmamış şüphelilerin varlığı her zaman olasıdır.

Bir diğer kategori yazar doğrulaması, iki metnin aynı kişi tarafından üretilip üretilmediğini sorgulayan bir görev içerir; bu görev, özellikle anonim tehdit mektupları, sahte itiraflar ve çevrimiçi sahte kimlik vakalarında öneme sahiptir. Koppel ve Schler (2004), yazarlık doğrulamasını tek-sınıf sınıflandırma problemi olarak oluşturmuştur, bilgisayar destekli model yalnızca bilinen yazarın örnekleriyle eğitilir ve tartışmalı metnin bu yazarın dilsel davranış aralığına düşüp düşmediği değerlendirilir. Böylece analiz, "Bu metin X'e mi ait?" sorusuna odaklanır; aday karşılaştırması yapmak yerine, tek bir yazarın dilsel sınırlarını belirlemeye çalışır. Grant ve MacLeod'un (2018) vurguladığı "kısıtlamalar" kavramı burada işlevsel hale gelir çünkü bireyin bilişsel ve dilsel rutinleri, performans varyasyonuna rağmen belirli sınırlar içinde kalır ve bu sınırların dışına çıkan metinler farklı bir yazara işaret edebilir.

Son kategori yazar profillemeye ise yazarın kimliğini doğrudan belirlemeye çalışmak yerine, metinden demografik veya sosyodilbilimsel özellikler çıkarmayı amaçlar. Argamon ve arkadaşları (2009), cinsiyet, yaş ve kişilik özelliklerinin metin içi izlerini araştırmış, belirli dilsel değişkenlerin bu sınıflarla ilişki gösterdiğini ortaya koymuştur. Öte yandan, profillemeye sonuçlarının bireysel düzeyde öngörü gücü sınırlıdır. Bucholtz ve Hall'un (2005) vurguladığı gibi kimlik, dilsel pratiklerin kaynağı değil sonucudur dolayısıyla sosyal kategorilerin mekanik bir şekilde dilsel özelliklere indirgenmesi, hem teorik hem de metodolojik riskler taşır. Eckert (2012), toplumsal cinsiyetin sabit bir değişken değil, söylem içinde performatif olarak üretilen bir kimlik boyutu olduğunu belirtir; bu durum, profillemeye sonuçlarının gerekircilik olgusundan kaçınmasını gerektirir.

Yazar analizi ve yazar doğrulaması görevlerinin başarıyla yerine getirilmesi, ortak bir metodolojik ilkeye bağlıdır: karşılaştırılabilirlik. Coulthard'ın (2004) işaret ettiği gibi, karşılaştırılan metinlerin tür, bağlam, uzunluk ve üretim koşulları açısından yeterince benzer olması gerekir. Bir iş e-postası ile kişisel bir günlük yazısı arasındaki stilistik farklılıklar, yazarlık değil bağlama göre uyum gösterebilir. Wright'ın (2017) bulguları ise bu varyasyonun belirli özellikler için sınırlı kaldığını gösterir: Resmi ve gayri resmi bağlamlarda sözcüksel yoğunluk değişse bile, işlev sözcüklerinin kullanım örüntüleri görece tutarlı kalır. Bu bulgu, adli dilbilimcinin hangi

dilsel düzeylere odaklanması gerektiği konusunda metodolojik bir yönlendirme sunar.

Ancak karşılaştırılabilirlik ilkesinin sağlanması ve görevin doğru tanımlanması tek başına yeterli değildir; bu verileri analiz etmek için sağlam yöntemlere de ihtiyaç vardır. Adli dilbilim, bu amaçla iki tamamlayıcı metodolojik yaklaşım geliştirmiştir: uzman yorumuna dayanan stilistik analiz ve sayısal kanıta dayanan stilometrik ölçümler.

2.1 Metodolojik yöntemler: Stilistik ve stilometrik yaklaşımlar

Yazarlık analizinde iki temel metodolojik gelenek birbirini tamamlar: nitel stilistik analiz ve nicel stilometrik ölçümler. Bu iki yaklaşımın kendi içlerindeki gelişimi, adli dilbilimin yöntemlerini şekillendirmiş ve günümüzde uyum arayışlarına yol açmıştır.

Stilistik analiz, uzman dilbilimcinin metni yakın okuma yoluyla incelemesine ve dilsel seçimlerin örüntülerini yorumlamasına dayanır. Crystal ve Davy'nin (1969) klasik tanımıyla stil (üslup), belirli bir bağlamda yapılan dilsel seçimlerin toplamıdır ve bu seçimler rastgele değildir, belirli işlevleri yerine getirip bireysel alışkanlıkları yansıtır.

McMenamin (2002), adli stilistik analiz için merkezinde stil-işaretleyicileri kavramı olan bir çerçeve sunmuştur. McMenamin, bunları sınıf-işaretleyicileri ve bireysel-işaretleyiciler olarak ikiye ayırır: sınıf-işaretleyicileri, yazarın ait olduğu sosyal grupların dilsel özelliklerini yansıtır (bölgesel lehçe özellikleri, mesleki jargon, eğitim düzeyi göstergeleri), bireysel-işaretleyiciler ise o yazara özgü örüntüleri gösterir (belirli bağlaçların tercihi kullanımı, karakteristik noktalama alışkanlıkları, sözdizimsel yapı tekrarları). Adli açıdan bireysel-işaretleyiciler daha değerlidir çünkü bunlar yazarı sosyal kategorisinden ayırt etmeyi sağlar.

Olsson (2004), stilistik özellikleri "temel" ve "katma" olarak sınıflandırır. Temel özellikler, yazarın farkında olmadığı veya kontrol edemediği dilsel alışkanlıklardır ve bilinçli yönlendirmeye karşı görece dirençlidirler. Katma özellikler ise belirli iletişim durumlarında kasıtlı olarak eklenen stilistik unsurlardır. Bu ayrım, Grant ve MacLeod'un (2018) "kısıtlamalar" kavramıyla örtüşür: temel özellikler, bireyin bilişsel ve dilsel rutinlerinin yarattığı kısıtlamaların ortaya çıkmasıdır ve performans varyasyonuna rağmen tutarlılık gösterir. Özellikle kimliğini gizlemeye çalışan yazarların gerçek dilsel alışkanlıklarının sahte kimliğin altından sızması, bu kısıtlamaların somut bir göstergesidir.

Stilistik analizin gücü, dilsel seçimlerin işlevsel yorumunu mümkün kılmasıdır. Sadece sıklık sayımları bir özelliğin ne sıklıkta görüldüğünü söyleyebilir; ancak bu özelliğin neden o bağlamda tercih edildiğini açıklayamaz. Bir metinde "ancak" bağlacının yüksek kullanım oranı, yazarın karşıtlık ilişkilerini vurgulama eğilimini gösterebilir ancak bu eğilimin söylem stratejisi olarak nasıl işlev gördüğü ise nitel yorumla anlaşılır.

Bununla birlikte, stilistik analiz bazı sınırlılıklar taşır. Birincisi, uzman yargısına dayandığı için öznel yanlılığa açıktır ve farklı analistler aynı metinde farklı stil-işaretleyicileri tespit edebilir. İkincisi, aynı metni inceleyen

iki uzmanın aynı sonuçlara ulaşacağı garanti değildir o yüzden tekrarlanabilirlik sorunu vardır. Sonuncusu, kısa metinlerde yeterli stilistik verisi bulunmayabilir o yüzden, çevrimiçi iletişimin yaygınlaştığı günümüzde sorunludur. Olsson'un (2008) vurguladığı gibi, adli dilbilimcinin ancak "olasılıksal bir değerlendirme" sunabileceği gerçeği, stilistik analizin sonuçlarının sunumunda kesinlikten kaçınarak, hata payını kabul ederek, yöntemsel şeffaflık sunmak gerektirir.

Öte yandan stilometri de dilsel özelliklerin istatistiksel analizine dayanan nicel bir yaklaşım geleneğidir. Mosteller ve Wallace'ın (1964) *The Federalist Papers* çalışması, modern stilometrinin kurucu metni sayılır. Hamilton ve Madison arasında tartışmalı olan makalelerin yazarlığını belirlemek için işlev sözcüklerinin sıklık dağılımlarını kullandılar. İşlev sözcükleri içerik sözcüklerine kıyasla konudan bağımsızdır ve bilinçdışı kullanıma daha yatkındır varsayımı temel alınmıştır. Örneğin, "upon" ile "on" arasındaki tercih, yazarın bilinçli bir kararından çok otomatik bir alışkanlık yansıtabileceği düşünülmüştür.

Burrows (2002), Delta yöntemini geliştirerek stilometrik karşılaştırmayı sistemleştirmiştir. Delta, tartışmalı metnin özellik vektörü ile aday yazarların özellik vektörleri arasındaki standartlaştırılmış mesafeyi ölçer ve en düşük değer, en olası yazarlığa işaret eder. Bu yöntem, derlem dilbiliminde geniş kabul görmüş ve farklı dillerde başarılı uygulamalar bulmuştur. Ancak Delta'nın performansı, derlem büyüklüğüne ve metin uzunluğuna (Eder, 2015; Hoover, 2004), ayrıca yazar havuzunun türdeşliğine ve büyüklüğüne (Evert vd., 2017; Rybicki & Eder, 2011) bağlıdır. Bu değişkenler kontrol edilmediğinde ve parametreler iyileştirilmediğinde sonuçların güvenilirliği azalmaktadır (Smith & Aldridge, 2011).

Karakter n-gram'ları da stilometrik çalışmalarda kullanılarak ardışık karakter dizilerinin frekans profillerinin yazar analizinde yüksek doğruluk oranlarına erişmeye katkı sağlamıştır (Keselj vd., 2003). Makine öğrenmesiyle birlikte destek vektör makineleri (SVM), rastgele ormanlar ve derin öğrenme modelleri, yüksek boyutlu özellik uzaylarında örüntü tanıma amacıyla kullanılmaya başlanmış; Stamatatos (2009) bu yöntemlerin karşılaştırmalı değerlendirmesini sunmuştur.

Ancak makine öğrenmesi tabanlı yaklaşımlar, adli bağlamda 'kara kutu' sorunuyla karşı karşıyadır. Bu modeller yüksek doğruluk oranları sunsa da, Sousa-Silva'nın (2024) belirttiği gibi, analiz edilen olgular için teorik olarak temellendirilmiş dilbilimsel açıklamalar sunmakta yetersiz kalabilmektedir. Bir model, belirli bir yazarlık kararına nasıl ulaştığını insan muhakemesine uygun biçimde açıklayamıyorsa, mahkemede kabul edilebilirliği sorgulanabilir (Grant, 2022).

Stilistik ve stilometrik yöntemler arasındaki bu durum, aslında bizi zorunlu bir senteze götürüyor çünkü istatistiksel veriler tek başına anlamı vermezken, sadece uzman yorumu da nesnellikten uzak kalabiliyor. Bu noktada Nini (2023), özellikleri geleneksel stil katmanları üzerinden birleştirmek yerine, Bilişsel Dilbilim ilkelerine dayanan ve yazarlığı bireyin belleğinde depolanan, otomatik olarak üretilen dilsel birimler kümesi olarak

tanımlayan ‘Dilsel Bireysellik Teorisi’ni öne sürmektedir. Nini, bu teoride karakter, sözcük, sözcük türü (POS) ve n-gram gibi farklı düzeylerini kullanarak, istatistiksel profillerin aslında bireyin bireydilinin istatistiksel bir yaklaşımı olduğunu savunur. Bu bilişsel tabanlı çerçeve, analiz yöntemlerinin yorumlanamayan birer 'kara kutu' olmaktan çıkmasını sağlayarak, Grant’ın (2022) üzerinde önemle durduğu yöntemlerin dilbilimsel açıdan açıklanabilir ve bilimsel olarak gerçekçi olmasını temin etmektedir

Ancak bir yöntemin bilimsel olarak açıklanabilir olması tek başına yeterli değildir; bu açıklamanın hukuk sistemi tarafından da geçerli ve güvenilir bulunması gerekir. İşte bu noktada, bilimsel kanıtın mahkeme salonunda delil niteliği kazanmasını belirleyen küresel ölçekte referans kabul edilen belirli hukuki standartlar, yani Daubert kriterleri devreye girmektedir.

2.2 Kanıt standartları: Daubert kriterleri ve bilimsel gerçeklik

Adli dilbilimsel bir analizin delil olarak değeri, teknik başarısından ziyade hukuk sistemi tarafından geçerli ve güvenilir bulunmasına bağlıdır. ABD Yüksek Mahkemesi’nin 1993 tarihli *Daubert v. Merrell Dow Pharmaceuticals* kararı, daha önceki *Frye* standardını düzenleyerek bilimsel kanıtın kabul edilebilirliği için uluslararası düzeyde referans alınan beş temel ölçüt sunmuştur.

Bu ölçütlerin ilki ve bilimselliğin temel şartı, test edilebilirlik ve yanlışlanabilirlik ilkesidir. Bir yöntemin mahkemede kabul görmesi için, dayandığı hipotezlerin deneysel olarak sınanabilir olması gerekir. Örneğin Chaski (2001), geliştirdiği ALIAS yöntemiyle bu standardı karşılamayı hedeflemiş ve sistemin performansını kör testlerle raporlamıştır. Ancak adli stilistikteki birçok uygulama açıkça formüle edilmiş ve test edilebilir hipotezlerden ziyade, uzmanların kişisel yorumlarına dayanmaktadır, bu durum da yöntemin bilimsellik iddiasını zayıflatan bir olabilmektedir.

Bu test süreçlerinin şeffaflığı, zorunlu olarak ikinci kriter olan hakemli yayın ve bilimsel inceleme sürecini beraberinde getirir. Yöntemin kapalı kapılar ardında değil, bilim topluluğunun denetimine açık mecralarda tartışılması beklenir. *Forensic Linguistics* ve *International Journal of Speech, Language and the Law* gibi alan dergileri bu işlevi üstlense de, Solan ve Tiersma’nın (2005) belirttiği gibi, adli dilbilim literatürü geleneksel adli bilimlere kıyasla hala sınırlıdır ve bazı yöntemler yeterli derecede bağımsız çoğaltma süzgecinden geçmemiştir.

Bilimsel denetimin doğal bir sonucu olarak mahkeme, üçüncü adımda yöntemin bilinen veya potansiyel hata oranını talep eder ancak yazar analizi için evrensel bir hata oranı belirlemek, yöntemin doğası gereği oldukça zordur çünkü başarı oranı metnin uzunluğuna, türüne ve yazar havuzunun büyüklüğüne göre değişiklik gösterir. Juola (2006), farklı stilometrik yöntemlerin performansını karşılaştırsa da bu testler kontrollü laboratuvar koşullarında yapılmıştır. Oysa gerçek adli vakalar, eksik veriler, kasıtlı kandırmacalar ve bağlamsal belirsizliklerle doludur. Bu belirsizlik, Olsson’un (2008) da vurguladığı gibi, adli dilbilimcinin "makul şüphenin ötesinde"

kesinlikten ziyade, ancak "dengeli olasılıklar" çerçevesinde görüş bildirebileceği gerçeğini doğurur.

Hata paylarını azaltmak ve uygulama birliğini sağlamak adına dördüncü kriter, standartların varlığı ve uygulanmasını gözetir. *International Association of Forensic Linguists (IAFL)* etik kodlar ve uygulama kılavuzları geliştirmiş olsa da, yazar analizi için üzerinde mutabık kalınmış tek standartlar mevzuatı henüz mevcut değildir. Bu metodolojik çoğulculuk hali, mahkemelerde sıkça rastlanan uzmanların çatışması senaryosuna zemin hazırlayabilir çünkü aynı metin üzerinden farklı yöntemlerle çalışan uzmanlar, birbirine zıt sonuçlara ulaşabilir.

Sonuçta mahkeme, eski *Frye* standardından miras kalan genel kabul ilkesine bakar. Grant (2022), adli yazarlık analizinin bir disiplin olarak dilbilim topluluğunda kabul gördüğünü, ancak belirli yöntemlerin geçerliliği konusunda hala görüş ayrılıkları bulunduğunu belirtir. Özellikle makine öğrenmesi tabanlı yaklaşımların şeffaflıktan uzak "kara kutu" yapısı, bu genel kabul sürecini hukuki düzlemde daha da karmaşıktırmaktadır.

Bu kriterlerin ötesinde, kanıtın sunum biçimi de önem taşır. Adli dilbilimci, bulgularını sunarken "Kesinlikle bu yazar" gibi keskin ifadelerden kaçınmalı; bunun yerine Shuy'un (2006) ilkesi doğrultusunda "Eldeki veriler ışığında en olası yazar" ifadesini tercih ederek yöntemin sınırlılıklarını ve alternatif senaryoları da açıkça belirtmelidir. Bu bakış açısı, Tiersma ve Solan'ın (2002) dikkat çektiği, jüri üyelerinin olasılıksal ifadeleri yanlış yorumlama riskini (örneğin "%90 olasılık" ifadesini "kesin suçluluk" olarak algılama eğilimini) yönetmek için bir zorunluluktur. Sonuçların ne anlama geldiği hatta daha önemlisi ne anlama gelmediğini açıklamak, uzmanın sorumluluğudur.

Bu noktada Morrison (2009), adli kıyaslama bilimlerinde bilimsel güveni ve mantıksal tutarlılığı sağlamak için Olasılık Oranları (*Likelihood Ratios - LR*) çerçevesini önermektedir (s. 298). Bu yaklaşımda uzman, doğrudan suçluluk veya masumiyet iddiasında bulunmak yerine kanıtın olasılığını değerlendirir (Morrison, 2009). Uzman, mevcut kanıtın (örneğin dilsel bir eşleşmenin), iddia makamının hipotezi (aynı kaynak) altında mı yoksa savunmanın hipotezi (farklı kaynak) altında mı ortaya çıkma olasılığının daha yüksek olduğunu karşılaştırır (Morrison, 2009, s. 299). Örneğin, "Bulgular, bu kanıtın aynı köken hipotezi altında gözlemlenme olasılığının, farklı köken hipotezine kıyasla 100 kat daha fazla olduğunu göstermektedir" şeklindeki bir ifade, kanıtın gücünü istatistiksel bir zemine oturtur (Morrison, 2009, s. 300). Bu yöntem, Shuy (2006) tarafından detaylandırılan adli dilbilimcinin etik sorumluluğunu ve bilimsel tarafsızlığını, ölçülebilir ve kurumsal bir standarda dönüştürmektedir (s. 126).

Yapay zekâ tarafından üretilen metinlerin tespitinde ise durum çok daha farklılaşır. Wright ve Coulthard'ın (2017) metodolojik geçerlilik üzerine yaptıkları ayırım, bugün YZ tespit araçları için de geçerliliğini korumaktadır: Bir tekniğin bilimsel olarak "ilginç" olması ile hukuken "kabul edilebilir" olması tamamen farklı şeylerdir.

Günümüzde yapay zekâ yazarlığını tespit etmek amacıyla kullanılan araçlar DetectGPT (Mitchell vd., 2023) ve GLTR (Gehrmann vd., 2019) gibi yaklaşımlar şu an için yalnızca bilimsel olarak ilginç kategorisindedir ancak Daubert kriterlerini karşılamada metodolojik sınırlılıklar barındırmaktadır. Bu araçların test edilebilirlikleri her model güncellemesinde sıfırlanır, hata oranları yüksek ve evrensel standartları belirsizdir. Dolayısıyla, bir adli raporda yer alacak "Bu metni yapay zekâ yazmıştır" yönündeki bir tespit, kullanılan araçlar metodolojik olgunluğa erişmediği sürece, mahkeme nezdinde geçerli bir delil hükmü kazanamaz.

Hukuki düzlemde yaşanan bu sorunun asıl sebebi, kullanılan araçlardaki teknik bir eksiklikten daha çok fail konumundaki yapay zekânın metin üretim şeklinin insan yazarlığından doğası gereği tamamen farklı olmasıdır. Bu farkı anlamak ve doğru tespit yöntemleri geliştirebilmek için, öncelikle yapay zekânın yazarlık kavramını nasıl dönüştürdüğüne ve insan bireydidinden ayrılan yapay bireydidil olgusuna yakından bakmak gerekmektedir.

3. Yapay Zekâ ve Yazarlık: Yapay Bireydidil

Adli dilbilimin üzerine inşa edildiği "yazar-îçi tutarlılık" ve "yazarlar-arası ayırt edicilik" varsayımları, insan zihninin ve yaşanmış deneyimin ürünü olan dilsel izleri sürmek için geliştirilmiştir. Buna karşın, büyük dil modellerinin metin üretim kapasitesi, bu varsayımların geçerlilik alanını sarsmaktadır. Love'ın (2002) işlevsel modelinde uzun süre yalnızca insana özgü bir yetenek olarak görülen ve sözcükleri bir araya getirme eylemini tanımlayan "yürütücü yazarlık" artık algoritmalar tarafından da gerçekleştirilebilmektedir.

Bu dönüşüm, adli dilbilimcinin soruşturduğu soruyu geleneksel "Bu metni kim yazdı?" sınırlarının ötesine taşımıştır. Artık karşımızda çok daha katmanlı ve karmaşık sorular durmaktadır: "Bu metin insan üretimi mi yoksa makine çıktısı mı?" ve daha da önemlisi, "İnsan-makine işbirliğiyle üretilen hibrit metinlerde sorumluluk kime aittir?". Bu bölüm; yapay zekâ tabanlı metin üretiminin altında yatan mekanizmaları, bu üretimin niyet ve sorumluluk kavramlarıyla gerilimli ilişkisini ve insan bireydidilinden kategorik olarak ayrılan "yapay bireydidil" olasılığını incelemektedir.

3.1 Yapay zekânın metin üretim mekanizması: "Stokastik papağanlar"

Hukuki ve yapısal ayrımın temelinde, büyük dil modellerinin çalışma prensibinin insan bilişinden kategorik olarak ayrışması yatar. GPT serisi ve BERT gibi transformer tabanlı modeller, devasa metin derlemelerinden öğrendikleri istatistiksel örüntülerle bir sonraki sözcüğü tahmin ederek ilerler. Vaswani ve arkadaşlarının (2017) geliştirdiği dikkat mekanizması, modelin uzun menzilli bağlamsal tutarlılık sağlamasına olanak tanısa da, üretim süreci özünde olasılıksaldır. Bender ve arkadaşları (2021), bu mekanizmayı "stokastik papağanlar" metaforuyla tanımlar: Modeller, eğitim verilerindeki dilsel formları yeniden üretir ancak bu formların anlamını veya dış dünyayla gönderimsel ilişkisini kavramaz.

Floridi ve Chiriatti'nin (2020) belirttiği gibi, bu sistem bir "anlam motoru" değil, istatistiksel bir "sözdizimsel motor"dur. Model, "yangın" sözcüğü ile "tehlike" sözcüğünün istatistiksel birlikteliğini hesaplar ancak yangının fiziksel gerçekliğini deneyimlemez. Bu hal, Austin'in (1962) Söz Eylem Kuramı bağlamında önemli bir sonuç doğurur yani model, biçimsel olarak kusursuz cümleler üretebilir ancak bu sözlerin arkasında bir iletişimsel niyet veya sorumluluk bilincine sahip değildir.

Ancak adli dilbilim açısından, buradaki temel mesele modellerin anlayıp anlamadığı felsefi sorunu değildir; asıl önemli olan, ürettikleri metinlerin insan yazarlığından ayırt edilebilir olup olmadığıdır. Kumarage ve arkadaşlarının (2024) işaret ettiği gibi, modellerin taklit yeteneği geliştikçe stilometrik tespit zorlaşmaktadır. Bununla birlikte, modellerin istatistiksel üretim doğası, insan bilişinin organik süreçlerinden yapısal olarak ayrıştığı için tespit edilebilir izler bırakır.

Sousa-Silva'nın (2024) güncel araştırmaları, bu ayrımı somut ve ölçülebilir bir zemine oturtmaktadır. Sousa-Silva, üretken yapay zekâ sistemleri tarafından oluşturulmuş öğrenci metinlerinden meydana gelen yaklaşık 31.500 sözcüklük *NewGenerAltion* derlemi (2024, s. 168-169) incelemiş ve yapay zekâ metinlerinin insan yazımındaki doğal çeşitliliğin aksine yapay bir aşırı düzenlilik sergilediğini raporlamıştır. İnsan yazarlar cümle uzunluklarında ve sözcük çeşitliliğinde dalgalanmalar gösterirken, yapay zekânın istatistiksel bir homojenlik sergilediği saptanmıştır. Modeller, hiyerarşik düşünceyi yansıtan yan cümle yapıları yerine, bilişsel yükü daha az olan "ve" bağlacıyla kurulmuş sıralı cümleleri sistematik olarak tercih etmektedir. Sözcüksel düzeyde ise *delve*, *underscore*, *nuanced* ve *paramount* gibi belirli sözcüklerin yapay zekâ tarafından insan ortalamasının çok üzerinde bir sıklıkla kullanılması, stokastik üretim doğasının parmak izleri olarak değerlendirilmektedir. Bu sözdizimsel ve sözcüksel aşırı düzenlilik kalıpları, niyetsiz üretimi tespit etmek için metodolojik bir dayanak noktası sunar.

Modelin iletişimsel niyetten yoksun bu stokastik doğası, adli dilbilimin üzerine kurulu olduğu niyet, eylem, sorumluluk üçlüsünde bir kırılma yaratmaktadır. Eğer bir metin, arkasında bilinçli bir zihin olmadan da üretilebiliyorsa, o metnin doğurduğu suçun faili kimdir sorusu bizi yapay zekânın fail olamama durumunu ve hibrit metinlerdeki sorumluluk karmaşasını incelemeye götürür.

3.2 Niyet, eylem, sorumluluk

Austin'in (1962) Söz Eylem Kuramı'na göre "söylemek" aynı zamanda "yapmak"tır ve bu edim, konuşanın sorumluluğunu doğurur (s. 12). Ancak bir dil modeli edimsel güce sahip sözler (örneğin bir tehdit, iftira veya intihal) ürettiğinde, bu eylemin hukuki muhatabını tayin etmek teorik bir soruna dönüşür. Hallevy (2010), yapay zekânın ceza hukuku açısından sorumlu bir fail olarak kabul edilmesinin önündeki en büyük engelin mens rea (suç işleme kastı) olduğunu vurgular (s. 175). Ceza hukuku, suçun maddi unsuru (*actus reus*) kadar manevi unsurunu da şart koşar ancak dil modelleri

bilinç, irade veya niyet taşımadığından, ürettikleri metinlerin doğrudan faili sayılamazlar.

Bu bağlamda sorumluluk zinciri, kaçınılmaz olarak modeli kullanan kişiye veya geliştiren kuruma yönelmektedir. Ancak eylemin sınırlarının belirsizliği, geleneksel sorumluluk atfını güçleştirmektedir. Floridi (2013), bu tür durumların yalnızca insan odaklı bir yaklaşımla açıklanamayacağını savunarak "dağıtılmış ahlak" kavramını önerir (s. 728). Yazara göre, gelişmiş bilgi toplumlarında ortaya çıkan ahlaki eylemler ve sorumluluklar, tek bir bireye indirgenemez aksine, bu eylemler çoklu etmen sistemlerindeki etkileşimlerin bir sonucu olarak "bireysel olmayan sorumluluklar" doğurur (Floridi, 2013, s. 728-729). Floridi'nin (2013) tanımladığı bu dağınık yapı, suçu tek bir bireyin bireyiline atfetmeye ve faili tekil olarak belirlemeye çalışan adli dilbilimci için ciddi bir metodolojik engel teşkil etmektedir.

Bu karmaşıklığı çözümlmek adına Gunkel (2016), geleneksel "tekil yazar" anlayışının yerine "remix" kavramını önerir (s. 15). Gunkel'e göre dijital çağda ve özellikle yapay zekâ üretimlerinde özgün bir yaratımdan ziyade, mevcut verilerin yeniden derlendiği dağınık bir süreç söz konusudur. Bu bakış açısı, Love'ın (2002) işlevsel yazarlık modelinin insan ve makine arasında paylaşıldığı hibrit bir yapıyı teorik olarak destekler (s. 34). Ancak bu işlevsel dağılım, uygulamada üç farklı hibrit senaryo doğurur ve her biri, Love'ın tanımladığı yazarlık işlevlerinin (öncül, yürütücü, düzeltici, beyan edici) farklı biçimlerde dağılmasına neden olur:

Yapay Zekâ Destekli Yazarlık (Senaryo 1): İnsan bir istem yazar, büyük dil modelleri metni üretir, insan düzenler ve yayınlar. Bu senaryoda insan; öncül ve beyan edici (sorumluluk üstlenme) yazarlık işlevlerini elinde tutarken, sözcükleri dizme eylemi olan yürütücü işlevi makineye devreder. Ancak insan, nihai aşamada "düzeltici" işleviyle metne son şeklini verdiği için faillik hukuken insana atfedilebilir. Burada adli analiz, kullanıcının bu düzeltici izlerine odaklanmak durumundadır.

Yapay Zekâ Aracılı Geliştirme (Senaryo 2): İnsan bir taslak yazar, büyük dil modelleri bu taslağı genişletir. Bu süreçte yürütücü yazarlık işlevi insan ve makine arasında erir; hangi cümlelerin insanın öncül taslağından geldiği, hangisinin makine tarafından yürütüldüğü belirsizleşir. Adli dilbilimsel analiz, iç içe geçmiş bu iki yürütücü katmanı ayırtmakta zorlanır.

Stil Transferi ve Yanıltma (Senaryo 3): İnsan, metni yapay zekâyı yazdırırken başka bir yazarın üslubunu taklit etmesini ister. Burada öncül yazarlık katmanlaşır kullanıcının niyeti birincil öncülken taklit edilen yazarın üslubu ikincil ve pasif bir kaynak haline gelir. Makine ise bu verileri işleyen yürütücü konumundadır. Bu üç katmanlı yapı, geleneksel analiz yöntemlerinin tespit gücünü zayıflatır. En kritik durum ise kullanıcının kasıtlı olarak kendi birey dilini gizlemek için yapay zekâyı araç olarak kullanmasıdır.

Bu senaryoların tümünde, adli dilbilimsel analiz için ortak bir metodolojik zorluk ortaya çıkar. McMenamin'in (2002) stilistik analiz çerçevesi, metindeki tutarlılık kırılmalarını inceleyerek bu katmanları ayırtmayı hedeflese de, bu yorumlar öznel kalma riski taşır. Nini'nin (2023, s. 45) istatistiksel modelleri birey-içi ve bireyler-arası varyasyonu ayırt

edebilmektedir. Denklemi "birey-makine arası varyasyon" eklendiğinde analiz karmaşıklaşır çünkü yapay zekâ modelleri insan yazımını taklit eden yüksek kaliteli metinler üretse de (Kumarage et al., 2024), bu üretim bilişsel bir dilbilgisi veya niyetten ziyade (Nini, 2023), eğitim verisindeki olasılıksal dağılımlara dayanır (Kumarage et al., 2024)

En önemli sorun ise kasıtlı yanıltma durumlarında ortaya çıkar. Bir kullanıcı, yapay zekâyı kendi birey dilini gizleyen bir araç olarak kullandığında, üretilen metnin yazarının kimliklendirilmesi problemi derinleşir. Bu noktada Grant ve MacLeod'un (2018) kısıtlamalar teorisinin geçerliliği tartışılabilir: insan yazarın bilişsel ve dilsel rutinlerinin yapay zekâ aracılığıyla tamamen maskelenip maskelenemeyeceği veya kullanıcının istemlerinin yeni bir iz katmanı oluşturup oluşturmayacağı, adli dilbilimin yanıltması gereken temel mesele haline gelir.

İnsan yazarın kendi birey dil özelliklerini gizleme olasılığı, adli dilbilimcinin dikkatini kaçınılmaz olarak metnin diğer kurucusuna, yani algoritmaya çevirmesini zorunlu kılar. Zira insan birey dilinin silindiği veya silikleştiği noktada, sahneye istatistiksel verilerden inşa edilmiş 'yapay birey dil' çıkmaktadır.

3.3 Yapay birey dil

İnsan yazarın dilsel kısıtlamalarını yapay zekâ aracılığıyla aşabilmesi, metinsel boşluğu dolduran dilsel örüntünün kaynağına dair temel bir sorunu işaret eder. Bu noktada, insan deneyiminden bağımsız gelişen ve adli analizde yeni bir değişken olabilecek "yapay birey dil" kavramını tartışmak zorunludur.

Geleneksel birey dil kavramı, her insanın benzersiz yaşam deneyiminin benzersiz bir dilsel repertuvar yarattığı varsayımına dayanır (Coulthard, 2004, s. 432). Ancak yapay zekâ üretimlerinde bu benzersiz deneyim yerini kolektif veriye bırakır. GPT-3 modeli, binlerce farklı kaynaktan, yazardan ve dönemden derlenen yaklaşık 570 gigabaytlık devasa bir metin yığınıyla eğitilmiştir (Brown vd., 2020, s. 9). Bu halde modelin dili, herhangi bir bireyin diline benzemez daha ziyade, eğitim verilerinin istatistiksel bir ortalamasını ve kolektif bir dilsel belleği temsil eder.

Bu bağlamda Sousa-Silva (2024), yapay zekâ üretimlerinin doğasına ilişkin önemli bir tespitte bulunur: Büyük dil modelleri olasılıksal hesaplamalara dayandığı için, ürettikleri metinler "bireysel yazarlık özelliklerinden yoksundur" (s. 1028) der. İnsan birey dili, niyetlerin, bilişsel süreçlerin ve sosyal etkileşimlerin dinamik bir ürünüyken Sousa-Silva'nın vurguladığı gibi yapay zekâ üretimi, bireysel izlerden arındırılmış istatistiksel bir çıktı olabilmektedir. Bu durum, yapay birey dil kavramını varoluşsal bir paradoksa dönüştürür çünkü karşımızda dilsel bir üretim vardır, fakat bu üretimin arkasında sosyodilbilimsel bir birey bulunmamaktadır. Bu paradoksun yanıtı, yapay birey dil kavramının insan birey dilinden kategorik olarak farklı bir nitelik taşımasında yatar.

Yapay birey dil, insan birey dilinin aksine sosyolojik bir geçmişe veya bilişsel bir niyete sahip olmayan (Bender vd., 2021; Nini, 2023), büyük dil modellerinin eğitim verilerindeki hâkim dilsel örüntüleri istatistiksel olarak

yeniden üretmesiyle oluşan (Bender vd., 2021), aşırı düzenli, düşük değişkenli (O'Sullivan, 2025; Sousa-Silva, 2024) ve model tabanlı ayırt edici özellikler taşıyan (Uchendu vd., 2020) bireysel yazarlık özelliklerini taşımayan sentetik bir dilsel repertuvar olarak tanımlanabilir.

Yine de bu bireysizlik hali, metinlerin tamamen ayırt edilemez olduğu anlamına gelmez. Uchendu ve arkadaşları (2020), farklı mimarilere sahip modellerin (örneğin GPT-2 ile GROVER) üretimlerinin birbirinden ayırt edilebileceğini göstererek, model bireydili kavramını düşünmemizi sağlamıştır (s. 8388). Bu bulgular, her model, eğitim verisinin, mimari tasarımının ve değişkenlerin birleşimi olarak kendine özgü, tutarlı fakat yapay bir stilistik bir imza taşıdığını gösterip model birey dil kavramının varlığını desteklemektedir.

Adli analizi en çok zorlayan senaryo ise insan ve model izlerinin birbirine geçtiği hibrit metinlerdir. Nini'nin (2023) "Dilsel Bireysellik Teorisi", bireyin dilsel sisteminin tutarlılığını savunsa da, hibrit metinlerde bu sistem modelin algoritmik yapısıyla bozulur (s. 4-5). Ancak Wright'ın (2017) bulguları burada metodolojik bir dayanak sunar: İşlev sözcükleri ve eşdizim örüntüleri, bilinçli manipülasyona karşı dirençlidir (s. 227). Kullanıcı metni yapay zekâya yazdırsa bile, yaptığı küçük düzenlemelerde kendi işlev sözcüğü kullanım alışkanlıklarını sızdırabilir.

Ne var ki, eğer kullanıcı model çıktısını hiç değiştirmeden kullanıyorsa, insan izleri tamamen silinir. Solaiman ve arkadaşları (2019), bu durumu model çıktılarının yazar atıf edilemezliği sorunu olarak tanımlar (s. 13). Modelin stokastik doğası, aynı istemle her seferinde farklı çıktılar üretmesine neden olur bu da belirli bir metnin belirli bir modele ve dolayısıyla o modeli o anda kullanan belirli bir kullanıcıya kesin olarak atfedilmesini, geleneksel yöntemlerle imkânsız hale getirebilir.

İnsan algısının ve klasik stilistik ve stilometrik yöntemlerin yetersiz kalabildiği bu noktada, çözüm arayışı makineyi makineyle yakalamayı hedefleyen otomatik tespit sistemlerine kaymıştır.

3.4 Yapay zekâ tespit sistemleri ve sınırlılıklar

Yapay zekâ tarafından üretilen metinlerin tespiti, hem akademik hem de ticari alanda yoğun ilgi gören bir araştırma alanı haline gelmiştir. Bu sistemler, insan ve makine yazımını istatistiksel olarak ayırt etmeyi hedeflerler ancak elde edilen çıktılar ve yazar yorumları adli bağlamda kabul edilebilirliğin sınırlarını zorlayan engeller içerebilir.

Bu alandaki ilk girişimlerden biri olan Gehrmann ve arkadaşları (2019), geliştirdikleri GLTR (*Giant Language Model Test Room*) aracıyla istatistiksel tespit yaklaşımını somutlaştırmıştır. GLTR, bir metindeki sözcüklerin model tarafından tahmin edilme olasılığını görselleştirir; yeşil ile işaretlenen yüksek olasılıklı sözcüklerin yoğunluğu, yapay zekâ üretiminin göstergesi olarak yorumlanır. Bu yaklaşımın mantığı, dil modellerinin güvenli yani yüksek olasılıklı seçimlere yönelme eğilimine dayanır; insan yazarlar ise daha az tahmin edilebilir ve daha şaşırtıcı seçimler yapabilir.

Daha güncel bir yaklaşım olan Mitchell ve arkadaşları (2023), *DetectGPT* sistemini geliştirerek "sıfır-atış" tespit yöntemini sunmuştur. Bu yöntem, metnin küçük değişikliklere karşı olasılık değişimini analiz eder. Modellerin ürettiği metinler, kendi üretim dağılımının merkezinde yer aldığından, yapılan değişikliklerle olasılık eğrileri düşüş gösterir, insan metinleri ise bu dağılımın dışında olduğundan farklı bir davranış sergiler.

Ancak bu teknik ilerlemelere rağmen, tespit sistemleri adli kullanım için önemli sınırlılıklarla karşı karşıyadır:

Direnç: Tespit sistemleri geliştikçe, bunları atlatmak için kullanılan karşı stratejiler de gelişmektedir. Sadasivan ve arkadaşları (2023), basit yeniden yazım tekniklerinin bile tespit oranlarını ölçüde düşürdüğünü ileri sürmüşlerdir.

Önyargı: Tespit sistemleri, özellikle anadili İngilizce olmayan yazarları yanlışlıkla yapay zekâ olarak etiketleme eğilimindedir. Liang ve arkadaşları (2023), bu durumun akademik ve adli bağlamda masum bireyleri hedef alan ciddi adaletsizliklere yol açabileceğini, modellerin metin karmaşıklığı düşüklüğünü doğrudan YZ üretimiyle karıştırdığını vurgular.

Hibrit Metinler: Mevcut sistemler insan vs. makine ikilemi üzerine kuruludur, ancak gerçek adli vakalar genellikle hibrit metinlerdir. Bir tehdit mektubunun yüzde kaçının yapay zekâ, yüzde kaçının insan olduğu sorusu, mevcut araçlarla yanıtlanamayabilir.

Model Çeşitliliği: Tespit sistemleri genellikle belirli modeller üzerinde eğitilir. Açık kaynak modellerin yaygınlaşması ve kullanıcıların kendi modellerini özelleştirmesi farklı tespit sistemlerinin kapsama alanını daraltmaktadır.

Adli dilbilim bakış açısından bakıldığında, bu sistemlerin Daubert kriterleri karşısındaki durumu kritiktir çünkü Coulthard, Johnson ve Wright'ın (2017) metodolojik çerçevesi düşünüldüğünde bu araçlar şu an için bilimsel olarak ilginç olsalar da, yüksek hata oranları ve test edilebilirlik sorunları nedeniyle hukuken kabul edilebilir olmaktan uzaktır.

Son olarak, meselenin etik boyutu Olsson'a (2008) göre uzman, raporunun mutlak bir gerçek değil, yalnızca bir görüş olduğunu asla unutmamalıdır. Yanlış pozitif bir tespit sonucu masum bir öğrenciyi veya sanığı sahtekârlıkla suçlamak, kanıt yükünü haksız bir şekilde tersine çevirerek kişiyi kendi yazarlığını ispatlamaya zorlar. Bu risk, tespit araçlarının adli bir kesinlik aracı olarak kullanılmasının önündeki en büyük etik engeldir. Mevcut tespit teknolojilerinin karşılaştığı bu metodolojik ve etik sınırlar, çözümün yalnızca teknik araçlarda aranmaması gerektiğinin altını çizmektedir. Bu bağlamda, teknik araçların sağladığı verileri dilbilimsel analizle destekleyebilecek güncel yaklaşımların incelenmesi, alanın gelişimi açısından önem taşımaktadır.

4. Yapay Zekâ Çağında Adli Dilbilimsel Yazarlık Analizine Yönelik Yeni Çerçeveler

Büyük dil modellerinin metin üretimi, yapay zekâ kaynaklı sorunları tanımlamak için gerekli kuramsal zemini sunmuş olsa da adli dilbilimin bu

yeni gerçekliğe metodolojik olarak nasıl yanıt vereceği sorusu güncelliğini korumaktadır. Dolayısıyla tanımlamanın ötesine geçerek çözüm üretmek, hem kavramsal hem de uygulamalı düzeyde yeni çerçevelere ihtiyaç vardır.

Bu bağlamda geliştirilecek yeni çerçevelerin, sadece metnin kaynağını tespit etmeye değil, aynı zamanda metnin insan-makine etkileşimi içinde nasıl kurgulandığını ve sorumluluğun bu süreçte nasıl dağıldığını anlamlandırmaya odaklanması önem arz etmektedir. Bu doğrultuda, insan-yapay zekâ işbirliğine dayalı hibrit yazarlık biçimlerini, bu süreçte insan izinin nasıl dönüştüğünü ve adli dilbilimsel analizin bu dönüşümü okumak için ihtiyaç duyabileceği analitik araçları incelenmeye ihtiyaç duymaktadır.

4.1 İnsan-yapay zekâ işbirliğine dayalı yazarlık analizi: yeni yazarlık biçimleri

Yapay zekâ destekli metin üretimi, yazarlığı geleneksel insan ve makine ikiliğinden çıkarak karmaşık bir yelpaze üzerine yerleştirir. Bu yelpazenin bir ucunda tamamen insan üretimi metinler, diğer ucunda ise tamamen model üretimi metinler bulunsun da günümüzdeki gerçek kullanım senaryolarının büyük çoğunluğu bu iki uç arasında, hibrit bir düzlemde konumlanır. Adli dilbilimcinin karşılaştığı temel zorluk, bu yelpaze üzerinde belirli bir metnin tam olarak nereye düştüğünü tespit etmek ve sorumluluğun fail ile araç arasında nasıl dağıldığını belirlemektir.

Grant (2022), adli yazarlık analizinin ilerleme kavramını tartışırken, alanın yeni zorluklarla karşılaştıkça metodolojik uyum sağlama kapasitesini vurgular. Yapay zekâ çağında bu uyum, yalnızca mevcut yöntemlerin iyileştirilmesini değil, aynı zamanda yeni analitik çerçevelerin geliştirilmesini gerektirir. Hibrit yazarlık senaryoları için potansiyel ve güçlü bir yaklaşım, yazarlık katmanları analizidir. Bu analiz, metnin farklı düzeylerinde içerik seçimi, yapısal organizasyon, dilsel form, sözcüksel tercihler insan ve makine katkılarını ayrı ayrı değerlendirmek olarak düşünülebilir. Bu yaklaşım, Love'ın (2002) işlevsel yazarlık modeliyle uyumludur ve hibrit metinlerin karmaşıklığını, tek boyutlu bir analizden çok daha etkili bir biçimde yakalayabilir.

Hibrit metinlerin analizi, öncelikle yazarın teknoloji karşısındaki konumunun teorik olarak yeniden tanımlanmasını gerektirir. Bajohr (2024), bu yeni durumu açıklamak için "nedensel yazarlık" modelini önermektedir. Bu modele göre, geleneksel "birincil yazarlık" ile kod veya veri seti aracılığıyla dolaylı yoldan üretilen "ikincil" ve "üçüncül" yazarlık biçimleri arasına, günümüz büyük dil modelleri ile ortaya çıkan "dördüncül yazarlık" eklenmiştir. Bu katmanda yazar, ne kodu yazar ne de veri setini seçer sadece bir arayüz üzerinden sisteme bir istem girer. Bu durum, yazarın metin üzerindeki kontrolünü sistemin parametreleriyle sınırlar ve yazarlığı bir tür editörlük veya istem mühendisliği sürecine dönüştürür (Bajohr, 2024).

Bu süreçteki insan-makine etkileşimi, Hancock ve arkadaşları (2020) tarafından "Yapay Zekâ Aracılı İletişim" çerçevesinde ele alınmaktadır. Burada yapay zekâ, iletişimsel hedeflere ulaşmak için mesajları değiştiren, genişleten veya üreten bir ajan olarak hareket ederken, insan asıl fail

konumundadır. Ancak bu teorik tanımlamalar, adli bir vakada somut delil sunmak için tek başına yeterli değildir. Soyut faillik tartışmalarını uygulamalı bir analiz yöntemine dönüştürmek amacıyla, metnin farklı üretim düzeylerini hedefleyen yapısal bir çerçeveye geçilmesi elzemdir.

4.2 Yazarlık katmanları analizi

Grant'in (2022) vurguladığı metodolojik ilerleme gereği, hibrit metinlerin analizi, metni tek bir bütün olarak ele almaktan ziyade, yazarlık katmanları üzerinden detaylandıran bir yaklaşımı uygulanabilir. Bu katmanlar; içerik üretimi, yapısal organizasyon ve dilsel form olmak üzere üç ana başlıkta incelenebilir.

İçerik Üretimi ve Fikirsal Katkı: Analizin en soyut katmanı olan fikir geliştirme sürecinde, insan ve makine rolleri değişkendir. Lee ve arkadaşları (2022), *CoAuthor* veri seti üzerindeki analizlerinde, yazarların yapay zekâyı kullanım amaçlarının fikir bulma ve metne dökme ekseninde farklılaştığını ortaya koymuştur. Çalışma, yazarların bir kısmının aracı yalnızca dilsel akıcılık için kullandığını; diğer bir kısmının ise olay örgüsünü ilerletmek veya karakter ismi türetmek gibi somut içerik üretimi için modele başvurduğunu göstermektedir (s. 12-13). Özellikle yapay zekânın önerdiği isimlendirilmiş varlıkların sonraki metinlerde %20 oranında yeniden kullanılması, modelin pasif bir daktilo olmaktan çıkıp, içeriğin gelişimine yön veren aktif bir fikir ortağına dönüştüğüne işaret eder. Ancak Lee ve arkadaşlarının bulgularına göre, metindeki sahiplik algısı, yazarın üretim sürecine ne kadar aktif katıldığıyla doğru orantılıdır yalnızca makine çıktılarını düzenleyen kullanıcıların tam bir sahiplik hissi geliştiremediği gözlemlenmiştir.

Yapısal Organizasyon ve Tür Analizi: Metnin retorik yapısının incelendiği bu katmanda, insan ve makine üretimi arasında belirgin türsel farklılıklar dikkat çekilebilir. Khalid ve arkadaşları (2025), akademik özetler üzerine yaptıkları karşılaştırmalı analizde, insan yazarların ve *ChatGPT* gibi yapay zekâ modellerinin retorik hamle kullanımı ve metin derinliği açısından belirgin şekilde ayrıştığını tespit etmiştir, buna göre yapay zekâ, metinlerinde "Çalışmanın Amacı" bölümüne orantısız bir ağırlık verip (%24,1'e karşı %11,8) okuyucuya doğrudan hedefi aktarmaya odaklanırken, insan yazarlar çalışmayı literatürle ilişkilendiren "Bağlamsallaştırma" ve somut bulguları içeren "Sonuçlar" bölümlerine çok daha dengeli ve yoğun bir yer vermektedir.

Bu yapısal durum, Maporn ve arkadaşları (2024) tarafından da desteklenmektedir, yapay zekâ çıktıları genellikle "Giriş" bölümünü atlayarak P-M-Pr-C (Amaç-Yöntem-Ürün-Sonuç) yapısını benimsemekte ve böylece çalışmanın bilimsel mirasını aktaran bağlamdan yoksun kalmaktadır. Benzer şekilde Hsu ve arkadaşları (2024), insan yazarların "Sonuçlar" ve "Yöntem" bölümlerine odaklandığını, yapay zekânın ise "Önem" ve "Sonuç/Yorum" gibi daha genel geçer bölümlerde tekrara düşme eğiliminde olduğunu saptamıştır.

Retorik eksikliklerin ötesinde, Amirjalili ve arkadaşları (2024) ile Bagdasarov ve Alves (2025)'in bulguları, yapay zekâ metinlerinin "yazar sesi" ve sözdizimsel karmaşıklık açısından da insanlardan ayrıştığını; yapay zekânın robotik, kalıplaşmış ve tekrarlayan yapılar kullanırken, insan yazarların daha

yüksek sözdizimsel değişkenlik ve özgün bir anlatı sergilediğini ortaya koymaktadır.

Özetle, yapay zekâ akademik yazımın şekilsel şartlarını yerine getirirse de, bağlam kurma, özgün bulguları vurgulama ve entelektüel bir imza taşıma konusunda insan derinliğine ulaşamamaktadır. İnsan yazarlar, Swales (1990) ve Bhatia'nın (1993) klasik tür modellerine uygun olarak bağlamsallaştırma ve sonuç tartışması bölümlerine ağırlık verirken; yapay zekâ metinleri orantısız bir biçimde çalışmanın amacı üzerine odaklanmakta, bağlam ve sonuç kısımlarını ise yüzeysel geçmektedir. Bu yapısal dengesizlikler, adli dilbilimci için metnin makro yapısında algoritmik bir formülizasyonun varlığına dair ipuçları sunabilir.

Dilsel Form ve Stilometrik Özellikler: Bu analiz katmanında, istatistiksel yöntemler devreye girer. O'Sullivan (2025), Burrows'un Delta yöntemiyle yaptığı incelemelerde, insan ve büyük dil modelleri üretimi yaratıcı metinler arasında istatistiksel olarak ayırt edilebilir farklar bulunduğunu öne sürer. O'Sullivan'a göre, insan yazarlar stilistik açıdan heterojen bir küme oluştururken, modeller kendi içlerinde yüksek tutarlılığa sahip homojen kümeler meydana getirmektedir. Özellikle GPT-4'ün stilistik tutarlılığı artmış olsa da, insan yazımının doğasında var olan değişkenlikten yoksun olduğu savunulmaktadır.

Holtzman ve arkadaşlarının (2019) verileri, insan dilinin istatistiksel olarak yüksek olasılık bölgesinde sabit kalmadığını fakat anlamı zenginleştirmek için şaşırtıcı ve düşük olasılıklı seçimlere yöneldiğini göstermektedir (s. 4-5). Dil modelleri, doğal olmayan derecede düşük şaşırtıcılık sergileyen metinler üretmektedir. Bu süreç, insana özgü bağlamsal sapmaları ve bilgilendirici sürprizleri eleyerek, metinleri güvenli bir hale getirir.

André ve arkadaşları (2023), 2.100 akademik özet üzerinde gerçekleştirdikleri karşılaştırmalı analizde, insan yazarların metinlerinde dilsel çeşitlilik ve karmaşıklığın daha yüksek olduğunu, yapay zekâ modellerinin ise n-gram dağılımlarında daha kalıplaşmış ve tekrara dayalı yapılar sergilediğini tespit etmiştir. Çalışma, insan yazarların ifade biçimlerinde özgünlüğe ve çeşitliliğe yöneldiğini buna karşın ChatGPT gibi modellerin, istatistiksel olarak güvenli kabul edilen kelime dizilerini (yüksek dereceli n-gramlar) daha sık kullanarak daha resmi ancak tekdüze bir yapı oluşturduğunu ortaya koymaktadır.

Öte yandan, Ippolito ve arkadaşları (2020), bu istatistiksel izlerin insan müdahalesiyle silinebileceğine dikkat çeker. Modern kod çözme stratejileri insanları yanıltmak üzere optimize edilse de, metin uzadıkça modellerin istatistiksel gariplikleri belirginleşir. Hibrit metinlerin analizindeki en temel metodolojik engel, yazarlık tutarlılığının sık sık kesintiye uğramasıdır. Zeng ve arkadaşları (2024), hibrit metinlerin önce insan, sonra makine şeklinde basit bloklardan oluşmadığını, aksine çok sayıda etkileşimlerle örülmüş, sınırların geçirgen olduğu karmaşık yapılar olduğunu vurgular.

Bu karmaşıklık, yazarın her cümlede müdahale ettiği senaryolarda parça uzunlukları kısaltıldıkça, dedektörlerin güvenilirliği azalmakta ve ortaya yazarlık tutarlılığı olmayan segmentler çıkmaktadır. Bu durum, tespit stratejisinin durağan bir yaklaşımdan ziyade, metnin yoğunluğuna göre dinamik olarak belirlenen bir katman analizine evrilmesini gerektirir.

Buna karşılık Zou ve arkadaşları (2024), alanın çözüm arayışının alana özgü dedektörlerden, Çok Modlu Büyük Dil Modellerine (MLLM) dayalı genel amaçlı sistemlere kaydığını belirtmektedir. MLLM tabanlı yeni nesil araçlar, sadece gerçek/sahte ayrımı yapmakla kalmayıp; aynı zamanda açıklanabilirlik ve sahteciliğin metin içindeki yerini saptama yetenekleri sunarak, adli dilbilimcilere kara kutuyu aralayacak bir analiz seti vaat etmektedir. Teknolojik araçların gelişimi 'metnin nasıl üretildiği' sorusuna teknik yanıtlar sunsa da, ortaya çıkan içeriğin hukuki ve ahlaki yükümlülüğünün kime ait olduğu sorusu, teknik değil, etik bir düzlemde tartışılmayı gerektirir.

Sonuç olarak, yapay zekâ çağında adli yazarlık analizi, Grant'in (2022) öngördüğü gibi, Love'ın (2002) işlevsel modelini de kapsayacak şekilde genişlemelidir. Analiz, metni tek bir bütün olarak ele almak yerine; (1) içerik üretimi ve fikrîsel katkı, (2) yapısal organizasyon ve tür analizi ve (3) dilsel form ve stilometrik özellikleri ayrı ayrı inceleyen katmanlı bir yaklaşımı benimsemelidir. Bu yaklaşım, sadece metnin kim tarafından yazıldığını değil, insan ve makine failliğinin metin içinde nasıl iç içe geçtiğini ve sorumluluğun nerede başladığını görmek için gereklidir. Ancak bu teknik analiz çerçevesi, sorumluluğun belirlenmesi için gerekli olsa da yeterli değildir, üretim sürecinin etik ve hukuki boyutları da dikkate alınmalıdır.

4.3 Etik sorumluluk

Teknik analizlerin ötesinde, hibrit metinler ciddi etik ve hukuki soruları da beraberinde getirir. Yousaf (2025), bilimsel yayıncılık bağlamında, yapay zekânın bir yazar olarak kabul edilemeyeceğini, çünkü ürettiği içeriğin doğruluğu veya bütünlüğü konusunda sorumluluk alamayacağını net bir şekilde ifade eder. Yapay zekâ araçları, var olmayan alıntılar üretmek veya intihal benzeri içerikler oluşturmak gibi riskler taşır. Bu nedenle, adli bağlamda nihai sorumluluk, metnin üretiminde yapay zekâyı ne kadar yoğun kullanmış olursa olsun, çıktılarını denetleyen ve nihai onayı veren insan failde kalmaktadır (Yousaf, 2025).

Ayrıca, bu sorumluluk sadece nihai onayla sınırlı kalmayıp, şeffaflık ilkesini de kapsamalıdır. Maporn ve diğerleri (2024), özellikle akademik bağlamda, yapay zeka tarafından oluşturulan içeriklerin tespit edilebilir retorik kalıplar sergilediğini ve bu durumun insan yazarlardan belirgin şekilde ayrıştığını ortaya koymuştur. Ancak Lafren ve diğerleri (2025), öğrencilerin yapay zeka araçlarını kullanırken etik sınırları ihlal etme eğiliminde olabildiklerini, bu nedenle yapay zeka okuryazarlığı eğitiminin, araçların sınırlılıklarını ve sahte akademik kaynak risklerini anlamaları açısından kritik olduğunu vurgular. Bu bağlamda, adli analiz sadece metnin kaynağını değil, aynı zamanda üretim sürecindeki şeffaflığı ve yazarın araç üzerindeki denetim

düzeyini de sorgulamalıdır çünkü sorumluluk, aracın kullanım şekli ve sonuçların doğrulanması sürecindeki insan müdahalesi ile doğrudan ilişkilidir (Laflen vd., 2025).

Teknolojik araçların gelişimi 'metnin nasıl üretildiği' sorusuna teknik yanıtlar sunsa da, ortaya çıkan içeriğin hukuki ve ahlaki yükümlülüğünün kime ait olduğu sorusu adli dilbilimin merkezinde kalmaktadır. Yapay zekâ çağında, faillik tespiti yalnızca dilsel izleri sürmek değil, aynı zamanda insan iradesinin ve denetiminin metinde ne ölçüde bulunduğunu belirlemektir.

4.4 Adli süreç ve hukuki etkiler

Yapay zekâ tabanlı metin üretimi, adli dilbilimin yalnızca metodolojik değil, aynı zamanda hukuki ve kurumsal boyutlarını da etkilemektedir. Mahkemelerde dilbilimsel kanıtın sunumu, uzman tanıklığının güvenilirliği ve sorumluluğun hukuki çerçevesi, bu yeni gerçeklik karşısında yeniden değerlendirilmek zorundadır. Daubert kriterleri, bilimsel kanıtın kabul edilebilirliği için bir çerçeve sunmaktadır ancak yapay zekâ çağında bu kriterlerin uygulanması, yeni sorular ve sorunlar doğurur. Test edilebilirlik ve yanlışlanabilirlik açısından, hibrit metinlerin analizi için geliştirilen yöntemlerin deneysel olarak test edilmesi gerekir. Ancak mevcut literatür çalışmaları, bu alanda yeterli deneysel çalışma içermemektedir.

Test edilebilirlik ve yanlışlanabilirlik açısından hakemli yayın ve bilimsel inceleme açısından, yapay zekâ ve adli dilbilim kesişimindeki literatür hızla büyümektedir ancak bu literatürün önemli bir bölümü henüz kapsamlı hakemli değerlendirmeden geçmemiştir. Bu durum, mahkemelerde sunulan kanıtın bilimsel temelini sorgulanmasına yol açabilir. Bilinen hata oranları açısından, hibrit metin analizi için güvenilir hata oranları henüz belirlenmemiştir. Liang ve arkadaşlarının (2023) gösterdiği gibi, mevcut tespit sistemleri özellikle anadili farklı olan yazarlar için yüksek yanlış pozitif oranları üretebilmektedir. Bu durum, adaletin tesisi açısından ciddi riskler içerir bu yüzden masum bireyler, yanlış tespit sonuçlarına dayanılarak suçlanabilir.

Standartların varlığı ölçütü açısından, yapay zekâ destekli metin analizi için uluslararası kabul görmüş standartlar henüz mevcut değildir. International Association of Forensic Linguists (IAFL) ve benzeri kuruluşlar, bu alanda kılavuzlar geliştirme sürecindedir ancak bu süreç, teknolojik gelişmelerin gerisinde kalmaktadır. Son kriter olan genel kabul açısından, yapay zekâ tespit yöntemlerinin dilbilim ve bilgisayar bilimleri topluluklarında ne ölçüde kabul gördüğü tartışmalıdır. Bazı araştırmacılar, mevcut yöntemlerin güvenilir tespit için yetersiz olduğunu savunurken diğerleri de belirli koşullarda bu yöntemlerin yararlı olabileceğini ileri sürmektedir.

Bu kriterler ışığında, adli dilbilimcinin mahkemede sunacağı uzman görüşünün kapsamı ve ifade biçimi yapay zekâ çağında daha da önem kazanır. Uzman, hibrit metin analizi sonuçlarını sunarken, yöntemin sınırlılıklarını açıkça belirtmeli ve kesinlik iddiasından kaçınmalıdır. "Bu metin yapay zekâ tarafından üretilmiştir" gibi kategorik ifadeler yerine, "eldeki veriler ışığında, bu metnin yapay zekâ katkısı içermesi olasıdır" gibi ölçülü sonuçlar tercih

edilmelidir. Bu ölçülülük, Tiersma ve Solan'ın (2002) dikkat çektiği iletişim sorunu da yeniden gündeme gelir. Jüri üyeleri, olasılıksal ifadeleri yanlış yorumlama eğilimindedir; yapay zekâ tespit sonuçları, bu riski artırabilir. "Yüzde 85 olasılıkla yapay zekâ üretimi" ifadesi, teknik olmayan dinleyiciler tarafından "kesinlikle yapay zekâ üretimi" olarak algılanabilir. Adli dilbilimcinin pedagojik sorumluluğu, bu bağlamda daha da önemli hale gelir.

Öte yandan, sorumluluk atfının hukuki çerçevesi de yeniden düşünülmelidir çünkü geleneksel yazarlık analizi, metnin arkasındaki tek bir bireyi yani faili belirlemeyi hedefler ancak hibrit yazarlık senaryolarında, sorumluluk tek bir bireye atfedilemeyebilir. Yapay zekânın kendisi hukuki fail olamaz, sorumluluk, kullanıcıya, geliştiriciye, dağıtıcıya yönelmek zorundadır ancak bu yönelimin hukuki temelleri ve ölçütleri henüz netleşmemiştir.

Bu boşluğu doldurmak için uluslararası ve ulusal düzeyde düzenleyici girişimler başlatılmıştır. Avrupa Birliği'nin Yapay Zekâ Yasası (European Commission, AI Act, 2024), bu alanda düzenleyici bir çerçeve sunma girişimindedir. Yasa, yapay zekâ sistemlerini risk kategorilerine ayırır ve yüksek riskli uygulamalar için şeffaflık ve hesap verebilirlik gereksinimleri belirler. Adli bağlamda kullanılan yapay zekâ tespit sistemleri, bu düzenlemelerin kapsamına girebilir; bu durum, sistemlerin geliştirilmesi ve kullanımı için yeni standartlar getirebilir. Amerika Birleşik Devletleri'nde ise düzenleyici yaklaşım eyalet düzeyinde çeşitli girişimler bulunmakla birlikte, federal düzeyde kapsamlı bir yapay zekâ mevzuatı henüz yoktur. Bu düzenleyici belirsizlik, bu anlamdaki adli dilbilim uygulamalarının hukuki temelini de belirsiz kılmaktadır.

Türkiye bağlamında, yapay zekâ ve metin üretimi konusunda "Yapay Zekâ Kanun Teklifi" gibi düzenleme girişimleri 2025 yılında TBMM'ye sunulmuş olsa da, henüz yürürlükte olan ve yapay zekâ destekli metin üretimini doğrudan kapsayan henüz kapsamlı bir hukuki çerçeve bulunmamaktadır. Mevcut 5237 sayılı Türk Ceza Kanunu (TCK) ve 5846 sayılı Fikir ve Sanat Eserleri Kanunu (FSEK), yapay zekânın otonom yapısından kaynaklanan özgün sorunları karşılamakta yetersiz kalmaktadır. Özellikle FSEK, bir çıktının eser sayılabilmesi için eser sahibinin hususiyetini taşıması ve yaratıcısının gerçek bir kişi olması şartını aradığından, yapay zekâ tarafından üretilen metinlerin mülkiyeti ve sorumluluğu konusunda yasal bir boşluk oluşmaktadır. Bu durum, adli dilbilimcinin uzman görüşünün hukuki bağlamda nasıl değerlendirileceği konusunda belirsizlik yaratmaktadır.

Benzer şekilde TCK, "Deepfake" gibi manipülatif içeriklerin oluşturulması veya yapay zekânın suç aracı olarak kullanılması durumunda failin rolünün belirlenmesinde net hükümler içermemektedir, bu durum yeni kanun teklifi ile giderilmeye çalışılmaktadır. Öte yandan, Türk yargı uygulamalarında, yapay zekâ tarafından üretilen veya manipüle edilen dijital içeriklerin delil niteliği (hukuka uygun elde edilip edilmediği ve güvenilirliği) konusunda Yargıtay'ın yerleşik bir içtihadı henüz bulunmamaktadır. Sonuç olarak, yapay zekâ çıktılarının gerçeklik ve kaynak sorunu, adli dilbilimcinin sunacağı uzman görüşünün hukuki bağlamda nasıl konumlandırılacağı konusunda ciddi bir belirsizlik yaratmaktadır.

Adli dilbilim eğitimi ve mesleki gelişim de bu dönüşümden etkilenmektedir. Gelecek kuşak adli dilbilimciler, geleneksel stilistik ve stilometrik yöntemlerin yanı sıra, yapay zekâ sistemlerinin işleyişi, tespit yöntemleri ve sınırlılıkları konusunda da yetkinlik kazanmak zorundadır. Bu durum, eğitim tasarımıyla içerik tasarımına kadar geniş bir yelpazede kurumsal değişiklikler gerektirmektedir.

5. Sonuç ve Değerlendirme

Bu çalışma, adli dilbilimsel yazarlık analizinin yapay zekâ çağındaki dönüşümünü; kavramsal temeller, metodolojik yenilikler ve hukuki sınırlar ekseninde incelemiştir. Barthes'ın "yazarın ölümü" metaforundan Foucault'nun "yazar işlevi" kavramına, Austin'in söz eylem kuramından Love'ın işlevsel yazarlık modeline uzanan teorik zemin; büyük dil modellerinin metin üretim kapasitesiyle birlikte farklı bir açıdan ele alınması gereklilik haline dönüşmüştür. Bajohr'un (2024) nedensel yazarlık olarak tanımladığı bu yeni durum, yazarlığı doğrudan bir üretim eyleminden çıkarıp, istem aracılığıyla yürütülen bir yaratım sürecine dönüştürmüştür.

Çalışmanın temel bulgusu, adli dilbilimin üzerine inşa edildiği bireysel tutarlılık ve yazarlar arası ayırt edicilik varsayımlarının, yapay zekâ destekli metin üretim süreçlerinde geçerliliğini yitirmeye başladığıdır. Holtzman ve arkadaşlarının (2019) verileri ışığında görüldüğü üzere dil modelleri, üretim sürecinde insana özgü ve bağlama dayalı dilsel sapmaları eleyerek, metinleri istatistiksel açıdan en olası ve standart forma indirgeme eğilimindedir. Bu durum, yazarın kendine has stil ve bireydir özelliklerinin silikleşmesine ve metinlerin tekdüzeleşmesine yol açmaktadır.

İnsan yazarın bilişsel sınırlarını yapay zekâ aracılığıyla aştığı bu süreçte, ortaya çıkan metin artık biyolojik bir bireydirin doğrudan yansıması değildir. Hancock ve arkadaşları (2020) bu durumu, niyeti belirleyen insan ile üretimi gerçekleştiren algoritma arasındaki bir süreç olarak tanımlar. Dolayısıyla hibrit metinler, tekil bir yazarın sesi olmaktan çıkıp, insan iradesi ile algoritmik olasılıkların iç içe geçtiği karmaşık bir yapıya dönüşmektedir.

Bu dönüşüm, metodolojik düzlemde tek boyutlu analizden katmanlı analize geçişi zorunlu kılmaktadır. Grant'in (2022) vurguladığı metodolojik uyum çerçevesinde önerilen "Yazarlık Katmanları Analizi", hibrit metinlerin karmaşıklığını çözümlemek için en elverişli model olarak öne çıkmaktadır. Bu çalışmada gösterildiği üzere, Love'ın (2002) teorik işlevsel ayrımları, somut metriklerle birleştirildiğinde metnin içerik, yapı ve form katmanlarındaki insan-makine katkısını haritalandırmak mümkün hale gelmektedir. Analiz artık "Yazar kim?" sorusundan ziyade, "Hangi katmanda insan niyeti, hangi katmanda algoritmik olasılık baskındır?" sorusuna odaklanmaktadır.

Hukuki boyutta, Daubert kriterlerinin yapay zekâ çağındaki uygulanabilirliği ciddi kısıtlılıklar barındırmaktadır. Tespit sistemlerinin kırılganlığı ve önyargı riskleri, algoritmik çıktıların mahkemede tek başına delil olarak sunulmasını tartışılmalı hale getirmektedir. Bu çalışma, hibrit metin analizinin test edilebilirlik, hakemli inceleme, hata oranları, standartlar ve genel kabul kriterleri açısından henüz yeterli olgunluğa erişmediğini ortaya

koymuştur. Özellikle Liang ve arkadaşlarının (2023) gösterdiği gibi, mevcut tespit sistemlerinin anadili farklı olan yazarlara karşı yüksek yanlış pozitif oranları üretmesi, adaletin tesisi açısından kabul edilemez riskler taşımaktadır. Bu belirsizlik karşısında, adli dilbilimci tanrısal bir kesinlik iddiasından kaçınarak bulgularını olasılıksal bir çerçevede ve yöntemin hata paylarını şeffaflıkla belirterek sunmakla yükümlüdür.

Çalışmanın sınırlılıkları da göz ardı edilmemelidir. İnceleme büyük ölçüde İngilizce literatüre ve modellere dayanmaktadır, Türkçe gibi morfolojik açıdan zengin dillerdeki hibrit üretim örüntüleri, ayrı ve kapsamlı araştırmaları gerektirmektedir. Ayrıca, tespit sistemlerinin performansı hızla değişen teknolojik dinamiklere bağlı olduğu için, buradaki teknik değerlendirmelerin geçici doğası da kabul edilmelidir.

Gelecek araştırmalar için kullanıcıların istem yazma stillerindeki bireysel farklılıkların adli bir iz olarak güvenilirliği ve hibrit metinlerdeki düzenleme aşamalarının izini süren veri setleri üzerinde süreç odaklı analizler yapılabilir. Öte yandan, Yapay Zekâ Yasası (*AI Act*) gibi uluslararası düzenlemelerin, Türk hukuk sistemi ve adli bilirkişilik pratiklerine uyumu üzerine karşılaştırmalı çalışmalar yapılabilir. Türk Ceza Kanunu'nda dijital delillerin kaynağının doğrulanması prosedürleri, AB'nin şeffaflık standartlarını karşılayacak teknik altyapıya sahip olup olmadığı denetlenebilir.

Son değerlendirmede, adli dilbilimde yazar analizi yaşayan ve sorumlu bir fail arayışında olduğu için Foucault'nun "Kimin konuştuğunun ne önemi var?" sorusunun aksine adli bir gereklilik kazanmıştır. Yapay zekâ çağında failin sesi algoritma sesiyle karışmış olsa da, adli dilbilimcinin görevi değişmemiştir: Bu seslerin içindeki insan niyetini, var oluşsal sınırlarının bilincinde olarak, titizlikle ayırtmak ve adaletin tesisi için kullanıma sunmaktır. Grant'ın (2022) ifade ettiği gibi; alandaki gerçek "ilerleme", sadece doğruluk oranlarını artırmak değil, değişen yazarlık doğasına uyum sağlayacak esnek ve sağlam çerçeveler inşa edebilmektir.

Kaynaklar

- Acemoğlu, D. (2021). Harms of AI. In M. Dubber, F. Pasquale, & S. Das (Eds.), *The Oxford handbook of ethics of AI* (pp. 595–616). Oxford University Press.
<https://doi.org/10.1093/oxfordhb/9780190067397.013.36>
- Amirjalili, F., Neysani, M., & Nikbakht, A. (2024). Exploring the boundaries of authorship: A comparative analysis of AI-generated text and human academic writing in English literature. *Frontiers in Education*, 9, Article 1347421. <https://doi.org/10.3389/feduc.2024.1347421>
- Argamon, S., Koppel, M., Pennebaker, J. W., & Schler, J. (2009). Automatically profiling the author of an anonymous text. *Communications of the ACM*, 52(2), 119–123.
<https://doi.org/10.1145/1461928.1461959>
- Austin, J. L. (1962). *How to do things with words* (J. O. Urmson, Ed.). Harvard University Press.

- Bagdasarov, S., & Alves, D. (2025, February). *Like a human? A linguistic analysis of human-written and machine-generated scientific texts* [Conference presentation]. Saarland University, Saarbrücken, Germany.
- Bajohr, H. (2024). *Writing at a distance: Notes on authorship and artificial intelligence*. transcript Verlag.
- Barthes, R. (1967). The death of the author. *Aspen Magazine*, (5–6).
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610–623). <https://doi.org/10.1145/3442188.3445922>
- Bhatia, V. K. (1993). *Analysing genre: Language use in professional settings*. Longman.
- Bloch, B. (1948). A set of postulates for phonemic analysis. *Language*, 24(1), 3–46. <https://doi.org/10.2307/410284>
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., ... Amodei, D. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901.
- Bucholtz, M., & Hall, K. (2005). Identity and interaction: A sociocultural linguistic approach. *Discourse Studies*, 7(4–5), 585–614. <https://doi.org/10.1177/1461445605054407>
- Burrows, J. (2002). ‘Delta’: A measure of stylistic difference and a guide to likely authorship. *Literary and Linguistic Computing*, 17(3), 267–287. <https://doi.org/10.1093/lilc/17.3.267>
- Chaski, C. E. (2001). Empirical evaluations of language-based author identification techniques. *Forensic Linguistics*, 8(1), 1–65. <https://doi.org/10.1558/sll.2001.8.1.1>
- Chaski, C. E. (2012). Best practices and admissibility of forensic author identification. *Journal of Law and Policy*, 21(2), 333–376.
- Coulthard, M. (2004). Author identification, idiolect, and linguistic uniqueness. *Applied Linguistics*, 25(4), 431–447. <https://doi.org/10.1093/applin/25.4.431>
- Coulthard, M., Johnson, A., & Wright, D. (2017). *An introduction to forensic linguistics: Language in evidence* (2nd ed.). Routledge. <https://doi.org/10.4324/9781315621811>
- Crystal, D., & Davy, D. (1969). *Investigating English style*. Longman.
- Eckert, P. (2012). Three waves of variation study: The emergence of meaning in the study of sociolinguistic variation. *Annual Review of Anthropology*, 41, 87–100. <https://doi.org/10.1146/annurev-anthro-092611-145828>

- Eder, M. (2015). Does size matter? Authorship attribution, small samples, big problem. *Digital Scholarship in the Humanities*, 30(2), 167–182. <https://doi.org/10.1093/llc/fqt066>
- European Parliament. (2024). *Regulation (EU) 2024/1689 on artificial intelligence (Artificial Intelligence Act)*. *Official Journal of the European Union*, L 2024/1689. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:32024R1689>
- Evert, S., Proisl, T., Jannidis, F., Reger, I., Pielström, S., Schöch, C., & Vitt, T. (2017). Understanding and explaining Delta measures for authorship attribution. *Digital Scholarship in the Humanities*, 32(Suppl. 2), ii4–ii16. <https://doi.org/10.1093/llc/fqx023>
- Floridi, L. (2013). Distributed morality in an information society. *Science and Engineering Ethics*, 19(3), 727–743. <https://doi.org/10.1007/s11948-012-9413-4>
- Floridi, L., & Chiriatti, M. (2020). GPT-3: Its nature, scope, limits, and consequences. *Minds and Machines*, 30, 681–694. <https://doi.org/10.1007/s11023-020-09548-1>
- Foucault, M. (1977). What is an author? (D. F. Bouchard & S. Simon, Trans.). In D. F. Bouchard (Ed.), *Language, counter-memory, practice: Selected essays and interviews* (pp. 113–138). Cornell University Press. Original work published 1969.
- Gehrmann, S., Strobelt, H., & Rush, A. M. (2019). GLTR: Statistical detection and visualization of generated text. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics: System Demonstrations* (pp. 111–116). <https://doi.org/10.18653/v1/P19-3019>
- Grant, T. (2022). The idea of progress in forensic authorship analysis. In M. Coulthard, A. May, & R. Sousa-Silva (Eds.), *The Routledge handbook of forensic linguistics* (2nd ed., pp. 15–30). Routledge. <https://doi.org/10.4324/9780429030581-2>
- Grant, T., & MacLeod, N. (2018). Resources and constraints in forensic authorship analysis. In *Language and law* (pp. 64–79). Cambridge University Press. <https://doi.org/10.1017/9781316585136.004>
- Grant, T., & MacLeod, N. (2020). *Language and online identities: The undercover policing of internet sexual crime*. Cambridge University Press. <https://doi.org/10.1017/9781108766326>
- Gunkel, D. J. (2016). *Of remixology: Ethics and aesthetics after remix*. MIT Press. <https://doi.org/10.7551/mitpress/10294.001.0001>
- Hallevy, G. (2010). The criminal liability of artificial intelligence entities: From science fiction to legal social control. *Akron Intellectual Property Journal*, 4(2), 171–201.
- Hancock, J. T., Naaman, M., & Levy, K. (2020). AI-mediated communication: Definition, research agenda, and ethical considerations. *Journal of Computer-Mediated Communication*, 25(1), 89–100. <https://doi.org/10.1093/jcmc/zmz022>

- Holtzman, A., Buys, J., Du, L., Forbes, M., & Choi, Y. (2019). The curious case of neural text degeneration. *International Conference on Learning Representations*. <https://openreview.net/forum?id=rygGQyrFvH>
- Hoover, D. L. (2004). Testing Burrows's Delta. *Literary and Linguistic Computing*, 19(4), 453–475. <https://doi.org/10.1093/lilc/19.4.453>
- Hsu, T.-W., Tseng, P.-T., Tsai, S.-J., Ko, C.-H., Thompson, T., Hsu, C.-W., Yang, F.-C., Tsai, C.-K., Tu, Y.-K., Yang, S.-N., Liang, C.-S., & Su, K.-P. (2024). Plagiarism, quality, and correctness of AI-generated vs. human-written abstracts for the same psychiatric research paper. <https://doi.org/10.1016/j.psychres.2024.116145>
- Ippolito, D., Yuan, D., Coenen, A., & Burnam, S. (2020). Automatic detection of generated text is easiest when humans are fooled. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics* (pp. 1808–1822). <https://doi.org/10.18653/v1/2020.acl-main.164>
- Johnstone, B. (1996). *The linguistic individual: Self-expression in language and linguistics*. Oxford University Press.
- Juola, P. (2006). Authorship attribution. *Foundations and Trends in Information Retrieval*, 1(3), 233–334. <https://doi.org/10.1561/1500000005>
- Keselj, V., Peng, F., Cercone, N., & Thomas, C. (2003). N-gram-based author profiles for authorship attribution. In *Proceedings of the Pacific Association for Computational Linguistics* (pp. 255–264).
- Khalid, P. Z. M., Kussin, H., Uri, N. F. M., Bahari, A. A., Zaini, K., & Kasuahdi, N. A. A. (2025). A comparative genre analysis of human-written and AI-generated research abstracts. *International Journal of Research and Innovation in Social Science*, 9(1). <https://doi.org/10.47772/IJRIS.2025.910000044>
- Koppel, M., & Schler, J. (2004). Authorship verification as a one-class classification problem. In *Proceedings of the 21st International Conference on Machine Learning* (pp. 62–69). <https://doi.org/10.1145/1015330.1015448>
- Kredens, K., & Coulthard, M. (2012). Corpus linguistics in authorship identification. In *Corpus linguistics and the web* (pp. 209–222). Brill. https://doi.org/10.1163/9789401208639_015
- Kumarage, T., Agrawal, G., Sheth, P., Moraffah, R., Chadha, A., Garland, J., & Liu, H. (2024). *A survey of AI-generated text forensic systems: Detection, attribution, and characterization*. <https://doi.org/10.48550/arXiv.2403.01152>
- Labov, W. (1972). *Sociolinguistic patterns*. University of Pennsylvania Press.
- Lee, M., Liang, P., & Yang, Q. (2022). CoAuthor: Designing a human–AI collaborative writing dataset for exploring interplay of user initiatives and AI suggestions. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems* (Article 393). <https://doi.org/10.1145/3491102.3502030>

- Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(7), 100779. <https://doi.org/10.1016/j.patter.2023.100779>
- Lima, D. (2018). Could AI agents be held criminally liable: Artificial intelligence and the challenges for criminal law. *South Carolina Law Review*, 69(3), Article 8.
- Love, H. (2002). *Attributing authorship: An introduction*. Cambridge University Press. <https://doi.org/10.1017/CBO9780511483166>
- Maporn, S., Burinprakhon, B., Chaiyasuk, I., & Nattheeraphong, A. (2024). Rhetorical moves of AI-generated research article abstracts and published research article abstracts in applied linguistics. *Academic Journal of Faculty of Humanities and Social Sciences Phranakhon Rajabhat University*, 8(1), 162–176.
- McMenamin, G. R. (2002). *Forensic linguistics: Advances in forensic stylistics*. CRC Press.
- Mitchell, E., Lee, Y., Khazatsky, A., Manning, C. D., & Finn, C. (2023). DetectGPT: Zero-shot machine-generated text detection using probability curvature. In *Proceedings of the 40th International Conference on Machine Learning* (pp. 24950–24982).
- Morrison, G. S. (2009). Forensic voice comparison and the paradigm shift. *Science & Justice*, 49(4), 298–308. <https://doi.org/10.1016/j.scijus.2009.09.002>
- Mosteller, F., & Wallace, D. L. (1964). *Inference and disputed authorship: The Federalist*. Addison-Wesley.
- Nini, A. (2023). *A theory of linguistic individuality for authorship analysis*. Cambridge University Press. <https://doi.org/10.1017/9781009127391>
- O’Sullivan, J. (2025). Stylometric comparisons of human versus AI-generated creative writing. *Humanities and Social Sciences Communications*, 12(1), Article 1708. <https://doi.org/10.1057/s41599-025-05986-3>
- Olsson, J. (2004). *Forensic linguistics: An introduction to language, crime and the law*. Continuum.
- Olsson, J. (2008). *Forensic linguistics* (2nd ed.). Continuum.
- Rybicki, J., & Eder, M. (2011). Deeper Delta across genres and languages: Do we really need the most frequent words? *Literary and Linguistic Computing*, 26(3), 315–321. <https://doi.org/10.1093/lilc/fqr031>
- Sadasivan, V. S., Kumar, A., Balasubramanian, S., Wang, W., & Feizi, S. (2023). *Can AI-generated text be reliably detected?* <https://doi.org/10.48550/arXiv.2303.11156>
- Shuy, R. W. (2006). *Linguistics in the courtroom: A practical guide*. Oxford University Press.
- Smith, P. W. H., & Aldridge, W. (2011). Improving authorship attribution: Optimizing Burrows’ Delta method. *Journal of Quantitative Linguistics*, 18(1), 63–88. <https://doi.org/10.1080/09296174.2011.533591>
- Solaiman, I., Brundage, M., Clark, J., Askill, A., Herbert-Voss, A., Wu, J., ... Wang, J. (2019). *Release strategies and the social impacts of language*

- models* (Report No. OpenAI-2019). OpenAI. <https://arxiv.org/abs/1908.09203>
- Solan, L. M., & Tiersma, P. M. (2005). *Speaking of crime: The language of criminal justice*. University of Chicago Press.
- Sousa-Silva, R. (2024). The new challenge for forensic linguistics: Generative AI and the ‘black box’ of authorship. *International Journal of Speech, Language and the Law*, 31(1), 1–22. <https://doi.org/10.1558/ijssl.26252>
- Stamatatos, E. (2009). A survey of modern authorship attribution methods. *Journal of the American Society for Information Science and Technology*, 60(3), 538–556. <https://doi.org/10.1002/asi.21001>
- Swales, J. M. (1990). *Genre analysis: English in academic and research settings*. Cambridge University Press.
- Tiersma, P. M., & Solan, L. M. (2002). The linguist on the witness stand: Forensic linguistics in American courts. *Language*, 78(2), 221–239. <https://doi.org/10.1353/lan.2002.0126>
- Turell, M. T. (2010). The use of textual, grammatical and sociolinguistic evidence in forensic text comparison. *International Journal of Speech, Language and the Law*, 17(2), 211–250. <https://doi.org/10.1558/ijssl.v17i2.211>
- Türk Ceza Kanunu. (2004). *Türk Ceza Kanunu* (Kanun No. 5237). *Resmî Gazete* (Sayı 25611).
- Fikir ve Sanat Eserleri Kanunu. (1951). *Fikir ve Sanat Eserleri Kanunu* (Kanun No. 5846). *Resmî Gazete* (Sayı 7981).
- Uchendu, A., Le, T., Zhang, K., & Lee, D. (2020). Authorship attribution for neural text generation. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)* (pp. 8384–8395). <https://doi.org/10.18653/v1/2020.emnlp-main.673>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30.
- Wright, D. (2017). Using word n-grams to identify authors and idiolects: A corpus approach to a forensic linguistic problem. *International Journal of Corpus Linguistics*, 22(2), 212–241. <https://doi.org/10.1075/ijcl.22.2.03wri>
- Yousaf, M. N. (2025). Practical considerations and ethical implications of using artificial intelligence in writing scientific manuscripts. *ACG Case Reports Journal*, 12(2), e01629. <https://doi.org/10.14309/crj.0000000000001629>
- Zeng, Z., Liu, S., Sha, L., Li, Z., Yang, K., Liu, S., Gašević, D., & Chen, G. (2024). Detecting AI-generated sentences in human–AI collaborative hybrid texts: Challenges, strategies, and insights. In *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence (IJCAI-24)*.
- Zou, Y., Li, P., Li, Z., Huang, H., Cui, X., Liu, X., Zhang, C., & He, R. (2024). *Survey on AI-generated media detection: From non-MLLM to MLLM*. <https://doi.org/10.48550/arXiv.2411.09259>

IDIOM STUDIES IN THE LAST DECADE: A BIBLIOMETRIC MAPPING OF RESEARCH TRACES BASED ON WEB OF SCIENCE

Nurbanu KORKMAZ

Asst. Prof., Ağrı İbrahim Çeçen University, Department of English Language and Literature,
nkorkmaz@agri.edu.tr, ORCID: 0000-0002-7254-781X

Korkmaz, N. (2025). Idiom studies in the last decade: A bibliometric mapping of research traces based on Web of Science. In O. Çınar, F. Başbuğ & H. Aydemir. (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 244–262). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X012>

ABSTRACT

This study presents a comprehensive bibliometric analysis of idiom research published between 2015 and 2025, utilising data from the Web of Science Core Collection. Despite increasing academic interest in idioms across linguistics, psycholinguistics, education, and applied linguistics, no previous study has systematically explored the intellectual structure, thematic development, and collaborative patterns of this field. To fill this gap, VOSviewer software was used to analyse 2,260 articles with the phrase "idiom" in their titles. Disciplinary distribution, research and citation trends over time, authorship and co-authorship networks, author, country, organisation, and document citation performance, bibliographic coupling of documents and authors, and co-occurrence of author keywords are all included in the study. The results demonstrate the multidisciplinary nature of idiom studies, which are prevalent in linguistics, education, cognitive science, and cultural studies. In terms of publication output and citation impact, the United States, England, and China rank at the top. *Idiom*, *phraseology*, *figurative language*, and *translation* are common keywords that illustrate the field's wide conceptual scope and its connections to applied linguistics and figurative meaning. Significant conceptual coherence is revealed by bibliographic coupling analyses, suggesting a solid intellectual foundation backed by common ideas and approaches. This study offers the first comprehensive, data-driven review of idiom research over the last ten years, providing insightful information on its major contributors, structural dynamics, and developing topics. Although limited to the Web of Science database and a single analytical tool (VOSviewer), the results establish a solid foundation for future cross-database and multi-tool research. The results can also help scholars identify important works, promote global cooperation, and investigate underrepresented fields such as idioms in digital conversations and cross-linguistic situations.

Keywords: Idiom, Bibliometric analysis, VOSviewer, Citation, Web of science

ÖZ

Bu çalışma, Web of Science Core Collection veritabanından elde edilen veriler kullanılarak 2015–2025 yılları arasında yayımlanan deyim araştırmalarına yönelik kapsamlı bir bibliyometrik analiz sunmaktadır. Dilbilim, psikodilbilim, eğitim ve uygulamalı dilbilim alanlarında deyimlere yönelik akademik ilginin giderek artmasına rağmen, bu alanın entelektüel yapısını, tematik gelişimini ve iş birliği örüntülerini sistematik biçimde inceleyen bir çalışma bugüne kadar yapılmamıştır. Bu boşluğu doldurmak amacıyla, başlığında “deyim” ifadesi geçen 2.260 makale VOSviewer yazılımı kullanılarak analiz edilmiştir. Çalışma kapsamında disiplinlere göre dağılım, zaman içindeki araştırma ve atıf eğilimleri, yazarlık ve ortak yazarlık ağları, yazar, ülke, kuruluş ve belge düzeyinde atıf performansı, belge ve yazarların bibliyografik eşleştirmesi ile yazar anahtar sözcüklerinin birlikte görülmeye örüntüleri incelenmiştir. Bulgular, deyim çalışmalarının dilbilim, eğitim, bilişsel bilim ve kültürel çalışmalar başta olmak üzere çok disiplinli bir yapıya sahip olduğunu göstermektedir. Yayın sayısı ve atıf etkisi açısından Amerika Birleşik Devletleri, İngiltere ve Çin ön sıralarda yer almaktadır. *Deyim*, *öbekbilim*, *imgesel dil* ve *çeviri* gibi yaygın anahtar sözcükler, alanın geniş kavramsal etki alanını ve uygulamalı dilbilim ile mecaz anlam çalışmalarına olan bağlantılarını ortaya koymaktadır. Bibliyografik eşleştirme analizleri, ortak kavram ve yaklaşımlarla desteklenen sağlam bir entelektüel temele işaret eden belirgin bir kavramsal tutarlılığı ortaya çıkarmaktadır. Bu çalışma, son on yıldaki deyim araştırmalarına ilişkin ilk kapsamlı ve veriye dayalı incelemeyi sunarak, alana önemli katkıda bulunanlar, alana dair yapısal dinamikler ve gelişmekte olan araştırma konuları hakkında önemli bilgiler sağlamaktadır. Web of Science veritabanı ve tek bir analiz aracıyla (VOSviewer) gerçekleştirilmesi yönüyle sınırlı olmasına rağmen, elde edilen bulgular yapılacak farklı veritabanları ve birden fazla araçla dizayn edilecek araştırmalar için bir temel oluşturmaktadır. Araştırmanın sonuçları, araştırmacıların önemli çalışmaları belirlemesine, küresel iş birliğini teşvik etmesine ve dijital iletişimlerde deyimler ile diller arası bağlamlar gibi yeterince temsil edilmeyen alanların araştırılmasına da katkı sağlayabilmektedir.

Anahtar Sözcükler: Deyim, Bibliyometrik analiz, VOSviewer, Atıf, Web of science

1. Introduction

Idioms are conventional multiword expressions whose meanings cannot be entirely inferred from the literal meanings of their parts. They serve as fixed linguistic units that convey figurative, culture-specific meanings, reflecting the cognitive and cultural frameworks of a linguistic community. According to Fernando (1996), idioms are “conventionalised multiword expressions that function as single units of meaning,” while Nunberg, Sag, and Wasow (1994) define them as “phrases whose meanings are not predictable from the meanings of their parts.” From a figurative language perspective, Glucksberg and McGlone (2001) emphasise that idioms function as metaphorical constructs that efficiently communicate complex meanings within discourse.

Idioms are of vital importance for communication skills and language proficiency. Idioms embody shared cultural knowledge, values, and experiences, functioning as linguistic carriers of culture (Kovecses, 2010). In second language acquisition (SLA), the skilful use of idioms is closely related to fluent speech and figurative ability, similar to the native language (Wray, 2002; Boers & Lindstromberg, 2008). As Heidari and Aliyar (2025) emphasise, idioms should be regarded not as peripheral elements but as “indispensable linguistic units whose mastery is essential for learners to attain communicative competence and native-like fluency”. Furthermore, idioms are found not only in informal language but also in academic speech and writing, where they perform rhetorical and evaluative roles (Miller, 2020), showing their importance across various genres and contexts.

Despite their importance, idioms remain challenging for both students and teachers due to their rigid structures and metaphorical, culturally dependent meanings (Liontas, 2017; Eyckmans & Lindströmberg, 2017). Over the past decade, research on idioms has rapidly expanded across the fields of linguistics, psycholinguistics, education, and applied linguistics, focusing on comprehension, teaching strategies, processing mechanisms, and cross-linguistic comparisons. Numerous bibliometric studies have been conducted to map research landscapes and track advances in related domains of figurative language study. For instance, Yuan and Sun (2023) used CiteSpace and Web of Science data to analyse metaphor research from 2010 to 2020 and identify key authors, organisations, and emerging hotspots. Similarly, Zhao, Zheng, and Zhao (2023) performed a scientometric review of conceptual metaphor research (2002–2022), revealing collaboration networks and evolving research themes, while Han, Peng, and Chen (2022) examined publications on Conceptual Metaphor Theory (1980–2022), highlighting key knowledge bases and research frontiers. Beyond metaphors, Yi, Zhao and Wang (2024) conducted a bibliometric analysis of multiword expression studies (1983–2024) by using CiteSpace, which includes idioms as a subcategory, reporting increasing interest in computational and applied linguistics contexts. A recent study conducted by Sun and Lin (2025) conducted a bibliometric analysis of the evolution and future direction of metonymy research between the years 2000 and 2023.

Additionally, Tauhid and Rohmah (2024) have explored the teaching of idioms within EFL contexts over the last two decades, providing insights into instructional trends and pedagogical approaches. However, these studies are either limited to specific subfields (e.g., metaphor or idiom pedagogy) or treat idioms as part of broader constructs, such as multiword expressions. To our knowledge, no cross-disciplinary, Web of Science–based bibliometric mapping has systematically examined idiom research as a distinct domain across its diverse subfields.

Bibliometric mapping offers a potential method to address this gap. Bibliometric analysis uses quantitative techniques to assess and

visualise scientific outputs, helping identify theme clusters, collaboration networks, publishing patterns, and important contributors in an area of study (Pritchard, 1969; Donthu et al., 2021). Unlike traditional narrative reviews, bibliometric mapping provides an objective, reproducible overview of a discipline's development by analysing citation patterns, co-authorship networks, and keyword co-occurrences. Visualisation tools, such as VOSviewer, facilitate the interpretation of these patterns, enabling researchers to detect emerging topics, influential works, and collaboration structures.

Bibliometric mapping can be used to track the development of multidisciplinary connections, reveal recurring themes in the literature, highlight important authors and institutions, and show how idiom studies have evolved over time. In addition to improving comprehension of the intellectual framework of idiom research, such a methodical summary points researchers toward uncharted territory and potential future paths. More important advantages of bibliometric analysis include recognising the knowledge-intensive nature of scientific research and enabling time-series evaluations of research subjects within predefined parameters (Van Raan, 2005). Therefore, this study aims to conduct a comprehensive bibliometric analysis of idiom-related research published between 2015 and 2025, using data retrieved from the Web of Science Core Collection. By employing VOSviewer, the study visualises co-authorship networks, country collaborations, citation patterns, and keyword co-occurrences to map the research landscape of idiom studies in the last decade. The findings of the study will show significant research trends and influential works, advance a meta-scientific knowledge of the field's evolution, and offer practitioners, educators, and academics interested in figurative language and applied linguistics insightful information.

The following questions are the focus of this study:

1. Which writers have produced the most in the last decade in the field of idiom studies?
2. Which countries have generated the greatest number of idiom studies?
3. Which studies have received the highest number of citations in the field?
4. Which keywords have been most frequently used by authors in idiom research?

2. Methodology

Bibliometric analysis, as originally defined by Pritchard (1969), involves using statistical and mathematical methods to examine books and other forms of written communication to identify patterns of publication, citation, and research progress. In recent years, bibliometric approaches have been extensively employed to map scientific knowledge, identify key

research areas, and reveal intellectual structures across disciplines (Donthu et al., 2021; Zupic & Čater, 2015).

This study employs a bibliometric mapping method to systematically examine the development and characteristics of idiom-related research over the past decade. The analysis examines publications indexed in the Web of Science Core Collection (WoSCC) from 2015 to 2025 that include the term “idiom*” in their title, ensuring that all selected documents explicitly focus on idioms as their main topic. The asterisk (*) enables the retrieval of records containing both singular and plural forms (e.g., idiom, idioms).

The dataset includes all available citation indexes within WoSCC to ensure comprehensive coverage of the relevant literature: Science Citation Index Expanded (SCIE, 1980–present), Social Sciences Citation Index (SSCI, 1980–present), Arts & Humanities Citation Index (A&HCI, 1975–present), Conference Proceedings Citation Index – Science (CPCI-S, 1990–present), Conference Proceedings Citation Index – Social Sciences and Humanities (CPCI-SSH, 1990–present), Book Citation Index – Science (BKCI-S, 2005–present), Book Citation Index – Social Sciences and Humanities (BKCI-SSH, 2005–present) and Emerging Sources Citation Index (ESCI, 2005–present). The retrieved records were exported in plain text format, complete with comprehensive bibliographic details, including authors, titles, abstracts, keywords, source titles, publication years, and citation counts. Duplicate and irrelevant entries were removed through manual screening.

The data collection process followed a systematic three-stage approach illustrated in Figure 1: (1) identification, (2) screening, and (3) inclusion. In the identification stage, all records containing the keyword “idiom” (n = 4183) were retrieved from WoSCC for the period 1975-2025. During the screening phase, records were filtered by publication year (2015–2025) and relevance, resulting in 2,260 documents. This refinement stems from our goal of documenting the mapping patterns observed in all papers published over the last ten years. In the inclusion phase, these 2,260 publications were selected for bibliographic analysis.

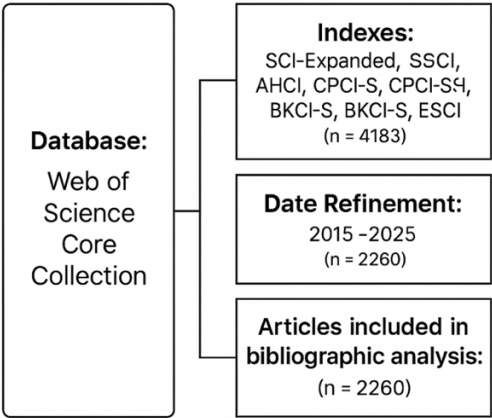


Figure 1. Data collection process for idiom studies

To analyse and visualise the bibliometric data, the software VOSviewer was used. Developed by Van Eck and Waltman (2010), VOSviewer is a widely recognised tool for constructing and visualising bibliometric networks, including co-authorship, co-citation, and keyword co-occurrence maps. The programme enables the identification of clusters representing thematic concentrations, research collaborations, and citation relationships among publications. (Dereli, 2024).

In this study, VOSviewer was used to create visual maps for co-authorship analysis (including authors, institutions, and countries), co-citation analysis (authors and journals), and keyword co-occurrence analysis (research themes and areas of interest). The bibliometric mapping approach provides an objective overview of the structure and evolution of idiom studies by highlighting influential contributors, key topics, and emerging trends. The methodology ensures transparency and reproducibility by employing a systematic search strategy, a widely recognised database (WoSCC), and established analytical tools (Zupic & Čater, 2015; Donthu et al., 2021).

3. Data Analysis

This section presents the results of bibliometric analyses conducted on studies related to terms obtained from the Web of Science Core Collection for the period 2015–2025. The analyses are structured sequentially to provide a comprehensive overview of the research environment. First, the disciplinary distribution of publications in Web of Science categories is explained, followed by an examination of the annual trends in publications and citations. In subsequent sections, dominant themes and emerging research focal points are identified by examining authors' productivity and collaboration networks, citation analyses for

authors, countries, organisations and documents, bibliographic matching models, and co-occurrence analyses of author keywords.

3.1 Web of Science categories

The dataset comprises 4,183 documents, spanning a wide range of disciplines. The most represented categories include “Language & Linguistics”, “Linguistics”, “Humanities Multidisciplinary”, “Education Research”, “Anthropology”, “Religion”, “History”, “Literature”, “Psychology Experimental”, “Computer Science Theory Methods”, and “Music”.

This broad disciplinary spread reflects the interdisciplinary nature of idiom studies, which draw from cognitive linguistics, cultural studies, psycholinguistics, education, and computational linguistics. Such diversity indicates that idioms are not confined to linguistic inquiry alone but are also studied as cognitive, cultural, and communicative constructs (Figure 2).

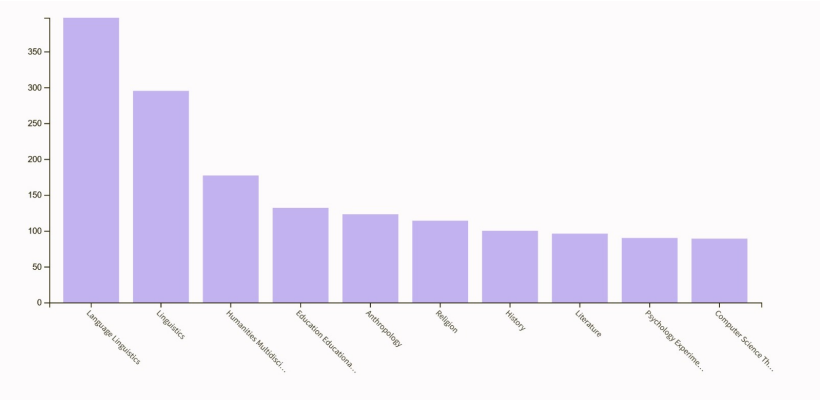


Figure 2. Web of Science categories

3.2 Annual trends in publications and citations

Analysis of publication dynamics reveals a steady rise in scholarly attention to idioms over the past decade, accompanied by fluctuations in citation counts. Certain years stand out with disproportionately high citation activity relative to publication volume, suggesting the influence of seminal works that shaped subsequent research directions. For instance, citation levels peaked around 2014, exceeding 2,000 citations, despite a relatively modest number of publications.

Overall, these temporal patterns illustrate the growing maturity and recognition of idiom studies as an independent research area, where foundational contributions continue to inform newer works (Figure 3).

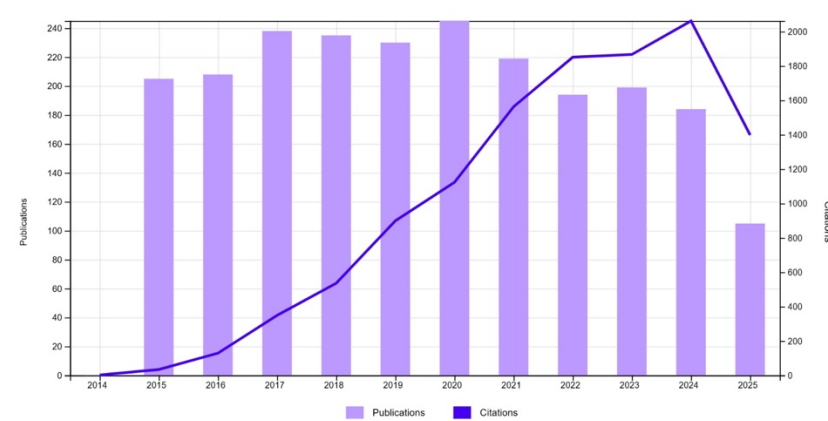


Figure 3. Distribution of the publications and citations by year

3.3 Authorship and co-authorship networks

The authorship analysis identified 3,941 contributors, of which 2,636 met the inclusion criteria of at least one publication and one citation. The co-authorship network constructed from the 50 most relevant authors reveals 18 central figures exhibiting frequent collaborative ties and shared thematic interests. Prominent nodes within the network include Bethel, Kelly J.; Bjork, Robert L.; Bloss, Cinnamon S.; Boelt, Debra L.; and Darst, Burcu F., each of whom is associated with publications that have received more than 26 citations.

Co-authorship patterns serve as a robust indicator of intellectual collaboration and the diffusion of knowledge across disciplinary and institutional boundaries. Dense interconnections among key authors indicate cohesive research clusters and collaborative research agendas (Figure 4).

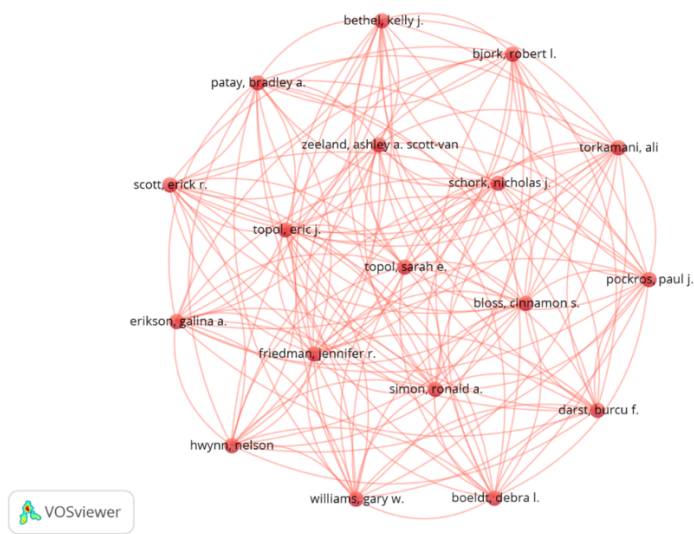


Figure 4. Co-authorship of authors

3.4 Citation analyses

According to Appio et al. (2014) and Donthu et al. (2021), citation analysis is a fundamental method for scientific mapping based on the idea that it represents the intellectual connections between publications arising from one journal citing another. To understand the intellectual dynamics of a study topic, citations can be used to examine the most significant articles in that field. In the present section, citation of authors, countries, organisations and documents will be included in turn.

3.4.1 Citation of authors

Citation analysis provides insight into the most influential contributors in idiom-related research. Among the 2,636 eligible authors, the top-cited researchers include Hinton, Devon E. (245 citations), Kohrt, Brandon A. (226), Sharma, Vivek (218), and Bass, Judith K. (202) (see Figure 5 and Table 1).

Table 1 provides a detailed overview of the 10 selected authors, including their publications and citations.

Table 1. Top 10 most cited authors

Author	Documents	Citation
Hinton, Devon E.	3	245
Kohrt, Brandon A.	4	226
Sharma, Vivek	13	218
Bass, Judith K.	3	202
Kaiser, Bonnie N.	2	202

Bolton, Paul A.	1	190
Haroz, Emily E.	1	190
Vulchanov, Valentin	6	184
Vulchanova, Mila	6	184
Zhang, Yiran	8	166

These highly cited authors serve as intellectual anchors in the field, producing seminal works that are frequently referenced in studies addressing idiom comprehension, cross-cultural interpretation, and applied pedagogy.

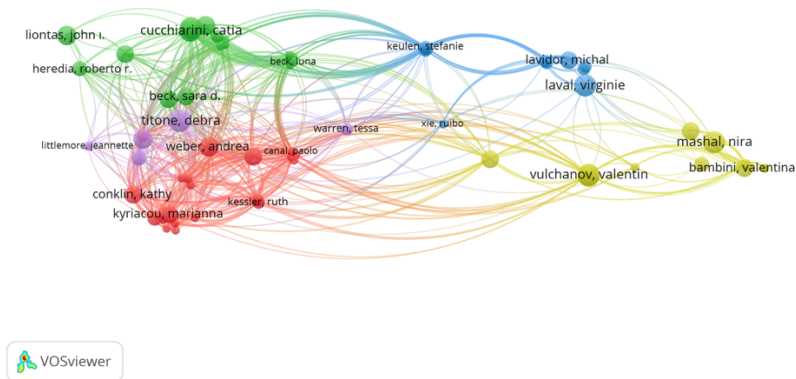


Figure 5. Citation of authors

Author citations reveal relationships among researchers based on how often they cite one another or are mentioned jointly (Figure 5). Each dot symbolises an author, and the size of the dot denotes the author's number of publications or citation effect; the larger the dot, the more significant the author. Colours indicate thematic or cooperative groupings within the field, such as groups of writers that regularly mention one another or study related subjects.

The red cluster means specific authors who focus on idioms and cognitive linguistics. While authors in the blue cluster are studying figurative language and learning a second language, authors in the green cluster are focused on the neurocognitive aspects of idiom comprehension.

3.4.2 Citation of countries

Citation analysis by country provides valuable insights into global research networks, geographical influence, and collaborative partnerships in idiom-related scholarship. Among the 108 countries represented in the dataset, each with at least one document and one citation, 98 countries met the inclusion criteria and were included in the citation network analysis.

The resulting network demonstrates a complex structure of international collaboration and citation linkages. The United States holds

a central position with the highest citation count and the strongest international connectivity, forming extensive links with Chile, Poland, Russia, Azerbaijan, Bangladesh, Georgia, Liberia, Mozambique, and Uganda. Additional regional clusters are identified across the network: a blue cluster including England, Denmark, Norway, Australia, and Türkiye; a green cluster including Germany, the Netherlands, Chile, Azerbaijan, and Switzerland; a purple cluster including China, Iran, Brazil, and Scotland; and a red cluster including Russia, Croatia, Ukraine, and Poland. These clusters indicate not only regional proximity but also shared thematic interests and co-citation patterns (Figure 6).

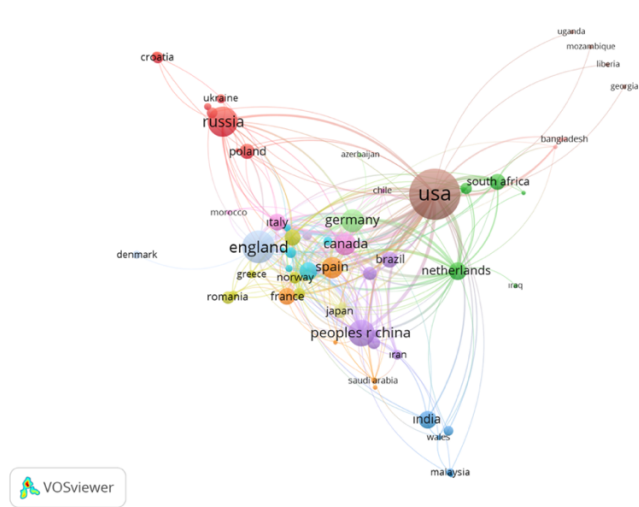


Figure 6. Citation of countries

Quantitative results further reveal that the countries with the highest total citation counts are the United States (4,853 citations), England (1,669), the Netherlands (1,211), China (1,163), and Canada (699). Other highly cited contributors include Germany (580), Spain (482), Norway (458), Italy (453), and Australia (404) (Table 2). Türkiye ranks 26th globally, with 38 publications, 41 citations, and a total link strength of 16, highlighting its emerging role in the international research community on idioms.

Table 2. Citation of the top 5 countries

Country	Documents	Citations	Total link strength
USA	497	4853	385
England	209	1669	223
Netherlands	63	1211	135
China	134	1163	158
Canada	98	699	165

Countries with a high volume of publications and strong citation linkages often serve as collaboration hubs, promoting cross-national knowledge exchange and advancing both theoretical and applied research. The presence of multiple clusters within the citation network indicates that studies on idioms have evolved through a wide range of international contributions, rather than being concentrated in a single region.

3.4.3 Citation of organisations

Institutional analysis shows that 449 organisations produced at least one publication with at least one citation. A total of 26 clusters, 2,363 links, and an overall connection strength of 2,954 were identified.

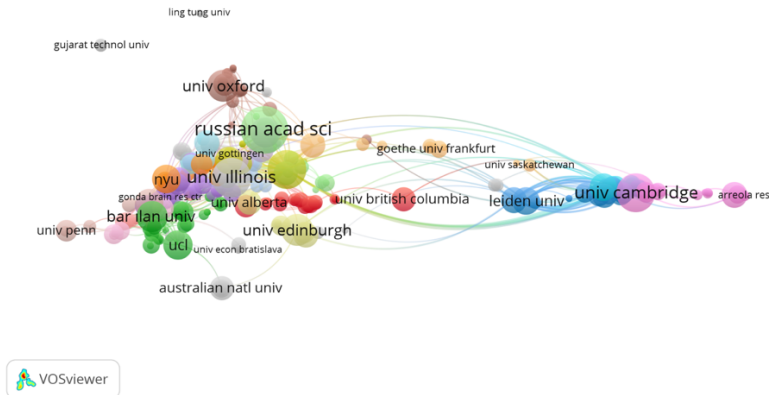


Figure 7. Citation of organisations/institutions

The most influential institutions are Harvard University (13 documents; 793 citations), University of Amsterdam (12; 615), Chinese Academy of Sciences (4; 469), McGill University (26; 443), and University of Illinois (27; 351). In Türkiye, Koç University is among the contributors, with 2 documents and 7 citations (Figure 7).

3.4.4 Citation of documents

The analysis of citation patterns at the document level provides insights into the intellectual foundations and core literature shaping idiom-related research. By examining the interconnections among publications, it is possible to identify the studies that have had the most significant influence on subsequent research and served as pivotal reference points in the field.

A total of 208 documents met the inclusion criteria of having at least one citation and were found to be interconnected within the citation network. These works collectively formed 18 clusters and 439 citation links, indicating a moderately dense network of mutual referencing and shared conceptual bases (Figure 8). The clustering structure reflects

distinct thematic areas, each anchored by seminal publications that have guided theoretical frameworks, methodological approaches, or applied research directions within idiom studies.

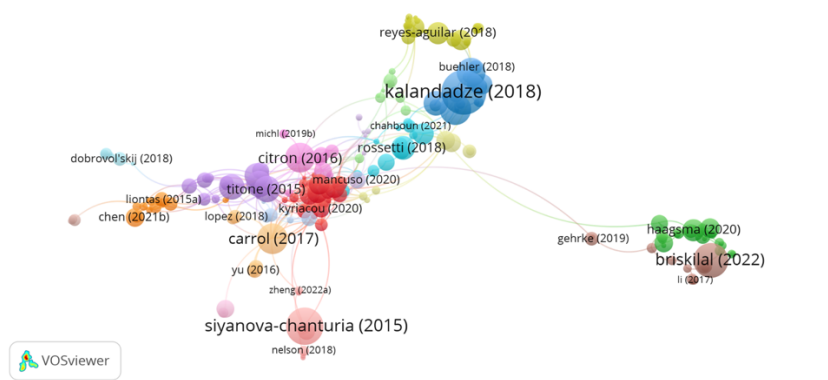


Figure 8. Citation of documents

The most frequently cited documents occupy a central position in this network and constitute the conceptual backbone of the field. As summarised in Table 3, McNally et al. (2015) ranks first with 462 total citations and an annual average of 42 citations, reflecting its sustained impact over time. Polinsky (2018) has 229 citations (28.63 per year), while Kaiser et al. (2015) has 190 citations (17.27 per year). Other highly influential publications include Drayton (2018) with 179 citations and Ray (2016) with 170 citations.

Table 3. Top 5 most cited documents and their average and total citations

Top 5 most cited documents	Citations	
	Average per year	Total
McNally et al. (2015)	42	462
Polinsky (2018)	28.63	229
Kaiser et al. (2015)	17.27	190
Drayton (2018)	22.38	179
Ray (2016)	17	170

3.5 Bibliographic coupling

Bibliographic coupling analysis measures the conceptual similarity between studies by examining the shared references between them. This method reveals intellectual connections and thematic coherence within a research field. When two documents cite the same references, they are considered bibliographically coupled, indicating that they share a common theoretical foundation or methodological approach (Dirik et al., 2023). This approach is particularly valuable for exploring research frontiers and core conceptual frameworks in emerging or interdisciplinary

domains such as idiom studies. In the following section, bibliographic coupling of documents and authors will be provided in turn.

3.5.1 Documents

The bibliographic coupling analysis of documents was conducted on 993 publications that met the inclusion criteria of having at least one citation and being connected within the network. The results revealed the formation of 16 distinct clusters, with 22,368 links and a total connection strength of 67,973, reflecting a moderate to high level of intellectual integration across the literature (Figure 9).



Figure 9. Bibliographic coupling of documents

The publications with the highest numbers of bibliographic couplings were Kaiser (2015) with 190 citations, Polinsky (2018) with 228 citations, and McNally (2015) with 459 citations. In terms of overall connection strength, the leading works were Senaldi (2022), Cieslicka (2015), and Senaldi (2024).

3.5.2 Authors

At the author level, bibliographic coupling analysis encompassed 2,095 researchers, resulting in 42 clusters and an overall link strength of 813,860, indicating extensive conceptual intersections and shared citation practices among scholars (Figure 10). Among the top five authors with the highest number of bibliographic couplings, citation counts and overall link strengths were relatively similar. Specifically, Denny Brosbroom received 459 citations with a link strength of 606; Marie K. Deserno, 459 citations with a link strength of 606; Richard J. McNally, 3,473 link strength; Donal J. Robinaugh, 2,282 link strength; and Li Wang, 604 link strength.

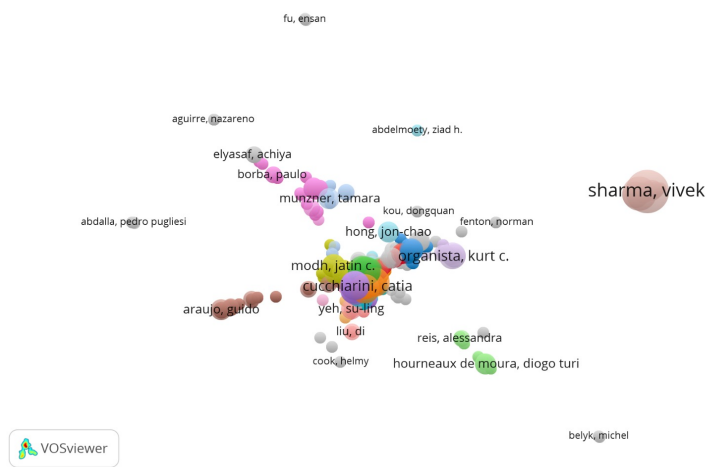


Figure 10. Bibliographic coupling of authors

Within this network, several authors stand out for their strong bibliographic linkages, indicating substantial alignment in the literature they cite. Notably, Richard J. McNally, Donald J. Robinaugh, and Li Wang exhibit the highest coupling strengths, suggesting that their work consistently draws on comparable conceptual foundations. This pattern implies that these authors contribute to related thematic domains and collectively shape the core intellectual structure of idiom research.

3.6 Co-occurrence of author keywords

Keywords are fundamental components of scientific articles that represent the core themes, areas of interest, and new developments within a field (Yan Xiang-Yun et al., 2022). By examining how frequently specific phrases appear together in the literature, the author's analysis of keyword co-occurrence enables the identification of conceptual connections. Using VOSviewer, keywords are automatically grouped into clusters that represent thematic structures and research hotspots within the field.

The analysis revealed 447 interconnected keywords that appeared at least three times across the dataset. These terms formed 20 clusters, 2,071 links, and a cumulative link strength of 2,598, indicating a well-structured, thematically diverse research field (Figure 11). Among the most frequently used keywords, “*idiom*” occurred 137 times, followed by “*phraseology*” (93), “*translation*” (51), “*figurative language*” (47), and “*metaphor*” (41). In terms of total link strength, “*idiom*” ranked highest (280), demonstrating its central role as the primary research focus, while “*phraseology*” (182), “*figurative language*” (116), “*translation*” (105), and “*metaphor*” (98) also exhibited strong interconnections.

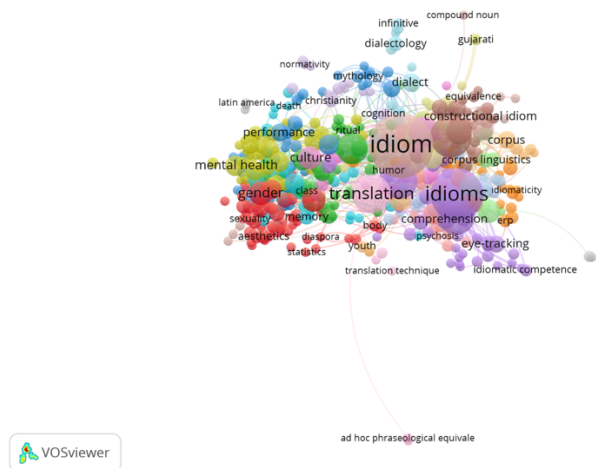


Figure 11. Co-occurrence of author keywords

These findings indicate that idiom-related research is closely integrated with broader inquiries into phraseological studies, figurative meaning, and translation practices, highlighting its multidimensional character across applied linguistics, cognitive linguistics, and education.

If we are to consider additional key terms, the *comprehension* of idioms, *constructional idioms*, *gender* differences, and the relationship between idioms and *mental health* are particularly noteworthy.

Overall, the keyword co-occurrence network demonstrates a cohesive and evolving intellectual landscape. The dominance of terms such as *idiom* and *phraseology* suggests a stable core of established themes, while the presence of terms like *figurative language* and *translation* reflects the field’s interdisciplinary expansion. This pattern highlights the increasing conceptual integration of idiomatic studies within broader linguistic and cognitive frameworks.

4. Conclusion

This bibliometric study provides the first comprehensive mapping of idiom research published between 2015 and 2025, based on the Web of Science Core Collection. Using quantitative analysis and visualisation techniques through VOSviewer, the study systematically identifies the main contributors, research trends, and thematic structures within the field. The findings offer valuable insights into the intellectual landscape and developmental trajectory of idiom studies across various disciplinary contexts.

The findings of this study highlight several important points regarding the research questions. Scholars like Hinton, Kohrt, Sharma, Bass, and Kaiser stand out with the greatest citation counts in terms of author productivity and influence, indicating their substantial contributions

to idiom-related studies. Secondly, the geographic distribution of research outputs reveals that the top nations in terms of publication volume and citation impact are the United States, England, the Netherlands, China, and Canada. These findings indicate a strong academic infrastructure and global collaborative networks. Türkiye's position, ranked 26th in terms of increasing scholarly engagement, indicates its growing participation in global idiom research. Thirdly, at the document level, highly cited publications such as Polinsky (2018) and McNally et al. (2015) make fundamental contributions that influence the field's theoretical and methodological approaches. Finally, the keyword co-occurrence analysis highlights the dominance of terms such as *idiom*, *phraseology*, *figurative language*, and *translation*, suggesting that idiom research is closely connected with broader inquiries into figurative meaning, phraseological studies, and applied linguistics. These findings collectively demonstrate the interdisciplinary and evolving nature of idiom scholarship.

Significant conceptual coherence throughout the subject was further demonstrated by the bibliographic coupling studies of documents and authors, which identified overlapping theoretical frameworks and common reference bases. Strong coupling clusters show that idiomatic studies remain open to new subfields such as computational linguistics and translation studies while coming together around fundamental topics in cognitive linguistics, psycholinguistics, and education.

From a meta-scientific perspective, this study not only documents the structural patterns of existing research but also points towards future directions. Linguistic comparisons, idioms in digital discourse, and pedagogical applications in various linguistic and cultural contexts are among the underrepresented areas that could be given greater emphasis. Furthermore, strengthening international collaboration, particularly encompassing new research areas, will further enrich the diversity and impact of this field.

Finally, by providing a systematic, data-driven summary of idiom research carried out during the last ten years, this bibliometric mapping adds to the body of literature. To the best of the researcher's knowledge, no bibliometric study has yet been done on idioms. This study offers important insights into the scope and development of idiom research over the past 10 years, despite several limitations detailed below. By bringing together diverse findings within a coherent framework, it lays the foundation for more focused and collaborative work in the fields of idiom studies and figurative language. Furthermore, it is an indispensable tool for academics seeking to understand the conceptual framework of this field, identify significant contributions, and explore new avenues for future research.

5. Limitations of the Study

The limitations of this study can be discussed in two main aspects: the process of constructing the dataset and the software employed for analysis.

Firstly, the dataset is restricted to publications indexed in the Web of Science database. Articles from other databases, such as Scopus, TRDizin, YÖKTEZ, PubMed, or EBSCO, were not included, and their inclusion could potentially influence the findings. Secondly, the analysis was conducted using VOSviewer, a free and user-friendly tool. While this software is widely used, alternative tools such as CiteSpace, Bibliometrix(R), or Litmap could be employed, and the results may differ depending on the chosen method to some extent.

Future research could benefit from using datasets from multiple databases and conducting comparative analyses across different software tools to achieve a more comprehensive understanding of the field.

References

- Appio, F. P., Cesaroni, F., & Di Minin, A. (2014). Visualizing the structure and bridges of the intellectual property management and strategy literature: A document co-citation analysis. *Scientometrics*, 101(1), 623–661. <https://doi.org/10.1007/s11192-014-1329-0>
- Boers, F., & Lindstromberg, S. (Eds.). (2008). *Cognitive linguistic approaches to teaching vocabulary and phraseology* (Vol. 6). De Gruyter.
- Dereli, A. B. (2024). Bibliometric analysis with VOSviewer. *Communicata*, 28, 1–7.
- Dirik, D., Eryılmaz, İ., & Erhan, T. (2023). Post-truth kavramı üzerine yapılan çalışmaların VOSviewer ile bibliyometrik analizi. *Sosyal Mucit Academic Review*, 4(2), 164–188.
- Donthu, N., Kumar, S., Mukherjee, D., Pandey, N., & Lim, W. M. (2021). How to conduct a bibliometric analysis: An overview and guidelines. *Journal of Business Research*, 133, 285–296. <https://doi.org/10.1016/j.jbusres.2021.04.070>
- Eyckmans, J., & Lindstromberg, S. (2017). The power of sound in L2 idiom learning. *Language Teaching Research*, 21(3), 341–361. <https://doi.org/10.1177/1362168816652411>
- Fernando, C. (1996). *Idioms and idiomaticity*. Oxford University Press.
- Glucksberg, S., & McGlone, M. S. (2001). *Understanding figurative language: From metaphor to idioms*. Oxford University Press.
- Han, Y., Peng, Z., & Chen, H. (2022). Bibliometric assessment of world scholars' international publications related to conceptual metaphor. *Frontiers in Psychology*, 13, Article 1071121. <https://doi.org/10.3389/fpsyg.2022.1071121>
- Heidari, K., & Aliyar, M. (2025). Thirty-five years of research on idioms in second language acquisition: A methodological review. *Research Synthesis in Applied Linguistics*, 1–23.

- Kövecses, Z. (2010). *Metaphor: A practical introduction*. Oxford University Press.
- Liontas, J. I. (2017). Why teach idioms? A challenge to the profession. *Iranian Journal of Language Teaching Research*, 5(3), 5–25.
- Miller, J. (2020). The bottom line: Are idioms used in English academic speech and writing? *Journal of English for Academic Purposes*, 43, Article 100810. <https://doi.org/10.1016/j.jeap.2019.100810>
- Nunberg, G., Sag, I. A., & Wasow, T. (1994). Idioms. *Language*, 70(3), 491–538. <https://doi.org/10.1353/lan.1994.0007>
- Pritchard, J. (1969). Statistical bibliography or bibliometrics? *Journal of Documentation*, 25(4), 348–349.
- Sun, Y., & Lin, M. (2025). A bibliometric analysis of metonymy in SSCI-indexed research (2000–2023): Retrospect and prospect. *Frontiers in Psychology*, 16(1), 1–15.
- Tauhid, A. B., & Rohmah, Z. (2024). Global insights into idiom teaching for English language learners: A 20-year bibliometric review. *Jurnal Pendidikan Bahasa Inggris Undiksha*, 12(2), 139–149.
- Van Eck, N. J., & Waltman, L. (2010). Software survey: VOSviewer, a computer program for bibliometric mapping. *Scientometrics*, 84(2), 523–538. <https://doi.org/10.1007/s11192-009-0146-3>
- Van Raan, A. F. J. (2005). Fatal attraction: Conceptual and methodological problems in the ranking of universities by bibliometric methods. *Scientometrics*, 62(1), 133–143. <https://doi.org/10.1007/s11192-005-0008-6>
- Wray, A. (2002). *Formulaic language and the lexicon*. Cambridge University Press.
- Yan, X.-Y., Yao, J.-P., Li, Y.-Q., Zhang, W., Xi, M.-H., Chen, M., & Li, Y. (2022). Global trends in research on miRNA–microbiome interaction from 2011 to 2021: A bibliometric analysis. *Frontiers in Pharmacology*, 13, Article 974741. <https://doi.org/10.3389/fphar.2022.974741>
- Yi, W., Zhao, Y., & Wang, W. (2025). Four decades of research on multiword expressions: A bibliometric analysis. *Applied Linguistics*, amae085. <https://doi.org/10.1093/applin/amae085>
- Yuan, G., & Sun, Y. (2023). A bibliometric study of metaphor research and its implications (2010–2020). *Southern African Linguistics and Applied Language Studies*, 41(3), 227–247. <https://doi.org/10.2989/16073614.2023.2201897>
- Zhao, X., Zheng, Y., & Zhao, X. (2023). Global bibliometric analysis of conceptual metaphor research over the recent two decades. *Frontiers in Psychology*, 14, Article 1042121. <https://doi.org/10.3389/fpsyg.2023.1042121>
- Zupic, I., & Čater, T. (2015). Bibliometric methods in management and organisation. *Organizational Research Methods*, 18(3), 429–472. <https://doi.org/10.1177/1094428114562629>

YA DIŞINDASINDIR DİL BİLİMİN YA DA İÇİNDE

Emel KÖKPINAR KAYA

Doç. Dr., Hacettepe Üniversitesi, İngiliz Dilbilimi Bölümü, emelkokpinar@hacettepe.edu.tr, ORCID: 0000-0001-5319-6527

Kökpinar Kaya, E. (2025). Ya dışındasındır dilbilimin ya da içinde. İçinde O. Çınar, F. Başbuğ & H. Aydemir (Haz.), *Çağdaş dilbilim yazıları I: İMÜ Dilbilimi 15. yıl armağanı* (ss. 263-272) Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X013>

ÖZ

Söylem kavramının ele alınma biçimlerini ortaya koyduğum bu yazıda, söylemin incelenme biçimlerinin de altını çizmekteyim. Özellikle anti-pozitivist, yorumcu ve eleştirel bakış açısından ortaya konulan söylem çözümlemesi çalışmalarını irdelemeyi amaçlamaktayım. Bu çalışmaların felsefik ve kuramsal duruşlarını göstermek amacıyla öncelikle Söylem Çözümlemesi alanının tarihçesine değinmekteyim. Bunu takiben alandaki kendine has duruşu ile Eleştirel Söylem Çözümlemesi alanını, yöntembilimsel duruşunu ve alanın akademik dünyadaki kabul ediliş biçimlerine yönelik bir anlayış sunmaya çabalamaktayım. Son olarak, alanın Türkiye Dilbilim çevresindeki konumuna yönelik eleştirel duruşumu sunmaktayım.

Anahtar Sözcükler: Eleştirel söylem çözümlemesi, Söylem çözümlemesi, Anti-pozitivist paradigma, Yorumcu paradigma

ABSTRACT

In this article that I present the ways in which the concept of discourse is addressed, I also underline the various methods of analyzing discourse. I specifically aim to examine discourse analysis studies conducted from anti-positivist, interpretive, and critical perspectives. To reveal the philosophical and theoretical stances of these studies, I first address the history of Discourse Analysis. Then, I endeavour to provide an understanding of Critical Discourse Analysis—with its unique position in the field—its methodological stance and the ways that it has been accepted in the academic world. Finally, I present my critical position regarding the status of the field within the Turkish Linguistics circle.

Keywords: Critical discourse analysis, discourse Analysis, Anti-positivist paradigm, Interpretivist paradigm

1. Giriş

Ya dışındadır çemberin

Ya da içinde yer alacaksın.

Kendin içindeyken kafan dışındaysa...

Yazıma Murathan Mungan'ın dizileri ile başlamak istedim. Dizelerdeki çember bana bir eleştirel söylem çözümlemecisi olarak dilbilimi hatırlatır. “Ya dışındadır dilbilimin ya da içinde yer alacaksın” diye tekrarlar ve çalışma alanım olan söylem çözümlemesinin dilbilim içindeki tartışmasız yerini hem de tartışmalı varoluşunu düşünürüm. İşte bu bağlamda, bu yazıda eleştirel söylem çözümlemesinin Türkiye dilbilim çemberindeki yerini irdelemek istiyorum. Bunu yaparken çemberin dışındaki akademik gelişmelere bakmak ve bir söylem çözümlemecisinin kendini konumlandırma çabasını gözler önüne sermeyi de arzuluyorum.

Bu yazıya yoruma son derece açık dizeler ile başlamam ve ‘ben’ dilini kullanmam bir tesadüf ya da bir yazagelmışlik neticesinde oluşmadı. Bunları bir dil kullanıcısı olarak ve dil kullanımının toplumsallıktaki gücünü göz önünde bulundurarak belli bir farkındalık ile kullanıyorum. Bu yazıyı mümkün olduğunca güç ilişkilerinin bir nesnesi olarak değil, bu güç ilişkilerinde kendi eylemlerine karar verebilen ve bu eylemlerinin sorumluluklarını göz önünde bulundurabilen bir ‘özne’ konumundan yazmak istiyorum. Şöyle ki, seçtiğim bu yazma biçimini bilinçli olarak kullanıyorum; aklımdakileri ne öylesine harflere döktüm ne de akademik bir metni nasıl yazmam gerektiğini direten bir yazma biçiminin sonucu olarak belli şekillerde metinleştirdim. Aksine bu yazıyı toplumsal bir aktör olarak ‘ben’ yazıyorum ve toplumda var olan akademik ve bilimsel özellikteki yazmanın nasıl olması gerektiğine yönelik söyleme karşı bir duruş sergiliyorum. Ayrıca, yoruma açık dizeler ile başlamam da bu duruşum ile örtüşüyor. Ben bu dizeleri Eleştirel Söylem Çözümlemesi (ESÇ) ve dilbilim ilişkisi bağlamında yorumluyorum ve yorumlarımı okuyucuma bir bakış açısı ve bir anlayış olarak sunuyorum.

2. Söylem, Söylem Çözümlemesi, Eleştirel Söylem Çözümlemesi

Söylem çözümlemesinin dilbilim ile olan yadsınamaz ilişkisinin derinlemesine ele alınması gerektiğinin bir zorunluluk olduğunun farkındayım. Bu ilişki dil ve söylem arasındaki diyalektik ilişkinin doğası gereği çetrefilli bir ilişkidir. Bu ilişkinin anlaşılması ancak ve ancak toplum, biliş ve çok katmanlı bir özellik içeren ‘bağlam’ gibi öge ve kavramların anlaşılması ile gerçekleşebilir. Dilbilim ve Söylem Çözümlemesi arasındaki ilişki de son derece çok yönlü bir anlayış ile mümkün olabilecektir. Bu açıdan, yazımda, Söylem Çözümlemesi’nin Türkiye Dilbilimi içindeki yerini tartışmaya geçmeden önce Söylem Çözümlemesi’nin ne olduğunun, Söylem Çözümlemesi altındaki farklı

bakış açılarının ve bu bakış açıları çerçevesinde kavranabilecek paradigma ayrımlarının altını çizmek istiyorum.

Douglas Biber, Ula Connor ve Thomas A. Upton 2007’de yazdıkları *Discourse on the Move* adlı kitaplarının söylem çözümlemesini ve kapsamını tartıştıkları ilk bölümünde, Deborah Schiffrin, Deborah Tannen ve Heidi Hamilton tarafından 2001’de hazırlanan *Handbook of Discourse Analysis* kitabının girişine atıfta bulunarak söylem çözümlemesi çalışmalarını üç ulam altında toplamışlar ve tartışmışlardır: (i) dil kullanımının çalışılması, (ii) ‘tümce ötesi’ dilsel yapının çalışılması, (iii) dilsel ve dil-dışı öğeleri de kapsayan göstergesel kullanımlar üzerinden toplumsal pratiklerin ve ideolojik varsayımların çalışılması

En basit tanımıyla, Söylem Çözümlemesi söylemi çözümler. Bu çok basite indirgenmiş söylem çözümlemesi tanımı kendi içinde birçok tartışmaya gebe. Bu tartışmalardan en belirgin olanı ‘söylem’ kavramının ne olduğu üzerinedir. Yukarıda sunduğum ulamlamayı ele alırsak söylem, ‘dil kullanımı’ olarak düşünülebildiği gibi tümce yapısı ötesindeki metinselliği ele alır şekilde de anlaşılabilir. Hatta bunların da ötesine giderek söylemi dilsel öğeler gibi her türden göstergesel kaynak ile vücut bulan toplumsal bir eylem olarak da etiketlemek mümkündür. Bu tanımlarla çok yönlülüğü ortaya konulan söylemi tanımlama biçimi onu ele alış biçimlerini de içkin şekilde değiştirecektir. Bu nokta da bu tanımlara göz atmamız gerektiğini düşünüyorum çünkü ancak bu şekilde söylemin ve dolayısıyla Söylem Çözümlemesi’nin doğası hakkında bir kanıya varmamız mümkün olabilir.

Söylemi dil kullanımı olarak tanımlama eğilimi geleneksel yapısal dilbilim çalışmalarının dolaylı etkisi ile ortaya çıkmış ve tümce dilbilgisinin çok ötesine geçmiştir. Söylemi belirli dilsel öğelerin kullanımı ile özdeşleştiren bu tanım yapısal bir tanım olarak düşünülebilir. Ses gerçekleştirmeleri, tonlama, sözcükbirimler ve sözdizimsel yapılanmalar gibi dilbilgisel öğelerin kullanımına odaklanan bu tanıma göre söylem, dil kullanımlarını bağlam-temelli değişkenler olarak ele alır. Bu çerçevede bu tanım geleneksel dilbilgisi bileşenlerine odaklanmakta ve bağlama bağlı olarak ortaya çıkan farklı yapısal değişkenlerin betimlendiği ve açıklandığı bir çalışma çerçevesine işaret etmektedir. Burada Söylem Çözümlemesi’nin çok bilindik bir tanımı karşımıza çıkmaktadır: Söylem Çözümlemesi kendi toplumsal bağlamındaki dil kullanımının çalışılmasıdır. Bu tanımı burada bırakarak bir diğer tanıma geçeceğim ancak tanıma içerdiği toplumsallık nedeniyle ilerideki bölümlerde geri döneceğim.

Söylemi tanımlarken ele alınabilecek bir diğer kabul görmüş tanım tümce ötesindeki dilsel yapıdır. Söylemi tümce dilbilgisinden çıkaran bu bakış açısına göre, söylem daha geniş dilsel yapılar şeklinde düşünülebilir. Bu noktada, sözceler ve tümcelerin bir dizi halinde ve belli bir sistematik içinde düzenlenerek oluşturduğu daha geniş dil kullanımından söz edebilir. Geleneksel tümce dilbilgisinden, özellikle öbek ve tümcecik sözdiziminden uzaklaşan bu tanıma göre, Söylem Çözümlemesi üretilen

metinlerin dilsel yapısı ile ilgilenmektedir. Bu bağlamda, ilk tanımla ortak özellikler taşımaktadır çünkü her ikisinde de dilbilgisi bileşenlerine odaklanılmakta ve bunların insan iletişimde nasıl kullanıldığı ele alınmaktadır (Biber ve diğ., 2007, s. 2).

Üçüncü tanım ise diğerlerinden birçok yönüyle farklılaşmaktadır. Söylemin göstergesel kaynaklar aracılığı ile ortaya konulan toplumsal pratikler ve ideolojik varsayımlar olarak tanımlanması, gerek söylem çözümlemesi çalışmalarının odağı gerek ise çalışma nesnesini ele alış biçimi açısından ilk iki tanımdan farklı bir yaklaşım ortaya koymaktadır. Önceki tanımlarda temel odak, dilbilgisel yapıların kullanımı ya da metin biçimi üzerinedir. Söylemi bu bağlamda tanımlayan bu çözümleme geleneklerinin temel amacı ise belli dilbilgisel ve metinsel yapıları bağlam temelli olarak betimlemektir. İşte bu çalışma nesneleri ve buna yönelik betimlemeci duruşu ile önceki tanımlar, son tanımdaki toplumsal-kültürel yönelimler ve güdülen söylem katılımcılarının iletişim eylemlerini anlama amacı açısından son derece farklıdır. Bu bağlamda, söylemi toplumsal pratikler ve ideolojik varsayımlar bütünü olarak etiketleyen söylem çözümlemecilerinin temel odağı ne sadece dilbilgisel yapılar ne de metin biçimidir. Temel odak metinlerdeki güç ilişkilerine yöneltilmiş, metinlerde vücut bulan eşitsizlik, ayrımcılık ve baskı gibi toplumsal sorunların gösterilmesine ve hatta eleştirilmesine kaymıştır. Bu açıdan sadece betimlemeyi amaçlamayan ve dil kullanımı ile ortaya koyulan/çıkan toplumsal sorunlara yönelik bir eleştiri ve anlayış sunmayı amaç haline getirmiş bir söylem çözümlemesi geleneğinden söz edebiliriz.

Bu noktada, daha önce dile getirdiğim temel bir tanıma yeniden dönmek istiyorum: Söylem kendi toplumsal bağlamındaki dil kullanımıdır. Bu tanım söylemin dil kullanımı ile özdeşleştiğinin altını çizer. Dil kullanımını ‘dil’ kavramından ayırmak imkânsız gibi gözükse de ve hatta birçok bağlamda birbirlerinin yerine kullanıldıkları ile sıkça karşılaşılsa da, Söylem Çözümlemesi geleneğinde ‘dil kullanımı’ kavramı ön plandadır. Şöyle ki, Söylem Çözümlemesi bağlamdan bağımsız bir tümce dilbilgisi ile ilgilenmez ya da dil çalışmalarını zihinsel gerçekliğe indirgemez. Aksine, dilsel öğeleri en sınırlı bağlam türü olan ‘dilsel bağlamdan’ çok daha geniş toplumsal-kültürel, tarihsel ve politik bağlamlara bağlı olarak inceler. Dilin bağlamsal kullanımına dayanan bu çerçeveden baktığımızda, edinç değil edim temelli bir arayıştır. Dili kullanarak ortaya koyduğumuz her türlü etkileşimsel veriyi inceleyen Söylem Çözümlemesi bu yönüyle metin-temelli bir çalışma alanı olarak etiketlenebilir. Metinler dilin toplumsal edim ile örtüştüğü iletişimsel üretimlerdir ve her türlü söylem çözümlemesinin inceleme nesnesidir. Yukarıda bahsettiğim ilk tanımda metinlerdeki dilbilgisel öğeler çalışma odağındayken, ikinci tanımda metnin kendisi odağa yerleşmiş durumdadır. Son tanımda ise, metinlerin, güç ilişkisi sonucunda ortaya çıkan ve toplumda sorun teşkil edebilecek her türden eşitsizlik ve ayrımcılığın bünyesinde izlenebildiği yapılanmalar olduğu kabul edilir.

Dil kullanımı ve metin ilişkisine değinmek yukarıdaki tanımları anlamlandırmakta yeterli değildir. Burada dil kullanımının bağlam temelli olduğunun altını da çizmek gerekir. Söylem Çözümlemesi'nin dilsel öğeleri kendi bağlamı çerçevesinde incelediğinden daha önce söz etmişim. Buradaki tartışma ise 'bağlam' kavramı üzerinedir. Kendi bağlamı derken ne kastedilmektedir? Bu noktanın biraz irdelenmesi söylem çözümlemesinin ne olduğunu daha belirgin şekilde ortaya koyacaktır. Ruth Wodak ve Michael Meyer 2001'de yazdıkları ve Eleştirel Söylem Çözümlemesi yöntemlerini ele aldıkları kitaplarında bağlamı dört aşamada ele almışlardır: dil-içi ve metin-içi dilsel bağlam, sözceler, metinler ve söylemler arasındaki metinlerarası ve söylemlerarası ilişki, 'durumsal bağlam' olarak adlandırılan ve orta büyüklükte kuramlar ile açıklanan dil-dışı (toplumsal) bağlam ve daha geniş toplumsal-politik ve tarihsel bağlamlar. Eğer ki söylem dediğimiz kavram kendi bağlamı içindeki dil kullanımı olarak düşünülürse, burada hangi bağlamın göz önünde bulundurulacağı yine söylemin tanımları çerçevesinde ortaya çıkmaktadır. İlk tanımları düşünerek yorumlarsak, dilsel öğelerin içinde bulundukları anlık duruma göre betimlenmesi durumunda işaret edilen durumsal bağlam iken, tümce ötesi metinselliğe odaklanan ikinci tanımda ise en basit şekilde dilsel bağlam olarak düşünülebilir. Son tanımda ise, baskın bağlam anlayışı metinlerarası ve söylemler arası ilişkiyi belirleyen bağlam ile çok daha geniş toplumsal, kültürel, tarihsel ve politik bağlamları da kapsar. Bu yönüyle baktığımızda, her durumda yapılan çözümlemeyi dilsel yapının ve bağlamın ortaklaşa çözülmesi olarak düşünülebiliriz.

Geldiğimiz noktada, bu üç tanımın çalışma nesnesi olan söylemi tanımlama biçimleri, temel çalışma odakları ve dil kullanımını bağlantılı olarak inceledikleri bağlama yönelik anlayışlarını irdedim. Bundan sonra, daha önceki bölümlerde yer vermediğim özellikle ilk iki tanımları son tanımdan ayıran yöntembilimsel yönelimleri tartışmak istiyorum. Böylelikle yazımın en başında bahsettiğim duruş farkını da daha anlaşılır kılmayı amaçlıyorum. Son tanımın yöntembilimsel duruşunu anlayabilmek, öncekilerden nasıl ayrıştığını görebilmekle mümkündür. Bu açıdan önceden tartıştığım odak ve bağlama yönelik anlayış farkı bir bakış açısı sunmaktadır. Yine de Söylem Çözümlemesi'nin kısa bir tarihçesi bu farkı ortaya koyabilmek ve yöntembilimsel ayrışmaları daha görünür kılmak açısından önemlidir. Bu amaçla yazıma Söylem Çözümlemesi'nin dilbilim ve Toplumsal Bilimler alanı içinde nasıl ortaya çıktığına ve tarihsel süreçlerine değinerek devam etmek istiyorum.

Teun A. van Dijk (1985, s.1) Söylem Çözümlemesi'nin kökenini Antik Yunanlıların sözbilim (İng. rhetorics) geleneğine dayandırır. En basit haliyle dilin ikna için kullanımı olarak da düşünebileceğimiz sözbilim, Antik Yunan demokrasisi gereğince politik ve hukuki mercilerde ikna etme amacı güderek dilin daha planlı ve etkili olarak kullanılması çalışmalarını ve edimlerini kapsar. Bu noktada, modern dilbilim alanlarından biçimbilim ve söylemin ilk tanımını oluşturabilecek dil kullanımının daha yapısal çözümlendiği çalışmalara öncülük eder.

Zaman içerisinde kardeşi dilbilgisine yerini kaptıran sözbilim, özellikle 19. Yüzyıl dilbilimi ile ortaya çıkmaya başlayan ve Ferdinand de Saussure ile zirveye ulaşan dilin yapısal özelliklerinin incelenmesi geleneği ile yerini daha dizgeleşmiş ve toplumsallaşmış bir dil çalışmaları yönelimine bırakır (Kocaman, 2003, s.1). Bu bağlamda ortaya çıkan başat çabalardan birini 1920'li yıllarda Rus Biçimcileri'nin çalışmaları olarak düşünebiliriz. İnsan iletişimini ve iletişimin metinsel boyutunu irdeleyen Vladimir Propp'un başını çektiği bu dil çalışması geleneği yapısalcı anlatı çözümlemesi geleneği ile örülmüştür. Bu özelliği ile Rus Biçimcileri'ni belli bir metin türünün yapısal özelliklerine odaklanması ile modern Söylem Çözümlemesi çalışmalarının öncüleri arasında kabul edebiliriz.

Yapısalcı gelenek söz konusu olduğunda akıllara Fransız yapısalcı antropolog Claude Levi Strauss'un 1940'lardaki söylence çalışmaları da gelmektedir. Söylence çalışmaları 1960'larda yine Fransız geleneğinde kendini göstermeye devam etmiş, özellikle Roland Barthes ve Algirdas Greimas'ın göstergebilimsel çözümlemelerine temel oluşturmuştur. Rus Biçimcileri'ne benzer şekilde yazınsal metinlerin ele alındığı bu gelenek dil kullanımı gibi göstergesel kaynakların kültür temelli olarak çözümlemesini gerçekleştirmiştir. Bu özelliği ile sembolik üretimler olan metinlerin kültürel ve toplumsal bağlamda bir okumasını sunması ile Fransız göstergebilimciler, Söylem Çözümlemesi geleneğinin ortaya çıkmasında belirgin bir adım teşkil etmiştir.

'Söylem Çözümlemesi' ifadesi ilk olarak 1952'de yapısalcı dilbilim geleneği içinde Zellig Harris tarafından kullanılmıştır. Saussure geleneği içinde yoğunlaşmış görüşleri ile Harris söylemi, tümce bazında çözümlemelerin anlamsal olarak işlemlenebilmesi için tümce ötesinde bir yapı birimi olarak ele almıştır. Bu yaklaşımı ile Söylem Çözümlemesi alanının ortaya çıkmasında katkısı bulunmaktadır. Yine yapısalcı çerçeve içinde Prag Dilbilim Çevresi'nin, özellikle Vilém Mathesius'un tümce işlevine yönelik dizgeci bakış açısı ve Çekya'daki işlevsel dilbilimcilerin çalışmaları da Söylem Çözümlemesi'nin bir alan olarak ortaya çıkmasında rol sahibidir. İşlevsel bakış açısı söz konusu olduğunda Michael A. K. Halliday'ın Dizgeci İşlevsel Dilbilgisi modeli ve aynı gelenek içerisindeki dilbilimcilerin Eleştirel Dilbilim çalışmaları da Söylem Çözümlemesi'ne kapı açmıştır. Dilsel öğelerin bağlam temelli işlevlerine odaklanan bu çalışmalar, günümüzün baskın Söylem Çözümlemesi anlayışının da bebek adımları olarak düşünülebilir.

Söylemin yapısal özelliklerine odaklanan en belirgin çalışmalar ise yukarıdaki entelektüel ve analitik birikimin bir getirisi olarak 1970'lerde Avrupa'da ortaya çıkmıştır. Bu söylem çözümlemesi çalışmalarını anlayabilmek için yukarıdakilere ek olarak Noam Chomsky tarafından ortaya konulan Üretici Dilbilgisi modelini de göz önünde bulundurmanın yerinde olacağı kanısındayım. Dilin, edinç bazında sınırlı sayıda bileşenden sınırsız sayıda üretim yapabileceği temel fikrine dayanan bu dilbilgisi modeli, sınırsız üretimi ortaya koyacak evrensel dilbilgisi kurallarının arayışındadır. Bu bakış açısını benzer şekilde dilin edimsel

boyutu olarak düşünölebilecek metinlere uyarlama arayışı, metinselliğin ölçütlerinin belirlenmesi ile Söylem Çözömlemesi'ne katkı sağlamıştır. Metindilbilim (İng. Text Linguistics) adıyla bildiğimiz bu alan, tümce ötesi yapısal ilişkilerin çalışılmasına olanak yaratmıştır. Dilbilimin bir alt alanı olarak kabul edilebileceğimiz bu alan Söylem Çözömlemesi geleneği içinde büyük bir öneme sahiptir. Bu alan özellikle metinsellik ölçütlerine odaklanması ile tümce ötesi yapısal ilişkilere dikkat çekmiş ve Söylem Çözömlemesi alanının başat çalışma alanlarından biri olarak ortaya çıkmıştır. Bu özelliği ile Metindilbilim söylem çözömlemesi gerçekleştiren araştırmacılara gerek yaklaşımsal gerek ise kavramsal bir katkı sunmuştur. Bu katkı ile söylem tanımlarından ikincisi belirlemeye başlamış ve tümce ötesinden metin yapısına yönelen bakış açısı ile sayıca çok metinlerin söylem çözömlmeleri gerçekleştirmeye başlamıştır.

Daha önce söz ettiğim üçüncü söylem tanımının – söylemin ideolojik varsayım ve toplumsal pratik olarak kabul görmesinin – ortaya çıkışını değerlendirdiğimizde de Metindilbilim çalışmalarının katkılarını yadsıyamayız. Yapısal gelenek içinde beliren Dizgeci İşlevsel Dilbilgisi'nin katkılarına ek olarak ESÇ'nin önemli kuramcılarından van Dijk'in metinsellik ölçütlerinden biri olan bağdaşıklık üzerine çalışmalar da söyleme bakış açısını değiştirmiştir. Van Dijk metinlerin mantıksal yapısı fikrine odaklanan bağdaşıklık üzerine getirmiş olduğu metindilbilimsel çalışmalarını söylem ve biliş ilişkisine yönelmiş ve Söylem Çözömlemesi geleneğinde yeni bir söylem tanımının ortaya çıkmasına da ön ayak olmuştur. Bu yeni tanım, gelenek içinde eleştirel çalışmaların var olmasına olanak sağlamış ve Eleştirel Söylem Çözömlemesi (ESÇ) alanı ortaya çıkmıştır.

ESÇ'nin tarihçesini irdelemek istediğimizde yukarıda bahsini ettiğim gelişmeleri göz önünde tutmamız önemlidir. Bu gelişmeler ESÇ'nin ortaya çıkışını anlamak için elzemdir ancak yeterli değildir: ESÇ birçok farklı kanalda gelişen entelektüel birikimin sonucudur. Bu bağlamda, hem Dilbilim'de işlevsel çalışmaların ve Dizgeci İşlevsel Dilbilgisi'nin hem de metindilbilim alanlarının metini ve metinsel olanı anlamak üzerine ortaya koyduklarını değerlendirmeliyiz. Ancak ESÇ'yi ve ortaya koyduğu söylem tanımını anlamak için bu çabalar yeterli olmayacaktır. ESÇ ancak Çatışma Kuramı, bu kuram içindeki Marksçı geleneğin ve Frankfurt Okulu'nun düşünce ve ilkeleri ve Toplumsal Bilimler alanındaki Pozitivizm karşıtı yaklaşımlar değerlendirildiğinde anlaşılabilir.

ESÇ, 1990'ların başlarında Teun A. van Dijk, Norman Fairclough, Gunther Kress, Theo van Leeuwen ve Ruth Wodak'ın Amsterdam'da bir araya gelerek Söylem Çözömlemesi'nin güncel durumunu, yönelimlerini ve ilkelerini tartışmışlardır. Bu toplantının Söylem Çözömlemesi tarihinde özellikle ilkesel duruşun belirlenmesi ve ESÇ'nin söylem tanımının belirginleşmesi açısından önemli bir yeri vardır. Bu toplantıyı van Dijk tarafından editörlüğü yürütölen Discourse & Society Dergisi'nin 1993'teki özel bir sayısında eleştirel çözömlmelere ve farklı eleştirel yaklaşımlara

yer vermesi izlemiştir. ESÇ'nin ilkelerinin ortaya konulduğu temel çalışma ise, Fairclough ve Wodak (1997)'ye aittir. Bu makalede, şu ilkeler ileri sürülmüştür:

- (i) ESÇ toplumsal sorunlara değinir.
- (ii) Güç ilişkileri söylemseldir.
- (iii) Söylem toplum ve kültürü inşa eder.
- (iv) Söylem ideolojik bir işlev görür.
- (v) Söylem tarihseldir.
- (vi) Toplum ile metin arasındaki bağ aracılıdır.
- (vii) Söylem Çözümlemesi yorumlayıcı ve açıklayıcıdır.
- (viii) Söylem bir toplumsal eylem biçimidir.

Bu ilkeleri göz önüne aldığımızda, ESÇ'nin sorun odaklı olduğunu görmekteyiz. Ayrıca ESÇ'nin söylemi toplumsal, tarihsel ve metinsel bir olgu olarak ele aldığını da belirtebiliriz. Özetlemek gerekirse bu ilkelere, söylemin ne olduğuna ve güç gibi diğer toplumsal olgular ile nasıl bir bağlantı kurduğuna yönelik bir bakış açısı görebiliriz. Bunlara ek olarak, söylemin ele alınma biçimine yönelik yöntembilimsel bir çerçeve sunduğunun da altını çizebiliriz. Bu bağlamda, ESÇ'ye göre Söylem Çözümlemesi alanının açıklamacı ve yorumlamacı olduğunu ve toplumsal sorunlara yönelik eleştirel bir duruş sergilediğini vurgulayabiliriz.

Burada bu yazının temel amacının tekrar altını çizmek istiyorum. Özellikle üçüncü söylem tanımıyla örtüşen kuramsal ve analitik duruşu ortaya koymaya çabaladığım bu yazıda aynı zamanda bu tanımın söylemin alışlagelmiş kabul ve ele alış biçimlerinden ayrışma durumunu da irdelemeye çalışıyorum. Bu bağlamda, ESÇ'nin felsefik ve yöntembilimsel arka planını daha derinlemesine değerlendirmek istiyorum. ESÇ adını ve temel duruşunu Frankfurt Okulu olarak da bilinen Alman Neo-Marksist düşünceyle özdeşleşen Eleştirel Kuram'dan almaktadır. Eleştirel Kuram'ın temel ilke ve duruşu toplumu ve toplumsal olan bütün olguları çatışma kavramı çerçevesinde açıklayan Çatışma Kuramı'na dayanmaktadır. Bu kuram toplumsal düzenin varlığını toplumu oluşturan gruplar arasındaki çıkar çatışmalarına göre açıklamaktadır. Bu özelliğiyle toplumsal olandaki güç ilişkilerine odaklanmaktadır. Bu yaklaşım ESÇ'ye şu şekilde nüfus etmiştir: ESÇ söylem çözümlemesini güç ilişkilerinin çözümlemesi olarak algılar. Bu açıdan söylem hem bir çatışma mercei hem de çatışma aracıdır. İşte bu görüş ile ESÇ, güç odaklarının düşünce sistemlerinde örtüklediği bileşenleri açığa çıkarmayı amaçlamaktadır. Bu açıdan, güç odaklarının baskıladığı ve bu gücün nesnesi konumundaki bireylerin, özne olduklarının farkına varmasının sağlanması ile ortaya konulan bir bilinçlendirmenin altını çizmektedir.

ESÇ'nin toplumsal sorun odaklı yaklaşımı ve farkındalık yaratma motivasyonunu göz önüne aldığımızda toplumsal bilimler alanındaki

baskın paradigma olan pozitivistizmin sınırlılıklarına karşıt bir duruş sergilediğini de görürüz. ESÇ, temel yaklaşımını aldığı Eleştirel Kuram'ın da belirlediği gibi doğa bilimlerinin yöntemliliği olan gözlem ve deneyin toplum ve toplumsal olguları çalışırken yetersiz kalacağını düşünür. Bu çerçevede, betimleyici nesnel sosyal bilim anlayışının toplumun ve toplumsalın özneler tarafından oluşturulduğunu göz ardı ettiğini ortaya koyar. ESÇ'ye göre, pozitivistizm ve pozitivist bilim anlayışı, toplumu ve toplumsal olguları nesneleştirir, bununla da kalmaz aynı zamanda araştırmacıyı da nesneleştirir. Ancak toplumsal bilimlerin odağı özneler ve öznelerin oluşturduğu toplumsal gruplardır. Bu bağlamda, öznelerin ve toplumun betimlenmesi yeterli değildir; anlaşılmalı gerekmektedir. Değerden bağımsız nesnel bir toplum çalışması bu noktada yetersiz kalmaktadır. Toplumsal bilimlerde, toplumsal olgular araştırmacının yorumlarına açıktır; yorumcu ve açıklayıcıdır. ESÇ çerçevesinde yürütülen çalışmaların yöntemliliği bu açıdan yorumlamaya dayanmaktadır ve araştırmacının bir özne olarak diğer özne birliktelik ve üretimleri anlaması amacını taşımaktadır.

ESÇ'nin bu paradigma duruşu ve yöntembilimsel yönelimine ek olarak Eleştirel Kuram'ın eleştirim (İng. critique) vurgusunu da göz önüne almak istiyorum. ESÇ'nin yorumcu duruş ve nesnel bilim karşıtı yaklaşımı, anlayış sunma ve bilinç yaratma amacı ile de örtüşmektedir. Böylesi bir amaç ancak eleştirim getirilmesi ile vuku bulabilecektir. Bu doğrultuda, ESÇ'nin anlayışına göre toplumsal bir sorun ancak araştırmacının belli duruşu üzerinden irdelenebilir. Bu da belirli bir bakış açısından ortaya konulan özne bir duruşu gerektirir. Soruna yönelik özne duruş, eleştirim ile farkındalık yaratmayı gerçekleştirir. İşte bu noktada, ESÇ'nin anlama amacı ve eleştirim motivasyonu, ESÇ'yi diğer betimleyici ve yapısal söylem çözümlemesi çalışmalarından ayırır. En özet haliyle, ESÇ araştırmacısı metinleri yorumlar, metinlerdeki toplumsal sorunları ve güç ilişkilerini anlamaya ve bunlara yönelik bir eleştirim sunmaya çabalar.

ESÇ yukarıda sözünü ettiğim özellikleri ile dilbilim içerisinde farklılaşmaktadır. Bu farklılıklara ek olarak ESÇ, disiplinlerarası ve çoklu disiplinli olması nedeniyle toplumsal çalışma alanları ve toplumsal kuramlardan beslenmektedir. Ancak metin ve konuşma üzerine yapılan dilbilim içinden çalışmalar ile toplumsal bilimler içinden yapılan çalışmalar arasında bir kopukluk mevcuttur; ilki sosyoloji ve siyaset bilimindeki güç istismarı ve eşitsizlik üzerine getirilmiş kuramları göz ardı etmekte, ikincisi ise belirgin bir dilsel çözümleme sunmamaktadır (Van Dijk, 2001, s. 363). Bu noktada, söylem çözümlemecilerinin, özellikle eleştirel bakış açısına sahip olanların akademik çevrelerde yaşadığı alışılmalı bir sorunun altını çizmek istiyorum. Her iki durumda da, bu araştırmacılar kendilerine yer bulmak konusunda sorun yaşamaktadırlar. Çalışmalarını sosyoloji gibi toplumsal bilimler içinde gerçekleştirmek istediklerinde, ortaya koydukları çözümlemeler çok dilsel özellik taşımaları gerekçesi ile öteki olarak etiketlenirken, tam tersi durumda dilsel çözümleme yapan dilbilim içindeki araştırmacılar sosyolojik çözümleme

yaptıkları yönünde eleştirilmektedirler. İşte burada, çalışmanın en başındaki çember metaforuna dönmek istiyorum. Gerek paradigma ve yöntembilimsel duruşları gerek ise araştırma motivasyonu ile farklılaşan ESÇ araştırmacıları çemberin dışında oldukları belirtilen gruptur. Çemberin içinde olmak ancak daha nesnel, pozitivist ve betimleyici, belki de daha deneyci bir çalışma alanı olmak ile mümkün görünmektedir. Ancak toplumsal sorunlar ve bu toplumsal sorunlara eleştirel bir duruş söz konusu olduğunda, çemberin içinde olmak yeterli değildir; çemberin dışına çıkmak gerekmektedir. Çemberin dışında olmak ise akademik olarak kabul görmemek ile özdeş düşünülebilir.

Bu sözlerimi Türkiye bağlamından bazı gözlemlerimle desteklemek istiyorum. Bu gözlemlerime dayanarak Türkiye Dilbilimi'nde ESÇ çalışanların konumuna yönelik bir anlayış ve eleştirim sunmak istiyorum. Bu gözlemlerim çerçevesindeki yorumlarım ve iddialarım bir ESÇ araştırmacısı olarak kendi deneyimlerim ile yaşamış olduklarım bağlamında ortaya çıkmıştır. Dilbilim alanı içinde, eleştirel söylem çözümlemesi içeren üretimlerim sosyolojik çözümleme olarak etiketlenirken, sosyoloji ve iletişim bilimleri gibi bünyelerinde ESÇ çalışanları barındıran çevrelerde dil odaklı olarak itham edilmiştir. Böylelikle her koşulda yine çember dışına itilmişliğim olmuştur. Baskın araştırma önyargılarının ve ideolojilerinin bir sonucu olarak ortaya çıktığını düşündüğüm bu durumu bu alanların akademik dergilerinde de gözlemleyebileceğimizi düşünüyorum. Bu noktada, dilbilim doktorasına sahip bir akademisyen olarak Türkiye'de kendi alanımda yayın yapan başlıca akademik mercilerinden biri olan Dilbilim Araştırmaları Dergisi'ni ele almak istiyorum. Dergi temel hedefini Türkiye'deki dilbilim çalışmalarının gelişimine destek vermek ve çağdaş dilbilim ilkeleri çerçevesinde nitelikli dilbilim çalışmalarını yayınlamak olarak belirtmiştir. Bu çerçevede, derginin ilk yayına geçmiş olduğu 1990'dan bu yana dergide 351 çalışma yayınlanmıştır. Bu çalışmalardan ESÇ alanındaki çalışmalara göz attığımızda, dergide toplam altı tane ESÇ çalışmasının bulunduğunu görebiliriz. Bunların iki tanesi 1999 sayısında yer almaktadır. Diğerleri sırası ile 2016, 2019, 2020 ve 2022 sayılarında yer almaktadır. İlk iki çalışma Türkiye'deki başat ESÇ çalışmaları olarak düşünülebilir. Kalan çalışmaların çoğunluğu ise geçtiğimiz son beş yıl içerisinde yayınlanmıştır. Ancak dergide 2022'den bu yana ESÇ çalışması yayınlanmamıştır. Bu akademik bir önyargının doğrudan bir sonucu mudur yoksa bu önyargıların ESÇ araştırmacıları üzerindeki motivasyon kırıcı etkisinden midir bu konuda yorum yapabilmenin çok kolay olmadığını düşünüyorum. Bu ancak çok daha kapsamlı doküman incelemesi ve etnografik gözlemler ile mümkün olacaktır kanısındayım. Bu çerçevede gelecekte, ESÇ'nin Türkiye'deki durumuna yönelik bir meta-çözümleme yapılabilir ve durum daha belirgin şekilde tartışılabilir.

3. Sonuç Gözlemleri

Yazıma son vermeden önce yazıma başladığım sözlerin sahibi Mungan'ın bu sözler ile ilgili söylediklerini hatırlatmak istiyorum. Mungan'a göre "çember, çemberi çizen için komedi, içindeki içince dramdır ama eğer biri kendini çember içine almışsa bu trajedidir." Baskın akademik ideolojinin gerektirdikleri ve akademik önyargılar çerçevesinde alışlagelmişin dışında kalan araştırmacılar bir karar vermek durumundadır. Özellikle akademik çevrede kabul görmek, kendine yer bulabilmek ve yayın yapabilmek için çalışma alanını çember içine kaydırmış olan araştırmacılar için yaşadıkları benim anlayışıma göre akademik bir trajedidir. Kendileri çemberin içindeyken kafaları dışarıda olmaya mahkumdur. Bu bağlamda, çemberin hangi yaklaşımsal, kuramsal ve yöntembilimsel duruşları içine alacağını, çağdaş dilbilim yönelimlerini ve evrensel paradigma dönüşümlerini göz önünde bulundurarak tekrar değerlendirmenin Türkiye Dilbilimi özelinde son derece elzem olduğunu belirterek sözlerime son veriyorum.

Kaynaklar

- Biber, B., Connor, U., & Upton, T. A. (2007). *Discourse on the move*. John Benjamins Publishing Company.
- Fairclough, N., & Wodak, R. (1997). Critical discourse analysis. In T. A. van Dijk (Ed.), *Discourse studies: A multidisciplinary introduction* (pp. 258–284). Sage.
- Kocaman, A. (2003). Dilbilim söylemi. In A. Kocaman (Ed.), *Söylem üzerine* (pp. 1–11). METU Press.
- Schiffrin, D., Tannen, D., & Hamilton, H. E. (Eds.). (2001). *The handbook of discourse analysis*. Blackwell Publishers.
- van Dijk, T. A. (1985). *Handbook of discourse analysis*. Academic Press.
- Wodak, R., & Meyer, M. (2001). *Methods of critical discourse analysis*. Sage.

REMARKS ON TURKISH INTERROGATIVE COMPLEMENT CLAUSES AND VERB SUBCATEGORIZATION

Tai MA

M.A., Graduate Student of Carleton University, Department of Linguistics, taima199506@gmail.com,
ORCID: 0000-0002-7657-9069

Ma, T. (2025). Remarks on Turkish interrogative complement clauses and verb subcategorization. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 273–288). Artsürem. <https://doi.org/10.7816/imuling-15-25-01X014>

ABSTRACT

This chapter examines the interpretations of interrogative complement clauses in Turkish, with a focus on the embedded wh-words. Verbs such as *unut-* (*forget*) and *hatırla-* (*remember*) do not constitute an interrogative environment for the embedded wh-words, and the embedded wh-words are non-interrogative. Interestingly, the two verbs can yield either an interrogative or an indefinite reading. Meanwhile, *san-* (*assume*) and *şüphelen-* (*suspect*), behaving like desiderative and jussive verbs, do not license an embedded wh-word in their complement clauses. In contrast, *düşün-* (*think*), *karar ver-* (*decide*) and *anla-* (*understand*) yield ambiguous readings for embedded wh-words. Moreover, unlike the English counterparts, the embedded wh-words in complement clauses of *sor-* (*ask*) and *merak et-* (*wonder*) can obtain different scopes, but their interpretations remain interrogative in Turkish. The evidence suggests that the interpretations of wh-words depend on the embedding environments, and this suggests reevaluating the verb subcategorization frame based on the observation that wh-words are not consistently interrogative but also indefinite conditionally (Kratzer & Shimoyama, 2002). The problem relates to the syntax-semantics interface. In conclusion, there are three different types of verbs based on their attitudes towards the interpretations of the embedded wh-words.

Keywords: Turkish, Verb subcategorization, Wh-words, Interrogative complement clauses, Syntax-semantics interface

ÖZ

Bu bölüm Türkçede içerik soru tümcelerinin okuma biçimini denetimlemekte ve değişken olan ne-öbeklerine odaklanmaktadır. Türkçede *unut-* ve *anla-* hiç belirsizlik göstermeyip içerik tümcelerini soru biçimine dönüştürmemektedir fakat ikisi de ne-öbeklerine ya soru ya belgisiz okuma kazandırabilir. Aynı anda, *san-* ve *şüphelen-* dilek ve emir belirten eylemler gibi hiç soru biçiminde içerik tümceleri almamaktadır. Öte yandan, *düşün-*, *karar ver-* ve *anla-* içerik tümcelerinde bulunan ne-öbeklerine belirsiz

okuma yaratabilir. Ayrıca, İngilizcedeki karşılıklarından farklı olarak, *sor-* ve *merak et-* de kendi içerik tümcelerindeki ne-öbeklerine belirsiz açı yaratabilir fakat sürekli soru okumadır. Bu çalışmaya göre içerik tümcelerinin ne-öbeklerinin okumaları ana eylemlerin özelliklerine göre değişmektedir. Bu bulgu nedeniyle eylemi sınıflandırma dizgesini gözden geçirmek gerek. Ne-öbekleri her zaman soru okuma kazanmamakta olup belgisiz olarak da davranabilir (Kratzer & Shimoyama, 2002). Dolayısıyla, bu çalışma bu sözdizim-anlambilim arayüzüyle ilgili problem doğrultusunda eylemi sınıflandırma dizgesine karşı yeni bir bakış açısını sunmaktadır. Sonuç olarak, içerik tümcelerindeki ne-öbeklerine yarattıkları okuma biçimlerine göre üç farklı çeşit eylem vardır.

Anahtar Sözcükler: Türkçe, Eylem altbirimleri, Ne-öbekleri, Sözdizim-anlambilim arakesiti

1. Introduction

In § 1.1, I will briefly discuss the particular $\bar{A}$ -movement strategy of wh-in-situ languages (Aoun & Li, 1992). Afterwards, in § 1.2, I will introduce the well-attested subcategorization model that unveils the semantic selectional restrictions of verbs for interrogative complement clauses (containing wh-words). Please note that I do not discuss the wh-adjuncts in this paper since Görgülü (2006: 63-64) claims that the extractions of wh-adjuncts in embedded clauses lead to ungrammaticality (cf. Çakır, 2016: 79). In § 1.3, I will adopt the definitions of *-mA*, *-DIk* and *-AcAk* made by Predolac (2017) and Ma (2025) for Turkish clausal complementation strategy.

§ 2 primarily focuses on the interpretations of the wh-words in Turkish complement clauses, excluding the readings of the interrogative complement clauses as exclamatory sentences or echo questions. In light of the syntactic account of clausal embedding facilitated by *-mA*, *-DIk* and *-AcAk*, I inquire about the variable interpretations of wh-words embedded under different verbs, and I will evaluate the data and reassess the subcategorization frame regarding the variable wh-words in complement clauses as far as § 3.

Concerning that *-mA* clauses yield only the interrogative reading, whereas *-DIk* and *-AcAk* clauses allow either the interrogative or the non-interrogative reading, the semantic approach is paid much attention to in this study as a result of the underlying semantic purposes associated with the wh-words embedded under different verbs. In short, this paper examines how wh-phrases in complement clauses interact with different embedding environments and yield different readings, beginning from the question of why the subcategorization is away from adequacy. For the sake of concreteness, the concern of this study can be paraphrased as to what

extent verbs taking tensed complement clauses differ by yielding different readings for the embedded wh-words.¹

Empirically, if we merely consider whether the interrogative complement clause is grammatical when embedded by a verb without investigating the interpretation strategies of the embedded words in connection with the embedding predicates, as shown in (1), we may temporarily extrapolate that non-factive predicates allow the extraction of embedded wh-words in the LF (see 1a), whereas factive predicates yield ambiguous interpretations (see 1b). However, the assumption is too problematic because it does not cover all examples (see 1c). Therefore, this study heuristically accounts for the various behaviours of wh-words without directly appealing to the verb subcategorization frame.

- (1) a. O Ali'nin ne-yi kazan-dığ-ı-nı san-dı?
He Ali-GEN what-ACC win-DIK-POSS-ACC assume-PERF-3.SG.
'What is x such that Ali assumed that Ahmet has won?'
- b. O Ali'nin ne-yi kazan-dığ-ı-nı bil-yor./?
He Ali-GEN what-ACC win-DIK-POSS-ACC know-PERF-3.SG.
'Ali knows what Ahmet has won.'
'What is x such that Ali knows that Ahmet has won?'
- c. O Ali'nin ne-yi kazan-dığ-ı-nı tahmin ed-iyor./?
He Ali-GEN what-ACC win-DIK-POSS-ACC guess-PROG-3.SG.
'Ali is guessing what Ahmet has won.'
'What is x such that Ali knows that Ahmet has won?'

1.1 The Qu-movement instead of invisible LF movement in wh-in-situ languages

From a syntactic perspective, instead of wh-movement, wh-in-situ languages exhibit the Qu-movement after the co-indexation of embedded wh-words with the Qu operator that is base-generated in the complement C⁰ (Aoun & Li, 1992; cf. Arslan, 1999; Görgülü, 2006; Çakır, 2016). The difference between relative clauses and interrogative complement clauses is that an empty category in the former replaces the overt wh-word. However, I will consistently view both types as A-bar movement in Turkish, even though this study will not involve Turkish (free)-relative clauses.

As a result of the implausibility of the assumption brought up based on (1), intuitively, the questions arise as to whether the subcategorization frame should be revisited based on the Qu-movement, and to what extent if the consideration is non-trivial. To be specific, verbs such as *sor-* (*ask*), *merak et-* (*wonder*) may license the Qu-movement or not, but the interpretations of the embedded wh-words are consistently

¹ Please note that the discussion will only involve the wh-words in tensed complement clauses because subjunctive and other tenseless complement clauses do not allow embedded wh-words (Predolac, 2017; Ma, 2025).

interrogative. In contrast, the interrogative reading is blocked when the wh-words are embedded under verbs such as *bil-* (know), *hayal et-* (imagine) and *unut-* (forget). Verbs such as *san-* (assume), *varsay-* (hypothesize) unconditionally have the Qu operator move to the matrix clause in the LF. I will proceed with these concerns in the following discussion.

1.2 The verb subcategorization for interrogative complement clauses

Technically, embedding verbs are annotated with $\pm wh$ according to the clause types that the embedded clauses can exhibit as a result of their choices of wh-words. Specifically, if a verb can embed a finite tensed clause that contains a wh-word, the verb is annotated with $+wh$, which indicates that it can take an interrogative clause as a complement, and vice versa. Wen (2002: 292) provides a list for subcategorization that descriptively annotates verbs with their licensing criteria for wh-words (see 2) (cf. Huang, 1982a: 371). In (2), *assume* demands the wh-movement, and *know* can select either interrogative or indicative complement clauses, whereas *ask* and *wonder* are compatible only with interrogative complement clauses. This subcategorization follows the basic understanding that verbs can choose the type of their complement clauses, such as indicative, interrogative, or subjunctive, which is understood to be their semantic selectional restriction.

- (2) a. Assume: $[-_{WH} CP]$;
 b. Know: $[\pm_{WH} CP]$;
 c. Ask: $[+_{WH} CP]$.

Based on the semantic selectional restrictions, the clausal complementation brings up a new question of why embedded wh-operators receive the identical scope reading from subjunctive complement clauses and complement clauses of verbs like *assume*. For instance, in Turkish, desiderative verbs like *iste* (want) and *bekle* (expect) take subjunctive complement clauses, and they are incapable of licensing embedded wh-words for the lack of overt genuine tense information (Ma, 2025). In complement clauses of *san-* (assume), wh-elements are interrogative and undergo wh-movement. Particularly, it is interesting that the verb *san-* (assume), which takes propositional complement clauses, does not behave as non-factive verbs like *düşün-* (think) and *hayal et-* (imagine) that also take propositional complement clauses but may assign either the interrogative or the non-interrogative reading to the embedded wh-words (see 3) (cf. Li, 1992: 125-155). Please also note that, unlike English, rogative verbs such as *ask* exhibit the $\pm wh$ feature in Turkish, which means that Turkish rogative verbs allow wh-movement.

- (3) O Ali'nin ne-yi kazan-dığ-ı-nı hayal ediyor./?
 He Ali-GEN what-ACC win-DIK-POSS-ACC imagine-PROG-3.SG.
 'He imagines what Ali has won.'
 'What is such x does he imagine that Ali has won?'

As shown in (3), the embedded wh-complement clauses embedded under the verb *imagine* may obtain a non-interrogative reading, behaving differently from when embedded under the verb *san-* (*assume*). In this connection, it is felicitous to extrapolate that despite the identical propositional structure, embedded wh-words act differently according to the embedding predicates. Non-factive predicates are also distinguished from each other in terms of their behaviours in assigning different readings to the embedded wh-words. Meanwhile, regardless of their syntactic and semantic differences, non-factive predicates, such as *san-* (*assume*), resemble subjunctive governors (e.g., *suggest*) due to the identical behaviour against the capacity for an embedded wh-word.

To summarize, embedded wh-words are ambiguous in complement clauses of factive predicates, such as *unut-* (*forget*), *hatırla-* (*remember*) and *bil-* (*know*), but their readings are unambiguous in rogative verbs (*ask* and *wonder*), albeit with different scopes. Some non-factive verbs, such as *hayal et-* (*imagine*), overlap with factive verbs *unut-* (*forget*) and *hatırla-* (*remember*) in licensing either the interrogative or the indefinite reading to the embedded wh-words. In spite of the contrast of factive and non-factive features, some non-factive verbs may overlap with factive verbs in licensing ambiguous readings to the embedded wh-words. Thus, there must be an additional feature involved in the verb classification for the semantic treatment of the embedded wh-words. However, in this study, instead of reassessing the subcategorization model exclusively for Turkish, I will attempt to attribute the licensing of (non-)interrogative features to the embedded wh-words to the embedding environments.

1.3 The brief information of *-mA* *-DIk* and *-AcAk* heading Turkish complement clauses

In this subsection, I will discuss the embedding strategies in Turkish. This may be additional, but it is necessary for illustrating the syntactic differences between the clausal complementation scenarios in Turkish, which at the same time serves as a brief introduction to the interpretations of the embedded wh-words. Albeit with the nominal morphology, I will regard the *-mA*, *-DIk* and *-AcAk* constructions as complement clauses that are embedded under mood governors (Farkas, 1992), following Kural (1993), Kornfilt (2003, 2007), Predolac (2017) and Ma (2025).

In Turkish, *-mA*, *-DIk* and *-AcAk* are followed by the nominal morphology (e.g., the possessive case and the case marking) but facilitate the sentential complementation. *-DIk* is for independently tensed complement clauses (please note that *-AcAk* is the prospective variant of *-DIk*, and the alternation of *-DIk* with *-AcAk* is grammatically licit) and *-mA* for subjunctive complement clauses (Predolac, 2017; Ma, 2025). Regarding the functions of the three suffixes, according to Demirok (2019), *-DIk* clauses indicate a propositional complex in the sense that they

manifest a tensed finite clause, whereas *-mA* does not represent any genuine tense information.

As for the behaviours of the embedded wh-words in the two different embedding environments (overtly tensed *-DIk* and *-AcAk* and untensed *-mA*), it is possible to adapt the complementary distribution of mood categories to the syntactic features of *-DIk*, *-AcAk* and *-mA* clauses, which are depicted as follows: indicatives are [-WH, +T], interrogatives are [+WH, +T], and subjunctives are [-WH, -T] (Ma, 2025). As shown in (4), *-DIk* and *-AcAk* can manifest the indicative and interrogative moods (see 4a & 4b), whereby they may yield either interrogative or non-interrogative readings for embedded wh-words depending on the embedding environments. By contrast, *-mA* subjunctives are not concerned with an embedded wh-word due to the lack of overt tense morphemes (see 4c).

- (4) a. O Ali'nin ne-yi ye-diğ-i-ni bil-iyor./?
 He Ali-GEN what-ACC eat-DIK-POSS-ACC know-PROG-3.SG.
 'He knows what Ali ate.'
 'What is such x that he knows Ali ate?'
 b. O Ali'nin ne-yi yi-yeceğ-i-ni unut-tu./?
 He Ali-GEN what-ACC eat-PROS-POSS-ACC forget-PERF-3.SG.
 'He forgot what Ali would eat.'
 'What is such x that he forgot Ali would eat?'
 c. O Ali-nin ne-yi ye-me-si-ni iste-di?
 He Ali-GEN what-ACC eat-MA-POSS-ACC want-PERF-3.SG.
 'What does he want Ali to eat?'

In other words, I will not relate the nominal morphology to the embedded wh-words in question. Suppose that the verb *iste* (*want*) will have a subjunctive complement clause manifested by the *-mA* construction. It always yields a wide scope and the interrogative readings for the embedded wh-operator. Meanwhile, the verb *karar ver-* (*decide*) can either take a subjunctive complement clause or a tensed complement clause (interrogative or indicative). However, we cannot conclude that the reading yielded for an embedded wh-operator depends on whether the nominalization is driven by *-mA* or *-DIk*, or *-AcAk*, and then maintain that the wide reading is associated with *-mA* but ambiguous with *-DIk* and *-AcAk*. This is too problematic and empirically misleading, given that the semantic evidence deviates from the preconceived conclusion. The manifestation of complement clauses by *-mA*, *-DIk* and *-AcAk* in Turkish is nonetheless trivial, and this study will only focus on the *-DIk* and *-AcAk* cases for the various interpretations that their embedding predicates assign to the embedded wh-words.

2. Wh-words are variables

In this section, I will discuss the interpretations of *wh*-words in various contexts and explore the interrelationship between the embedding predicates and the readings of the embedded *wh*-words. A large body of research shows that *wh*-words have been treated as quantifiers, indefinites or variables (Chomsky, 1981: 55–60; Wen, 2002: 290; Nishigauchi, 1990; Cheng, 1997).²

Due to the mediation of negation, modality and mood, *wh*-words may be multiply ambiguous and interpreted as the universal quantifier, the negative quantifier or the existential quantifier (cf. Huang, 1982b: 242; Lin, 1996: 230). Huang (1982b: 242) notes that the non-interrogative readings stem from contextual information that provides a less-than-positive meaning. Furthermore, as Lin (1996) notes, *wh*-words are interpreted in line with an indeterminate expression when the truth value of the proposition is associated with uncertainty or negation.

It is also well-attested that *wh*-words are not consistently interrogative in Turkish and their readings vary according to the quantificational forces stemming from the embedding environments (Görgülü, 2006). For instance, in negative environments, *wh*-words typically refer to a quantifier such as *none* and *nobody*, whereby the double negation undergoes an affirmative paraphrasing, and they obtain a reading as the universal quantifier (see 5a). In conditional clauses, *wh*-words denote a referentiality and behave as the existential quantifier (see 5b). By contrast, by means of the aorist aspect, *wh*-words may refer to a negative quantifier (see 5c). Empirically, the non-interrogative readings of the embedded *wh*-words are pragmatically by-products of the contexts, and the non-interrogative reading is dispensable as a result of the run-of-the-mill function of *wh*-words to express questions.

- (5) a. Ali ne-yi bil-mez.
 Ali what-ACC know-NEG.PRES.3.SG.
 ‘What does Ali not know?’
 (Intended: Ali knows everything.)
- b. Ali kim-i tanı-r-sa.
 Ali who-ACC know-AOR-COND.3.SG
 ‘If Ali knows who.’
 (Intended: If Ali knows someone.)
- c. Ali ne-yi bil-ir.
 Ali what-ACC know-AOR-3.SG.
 ‘What does Ali know?’
 (Ali knows nothing.)

(Adapted from Görgülü, 2006)

² In Cheng’s (1997) analysis, *wh*-words are analogized as *polarity-sensitive items*, such as negative polarity items and free choice items (e.g. *any*).

On the other hand, interrogative complement clauses differ from matrix interrogative clauses in that the former denote a set of propositions (Uegaki, 2015: 34). To be specific, *wh*-words in complement clauses are liable to provide a set of potential answers via the retrieval in the context. Interrogative complement clauses are ultimately analogized to indicative/propositional ones, and embedded *wh*-words are not essentially interrogative (cf. Uegaki, 2020: 15). Put differently, interrogative complement clauses technically are not distinguished from tensed indicative clauses in terms of the essential propositional modal (cf. Palmer, 2002: 52-54).

From that position, embedded *wh*-words in a proposition obtain different readings as a result of the adaptation of the complement clauses to the embedding predicates. Suppose that an embedding verb is a rogative predicate; an interrogative reading is licensed to the embedded *wh*-word. If the embedding predicate is altered with one that can block the interrogative reading of the embedded *wh*-word and take *bil-* (*know*) for example, its complement clause, in the circumstance that the non-interrogative reading is expected other than the interrogative one (see 6), must contain such information that the speaker has acquired the knowledge or the intuition that Ali has bought something.³ This evidence suggests that some verbs may provide or relate the utterance to contextual information so that the embedded *wh*-words yield non-interrogative readings. Therefore, I will compare the syntactic and semantic approaches to the embedded *wh*-words in § 2.1.

- (6) Ben Ali'nin ne-yi al-dığ-ı-nı bil-iyor-um.
 I Ali-GEN what-ACC buy-DIK-POSS-ACC know-PRO-1.SG.
 'I know what Ali bought.'

Incidentally, verbs that can take complement clauses are generally known to be cognition verbs (including mental and psych verbs, e.g., *think*, *believe* and *remember*) and desiderative verbs (e.g., *want*), etc. Still, some action verbs can embed finite complement clauses (e.g., *show*). From a semantic perspective, embedding a clause in the argument structure relates to intensional semantics (Farkas, 1992: 72). Factive verbs essentially tend to convey veridicality-based truth judgment, whereas non-factive verbs do not. In other words, as a result of the factive feature, different from non-factive verbs, the embedded clauses of factive verbs are subject to truth evaluation. However, both non-factive and factive verbs in this discussion are limited to the ones that take propositional complement clauses, and particularly, the factivity-based verb classification does not participate in the interpretations of the embedded *wh*-words essentially. In § 2.2, I will depart from the feature modification of verbs but illustrate the interconnection of the embedding verbs and the embedded *wh*-words.

³ Raj Singh, personal contact.

2.1 The mismatch between the syntactic and semantic analyses of wh-words

The section focuses on the syntactic and semantic approaches to the embedded wh-words. Within the Minimalist framework, wh-words in complement clauses may yield either a wide or a narrow scope. The variations of the wh-features associated with different verbs in one language constitute a subcategorization model, as shown in (2). This syntactic approach to the behaviours of variable wh-words accounts for the binary scope (either wide or narrow) resulting from wh-movement from the embedded clauses to the matrix clauses. Syntax, being concerned with the derivation, accommodates the subcategorization model with the wh-movement criterion. For instance, the wh-operator yields a narrow scope in (7a & 7b) (please note that I temporarily ignore the ambiguous interpretations of these embedded wh-words in Turkish).

- (7) a. John asked [_{interrogative} what Mary said].
 ‘John asked what x is such that Mary said.’
 (John did not know what Mary said.)
- b. John knows [_{non-interrogative} what Mary said].
 ‘John knows what x is such that Mary said.’
- c. *John assumes what Mary said.
 *‘There is x such that John assumed that Mary said.’

However, the complement clause in (7a) is essentially interrogative because the reply to the wh-variable is yet unknown (which can be cancelled by saying *Mary said nothing*). On the contrary, in (7b), not only does the wh-operator receive a narrow reading, but also the embedded wh-operator is non-interrogative (the presupposed information cannot be cancelled by saying *Mary said nothing*).⁴ The comparison of (7a) with (7c) indicates that the rogative predicate *ask* overlaps with the non-factive predicate *assume* in licensing the interrogative reading to the embedded wh-word regardless of the scope difference. Overall, (7a), (7b) and (7c) are distinguished from each other according to the extrapolation that embedded wh-words receive indefinite readings when the embedding predicates do not require their embedded interrogative clauses to provide a set of potential answers as genuine questions do. Next, I will illustrate the analysis of the embedded wh-words that hinge on the Minimalist syntactic approach, which also shows the challenging aspects of this problem.

From a syntactic perspective, Aoun and Li (1992: 232-233) maintain that the Qu operator is base-generated in the embedded clauses, whereby it moves to the matrix clauses if the embedding environment blocks the Qu operator. Meanwhile, the wh-words are also associated with two different values, [-wh] and [+wh]. Therefore, a question will be generated with the combination of [+Qu] and [+wh], [-Qu] and [+wh] yield

⁴ Please note that the use of the term *presuppose* does not relate to presuppositional predicates.

exclamatory sentences, [+Qu] and [-wh] yield yes/no questions, and [-Qu] and [-wh] yield declarative sentences.

Empirically, the rogative predicates license the Qu feature down to the C of their complement clauses, by which the embedded wh-word checks the Qu feature, yielding an interrogative complement clause. On the contrary, some verbs overlap in licensing either interrogative or non-interrogative readings to the embedded wh-words, while some block the non-interrogative reading and the existence of the embedded wh-word. It is not a welcome approach to connect the non-interrogative readings to embedded questions despite the presence of wh-words. From this position, this approach still leaves the question open as to why the embedded wh-words behave differently depending on the embedding predicates and how the problem can be addressed uniformly without appealing to a subcategorization frame that appears redundant.

Within the Minimalist architecture, the mapping of the wh-words to the interrogative reading can be understood as an agreement, namely, the complementizer has an interpretable Qu feature (iQ feature) and an uninterpretable wh-feature (uw-features), and the wh-phrase has an interpretable wh-feature (iwh) and an uninterpretable Qu feature (uQ) (Citko, 2014: 19-20). In the Minimalist fashion, uninterpretable features should be removed and evaluated prior to the LF, and interpretable features are handed over to the LF (Richards, 2007: 566). Overall, the interrelation of the Qu with wh does not account for the non-interrogative reading of the wh-words embedded under verbs such as *unut-* (*forget*), *hayal et-* (*imagine*), and *bil-* (*know*).

To summarize, the mismatch originates from the inadequacy of the traditional syntactic approach to the attitudes of the embedding predicates towards embedded wh-words. This syntactic-semantic mismatch can be ameliorated by adding the consideration of whether the embedded wh-words need to denote a set of answers to the verb subcategorization frame. In § 2.2, I will show the evidence that the embedding predicates play a significant role in determining the non-interrogative reading for the embedded wh-words or not.

2.2 Wh-words in different embedding contexts

In this section, I will implement the treatment of wh-words as variables and develop the argument that the interpretations of wh-words are affected by the embedding predicates in terms of whether they demand the embedded wh-words to denote a set of possible answers as genuine interrogative wh-words do. Benefitting from the exhaustive study of the variable feature of wh-words, we now reach a consensus that wh-words do not primarily serve the interrogative mood. The same interrogative complement clauses are possible to acquire different LFs in terms of the interpretations of the embedded wh-words as a result of the embedding predicates.

Incidentally, the verb subcategorization provided by Huang (1982a) is borne out by English and Mandarin evidence with respect to the syntactic approach to wh-movement. However, the mismatch of the syntactic and semantic approaches to the different interpretations of the embedded wh-words remains puzzling. Besides, the same verb subcategorization frame does not apply to Turkish. Technically, the parametric discrepancies in the subcategorization models for different languages are nonetheless acceptable to the extent that the irregularity does not refute the universality. Turkish rogative verbs may be exceptions when compared to English and Mandarin rogative predicates that preclude the movement of the embedded wh-words either in the LF or the PF.

On the other hand, the slight parametrization marginally pushes forward a comparison of the interrogative interpretations with the non-interrogative interpretation of the embedded wh-words. In English and Mandarin, the wh-words embedded under rogative predicates are consistently interrogative but narrow-scoped. In Turkish, the wh-words embedded under rogative predicates can move to the matrix clauses in the LF, and they remain interrogative. Comparatively, the wh-words embedded under verbs such as *bil-* (know), *hayal et-* (imagine) and *unut-* (forget) receive a narrow scope and the non-interrogative reading when the movement is blocked in the LF, but they can be assigned a wide scope and the interrogative reading when they move to the matrix clauses in the LF. Interestingly, verbs such as *san-* (assume) and *zannet-* (suppose) are limited to assigning a wide scope and the interrogative reading to the embedded wh-words. Apart from the reality that Turkish rogative predicates allow wh-movement in the LF, what feature makes different types of predicates differ by whether they license the non-interrogative interpretation to the embedded wh-words or not?

In the preceding sections, it has been noted that it is not on the right track to maintain the distinction based on the factive and non-factive features of verbs to account for this phenomenon. It is also certain that verbs are not arbitrarily associated with wh-features when taking tensed complement clauses that are manifested by *-Dik* and *-AcAk* constructions in Turkish. Moreover, they systematically select wh-complement clauses. Take *söyle-* (tell) for instance; the verb *tell* is a dual mood governor, which means that it is compatible with both tensed complement clauses (see 8a) and subjunctive complement clauses (see 8b) (Cornilescu, 2003: 172; Ma, 2025: 96). Thus, as an inference from the underlying information in the utterances connected to the embedding predicates, there must be a corresponding feature associated with these verbs that facilitate the assignment of the non-interrogative reading to the embedded wh-words.

In (8a), *söyle-* takes a tensed complement clause containing a wh-word, and the wh-word yields ambiguous interpretations. Based on the non-interrogative reading of (8a), the speaker's knowledge that Ali has said he would work with someone serves as an affective context so that the embedded wh-word refers to an existential/indefinite expression. In other

words, the underlying context rules out the interrogative reading of the *wh*-word, and the embedded clause is no longer a set of propositions, and the embedded *wh*-word is, in the same fashion, no longer a set of possible answers.

- (8) a. Ali *pro* kim ile çalış-acağ-ı-nı söyle-di./?
 Ali who with work-ACAK-POSS-ACC tell-PERF-3.SG.
 ‘Ali told (us) whom he would work with.’
 ‘Whom did Ali tell he would work with?’
 b. O Ali’nin kim ile çalış-ma-sı-nı söyle-di.
 He Ali-GEN who with work-MA-POSS-ACC tell-PERF-3.SG.
 ‘Whom did he tell that Ali should work with?’

On the other hand, the interrogative interpretation of (8a) overlaps with the subjunctive counterpart of the complement clause of *söyle-* (*tell*). However, this does not indicate that the tensed complement clause is paraphrased in a subjunctive scenario. In this connection, we can infer that the embedded *wh*-words in both the tensed complement clause and the subjunctive complement clause yield the interrogative interpretation because the same embedding predicate does not provide any affective context for the *wh*-word to obtain the non-interrogative reading. Please note that the analysis does not relate the interpretation to the scope of *wh*-words.

Incidentally, even this paper does not pay attention to conditional clauses embedding *wh*-words, the evidence in (9) nonetheless proves that the embedded *wh*-words are interpreted as interrogative or non-interrogative depending on the embedding environments. Compared to the ungrammatical interrogative translation of (9), the non-interrogative reading stems from the conditional marker *-(y)sa*. The conditional function triggers the underlying information that, as far as the speaker considers, he has decided that he would work with someone.⁵

- (9) O_i *pro*_i kim ile çalış-tığ-ı-na karar ver-di-yse./*?
 He who with work-DIK-POSS-DAT decide-PERF-COND-3.SG.
 ‘If he decided whom he would work with.’

The verb *karar ver-* (*decide*) is a dual-mood governor as well. It allows the embedded *wh*-word in its indicative complement clause and can license either the non-interrogative or the interrogative reading. The conditional clause provides the utterance with an affective context that rules out the interrogative interpretation but allows only the non-interrogative reading as *someone* (Huang, 1982b; Lin, 1998: 219-255). Even though the contextual information may be beyond the speaker’s

⁵ Suppose that the conditional type is associated with the present tense or the optative mood. The embedded *wh*-word still receives the non-interrogative reading since the conditional context (a.k.a. affective context) implies that he will and he is expected to work with someone as far as the speaker considers (this note is owing to Oktay Çınar’s feedback).

knowledge, it must be available for the utterance to be grammatically valid and meaningfully appropriate. Briefly, the non-interrogative reading is not arbitrarily assigned because of the involvement of the underlying contextual information but is dependent on the embedding predicates. In the next section, I will discuss the verbs that may license the non-interrogative interpretation to the embedded wh-words.

3. The embedding environments for the different interpretations of embedded wh-words

In this section, I restrict the environments to the embedding predicates instead of extending them to cover the contextual information or moods. Uegaki (2020) utilizes a mechanism of question-to-proposition reduction, which facilitates the novel treatment of interrogative complement clauses as a proposition structure and uniformly accounts for the non-interrogative readings of embedded wh-words. Specifically, Uegaki (2020) ascribes this phenomenon to the predicates providing the implication/presupposition that the embedded clause is semantically true, whereby the references of embedded wh-words have implicit indices.

3.1 Predicates that may not require the embedded interrogative clauses to denote a set of answers

Take (10) for instance. The interpretation of the embedded interrogative clause varies in terms of whether it marks a question or implies an answer. In the former scenario, it can be the case that nobody ever danced, whereas, under the presupposed information stemming from Ali's motivation for thinking, there might be at least one person who danced. Uegaki (2015) refers to this presupposed reference of wh-questions as an existential presupposition in interrogative complement clauses. The existential semantics is available because the embedding predicates create space for non-veridicality and uncertainty.

- (10) Ali kim-in dans et-tiğ-i-ni düşün-dü?/.
 Ali kim-GEN dance-DIK-POSS-ACC think-PERF-3.SG.
 'Ali thought who danced.'
 'Who did Ali think danced?'

The embedded wh-words may refer to only one true informative answer as a consequence of predicates of veridicality or certainty; for instance, *bil-* (*know*), *belli* (*be certain*) (see 11) (Uegaki, 2015). In the context of (11) for the non-interrogative reading, benefitting from the predicate *belli* (*certain*), the wh-word refers to a very specific entity, and it originates from the underlying background that Ali has exactly won something. The presupposed answer to the embedded wh-word is unique in (11) as a response to the embedding predicate *belli*.

- (11) Ali'nin ne-yi kazan-dığ-ı belli.
 Ali-GEN what-ACC win-DIK-POSS certain-PRES-3.SG.

‘It is certain what Ali has won.’

Nevertheless, I do not distinguish in a detailed manner the existential presupposition from the uniqueness presupposition for embedded *wh*-words. The slight difference between the two non-interrogative readings will not be considered in this study. A general repercussion of this approach is that if we decompose or define interrogative clauses as a set of propositions, certain verbs may rule out the interrogative readings for their embedded interrogative clauses, and other verbs fail to yield the non-interrogative readings so that the *wh*-words must move and obtain the semantic interpretations from higher positions (matrix tenses, moods, etc.).

3.2 Predicates that assign the interrogative reading to the embedded *wh*-words

To begin with, two different types of verbs require the embedded *wh*-words to be interrogative. First, rogative predicates such as *ask*, *wonder* and *inquire* in Turkish can allow *wh*-movement in the LF to have the interrogative feature evaluated in the matrix clauses. Therefore, it is possible to posit that interrogative complement clauses may not incur any violation when moving the embedded *wh*-words to the matrix clauses in the LF in Turkish.

Second, to contrast with the embedding predicates that can license the non-interrogative reading to the embedded *wh*-words, technically, the embedding predicates that allow *wh*-movement do not reduce the questions to propositions because their complement clauses must not contain a *wh*-word (please remember that interrogative clauses are considered to be a set of propositions from a semantic perspective) (cf. Uegaki, 2015: 74). These verbs consistently take declarative complement clauses and preclude the presence of a *wh*-word in their complement clauses, e.g., *inan* (*believe*) and *san-* (*assume*).

Intuitively, the distinction based on factivity does not account for this problem, and the cognitive and mental features lexically do not distinguish verbs like *believe* and *assume* from others that can take interrogative complement clauses. This problem represents an interaction between syntax and semantics, and the semantic feature of these verbs remains an open question for further studies in two respects: why cannot verbs like *believe* and *assume* take interrogative complement clauses, and why do verbs like *imagine*, *know* and *forget* embed non-interrogative *wh*-words?

4. Concluding Remarks

In this study, I compared the interpretation strategies of the embedded *wh*-operators in different embedding environments. By virtue of the finding, the applicability of the traditional verb subcategorization is questioned and is expected to be remedied. Meanwhile, regarding the

ambiguity of interpretations of the embedded wh-words, I heuristically attributed the phenomenon to the embedding predicates and discussed the problem from both syntactic and semantic perspectives.

Incidentally, this study does not intend to remedy the verb subcategorization for Turkish evidence. However, the traditional verb subcategorization designed for the syntax of the question-embedding phenomenon underpredicts the non-interrogative interpretations of the embedded wh-words compared to the interrogative readings.

The Minimalist approach does not cover the possible indefinite interpretations of the embedded wh-words. In this connection, this study tentatively implements a semantic approach to embedded wh-words to consider their interrogative and non-interrogative interpretations. Originally, the embedding environments are taken into consideration based on which readings they assign to the embedded wh-words and whether they allow wh-movement in the LF in Turkish.

As a result, three types of verbs are identified according to their licensing criterion for the embedded wh-words. Rogative verbs in Turkish consistently assign the interrogative reading to the embedded wh-words and allow the wh-movement in the LF. Verbs such as *hayal et-* (*imagine*), *unut-* (*forget*), *düşün-* (*think*) and *bil-* (*know*) take either the non-interrogative or the interrogative counterparts of the wh-complement clauses. By contrast, verbs such as *inan-* (*believe*), *san-* (*assume*) and *zannet-* (*suppose*) overlap with subjunctive mood governors and preclude the existence of the embedded wh-words. In this study, I tentatively argue that the non-interrogative reading of the embedded wh-words results from the indefeasible underlying premises associated with the embedding verbs which pragmatically implies that the intended propositions are true.

References

- Aoun, J., & Li, Y. A. (1993). Wh-elements in situ: Syntax or LF? *Linguistic Inquiry*, 24(2), 199–238.
- Arslan, Z. C. (1999). *Approaches to wh-structures in Turkish* (Doctoral dissertation, Boğaziçi University).
- Citko, B. (2014). *Phase theory: An introduction*. Cambridge University Press.
- Cheng, L. L.-S. (1997). *On the Typology of wh-Questions*. Taylor & Francis.
- Chomsky, N. (1981). *Lectures on Government and Binding*. Dordrecht: Foris Publication.
- Cornilescu, A. (2003). *Complementation in English: A minimalist approach*. Editura Universității din București.
- Çakır, S. (2016). Wh-island constraints in Turkish. *Dilbilim Araştırmaları Dergisi*, 2017(2), 73–91. <https://doi.org/10.18492/dad.295018>
- Demirok, Ö. (2019). A Semantic Characterization of Turkish Nominalizations. In *Proceedings of the 36th West Coast Conference on Formal Linguistics*, ed. Richard Stockwell et al., 132–142. Somerville, MA: Cascadilla Proceedings Project.

- Görgülü, E. (2006). *Variable wh-words in Turkish* (Master's thesis, Boğaziçi University).
- Huang, C.-T. J. (1982a). Move wh in a language without wh-movement. *The Linguistic Review*, 1, 369–416.
- Huang, C.-T. J. (1982b). *Logical relations in Chinese and the theory of grammar* (Doctoral dissertation, Massachusetts Institute of Technology).
- Kornfilt, J. (2003). Subject case in Turkish nominalized clauses. In U. Junghanns & L. Szucsich (Eds.), *Syntactic structures and morphological information* (pp. 129–215). Mouton de Gruyter.
- Kornfilt, J. (2007). Verbal and nominalized finite clauses in Turkish. In I. Nikolaeva (Ed.), *Finiteness: Theoretical and empirical foundations* (pp. 305–332). Oxford University Press.
- Kural, M. (1993). V-to(-I-to)-C in Turkish. In F. Beghelli & M. Kural (Eds.), *UCLA Occasional Papers in Linguistics* (Vol. 11, pp. 1–37). UCLA Department of Linguistics.
- Kratzer, A. & Shimoyama, J. (2002). Indeterminate Pronouns: The View from Japanese, in Yukio Otsu (ed.): *The Proceedings of the Third Tokyo Conference on Psycholinguistics*. Tokyo (Hituzi Syobo), 2002, 1-25.
- Li, Y.-H. A. (1992). Indefinite wh in Mandarin Chinese. *Journal of East Asian Linguistics*, 1, 25–155.
- Lin, J.-W. (1998). On existential polarity wh-phrases in Chinese. *Journal of East Asian Linguistics*, 7, 219–255.
- Ma, T. (2025). *Nominal morphology on Turkish complement clauses* (Master's thesis, Carleton University).
- Nishigauchi, T. (1990). *Quantification in the Theory of Grammar*. Kluwer, Dordrecht.
- Palmer, F. R. (2001). *Mood and Modality*. Cambridge: Cambridge University Press.
- Predolac, E. (2017). *The syntax of sentential complementation in Turkish* (Ph.D. dissertation, Cornell University).
- Richards, M. D. (2007). On feature inheritance: An argument from the Phase Impenetrability Condition. *Linguistic Inquiry*, 38(3), 563–572.
- Uegaki, W. (2015). *Interpreting questions under attitudes* (Doctoral dissertation, Massachusetts Institute of Technology).
- Uegaki, W. (2020). The existential/uniqueness presupposition of wh-complements projects from the answers. *Linguistics and Philosophy*, Advance online publication, 1-41. <https://doi.org/10.1007/s10988-020-09309-4>
- Wen, Binli. (2002). *An introduction to syntax*. Beijing: Foreign Language Teaching and Research Press.

SELF-PACED READING TESTS WITH PSYCHOPY: A VERSATILE TOOL FOR LINGUISTIC DATA COLLECTION

Engin Evrim ÖNEM

Dr., Erciyes University, Department of Basic English, School of Foreign Languages, eonem@erciyes.edu.tr,
ORCID: 0000-0002-2711-7511

Önem, E. E. (2025). Self-paced reading tests with PsychoPy: A versatile tool for linguistic data collection. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 289–304). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X015>

ABSTRACT

This chapter examines the application of self-paced reading (SPR) tests using PsychoPy software in linguistic research, highlighting the method's ability to provide precise measurements of the cognitive processes involved in language comprehension. It reviews existing studies utilizing PsychoPy to investigate phenomena such as syntactic ambiguity resolution and discourse processing. Furthermore, the chapter proposes a methodological framework for conducting SPR experiments, offering specific guidelines for task design, stimulus presentation, and data analysis. By evaluating the advantages and limitations of this approach, this work aims to contribute to the methodological toolkit of linguists and researchers interested in real-time language processing.

Keywords: Self-paced reading, PsychoPy, Linguistics

ÖZ

Bu bölümde, PsychoPy yazılımı kullanılarak uygulanan kendi hızında okuma testlerinin dilbilimsel araştırmalarda kullanımı ve bu yöntemin dil anlama sürecindeki bilişsel süreçlerin ölçülmesine olan katkıları incelenmektedir. Ayrıca bu bölüm, sözdizimsel belirsizliklerin çözümü ve söylem işleme gibi çeşitli olguları araştırmak için PsychoPy kullanan mevcut çalışmaları ortaya koymaktadır. Ek olarak, test tasarımı, uyaran sunumu ve veri analizi konularında rehberlik sunarak, PsychoPy ile kendi hızında okuma deneyleri yürütmek için yöntemsel bir çerçeve önerilmektedir. PsychoPy yazılımının faydalarını ve sınırlılıklarını sunan bu çalışma, dil işleme ile ilgilenen dilbilimcilerin ve araştırmacıların yöntemsel araç setine katkıda bulunmayı amaçlamaktadır.

Anahtar Sözcükler: Kendi hızında okuma, PsychoPy, Dilbilim

1. Introduction

Language is a system reflecting various interconnected cognitive processes in the human mind. Due to the close connection and relationship between language production and comprehension in the cognitive system, “selection and inhibition of verbal or non-verbal input stimuli are comparable to the executive mechanisms involved in the activation of semantic and/or lexical representations” (Turgeon & Macoir, 2008: 10). Although some early as well as some modern psycholinguistic studies focus on language deficiencies, modern methods have been proposed to analyze these mechanisms. As Ullman (2013) summarized, these methods include, but are not limited to, a wide variety of data collection tools, from medical imaging equipment to written tests, questionnaires, and judgment calls. All these methods actually focus on reaction, which is the voluntary or involuntary response to an external stimulus requiring complex processes. The duration spent reacting to any stimulus can be used to measure complex cognitive processes. As Balakrishnan et al. (2014) stated, reaction time is the difference in time between the presentation of a stimulus and the inception of the response to that stimulus. In other words, it is the duration from stimulus onset to voluntary response.

Francis Galton is considered the first to study simple reaction time, as he investigated teenagers’ reaction times to light and sound stimuli in the 19th century (Woods et al., 2015). Since Galton, psycholinguistic studies have developed various methodologies utilizing reaction time, including simple reaction time, recognition, choice, and serial reaction time studies to study language processing (Baayen & Milin, 2010; Woods et al., 2015). Early experimental psychology utilized reaction time (RT) measurements to infer cognitive processes, as it was widely accepted that longer RTs indicated more complex or demanding processing, or vice versa. This established a crucial methodological framework for measuring temporal aspects of cognition. While basic RT tasks often involved simple stimuli and responses, reading studies required more sophisticated methods to track the processing of continuous text. As Jegerski (2013) stated, along with the advances in computers in the 1970s, researchers applied the principles of RT studies to language processing, and the self-paced reading (SPR) technique emerged to measure the time course of language processing. Also, in line with the development of psychometrics in the realm of psycholinguistic research, efficient and flexible tools have become important. In this sense, PsychoPy has stood out as a powerful and open-source software suite designed to construct and execute a wide array of psychological experiments, including those crucial to understanding language processing. Its versatility has cemented its place in the academic toolkit as it offers a platform for creating and running self-paced reading experiments. In essence, PsychoPy provides a comprehensive toolkit for psycholinguists, enabling the precise presentation of linguistic stimuli and the accurate recording of reaction and reading times, empowering researchers to delve deeper into the intricacies of language processing.

This chapter aims to explore the application, advantages, and disadvantages of SPR with PsychoPy, review existing studies, and provide a methodological framework for future research. However, this chapter should not be confused with a guide to designing and running self-paced reading (SPR) experiments on PsychoPy. However, if needed, researchers can find detailed tutorials and documents on creating and running self-paced reading tests with PsychoPy on the dedicated website, <https://www.psychopy.org/>.

2. Self-Paced Reading: Principles and Methodology

2.1 The underlying principles of self-paced reading

The self-paced reading (SPR) technique is an effective way to deconstruct the intricate, moment-by-moment processes of language comprehension through the presentation of texts or segments in a serial manner. SPR is not just a question of what one knows but how one knows it at the moment. As Marsden et al. (2018) discuss, SPR is an online, computer-based experimental method that presents sentences divided into single words or meaningful phrases. Participants, being active controllers of the reading pace, move through the text by pressing a designated key. Notably, the system automatically records the time intervals between every key press and provides a precise measure of the cognitive effort exerted at each section.

The foundational philosophy of SPR, which involves deducing cognitive processing from reading time patterns, is straightforward yet deeply perceptive. As Jegerski (2013) puts it, "the eyes can be a window on cognition" (4), which highlights that individuals' reading activity, and more particularly the amount of time spent processing each word or sentence, has a direct correlation to the cognitive load imposed on their cognition. As participants control the presentation rate of text, SPR paradigms heavily rely on reading time data as a direct index of the cognitive processing load (Marinis, 2010). This data reveals a participant's sensitivity to and awareness of the linguistic phenomenon being researched, as well as their working memory capacity, which is responsible for holding and manipulating information during reading (Marsden et al., 2018). In fact, SPR's strongest point is its ability to detect processing issues through fluctuations in reading time. Measuring reading times for each segment marks trends in relation to language processing. Decreased reading times typically suggest processing ease, which means the linguistic input is aligned with the reader's expectations or requires minimal cognitive effort. Conversely, longer reading times reflect processing difficulty, indicating that the reader is experiencing difficulty with unexpected or complex linguistic information. It is this sensitivity that renders SPR particularly handy for investigating the impact of ambiguities, anomalies, or complicated linguistic dependencies on language processing. For example, a sudden increase in reading time on a specific word may

indicate that the word creates an unexpected syntactic structure or semantic deviation that requires more cognitive resources to process. By closely observing reading time patterns, researchers can gain a fine-grained understanding of the cognitive operations involved in syntactically parsing and interpreting language, which enables them to draw firm conclusions about the nature of online language comprehension.

2.2 Paradigms for self-paced reading

As the designation suggests, the focus of the study is to ascertain the time required to read the linguistic data displayed on the monitor. A number of paradigms of self-paced reading are available, on which researchers can rely to present and analyze linguistic stimuli. These are most often cumulative and non-cumulative. As Luke & Christianson (2013) describe, under the cumulative paradigm, parts of the linguistic stimuli (clause, phrase, or sentence) are displayed as the participant types any key, and each part is shown until the entire stimulus is displayed on the screen. Conversely, under the non-cumulative model, a new segment appears on the screen when the button is pressed, and the prior segment is hidden. Another distinction in how linguistic stimuli are provided is based on the segment's position on the screen (Jegerski, 2013). The segment(s) can be positioned at the middle of the screen (centered) or presented in a sequence format, from left to right (linear). According to Marinis (2010), when segments are centrally presented, participants are unable to estimate the length of the sentence, which provides a clearer view of the processing. Furthermore, according to Jegerski (2013), since linear presentation is similar to reading real sentences, it may be better to present stimuli in this manner. These advantages have led to the development of the moving window method, which presents segments in a linear and noncumulative manner, similar to conventional reading. Such differences in presentation and procedural considerations also underlie the flexibility inherent in self-paced reading paradigms, which enable the investigation of the temporal dynamics of language processing.

With this range of methodological choices, researchers can tailor self-paced reading paradigms by combining presentation methods and other measures to address specific research questions on language processing. For example, in their research, Luke & Christianson (2013) created a paradigm that combined self-paced reading with masked priming to enable the investigation of sentence context effects on early word recognition. They explained in the conclusion of their paper that their paradigm accurately mimicked results from masked-priming and self-paced reading literature. Gong et al. (2023) introduced a new paradigm, a forward-and-backward mode, by combining the self-paced reading technique, which typically requires reading forward, with the eye-tracking technique, which enables participants to move forward and backward as needed, resulting in successful outcomes in psycholinguistics. Ultimately, it can be affirmed that SPR paradigms are easily adaptable to the needs of

researchers, and this is the primary reason why the self-paced reading technique has been adopted in the majority of studies, as Jegerski (2013) notes.

2.3 Advantages and disadvantages of SPR

Since it is one of the online (or real-time) data gathering techniques, there are some distinct advantages of self-paced reading compared to traditional offline psycholinguistic data collection techniques, such as questionnaires or sentence completion tasks. Firstly, online tests are implicit. For instance, live collection tools are resistant to the metalinguistic abilities of participants because such tools record the unconscious and habitual responses of the participants to what they are engaged in (Marinis, 2010; Jegerski, 2013). However, participants would demonstrate controlled and conscious decisions by performing a slow deliberation over an offline or non-real-time activity at hand, utilizing their explicit knowledge of language and metalinguistic ability (Marinis, 2010; Jegerski, 2013). In fact, implicit grammar rules are limited in SPR (Jegerski, 2013), suggesting that SPR is a more accurate reflection of language processing. Secondly, SPR offers a finer measurement of processing. Marinis (2010) claims that because the reading time of each word, phrase, or clause can be independently measured, SPR easily furnishes information on processing difficulty or ease in individual segments. Greater reading time, as mentioned above, points to processing difficulty, and vice versa. Thirdly, SPR is very simple to administer and inexpensive. In particular, no elaborate participant training or ongoing researcher supervision is required, as it is pretty simple in most cases for the participants: read in chunks and press buttons to proceed. There is also no need for experimental equipment, as only a computer is needed. If other psycholinguistic measures are considered, such advantages are apparent.

Conversely, the self-paced reading technique has some drawbacks, particularly when compared to traditional reading. Firstly, the potential for exhibiting pseudo-reading patterns and their effect on reading times is an issue of ecological validity (Jegerski, 2013). As mentioned above, although items are presented linearly, as in regular reading, they are actually offered in a segmented form, which is far from the normal reading experience. Therefore, this can result in unintended processing load and can interfere with the data of reading time. Habit formation is another potential drawback. Participants in an SPR experiment must repeatedly press a button, and this physical response can lead to automatic reactions over time, which is unnatural in the context of regular reading. Jegerski (2013) also emphasizes the importance of participant fluency in reading for SPR paradigms. Unless readers are native or native-like speakers in the language being tested, reading times may be misleading as to the amount of time spent on reading.

Yet, all these drawbacks are remediable. For example, Jegerski (2013) reported that barely any of the subjects in one of her experiments

had trouble reading in chunks, and button pressing was already a habit for most young subjects. Additionally, breaking after some time and/or limiting the time for SPR experiments can solve fatigue and habit formation problems. Finally, although reading fluency is required for SPR experiments, differences in the native language and second/foreign language processing have been studied (for an extensive review of studies, see Marsden et al., 2018). In other words, non-fluency does exist in linguistic research carried out through the self-paced reading method. In conclusion, although SPR has its limitations, it remains a valuable and widely applied methodology in psycholinguistic research, as it offers high temporal resolution and provides fine-grained information about the cognitive processes underlying language comprehension.

3. PsychoPy: A Powerful Tool for SPR Implementation

PsychoPy is open-source software that allows researchers to simplify and extend it, as it is programmed in Python, a free programming language, and freely distributed software. The free computer language simplifies personalizing the software to integrate with other software and websites, supporting online data collection. The software originated as a psychophysics tool but has since expanded significantly to support complex language-based paradigms. Early releases, as seen by Peirce (2007), laid the groundwork for what it is now, providing scientists with a flexible platform that enables fine control of experimental stimuli and data acquisition. At the same time, the software offers an open space for collaboration, where scientists continuously contribute to its development, maintaining it at the forefront of experimental psychology software. That is, PsychoPy is an ever-changing software that continues to expand, providing researchers more time to ponder their experimental design rather than spending time struggling with complex programming.

Some intrinsic features of PsychoPy make it suitable for self-reading experiments. Firstly, PsychoPy offers precise timing records and flexible control of display. The time taken between the appearance of each linguistic stimulus on the screen and the target key being pressed is recorded in milliseconds. PsychoPy's accuracy in measuring reaction times is also crucial for language experiments, such as lexical decision, ambiguity resolution, or sentence judgment tasks. The software captures participant responses, whether button presses or vocal, along with the duration from stimulus onset to response. Additionally, the segment provided can be easily manipulated in PsychoPy. In the context of self-paced reading (SPR) experiments, PsychoPy is especially well-suited to present sentence fragments with precision. It can be configured to present text word by word or phrase by phrase, allowing participants to read at their own pace. This allows for reading times for each fragment to be investigated in minute detail, providing invaluable information on cognitive processes in language understanding.

Second, PsychoPy's use extends beyond its accurate control and timing, which are critical because it is a powerful tool for presenting stimuli in psycholinguistic research. According to Peirce (2007), PsychoPy enables the creation and manipulation of various types of stimuli, including text, images, and sound files. This is particularly crucial in language research, where researchers often need to present rich textual stimuli, manipulate linguistic factors, and control stimulus presentation timing to the millisecond. Peirce (2009) describes the software's capacity to generate diverse stimuli, a critical feature that allows researchers to create dynamic and interactive testing environments.

Another critical feature of PsychoPy is its convenience. The relaxed or adaptable nature of PsychoPy makes it easier for researchers to design experiments. For example, while experienced computer users or "people who know about programming" would be most advantaged by PsychoPy, there is another mode called the Builder view (Peirce et al., 2019), where a researcher can choose the type of stimulus, the input type for the response that is expected and alter preferences such as the time for any response or what kind of data to save or not save just like they would on Windows® or Mac® OS. Therefore, the simplicity of using PsychoPy to design complex experiments is apparent (Peirce et al., 2019). In fact, this feature is also discussed in detail in Peirce et al. (2022), which provides a step-by-step tutorial on designing experiments with the PsychoPy system.

Another fundamental feature of PsychoPy is the ability to export and log data. Since both the presentation of the stimulus and the recording of reaction time are critical in self-paced reading, PsychoPy takes accurate logs of data and renders them exportable across a range of formats (e.g., PDF, XLSX, CSV, etc.) so that researchers can analyze them freely with their preferred programs. As Peirce et al. (2019) continue to note, the power of PsychoPy lies in its ability to bridge the gap between high-level experimental design and implementable code that can be tailored to meet the specific needs of the researchers.

Finally, note should be taken of the observation that PsychoPy is not a single software to be used in psycholinguistic research. There are a relatively large number of computer programs available for use in psycholinguistic research, including E-Prime®, NBS Presentation®, Psychophysics Toolbox, OpenSesame, Expyriment, Gorilla, jsPsych, Lab.js, and Testable. However, when comparing the precision of visual and auditory stimulus timing and response times provided by various software, Bridges et al. (2020) asserted that PsychoPy performed much more reliably in both laboratory and online studies than other toolkits. However, they also highlighted researchers' adoption of specific timing validation procedures for their experimental design. In addition, Gallant & Libben (2019) provided the primary benefits of PsychoPy, including greater research efficacy, higher ecological validity, and the ability to recruit more representative and historically underrepresented participant samples. They considered the central aspect of millisecond timing in this regard and

offered helpful technical details and recommendations. Therefore, although numerous software packages are available for psycholinguistic experiments, the core functionality and timing accuracy demonstrated by PsychoPy make it a highly trustworthy option, albeit one that still needs to be tested on an individual basis for specific experimental paradigms.

4. Applications of SPR with PsychoPy in Linguistic Research

Many studies employ SPR with PsychoPy to investigate various phenomena in linguistics. Some studies examined syntactic ambiguity resolution and semantic interpretation. For example, Hassim & Kukona (2024) examined the effect of extended cognitive control on processing syntactically ambiguous "garden path" sentences using two experiments on PsychoPy and concluded that a high proportion of incongruent Stroop trials enhanced extended cognitive control, which in turn had a causal influence on improving sentence comprehension, whether the sentence was written or heard. In a further study, Kyriacou et al. (2023) emphasized the resolution of passivization ambiguity in passivized idioms using an eye-tracking task. In addition to the eye-tracing task, they also conducted a lexical decision task using PsychoPy with keywords and the same number of pseudowords, all of which were the same length as the keywords, to ensure that any differences found in control keyword processing were attributed to their contextual inconsistency rather than to other lexical features. In the end, they discovered that keywords were responded to significantly faster than pseudowords, and response times for the keywords did not have a substantial impact on the condition (whether they were in a figurative, literal, or control situation). Jansen et al. (2017) used PsychoPy and a self-paced reading test to study semantic priming on homonymous nouns and found that reaction time and error rate were identical and unaffected by semantic association, whether absent or present, in both ambiguous and unambiguous conditions. They argued that the non-selective access hypothesis was prejudiced in favor of both senses of the processed word being activated.

Some other studies utilized PsychoPy and SPR to examine semantic integration and anomaly detection. For instance, Berger et al. (2022) used PsychoPy to investigate how merely seeing a task cue, compared to performing the cued task, affects unconscious semantic priming based on task-set dominance. The findings revealed that the mere presence of a cue resulted in inhibition of the reversed priming task sets compared to performance on the cued task, indicating that the particular timing of activation and inhibition of the task set following cue presentation could be controlled by certain experimental conditions. In a further study, Mifka-Profozic et al. (2020) examined the influence of syntactic violation and semantic ambiguity on the processing of the modal auxiliaries "can" and "may" in agent-oriented, epistemic, and speaker-oriented contexts. Reading penalties following incongruent modal use were observed in agent-oriented and epistemic contexts, but speaker-oriented

modality processing remained uninfluenced by the inconsistency of formality. PsychoPy is also used in discourse processing research. For example, Wetzel et al. (2022) investigated the processing of native speakers of French causal and concessive sentences with appropriate or inadequate discourse connectives, as well as how connective frequency and polyfunctionality affected processing, based on three experiments. The results showed that processing was affected by incoherent elements irrespective of the connective and that reading fluency decreased with less frequent connectives, especially polyfunctional ones. Studies have shown that frequent connectives are convenient for reading; however, processing discourse is a robust process that allows readers to extract meaning despite infrequent connectives. Crible et al. (2021) studied whether native French speakers and French learners processed contrastive relations indicated by syntactic parallelism with and without discourse connectives in the same way in three experiments using PsychoPy. They examined the impact of parallelism with or without a connective, as well as with often versus rarely and unclear versus clear connectives. The findings stressed that while parallelism was essential to both groups, it was a more potent cue for non-native speakers and was influenced by task difficulty in native speakers.

It can be said that studies identify PsychoPy's capacity to facilitate rich explorations of language processing online, providing valuable insights into the cognitive processes underlying various linguistic phenomena in both native and non-native speakers, across different experimental settings, and even in second language studies. As stated by Hamrick (2022), since second language studies require extremely sensitive, reliable chronometric measures, PsychoPy and SPR are also relevant to second-language studies. Among many studies, Berghoff (2020) explored the processing of object-subject ambiguities by early second-language acquirers in a self-paced reading test administered on PsychoPy. The study reported that early childhood second-language acquirers in the L2 group relied more heavily on online reanalysis to process object-subject ambiguities. Banerjee et al. (2025) used PsychoPy and SPR to examine the effect of word priming on the response time of Bengali-speaking bilingual dyslexic children in identifying the target word in another study, reporting a positive impact of word priming on processing. Patterson & Nicklin (2023) analyzed the reliability and quality of online reaction time data in L2 studies by comparing the self-paced reading responses of crowdsourced workers with those of online students, using an in-person measure as a baseline. The outcome was that strong L1 effects were reproduced across both online populations; however, the data quality of online students was lower compared to that of crowdsourced participants. In conclusion, the various applications outlined above, along with many more, demonstrate PsychoPy's central role as a versatile and reliable tool for collecting self-paced reading data across a broad spectrum of linguistic research, ranging from sentence comprehension and lexical

ambiguity resolution to semantic processing, discourse analysis, and even second language acquisition.

5. Methodological Framework for SPR Experiments with PsychoPy

This section outlines the crucial steps in designing an SPR experiment with PsychoPy, including selecting an appropriate paradigm, presenting the stimuli, and discussing data analysis, interpretation, and potential methodological challenges.

First and foremost, selecting an appropriate SPR paradigm is crucial for eliciting the target cognitive processes and obtaining interpretable reaction times. The variations, as noted earlier, involve the moving window or cumulative paradigm. A researcher needs to select an appropriate paradigm for the linguistic structures with which they are working. For example, Ferreira & Clifton (1986) note that the moving window paradigm allows for high temporal resolution in experiments testing the fine-grained temporal properties of processing, such as the short-term influence of syntactic or semantic variation. Another critical feature of SPR design is controlling the amount and complexity of the text segments. Because reading time within every segment is being measured, sentences must be of equal length and makeup to avoid introducing inequalities in terms of complexity and length (Marinis, 2010). Target sentences and control sentences must be constructed in a way that they are syntactically similar, with similar word or character lengths per segment, so that variation in processing load due to variation in grammaticality is reduced. For example, longer sections of course take longer to read, and greater lexical frequency or syntactic complexity also influence reading times, independent of experimental manipulation (Rayner, 1998). Therefore, target and control items must also be equated on word frequency measures and psycholinguistic properties such as the number of syllables. Third, filler item development is also necessary for any SPR experiment with PsychoPy, as it prevents participants from becoming overly sensitive to the experimental manipulation, reduces the predictability of the target stimuli, and allows for an equal task demand throughout the experiment. According to Goodall (2021), fillers and the experimental items are also required to be counterbalanced on a ratio of 1:1. The filler items, being efficacious, need to be cleverly designed and semantically equatable, need to have an array of semantic content and syntactic structures to avoid allowing participants to develop target-specific processing procedures and need to be of identical length and difficulty as the target items.

Second, planning is required with care about stimulus presentation in SPR using PsychoPy. As noted above, the fundamental method of SPR is to present text section by section, typically after a participant presses a key. Researchers should ensure that the presentation remains uniform across trials and participants, which involves maintaining a consistent location for the text on the screen and providing distinct instructions for

participants to follow when executing the task (Jegerski, 2013). Another aspect to mention is the display of comprehension question(s) suited to linguistic stimuli. As Marinis (2010) argues, comprehension questions should be placed at the end of a segment or sentence so that participants do not form habitual or automatic answers by maintaining attention and readiness. Simultaneously, trials on the SPR should also be of only 15-20 minutes, such that fatigue cannot accumulate in the participants, with at least a week maintained between trials to prevent long-term memory recalls (Marinis, 2010; Jegerski, 2013). Another element that requires consideration is adjusting the text display. According to research, font characteristics such as font size, typeface, and spacing can have a significant impact on reading speed and comprehension. In an experiment, Banerjee et al. (2011) investigated the effect of font size and style on reading performance on computer screens and recommended using font sizes of 14-point for on-screen reading. They suggested Courier New for optimal speed and Verdana for visual quality and reduced mental load. Therefore, for SPR tasks, sans-serif fonts such as Arial, Verdana, Courier New, or Calibri, with a suitable font size (e.g., 12-14 points), are typically recommended for optimal readability on computer screens. A further consideration here is minimizing visual distraction for the SPR task. This includes having a plain background color (e.g., white or gray), excluding unnecessary items on the screen, and ensuring the test environment is free from external distractions (Jegerski, 2013). Finally, researchers should ensure adequate timing and synchronization by running a pilot test and factoring in the monitor's frame rate.

Data analysis and interpretation in SPR with PsychoPy involves several crucial steps: initial data cleaning, preprocessing, applying appropriate statistical methods, and interpreting reading time patterns in the context of linguistic hypotheses. Raw data exported from PsychoPy often require preprocessing and cleaning before they are ready for statistical analysis. This pre-processing involves the identification and exclusion of trials with technical errors (e.g., anticipatory key presses) or participant errors (e.g., incorrect answers to comprehension questions), as defined by researchers before data collection (Ferreira & Clifton, 1986). Outliers in the data, which could result from fatigue, transient lapses, or technical errors, must be excluded. Standard methods include visual inspection of the data, application of standard deviation-based cut-offs (e.g., truncating data points more than 2.5 or 3 standard deviations from a subject's mean under a condition), or using more robust methods, such as the median absolute deviation (Baayen & Milin, 2010). The final step in preprocessing data is the optional aggregation process, which depends on the question being asked. For instance, reading times to specific regions of interest within clauses or sentences might be summed up (e.g., adding reading times across several words within a significant area) to provide a more accurate overall estimate of processing effort for the area. SPR experiment reading time data tends to be investigated using statistical

methods that can account for the repeated measures nature of the data (multiple data points per participant and item). Linear mixed-effects models (LMMs), as Baayen et al. (2008) and Jegerski (2013) explain, have become increasingly popular and are well-suited to this type of data. LMMs enable the estimation of the effect of experimental manipulations (fixed effects) while accounting for variation between items and participants (random effects). Jaeger (2008) explains that fixed effects are experimental manipulations under investigation (e.g., syntactic ambiguity, semantic anomaly, or presence/absence of a discourse connective) and LMMs estimate the average effect of such manipulations on reading times. Random effects, in contrast, reflect non-independence at both the participant and item levels through random participant and item intercepts and slopes, allowing the model to adjust for individual differences in baseline reading speed and for sensitivity to differential amounts of experimental treatment across items (Jaeger, 2008). Parallel with aggregation, Jegerski (2013) states that if LMMs are utilized, computing aggregate means and conducting individual item analyses will no longer be required, and data trimming may be reduced or eliminated. On the other hand, when standard ANOVAs and t-tests are employed, the data need to be reduced and trimmed into aggregate means to account for subject and item variance, as necessitated by the need for item analysis in psycholinguistics.

Knowing the pattern of reading times is crucial for concluding cognitive processes and the linguistic phenomenon under study. First, the variation in reading times for particular segments can be attributed to the ease or difficulty of processing. As mentioned previously, longer reading times or a slowing down of reading speed for a specific segment indicate increased difficulty in processing, and vice versa. Additionally, processing difficulty at a point within the sentence may not be immediately apparent. Still, it can instead be detected in the reading times of subsequent words or phrases, known as spillover. This could indicate that the cognitive processes elicited by a manipulation continue to influence processing in following areas (Rayner, 1998). Furthermore, the effect of a linguistic manipulation could be confounded with other factors, such as between-participant differences in working memory capacity (Just & Carpenter, 1992) or usage frequency of a particular linguistic construction (Chater & Manning, 2006), and can be controlled using LMMs. Significantly, failing to find noticeable differences in reading times across conditions can be too revealing, as these findings could be informative, suggesting that the manipulated linguistic factor has no significant impact on online processing, as measured by reading times.

6. Addressing Potential Methodological Challenges

There are several key considerations to take into account when designing stable self-paced reading (SPR) experiments in PsychoPy, including addressing individual variation, controlling for practice effects

and fatigue, and ensuring adherence to ethical practices in participant handling. First, SPR data are marked by extreme individual variation in reading speed, which is highly expected. Nevertheless, as mentioned above, linear mixed-effects models can also be utilized to handle this variability by incorporating random effects for participants (Baayen et al., 2008), allowing researchers to model experimental condition manipulations while holding individual differences in baseline reading times constant. Another potential SPR difficulty is the impact of practice effects and fatigue. Following Jegerski's (2013) suggestion, practice trials should be included at the start of the experiment to allow participants to adapt to the task, and these steps from practice trials are typically excluded from analysis. In addition, as time elapses from the experiment's commencement, participants may start to tire or lose concentration, resulting in progressively slower and less consistent reading times. To minimize such effects, short breaks between blocks of trials (Marinis, 2010), maintaining the experiment duration between 15 and 20 minutes (Marinis, 2010; Jegerski, 2013), and using a sufficient number of filler items (Goodall, 2021) can be employed. Finally, as ethical considerations are essential in any study involving human subjects, researchers must take note of key ethical requirements, including participants' informed consent, confidentiality, anonymity, and clearance from an ethics committee.

7. Conclusion

The self-paced reading (SPR) technique is an advantageous method for examining the time course of language processing, providing information about cognitive processes that are not easily achieved with other methods. PsychoPy, with its precise timing, flexible stimulus presentation, and easy-to-use interface, provides a solid foundation for conducting SPR experiments. This chapter has discussed the SPR method, as implemented in the PsychoPy program, with an emphasis on its utility as a functional tool for linguistic studies. This chapter outlines the basic principles of SPR, explaining its various paradigms and emphasizing the importance of factors such as stimulus control, data analysis, and the investigation of potential artifacts. A survey of experiments conducted with SPR in PsychoPy to demonstrate the wide range of the technique's potential contributions to different areas of linguistic research, such as syntactic ambiguity resolution, semantic interpretation, discourse processing, and second language acquisition. Finally, the chapter's methodological framework offers researchers a blueprint for designing successful SPR experiments in PsychoPy. By paying close attention to paradigm choice, stimulus design, data analysis techniques, and control of potential confounds, researchers can leverage the strengths of SPR and PsychoPy to inform the complexities of language processing. In short, SPR, as facilitated by PsychoPy, is a valuable and growing tool for linguists seeking to unravel the mystery of how people process language in real-time.

References

- Baayen, R. H., Davidson, D. J., & Bates, D. M. (2008). Mixed-effects modeling with crossed random effects for subjects and items. *Journal of Memory and Language*, 59(4), 390–412. <https://doi.org/10.1016/j.jml.2007.12.005>
- Baayen, R. H., & Milin, P. (2010). Analyzing reaction times. *International Journal of Psychological Research*, 3(2), 12–28. <https://doi.org/10.21500/20112084.807>
- Balakrishnan, G., Uppinakudru, G., Singh, G. G., Bangera, S., Raghavendra, A. D., & Thangavel, D. (2014). A comparative study on visual choice reaction time for different colors in females. *Neurology Research International*, 2014, Article 301473. <https://doi.org/10.1155/2014/301473>
- Banerjee, J., Majumdar, D., Pal, M. S., & Majumdar, D. (2011). Readability, subjective preference and mental workload studies on young Indian adults for selection of optimum font type and size during onscreen reading. *The American Journal of the Medical Sciences*, 4(2), 131–143.
- Banerjee, J., Nithin, K. S., Sai, U. H. K., Chakraborty, B., & Basu, A. (2025). Exploration of semantic priming studies in native language of children: Study on Bengali L1 and English L2 bilingual children. *SN Computer Science*, 6, Article 119. <https://doi.org/10.1007/s42979-025-03664-4>
- Berger, A., Kunde, W., & Kiefer, M. (2022). Task cue influences on lexical decision performance and masked semantic priming effects: The role of cue–task compatibility. *Attention, Perception, & Psychophysics*, 84, 2684–2701. <https://doi.org/10.3758/s13414-022-02568-2>
- Berghoff, R. (2020). The processing of object–subject ambiguities in early second-language acquirers. *Applied Psycholinguistics*, 41(4), 963–992. <https://doi.org/10.1017/S0142716420000314>
- Bridges, D., Pitiot, A., MacAskill, M. R., & Peirce, J. W. (2020). The timing mega-study: Comparing a range of experiment generators, both lab-based and online. *PeerJ*, 8, e9414. <https://doi.org/10.7717/peerj.9414>
- Chater, N., & Manning, C. D. (2006). Probabilistic models of language processing and acquisition. *Trends in Cognitive Sciences*, 10(7), 335–344. <https://doi.org/10.1016/j.tics.2006.05.006>
- Crible, L., Wetzel, M., & Zufferey, S. (2021). Lexical and structural cues to discourse processing in first and second language. *Frontiers in Psychology*, 12, Article 685491. <https://doi.org/10.3389/fpsyg.2021.685491>
- Ferreira, F., & Clifton Jr., C. (1986). The independence of syntactic processing. *Journal of Memory and Language*, 25(3), 348–368. [https://doi.org/10.1016/0749-596X\(86\)90006-9](https://doi.org/10.1016/0749-596X(86)90006-9)
- Gallant, J., & Libben, G. (2019). No lab, no problem: Designing lexical comprehension and production experiments using PsychoPy3. *The Mental Lexicon*, 14(1), 152–168. <https://doi.org/10.1075/ml.00002.gal>

- Gong, T., Gao, X., & Jiang, T. (2023). A “dummy’s” program for self-paced forward and backward reading. *Behavior Research Methods*, 55, 4419–4436. <https://doi.org/10.3758/s13428-022-02025-w>
- Goodall, G. (2021). Sentence acceptability experiments: What, how, and why. In G. Goodall (Ed.), *The Cambridge handbook of experimental syntax* (pp. 7–38). Cambridge University Press.
- Hamrick, P. (2022). Conducting reaction time research in second language psycholinguistics. In A. Godfroid & H. Hopp (Eds.), *The Routledge handbook of second language acquisition and psycholinguistics*. <https://doi.org/10.4324/9781003018872>
- Hassim, N., & Kukona, A. (2024). Linking cognitive control to language comprehension: Proportion congruency effects in syntactic ambiguity resolution. *Language, Cognition and Neuroscience*, 39(4), 431–447. <https://doi.org/10.1080/23273798.2024.2314027>
- Jaeger, T. F. (2008). Categorical data analysis: Away from ANOVAs (transformation or not) and towards logit mixed models. *Journal of Memory and Language*, 59(4), 434–446. <https://doi.org/10.1016/j.jml.2007.11.007>
- Jansen, M. T., Jansen, N. C., Weber, A., Gadea, G. H., Ansari, E., & Scheren, P. (2017). Semantic priming with homonymous nouns: Hints of clarifying the issue of selective vs. non-selective priming. *Journal of European Psychology Students*, 8(1), 15–29. <https://doi.org/10.5334/jeps.408>
- Jegerski, J. (2013). Self-paced reading. In J. Jegerski & B. VanPatten (Eds.), *Research methods in second language psycholinguistics* (pp. 20–49). Routledge.
- Just, M. A., & Carpenter, P. A. (1992). A capacity theory of comprehension: Individual differences in working memory. *Psychological Review*, 99(1), 122–149. <https://doi.org/10.1037/0033-295X.99.1.122>
- Kyriacou, M., Conklin, K., & Thompson, D. (2023). Ambiguity resolution in passivized idioms: Is there a shift in the most likely interpretation? *Canadian Journal of Experimental Psychology / Revue canadienne de psychologie expérimentale*, 77(3), 212–226. <https://doi.org/10.1037/cep0000300>
- Luke, S. G., & Christianson, K. (2013). SPaM: A combined self-paced reading and masked-priming paradigm. *Behavior Research Methods*, 45, 143–150. <https://doi.org/10.3758/s13428-012-0239-4>
- Marinis, T. (2010). Using on-line processing methods in language acquisition research. In E. Blom & S. Unsworth (Eds.), *Experimental methods in language acquisition research* (pp. 139–162). John Benjamins. <https://doi.org/10.1075/llt.27.09mar>
- Marsden, E., Thompson, S., & Plonsky, L. (2018). A methodological synthesis of self-paced reading in second language research. *Applied Psycholinguistics*, 39(5), 861–904. <https://doi.org/10.1017/S0142716418000036>

- Mifka-Profozic, N., O'Reilly, D., & Guo, J. (2020). Sensitivity to syntactic violation and semantic ambiguity in English modal verbs: A self-paced reading study. *Applied Psycholinguistics*, 41(5), 1017–1043. <https://doi.org/10.1017/S0142716420000338>
- Patterson, A. S., & Nicklin, C. (2023). L2 self-paced reading data collection across three contexts: In-person, online, and crowdsourcing. *Research Methods in Applied Linguistics*, 2(1), 100045. <https://doi.org/10.1016/j.rmal.2023.100045>
- Peirce, J., Gray, J. R., Simpson, S., MacAskill, M., Höchenberger, R., Sogo, H., Kastman, E., & Lindeløv, J. K. (2019). PsychoPy2: Experiments in behavior made easy. *Behavior Research Methods*, 51(1), 195–203. <https://doi.org/10.3758/s13428-018-01193-y>
- Peirce, J. W. (2007). PsychoPy—Psychophysics software in Python. *Journal of Neuroscience Methods*, 162(1–2), 8–13. <https://doi.org/10.1016/j.jneumeth.2006.11.017>
- Peirce, J. W. (2009). Generating stimuli for neuroscience using PsychoPy. *Frontiers in Neuroinformatics*, 2, Article 10. <https://doi.org/10.3389/neuro.11.010.2008>
- Peirce, J. W., Hirst, R., & MacAskill, M. (2022). *Building experiments in PsychoPy* (2nd ed.). Sage.
- Rayner, K. (1998). Eye movements in reading and information processing: 20 years of research. *Psychological Bulletin*, 124(3), 372–422. <https://doi.org/10.1037/0033-2909.124.3.372>
- Turgeon, Y., & Macoir, J. (2008). Classical and contemporary assessment of aphasia and acquired disorders of language. In B. Stemmer & H. A. Whitaker (Eds.), *Handbook of the neuroscience of language* (pp. 3–11). Elsevier.
- Ullman, M. T. (2013). Language and the brain. In R. Fasold & J. Connor-Linton (Eds.), *An introduction to language and linguistics* (6th ed., pp. 235–274). Cambridge University Press.
- Wetzel, M., Zufferey, S., & Gygax, P. (2022). How robust is discourse processing for native readers? The role of connectives and the coherence relations they convey. *Frontiers in Psychology*, 13, Article 822151. <https://doi.org/10.3389/fpsyg.2022.822151>
- Woods, D. L., Wyma, J. M., Yund, E. W., Herron, T. J., & Reed, B. (2015). Factors influencing the latency of simple reaction time. *Frontiers in Human Neuroscience*, 9, Article 131. <https://doi.org/10.3389/fnhum.2015.00131>

YENİSEY TÜRKÇESİ RUNİK YAZITLARININ BELGELENMESİ: 2024 YILI EPIGRAFIK SAHA ÇALIŞMASI¹

Alexander V. SAVELYEV

Dr., Kıdemli Araştırmacı, Rusya Bilimler Akademisi Dilbilim Enstitüsü ve Ulusal Araştırma Üniversitesi “Yüksek Ekonomi Okulu” (Moskova, Rusya), a.savelyev@iling-ran.ru, ORCID: 0000-0002-8343-2057

Yuriy M. SVOYSKIY

Laboratuvar Müdürü, RSSDA Laboratuvarı / Araştırma Mühendisi, Ulusal Araştırma Üniversitesi “Yüksek Ekonomi Okulu” (Moskova, Rusya), rutil28@gmail.com, ORCID: 0000-0001-6256-4299

Anastasiya A. CHERNUKHINA

Laboratuvar Personeli, RSSDA Laboratuvarı / Araştırma Asistanı, Rusya Bilimler Akademisi Arkeoloji Enstitüsü (Moskova, Rusya), nastya20022105@gmail.com, ORCID: 0009-0001-1645-9487

Savelyev, A. V., Svoyskiy, Y. M., & Chernukhina, A. A. (2025). Yenisey Türkçesi runik yazıtlarının belgelenmesi: 2024 yılı epigrafik saha çalışması. İçinde O. Cınar, F. Başbuğ & H. Aydemir (Ed.), *Çağdaş dilbilim çalışmaları I: İMÜ Dilbilimi 15. yıl armağanı* (ss. 305–310). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X016>

ÖZ

Yenisey Nehri havzasına ait Türk runik yazıtları, Eski Türk epigrafisinin en kapsamlı ve ayrıntılı biçimde incelenmiş derlemelerinden birini oluşturmaktadır. Üç yüzyılı aşkın bir araştırma geçmişine ve son derece kapsamlı bir bibliyografyaya rağmen, terminoloji, indeksleme ve kaynak belgelenmesine ilişkin temel sorunlar alandaki ilerlemeyi sınırlamıştır. Yenisey yazıtları indeksi; tutarsız kodlama ilkeleri, paralel adlandırmalar, metin içermeyen anıtların kapsama alınması ve Güney Yenisey’e ait ilgisiz yazıtlardan bahsedilmesi gibi nedenlerle zaman içinde düzenleyici işlevini büyük ölçüde yitirmiştir. Buna ek olarak, mevcut belgelerin önemli bir bölümü öznel ya da doğrulanması güç materyallere dayanmaktadır; yayımlanmış fotoğraflar ise çoğu zaman rötuşlanmış ya da güvenilir okumalara olanak tanımayacak ölçüde teknik açıdan yetersizdir. Bu kitap bölümü, runik epigrafik verilerin sistematik biçimde düzenlenmesi ve anıtların modern dijital teknikler kullanılarak yeniden belgelenmesi amacıyla başlatılan geniş ölçekli bir projenin amaçlarını, yöntemini ve bulgularını sunmaktadır. Çalışmada özellikle fotogrametrik üç boyutlu belgeleme yöntemine odaklanılmakta; bu yöntemle, kötü korunmuş yazıtların algoritmik olarak iyileştirilmesine olanak tanıyan yüksek çözünürlüklü çokgen modeller üretilmektedir. Bölümde, 2024 yılında Hakasya

¹ Bu kitap bölümü, Rusya Bilim Vakfı’nın 24-28-01042 numaralı “Eski Türk runik yazıtlarından dijital bir derlemin oluşturulmasına yönelik bilgi sisteminin geliştirilmesi” başlıklı araştırma projesi kapsamında hazırlanmıştır (Destekleyen Kurum: Rusya Bilimler Akademisi Dilbilim Enstitüsü; URL: <https://rscf.ru/project/24-28-01042/>).

ve Tuva’da gerçekleştirilen saha çalışmasının sonuçları özetlenmekte; bu kapsamda, müzelerde ve yerinde (*in situ*) olmak üzere Yenisey runik yazıtlarının 76’sı — bilinen toplam derlemin yaklaşık yarısı — belgelenmektedir. Çalışma bu açıdan, üç boyutlu modellerin yazıt okumalarının yeniden değerlendirilmesi ve gelecekteki araştırmalar için güvenilir ve doğrulanabilir bir kaynak tabanının oluşturulması açısından vazgeçilmez olduğunu ortaya koymaktadır.

Anahtar Sözcükler: Eski Türk runik yazıtları, Yenisey yazıtları, Türk epigrafisi, Runoloji, Belgeleme, Fotogrametri, Üç boyutlu modelleme

ABSTRACT

Turkic runiform inscriptions of the Yenisei River basin constitute one of the largest and most intensively studied corpora of Old Turkic epigraphy. Despite more than three centuries of research and an extensive bibliography, further progress has been impeded by fundamental problems in nomenclature, indexing, and source documentation. The Yenisei index has gradually lost its organizing function due to inconsistent principles of codification, parallel naming, inclusion of non-textual monuments, and the attribution of unrelated South Yenisei inscriptions. In addition, much of the available documentation relies on subjective or poorly verifiable materials, while published photographs are often retouched or technically insufficient for reliable readings. This paper presents the aims, methodology, and results of a large-scale project launched to systematize runiform epigraphic data and to re-document monuments using modern digital techniques. Special emphasis is placed on photogrammetric 3D documentation, producing high-resolution polygonal models that allow algorithmic enhancement of poorly preserved inscriptions. The article summarizes the results of the 2024 fieldwork in Khakassia and Tuva, during which 76 Yenisei runiform inscriptions—approximately half of the known corpus—were documented in museums and *in situ*. The study demonstrates that three-dimensional models are essential for revising readings and establishing a reliable, verifiable source base for future research.

Keywords: Old Turkic runiform inscriptions, Yenisei inscriptions, Turkic epigraphy, runology, Documentation; Photogrammetry, 3D modeling

Sibirya’da Yenisey Nehri havzasını (Krasnoyarsk Krayı, Hakasya ve Tuva bölgelerini kapsayan alan) oluşturan bölgedeki Türk runik yazıtları üç yüzyılı aşkın bir süredir araştırılmaktadır. Bugün itibarıyla bu geniş bölgede, dikilitaşlar, mezar taşları, kaya yüzeyleri ve taşınabilir nesneler üzerine kazınmış 150’den fazla yazıt tespit edilmiştir. Bilimsel araştırmalar süresince, çizimler, estampajlar, grafit sürtme kopyaları, mika kopyaları, analog ve dijital fotoğraflar dahil olmak üzere zengin bir belge birikimi oluşmuştur. Yazıtların büyük çoğunluğuna ilişkin çalışmalar yayımlanmıştır. Mevcut değerlendirmelerimize göre, Yenisey Türk runik epigrafisine ilişkin çalışmaların sayısı en az bindir.

Biriken zengin malzemeye rağmen, Yenisey Türk runik yazıtları üzerine yapılacak ileri araştırmaları engelleyen birkaç ciddi sorun bulunmaktadır. Bunların başında, yazıtların ilişkin adlandırma sisteminin tutarsız olması gelmektedir; aynı anıt birden fazla adla anılabilmektedir (örneğin *Ottuk Daş I* ile *Tuva Monument B*; *Aldı Bel I* ile *Uluğ Kem Kuli*

Kem; Haya Uju ile Haya Bajı; Kres Haya ile Fırkalı). Buna karşılık tek bir ad da birden çok yazım biçimiyle karşımıza çıkabilmektedir (örneğin *Uyuk Tarlağ* ve *Uyuk Tarlak; Çaa Höl* ve *Çakul*’; *Yur Sayır, Yır Sayır* ve *Yir Sayır*).

Paralel adlandırma sorununu kısmen gidermenin bir yolu, yazıtların dizinlenmesidir. Yenisey Türk runik yazıtları için hazırlanan ilk dizin, S. Ye. Malov tarafından oluşturulmuştur (Malov 1952’de E1–E51); buradaki ilk 40 sayı, Radloff (1894–1895) tarafından belirlenen sırayı takip etmektedir. Bu dizin, *Eski Türkçe Sözlüğü*’nde (DTS) E85’e kadar genişletilmiştir. Ardından E86–E145 numaralı yazıtlar Vasilyev (1976; 1983) tarafından dizinlenmiştir. Kormuşin’de (1997; 2008) ise kodlama E154’e kadar devam ettirilmiştir. Ancak bu aşamaya gelindiğinde, Yenisey dizininin yayımlanmış çalışmaları düzenleme işlevini büyük ölçüde yitirdiği artık açıkça görülmekteydi. Yeni keşfedilmiş ancak henüz yayımlanmamış yazıtlara daha sonra numara verebilmek amacıyla dizin numaralarının boş bırakılması yaygın bir uygulama hâline gelmişti; böylece bu durum, belirli bir dizin numarasının hangi yazıtla karşılık geldiğinin kesin biçimde ortaya konmasını imkânsız hâle getirdi.² Ayrıca dizine neyin dâhil edilmesi gerektiğine ilişkin tutarlı bir ilke de bulunmamaktadır — epigrafik malzemeyi taşıyan nesnelerin mi, yoksa bizzat yazıtların kendisi mi dizinlenmelidir? Örneğin Tepsey kayasındaki her bir yazıt, E111–117 ve E123–126 aralığında ayrı dizin numarası alırken, Haya Bajı kayasındaki tüm yazıtlar (sayısı en az 16) tek bir numara altında, E24 olarak toplanmıştır. Dahası, Yenisey dizinine metin içermeyen bazı tamga işaretleri de eklenmiştir (örneğin E89 ve E90). Buna ek olarak, Kızlasov’un (1994) gösterdiği gibi, Yenisey dizin numarası taşıyan bazı yazıtların — örneğin E83, E87, E141, E142 ve E39 olarak listelenen Sulek petroglif alanının büyük bölümü — gerçekte Yenisey Türk runik yazısına değil, henüz çözümlenmemiş ayrı bir sistem olan *Güney Yenisey runik yazısına* ait olduğu anlaşılmıştır.

Bir diğer önemli sorun ise kaynak temelinin durumu ile ilgilidir. Runik epigrafiyle ilgili, doğrulanması oldukça zor olan belgelerin önemli bir bölümü (saha çizimleri vs.) ise öznel bir yorum niteliği taşımaktadır. Estampajlar, grafit sürme kopyaları ve mika kopyaları gibi daha nesnel belgeleme yöntemleri ise genellikle erişimi zor malzemelerdir ve çoğu zaman kötü durumdadır. Fotoğraflar — özellikle basılı yayınlarda yer alanlar — çoğu durumda yazıtın güvenilir biçimde okunmasına olanak sağlamamaktadır; bu durum özellikle yazıtın kendisinin iyi durumda olmadığı örneklerde daha belirgindir. Ayrıca Türk runik yazıtbilimi (runology) araştırmalarında, yazıt fotoğraflarının rötuşlanmış hâllerinin yayımlanması ve çoğu zaman rötuş yapıldığının belirtilmemesi gibi

² Bu konuya ilişkin olarak bkz. Ondar 2020. Ayrıca, “ayrılmış” dizin numaraları E-148 ve E-150 olan iki anıta yönelik bir yorum önerisi sunan yakın tarihli çalışma için bkz. Şabdanov 2022.

yerleşik bir uygulama da bulunmaktadır; ki bu durum sonraki araştırmacıları yanıltmaktadır.³

Biriken sorunları çözmek amacıyla, A. V. Savelyev (Dilbilim Enstitüsü, Rusya Bilimler Akademisi) ve Yu. M. Svoyskiy (RSSDA Laboratuvarı) liderliğinde bir araştırma ekibi, Eski Türk (Orhun, Yenisey ve Talas) runik yazısına ait anıtlara ve Avrasya bozkırlarının diğer runik yazı sistemlerine ilişkin bilgileri sistematik hâle getirmeyi hedefleyen bir proje başlattı. Projenin temel amaçları, bibliyografik bir veri tabanının oluşturulması ve bilinen runik yazıtların listesinin gözden geçirilerek kesinleştirilmesidir. Bununla eş zamanlı olarak, runik yazıtların modern teknik yöntemlerle yeniden belgelenmesine yönelik saha çalışmaları yürütülmektedir. Bu yeniden belgeleme, her bir yazıtın yaklaşık 0,05 mm çözünürlüğe sahip üç boyutlu poligonal bir modelinin üretilmesine imkân veren fotogrametrik 3D modelleme yoluyla gerçekleştirilmektedir. Bu dijital gösterim daha sonra matematiksel algoritmalarla işlenerek yazıtın okunabilirliğinin artırılması sağlanabilmektedir.

2024 yılındaki saha çalışmalarının ana odağı Hakasya ve Tuva bölgeleri olmuştur. Bu çalışmalar sırasında hem müze koleksiyonlarında muhafaza edilen yazıtlar hem de bulundukları yerde (*in situ*) korunan anıtlar belgelenmiştir.⁴

2024 saha çalışmalarının en önemli sonucu, Tuva Cumhuriyeti'nin Kızıl kentinde bulunan Aldan-Maadır Ulusal Müzesinde muhafaza edilen runik yazıtların belgelenmesinin tamamlanması olmuştur. Müzede toplam 41 anıt incelenmiş ve belgelenmiştir — bunlar arasında dikilitaşlar, yeniden kullanılmış geyik taşları, mezar taşları ve işlenmemiş taş bloklar bulunmaktadır: Uyük Turan (E3), Ottuk Daş III (E54), Barık I (E5), Barık II (E6), Barık III (E7), Barık IV (E8), Kara Suğ (E9), Elegest II (E52), Elegest III (E53), Elegest IV (İr-Höl olarak da bilinir, E70), Çaa Höl II (E14), Çaa Höl V (E17), Çaa Höl IV (E18), Çaa Höl VIII (E20), Çaa Höl IX (E21), Çaa Höl X (E22), Çaa Höl XI (E23), Kızıl Çıraa I (E43), Kızıl Çıraa II (E44), Köjeelig Hovu (E45), Tele (E46), Saygın (Barbak Arığ, Bobak Arık veya Bengü Çor'un yazıtı olarak da bilinir, E57), Kezek Hüree (E58), Herbis Baarı (E59), Kara Bulun I (E65), Kara Bulun II (E66), İyme (E73), Övür II (bir yazıt değildir ancak geleneksel olarak Yenisey dizinine E90 olarak dâhil edilmiştir), Bedelig (E91), Demir Suğ (E92), Yur Sayır II (E94), Hemçik Boom II (E96), Ortaa Tey (E99), Bayan Kol (E100), “Kızıl Müzesi Dikilitaşı” (E105), Çerbi (E106), Uyük Oorzak I (E108), Uyük Oorzak II (E109), Uyük Oorzak III (E110), Ongar Hovu I (Eerbek I olarak da bilinir, E147) ve Ongar Hovu II (E148). Bu anıtlardan sekizi

³ Burada ele alınan güçlükler, Avrasya bozkırlarının diğer runik yazı sistemleri için de geçerlidir; bkz. Kızlasov 1994. Runik yazıtların incelenmesi ve bunların yayımlanmasına ilişkin belgeleme sorunlarının güncel durumuna dair genel bir değerlendirme için bkz. Svoyskiy vd. (yayına hazırlık aşamasında).

⁴ Bu makalede anılan runik yazıtlı yüzeylerin belgelenmesi ve üç boyutlu modellenmesi, Yu. M. Svoyskiy ve E. V. Romanenko'nun gözetiminde M. D. Dynin, A. A. Chernukhina, A. A. Pichugina ve A. P. Girich tarafından gerçekleştirilmiştir.

parçalanmış durumda olup iki ya da üç parçaya ayrılmıştır. Bu tür durumlarda her bir parça ayrı ayrı modellenmiş, ardından bireysel modellerden tam taşın bütüncül bir birleşik modeli oluşturulmuştur. Kızıl Müzesi'ndeki yazıtların belgelenmesi için toplam 55.441 fotoğraf çekilmiştir.

Abakan kentindeki L. R. Kızlasov Hakasya Ulusal Yerel Tarih Müzesinde, Yenisey ve Güney Yenisey yazıtları taşıyan yedi anıt belgelenmiştir. Bunların üçü Hakasya kökenlidir: Oznaçennoye II (E104), Uytağ (Yenisey dizin numarası olmayan) ve Uybat VI (E98). Üç anıt ise Tuva bölgesinden Abakan'a getirilmiştir: Edegey I (EYu1), Edegey II (Yu12) ve Edegey III (Yu13). Yenisey Türk runik yazısı içeren bir başka taşın ise kökeni henüz belirlenememiştir.

Hakasya'nın Poltakov köyündeki Kaya Sanatı Müzesi'nde, Apışşire mezarlığından getirilen ve Yenisey Türk runik yazısı içeren büyük bir dikilitaş da belgelenmiştir. Bu yazıt, bilinmeyen nedenlerle bugüne kadar yayımlanmamıştır.⁵

Hakasya Cumhuriyeti sınırlarında yer alan aşağıdaki yazıtlar bulundukları orijinal konumlarında incelenmiş ve biri hariç tamamı belgelenmiştir:

- i. Kara Kurgan dikilitaşı (Uybat V, E34 olarak da bilinir). Bu dikilitaş arazide tespit edilmiş ancak üzerindeki likenlerin temizlenmesi gerektiğinden henüz belgelenmemiştir.
- ii. Uybat Çaatas dikilitaşı (Uybat VIII, E132 olarak da bilinir).
- iii. Ust'-Sos mezar taşı (E134).
- iv. Kök Haya kaya yazıtı (Yenisey dizin numarası yok) — Sibiry Türkçesi sahasında benzeri bulunmayan, kazıma tekniğiyle değil mineral pigment kullanılarak yapılmış eşsiz bir runik yazıttır. Aynı Kök Haya kayasında, bu yazıtın yakın çevresinde ikinci bir runik yazıt daha tespit edilmiştir; ancak bunun sahte olduğunu düşünmekteyiz.

Ayrıca, Uzun Oba dikilitaşının (Uybat IV olarak da bilinir) ve Uybat Çaatas'tan (Uybat VII olarak da bilinir) runik yazıtlı gümüş kabin buluntu yerleri arazide tespit edilmiştir.

Tuva Cumhuriyeti'nde ise 2024 yılı içerisinde müze koleksiyonları dışında yalnızca bir nesne bulunmuş ve belgelenmiştir: Kanmıldığı Hovu dikilitaşı (Şançı II, E62 olarak da bilinir).

Tüm dikilitaşlar ve kaya yazıtları için uydu jeodezisi kullanılarak kesin coğrafi koordinatlar belirlenmiştir.

2024 saha çalışmaları sonucunda, modern yöntemlerle belgelenmiş Yenisey runik yazıtlarının toplam sayısı 76'ya ulaşmıştır; bu sayı, bu grupta bilinen tüm yazıtların yaklaşık yarısına karşılık gelmektedir.⁶

⁵ Bu taşın kökeni ve orijinal koordinatları, E. A. Miklaşevič tarafından belirlenmiştir.

⁶ Avrasya bozkırlarının diğer runik yazı sistemleriyle yazılmış yazıtlar da dikkate alındığında (bkz. I. L. Kızlasov tarafından önerilen "Asya" ve "Avrasya" runik yazı grupları ayrımı), ekibimiz tarafından

Özellikle kötü korunmuş yazıtlar söz konusu olduğunda, toplanan verilerin ön analizleri, üç boyutlu modellerin runik yazıt okumalarının iyileştirilmesi için oldukça gerekli olduğunu göstermektedir; zira bu tür yazıtların geleneksel yöntemlerle sağlıklı biçimde yorumlanması neredeyse imkânsızdır. Araştırması yapılan tüm yazıtlara ilişkin belgeleme materyalleri ile bunların büyük çoğunluğu için gözden geçirilmiş okuma önerileri şu anda yayıma hazırlanmaktadır. Ara sonuçlara proje internet sitesinden ulaşılabilir: <https://rssda.su/workdata/rsf-24-28-01042/>

Kaynaklar

- Nadelyayev, V. M., Nasilov, D. M., Tenişev, E. R., & Ščerbak, A. M. (1969). *Drevnetyurkskiy slovar*. Nauka.
- Kızlasov, I. L. (1994). *Runičeskiye pis'mennosti yevraziyskikh stepey* [Avrasya bozkırlarının runik yazı sistemleri]. Vostočnaya literatura.
- Kormuşin, I. V. (1997). *Tyurkskiye yeniseyskiye epitafii: Teksty i issledovaniya* [Turkic runiform epitaphs: Texts and studies]. Nauka.
- Kormuşin, I. V. (2008). *Tyurkskiye yeniseyskiye epitafii: Grammatika, tekstologiya*. Nauka.
- Malov, S. Ye. (1952). *Yeniseyskaya pis'mennost' tyurkov*. Izdatel'stvo Akademii nauk SSSR.
- Ondar, Č. G. (2020). Perečen' pamyatnikov drevnetyurkskoy pis'mennosti, naydennykh na territorii Tuvy. *Yermolayevskiy čteniya*, (4), 95–101.
- Radloff, F. W. (1894–1895). *Die alttürkischen Inschriften der Mongolei* (Liefer. 1–3). St. Petersburg.
- Svoyskiy, Y. M., Savelyev, A. V., Dynin, M. D., & Romanenko, E. V. (in preparation). *Specifics and technical aspects of applying 3D modeling in the study of runiform inscriptions of the Eurasian Steppes*.
- Şabdanov, A. (2022). E 147b, E 148 ve E 150 numaralı Yenisey yazıtları üzerine. *Dil Araştırmaları*, 16(30), 213–223.
<https://doi.org/10.54316/dilarastirmalari.1079883>
- Vasilyev, D. D. (1976). Pamyatniki tyurkskoy runičeskoy pis'mennosti aziatskogo areala [Asya'daki Türk runik yazıtları]. *Sovetskaya tyurkologiya*, (1), 71–81.
- Vasilyev, D. D. (1983). *Korpus tyurkskikh runičeskikh pamyatnikov basseyna Yeniseya* [Yenisey Vadisi'ndeki Türk runik yazıtlar derlemi]. Nauka.

MORPHOLOGICAL PROCESSING OF TURKISH DERIVED WORDS: DOES BILINGUALISM AFFECT THE PROCESSING ROUTE?¹

Mürselin TAŞAN

Res. Asst., Istanbul Medipol University, Department of English Language Teaching, mtasan@medipol.edu.tr, ORCID: 0000-0003-4677-0927

Serkan UYGUN

Asst. Prof., Bahçeşehir University, Department of English Language Teaching, serkan.uygun@bau.edu.tr, ORCID: 0000-0002-0880-9280

Taşan, M. & Uygun, S. (2025). Morphological processing of Turkish derived words: Does bilingualism affect the processing route? In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 311–334). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X017>

ABSTRACT

It has been suggested that native speakers may develop different processing patterns in their first language as they become proficient second language users. While most of the studies are conducted with heritage speakers whose first language is the minority language in the society they live in, the number of studies that investigate first language as the majority language remains scarce. These studies have shown that high proficiency in a second language can influence first language processing even in the majority language context (van Hell & Dijkstra, 2002; Uygun & Gürel, 2020). The aim of the present study is to explore how proficient Turkish-English late bilinguals process Turkish derived words. 61 monolingual Turkish speakers and 46 proficient Turkish-English late bilingual speakers were tested via a masked priming experiment. The stimuli consisted of (i) transparent words (*dalga* “wave”, *dal* “dive” and *-ga* is the derivational suffix), (ii) opaque words (*karga* “crow”, *kar* “snow” but *-ga* does not function as a derivational suffix), and (iii) form/control words (*devre* “period”, *dev* “giant”, *-re* is not an existing derivational suffix). The results showed no significant group differences in the morphological processing of Turkish derived words. While both monolingual and bilingual speakers employed decomposition for transparent and opaque words, no morphological parsing was observed for the form/control words.

¹ This study was supported by Scientific and Technological Research Council of Türkiye (TUBITAK) under the Grant Number 223K342. The authors thank to TUBITAK for their supports. We also thank Ayşe Gürel for her help and recommendation in designing the stimuli list and Ece Nur Temuçin for her help with participant recruitment and data collection.

These results suggest that not only monolingual but also bilingual speakers decompose derived words regardless of their transparency, suggesting that high proficiency in a second language does not affect the morphological processing route of the first language.

Keywords: Turkish derivation, Morphological processing, Semantic transparency, First language, Proficient late bilinguals

ÖZ

Dil konuşucularının ikinci bir dilde yetkin hale gelince ana dilleriyle ilgili farklı işleme örüntüleri geliştirebilecekleri öne sürülmüştür. Bu alandaki çalışmaların çoğu konuştukları ana dil yaşadıkları toplumda bir azınlık dil olan miras dil konuşucuları ile yürütülürken, konuştukları ana dil yaşadıkları toplumdaki çoğunluk dil olan katılımcılar ile yapılan çalışmaların sayısı çok azdır. Az sayıdaki bu çalışmalar ise, ikinci bir dilde yüksek düzeydeki yetkinliğin ana dilin çoğunluk dil olduğu ortamlarda bile ana dilin işlenmesini etkileyebileceğini göstermektedir (van Hell & Dijkstra, 2002; Uygun & Gürel, 2020). Bu çalışmanın amacı ileri düzeyde İngilizce bilen ve İngilizceyi ileri yaşta öğrenen Türkçe-İngilizce iki dilli konuşucularının Türkçede yapım ekleriyle üretilmiş sözcükleri nasıl işlemlediklerini araştırmaktır. 61 tek dilli Türkçe konuşucusu ile 46 Türkçe-İngilizce iki dilli konuşucusu maskelenmiş hazırlama deneyi kullanılarak test edilmiştir. Çalışmada kullanılan sözcükler (i) anlamsal olarak geçirimli (*dalga* – *DAL*, *-ga* yapım eki olarak işlev göstermektedir), (ii) anlamsal olarak geçirimsiz (*karga* – *KAR*, *-ga* yapım eki gibi görünmesine rağmen bu işlevi göstermemektedir) ve (iii) örtüşük yazımlı sözcüklerden (*devre* – *DEV*, *-re* mevcut bir yapım eki değildir) oluşmaktadır. Sonuçlar Türkçede yapım ekleriyle türetilmiş sözcüklerin biçimbilimsel işlenmesinde iki grup arasında bir fark olmadığını göstermektedir. Her iki grupta yer alan katılımcılar anlamsal olarak geçirimli ve geçirimsiz sözcükleri biçimbirim bileşenlerine ayırarak işlemlerken, örtüşük yazımlı sözcüklerde ise bu çözümleme gözlenmemiştir. Bu sonuçlar, hem tek dilli hem de iki dilli katılımcıların yapım ekleriyle üretilmiş sözcükleri biçimbirim bileşenlerine ayırırken anlamsal geçirimlilikten etkilenmediklerini göstermektedir. Bu durum, ikinci bir dilde yüksek düzeyde yeterliliğe sahip olmanın ana dilin biçimbilimsel işlenmesini etkilemediği anlamına gelmektedir.

Anahtar Sözcükler: Türkçede yapım ekleri, Biçimbilimsel işleme, Anlamsal geçirimlilik, Ana dil, İkinci dili geç öğrenen ve bu dili ileri düzeyde bilen iki dilliler

1. Introduction

A fundamental question in bilingualism concerns how two languages interact in the bilingual mind and whether frequent exposure to a second language (L2) leads to changes in first language (L1) processing. Early work suggested that exposure to more than one language could lead to deviations in linguistic forms, particularly in phonology (Nagle et al., 2023), lexicon (Fridman & Meir, 2023), and syntax (D'Alessandro et al., 2025). When the representation and processing of morphologically complex words are considered, some studies have documented L1-to-L2 influence (Basnight-Brown et al., 2007; Portin et al., 2008), yet recent

research highlights that L2 proficiency can also shape L1 representation and its processing mechanisms (Uygun & Gürel, 2020). This bidirectional influence suggests that bilinguals constantly negotiate between their two linguistic systems, which may impact how they process words in their L1.

The view that even L1 speakers can develop nonnative-like processing strategies as they learn an L2 and become proficient L2 users has been considered relatively counter-intuitive (Clahsen & Felser, 2006). This perspective has given rise to a substantial body of research into heritage speakers (HS), examining whether and how prolonged L2 exposure and L2 proficiency—as well as factors such as the quality and quantity of L1 input, the amount of L1 use, aptitude, and motivation—may reshape the way their L1 is represented and processed. In doing so, researchers have investigated a range of linguistic domains such as phonology, morphology, syntax, lexicon, and semantics, highlighting the complex interplay between L1 and L2 systems (Kasparian & Steinhauer, 2017; Montrul, 2008; Polinsky, 2018; Polinsky & Putnam, 2024; Schmid & Yılmaz, 2021). The L1 of HS has traditionally been studied in immigrant populations, where L1 is no longer the dominant language, and speakers shift towards using L2 in their daily interactions (Isurin & Wilson, 2022). However, there is still no consensus regarding the extent and nature of heritage language effects, leaving open the question of whether they primarily reflect limitations in accessing existing L1 representations, structural shifts prompted by L2 proficiency, or increased cognitive demands during real-time L1 processing (Felser & Uygun, 2022; Kasparian & Steinhauer, 2017; Polinsky & Scontras, 2020; Schmid & Köpke, 2017; Shin, 2024).

The present study focuses on the key issue of whether highly proficient L2 users process derived words in their L1 differently from monolinguals. Understanding this phenomenon is crucial for bilingual morphological processing, as prior research has demonstrated that even when an L1 remains dominant, L2 proficiency can significantly influence the processing of morphologically complex words (Uygun & Gürel, 2020). Much of the existing research has examined HS, whose L1 is a minority language and who receive limited L1 input over time. However, little attention has been paid to how L2 proficiency affects L1 morphological processing when individuals continue to reside in an environment where their L1 serves as the majority language. This gap in the literature raises an important question: To what extent does high proficiency in an L2 reshape L1 morphological processing even when the L1 continues to be dominant?

1.1 L2 effects on L1 in the L2 setting: Morphological processing studies on heritage Turkish

HS are described as a subset of bilinguals raised in homes where a language other than the dominant community language was spoken (Scontras et al., 2018). This is usually the case for immigrants who move

to a country where another language is spoken and for their children who were exposed to a language in their childhood but then were exposed to the language of the host country (Schmid, 2010; Scontras et al., 2015). The study of HS contributes to our understanding of language acquisition, storage, and language processing. By examining the linguistic knowledge of HS, researchers investigate the linguistic system developing under reduced exposure/input conditions to explore how HS acquire their L1, how their L1 develops or fails to develop, how their L1 is represented and processed, which linguistic domains are fully acquired, and which are susceptible to change and deviate from the baseline (Montrul, 2013).

Recently, HS studies have started to investigate the morphological processing of complex word forms by using time-sensitive tools with an attempt to explore the real-time processing of the morphologically complex words. In one of these studies, Jacob and Kırkıcı (2016) compared the morphological processing of Turkish HS living in Germany with Turkish native speakers based on the materials from Kırkıcı and Clahsen's (2013) study, which tested the regular aorist *-Ar* and the deadjectival *-Ilk* nominalizations. Both groups showed significant morphological priming effects not only for the regular aorist (e.g., *sorar* – *SOR* “asks – ASK”) but also for the deadjectival suffix (e.g., *sağlık* – *SAG* “health – HEALTHY”); however, significant priming effects were also observed in the orthographic control condition (e.g., *devre* – *DEV* “period – GIANT”) only for HS, suggesting that, unlike Turkish native speakers, HS rely more on orthographic (surface form) properties of the words as they fail to access morphemic constituents in online processing. In another study, Jacob et al. (2019) investigated another inflected form, the reported past suffix *-mİş* (e.g., *satmış* – *SATMAK* “s/he sold – SELL”) together with the deverbal nominalization suffix *-(y)İcİ* (e.g., *satıcı* – *SATMAK* “seller – SELL”) with native Turkish speakers and HS living in Germany. The results displayed significant priming magnitudes in both groups for the inflected and derived forms. Contrary to Jacob and Kırkıcı (2016), the results of the semantic (e.g., *postane* – *MEKTUP* “post office – LETTER”) and orthographic (e.g., *haziran* – *HAZİNE* “June – TREASURE”) conditions indicated that the observed priming effects were purely morphological in nature because no significant priming effect was obtained for these conditions. The authors concluded that despite acquiring Turkish under disadvantaged conditions, HS can develop native-like processing mechanisms. Uygun and Clahsen (2021) tested the morphological processing of regular (*-Ar*) and irregular (*-Ir*) aorist forms in HS living in Germany and compared their processing routes to Turkish native speakers. While the native speakers and HS exhibited priming effects for both regular (e.g., *duyar* – *DUY* “hears – HEAR”) and irregular (e.g., *gelir* – *GEL* “comes – COME”) aorist forms, the HS also displayed priming effects for the semantic condition (e.g., *kafa* – *BAŞ* “head – HEAD”), which involved synonyms and antonyms. The researchers concluded that HS showed more reliance on meaning-based associations in their processing. Recently, Eldem-Tunç and Fuchs (2024)

sought to describe the morphological processing of the denominal adjective marker *-sİz* (e.g., *sınırsız* – *SINIR* “limitless – LIMIT”) and the deadjectival nominalizer *-lık* (e.g., *hastalık* – *HASTA* “sickness – SICK”) in Turkish HS living in the US. The researchers observed similar processing patterns in Turkish native speakers and HS and proposed that morphological decomposition was not affected under heritage language conditions.

As can be seen, more research is needed on the morphological processing patterns of the HS to reach any generalizable conclusions because of the diverse observations made. Research on HS focuses on the effects of L2 (the majority language of the society) on L1, which is the minority language and functions as the first language of the immigrants. However, only a handful of studies have questioned the effect of proficient L2 knowledge on L1 speakers, who reside in a country where L1 is the dominant language.

1.2 L2 effects on L1 in the L1 setting: Morphological processing studies on L1 Turkish

Since there is only one study that explores the effect of L2 knowledge on the morphological processing of L1 Turkish, we will briefly summarize some studies that focused on the same concept but tested different phenomena, namely word fluency and word recognition. In one of these studies, Ayçiçeği-Dinn et al. (2017) questioned whether an L1 would suffer when students took all their classes in a foreign language even if they lived in their home country. They compared students, who were studying psychology, economics, and English literature at a university in Türkiye. The researchers found significantly lower word fluency scores for the English literature students when compared to psychology and economics students. The self-reported L1-Turkish speaking and writing abilities were also found to be smaller in this particular group. The researchers concluded that prolonged engagement with an L2 might put L1 at risk for temporary attrition. More recently, Ma and Vanek (2024) investigated how learning an L2 might influence a person's L1 even when they are surrounded by the L1 environment. Specifically, they recruited two groups of participants: 25 Chinese teachers of English and 25 Chinese teachers who did not teach English from the same school. They employed a time-sensitive word decision task where the participants had to quickly decide if the presented words were real Chinese words. They reported that Chinese teachers of English were slower to respond to common Chinese words in the word decision task. These findings are taken to support the view that extended use of and exposure to L2 can result in some weakening of a person's L1 spoken vocabulary, even while they are mainly using their L1 in everyday life.

The only research that most closely aligned with the context of the present study is that of Uygun and Gürel (2020). The researchers examined the potential changes in the processing of morphologically complex words, namely Turkish compounds in monolingual Turkish speakers and

proficient Turkish-English late bilinguals, by employing a masked priming experiment. Their stimuli consisted of fully transparent compounds, in which the meanings of both words were related to the meaning of the whole word (e.g., *kuzeydoğu* “northeast”, *kuzey* “north”, *doğu* “east”), and partially opaque compounds, in which the meaning of one word was not related to the meaning of the whole word (e.g., *büyükelçi* “ambassador”, *büyük* “big”, *elçi* “delegate”). The results showed that monolingual participants decomposed Turkish compounds regardless of transparency. While they were able to activate both constituents in fully transparent compounds, they only activated the second constituent, which served as the head, in partially opaque compounds. Conversely, the bilingual group could not activate any of the constituents in partially opaque compounds and was able to activate only the first constituent (i.e., nonhead) in fully transparent compounds. These findings suggest qualitatively and quantitatively different processing patterns between monolinguals and bilinguals, even though the bilinguals resided in the L1 country.

In summary, the number of studies that deal with the question as to whether proficient L2 knowledge affects L1 in speakers residing in the L1 country is scarce, and the available studies have employed various data sets and data collection techniques. The aim of the present study is to explore how proficient Turkish-English late bilinguals living in Türkiye process Turkish derived words in comparison to Turkish monolinguals. Before moving on to the study, we will briefly introduce different accounts for the processing of derived words and the previous studies conducted with L1 and L2 speakers.

1.3 Previous research on the morphological processing of derivation in L1 and L2

For over the past fifty years, psycholinguistic research has been focusing on the processing of morphologically complex words (e.g., *singing*, *singer*) with the goal of developing a model applicable across languages. Three main models have been proposed for the morphological processing of complex words. According to the Decomposition Model (Taft, 1979, 2004; Taft & Forster, 1975), morphologically complex words are parsed into their morphemic constituents (*sing+ing*, *sing+er*) prior to lexical access, suggesting that affixes and roots are represented separately in the lexicon. In contrast, the Full-Listing Model (Butterworth, 1983) posits that morphologically complex words are stored as whole units and are not decomposed into their constituents, assuming no morphological computation during word recognition. In contrast to these two views, research has also revealed that several factors such as, familiarity (e.g., Caramazza et al., 1988), frequency (e.g., Schreuder & Baayen, 1995), regularity (e.g., Clahsen et al., 2010), and semantic transparency (e.g., Schreuder & Baayen, 1995) play a crucial role in determining the processing route, leading to the proposal of the dual-route model. Although the findings of studies using masked priming experiments have shown that

morphologically complex words are processed by decomposing them into their morphemic constituents, debate still continues about how exactly the process of decomposition takes place and what factors affect this process.

Within this broader interest and ongoing discussions, research on derived word processing has expanded because words constructed with derivational suffixes may differ in their semantic transparency. Semantic transparency refers to the degree to which the meaning of a morphologically complex word can be inferred from its root form. In semantically transparent words, the meaning of the derived word can be computed from its root (e.g., *teacher* – *TEACH*, *singer* – *SING*). For example, *-er* is a derivational suffix in English that is added to verb forms and refers to the agent who does the action of the verb as a profession (i.e., a *teacher* is someone who *teaches* and does this as a profession). In contrast, semantically opaque words consist of a root and a pseudo-suffix that looks like a real derivational suffix, and the relationship between the derived word form and the root is not genuinely morphological (e.g., *corner* – *CORN*, *number* – *NUMB*). In these examples, *-er* resembles the derivational suffix in its form but does not function as a real derivational suffix; therefore, these words are also classified as pseudo-derived words (i.e., a *corner* is not someone who *corns*). Whether semantic transparency influences the morphological processing of derived words has been debated extensively, and the design of these studies usually employs three types of word forms: (i) semantically transparent (e.g., *teacher* – *TEACH*), (ii) semantically opaque (e.g., *corner* – *CORN*), and (iii) form/control words (e.g., *freeze* – *FREE*). The words in the form/control condition do not have a morphological relationship because *-ze* does not exist as a suffix in English; therefore, the word pairs in this condition resemble each other only in terms of orthography.

Different theories have been proposed for the morphological processing of derived words. The first one is the form-then-meaning account (Rastle et al., 2000, 2004). This account posits that both semantically transparent and opaque words are decomposed into their morphemic constituents, and semantic transparency does not play a role in the early stages of morphological processing. A substantial body of research investigating L1 processing through the masked priming paradigm found support for the form-then-meaning account (Boudelaa & Marslen-Wilson, 2005, in Arabic; Beyersmann et al., 2016; Rastle et al., 2000, 2004, in English; Longtin & Meunier, 2005, in French; Fleischhauer et al., 2021, in German; Frost et al., 1997, in Hebrew; Marelli et al., 2013, Experiment 1, in Italian; Heyer & Kornishova, 2018, in Russian; Lázaro et al., 2016, 2021, in Spanish). In one influential study, Rastle et al. (2004) examined the role of semantic transparency in derived word recognition using a masked priming design with 42 ms stimulus-onset asynchrony (SOA). The experiment included prime-target pairs for semantically transparent (e.g., *cleaner* – *CLEAN*), semantically opaque (e.g., *corner* – *CORN*), and form/control (e.g., *brothel* – *BROTH*) pairs. The results revealed

significant priming effects in both the transparent and opaque conditions, with priming magnitudes of 27 ms and 22 ms, respectively, indicating decomposition. However, the form/control condition yielded only a 4 ms effect, which was not significant and was interpreted as whole-word access. These findings were taken as support for the form-then-meaning account because in the early stages of visual word recognition, the parser uses morpho-orthographic information to process morphologically complex words and decomposes all derived words irrespective of semantic transparency.

Conversely, the form-with-meaning account suggests that semantic information plays a crucial role in the processing of complex words, and transparent words are processed easier than opaque words (Feldman et al., 2009, 2012, 2015). For example, Feldman et al. (2009) conducted a masked priming study with 50 ms SOA to explore the role of semantic transparency in processing English derived words. The results showed a significant priming magnitude difference favoring semantically transparent primes (30ms) over opaque primes (4 ms). The authors concluded that morpho-semantic information influenced the early stages of processing derived words. Most of the research supporting the form-with-meaning account indicated that in semantically opaque prime-target word pairs, prime words either showed no impact or had a weaker effect on their processing when compared to semantically transparent prime-target word pairs, which exhibited significantly larger effects.

Certain models have been proposed to interpret the obtained results. The most outstanding model is the Hybrid Model (Diependaele et al., 2005, 2009, 2011), which includes a dual-route processing model. In one of these studies with L1 English speakers, Diependaele et al. (2009) sought to describe the processing route of transparent and opaque English derived words. They observed a priming magnitude of 36 ms for transparent, 15 ms for opaque, and 1 ms for form/control conditions. These results showed that both morpho-orthographic and morpho-semantic information get activated in parallel upon encountering a morphologically complex word. Semantically transparent prime-target pairs use both morpho-orthographic and morpho-semantic information, resulting in much larger facilitation effects, whereas semantically opaque pairs only receive morpho-orthographic information, leading to less facilitation effect.

To the authors' knowledge, only two studies have explored this topic with native Turkish speakers. In one of these studies, Oğuz and Kırkıcı (2023) examined the processing of complex words by Turkish-speaking 2nd and 4th graders via a masked priming experiment by using four types of prime-target pairs: transparent (e.g., *gururlu* – *GURUR* “proud – PRIDE”; *gururla* – *GURUR* “with proud – PRIDE”), opaque (e.g., *elma* – *EL* “apple – HAND”), formal overlap (e.g., *hapis* – *HAP* “prison – PILL”), and semantic relation (e.g., *çatı* – *EV* “house – ROOF”). The transparent condition included both inflected and derived forms. A significant priming effect emerged only for transparent pairs, indicating

that Turkish children were sensitive only to suffixes in the early stages of word processing. In another study, Çağlar (2019) investigated the processing of transparent (e.g., *çizim* – *ÇİZ* “drawing – DRAW”), opaque (e.g., *tuzak* – *TUZ* “trap – SALT”), and form-overlap (e.g., *kasap* – *KAS* “butcher – MUSCLE”) word pairs among adult Turkish speakers. Results showed significant priming for both transparent (26 ms) and opaque (13 ms) words, but not for form overlap (7 ms). However, since the study primarily aimed to explore the transposed letter effect at morpheme boundaries, it was not fully suited to comparing the obtained priming magnitudes between transparent and opaque forms. Therefore, it remains an open question whether Turkish native speakers are influenced by semantic transparency when processing derived words.

Only a handful of studies have examined how semantic transparency affects derived word processing in late bilingual speakers. Diependaele et al. (2011) compared English monolinguals and Spanish-English and Dutch-English bilinguals using a masked priming task with 53 ms SOA. Both groups showed a similar graded pattern of priming effect, with the largest priming magnitude for transparent pairs and the weakest for form overlap. The opaque pairs elicited a priming magnitude in between the transparent and form overlap pairs. These results suggest that both monolinguals and bilinguals employ the same processing strategies and rely on both morpho-orthographic and morpho-semantic information for transparent pairs and only on morpho-orthographic information for opaque pairs. In another study, Zhang et al. (2017) examined Chinese-English bilinguals using transparent, opaque and semantic pairs. In Experiment 1, they used a 40 ms SOA and found significant priming effect only for the transparent condition. By using an 80 ms SOA, they observed significant priming effects for opaque primes only. Semantically related pairs did not produce any priming effects. However, the absence of a monolingual control group and form/control condition together with the discrepancies of the results limited the conclusions of the study. By employing inflected, derived, pseudo-complex, and stem-embedded forms, Lõo et al. (2022) investigated morphological processing in monolingual and bilingual English speakers. While the monolingual speakers showed priming only for true morphological relations (i.e., inflected and derived), the bilingual speakers showed priming across all conditions. However, the bilingual participants varied in L1 backgrounds, their proficiency was self-rated, and semantic transparency was not controlled; therefore, the results cannot be generalized. Finally, Viviani and Crepaldi (2022) found significant priming for transparent and opaque forms both in monolingual and bilingual English speakers together with a graded priming pattern favoring transparent items. Interestingly, they also observed significant priming effect in the form/control condition but only in the bilingual group. As can be seen, studies on L2 English speakers remain limited, and the results of these studies are not sufficient to give a general idea about the

morphological processing patterns of L2 English speakers and if semantic transparency affects their morphological processing pattern.

2. Method

2.1 Sample

We tested 61 monolingual Turkish speakers (31 females, between ages 18-55, $M = 29.60$, $SD = 9.86$), who were residing in Türkiye. 29 of them were high school graduates, 24 were university graduates, 7 were graduates with a master's degree, and 1 had a doctoral degree. This group will be referred to as monolingual speakers (MS). In addition, 46 highly proficient Turkish-English late bilingual speakers (27 females, between ages 20-43, $M = 26.00$, $SD = 6.54$) residing in Türkiye were tested. This group will be referred to as Turkish-English bilingual speakers (TEBS) and included 21 individuals with a high school diploma, 15 with a university degree, 8 with a master's degree, and 2 with a doctoral degree. All participants in the TEBS group had learned English when they started school (between ages 7-11, $M = 8.93$, $SD = 1.11$) and attended English lessons for a long time (between years 9-15, $M = 12.48$, $SD = 1.47$). To evaluate their English proficiency, all participants completed the Oxford Quick Placement Test, which consisted of 40 multiple choice questions that measured the grammatical and vocabulary knowledge of the participants. The results of the test indicate high scores in English proficiency ($M = 34.98$, $SD = 2.61$, *Maximum Score* = 40). The self-rating scores of the participants also showed a high proficiency in their English (Speaking: $M = 6.98$, $SD = 1.39$; Listening: $M = 7.98$, $SD = 1.29$; Reading: $M = 8.30$, $SD = 1.23$; Writing: $M = 7.26$, $SD = 1.42$, *Maximum Score* = 10). The TEBS group also used English actively in their daily lives (English: $M = 53.48$, $SD = 11.65$; Turkish: $M = 46.52$, $SD = 11.65$, *Maximum Score* = 100). All participants in both groups were native speakers of Turkish and spoke standard Turkish. All were right-handed and did not have any health concerns such as issues with sight, hearing, learning, or language. All participants were informed about the study, and their consent forms were collected.

2.2 Instrument

The stimuli list consisted of three conditions: (i) semantically transparent words (e.g., *dalga* – DAL “wave – DIVE”), (ii) semantically opaque words (e.g., *karga* – KAR “crow – SNOW”), and (iii) form/control words (e.g., *devre* – DEV “period – GIANT”). In semantically transparent words, the meaning of the derived word can be computed from its root (e.g., *dalga* – DAL “wave – DIVE”). For example, *-ga* is a derivational suffix in Turkish that is added to verb forms to derive a noun. Semantically opaque words, however, consist of a root and a (pseudo)suffix that looks like a real derivational suffix, and the relationship between the derived word form and the root is not genuinely morphological (e.g., *karga* – KAR

“crow – SNOW”). As can be seen in the example, *-gA* resembles the derivational suffix orthographically but does not function as the real derivational suffix *-gA*. Finally, form/control words do not involve an existing derivational suffix; therefore, these word pairs do not involve any morphological relationship (e.g., *devre – DEV* “period – GIANT”). To illustrate, the word *dev* “giant” is completely involved in the word *devre* “period”, but *-re* is not an existing suffix in Turkish; therefore, these word pairs do not have any morphological or semantic relation but resemble each other only in terms of orthography. 22 prime-target pairs were created for each condition, making a total of 66 pairs. Table 1 presents a sample set of the prime-target word pairs used across three conditions in the masked priming experiment.

Table 1. Sample prime-target word pair sets across three conditions

Condition	Prime Word	Target Word
Transparent	<i>dalga</i>	DAL
Opaque	<i>karga</i>	KAR
Form/Control	<i>devre</i>	DEV

We started with a total of 721 prime-target word pairs across three conditions. Çotuksöken (2011) was used as a reference for checking the existence of the Turkish derivational suffixes. The spelling and etymology of the selected words were examined and verified through the online Turkish dictionary of the Turkish Language Association (<https://sozluk.gov.tr/>). Following this stage, a transparency judgment questionnaire was administered to determine the level of semantic transparency between the root form and the derived form (e.g., *DAL – DALGA* “DIVE – WAVE”). The questionnaire was prepared via Google Forms and implemented online. 70 L1 Turkish speakers, who did not participate in the masked priming experiment, were asked to rate the level of semantic transparency of the root-derived form pairs on a seven-point scale (1: strongly unrelated, 7: strongly related). Based on the results of the judgment questionnaire, semantically transparent and opaque root-derived form pairs were determined. Root-derived form pairs with a mean score of 5.5 and over were classified as semantically transparent, while the pairs with a mean score of 2.5 and below were categorized as semantically opaque, and these pairs were used as the prime-target pairs in the masked priming experiment.

In the next step, several elimination criteria were applied to semantically transparent and opaque prime-target pairs. The first elimination criterion was syllable count. Target words with more than one syllable and prime words with more than three syllables were excluded from the list. The second elimination criterion was the number of letters. Target words with more than four letters and prime words with more than seven letters were removed. The next elimination criterion was interlingual homographs. Interlingual homographs are words that exist in multiple languages but do not carry the same meaning (Poort, 2018). For instance, the Turkish word *boy* means height, while in English, the same word refers

to a male child. The final elimination criterion was the root-root relationship in the prime words. For example, the word *etçil* “carnivore” is formed by adding the suffix *-Cil* to the root *et* “meat”. The suffix *-Cil* carries the meaning of ‘inclined to or related to something specific’. It should be noted, however, that *çil* “freckles” is also a noun in Turkish. Therefore, all prime words with a root-root relationship were removed from the list of words to be used in the experiment. As a result of this elimination process, a total of 44 prime-target word pairs were determined, with 22 prime-target pairs in each semantically transparent and semantically opaque category. The same elimination criteria were applied to the 224 form/control prime-target word pairs, and as a result, 22 prime-target pairs were selected.

These 66 prime-target word pairs were classified as related because the primes were morphologically, orthographically, and/or semantically related to the target words. These prime and target words were matched with each other in terms of frequency, length, and orthographic overlap. The frequency of the prime and target words was checked using TS Corpus (Sezer, 2017). Length was controlled in terms of the number of letters. The orthographic overlap was also calculated because target words were fully contained in the related prime words. The mean values of these variables for three different word groups were compared using a one-way ANOVA test, and no statistically significant differences were obtained for any of the variables (target word frequency $F(2, 65) = 0.49$; $p = .612$; prime word frequency $F(2, 65) = 0.61$; $p = .546$; target word length $F(2, 65) = 0.77$; $p = .464$; prime word length $F(2, 65) = 2.35$; $p = .104$; orthographic overlap $F(2, 65) = 2.85$; $p = .065$).

Table 2. Sample related and unrelated prime-target word pair sets across three conditions

Condition	Related Prime Word	Unrelated Prime Word	Target Word
Transparent	dalga	tatlı	DAL
Opaque	karga	evlat	KAR
Form/Control	devre	açlık	DEV

In addition to the 66 related primes, we added 66 unrelated primes. These unrelated primes were statistically matched with the related primes in terms of length ($t(65) = 0.13$, $p = .894$) and frequency ($t(65) = 0.56$, $p = .575$). The unrelated prime words were entirely independent from the target words they were paired with by means of morphology, semantics, and orthography (Table 2). In addition, filler words were added to the experiment to prevent participants from realizing the purpose of the experiment. The prime words used in fillers comprised simple past tense, reported past tense, and the locative and ablative case suffixes (e.g., *bindi* – *SÜT* “got on – MILK”). In the filler word pairs, the target words ($F(3,128) = 2.19$, $p = .093$) and the prime words ($F(3,128) = 1.47$, $p = .234$) were statistically matched in terms of length across three different word

conditions. Finally, 132 nonwords were added to the stimuli list so that half of the words in the experiment required a ‘no’ response. Nonwords were created by changing 1-2 letters of an existing word without violating the phonotactic rules of Turkish (e.g., *ZİÇ*). The existence of the nonwords was checked via the online Turkish dictionary (<https://sozluk.gov.tr/>).

2.3 Design

By employing E-Prime 3.0 (Psychology Software Tools, Pittsburgh, PA), a masked priming experiment was conducted to accurately measure the participants’ reaction times (RTs) and accuracy. Because each target word had a related and an unrelated prime, the prime-target word pairs were distributed over two presentation lists by following a Latin Square design to ensure that no participant saw the same target word more than once. All participants completed the experiment in a quiet, dimly lit room with minimal distractions. Before the experiment started, all participants were asked to fill in the demographic background questionnaire and the consent form. After that, detailed instructions about the procedure of the experiment were provided. At the beginning of the experiment, a short practice session was conducted to familiarize the participants with the masked priming procedure. In the practice session, there were a total of 10 items. The experiment began when participants declared that they understood the procedure and were ready to proceed. During the main experiment, participants were asked to decide as accurately and quickly as possible whether the letter string shown on the computer screen was a Turkish word by pressing the designated ‘YES’ or ‘NO’ buttons. The ‘YES’ button was always on the side of the participant's dominant hand. Each trial started with a blank screen displayed for 500 milliseconds, followed by a mask (#####) that remained on the screen for 500 milliseconds. The number of hashes in the mask was equal to the number of letters that made up the prime word. Following the mask, the prime word in lowercase letters was shown to the participants for 42 milliseconds. After the prime word was displayed, the target word in capital letters appeared and remained on the screen until the participant made a response. The experiment lasted approximately 10 minutes. After the experiment, the bilingual participants completed the Oxford Quick Placement Test.

2.4 Data analysis

The analysis of participants’ RTs was carried out only on the correct responses. Thus, all incorrect answers (MS = 6.56%; TEBS = 5.53%) of the participants were removed before the data analysis. Then, the data was log-transformed, and RTs that were above or below 2.5 standard deviations (MS = 2.47%; TEBS = 2.51%) of a participant’s mean log RTs were removed. The remaining log-transformed RT data was analyzed using the R statistical program (R Development Core Team, 2021).

A linear multiple regression model was used in the analysis of the data, and *lme4* package was used in the model building process (Bates et al., 2015). In the data analysis, participants' RTs were the dependent variable, while *condition* (Transparent vs. Opaque vs. Form), *prime type* (Related vs. Unrelated) and *group* (MS vs. TEBS) were the independent variables. The models included item and participant as random effects, and condition, prime type, group, and their interactions as fixed effects. For single comparisons, treatment contrasts were employed, while sum-coded contrasts (-0.5, 0.5) were employed for the main effects and interactions. In simplifying the model (Barr et al., 2013), random slopes for item and participant were retained for each fixed effect only when they significantly enhanced model fit, as measured by the Akaike Information Criterion (AIC). The comparison made with the lowest AIC score aims to keep only significant variables in the model and penalize complexity (Venables & Ripley, 2002). The best-fit model only included the random slope for prime type by participant. The effect sizes are reported by using model coefficients in log odds (β), standard errors (*SE*), *t*-statistics, and *p* values. The *p* values in the data analysis were calculated using the *lmerTest* package (Kuznetsova et al., 2014).

In the analyses, participants' RTs for related and unrelated prime-target word pairs in each condition were compared. If the RT difference (i.e., priming) between the related vs. unrelated prime-target word pairs in a condition is significant, this suggests that the prime facilitates the recognition of the target words in this condition, indicating morphological decomposition. Conversely, a non-significant RT difference in a condition indicates no facilitation of the prime word, resulting in storing the target words of that condition as a whole unit in the mental lexicon. Following this, the three different word conditions were compared with each other to unveil whether differences exist in their processing patterns.

3. Findings

Table 3 provides the back-transformed mean RTs by condition, prime type, and group.

Table 3. Back-transformed mean RTs (and standard deviations) for conditions and groups

	MS			TEBS		
	Transparent	Opaque	Form	Transparent	Opaque	Form
Related	562.78 (156.77)	592.57 (185.97)	610.84 (173.81)	587.44 (157.59)	602.86 (170.99)	637.88 (196.06)
Unrelated	590.34 (143.81)	607.42 (192.81)	600.91 (149.09)	614.11 (159.78)	620.76 (170.20)	629.27 (151.33)
Priming	27.56 ms	14.85ms	-9.93 ms	26.67 ms	17.90ms	-8.61 ms

Regarding the RT data, Table 4 presents the results from the best-fit model testing the morphological priming effects. The results showed a main effect of prime type (β : -0.037, *SE*: 0.010, *t* = -3.607, *p* < .001) indicating that participants employed decomposition in general when

processing the words. However, the main effect of group was not significant, revealing no difference between the groups. In addition, we did not obtain a significant interaction of prime type and group, which showed that both groups used the same processing strategies when processing the words used in the experiment (see Table 4a).

Table 4. Linear mixed effects model output of the best-fit model

	β	<i>SE</i>	<i>t</i>	<i>p</i>
<i>(a) Overall Model</i>				
Prime type (related vs. unrelated)	-0.037	0.010	-3.607	.000
Group (MS vs. TEBS)	-0.035	0.026	-1.315	.188
Prime type*Group	0.002	0.009	0.156	.876
<i>(b) Revelled for Transparent</i>				
Condition1 (Transparent vs. Opaque)	-0.037	0.012	-3.084	.002
Condition2 (Transparent vs. Form)	-0.052	0.012	-4.299	.000
Prime type (related vs. unrelated)	-0.062	0.017	-3.623	.000
Group (MS vs. TEBS)	-0.040	0.027	-1.489	.136
Condition1*Prime type	0.015	0.022	0.730	.466
Condition2*Prime type	0.059	0.023	2.559	.010
Condition1*Group	0.019	0.012	1.688	.092
Condition2*Group	0.002	0.012	0.203	.838
Prime type*Group	0.003	0.016	0.159	.874
Condition1*Prime type*Group	0.009	0.024	0.387	.698
Condition2*Prime type*Group	0.006	0.024	0.247	.804
<i>(c) Revelled for Opaque</i>				
Condition3 (Opaque vs. Form)	-0.016	0.011	-1.432	.152
Prime type (related vs. unrelated)	-0.046	0.016	-2.894	.004
Group (MS vs. TEBS)	-0.021	0.027	-0.757	.450
Condition3*Prime type	0.044	0.022	2.015	.044
Condition3*Group	0.022	0.012	1.884	.060
Prime type*Group	0.006	0.017	0.383	.702
Condition3*Prime type*Group	0.015	0.024	0.629	.530
<i>(d) Revelled for Form</i>				
Prime type (related vs. unrelated)	0.003	0.016	0.169	.866
Group (MS vs. TEBS)	-0.043	0.027	-1.575	.116
Prime type*Group	0.008	0.017	0.505	.614

Formula in R: $\log(\text{RT}) \sim \text{condition} * \text{prime type} * \text{group} + (1 | \text{item}) + (1 + \text{prime type} | \text{participant})$.

As a next step, analyses were carried out by releveing each condition with an attempt to explore how the releveled condition was processed, if the groups differed from each other in processing that condition, and whether that condition was processed differently from the other two conditions. When the condition was releveled for transparent words (see Table 4b), a significant main effect of prime type (β : -0.062, *SE*: 0.017, $t = -3.623$, $p < .001$) was obtained. This effect showed that transparent words were decomposed into their constituent morphemes. In addition, we observed a significant main effect of condition, in which transparent words were processed significantly faster than opaque (β : -

0.037, $SE: 0.012$, $t = -3.084$, $p < .003$) and form ($\beta: -0.052$, $SE: 0.012$, $t = -4.299$, $p < .001$) conditions. Most importantly, no significant main effect or interactions involving the factor group were found, which indicated that both the MS and TEBS groups employed the same processing mechanism, namely decomposition, when processing transparent derived words. Finally, a significant two-way interaction of condition and prime type between transparent and form conditions ($\beta: 0.059$, $SE: 0.023$, $t = 2.559$, $p < .011$) demonstrated that while transparent words were accessed via decomposition, the words in the form condition were accessed by employing full-listing.

Table 4c displays the results of the analysis when the condition was relevelled for opaque words. The significant main effect of prime type ($\beta: -0.046$, $SE: 0.016$, $t = -2.894$, $p < .005$) reflected that conscious analysis of constituent morphemes was also employed for opaque words. Mirroring the transparent condition, no significant main effect of group and no significant interaction with the factor group were obtained. We also found a significant two-way interaction of condition and prime type between opaque and form conditions ($\beta: 0.044$, $SE: 0.022$, $t = 2.015$, $p < .045$). This significant interaction implied that participants employed different processing mechanisms when accessing these two conditions, namely decomposition for opaque words but whole-word access for words in the form condition.

Finally, the releveling was conducted to the form condition (see Table 4d). Unlike the transparent and opaque conditions, we did not observe a significant main effect of prime type ($p = .866$), indicating that morphological decomposition does not take place in accessing the words in the form condition, and both groups processed these words as unanalyzed whole forms. These results demonstrated that both the MS and TEBS groups used the same processing mechanisms when accessing not only transparent and opaque derived words but also the words in the form condition. Besides, the semantic transparency of the derived words did not influence the processing routes of both groups, providing further support for the form-then-meaning account.

4. Discussion

By employing a masked priming experiment, the aim of the present study was to explore how proficient Turkish-English late bilinguals residing in Türkiye process semantically transparent and opaque Turkish derived words and compare their processing patterns to the monolingual speakers.

The results of the RT data showed that both the MS and TEBS groups used the same processing routes when processing derived words. More specifically, they decomposed Turkish derived words regardless of transparency, which indicates that the transparency level of the derived words does not affect the processing mechanisms, supporting the claims of the form-then-meaning account. Despite obtaining larger priming

magnitudes for the transparent words ($MS = 27.56$ vs. $TEBS = 26.67$) when compared to opaque words ($MS = 14.85$ vs. $TEBS = 17.90$), this difference did not turn out to be significant and was only a numerical difference. In addition, both groups did not exhibit morpheme-based processing for the words in the control condition, and the obtained priming magnitudes for this condition did not indicate any facilitation effect but an inhibition effect for both groups ($MS = -9.93$ vs. $TEBS = -8.61$). These results clearly demonstrate that both groups employed a morpho-orthographic decomposition when processing derived words, and they used morpheme-based lexical access when they processed root + (pseudo)suffix forms but no morphological analysis for root + nonexistent forms. The results of the MS group are in line with many previous studies conducted in various languages with L1 speakers (see Hasenäcker, 2016). The results also resonate with the studies that compared L1 and L2 speakers because our results posit no morphological processing differences between the groups (see Diependaele et al., 2011).

However, the results of the present study contradict Uygun and Gürel (2020), which is the only study that is most closely aligned with the present study. While Uygun and Gürel (2020) reported different processing mechanisms between the MS and TEBS groups, our results did not confirm their findings. There are several reasons why the findings of the present study do not confirm the results of Uygun and Gürel (2020).

The first reason might be related to the certain characteristics that distinguish compounding and derivation from each other. For example, Booij (2005, p. 109) emphasized that compounds consist of the combination of two or more lexemes (e.g., *football* is formed by the stems of *foot* and *ball*), whereas derivation is characterized by the addition of an affix (i.e., bound morpheme) to a stem (e.g., *artist* is formed by the stem *art* and the bound morpheme *-ist*). In addition, ten Hacken (1994) proposed that compound elements can either be the head or non-head of the structure, but affixes are of a different nature. For instance, derivational suffixes are always placed in the final position of a derived word, so the position of the suffix is quite predictable. However, in compound words, the position of each constituent is unpredictable because, for example, the lexeme *book* is the first constituent in the English compound *bookstore* and the second constituent in *bankbook*. Finally, semantic transparency can play a more privileged role in compounds because compound words consist of root + root words rather than root + affix forms. While derived forms are divided into two in terms of transparency (transparent vs. opaque), compounds words are divided into four groups depending on the transparency of the constituents (Libben et al., 2003): transparent-transparent (e.g., *bedroom*), opaque-transparent (e.g., *nickname*), transparent-opaque (e.g., *shoehorn*) and opaque-opaque (OO) (e.g., *deadline*). These characteristic differences between compounds and derived forms might have led to the observed differences between the two studies since compounding is a more complex word formation process

when compared to the derived forms, and semantic transparency plays a more prominent role in the processing of compounds.

The second and more crucial reason of the divergent results is related to the design of the studies. While both studies used masked priming experiments, the presentation of the prime-target pairs differed from each other. Uygun and Gürel (2020) used the roots of the compound as primes and the compound word as the target (e.g., *kuzey* – *KUZEYDOĞU* “north – NORTHEAST” and *doğu* – *KUZEYDOĞU* “east – NORTHEAST”). The design of Uygun and Gürel (2020) provides more information on the representation of the complex words in the mental lexicon and investigates whether seeing the constituent of a compound word for a very short time facilitates the access and activation of the compound word in the mental lexicon, which is also known as mental representation. Conversely, the present study used the derived word as the prime and the root form as the target (e.g., *dalga* – *DAL* “wave – DIVE”). The design of the present study allows us to investigate if seeing the complex form of a word very shortly facilitates the access of its root form. If this facilitation works, this means the participants decompose the complex form into its morphemic constituents and can reach the root forms successfully. This design provides precise information on the morphological processing of the complex word forms in different groups with different language learning histories and is a widely accepted design in the masked priming experiments.

All in all, the differences between the processes of compounding and derivation, together with the differences in the presentation of the prime-target pairs, might be the reasons for the divergent results of the two studies that investigate the effect of high proficiency in L2 on L1 word processing.

5. Conclusion

The present study questioned the potential changes in proficient Turkish-English late bilingual speakers’ processing of Turkish derived words. By employing a masked priming experiment, which measures the accuracy and the RTs of the participants, the study focused on the extent to which high L2 proficiency can modulate the morphological processing of Turkish derived words in the bilingual speakers living in Türkiye when compared to monolingual speakers. The findings revealed that both groups employed decomposition for derived words regardless of semantic transparency. In addition, both groups also used whole-word access for the words in the form condition. These results indicate no qualitative nor quantitative differences between the bilinguals and monolinguals. However, more studies on different word formation processes that employ the same experimental design are needed to draw any safe conclusions on the morphological processing of bilinguals living in the L1 setting.

References

- Ayçiçeği-Dinn, A., Şişman-Bal, S., & Caldwell-Harris, C. L. (2017). Does attending an English-language university diminish abilities in the native language? Data from Turkey. *Applied Linguistics*, 38(4), 540-558. <https://doi.org/10.1093/applin/amv052>
- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. (2013). Random-effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of Memory and Language*, 68(3), 255-278. <https://doi.org/10.1016/j.jml.2012.11.001>
- Basnight-Brown, D. M., Chen, L., Hua, S., Kostić, A., & Feldman, L. B. (2007). Monolingual and bilingual recognition of regular and irregular English verbs: Sensitivity to form similarity varies with first language experience. *Journal of Memory and Language*, 57(1), 65-80. <https://doi.org/10.1016/j.jml.2007.03.001>
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67(1), 1-48. <https://doi.org/10.18637/jss.v067.i01>
- Beyersmann, E., Ziegler, J. C., Castles, A., Coltheart, M., Kezilas, Y., & Grainger, J. (2016). Morpho-orthographic segmentation without semantics. *Psychonomic Bulletin and Review*, 23, 533-539. <https://doi.org/10.3758/s13423-015-0927-z>
- Booij, G. (2005). Compounding and derivation: Evidence for construction morphology. In W. U. Dressler, D. Kastovsky, O. E. Pfeiffer & F. Rainer (Eds.), *Morphology and its demarcations* (pp. 109-132). John Benjamins Publishing. <https://doi.org/10.1075/cilt.264>
- Boudelaa, S., & Marslen-Wilson, W. D. (2005). Discontinuous morphology in time: Incremental masked priming in Arabic. *Language and Cognitive Processes*, 20(1-2), 207-260. <https://doi.org/10.1080/01690960444000106>
- Butterworth, B. (1983). Lexical representation. In B. Butterworth (Ed.), *Language production: Development, writing and other language processes* (pp. 257-294). Academic Press.
- Caramazza, A., Laudanna, A., & Romani, C. (1988). Lexical access and inflectional morphology. *Cognition*, 28(3), 287-332. [https://doi.org/10.1016/0010-0277\(88\)90017-0](https://doi.org/10.1016/0010-0277(88)90017-0)
- Clahsen, H., & Felser, C. (2006). Grammatical processing in language learners. *Applied Psycholinguistics*, 27(1), 3-42. <https://doi.org/10.1017/S0142716406060024>
- Clahsen, H., Felser, C., Neubauer, K., Sato, M., & Silva, R. (2010). Morphological structure in native and nonnative language processing. *Language Learning*, 60(1), 21-43. <https://doi.org/10.1111/j.1467-9922.2009.00550.x>

- Çağlar, O. C. (2019). *The effects of cross-morphemic letter transpositions on morphological processing in Turkish: A psycholinguistic investigation*. [Unpublished master's thesis]. Middle East Technical University.
- Çotuksöken, Y. (2011). *Yapı ve işlevlerine göre Türkiye Türkçesi'nin ekleri*. Papatya Bilim Yayınevi.
- D'Alessandro, R., Putnam, M. T., & Terenghi, S. (2025). *Heritage languages and syntactic theory*. Oxford University Press. <https://doi.org/10.1093/9780191987731.001.0001>
- Diependaele, K., Sandra, & D., Grainger, J. (2005). Masked cross-modal morphological priming: Unravelling morpho-orthographic and morpho-semantic influences in early word recognition. *Language and Cognitive Processes*, 20(1/2), 75-114. <https://doi.org/10.1080/01690960444000197>
- Diependaele, K., Sandra, & D., Grainger, J. (2009). Semantic transparency and masked morphological priming: The case of prefixed words. *Memory and Cognition*, 37(6), 895-908. <https://doi.org/10.3758/MC.37.6.895>
- Diependaele, K., Duñabeitia, J. A., Morris, J., & Keuleers, E. (2011). Fast morphological effects in first and second language word recognition. *Journal of Memory and Language*, 64(4), 344-358. <https://doi.org/10.1016/j.jml.2011.01.003>
- Eldem-Tunç, E., & Fuchs, Z. (2024). Morphological decomposition in heritage Turkish in the US. *Proceedings of the Workshop on Turkic and Languages in Contact with Turkic*, 9, 89-105. <https://doi.org/10.3765/v73aqx67>
- Feldman, L. B., O'Connor, P. A., & del Prado Martín, F. M. (2009). Early morphological processing is morphosemantic and not simply morpho-orthographic: A violation of form-then-meaning accounts of word recognition. *Psychonomic Bulletin and Review*, 16(4), 684-691. <https://doi.org/10.3758/PBR.16.4.684>
- Feldman, L. B., Kostić, A., Gvozdenović, V., O'Connor, P. A., & del Prado Martín, F. M. (2012). Semantic similarity influences early morphological priming in Serbian: A challenge to form-then-meaning accounts of word recognition. *Psychonomic Bulletin and Review*, 19(4), 668-676. <https://doi.org/10.3758/s13423-012-0250-x>
- Feldman, L. B., Milin, P., Cho, K. W., del Prado Martín, F. M., & O'Connor, P. A. (2015). Must analysis of meaning follow analysis of form? A time course analysis. *Frontiers in Human Neuroscience*, 9, 1-19. <https://doi.org/10.3389/fnhum.2015.00111>
- Felser, C., & Uygun, S. (2022). Optional plural agreement in heritage Turkish speakers' verb form choices. *Heritage Language Journal*, 19(1), 1-30. <https://doi.org/10.1163/15507076-bja10004>
- Fleischhauer, E., Bruns, G., & Grosche, M. (2021). Morphological decomposition supports word recognition in primary school children learning to read: Evidence from masked priming of German derived words. *Journal of Research in Reading*, 44(1), 90-109. <https://doi.org/10.1111/1467-9817.12340>

- Fridman, C., & Meir, N. (2023). Lexical production and innovation in child and adult Russian heritage speakers dominant in English and Hebrew. *Bilingualism: Language and Cognition*, 26(5), 880-895. <https://doi.org/10.1017/S1366728923000147>
- Frost, R., Forster, K. I., & Deutsch, A. (1997). What can we learn from the morphology of Hebrew? A masked priming investigation of morphological representation. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 23(4), 829-856. <https://doi.org/10.1037/0278-7393.23.4.829>
- Hasenäcker, J. (2016). *Learning to read complex words: Morphological processing in reading acquisition*. [Doctoral dissertation, Freie Universität Berlin]. Refubium. <https://refubium.fu-berlin.de/handle/fub188/13452>
- Heyer, V., & Kornishova, D. (2018). Semantic transparency affects morphological priming . . . eventually. *Quarterly Journal of Experimental Psychology*, 71(5), 1112-1124. <https://doi.org/10.1080/17470218.2017.1310>
- Isurin, L., & Wilson, H. (2022). First language attrition in bilingual immigrants. In G. L. Schiewer, J. Altarriba, & B. C. Ng (Eds.), *Handbook of language and emotion* (Vol. 2, pp. 1320-1339). De Gruyter Mouton. <https://doi.org/10.1515/9783110670851-031>
- Jacob, G., & Kırkıcı, B. (2016). The processing of morphologically complex words in a specific speaker group: A masked priming study with Turkish heritage speakers. *The Mental Lexicon*, 11(2), 308-328. <https://doi.org/10.1075/ml.11.2.06jac>
- Jacob, G., Şafak, D. F., Demir, O., & Kırkıcı, B. (2019). Preserved morphological processing in heritage speakers: A masked priming study on Turkish. *Second Language Research*, 35(2), 173-194. <https://doi.org/10.1177/0267658318764535>
- Kasparian, K., & Steinhauer, K. (2017). When the second language takes the lead: Neurocognitive processing changes in the first language of adult attriters. *Frontiers in Psychology*, 8, 389. <https://doi.org/10.3389/fpsyg.2017.00389>
- Kırkıcı, B., & Clahsen, H. (2013). Inflection and derivation in native and nonnative language processing: Masked priming experiments on Turkish. *Bilingualism: Language and Cognition*, 16(4), 776-791. <https://doi.org/10.1017/S1366728912000648>
- Kuznetsova, A., Bruun Brockhoff, P., & Haubo Bojesen Christensen, R. (2014). lmerTest: Tests for random and fixed effects for linear mixed effect models (lmer objects of lme4 package). R package version 2.0-11. <https://cran.r-project.org/web/packages/lmerTest/index.html>
- Lázaro, M., Illera, V., & Sainz, J. (2016). The suffix priming effect: Further evidence for an early morpho-orthographic segmentation process independent of its semantic content. *Quarterly Journal of Experimental Psychology*, 69(1), 197-208. <https://doi.org/10.1080/17470218.2015.1031146>

- Lázaro, M., García, L., & Illera, V. (2021). Morpho-orthographic segmentation of opaque and transparent derived words: New evidence for Spanish. *Quarterly Journal of Experimental Psychology*, 74(5), 944-954. <https://doi.org/10.1177/1747021820977038>
- Libben, G., Gibson, M., Yoon, Y. B., & Sandra, D. (2003). Compound fracture: The role of semantic transparency and morphological headedness. *Brain and language*, 84(1), 50-64. [https://doi.org/10.1016/s0093-934x\(02\)00520-5](https://doi.org/10.1016/s0093-934x(02)00520-5)
- Longtin, C. M., & Meunier, F. (2005). Morphological decomposition in early visual word processing. *Journal of Memory and Language*, 53(1), 26-41. <https://doi.org/10.1016/j.jml.2005.02.008>
- Lõo, K., Toth, A., Karaca, F., & Järvikivi, J. (2022). [Morphological processing is gradient not discrete in L1 and L2 English masked priming](https://doi.org/10.1075/ml.21008.loo). *The Mental Lexicon*, 17(1), 76-103. <https://doi.org/10.1075/ml.21008.loo>
- Ma, Y., & Vanek, N. (2024). First language lexical attrition in a first language setting: A multi-measure approach testing teachers of English. *Journal of Psycholinguistic Research*, 53(2), 23. <https://doi.org/10.1007/s10936-024-10068-7>
- Marelli, M., Amenta, S., Morone, E. A., & Crepaldi, D. (2013). Meaning is in the beholder's eye: Morpho-semantic effects in masked priming. *Psychonomic Bulletin and Review*, 20, 534-541. <https://doi.org/10.3758/s13423-012-0363-2>
- Montrul, S. (2008). *Incomplete acquisition in bilingualism: Re-examining the age factor*. John Benjamins Publishing. <https://doi.org/10.1075/sibil.39>
- Montrul, S. (2013). Bilingualism and the heritage language speaker. In W. Ritchie & T. Bhatia (Eds.), *The handbook of bilingualism* (pp. 163-189). Wiley-Blackwell. <https://doi.org/10.1002/9781118332382.ch7>
- Nagle, C., Baese-Berk, M. M., Diantoro, C., & Kim, H. (2023). How good does this sound? Examining listeners' second language proficiency and their perception of category goodness in their native language. *Languages*, 8(1), 43. <https://doi.org/10.3390/languages8010043>
- Oğuz, E., & Kırkıcı, B. (2023). The processing of morphologically complex words by developing readers of Turkish: A masked priming study. *Reading and Writing*, 36(8), 2053-2080. <https://doi.org/10.1007/s11145-022-10377-0>
- Polinsky, M. (2018). *Heritage languages and their speakers*. Cambridge University Press. <https://doi.org/10.1017/9781107252349>
- Polinsky, M., & Scontras, G. (2020). Understanding heritage languages. *Bilingualism: Language and Cognition*, 23(1), 4-20. <https://doi.org/10.1017/S1366728919000245>
- Polinsky, M., & Putnam, M. T. (2024). *Formal approaches to complexity in heritage language*. (Current Issues in Bilingualism 3). Language Science Press. <https://doi.org/10.5281/zenodo.12073160>
- Poort, E. D. (2018). *The representation of cognates and interlingual homographs in the bilingual lexicon*. [Doctoral dissertation, University College London]. UCL Discovery.

- <https://discovery.ucl.ac.uk/id/eprint/10064100/>
- Portin, M., Lehtonen, M., Harrer, G., Wande, E., Niemi, J., & Laine, M. (2008). L1 effects on the processing of inflected nouns in L2. *Acta Psychologica*, 128(3), 452-465.
<https://doi.org/10.1016/j.actpsy.2007.07.003>
- R Core Team. (2021). R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria.
<https://www.r-project.org/>
- Rastle, K., Davis, M. H., & New, B. (2004). The broth in my brother's brothel: Morpho-orthographic segmentation in visual word recognition. *Psychonomic Bulletin and Review*, 11(6), 1090-1098.
<https://doi.org/10.3758/BF03196742>
- Rastle, K., Davis, M. H., Marslen-Wilson, W., & Tyler, L. K. (2000). Morphological and semantic effects in visual word recognition: A time-course study. *Language and Cognitive Processes*, 15(4/5), 507-537.
<https://doi.org/10.1080/01690960050119689>
- Schmid, M. S. (2010). Languages at play: The relevance of L1 attrition to the study of bilingualism. *Bilingualism: Language and Cognition*, 13(1), 1-7. <https://doi.org/10.1017/S1366728909990368>
- Schmid, M. S., & Köpke, B. (2017). The relevance of first language attrition to theories of bilingual development. *Linguistic Approaches to Bilingualism*, 7(6), 637-667. <https://doi.org/10.1075/lab.17058.sch>
- Schmid, M. S. & Yılmaz, G. (2021). Lexical access in L1 attrition—competition versus frequency: A comparison of Turkish and Moroccan attriters in the Netherlands. *Applied Linguistics*, 42(5), 878-904.
<https://doi.org/10.1093/applin/amab006>
- Schreuder, R., & Baayen, R. H. (1995). Modeling morphological processing. In: Feldman, L. B. (Ed.), *Morphological aspects of language processing* (pp. 131-154). Erlbaum.
- Scontras, G., Fuchs, Z., & Polinsky, M. (2015). Heritage language and linguistic theory. *Frontiers in Psychology*, 6:1545.
<https://doi.org/10.3389/fpsyg.2015.01545>
- Scontras, G., Polinsky, M., & Fuchs, Z. (2018). In support of representational economy: Agreement in heritage Spanish. *Glossa: A Journal of General Linguistics*, 3(1), 1-29. <https://doi.org/10.5334/gjgl.164>
- Sezer, T. (2017). TS corpus project: An online Turkish dictionary and TS DIY corpus. *European Journal of Language and Literature*, 3(3), 18-24.
<https://doi.org/10.26417/ejls.v9i1.p18-24>
- Shin, G.-H. (2024). Good-enough processing, home language proficiency, cognitive skills, and task effects for Korean heritage speakers' sentence comprehension. *Frontiers in Psychology*, 15, 1382668.
<https://doi.org/10.3389/fpsyg.2024.1382668>
- Taft, M. (1979). Recognition of affixed words and the word frequency effects. *Memory and Cognition*, 7(4), 263-272.
<https://doi.org/10.3758/BF03197599>

- Taft, M. (2004). Morphological decomposition and the reverse base frequency effect. *Quarterly Journal of Experimental Psychology*, 57(4), 745-765. <https://doi.org/10.1080/02724980343000477>
- Taft, M., & Forster, K. I. (1975). Lexical storage and retrieval of prefixed words. *Journal of Verbal Learning and Verbal Behavior*, 14(6), 638-647. [https://doi.org/10.1016/S0022-5371\(75\)80051-X](https://doi.org/10.1016/S0022-5371(75)80051-X)
- ten Hacken, P. (1994). *Defining morphology: A principled approach to determining the boundaries of compounding, derivation, and inflection*. [Doctoral dissertation, Universität Hildesheim].
- Uygun, S., & Gürel, A. (2020). Does the processing of first language compounds change in late bilinguals? In M. Schlechtweg (Ed.), *The learnability of complex constructions: A cross-linguistic perspective* (pp. 63-89). De Gruyter Mouton. <https://doi.org/10.1515/9783110695113-004>
- Uygun, S., & Clahsen, H. (2021). Morphological processing in heritage speakers: A masked priming study on the Turkish aorist. *Bilingualism: Language and Cognition*, 24(3), 415-426. <https://doi.org/10.1017/S1366728920000577>
- van Hell, J. G., & Dijkstra, T. (2002). Foreign language knowledge can influence native language performance in exclusively native contexts. *Psychonomic Bulletin & Review*, 9(4), 780-789. <https://doi.org/10.3758/BF03196335>
- Venables, W. N., & Ripley, B. D. (2002). *Modern applied statistics with S*. Springer. <https://doi.org/10.1007/978-0-387-21706-2>
- Viviani, E., & Crepaldi, D. (2022). Masked morphological priming and sensitivity to the statistical structure of form-to-meaning mapping in L2. *Journal of Cognition*, 5(1), 30. <https://doi.org/10.5334/joc.221>
- Zhang, Q., Liang, L., Yao, P., Hu, S., & Chen, B. (2017). Parallel morpho-orthographic and morpho-semantic activation in processing second language morphologically complex words: Evidence from Chinese-English bilinguals. *International Journal of Bilingualism*, 21(3), 291-305. <https://doi.org/10.1177/1367006915624249>

UNDERSTANDING HUMOR IN TEMEL JOKES: INSIGHTS FROM CONVERSATION ANALYSIS, CONVERSATIONAL MAXIMS, AND SPEECH ACT THEORY

Abdullah TOPRAKSOY

Dr., Istanbul University, Department of Linguistics, abdullah.topraksoy@istanbul.edu.tr, ORCID: 0000-0003-3240-3915

Topraksoy, A. (2025). Understanding humor in Temel jokes: Insights from conversation analysis, conversational maxims, and speech act theory. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 335–349). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X018>

ABSTRACT

What we find pleasurable and funny, or in its broader sense, what we laugh at can often be regarded as ‘humor’. Therefore, humor represents one of the main aspects of everyday interactions. Consequently, the issues like the reasons of laughter and what constitutes humor have been questioned since Plato and Aristotle. One reason that humor has come out as a theatrical form is the pleasure people take in laughing. Indeed, one of the most significant reflectors of humor in language is jokes. This study aims at analyzing and describing the linguistic features of humor in Temel Jokes, a cultural figure in Anatolia, regarding Conversation Analysis, Conversational Maxims and Speech Act Theory. 20 randomly chosen jokes of Temel were determined as the database of the study. In terms of Conversation Analysis, these jokes were studied in terms of turn taking and adjacency pairs; in terms of Speech Act analysis, Austin’s performative speech acts and the framework of Finegan and Besnier based on Searle’s study were made use of. The results showed that turn-by-turn allocation of speeches are dominant in Temel jokes and the humor in these jokes can be said to be created mainly by means of the violation of maxims and less likely by obeying to the maxims. Moreover, answers to the questions; assertions; statements and descriptions were significantly observed in Temel jokes and together with illocutionary acts that are mostly seen in the jokes, these acts also play a significant role in creating humor.

Keywords: Humor language, Temel jokes, Conversation analysis, Conversational maxims, Speech act theory

ÖZ

Keyifli ve komik bulduğumuz ya da daha geniş anlamıyla güldüğümüz şeyler genellikle ‘mizah’ olarak kabul edilebilir. Bu nedenle mizah, günlük etkileşimlerin ana

yönlerinden birini temsil eder. Sonuç olarak, gülmenin nedenleri ve mizahı neyin oluşturduğu gibi konular Platon ve Aristoteles'ten bu yana sorgulanmaktadır. Mizahın teatral bir form olarak ortaya çıkmasının bir nedeni de insanların gülmekten aldıkları hazdır. Nitekim mizahın dildeki en önemli yansıtıcılarından biri de fıkralardır. Bu çalışma, Anadolu'da kültürel bir figür olan Temel Fıkralarındaki mizahın dilsel özelliklerini Konuşma Analizi, Konuşma Maksimleri ve Söz Edimi Kuramı açısından incelemeyi ve betimlemeyi amaçlamaktadır. Temel'in rastgele seçilen 20 fıkrası çalışmanın veri tabanı olarak belirlenmiştir. Konuşma Analizi açısından bu fıkralar sıra alma ve bitişiklik çiftleri açısından incelenmiş; Söz Edimi analizi açısından Austin'in edimsel söz edimleri ve Searle'ün çalışmasına dayanan Finegan ve Besnier'in çerçevesinden yararlanılmıştır. Sonuçlar, Temel fıkralarında konuşmaların sırayla dağılımının baskın olduğunu göstermiştir ve bu fıkralardaki mizahın büyük ölçüde özdeyişlerin ihlali yoluyla, daha az olasılıkla da özdeyişlere uyma yoluyla yaratıldığı söylenebilir. Ayrıca, Temel fıkralarında sorulara verilen cevaplar; iddialar; ifadeler ve betimlemeler önemli ölçüde gözlemlenmiştir ve fıkralarda çoğunlukla görülen söz edimleriyle birlikte bu edimler de mizah yaratmada önemli bir rol oynamaktadır.

Anahtar Sözcükler: Mizah dili, Temel fıkraları, Konuşma analizi, Konuşma maksimleri, Söz-Eylem Kuramı

1. Introduction

Individuals in a society communicate with one another via language, and humor can be regarded as a means of communication since the humorous meaning is transferred to the people through language. "The tendency of particular cognitive experiences to provoke laughter and provide amusement" is the definition of humor (Humor, 2015). Moreover, it is believed that humor does not occur on a single individual itself. To support this idea, Viktoroff (1953: 14) asserts that "one never laughs alone; laughter is always the laughter of a particular social group, and if one does not share the group's norms, feelings, and ideas, or if one is not part of it, one cannot associate oneself with it."

Jokes are significant indicators of humors. According to Berger (1976), humor in jokes is a particular kind of communication that creates an incongruent relationship or imaginative meaning and is delivered in a way that elicits laughter. A good many studies on the concept of humor in jokes have been carried out. However, studies on Turkish texts regarding humor in Turkish jokes have mostly been concentrated upon the aspects of literary language, and for this reason, the studies on different kinds of texts including the language of humor in terms of a linguistic perspective are relatively limited. This limitedness of linguistic studies on Turkish jokes in view of humor has been taken as an inspiration while beginning this study. In terms of conversation analysis, turn taking and adjacency pairs are two significant concepts for the present study. Turn taking refers both to the turns of different speakers in a conversation and what is said in each turn of the speakers. Put differently, turn-taking refers to the mechanism via which participants divide up the responsibility or right to engage in an

interactive activity (Jefferson, Schegloff, and Sacks, 1974). Adjacency pairs include the speeches of the speakers and are useful in determining the relationship between the speeches of different speakers in a conversation. This happens when a speaker's statement increases the likelihood of a specific form of response. Question-answer exchanges are the adjacency pairs that are used in conversations the most. Adjacency pairs have strong inherent expectations as a recognized component of conversational structure: inquiries are addressed, declarations are acknowledged, complaints are addressed, and greetings are reciprocated (Pridham 2001: 27). There are pre-allocated systems and turn-by-turn allocation in adjacency pairs. Pre-allocated systems in a conversation refer to long sentence structures uttered only one speaker at a single turn whereas turn-by-turn allocation reflects short and sequential sentence structures in which each person speaks respectively. In regard to conversational maxims, Grice, who introduced the Cooperative Principle, which is constituted by the maxims of conversation, states "Make your conversational contribution such as is required, at the stage at which it occurs, by the accepted purpose or direction of the talk exchange in which you are engaged." (Grice 1975: 45). The cooperative principle can be divided into four maxims. These maxims are those of quality, relation, manner, and quantity, being related to "the speaker's commitment to truth, relevance, clarity, and to providing the right quantity of information at any given time" (Attardo 1994:274).

As for speech acts, the classification of speech acts by Finegan and Besnier (1989) based on Searle's speech acts and Austin's (1962) locutionary, illocutionary and perlocutionary acts are greatly significant concepts to be mentioned for the current study. Finegan and Besnier (1989: 304) define speech acts as acts performed by means of language. They classified speech acts as follows: *representatives, commissives, directives, declarations, expressives and verdictives*. Affirmations, declarations, claims, theories, explanations, and recommendations are all mentioned by representatives. Generally speaking, they can be classified as true or untrue. Speech acts like vows, threats, pledges, and promises that bind the speaker to a certain course of action are known as commissives. The goal of directives is to elicit responses from the addressee in the form of requests, challenges, invitations, commands, and so on. The conditions that declarations describe—blessings, hirings, firings, marriages, and mistrial declarations—come to pass. Expressives like "hello," "sorry," "congratulations," and "thank you" reveal the speaker's mental condition or attitude. Verdictives render evaluations or conclusions by rating, evaluating, appraising, and etc. Austin divides speech acts into three categories, referring to them as "performatives": The term "locutionary act" refers to the act of performing or saying something; it includes the actual utterance and its apparent meaning, which includes phonetic, phatic, and rhetic acts that correspond to the verbal, syntactic, and semantic aspects of any meaningful utterance; illocutionary acts are the pragmatic "illocutionary force" of the utterance, meaning that its intended

significance is a socially valid verbal action; and perlocutionary acts are the real effects and outcomes of illocutionary acts on people, such as persuading, convincing, frightening, enlightening, inspiring, or in some other way getting someone to do or realize something, whether on purpose or not.

The present study aims to analyze and describe the linguistic features of humor language in Temel Jokes regarding turn-taking and adjacency pairs within Conversation Analysis which emerged under the field of Discourse Analysis; Conversational Maxims developed by Grice (1975) as constituting The Cooperative Principle; and with regard to Speech Act Theory initiated by Austin (posthumously in 1962) and later developed by Searle (1969).

1.1 Significance of the study and research questions

Although a good many studies on the concept of humor in jokes have been carried out, studies on Turkish texts regarding humor in Turkish jokes have mostly been concentrated upon the aspects of literary language, and for this reason, the studies on different kinds of texts including the language of humor in terms of a linguistic perspective are relatively limited. This lack of linguistic studies on Turkish jokes in view of humor has been taken as an inspiration while beginning this study.

Temel jokes in the current study were analyzed within the scope of these research questions:

1. What are the conversational features in Temel jokes in respect to turn-taking and adjacency pairs?
2. Do conversational maxims play a significant role in creating humour in Temel jokes? If so, in what ways?
3. What features can be seen in Temel jokes in regard to speech acts?

2. Review of Literature on Humor

In its many forms, humor seems to be one of the most delineative characteristics of humans. Although humor is commonplace in everyday life, it is thought by some people to be an elusive and indefinite theoretical concept. Lafollette and Shanks (1993) point that humor is a pervasive feature of human life, the nature of which is elusive, and they add that it has generated little theoretical interest among the scholars. Nonetheless, this has not been seen by the scholars of various disciplines as an obstacle in studying humor, which has resulted in “epistemological hairsplitting” rather than clarifying the issue (Attardo 1994:1). Scholars have been studying humor in many fields of research, like psychology, philosophy, linguistics, sociology, and literature. In sooth, studies on humor began to increase especially in the second half of the 20th century and Davies (1987), Raskin (1987), Apte (1988), and Özünlü (1991) carried out studies emphasizing semantic, pragmatic, social and cultural dimensions of

humorous items. According to Eastman (as cited in Raskin 1985: 19), there are Ten Commandments of humor:

a) Be interesting, b) Be unimpassioned, c) Be effortless, d) Remember the difference between cracking practical jokes and conveying ludicrous impressions, e) Be plausible, f) Be sudden, g) Be neat, h) Be right with your timing, i) Give good measure of serious satisfaction, j) Redeem all serious disappointments

As Attardo (1994:4) points out, in the field of literary criticism, for example, that there is a need for a fine-grained categorization, whereas linguists have often been happy with broader definitions, arguing that whatever evokes laughter or is felt to be funny is humor. At the turn of the 20th century, Grieg (1923) was able to list a total of eighty-eight separate theories of humor but many of these theories had borrowed heavily from one another, and they differ only in details. Keith-Spiegel (1972) classifies various humor theories into eight categories that can be summarized as: Superiority Theory, Biological Instinct and Evolution Theory, Incongruity Theory, Surprise Theory, Ambivalence Theory, Release and Relief Theory, Configurational Theory and Psycho-analytic Theory. Sometime later, Raskin (1985: 31- 36) divides the theories of humor into three categories: incongruity theories, hostility theories and release theories. For the sake of relevance, they will not be explained in this paper. The following section clarifies the linguistic perspective on humor theorization.

2.1 Linguistic theorization of humor

It is possible to divide humor into two categories: nonverbal and verbal humor. The former contains forms such as cartoons, gestures, facial expressions, and clumsiness observed in both humans and animals whereas the latter includes canned jokes, conversational jokes, witticisms, puns and every day or media discourses, both written and spoken. Such a diversified and broad field cannot be narrowed down to a discipline, and it is nearly impossible to come up with an all-embracing theory of humor. However, if one talks about verbal humor, it becomes a must to apply a linguistic insight to the field. Later, some questions such as given below have been listed by Ermida (2008: 41) as the questions that should be dealt with by linguists studying humor: “What is the language of humor? Is there one, to begin with? In what way, and to what extent, is humor conveyed linguistically? Does it conform to common language usage, or does it subvert its norms and conventions? What linguistic mechanisms do comedians exploit? Are there specific linguistic resources available to humorous use? If so, how do they work?”.

From above mentioned theories, linguistics has sided with incongruity theories since these theories are essentialist, i.e., the attempt to pinpoint what makes humor funny (Attardo, 1994).

Linguistic approaches to verbal humor have opened the path for more systematic and falsifiable humor theories and “the period from the 1980s to the present has witnessed sustained and progressively heightened

interest among language scholars in linguistically based accounts of verbal humor” (Simpson, 2003: 16). Since then, linguists have contributed to the field with their renewed theories that offers new hypotheses about the cognitive processes of humor. Raskin’s (1985) *Semantic Script Theory of Humor* and Attardo and Raskin’s (1991) *General Theory of Verbal Humor* are still considered as the leading and the most influential theories in linguistics.

2.1.1 Semantic script theory of humor (SSTH)

In Ritchie’s (2004: 69) words, “the Semantic Script Theory of Humor (hereafter, the SSTH) is one of the few attempts to approach verbally expressed humor in a systematic and theoretical fashion, and as such is to be welcomed”. General aim of the SSTH is to form information – processing system which is able to account for the humorousness of a text. For the SSTH, the central aspect of verbal humor was semantic/pragmatic, and it explains the funniness of a joke with scripts or frames defined as “an organized complex of information about some entity, in the broadest sense: an object (real or imaginary), an event, an action, a quality, etc.” (Attardo, 2001: 2).

2.1.2 General theory of verbal humor (GTVH)

Attardo and Raskin (1991) revised the SSTH and extended it into GTVH. GTVH has been further extended by Attardo (1997, 2001) to analyze not only jokes but also narratives, stand-up or longer humorous texts. Attardo (2001: 22) claims that “whereas the SSTH was a semantic theory of humor, the GTVH is a linguistic theory at large – that is, it includes other areas of linguistics as well, including, most notably, textual linguistics, the theory of narrativity, and pragmatics broadly conceived”. The extension has been made with six “knowledge resources” (KRs in short) which are: the script opposition (SO) - known from the SSTH, the logical mechanism (LM), the target (TA), the narrative strategy (NS), the language (LA) and the situation (SI).

For more detailed explanations and some logical mechanisms and examples in jokes, see Attardo (1997) and Attardo et al. (2002).

According to Raskin (1985), the text of a joke is clear up to the punchline. The punchline forces a change in script and opens the hearer's eyes to the fact that there are multiple ways to understand the text from the outset. Furthermore, covered by Raskin (ibid) is the concept of the "non-bona-fide mode of communication." He claims that the manner that non-bona-fide communication deviates from bona-fide communication—which is sincere, serious, and information-conveying—is by breaking one or more of the four conversational maxims of Grice's (1975) Cooperative Principle. Moreover, Raskin (1985:100) notes that the speaker may violate these maxims intentionally or inadvertently; in the former scenario, the speaker is conscious of the semantic ambiguity s/he has generated, but in the latter scenario, the speaker is unaware of it. Thus, even if the speaker in the later

instance is sincere, the hearer will see the statement as being dishonest, meaning they will pick up on the ambiguity.

There is a study carried out by Karahan (1997) which is similar to the present study in terms of the features analyzed. In that study, she analyzed Nasreddin Hodja jokes within the viewpoint of conversation analysis and speech acts. She found out that in Nasreddin Hodja's speech turns in the text of jokes, there are pre-allocated systems of turn taking since there is no pressure on Hodja while talking to others and there is no one to interrupt his speech; whereas the other people, whom Hodja talks to, feel pressure on themselves due to Hodja and thus their speech is relatively short and this refers to turn-by- turn allocation. In terms of adjacency pairs, the speech goes on in turns with different people's talking. An important point here is that Hodja does not permit others to talk a second time in a row. According to the classification of Finegan ve Besnier (1989); verdictives, directives and representatives have been found in the Hodja's jokes in regard to speech acts. In addition, illocutionary and perlocutionary acts have mostly been observed. Regarding the conversational maxims, Hodja sometimes obey the maxims when he tells the truth on the jokes and sometimes does violate one or more maxims intentionally so as to make the readers think on the jokes by means of implicature and later find the humor themselves.

In another study, Uçar (2014) studied the puns collected from the specialized authentic corpus which is built from the transcriptions of the video recordings of the television show *Komedi Dükkanı* 'Comedy Shop'. Puns can be defined shortly as the use of words in a way that suggests two or more meanings which are usually incompatible and thus the creation of humorous effect. The instances used in *Komedi Dükkanı* 'Comedy Shop' were categorized and from the analysis, Uçar determined that these puns have a different structure from other humor genres like stand-up and sitcoms and thereby puns can be defined as a new humor genre.

3. Method

3.1 Sample

To test the questions of the study, lots of Temel jokes were found. However, although humor is created through linguistic features in some jokes, it results from adult jokes in some others. Thus, with the aim of censoring unwanted meanings, this latter kind of jokes were excluded from the scope of the study and 20 randomly chosen jokes of Temel were determined as the database of the study. They were gathered from various websites since it is easier to get them from websites given in the references section.

3.2 Data collection and analysis

The framework of analysis in Temel jokes was adapted from General Theory of Verbal Humor (GTVH) based on three different levels.

In regard to Conversation Analysis, these jokes were studied in terms of turn taking which determines the turns of different characters in the jokes and adjacency pairs that include the speeches of these characters and their relationships and later conversational maxims in Temel jokes were analyzed in order to find out their importance in creating humor.

Finally, speech act analysis was carried out on Temel jokes. In this part, first, the framework of the classification of speech acts by Finegan and Besnier (1989) based on Searle's study was employed and then Austin's performative speech acts were investigated in the jokes.

Later on, the findings were described, and the results gathered from the analysis were based on descriptive statistics and given in tables with relevant numbers of instances and their frequencies.

It is crucial to note that this study adopts a text-based methodology, with a specific focus on delineating and interpreting linguistic features within Temel jokes. The intent is not to generalize humor language but to offer a detailed account of linguistic nuances specific to this particular genre. The analysis and numerical data were processed without the use of any software tool.

For convenience and clarity before entering the analysis of jokes, a short description of characters in Temel jokes is given: the main character in these jokes is mainly Temel. He is from Rize in some jokes and from Trabzon in some others. Dursun and/or İdris are secondary characters. Fadime is either Temel's wife or his lover. All the characters, primarily Temel, generally perform illogical or absurd behaviors in jokes to create humor and to make the readers laugh.

Notably, the textual scripts of the jokes are not included in this paper for brevity. Interested readers can access both the English and original scripts, each numerically aligned with the linguistic analysis section, through the provided link below:

<https://docs.google.com/document/d/1gQ325uls50CXuRXEru3xl vVT0KFMnw9c/edit>

4. Results

As the analyses were text based relevant to each joke and each category, it is a long file. Therefore, the linguistic analyses carried out can be reached at the link provided below:

https://docs.google.com/document/d/1jEORkiz4jFMsjBQ2F12qKoZMIdHTYAzVs7vHf_M2VeI/edit?usp=share_link

The results of the study are described in tables and explanations for tables are given beneath them. Table 1 shows the numbers related to turn-taking and adjacency pairs. Table 2 demonstrates conversational maxims occurred in the jokes. Table 3 shows speech acts based on Finegar and Besnier's classification and Table 4 reflects Austin's speech acts.

Table 1. Uses of turn takings and adjacency pairs in Temel jokes

Categories	Instances in the jokes (n=)	Percentage
Turn-by-turn allocation	18	90%
Pre-allocated systems	1	5%
Uncategorized	1	5%
TOTAL	20	100%

Out of 20 jokes, there is a huge dominance of turn-by-turn allocation of speeches over the other categories as seen in Table 1. This means that the talking-turns of the characters are sequential and that Temel has no pressure on the other characters in these jokes. There is one instance of a ‘pre-allocated system’ in the study and it results from the announce in the plane. The content of the announce is pre- determined and thus it can be regarded as a pre-allocated system. In addition, there is another instance which is ‘uncategorized’. This group occurs since there is no conversation in one of the jokes analyzed. Temel talks in a monolog, not in a dialog.

Table 2. Violations of maxims

Category of maxims	Violation of maxims(n=)	Percentage
Relevance	7	30,43%
Quantity	5	21,73%
Quality	2	8,69%
Manner	1	4,34%
No violation	8	34,78%
TOTAL	23	100%

There are 23 instances of maxims observed in total in the jokes. The humor in these jokes is created sometimes with the help of the violation of maxims and sometimes of obeying the maxims. Violation of maxims (n=15, 65,22% out of 23 total) is the one that is most observed in the jokes. This means that Temel generally gives responses to other characters but in an irrelevant (n=7, 30,43% out of 23 total) way and by doing so, the humor is created. Violation of quantity (n=5, 21,73% out of 23 total) follows it and humor is said to have been created in some instances through the violation of the maxim of quantity. However, there are also instances in which humor is created by means of obeying the maxims. There are 8 instances (34,78% out of 23 total) in this category.

Table 3. Speech acts in Temel jokes based on Finegar and Besnier' s classification

Categories	Instances(n=)	Percentage
Representatives	20	66,66%
Verdictives	6	20%
Directives	4	13,33%
Declarations	0	0%
Expressives	0	0%
Commissives	0	0%
TOTAL	30	100%

Out of 30 instances of speech acts given in Table 3, the most observed category is 'representatives' with 20 instances (66,66%) in the jokes. 'Verdictives' follow 'representatives' with 6 instances (20%) and lastly 'directives' (13,33%) have been observed with 4 instances. The interpretation of this result is that there are many answers given to the questions in the jokes, and assertions, statements, claims and descriptions are generally found in Temel jokes. Furthermore, the presence of verdictives in the jokes show that some assessments or judgments are made in the conversations of these jokes. Finally, the instances of directives in the jokes represent that commands and requests are seen during conversation in some jokes. Commissives, declarations and expressives have not been observed and this means that, in terms of humor in Temel jokes, there is no place for them since humor in these jokes is not based on these categories.

Table 4. Austin's speech acts in Temel jokes

Categories	Instances(n=)	Percentage
Illocutionary acts	26	43,33%
Locutionary acts	21	35%
Perlocutionary acts	13	21,66%
TOTAL	60	100%

It is clear from the Table 4 that illocutionary acts (n=26, 43,33% out of 60 total) are the mostly observed speech acts. This means that illocutionary acts in the jokes represent that there are questions and requests mostly used in the conversation of the jokes and humor is created mostly through questions and requests. In the second place, locutionary acts (n=21, 35% out of 60 total) follow illocutionary acts and this reflects that in many parts of the jokes, humor is behind the statements uttered by the characters. Finally, the small presence of perlocutionary acts (n=13, 21,66% out of 60 total) refers to the instances in which characters try to persuade or otherwise get someone to do or realize something in the jokes. The humor is created in few numbers of jokes by doing so.

5. Discussion

In regard to turn-taking and adjacency pairs, turn-by-turn allocation of speeches are dominant in Temel jokes since the main character Temel does not have any pressure on the other characters and the sequences of turn-taking and duration of talking are more or less similar for any character in the jokes. In addition, the humor in Temel jokes can be said to have been created sometimes by means of the violation of maxims and sometimes by obeying the maxims. However, the most violated maxim is the maxim of relevance. This means that Temel generally gives responses to other characters but in an irrelevant way and by doing so, the humor is created. Examining the classification of speech acts by Finegar and Besnier, a substantial preponderance of 'representative' speech acts is observed in Temel jokes. Responses to questions, assertions, statements, and descriptions emerge as prominent linguistic features, overshadowing the relatively sparse instances of assessments, judgments, and commands within these narratives. The overarching reliance on question-and-answer formats further aligns Temel jokes with Austin's notion of illocutionary acts, establishing a significant link between linguistic acts and humor creation. In parallel, the investigation into locutionary acts reveals that humor often emanates from the statements articulated by characters in Temel jokes. This underscores the centrality of locutionary acts as vehicles for humor, underscoring the linguistic intricacies that contribute to comedic effects. While perlocutionary acts, encompassing persuasion, convincing, and influencing actions, play a role in humor creation, they constitute a smaller proportion of instances within Temel jokes.

When compared to the results of the study on Nasreddin Hodja Jokes carried out by Karahan (1997), some points can be mentioned between Temel Jokes and Nasreddin Hodja Jokes. First of all, turn-by-turn allocation of speeches are dominant in Temel jokes in terms of turn-takings since Temel does not have any pressure on or dominance over the other characters while pre-allocated systems are seen in Hodja's jokes due to his dominance in his jokes. Secondly, moving on maxims, the humor in both characters' jokes are created through mostly the violation and, to a lesser degree, obeying to these maxims. Thirdly, while the Hodja is a dominant character and thus perlocutionary acts can be greatly encountered in his jokes, illocutionary and locutionary acts are frequently observed in Temel Jokes since he is not a dominant character giving commands to others or he does not have any pressure on the other characters in the jokes.

In summary, the exploration of Temel jokes reveals a nuanced interplay of linguistic elements, wherein turn-taking dynamics, maxim violations, and speech acts collectively contribute to the intricate fabric of humor. This study adds a distinctive layer to the understanding of humor in linguistic contexts, underscoring the unique comedic dynamics inherent in Temel jokes.

6. Conclusion

This study undertook a comprehensive analysis of Temel jokes, employing a multifaceted approach encompassing turn-taking and adjacency pairs, conversational maxims, and speech acts classified within the frameworks of Finegan-Besnier and Austin. The findings offer valuable insights into the dynamics of humor creation within the context of Temel jokes.

The dominant observation in turn-taking dynamics reveals a conspicuous preference for a turn-by-turn allocation of speeches over pre-allocated systems. This pattern suggests that Temel, as the main character, operates without exerting pressure on other characters, fostering a conversational landscape characterized by uniformity in turn-taking sequences and duration of speech for each character.

Furthermore, the analysis of humor in Temel jokes underscores a noteworthy reliance on the violation of conversational maxims as a primary mechanism for generating comedic effects. Notably, the maxim of relevance emerges as the most frequently violated, with Temel often responding to other characters in an irrelevant manner, thus constituting a foundational element in the creation of humor.

Delving into the realm of speech acts, the study reveals a marked preference for 'representative' acts, particularly answers to questions, assertions, statements, and descriptions within Temel jokes. Assessments, judgments, and commands, while present, constitute a minor proportion within the narrative. Given the prevalent question-and-answer format of Temel jokes, a predominance of illocutionary acts, as theorized by Austin, becomes evident, playing a pivotal role in humor creation.

Moreover, the study illuminates the nuanced role of locutionary acts, indicating that humor occasionally emanates from the statements made by characters in Temel jokes. In contrast, acts related to persuading, convincing, or influencing actions are infrequent within these narratives, suggesting a limited role in humor creation.

In summary, while the inherent illogical and nonsensical behaviors of characters contribute to the humor in Temel jokes, a deeper examination reveals the pivotal role played by linguistic features. The intentional violation of conversational maxims, coupled with a reliance on specific speech acts, emerges as a robust mechanism for humor creation within Temel jokes. These linguistic intricacies, outlined in the study, significantly contribute to our understanding of the unique comedic dynamics inherent in Temel jokes.

References

- Apte, M. (1988). Disciplinary boundaries in humorology: An anthropologist's ruminations. *International Journal of Humor Research*, 1(1), 5–26.
- Attardo, S. (1994). *Linguistic theories of humor*. Mouton de Gruyter.

- Attardo, S. (1997). The semantic foundations of cognitive theories of humor. *Humor: International Journal of Humor Research*, 10(4), 395–420.
- Attardo, S., & Raskin, V. (1991). Script theory revis(it)ed: Joke similarity and joke representation model. *Humor: International Journal of Humor Research*, 4, 293–347.
- Attardo, S., Attardo, C. H., & Ritchie, S. M. (2002). Script oppositions and logical mechanisms: Modeling incongruities and their resolutions. *Humor: International Journal of Humor Research*, 15(1), 3–46.
- Austin, J. L. (1962/2005). *How to do things with words* (2nd ed.). Harvard University Press.
- Barnes, C. (1992). Comedy in dance. In W. Sorrell (Ed.), *Pennington*. Capella Books.
- Berger, A. A. (1976). Anatomy of the joke. *Journal of Communication*, 26, 113–115.
- Davies, C. (1987). Language, identity, and ethnic jokes about stupidity. *International Journal of the Sociology of Language*, 65, 39–52.
- Ermida, I. (2008). *The language of comic narratives: Humor construction in short stories*. Mouton de Gruyter.
- Finegan, E., & Besnier, N. (1989). *Language: Its structure and use*. Harcourt Brace Jovanovich.
- Freud, S. (1960). *Jokes and their relation to the unconscious* (J. Strachey, Trans.). Norton. (Original work published 1905)
- Gregory, J. C. (1924). *The nature of laughter*. Harcourt, Brace.
- Grice, H. P. (1975). Logic and conversation. In P. Cole & J. Morgan (Eds.), *Syntax and semantics: Vol. 3. Speech acts* (pp. 41–58). Academic Press.
- Humour. (2015, December 13). In *Wikipedia*. <https://en.wikipedia.org/wiki/Humour>
- Karahan, F. (1997). Nasreddin Hoca'nın söylemi üzerine bir deneme. In H.Ü. *İngiliz Dilbilimi Bölümü 25. Yıl Yazıları (1972–1997)*. Bizim Büro Basımevi.
- Keith-Spiegel, P. (1972). Early conceptions of humor: Varieties and issues. In J. H. Goldstein & P. E. McGhee (Eds.), *The psychology of humor* (pp. 3–39). Academic Press.
- Koestler, A. (1964). *The act of creation*. Hutchinson.
- LaFollette, H., & Shanks, N. (1993). Belief and the basis of humor. *American Philosophical Quarterly*, 30(4), 329–339.
- Morreall, J. (1987). *The philosophy of laughter and humor*. SUNY Press.
- Özünlü, Ü. (1991). Fıkralar ve gülmece dili. *Çağdaş Türk Dili*, 37.
- Raskin, V. (1985). *Semantic mechanisms of humour*. D. Reidel.
- Ritchie, G. (2004). *The linguistic analysis of jokes*. Routledge.
- Sacks, H., Schegloff, E. A., & Jefferson, G. (1974). A simplest systematics for the organization of turn-taking for conversation. *Language*, 50, 696–735.
- Searle, J. R. (1969). *Speech acts: An essay in the philosophy of language*. Cambridge University Press.

- Simpson, P. (2003). *On the discourse of satire: Towards a stylistic model of satirical humor*. John Benjamins.
- Speech act. (2015, November 16). In *Wikipedia*. http://en.wikipedia.org/wiki/Speech_act
- Uçar, A. (2014). Gülmece ve gülme kaynağı olarak sözcük oyunları. *Journal of Linguistics and Literature*, 11(1), 17–43.
- Viktoroff, D. (1953). *Introduction to social psychology of laughter*. Presses Universitaires de France.

The sources from which Temel Jokes are retrieved

- ‘Hasan’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/hasan.html>.
- ‘Mazeret’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/mazeret.html>.
- ‘Önem almak’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/nlem-almak.html>.
- ‘Bakış’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/bak.html>.
- ‘Islaklık’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/islaklik.html>.
- ‘Çirkinlik’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/irkinlik.html>.
- ‘Yılan’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/ylan.html>.
- ‘Su derin mi?’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/su-derin-mi.html>.
- ‘Temelin Treni’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/temelin-treni.html>.
- ‘Tek Asker’. Retrieved October 8, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/tek-asker.html>.
- ‘Karne’. Retrieved October 5, 2015 from <http://www.fikracenneti.com/temel-fikralari/karne-2-2-2>.
- ‘Cinsiyet’. Retrieved October 5, 2015 from <http://www.gulum.net/fikra/bolumler.php?op=goster&id=292>.
- ‘Temel ve Zeka Gücü’. Retrieved October 5, 2015 from <http://www.gulum.net/fikra/bolumler.php?op=goster&id=1113>.
- ‘Kiminle Evli ?’. Retrieved October 5, 2015 from http://www.kubidik.com/fikralar/temel_fikralari.htm.
- ‘Temele Mercedes Lazım’. Retrieved October 5, 2015 from http://www.trsohbeti.net/pages/fikra_oku.htm.
- ‘Temel Uçakta’. Retrieved October 5, 2015 from http://www.erenet.net/fikralar.php?op=FikraOku&id=Temel_Ucakta_bc3.
- ‘Avcı Temel’. Retrieved October 5, 2015 from <http://temeldursunfikralari.blogspot.com/2008/09/avc-temel.html>.
- ‘Tanker’. Retrieved October 5, 2015 from

<http://www.fikrasi.com/tanker-10011.htm>.

‘Adi neydi ?’ . Retrieved October 5, 2015 from

<http://www.fikrasi.com/adi-neydi-513.htm>.

‘Temel ile Fadime parkta’ . Retrieved October 5, 2015 from

<http://www.fikrasi.com/temel-ile-fadime-parkta-10073.htm>.

OPTIONAL SUBJECT-VERB AGREEMENT IN HERITAGE TURKISH¹

Serkan UYGUN

Asst. Prof., Bahçeşehir University, Department of English Language Teaching, serkan.uygun@bau.edu.tr,
ORCID: 0000-0002-0880-9280

Uygun, S. (2025). Optional subject-verb agreement in heritage Turkish. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 350–371). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X019>

ABSTRACT

In Turkish, 3rd person plural subjects normally appear with verbs that are unmarked for number, rendering these verb forms indistinguishable from the singular form. The plural morpheme *–lar* is preferentially omitted from the verb, especially in spoken discourse, so as to avoid repeating the same morpheme that also marks plurality on nouns. Plural suffix omission in Turkish is optional and is affected by grammatical, surface-level, and semantic constraints such as subject animacy and subject position. The present study investigates to what extent Turkish heritage speakers are sensitive to these constraints and whether they differ from non-heritage speakers. 48 non-heritage Turkish speakers resident in İstanbul, Türkiye, and 58 heritage Turkish speakers resident in Berlin and Potsdam, Germany, were tested by using a scalar acceptability judgement task. The experimental stimuli were created by manipulating both subject animacy and subject position to test the effect of animacy and subject-verb distance on the acceptability of overt plural marking on the verb. Besides confirming the general preference for singular verb forms, participants' judgement patterns were affected both by subject animacy and by subject position. Significant differences were observed between heritage and non-heritage speakers in their acceptance of plural-marked verbs, suggesting that the relatively subtle interplay between subject animacy and subject position on optional subject-verb agreement marking is not always fully acquired under heritage language conditions.

Keywords: Subject-verb agreement, Turkish, Heritage speakers, Subject animacy, Subject position

ÖZ

Türkçede 3. çoğul şahıs özneleri normalde fiilde 3. çoğul şahıs eki olmadan kullanılır, bu durumda bu fiil biçimleri 3. tekil şahıs çekiminden ayırt edilemezler. Çoğul eki olan

¹ This research is funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) – Project Number 317633480 – SFB 1287, Project B04. I thank Çilem Çiçek for her help with participant recruitment and data collection.

–*Ar*, özellikle konuşma dilinde isimlerde de çoğul durumu belirten aynı eki tekrarlamaktan kaçınmak için tercihen fiilde kullanılmaz. Türkçede 3. çoğul şahıs ekinin fiilde kullanılmaması tercihe bağlıdır ve bu durum öznenin canlı olup olmaması ile öznenin cümle içindeki konumu gibi dilbilgisel, yüzeysel ve de anlamsal kısıtlamalardan etkilenmektedir. Bu çalışmanın amacı Türkçe miras dil konuşucularının bu kısıtlamalardan ne derecede etkilendiklerini ve D1 Türkçe konuşucularından farklılık gösterip göstermediklerini araştırmaktır. Türkiye’nin İstanbul şehrinde yaşayan 48 D1 Türkçe konuşucusu ile Almanya’nın Berlin ve Potsdam şehirlerinde yaşayan 58 miras dil Türkçe konuşucusu skaler bir kabul edilebilirlik yargı testi kullanılarak test edildiler. Deneyde kullanılan cümleler fiilde 3. çoğul şahıs ekinin bulunmasının kabul edilebilirliği üzerinde öznenin canlı olup olmaması ile öznenin cümle içindeki konumunun etkisini ölçmek amacıyla hem öznenin canlılığının hem de öznenin cümledeki konumuna manipüle edilmesiyle oluşturulmuştur. Fiilin 3. tekil şahıs olarak kullanılmasına yönelik genel tercihin doğrulanmasının yanı sıra, sonuçlar katılımcıların yargı kalıplarının hem öznenin canlı olup olmamasından hem de öznenin cümle içindeki konumundan etkilendiğini göstermektedir. Fiilin 3. çoğul şahıs eki ile kullanımı konusunda D1 ve miras dil Türkçe konuşucuları arasında önemli farklılıklar gözlemlenmiştir. Bu durum, tercihe bağlı olan 3. çoğul şahıstaki özne-fiil uyumunda etkili olan öznenin canlılığı ve konumu etkileşiminin miras dil koşulları göz önünde bulundurulduğunda tam olarak edinilemediğini göstermektedir.

Anahtar Sözcükler: Özne-fiil uyumu, Türkçe, Miras dil konuşucuları, Öznenin canlılığı, Öznenin cümledeki konumu

1. Introduction

In languages with subject-verb agreement (SVA) marking, the verb has to correspond in number with the subject for the sentence to be grammatical; that is, the verb has to be properly inflected. Previous research has reported that heritage speakers (HS) experience difficulties with inflectional morphology and morphosyntax, including SVA marking in their heritage language (HL) (see Benmamoun et al., 2013a, 2013b for reviews). HS are considered a subset of (unbalanced) bilinguals who were raised in homes where a language other than the dominant community language was spoken, resulting in some degree of bilingualism in both the heritage and the community language (Scontras et al., 2018). Because HS acquire the non-dominant language from reduced input and practice it less, they struggle a lot with inflectional processes and frequently make SVA marking errors, such as failing to produce the correct agreement or accepting agreement mismatches (Polinsky, 2018; Scontras et al., 2018). HS have also been reported to be inclined to use more singular forms (Prada Pérez & Pascual y Cabo, 2011), and Polinsky and Scontras (2020) claim that HS show a tendency to simplify agreement systems by using singular forms as a default. While most of the HS studies have investigated the categorical SVA marking in different languages such as Arabic, Russian, and Spanish, very few studies have focused on the optional agreement. It has been argued that optional SVA marking is more vulnerable and more likely to be affected under HL conditions when

compared to the categorical SVA marking (Benmamoun et al., 2013a, p. 161-166).

One example of the optional SVA marking languages is Turkish. Although Turkish is an agreement-marking language with specific person markers, it has been observed that under certain circumstances, the 3rd person plural marker *-lar* can be omitted from the verb. The omission of the plural marker *-lar* is affected by grammatical, surface-level, and semantic constraints. The omission of the plural marker makes the verb form indistinguishable from the 3rd person singular form, in which no person marker is added to the verb (e.g., Göksel & Kerslake, 2005; Kornfilt, 1997; Schroeder, 1999; Sezer, 1978). This situation is illustrated in (1).

- (1) Çocuk-lar kütüphane-den gel-di-(ler).
 *child-PL library-ABL come-PST-(PL)*²
 'Children came from the library.'

Although native Turkish speakers have the tendency to avoid multiple occurrences of plural markers within a single clause or phrase, prescriptive grammar books clearly prohibit the omission of plural agreement markers (Göknel, 2013, p. 424-425; Korkmaz, 2009, p. 27). According to Johanson (1998, p. 37), this phenomenon is a feature of Turkic languages, which is applied with an attempt to use morphological devices economically and avoid redundancy. Yet, some factors, such as subject animacy (Sezer, 1978) together with the subject position (Bamyacı, 2016) have been found to affect the possible omission of the plural marker from the verb.

The current study aims at exploring the effect of subject animacy and subject position in the acceptance of plural marker in the verb forms with 3rd person plural subjects in an acceptability judgment experiment and comparing heritage and non-heritage Turkish speakers in their optional agreement marking. As HS receive limited and qualitatively different input when compared to monolingually-raised individuals, this may result in divergent attainment in their HL (Scontras et al., 2015) and differences in the usage of the optional agreement marking due to their problems with morphosyntactic operations (Montrul 2008). The present study also investigates how and to what extent subject animacy interacts with subject position in the optional agreement marking. Finally, the design of the stimuli intends to find out how both groups' optional acceptance of plural verb forms is affected when the distance between the subject and the verb increases.

² Abbreviations used for morphemic glosses: 1ST PERSON PL = first person plural marker; ABL = ablative; ACC = accusative; AOR = aorist; DAT = dative; GEN = genitive; INABIL = inability marker; LOC = locative; PART = participle; PL = plural; POSS = possessive; PRS = present tense; PST = past tense; REL CL = relative clause; REP PST = reported past tense; SP = subject participle.

1.1 Subject-verb agreement in heritage languages

The agreement relation between the subject and the verb has been an ideal linguistic phenomenon to investigate with HS, and most of the previous studies have focused on the agreement errors during production. Polinsky (1997, 2006) found a clear proficiency effect for HS in their SVA marking; that is, highly proficient Russian HS had an accuracy rate of 66% in their correct SVA marking, whereas the low proficiency group's accuracy was only 30%. In another study with Egyptian and Palestinian Arabic HS, Albrini et al. (2013) found an error rate of 18% in SVA marking. It is crucial to note that non-heritage speakers also make SVA marking errors during production; however, their error rate is quite low, remaining below 5% (Poullisse, 1999). From the studies above, it can be inferred that there is a considerable difference in the error rates of heritage and non-heritage speakers, confirming that SVA marking is a vulnerable grammatical domain for HS.

Conversely, there are also studies that show SVA marking in HS is unaffected and native-like. For example, studies with Hungarian HS (De Groot, 2005; Fenyvesi, 2000) demonstrate no vulnerability in their SVA marking. Similar results were also obtained for English-Hungarian bilingual children living in the United States and whose Hungarian is the HL (Bolonyai, 2007). In another study, Montrul et al. (2012) used an oral production task with Hindi HS living in the United States and found that they had approximately the same accuracy level as the non-heritage speakers in their production of correct SVA marking (99.87% vs. 99.76%).

There are also studies that explore SVA marking errors in the comprehension of HS. By employing an untimed grammaticality task, Sherkina-Lieber (2011) and Sherkina-Lieber et al. (2011) investigated SVA marking with HS in an agglutinative language, Labrador Inuittitut. The results indicated 75% accuracy in HS with high proficiency and 33% accuracy in HS with low proficiency. These numbers suggest that SVA marking errors can be observed in comprehension as well. In another study, Foote (2011) used a word-by-word sentence reading task and tested SVA marking in Spanish non-heritage speakers, Spanish HS, and L2 Spanish learners. Foote observed that all participating groups were sensitive to number agreement errors since they displayed longer reading times for ungrammatical sentences. Finally, Scontras et al. (2018) replicated the design and materials of Fuchs et al. (2015) to investigate Spanish agreement marking with HS. The results of an auditory sentence rating task showed differences between heritage and non-heritage speakers, as HS were found to lose sensitivity to agreement errors.

To recapitulate, previous production and comprehension studies show that categorical SVA marking can be vulnerable in HS not only in production but also in comprehension. Yet, very little is known about the optional SVA marking, and few studies have focused on how HS would perform in this phenomenon, which serves as the main motivation of the present study.

1.2 Optional Subject-Verb Agreement in Turkish

Although Turkish is an SVA marking language, it is also a good example of the optional marking. While 1st and 2nd person singular subjects are marked with singular verb forms, their plural counterparts are marked with plural verb forms. Yet, 3rd person plural subjects have optional verb marking and can also be used with unmarked ($\emptyset$) singular verb forms, where the omission of the plural suffix *-lar* renders the verb form indistinguishable from the 3rd person singular form (Bamyacı et al., 2014; Bamyacı, 2016; Göksel & Kerslake, 2005; Kornfilt, 1997; Schroeder, 1999; Sezer, 1978); see (1) above. According to Kirchner (2001), the preferential omission of the plural marker from the verb is mostly seen in the spoken discourse but not in the formal written discourse. Previous theoretical and experimental studies have shown that the subject's degree of animacy plays a significant role in the usage of the plural marker on the verb for sentences with 3rd person plural subjects (Bamyacı et al., 2014; Bamyacı, 2016; Göksel & Kerslake, 2005; Kirchner, 2001; Kornfilt, 1997; Lago et al., 2019; Schroeder, 1999; Sezer, 1978; Uygun & Felser, 2023). The general consensus is that 3rd person animate plural subjects may take either the plural marker on the verb or remain unmarked (2), while 3rd person inanimate plural subjects usually remain unmarked (3), taken from Sezer (1978, p. 26).

- (2) *Öğrenci-ler* *gel-di-(ler).*
 student-PL come-PST-(PL)
 'Students came.'
- (3) *Mektup-lar* *gel-di-(*ler).*
 letter-PL come-PST-(*PL)
 'Letters arrived.'

The example with the 3rd person animate plural subject (2) shows optional SVA marking. It has been argued that the presence of the plural marker on the verb in sentences with 3rd person animate plural subjects depends on the speaker's stylistic preferences and does not cause a change in the meaning of the statement (Kirchner, 2001; Kornfilt, 1997; Sezer, 1978).

Another crucial factor that affects the optional SVA marking in Turkish is the subject's position in the sentence. Turkish is classified as a head-final and left-branching language whose canonical word order is Subject-Object-Verb (SOV) (Erguvanlı, 1984; Kornfilt, 1997; Göksel & Kerslake, 2005). However, Turkish word order is very flexible and allows all the possible orders. This means that constituents can occur in any order, and each order is equally grammatical (Göksel & Kerslake, 2005, p. 343). Erguvanlı (1984) claims that the grammatical role of a Turkish constituent does not depend on its position in the sentence since Turkish constituents are morphologically marked to signal their grammatical role, leading to

flexibility in word order. Johanson and Csátó (1998) claim that when a 3rd person plural subject appears close to the verb, especially in Object-Subject-Verb (OSV) sentences, the SVA marking shows optionality and occurs less frequently. Turkish speakers exhibit a general tendency to use morphological markers economically, and this tendency leads them to avoid duplicating plural markers that are very close to each other. This is illustrated by (4) taken from Schroeder (1999, p. 117-118).

- (4) *Türkiye'nin çeşitli yer-ler-in-den gel-en*
 Türkiye'GEN various region-PL-POSS-ABL come-SP
organize ol-ama-mış gemi adam-lar-ın-ı
 organized be-INABIL-REP PST ship man-PL-POSS-ACC
bu simsar-lar sat-ar.
 this agent-PL sell-AOR
 'It is these job agents who manage the sailors which come from different regions of Turkey and could not organize themselves.'

The information-structure dimension of sentence interpretation in Turkish has assigned three sentence positions: immediately preverbal, postverbal, and sentence-initial (Erguvanlı 1984, Kornfilt 1997, Kılıçaslan 2004, Göksel & Kerslake 2005). The immediately preverbal position is used to emphasize a particular constituent, which is also called *focusing*, to highlight the information provided by the constituent. Postverbal position is for de-emphasizing a particular constituent or constituents, which is also known as *backgrounding*. Finally, the sentence-initial position is used to make a particular constituent or constituents the pivot of the information that is provided in a sentence and is also referred to as *topic* (Göksel & Kerslake, 2005, p. 344). According to Bamyacı (2016, p. 120), if a subject appears in the non-initial position, this subject cannot be interpreted as topic, and this situation also affects the omission of the plural marker on the verb. A 3rd person plural subject that is in the sentence-initial position can trigger the use of the plural marker on the verb as it is the topic of the sentence; however, the same subject in the immediately preverbal position may trigger the use of a singular verb form because it is the focus of the sentence (Bamyacı, 2016, p. 127).

The final factor that affects the optional SVA marking is the linear distance between the subject and the verb. According to Göksel (1987), if there are many sentence materials between the 3rd person plural subject and the verb, then the verb takes the plural marker without considering the animacy of the 3rd person plural subject. The example with the animate subject (5) is taken from Schroeder (1999, p. 117), and the one with the inanimate subject (6) is taken from Göksel (1987, p. 71).

- (5) *Tabii kadın-lar toplum-umuz-da bizim için*
 of course woman-PL society-POSS-LOC our for
çok çok önemli bir yer teşkil ed-iyor-lar.
 very very important a place form-PRS-PL
 'Of course, for us, women are of great importance in our society.'
- (6) *Kitap-lar ya dün akşam-ki deprem-in*
 book-PL either last night-REL CL earthquake-GEN
şiddet-in-den ya da o eski kütüphane zaten
 force-POSS-ABL or that old library already
çürük ol-duğ-u için yer-e düş-müş-ler.
 rickety be-PART-POSS because floor-DAT fall-REP PST-PL
 'The books have fallen on the floor either owing to the force of
 last night's earthquake or because that old library was rickety
 anyway.'

To encapsulate, the animacy of the 3rd person plural subject, the word order, and the linear distance between the 3rd person plural subject and the verb have been found to be crucial factors that influence the optional SVA marking in Turkish.

1.3 Experimental approaches to optional subject-verb agreement marking in Turkish

Few experimental studies have been conducted to explore the optional SVA marking in Turkish. In one of these studies, Bamyacı et al. (2014) employed a magnitude estimation acceptability judgment task to investigate the interaction of the optional SVA marking and the subject's level of animacy³ in two-word sentences that consist of a plural subject noun phrase followed by a verb. The experiment was conducted with young (between ages 21-35) and old (between ages 40-54) non-heritage Turkish speakers living in Türkiye. The researchers observed an overall preference for the singular verb forms and concluded that these forms were considered as default forms by non-heritage Turkish speakers. The researchers also found a strong effect of the subject's level of animacy on the optional SVA marking. When the subject was animate, the omission of the plural marker on the verb was optional; however, the plural marker was preferentially omitted when the subject was inanimate. Subtle differences were also observed between the young and old participant groups such that older participants had a stronger preference to accept the plural marker on the verb when compared to younger participants. In addition, the subject's level of animacy affected the older participants differently than the younger participants. The researchers concluded that the observed differences between the younger and older participants indicate synchronic language change, in which the omission of the plural marker on the verb becomes the norm.

³ The levels of animacy included human, animal, quasi-animate, and inanimate.

In another study, Bamyacı (2016) explored the interaction of the optional SVA marking and the subject's level of animacy by comparing heritage and non-heritage speakers of Turkish. She used the same stimuli and task and obtained similar judgment patterns from both groups. A general preference for the omission of the plural marker on the verb in two-word sentences was also observed in the HS. In addition, HS accepted plural-marked verbs more when the subject was animate but preferred the omission of the plural marker for inanimate subjects. Yet, Bamyacı also observed some differences between the groups. For example, HS had a greater likelihood of accepting plural-marked verbs, and the subject's level of animacy affected them differently. That is, both groups exhibited similar judgment patterns for inanimate subjects, whereas HS had a preference to accept plural-marked verbs more when the subject was animate. These findings show that the modulation of animacy on the optional SVA marking was different in the grammar of HS.

Lago et al. (2019) used a binary-choice acceptability judgment task and investigated the optional SVA marking in heritage and non-heritage speakers only with animate subjects. The researchers obtained some group differences in the judgment patterns, such as while the non-heritage group preferred the omission of the plural marker on the verb, HS accepted the plural-marked and singular forms of the verb to similar extents (86% and 91%, respectively). The researchers also found a significant effect of the age of acquisition (AoA) of German; that is, HS with later AoA of German performed similar to non-heritage speakers, whereas HS with early AoA of German preferred the plural-marked verb forms more.

Uygun and Felser (2023) recently examined the effect of both subject animacy and subject position (SOV vs. OSV) on the optional SVA marking with heritage and non-heritage speakers in a set of 24 experimental items consisting of five words. The HS were divided into two groups based on their AoA of German: lower proficiency HS (LPHS) were born in Germany and had an AoA of German below age seven, and advanced proficiency HS (APHS) were born in Türkiye and had an AoA of German after age seven. The results showed that singular verb forms were generally rated more favourably than plural-marked forms by all participant groups. All groups rated plural-marked verbs significantly better when the subject was animate, while singular verb forms were rated numerically better when the subject was inanimate. However, the effect of animacy was found to be less in OSV sentences, in which the subject was directly adjacent to the verb. In OSV sentences, singular verb forms were rated better than plural-marked forms by all groups regardless of animacy. The researchers also found significant differences between the LPHS and the two comparison groups. The LPHS over-accepted plural-marked forms with animate subjects in SOV sentences whilst they over-accepted singular verb forms in OSV sentences, indicating that they are particularly keen to

avoid morpheme duplication in OSV sentences. They also contrasted the two animacy conditions more strongly than both comparison groups.

Taken together, these results indicate that Turkish HS are sensitive to subject animacy that plays a role in optional SVA marking with 3rd person plural subjects. Compared to non-heritage speakers, they have a greater likelihood of accepting plural-marked verbs, which is especially observed in HS with early AoA of German. Due to the paucity of the previous studies, it is not completely clear if the observed pattern is a general feature of heritage Turkish or if the obtained divergences are specific to grammatical phenomena that allow optionality (Bamyacı, 2016). The effect of subject position in optional SVA marking needs more empirical evidence to make any generalizations. Therefore, the current study aims to build on and extend previous research by examining the effect of subject animacy and subject position in HS by further exploring the effect of linear distance between the subject and the verb on the optional SVA marking.

1.4 Purpose of the present study

The present study investigates to what extent Turkish HS with an early AoA of German are sensitive to grammatical, surface-level, and semantic factors on optional Turkish SVA marking in longer (8-word) sentences with an attempt to explore the following research questions:

- RQ1: Do heritage and non-heritage speakers show any differences in accepting plural-marked vs. singular verb forms across different subject animacy and subject position conditions?
- RQ2: If so, how do subject animacy and subject position conditions affect heritage and non-heritage speakers' acceptance of different verb forms?

Based on the results of the previous studies, an overall preference for singular verb forms in both groups is expected. Regarding the effect of subject animacy, more singular verb forms are predicted for sentences with inanimate than animate subjects. When it comes to the effect of subject position, the most plural-marked verb acceptances are anticipated in SOV sentences, where there is a lot of sentence material between the subject and the verb, whereas singular verb forms are predicted to be accepted mostly in OSV sentences, when the subject and the verb are adjacent constituents. Finally, as the HS participants have an early AoA of German, significant differences with the non-heritage group are expected. The first difference is assumed to be observed in the acceptance of different verb forms. If HS accept more singular verb forms, this will show their tendency to simplify the optional SVA marking and use singular verb forms as default. Conversely, if HS accept more plural-marked verbs, this can be related to their tendency to regularize the optional SVA marking. The second difference is foreseen in the effect of subject animacy and subject position because previous studies have shown that HS are more strongly affected

by their effects and exhibited more divergent sensitivity to their effects in the optional SVA marking.

2. Method

2.1 Sample

48 non-heritage Turkish speakers were recruited and tested in İstanbul, Türkiye. All spoke the standard dialect of Turkish, and none of them reported any knowledge of German. This group (37 females, between ages 19-60, $M = 36.25$, $SD = 9.94$) will be referred to as non-heritage control speakers (CTR). All of them were either studying at the university or were university graduates when they were tested. The HS group involved 60 participants who spoke both Turkish and German daily and were recruited from the large Turkish community in Berlin and Potsdam. 2 participants from the HS group were excluded, one due to low Turkish proficiency and one due to high error rates ($> 30\%$) in the experiment. The remaining 58 HS (41 females, between ages 18-50, $M = 27.78$, $SD = 6.11$) were exposed to Turkish from birth and had an early age of acquisition of German (between ages 0-6, $M = 3.06$, $SD = 1.82$). The Turkish proficiency of the HS group was measured with the language structure section of the Turkish TELC (The European Language Certificates) test, which is designed for B2 level according to the CEFR (Common European Framework of Reference). The language structure section consists of two cloze tests comprising 20 questions in total. Participants who received less than 12 out of 20 were excluded because this number indicates a proficiency lower than B2 level. The HS group had a high Turkish proficiency score out of 20 ($M = 18.45$, $SD = 1.63$). The HS group also self-rated their Turkish skills on a 10-point scale for speaking ($M = 7.95$, $SD = 1.59$), listening ($M = 8.85$, $SD = 1.16$), writing ($M = 7.20$, $SD = 2.01$), and reading ($M = 8.16$, $SD = 1.73$). Both the TELC scores and the self-rating scores indicate that the HS group had a high Turkish proficiency, and they were using Turkish actively in their daily lives ($M = 60.58\%$, $SD = 23.28$).

2.2 Instrument

The materials of the present study were adapted from the 24 experimental sentence sets that were used by Uygun and Felser (2023)⁴. Uygun and Felser (2023) constructed five-word sentences in which they manipulated the subject's position (SOV vs. OSV), subject animacy (animate vs. inanimate), and verb marking (plural vs. singular). The subject of the experimental sentences was always a 3rd person plural subject, and all animate subjects referred to human entities. In SOV sentences, the subject appeared in the initial position of the sentence while it was adjacent

⁴ A full list of Uygun and Felser's (2023) experimental sentence sets can be found at the Center for Open Science Framework website at <https://osf.io/9caxb>.

to the verb in OSV sentences. In the present study, these 24 experimental sentences were made eight words long without changing their subjects and the manipulations described above. This resulted in eight different conditions of an experimental sentence, which is illustrated in (7a–h).

(7) a. SOV – ANIMATE – PLURAL (SOV-ANI-PL)

Dağcı-lar dün akşam yüksek ve
climber-PL last night high and
karlı dağ-dan düş-tü-ler.
snowy mountain-ABL fell-PST-PL
'Climbers fell from the high and snowy mountain last night.'

b. SOV – ANIMATE – SINGULAR (SOV-ANI-SG)

Dağcı-lar dün akşam yüksek ve
climber-PL last night high and
karlı dağ-dan düş-tü.
snowy mountain-ABL fell-PST-Ø
'Climbers fell from the high and snowy mountain last night.'

c. SOV – INANIMATE – PLURAL (SOV-INANI-PL)

Kaya-lar dün akşam yüksek ve
rock-PL last night high and
karlı dağ-dan düş-tü-ler.
snowy mountain-ABL fell-PST-PL
'Rocks fell from the high and snowy mountain last night.'

d. SOV – INANIMATE – SINGULAR (SOV-INANI-SG)

Kaya-lar dün akşam yüksek ve
rock-PL last night high and
karlı dağ-dan düş-tü.
snowy mountain-ABL fell-PST-Ø
'Rocks fell from the high and snowy mountain last night.'

e. OSV – ANIMATE – PLURAL (OSV-ANI-PL)

Dün akşam yüksek ve karlı dağ-dan
last night high and snowy mountain-ABL
dağcı-lar düş-tü-ler.
climber-PL fell-PST-PL
'Last night, climbers fell from the high and snowy mountain.'

f. OSV – ANIMATE – SINGULAR (OSV-ANI-SG)

Dün akşam yüksek ve karlı dağ-dan
last night high and snowy mountain-ABL
dağcı-lar düş-tü.
climber-PL fell-PST-Ø
'Last night, climbers fell from the high and snowy mountain.'

g. OSV – INANIMATE – PLURAL (OSV-INANI-PL)

Dün akşam yüksek ve karlı dağ-dan
 last night high and snowy mountain-ABL
kaya-lar düş-tü-ler.
 rock-PL fell-PST-PL

‘Last night, climbers fell from the high and snowy mountain.’

h. OSV – INANIMATE – SINGULAR (OSV-INANI-SG)

Dün akşam yüksek ve karlı dağ-dan
 last night high and snowy mountain-ABL
kaya-lar düş-tü.
 rock-PL fell-PST-Ø

‘Last night, climbers fell from the high and snowy mountain.’

By using the Latin-square design, eight different experiment lists were created, and each participant completed only one list. In each list, there were 24 experimental sentences that were pseudo-randomized. 48 filler sentences were created and mixed with the experimental sentences, resulting in 72 sentences in each list. All of the filler sentences were 8-word long and had the same structure as the experimental sentences. The subjects of the filler sentences were first person plural, second person plural, or third person singular, and in half of the filler sentences the verb was correctly conjugated, whereas in the other half the conjugation was not correct (8a–b).

(8) a. GRAMMATICAL FILLER

Biz geçen akşam yaşlı bir kadın-a
 we last night old one woman-DAT
yardım et-ti-k.
 help-PST-1ST PERSON PL

‘We helped an old lady last night.’

b. UNGRAMMATICAL FILLER

Biz geçen sabah kahvaltı için
 we last morning breakfast for
taze simit al-dı.
 fresh bagel buy-PST-Ø

‘We bought fresh bagel for breakfast yesterday morning.’

There were three main motivations for including these filler sentences in the experiment. The first motivation was to add some variety to the stimulus list with an attempt to reduce the likelihood of participants’ developing a response strategy for the experimental sentences. The second motivation was to assess whether the participants were sensitive to the obligatory SVA marking in Turkish. This was important to eliminate the possibility that the heritage group’s performance on the experimental sentences might indicate a general problem with the SVA marking system

of Turkish. And the final motivation was to check whether the participants did the experiment conscientiously.

2.3 Design

The acceptability judgment experiment was prepared using Google® Forms, which is also used as a web-based survey administration. The experiment was set up on a five-point scale where 1 refers to “absolutely acceptable” and 5 refers to “absolutely unacceptable”. Participants were instructed to rate each sentence they saw on the computer screen intuitively. Sentences were presented one by one, and there was no time limitation. Each sentence was presented separately, and the five-point rating scale with reference to what 1 and 5 refer to was placed below each sentence. Participants made their ratings via mouse click, and when they pressed the next button, the next sentence was presented to them. A bar on the page allowed the participants to keep track of their progress.

The participants in the CTR group were sent a web link to the experiment and completed it remotely on their personal computers. The HS group was invited to the laboratory room and did the experiment on the laboratory computer under supervision. Before starting the experiment, HS were asked to fill in a consent form and a detailed background questionnaire and to provide self-rating for their four skills in Turkish. The experiment started with an instruction page followed by a short bio form of the participants. After reading and completing these two pages, participants completed the experiment, which took approximately 20 minutes. After the experiment, the HS group completed the Turkish proficiency test.

Before starting the data analysis, the collected rating data were z-transformed with an attempt to eliminate possible bias such as scale compression or scale skew (Schütze & Sprouse, 2014). Statistical analyses on the rating data were conducted with R, which is an open-source programming language and environment for statistical computing (R Core Team, 2021). Linear mixed-effects regression models with crossed random effects for items and participants were used to analyze the z-transformed rating data (Baayen et al., 2008). The models included the participant-level variable *group* (CTR vs. HS) and item-level variables *subject animacy* (animate vs. inanimate), *subject position* (SOV vs. OSV), and *verb marking* (plural vs. singular) as fixed effects together with random slopes for item and participant. The models were fitted using the package *lme4* (Bates et al., 2015). For the main effects and overall interactions of the factors group, subject animacy, subject position, and verb marking, sum-coded contrasts (-0.5, 0.5) were employed, while treatment contrasts were employed for single comparisons. A model with maximum random effects and interactions was constructed as a starting point, and when the model did not converge, it was gradually simplified until convergence was reached (Barr et al., 2013). During the simplification process, random slopes by item and participant for each fixed effect in the model were retained only

if they improved the model fit significantly. This improvement was measured by using the Akaike Information Criterion (AIC), which provides a measure that penalizes complexity and leads to predictors being kept only when they substantially contribute to explaining variance in the data (Venables & Ripley, 2002). Each time during the simplification process, the model with the lower AIC was selected until no model with a lower AIC was produced. The final version of the model included random slopes for subject animacy and verb marking by item and the interaction of subject animacy and verb marking by participant. The effect sizes are reported by using model coefficients in log odds (β), standard errors (SE), t -statistics, and p values. P -values were computed by using the *lmerTest* package and the Satterthwaite's approximation for denominator degrees of freedom (Kuznetsova et al., 2014).

3. Results

Table 1 provides an overview of the group's acceptability ratings for each condition. Recall that ratings closer to 1 are considered as better ratings, indicating more acceptance.

Table 1. Means and standard deviations of acceptability ratings for both groups

Condition	CTR		HS	
	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>
SOV-ANI-PL	1.80	1.15	1.41	0.84
SOV-ANI-SG	1.89	1.20	2.02	1.42
SOV-INANI-PL	2.93	1.56	1.78	1.09
SOV-INANI-SG	1.88	1.22	1.98	1.31
OSV-ANI-PL	2.76	1.32	2.19	1.30
OSV-ANI-SG	2.09	1.09	1.99	1.25
OSV-INANI-PL	3.83	1.30	3.22	1.58
OSV-INANI-SG	1.94	1.10	1.73	1.01

SOV: Subject-Object-Verb sentence; OSV: Object-Subject-Verb sentence; ANI: Animate subject; INANI: Inanimate subject; PL: Plural-marked verb; SG: Singular (unmarked) verb form.

Table 2 presents the full results of the omnibus analysis. The main effect of subject position (β : -0.375, SE : 0.036, t = -10.478, p < .001) shows that SOV sentences (M = 1.96) were rated significantly better than OSV sentences (M = 2.47), the main effect of verb marking (β : 0.310, SE : 0.079, t = 3.921, p < .001) reflects that plural-marked verbs (M = 2.49) were rated significantly worse than singular verb forms (M = 1.94), the main effect of subject animacy (β : -0.310, SE : 0.061, t = -5.113, p < .001) indicates that animate subjects (M = 2.02) were rated significantly better than inanimate subjects (M = 2.41). Finally, the main effect of group (β : 0.222, SE : 0.073, t = 3.042, p < .003) reveals that the HS group's ratings (M = 2.04) were significantly better than the CTR group's ratings (M = 2.39), indicating more accepting responses for the HS group in general.

Table 2. Linear mixed effects model output of the omnibus analysis

	β	SE	t	p
Subject position (SOV vs. OSV)	-0.375	0.036	-10.478	.000
Verb marking (plural vs. singular)	0.310	0.079	3.921	.000
Subject animacy (animate vs. inanimate)	-0.310	0.061	-5.113	.000
Group (CTR vs. HS)	0.222	0.073	3.042	.002
Subject position*Verb marking	0.754	0.072	10.532	.000
Subject position*Subject animacy	-0.044	0.072	0.608	.544
Subject position*Group	0.034	0.072	0.473	.636
Verb marking*Subject animacy	0.668	0.088	7.608	.000
Verb marking*Group	-0.545	0.123	-4.419	.000
Subject animacy*Group	0.259	0.074	3.524	.000
Subj position*Verb marking*Subj animacy	0.360	0.143	2.514	.012
Subj position*Verb marking*Group	0.338	0.143	2.360	.018
Subj position*Subj animacy*Group	-0.232	0.143	-1.623	.104
Verb marking*Subj animacy*Group	-0.230	0.176	-1.312	.190
Subj position*Verb marking*Subj animacy*Group	0.582	0.287	2.033	.042

Formula in R: rating_z ~ subject position * verb marking * subject animacy * group + (1 + verb marking + subject animacy | item) + (1 + verb marking * subject animacy | participant).

A number of significant interactions were also observed; however, as the focal point of the research questions was on group differences, significant interactions without the fixed effect *group* are not reported. Instead, significant interactions involving *group* are reported and discussed in detail. Two significant two-way interactions were obtained. The significant verb marking and group interaction (β : -0.545, SE : 0.123, t = -4.419, p < .001) shows that while both groups provided worse ratings for plural-marked verbs than for singular verb forms (CTR: M = 2.83 vs. 1.95, p < .001; HS: M = 2.15 vs. 1.93, p = .698), the observed difference in plural-marked vs. singular verb forms was significant only for the CTR group. In addition, the HS group rated plural-marked verbs significantly better than the CTR group (M = 2.83 vs. 2.15, p < .001), while there was no statistical difference regarding the singular verb forms (M = 1.95 vs. 1.93). Another significant two-way interaction was obtained for subject animacy and group (β : 0.259, SE : 0.074, t = 3.524, p < .001). This interaction indicates that while both groups rated sentences with animate subjects significantly better than sentences with inanimate subjects (CTR: M = 2.14 vs. 2.65, p < .001; HS: M = 1.90 vs. 2.18, p < .013), in general, HS provided better ratings for both animacy conditions. The interaction reflects the fact that both groups provided similar ratings for animate

subjects ($M = 2.14$ vs. 1.90 , $p = .266$), yet the group difference was significant for inanimate subjects ($M = 2.65$ vs. 2.18 , $p < .001$).

There was also a significant three-way interaction of subject position, verb marking, and group ($\beta: 0.338$, $SE: 0.143$, $t = 2.360$, $p < .019$). To resolve this interaction, the fixed effects subject animacy and group were excluded from the best fit model of the analysis. The results show significant subject position and verb marking interactions for both the CTR ($\beta: 0.569$, $SE: 0.115$, $t = 4.953$, $p < .001$) and the HS ($\beta: 0.952$, $SE: 0.108$, $t = 8.793$, $p < .001$) groups. When the verb was plural-marked, both groups rated SOV sentences significantly better than OSV sentences (CTR: $M = 2.37$ vs. 3.30 , $p < .001$; HS: $M = 1.60$ vs. 2.71 , $p < .001$). Within SOV sentences, while the CTR group provided significantly worse ratings for plural-marked verbs ($M = 2.37$ vs. 1.89 , $p < .007$), the HS group performed differently and rated plural-marked verbs significantly better ($M = 1.60$ vs. 2.00 , $p = .011$). When it comes to the singular verb forms, the CTR group rated SOV sentences numerically better than OSV sentences ($M = 1.89$ vs. 2.02 , $p = .242$), whereas the HS group had a different preference and rated SOV sentences numerically worse than OSV sentences ($M = 2.00$ vs. 1.86 , $p = .174$). Within OSV sentences, both groups rated plural-marked verbs significantly worse than singular verb forms (CTR: $M = 3.30$ vs. 2.02 , $p < .001$; HS: $M = 2.71$ vs. 1.86 , $p < .001$). The HS group's rating of sentences in the opposite direction when compared to the CTR group is the primary source of the significant subject position, verb marking, and group interaction.

Finally, a significant four-way interaction of subject position, verb marking, subject animacy, and group was obtained ($\beta: 0.582$, $SE: 0.287$, $t = 2.033$, $p < .043$). The results indicate no significant subject position, verb marking, and subject animacy interaction for the CTR group ($\beta: 0.067$, $SE: 0.193$, $t = 0.347$, $p = .728$) but a significant three-way interaction for the HS group ($\beta: 0.674$, $SE: 0.205$, $t = 3.286$, $p < .003$). Therefore, further interaction resolution analysis was carried out only for the HS group. When the verb was plural-marked, the HS group rated SOV sentences significantly better than OSV sentences regardless of animacy (animate: $M = 1.41$ vs. 2.19 , $p < .001$; inanimate: $M = 1.78$ vs. 3.22 , $p < .001$). With singular verb forms, the HS group rated SOV sentences numerically worse than OSV sentences when the subject was animate ($M = 2.02$ vs. 1.99 , $p = .840$) and marginally worse with inanimate subjects ($M = 1.98$ vs. 1.73 , $p = .064$). In addition, in SOV sentences, plural-marked verbs receive significantly better ratings than singular verb forms only when the subject was animate ($M = 1.41$ vs. 2.02 , $p < .001$), but with inanimate subjects, the difference was numerical ($M = 1.78$ vs. 1.98 , $p = .266$). Conversely, in OSV sentences, plural-marked verbs receive numerically worse ratings than singular verb forms with animate subjects ($M = 2.19$ vs. 1.99 , $p = .264$) but the difference was significant when the subject was inanimate ($M = 3.22$ vs. 1.73 , $p < .001$).

4. Discussion

The present study employed an acceptability judgment task and investigated to what extent Turkish HS were sensitive to grammatical, surface-level, and semantic factors on optional Turkish SVA marking in long (8-word) sentences and compared their ratings to the CTR group.

The results of the judgment data showed that SOV sentences were rated significantly better than OSV sentences by both groups, which replicates Uygun and Felser's (2023) study regarding the effect of subject position. In addition, both groups preferred singular verb forms over plural-marked forms overall, and this finding is in line with previous studies (Bamyacı, 2016; Lago et al., 2019; Uygun & Felser, 2023). Furthermore, both groups rated sentences with animate subjects significantly better than inanimate subjects, and animate subjects received better ratings for plural-marked verbs compared to inanimate subjects. This finding provides support for the subject asymmetry in Turkish (Sezer, 1978) and replicates previous experimental findings as well (Bamyacı, 2016; Lago et al., 2019; Uygun & Felser, 2023).

The first research question investigated if HS and CTR groups display any differences in accepting plural-marked vs. singular verb forms across different subject animacy and subject position conditions. The results indicated that the HS group performed differently from the CTR group. First of all, the HS group gave significantly better ratings than the CTR group in general, indicating more accepting responses. In addition, the HS group had a greater likelihood of accepting plural-marked verbs and rated sentences with plural-marked verbs significantly better than the CTR group. The acceptance of the plural-marked verbs by HS was also observed in previous experiments (Bamyacı, 2016; Lago et al., 2019; Uygun & Felser, 2023). This pattern provides no evidence for the simplification of the optional SVA marking but implies that the HS group tries to regularize it by accepting sentences with plural-marked verbs more in comparison to the CTR group. In addition, the HS group was affected differently by subject animacy and subject position because many interactions involving the fixed effect *group* were obtained. This finding supports previous studies that also employed HS with an early AoA of German (Lago et al., 2019; Uygun & Felser, 2023) and obtained many differences between the HS and CTR groups.

The second research question explored how subject animacy and subject position conditions affect the HS and CTR groups' acceptance of different verb forms and several differences between the groups were observed. For example, in sentences with plural-marked verbs, both groups performed similarly and provided significantly better ratings for SOV sentences in comparison to OSV sentences. However, in sentences with singular verb forms, a crucial difference between the groups was observed. While the CTR group rated singular verb forms better in SOV sentences than in OSV ones, the HS group rated singular verb forms better in OSV sentences, indicating that they are keen to avoid morpheme duplication in

OSV sentences. The HS group did not favor *dağcılar düştüler* ‘mountaineers fell + PL’ but preferred *dağcılar düştü* ‘mountaineers fell + Ø’ much more, which was also observed by Uygun and Felser (2023). Another difference that was obtained was related to the effect of subject animacy. While both groups gave similar ratings for sentences with animate subjects, this difference was significantly different when the subject was inanimate. This result indicates that the HS group contrasted the animacy of the subject more strongly than the CTR group. This result is consistent with the previous studies that found higher levels of sensitivity to the contrast of different animacy levels in HS with an early AoA of German (Krause, 2020; Krause & Roberts, 2020; Uygun & Felser, 2023). Finally, the interaction of subject position, verb marking, and subject animacy did not affect the optional SVA marking of the CTR group, whereas this interaction played a decisive role for the HS group. This interaction showed that the HS group was affected differently by subject animacy and subject position when they had to rate different verb forms. When the verb was plural-marked, the HS group rated SOV sentences significantly better than OSV sentences regardless of animacy. With singular verb forms, the HS group rated OSV sentences numerically better than SOV ones in animate and marginally better in inanimate subjects, indicating a stronger animacy distinction in sentences with singular verb forms.

The HS group did not show a tendency to simplify the optional SVA marking in Turkish and use the singular verb forms as the default form. Instead, HS tried to regularize the optionality and accepted plural-marked verbs more than the CTR group. Meanwhile, subject animacy and subject position also affected their verb form preferences differently when compared to the CTR group. How can the obtained group differences be explained? One possibility might be related to the influence of German. However, it is important to note that SVA marking in German is obligatory and is not affected by subject animacy and subject position. Therefore, the observed pattern of the HS group can only be explained by the conditions in which heritage Turkish is acquired. Reduced input and exposure to the HL, together with the lack of formal education in it, made it more difficult for HS to choose the correct verb form in a structure that involves optionality. The HS group views the effects of subject animacy and subject position on optional SVA marking differently than the CTR group. The acquisition of the optional SVA marking is experience-based, and as HS lack this experience, they cannot use the critical cues provided by subject animacy and subject position as effectively as the CTR group, leading to significant differences in the optional SVA marking in Turkish.

5. Conclusion

The present study employed a scalar acceptability judgement task and investigated to what extent HS with an early AoA of German were sensitive to grammatical, surface-level, and semantic constraints such as

subject animacy and subject position on optional SVA marking in Turkish. While the results indicate several similarities with the CTR group, some crucial differences were also observed. The results indicate that HS try to regularize the optional SVA marking by over-accepting plural-marked verbs. In addition, subject animacy and subject position also affected the HS group in a different way that resulted in divergent rating patterns when compared to the CTR group. The observed group divergencies can be seen as a general feature of heritage Turkish, as these subtle differences are related to the limited and qualitatively different input HS with an early AoA of German received while acquiring their HL. In addition, HS with an early AoA of German also lack the experience of using the language effectively and did not receive formal education in their HL. All of these factors lead to crucial disparities in the usage of the optional SVA marking and the factors that affect the optionality when compared to non-heritage Turkish speakers.

References

- Albrini, A., Benmamoun, E., & Chakrani, B. (2013). Gender and number agreement in the oral production of Arabic heritage speakers. *Bilingualism: Language and Cognition*, 16(1), 1-18. <https://doi.org/10.1017/S1366728912000132>
- Baayen, R. H., Davidson, D. J., & Bates, D. (2008). Mixed-effects modelling with crossed random effects for subjects and items. *Journal of Memory and Language*, 59(4), 390-412. <https://doi.org/10.1016/j.jml.2007.12.005>
- Bamyacı, E. (2016). *Competing structures in the bilingual mind: A psycholinguistic investigation of optional verb number agreement*. Springer. <https://doi.org/10.1007/978-3-319-22991-1>
- Bamyacı, E., Häussler, J., & Kabak, B. (2014). The interaction of animacy and number agreement: An experimental investigation. *Lingua*, 148, 254-277. <https://doi.org/10.1016/j.lingua.2014.06.005>
- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. (2013). Random-effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of Memory and Language*, 68(3), 255-278. <https://doi.org/10.1016/j.jml.2012.11.001>
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67(1), 1-48. <https://doi.org/10.18637/jss.v067.i01>
- Benmamoun, E., Montrul, S., & Polinsky, M. (2013a). Heritage languages and their speakers: Opportunities and challenges for linguistics. *Theoretical Linguistics*, 39(3-4), 129-181. <https://doi.org/10.1515/tl-2013-0009>
- Benmamoun, E., Montrul, S., & Polinsky, M. (2013b). Defining an ideal heritage speaker: Theoretical and methodological challenges. Reply to peer commentaries. *Theoretical Linguistics*, 39(3-4), 259-294. <https://doi.org/10.1515/tl-2013-0018>

- Bolonyai, A. (2007). (In)vulnerable agreement in incomplete bilingual L1 learners. *The International Journal of Bilingualism*, 11(1), 3-23. <https://doi.org/10.1177/13670069070110010201>
- De Groot, C. (2005). The grammars of Hungarian outside Hungary from a linguistic-typological perspective. In A. Fenyesi (Ed.), *Hungarian language contact outside Hungary* (pp. 351-370). John Benjamins Publishing. <https://doi.org/10.1075/impact.20>
- Erguvanli, E. E. (1984). *The function of word order in Turkish grammar*. University of California Press.
- Fenyvesi, A. (2000). The affectedness of the verbal complex in American Hungarian. In A. Fenyesi & K. Sándor (Eds.), *Language contact and the verbal complex of Dutch and Hungarian: Working papers from the 1st Bilingual Language Use Theme of the Study Center on Language Contact*, November 11-13, 1999, Szeged, Hungary, (pp. 94-107). JGyTF Press.
- Foot, R. (2011). Integrated knowledge of agreement in early and late English-Spanish bilinguals. *Applied Psycholinguistics*, 32(1), 187-220. <https://doi.org/10.1017/S0142716410000342>
- Fuchs, Z., Polinsky, M., & Scontras, G. (2015). The differential representation of number and gender in Spanish. *The Linguistic Review*, 32(4), 703-737. <https://doi.org/10.1515/tlr-2015-0008>
- Göknal, Y. (2013). *Turkish grammar: Updated academic edition*. Ege Basım.
- Göksel, A. (1987). Distance restrictions on syntactic processes. In H. E. Boeschoten & L. T. Verhoeven (Eds.), *Studies on modern Turkish. Proceedings of the third on Turkish Linguistics* (pp. 69-81). Tilburg University Press.
- Göksel, A., & Kerslake, C. (2005). *Turkish: A comprehensive grammar*. Routledge.
- Johanson, L. (1998). The structure of Turkic. In L. Johanson & É. Á. Csátó (Eds.), *The Turkic languages* (pp. 30-66). Routledge.
- Johanson, L., & Csátó, É. Á. (1998). *The Turkic languages*. Routledge.
- Kılıçaslan, Y. (2004). Syntax of information structure in Turkish. *Linguistics*, 42(4), 717-765. <https://doi.org/10.1515/ling.2004.024>
- Kirchner, M. (2001). Plural agreement in Turkish. *Turkic Languages*, 5, 216-225.
- Korkmaz, Z. (2009). *Türkiye Türkçesi Grameri Şekil Bilgisi*. Türk Dil Kurumu.
- Kornfilt, J. (1997). *Turkish*. Routledge.
- Krause, E. (2020). High sensitivity to conceptual cues in Turkish speakers with dominant German L2: Comparing semantics-morphosyntax and pragmatics-morphosyntax interfaces. In B. Brehmer & J. Treffers-Daller (Eds.), *Lost in transmission: The role of attrition and input in heritage language development* (pp. 198-228). John Benjamins Publishing. <https://doi.org/10.1075/sibil.59>
- Krause, E., & Roberts, L. (2020). [Interlanguage cue competition at semantics-morphosyntax interface: Animacy effects on DOM in L1 Turkish](https://doi.org/10.1075/sibil.59)

- [speakers with dominant German L2](#). In A. Mardale & S. Montrul (Eds.), *The acquisition of differential object marking* (pp. 313-341). John Benjamins Publishing. <https://doi.org/10.1075/tilar.26>
- Kuznetsova, A., Bruun Brockhoff, P., & Haubo Bojesen Christensen, R. (2014). lmerTest: Tests for random and fixed effects for linear mixed effect models (lmer objects of lme4 package). R package version 2.0-11. <https://cran.r-project.org/web/packages/lmerTest/index.html>
- Lago, S., Gracani-Yukse, M., Şafak, D.F., Demir, O., Kırkıcı, B., & Felser, C. (2019). Straight from the horse's mouth: Agreement attraction effects with Turkish possessors. *Linguistic Approaches to Bilingualism*, 9(3), 398-426. <https://doi.org/10.1075/lab.17019.lag>
- Montrul, S. (2008). *Incomplete acquisition in bilingualism: Re-examining the age factor*. John Benjamins Publishing. <https://doi.org/10.1075/sibil.39>
- Montrul, S., Bhatt, R., & Bhatia, A. (2012). Erosion of case and agreement in Hindi heritage speakers. *Linguistic Approaches to Bilingualism*, 2(2), 141-176. <https://doi.org/10.1075/lab.2.2.02mon>
- Polinsky, M. (1997). Cross-linguistic parallels in language loss. *Southwest Journal of Linguistics*, 14(1-2), 87-123.
- Polinsky, M. (2006). Incomplete acquisition: American Russian. *Journal of Slavic Linguistic*, 14(2), 191-262.
- Polinsky, M. (2018). *Heritage languages and their speakers*. Cambridge University Press. <https://doi.org/10.1017/97811075252349>
- Polinsky, M., & Scontras, G. (2020). Understanding heritage languages. *Bilingualism: Language and Cognition*, 23(1), 4-20. <https://doi.org/10.1017/S1366728919000245>
- Poulisse, N. (1999). *Speech errors in first and second language production*. John Benjamins Publishing. <https://doi.org/10.1075/sibil.20>
- Prada Pérez, A., & Pascual y Cabo, D. (2011). Invariable gusta in the Spanish of heritage speakers in the US. In J. Hershenshon & D. Tanner (Eds.), *Proceedings of the 11th Generative Approaches to Second Language Acquisition* (pp. 110-120). Cascadia Proceedings Project.
- R Core Team. (2021). R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria. <https://www.r-project.org/>
- Schroeder, C. (1999). *The Turkish nominal phrase in spoken discourse*. Harrassowitz Verlag.
- Schütze, C., & Sprouse, J. (2014). Judgment data. In R. Podesva & D. Sharma (Eds.), *Research methods in linguistics* (pp. 27-51). Cambridge University Press.
- Scontras, G., Fuchs, Z., & Polinsky, M. (2015). Heritage language and linguistic theory. *Frontiers in Psychology*, 6:1545. <https://doi.org/10.3389/fpsyg.2015.01545>
- Scontras, G., Polinsky, M., & Fuchs, Z. (2018). In support of representational economy: Agreement in heritage Spanish. *Glossa: A Journal of General Linguistics*, 3(1), 1-29. <https://doi.org/10.5334/gigl.164>

- Sezer, E. (1978). Eylemlerin çoğul öznelere uyumu. *Genel Dilbilim Dergisi*, 1, 25-32.
- Sherkina-Lieber, M. (2011). *Comprehension of Labrador Inuttitut functional morphology by receptive bilinguals*. [Doctoral dissertation, University of Toronto]. TSpace. <https://utoronto.scholaris.ca/items/a3a33113-7545-4f61-b08d-58f32e208fa4>
- Sherkina-Lieber, M., Perez-Leroux, A. T., & Johns, A. (2011). Grammar without speech production: The case of Labrador Inuttitut heritage receptive bilinguals. *Bilingualism: Language and Cognition*, 14(3), 301-317. <https://doi.org/10.1017/S1366728910000210>
- Uygun, S., & Felser, C. (2023). Constraints on subject-verb agreement marking in Turkish-German bilingual speakers. *Linguistic Approaches to Bilingualism*, 13(2), 190-217. <https://doi.org/10.1075/lab.19081.uyg>
- Venables, W. N., & Ripley, B. D. (2002). *Modern applied statistics with S*. Springer. <https://doi.org/10.1007/978-0-387-21706-2>

PHASAL DIAGNOSTICS: A CRITICAL REVIEW

Kardelen Tuana YILMAZ

PhD Candidate, Dokuz Eylül University, Department of Linguistics, tuana879@gmail.com, ORCID: 0009-0002-4149-4580

Murat ÖZGEN

Prof., Dokuz Eylül University, Department of Linguistics, murat.ozgen@deu.edu.tr, ORCID: 0000-0001-7960-6627

Yılmaz, K. T. & Özgen, M. (2025). Phasal diagnostics: A critical review. In O. Çınar, F. Başbuğ & H. Aydemir (Eds.), *Contemporary studies in linguistics I: IMU Linguistics 15th anniversary commemorative volume* (pp. 372–401). Artsürem. <https://doi.org/10.7816/imuling-15-2025-01X020>

ABSTRACT

Phase theory posits that syntactic structures are not generated in a single step, but rather through multiple cyclic derivations. Relevant literature regards CP, vP, DP, PredP, and PP as phases. Conceptually, this theory assumes that language is divided into smaller processing units due to the limited capacity of working memory. Empirically, phases are regarded as interface segments that display independent distributions across the perceptual-motor and conceptual-intentional systems. The main criteria used to identify phases include agreement, uninterpretable features, case assignment, extraction, ellipsis, and wh-movement. Our study examined the applicability of these criteria to Turkish, finding that most are incompatible with its structural properties. Consequently, we suggest that the notion of phase should be approached more critically, and that more crosslinguistic data are needed to better evaluate the validity of the phase theory.

Keywords: Phase theory, Multiple spell-out, Agreement, Phasal diagnostics

ÖZ

Evre kuramı, sözdizimsel yapıların tek aşamada değil çoklu döngüler sonucunda üretildiğini savunan bir kuramdır. İlgili alanyazın TÖ, eÖ, BelÖ, YükÖ ve İÖ'yü evre olarak görmektedir. Kavramsal olarak bu kuramda işler belleğin sınırlı kapasitesi nedeniyle dilin küçük işlem birimlerine ayrıldığı öne sürülür. Görgül açıdan ise evreler, duyudevinim ve düşünce sistemlerinde bağımsız dağılımlar sergileyen arakesit bölütleri olarak değerlendirilir. Evreleri belirlemede başlıca uyum, yorumlanamaz özellikler, durum atama, çıkarma, silme ve ne-taşıma yapıları ölçüt olarak kullanılmaktadır. Çalışmamızda bu ölçütlerin Türkçeye uygulanabilirliği incelenmiş ve çoğunun Türkçe dil yapısına uygun olmadığı görülmüştür. Sonuç olarak evre olgusuna daha eleştirel bir biçimde yaklaşılması gerektiği ve kuramın geçerliliğinin değerlendirilebilmesi için dillerarası verilerin incelenmesi gerektiği önerilmiştir.

Anahtar Sözcükler: Evre kuramı, Çoklu dağıtım, Uyum, Evre tanıları

1. Introduction

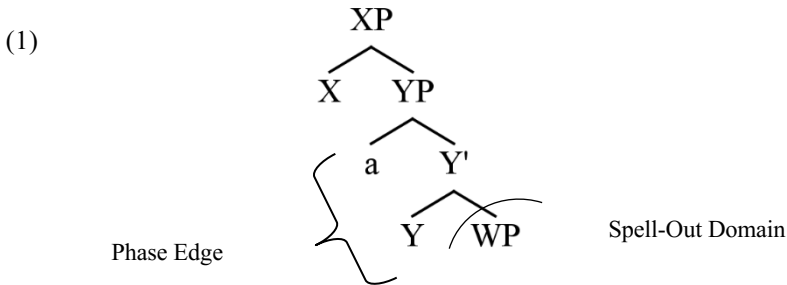
Phase theory is a linguistic framework that posits that syntactic structures are generated in a cyclic manner. While Phase Theory is broadly accepted today, there are also opposing views regarding the framework. This study re-evaluates the arguments supporting the theory and addresses its critics' claims. The primary aim of our work is to question the existence of the phase phenomenon itself. A review of the literature on Phase Theory reveals that most studies bypass the question "Do phases exist?" and instead focus on "Which phrases qualify as phases?"

Our study primarily addresses the definitions of phases, concepts related to phases, and the criteria for phase identification before shifting focus to testing the validity of the theory in the context of Turkish. Given that the theory incorporates various sub-theories and concepts, certain sections discuss these sub-theories and concepts with illustrative examples.

1.1 The phase phenomenon

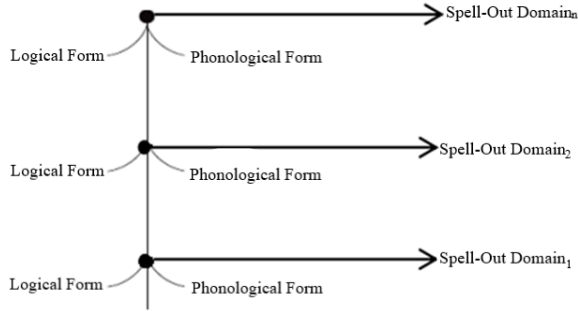
Chomsky (2000) asserts that syntactic structures are formed in phases, or groups of parts, from discrete lexical subarrays. Citko (2014) later offered a similar description, characterizing phases as domains where movement is identified and uninterpretable elements are localized.

According to Chomsky (2001), phase heads, specifier positions, and spell-out domains are the elements that make up a phase. We can represent Chomsky's definition in the tree as follows:



Spell-out domains are parts of phases that become inaccessible to subsequent operations. Phase Theory posits that multiple spell-out domains and, consequently, multiple transfers to logical and phonological forms are necessary for the realization of this cyclicity:

(2)



(Chomsky, 2000 and subsequent works)

As can be seen from the above multiple spell-out model, each point represents a spell-out domain, indicating that there must be multiple phases and spell-out domains during the derivation of a sentence.

The first fundamental assumption regarding Phase Theory, as stated by Chomsky (2000), is that phases operate under certain conditions. One of these conditions is known as the Phase Impenetrability Condition (PIC). According to this condition, after a phase is completed, its internal domains—i.e., the complements of the phase head—are sent to the interfaces. Consequently, these internal domains become inaccessible for any operations unless they are within the same phase. Chomsky (2001) later divided this condition into two versions: strong and weak. The main assumption of the strong version is that, except for the phase head and phase-specifier position located at the edge of a phase, all other units are inaccessible from outside the phase. The main assumption of the weak version is that the internal domain of a phrase which is a phase remains accessible until another phase is formed. Once another phase is formed, only the phase head and phase-specifier position of those domains are accessible. From these two versions, we can see that the strong version suggests that the transfer to the interfaces happens immediately after the phase is created. In contrast, the weak version posits that the transfer is delayed until another phase head is merged. During this interim period, the internal domain of a phrase, which constitutes a phase, remains accessible to other constituents.

2. Phasal Diagnostics and Turkish

There are various diagnostics that attempt to identify the borders of phasal domains (*see* Chomsky 2001, Citko 2014, Den Dikken 2006, Gallego 2010, Bošković 2014, Legate 2003, Matushansky 2005). Since phasal domains are convergent (PHON, SEM) objects, literature on phasal diagnostics dwell on interface questions. In this sense, agreement, case features, extraction, ellipsis, and sentential stress are basic test units that

indicate the phasal domains. Below, we focus on such tests and present contradictory cases with respect to Turkish data, applying the tests on renowned classic phases -i.e., CPs, DPs, *v*Ps, and PPs.

2.1 Agreement as a diagnostic for phases

Agreement occurs as follows (Chomsky 2000: 122):

- (3)
 - a. A probe *a* enters into an agreement relation with the closest matching goal *b*. The probe *a* carries at least one uninterpretable unvalued feature, while the goal *b* carries a matching interpretable valued feature.
 - b. *a* c-commands *b*.
 - c. *b* is the closest goal to *a*.
 - d. *b* carries uninterpretable unvalued features.

The uninterpretable ϕ -features present on the probe are what keep the probe active. The reason T inherits its uninterpretable ϕ -features from C (feature inheritance ‘FI’) is due to the closer structural position of the C head relative to the subject. Chomsky (2001) proposes that the probe-goal relationship is a property exclusive to phases, and that agreement plays a significant role in identifying phase heads as a phase diagnostic. In light of this reasoning, we should consider each agreeing phrase as a phase. Conversely, it implies that any phrase lacking an agreement relationship does not qualify as a phase.

Gallego (2010) argued that uninterpretable features identify phase boundaries, suggesting that agreement is an identifying factor in phasal domains. To test this claim, we should begin by examining Determiner Phrases (DPs) where agreement is clearly evident. The primary reason for this is that Turkish is an agglutinative language, meaning agreement is explicitly encoded within the language. First, let us look at the various types of DPs found in Turkish:

- (4)
 - a. DPs with possessive agreement

$[_{DP} \text{Ali-nin } [_{NP} \text{kitab-ı}]]$
 Ali-gen book-poss
 “Ali’s book”
 - b. Complex DPs

$[_{DP} \text{Doktor-un } [_{PredP} \text{hasta-yı} \quad \text{muayene-si}]]$
 Doctor-gen patient-acc examination-poss.3sg¹
 “the doctor’s examination of the patient”

¹ We employed Leipzig Glossing Rules throughout the paper:
<https://www.eva.mpg.de/lingua/resources/glossing-rules.php>

c. Sentential DPs

[_{DP} Ali-nin [_{CP} kereviz-i ye-diğ-i]]
 Ali-gen celery-acc eat-vnom-poss.3sg
 “the celery that Ali ate”

(Özgen, 2018: 7)

First, let us analyze a possessive agreement in a DP:

- (5) Ben-im kitab-ım
 pro-gen.1sg book-poss.1sg
 “My book”

In (5), there is a possessive agreement between "ben" (I) and "kitab" (book). Now, let us compare and contrast the following data:

- (6) Ben-im kitab-Ø
 pro-gen.1sg book-Ø
 “My book”

The possessive agreement morpheme is absent in example (6). One possible account of such optionality between (5) and (6) comes from Ouali (2008: 160), suggesting that the optionality arises due to the different outcomes of operations that heads with uninterpretable features apply to those features. In this sense, there are three sub-options of FI: *share*, *donate*, and *keep*.

To generate these possibilities, first, C donates its features. However, if the transfer results in a derivational conflict, the *keep* operation is applied, and Ouhalla (1993) states that the *keep* operation accounts for anti-agreement data in languages such as Tamazight Berber. Finally, if a conflict still occurs during the derivation, C's features are *shared*.

With the following examples, we can see the donate-keep-share triad more concretely:

- (7) a. Donate
 John drinks coffee.
- b. Keep
 ytsha wrba thamen.
 3sg.eat.PERF boy honey
 “The boy ate honey.”
- c. Share
 ma ag inna ali theila (*ilan) araw?
 who COMP 3sg.said ali 3sg.FEM.saw (*saw.PART) boys
 “Who did Ali say saw the boys?”

(Ouali, 2008)

Looking at the examples above, in the first one we see the donate case. For donate, Ouali (2008) states that the C head donates its

uninterpretable features to the T head without leaving a copy behind. As concrete evidence of this, we observe the subject–predicate agreement that occurs in T, namely the agreement between "John" and "drink," which results in the predicate "drinks."

When we examine the keep case, Ouali (2008) also refers to this option as reverse agreement. The uninterpretable features on C are not transferred to T. Therefore, no subject–predicate agreement arises here, and as a result, C functions as the probe.

Finally, when we look at the share case, C both shares its uninterpretable features with T and retains a copy of them on itself. As we can see from the example, agreement morphemes appear on both the object and the subject. This situation represents the share option.

Let us now examine how this system works:

(8) **Donate > Keep > Share**

a. Donate

C donates the ϕ -features it has to T without leaving a copy of them on itself.

b. Keep

C does not transfer the ϕ -features it has to T.

c. Share

C transfers its ϕ -features to T and retains a copy of them on itself.

(Ouali, 2008)

The first point that draws our attention is that, in principle, the *keep* and *share* options can only be active if the *donate* option causes a conflict in the derivation (Ouali, 2008: 162). If we assume that the *keep* option is active in example (6), we would anticipate that the derivation to cause a conflict due to the *donate* option. However, as we can see from example (5), such a situation does not occur. Therefore, two distinct issues arise here.

The first is the dependency issue. The *donate > keep > share* options could potentially be selected without being dependent on a conflict in the derivation. The second issue is the optionality issue. The apparent agreement structure in examples (5) and (6), which seems to be optional, creates a problem when aligned with Ouali's (2008) theory. Similarly, if we consider that the *donate* and *keep* options, and thus FI, are optional phenomena, they should be applicable to every DP containing a possessive:

- (9) a. *Öğretme-nin hala
teacher-gen aunt
"The teacher's aunt"

- b. *makale-nin başlık
article-gen title
"The title of the article"
- c. *bina-nın yıkım
building-gen demolition
"The demolition of the building"

(Öztürk & Taylan, 2015: 5)

If the *keep* were present in the "ben-im kitab-ım" (I-GEN book-POSS.1sg) example, examples in (9 a-b-c) would also be fine. Ouali's (2008) proposed trio of donate-keep-share does not explain the agreement in Turkish DP structures, and that agreement is influenced by factors independent of the presence of uninterpretable features on the DP head. Therefore, the lack of agreement in the "ben-im kitap" (I-GEN book-POSS.1sg) example is because the *keep* operation cannot apply. The presence of the same type of DP in the "bina-nın yıkım" (building-GEN demolition) and "ben-im kitab-ım" (I-GEN book-POSS.1sg) examples should display the same syntactic distribution if we consider DP to be a phase and if we employ agreement to mark the phase. However, while both DPs are identical in terms of uninterpretable features, the fact that agreement is mandatory in "bina-nın yıkım-ı" (building-GEN demolition-POSS.3sg) and optional in "ben-im kitab-ım" (I-GEN book-POSS.1sg) means that the agreement does not form a sound framework for discussing the phases of DPs. Alternatively, we can account for this structure without resorting to the phase concept. Contrary to Ouali's (2008) *keep* approach, Öztürk & Taylan (2016) suggest that in such structures in Turkish, the pronoun is an attachment to the DP, which explains the absence of agreement without relying on phases. Hence, we do not need phase and phasal features to account for this structure. Another issue here relates to Gallego's (2010) claim. Turkish is a language that clearly encodes agreement, and in order for agreement to occur, the head of the phrase must have uninterpretable ϕ -features. Using uninterpretable ϕ -features as phase diagnostics for DP structures would cause a problem, since we would conclude that "if there is agreement, uninterpretable features exist; if there is no agreement, uninterpretable features do not exist." Therefore, external factors that affect agreement in Turkish, independent of uninterpretable ϕ -features, would not offer a solid framework for discussing the phasal status of DPs.

In terms of CPs, Exceptional Case Marking (ECM) constructions in Turkish are a potential candidate for the agreement test. Let us first introduce a crash course on ECMs in Turkish. Şener (2008: 5) proposed that all ECM constructions are CPs and that within these structures, there can be a null C or an overt C:

- (10) Pelin [sen-i Timbuktu-ya git-ti-(n) diye] bil-iyor-muş.
 Pelin you-acc Timbuktu-dat go-past.(2sg) that know-prog-evid
 “Pelin thinks you went to Timbuktu.”

(Şener, 2008: 7)

The T head "diye" is overtly present in the sentence. Consider this example that involves a null T:

- (11) Pelin [sen-i Timbuktu-ya git-ti-(n)] san-ıyor.
 Pelin you-acc Timbuktu-dat go-past.2sg think-prog
 “Pelin thinks you went to Timbuktu.”

(Şener, 2008: 7)

We follow Şener (2008) and refer to the example in (10) as Type A ECM and the example (11) as the Type B ECM. Comparing the two examples, the absence of agreement in (10) shows us that the TP domain is defective. In (11), the presence of agreement indicates that T is not defective; however, the fundamental issue here is that, CP domains of both constructions are insensitive to tense constructions. In short, the CP domain is defective in both types of ECM:

- (12) a. A-type: Ali [sen-i İstanbul-a git-ti] san-ıyor.
 Ali you-acc İstanbul-dat go-past think-prog
 “Ali thinks you went to İstanbul.”
 b. B-type: Ali [sen-i İstanbul-a git-ti-n] san-ıyor.
 Ali you-acc İstanbul-dat go-past.2sg think-prog
 “Ali thinks you went to İstanbul.”

(Özgen & Aydın, 2016: 307)

Comparing the two examples, we observe that the absence of agreement in example (12a) indicates that the T domain is defective. In example (12b), the presence of agreement indicates that T is not defective; however, as we have noted earlier, the TP domains of the constructions are insensitive to tense. Accordingly, one can state that the TP domain is defective in both types of ECM. At this stage, the point that should draw our attention is that although there is a non-defective T in Type B ECMs, these structures are not phases. Özgen & Aydın (2016: 315), based on the tests they conducted on these structures, have stated that although there is a T head with ϕ features, the entire ECM domain is defective and, thus, the structure is not a phase. Accordingly, testing phasehood solely based on Chomsky's (2001) proposal of agreement and Gallego's (2010) proposal of uninterpretable features and obtaining the same result from both structures accepted as phases and those not accepted as phases in terms of CPs does not seem to provide us with a sound framework for discussion.

In our analysis within the CP framework, we noted that views suggest that uninterpretable features are essentially donated from the C head. As a result, it is concluded that the agreement between the verb

moved from *v* to *T* and the subject is positioned here. However, the fundamental problem in such an analysis is that when this view is adopted, no tests can be applied to *v*Ps in terms of agreement and uninterpretable features. This leads to the conclusion that uninterpretable features cannot be directly employed in the phase diagnostics of *v*Ps. Therefore, agreement and uninterpretable features exclude *v*Ps in defining the phase phenomenon.

Finally, let us evaluate the PPs in terms of agreement and uninterpretable features. See the example in (13) containing a bare PP:

- (13) O adam [PP para **için**] her şey-i yap-ar.
 the man money for everything-acc do-aor
 "The man does everything for money."

(Göksel & Kerslake, 2005: 155)

The preposition *için* (for) in (13) does not have any ϕ -features. PP is considered a phase (see Drummond, Hornstein & Lasnik, 2010: 690). In this case, either phase heads do not need to bear ϕ -features, or the situation we observed may be related to the preposition "için". To determine which option is plausible, we need to test with a different PP. Now, let us examine the following example:

- (14) [PP Saat iki-ye **kadar**] çalış-tı-k.
 o'clock two-dat until work-past-1pl
 "We worked until two o'clock."

(Göksel & Kerslake, 2005: 158)

The postposition *kadar* (until) in example (14) is also a P head, yet it does not bear any ϕ -features. Not only does it lack ϕ -features, but the P heads in the examples above do not enter into agreement either. However, there are also studies in the literature suggesting that there are P heads that do enter into agreement relationships. See (15) below:

- (15) Ben araba-nın_i dün [PP *t_i* önünde] dur-du-m.
 I car-gen yesterday in front of stand-past-1sg
 "I stood in front of the car yesterday."

(Şener, 2006: 87)

Bošković (2014: 9) claims that *önünde* (in front of) is a postposition and is in an agreement relationship. However, another issue that should draw our attention here is whether *önünde* is really a postposition. First of all, to see what types of words the -DA morpheme attaches to, let us examine the following example:

- (16) ev-de kal-dı
 home-loc stay-past

In the example above, the word "ev" (home) is a noun that has taken on the morpheme "-DA." This indicates that the "-DA" morpheme

can attach to nouns. Since case markers are added to nominals in a paradigmatic manner, we can conclude that the phrase "arabanın önünde" (meaning "in front of the car") in example (15) is not a postpositional phrase but rather a Determiner Phrase (DP).

Since the agreement diagnostic proposed by Chomsky (2001) and the uninterpretable features diagnostic proposed by Gallego (2010) are structures not present in PPs in Turkish, as can be seen from our PP analysis above, we cannot employ agreement and uninterpretable features as phase diagnostics concerning PPs. Accordingly, using agreement and uninterpretable features as phase diagnostics also creates a discussion framework governed by independent factors such as the status of the preposition as a noun, or when the word it accompanies is a pronoun or a noun. Another conclusion we can draw is that the concept of the phase may not be as strict as accepted, or there may not be a phase phenomenon at all. To say, there should be a phase phenomenon where agreement is performed optionally and where the non-valuation of uninterpretable features does not cause a crash in the derivation. However, such a phase phenomenon is contrary to the general understanding and the nature of the concept.

2.2 Case as a diagnostic for phases

Miyagawa (2011) argued that phase heads are determined by case features, proposing that the case feature is a more fundamental property in determining phase heads. In (17) below, if case assignment were a result of agreement, we would not expect any case assignment in DPs that we have identified as lacking agreement:

- (17) Ben-im ev
I-gen home
"my house"

The genitive case in (17) is assigned despite the absence of any agreement morpheme. See (18), where a complex DP exists:

- (18) a. [_{DP} Böcekler-in [_{PredP} ev-I istila-sı]]
bugs-gen home-acc invasion-poss.3sg
görenleri şok et-ti.
who.saw shock-past
"Home invasion of insects shocked those who saw it."
- b. [_{DP} Böcekler-in [_{Predp} ev istila-sı]]
bugs-gen home invasion-poss.3sg
görenleri şok et-ti.
who.saw shock-past
"Home invasion of insects shocked those who saw it."

In (18), the accusative case present within the PredP dominated by the DP is not obligatorily present. Its presence only provides a semantic difference in terms of specificity. If the DP head shares its uninterpretable features with the PredP head, and the PredP head licenses the accusative

case, then in the example (18b), the absence of the accusative case should have rendered the sentence ungrammatical. Another issue is that the accusative case makes a difference in terms of specificity. That is, if there is an accusative case licensed by the DP here, and if this case is licensed only when the word "ev" (home) refers to a specific house, this does not provide us with any conclusions in terms of the phasal status. Since the case assignment here varies depending on whether "house" is a definite or a specific house, independently of whether the DP is a phase or not, using case as a phase diagnostic does not provide us with a sound framework for discussion.

Now, in order for us to evaluate the use of verb phrases (vPs) as a phase diagnostic based on case, see the following example:

- (19) Ali bir öğrenci-yi başkan olarak seç-ecek.
 Ali a student-acc president as elect-fut
 "Ali will elect a student as president."

In the example above, the indefinite article "bir" (a) used before "öğrenci" (student) indicates the non-specificity of the student referred to; however, there is still an accusative case assigned to the object by the verb. Observe the contrast in (20a-b):

- (20) a. *Ali bir öğrenci başkan olarak seç-ecek.
 Ali a student president as elect-fut
 "Ali will elect a student as president."
 b. Ali bir öğrenci seç-ecek başkan olarak.
 Ali a student elect-fut president as
 "Ali will elect a student as president."

Comparing the two examples above, although in (20a) the example containing an object with an unlicensed accusative case is ungrammatical, in example (20b) the object with an unlicensed accusative case does not render the sentence ungrammatical. From this, we can deduce that altering the positions of words within a sentence affects the grammaticality independently of the case marking on the object. It indicates that if we employ case as a phase diagnostic, the grammaticality of sentences will change due to independent factors such as specificity, suggesting that the framework for such a discussion would not be reliable. From the perspective of vPs, case assignment, according to the data we have examined, is affected by independent factors such as agreement and syntactic position within the sentence. The fact that the grammaticality of the sentence changes with focus when case is absent, or becomes grammatical through a change of position, is an independent reason that affects case. The same issue applies to PredPs as well. Due to their inability to take complements, no case testing can be performed on PredPs. This shows that using case assignment as a phase diagnostic does not provide a sound framework for discussion.

Let us examine the case diagnostic from the perspective of CPs. The general consensus is the view that among Ouali's (2008) Donate>Keep>Share options, the *share* option operates for CPs, and uninterpretable features are transferred to T. Accordingly, T is accepted as a probe, and subjects are assigned nominative case through feature matching. At this point, there are essentially two problems. The first of these problems is that CPs cannot be tested with respect to Case. If the C phase head activates the T head after feature inheritance and is no longer active itself, all independent tests will be applied to the T head, not to C. Second, Miyagawa (2011) argues that Case and agreement are independent in Turkish. However, in the following example, no agreement occurs in the TP domain:

- (21) *Ben okul-a git-ti.
 I school-dat go-past
 "I went to school."

As we can see from the example above, when the agreement is absent in Turkish, the sentence is ungrammatical. Therefore, the only way we can test whether Case is dependent on agreement in terms of CPs is by using ECM constructions, since a CP structure that can be grammatical despite the absence of agreement is observed only in ECM clauses:

- (22) a. Ali [ben-i okul-a git-ti diye] bil-iyor.
 Ali I-acc school-dat go-past that know-prog
 "Ali thinks that I went to school."
 b. Ali [ben-i okul-a git-ti-m diye] bil-iyor
 Ali I-acc school-dat go-past-1sg that know-prog
 "Ali thinks that I went to school."

Regardless of the presence of agreement, the accusative case appears on "ben" (I). However, the fundamental issue here is that ECM constructions are not accepted as phases. We can test Miyagawa's (2011) claim only with the ECM structure, which is not considered a phase. If we apply a test with a root clause structure that is accepted as a phase, we obtain the following contradiction:

- (23) a. Ali [ben okul-a git-ti-m diye] bil-iyor.
 Ali I school-dat go-past-1sg that know-prog
 "Ali thinks that I went to school."
 b. *Ali [ben okul-a git-ti diye] bil-iyor.
 Ali I school-dat go-past that know-prog
 "Ali thinks that I went to school."

As we can see from the example above, the structure accepted as a phase also contains agreement. The example (23b), where agreement is absent, is ungrammatical. Therefore, from the perspective of CP, the fact that the case hypothesis can only be tested with phrases not considered as

phases poses problems in using this hypothesis as a diagnostic for phasal status. Additionally, the obligatory presence of agreement in Turkish hinders our examination within the CP framework. Again, if we consider this claim in terms of the Donate>Keep>Share options proposed by Ouali (2008), all the tests we conduct occur within the TP framework, and we cannot test the CP. However, case is a syntactic phenomenon that shows where case is valued in the syntax, not at the interfaces. Therefore, using case as a phase diagnostic does not provide us with a sound framework for discussion.

Note that case determines phase heads for another phrase accepted as a phase, i.e. the PP:

- (24) a. [_{PP} sınav için]
 exam for
 “for the exam”
- b. [_{PP} Bira kadar] güzel içecek yok.
 beer as good drink no
 “There is no drink as good as beer.”

Again, in the examples above, there is no case assigned by P; however, there can also be PPs that license case:

- (25) a. [_{PP} sen-in gibi]
 you-gen like
 “like you”
- b. [_{PP} sen-in için]
 you-gen for
 “for you”
- c. [_{PP} Ankara-ya kadar]
 Ankara-dat up to
 “up to Ankara”
- d. [_{PP} Ders-ten önce]
 ders-abl before
 “before the class”

There are multiple conceptual issues here. First, according to the phrases we have examined so far, in Turkish, case seems to be related to agreement. At the same time, PPs in Turkish do not enter into any agreement relations. Therefore, the genitive cases in examples (25a-b) must not have been licensed by the PP. Another issue is that the dative case present in the example (25c) is licensed due to semantic reasons. That is, it is not necessarily a case assignment. In the event of a change in context, the dative case also disappears:

- (26) [_{PP} Ankara kadar] kirli bir şehir yok.
 Ankara as dirty a city no
 “There is no city as dirty as Ankara.”

As we can see from the example above, there are instances where the case is not licensed, and this actually shows that in PPs, the case is licensed due to semantic and contextual factors. Further, see the distinction between demonstrative and personal pronouns from the perspective of PPs:

- (27) a. bu ev
this house
“this house”
- b. *bu-nun ev
this-gen house
“the house of this one”
- c. *bu için
this for
“for this”
- d. bu-nun için
this-gen for
“for his”

PPs behave differently from DPs regarding the case assigned to demonstrative pronouns. While demonstrative pronouns cannot receive cases in DPs, they obligatorily receive cases in PPs. In this context, the case hypothesis will give us different results depending on each sentence, context, and phrase we test. In line with the phrases we have examined above, the fact that case is not always assigned in the same way in every context, with respect to Miyagawa's (2011) claim that case should be a phase diagnostic, indicates that this claim does not provide us with a sound framework for discussion, or suggests that no phrase in Turkish can be a phase. Now, we will examine the examples of extraction constructions in Turkish as a phase diagnostic.

2.3 Extraction as a phase diagnostic

Bošković (2014) has proposed that the extraction process can be used to test phasehood and that, as a result of the test, the heads V, P, and N should also be accepted as phases. The reason underlying this view is that the edge positions of phases are considered escape hatches, and thus, extractions from phases should be allowed.

According to Chomsky (2000), the only way an element can be extracted from a phase is if that element either moves to or is merged at the edge position of the phase. This view dates back to Huang (1982: 505), who defines the Condition on Extraction Domain as follows:

- (28) *Condition On Extraction Domain*

An A-phrase can only be extracted from a domain B if B is properly governed.

Here, by proper government, Huang (1982) means that the extraction domain is a lexical head of A and that A is locally licensed. We can consider this licensing as the head V licensing the complement NP:

- (29) Who_i did you see [a picture of t_i]?

(Stepanov, 2007: 80)

As we can see from the example above, "a picture of who," which is the complement of the verb "see," is licensed by the verb. The fact that it is licensed by the verb shows that the complement "a picture of who" is properly governed, and therefore, extraction can take place. Moving "who" to the beginning of the sentence does not create any ungrammaticality. However, an extraction from other structures, namely subjects and adjuncts, apart from the object, renders the sentence ungrammatical:

- (30) a. ?*Who_i does [a picture of t_i] hang on the wall?
b. ?*Who_i did Mary cry [after Peter hit t_i]?

(Stepanov, 2007: 80)

In example (30a), extraction has been made from the subject, and in example (30b), extraction has been made from within an adjunct. The result of extracting from subjects and adjuncts renders the sentence ungrammatical. As for the extraction structure in Turkish, Jo & Palaz (2019: 22) state that in Turkish, the object can be extracted from the VP domain:

- (31) Ali [şarkı-yı güzel / *güzel şarkı-yı] söyle-di.
Ali the song-acc beautifully beautifully the song-acc sing-past
"Ali sang a song beautifully."

(Jo & Palaz, 2019: 23)

The adverbial "güzel" (beautifully) has been attached to the VP as an adjunct. In order for the DP "şarkı-yı" (the song-ACC) to receive the accusative case, it has been moved to the specifier position of *v*P. This prevents "güzel" (beautifully) from appearing before "şarkı-yı" (the song-ACC) because "şarkı-yı" (the song-ACC) is in a higher position. That is, it shows that it has been extracted from the VP. Although this provides an explanation in terms of phasehood, as we mentioned under the section on case as a phase diagnostic, ungrammaticalities arise due to independent factors in tests involving adverbs and objects that have received accusative case. For example, when an adverb intervening between the verb and an object that has not received the accusative case is focused, the sentence becomes grammatical:

- (32) Ali şarkı güzel söyle-r.
Ali the song beautifully sing-aor
"Ali sings a song beautifully."

As we can see from the example above, when the focused element is the adverb “güzel” (beautifully), the adverb intervening between the verb and an object that has not received the accusative case does not affect the grammaticality of the sentence. Therefore, the grammaticality of the sentence is also determined by focusing the adverb. This suggests that adverbs can occur between verbs and objects that have not received the accusative case, which is contrary to the idea that an object that has received the accusative case has been extracted from the VP. The fact that the sentence is affected by this independent factor means that if extraction is used as a test for the phasehood of the VP, factors such as focus will influence this test.

Ulutaş (2009:11), unlike Jo and Palaz (2019), has focused on extraction from the CP domain:

- (33) a. Dinleyici-ler [____ viski-yi iç-ti]
 listener-pl whiskey-acc drink-past
 san-ıyor-lar biz-i.
 think-prog-3pl we-acc
 “The listeners think that we drank the whiskey.”
- b. ?Dinleyici-ler [____ viski-yi iç-tiğ-imiz]-i
 listener-pl whiskey-acc drink-vnom-poss.1pl-acc
 san-ıyor-lar biz-im.
 think-prog-3pl we-gen
 “The listeners think that we drank the whiskey.”
- c. *Dinleyici-ler [____ viski-yi iç-ti-k]
 listener-pl whiskey-acc drink-past-1pl
 san-ıyor-lar biz.
 think-prog-3pl we
 “The listeners think that we drank the whiskey.”
- d. *Can [____ Manisa-ya gid-elim] isti-yor biz.
 Can ____ Manisa-dat go-sbjnc.1pl want-prog we
 “Can wants us to go to Manisa.”
- e. *Dinleyici-ler [[____ bütün viski-yi iç-tiğ-imiz]
 listener-pl all whiskey-acc drink-vnom-poss.1pl
 için] biz-i sarhoş san-ıyor-lar biz-im.
 for we-acc drunk think-prog-3pl we-gen
 “The listeners think we are drunk because we drank all the whiskey.”

As we can see from the examples above, extraction is permitted in embedded clauses where agreement is present, that is, when the C head is defective—examples (33a-b). However, extraction is not allowed in examples (33c-d-e), where the C head is non-defective. This suggests that the subject extracted from the embedded clause may have been sent to the Spell-Out domain by the complete C. Considering example (33a), since the

accusative case of "biz-i" (we-ACC) is assigned by the verb, it may be possible to extract it adjacent to the main verb:

- (34) *pro* sen-i_i [_{CP} t_i viski iç-ti] bil-iyor-du-m.
pro you-acc whiskey drink-past know-prog-past-1sg
 (Moore, 1998: 169)

In the analysis above, Moore (1998: 169) has stated that the subject "sen-i" (you-ACC) is raised to the main clause as an object. Therefore, subjects that receive the accusative case from the main clause, and thus can be raised to the main clause as objects, can be extracted. That is, they are actually within the main clause. This shows that the reason for the extraction of ECM subjects that have received the accusative case may not be due to phasehood, yet could stem from raising or movement to the main clause.

On the other hand, in none of the examples (33b-c-d-e) can the cases on the extracted subject "biz" (we) be licensed by the main verb. In example (33b), the case is licensed by the D head; in example (33c), by the C head of the embedded clause; and in example (33d), again by the C head of the embedded clause. The presence of cases licensed by these heads may have emerged independently of the main clause. Additionally, in sentence (33e), the extracted "biz-im" (we-GEN) is a DP branched from a PP. Therefore, the presence of the genitive case on "biz" (we) also creates ungrammaticality. Thus, extracting the object of the verb, which is the head that licenses the case, should not cause a problem. However, in examples (33b-c-d-e), the words "biz" ("we") are not in the object position of the verb. If the edge positions of phases are escape positions, we would expect that all the subjects in examples (33) could be extracted. Therefore, another reason affecting the extraction here is the Case. That is, for a word to be extracted to a different position, there needs to be case licensing by the unit in the position to which it is extracted. In examples (33b-c-d-e), the cases on the subjects that cannot be extracted are licensed within the embedded clause. However, in example (33a), the main clause verb licenses the accusative case. This shows us that in example (33a), extraction is triggered by the main clause verb, and that case licensing affects extraction. Therefore, regardless of whether the C head is defective or not, extraction is related to the case licensing by the main clause verb. This indicates that if extraction is accepted as a phase diagnostic, factors independent of phasehood, such as case and position, will influence this diagnostic.

Now, in the light of the information we have examined, see the first test the DPs in Turkish:

- (35) a. [_{DP} Ali-nin deri cüzdan-ı] çal-ın-dı.
 Ali-gen leather wallet-poss.3sg steal-pass-past
 "Ali's leather wallet was stolen."

- b. *Deri_i [DP Ali-nin *t_i* cüzdan-ı] çal-ın-dı.
 leather Ali-gen wallet-poss.3sg steal-pass-past
 “Ali’s leather wallet was stolen.”

Regarding the examples above, in example (35b), it seems that extraction is not possible; however, this may be due to an independent reason, namely, that the position to which it is moved is a position that does not allow extraction:

- (36) a. ?[DP Ali-nin *t_i* cüzdan-ı] çal-ın-dı deri_i.
 Ali-gen wallet-poss.3sg steal-pass-past leather
 “Ali’s leather wallet was stolen.”
- b. [DP Ali-nin yakışıklı arkadaş-ı] gel-di.
 Ali-gen handsome friend-poss.3sg come-past
 “Ali’s handsome friend arrived.”

Regarding the examples above, firstly, moving the adjective to a position after the verb, as in example (36a), is not as ungrammatical as in sentence (36b). Therefore, performing the extraction into a post-verbal position or to the sentence-initial position yields different grammaticality results. Similarly, another factor affecting extraction is agreement:

- (37) a. [DP Ali-nin sınıf] okul-dan kaç-tı.
 Ali-gen class school-abl escape-past
 “Ali’s class escaped from school.”
- b. *Sınıf_i [DP Ali-nin *t_i*] okul-dan kaç-tı.
 class Ali-gen school-abl escape-past
 “Ali’s class escaped from school.”
- c. *[DP Ali-nin *t_i*] okul-dan kaç-tı sınıf_i.
 Ali-gen school-abl escape-past class
 “Ali’s class escaped from school.”
- d. [DP Ali-nin *t_i*] okul-dan kaç-tı sınıf-ı_i.
 Ali-gen school-abl escape-past class-poss.3sg
 “Ali’s class escaped from school.”

There is again an independent factor affecting extraction. In examples (37b) and (37c), where agreement is not explicitly encoded, the sentence becomes ungrammatical regardless of the position to which the extraction is made. However, as we can see from example (37d), when the agreement is explicitly present, the sentence remains grammatical after extraction. Therefore, the explicit coding of agreement affects both extraction and the grammaticality of the sentence. This shows that independent factors, such as the explicit encoding of case and agreement, will influence the test if we use extraction as a phase diagnostic.

Now, let us perform the same tests within the CP framework. Since the ECM structure is used as a phase diagnostic, we need to examine the extraction phenomenon in both ECMs and root clauses:

- (38) a. Polis [_{CP} sen-i kimse-ye vur-du] bil-mi-yor.
 police you-acc anyone-dat hit-past know-neg-prog
 “The police do not know that you shot anyone.”
- b. Sen-i_i polis [_{CP} t_i kimse-ye vur-du] bil-mi-yor.
 you-acc police anyone-dat hit-past know-neg-prog
 “The police do not know that you shot anyone.”
- (Özgen & Aydın, 2016: 311)

In the examples above, sentence (38a) is an ECM sentence. As we can see from (38b), extraction is allowed. Observe the contrast in the grammaticality between ECM construction in (38b) & (39b):

- (39) a. Ali [_{CP} sen İstanbul-a gid-iyor-du-n] san-mış.
 Ali you İstanbul-dat go-prog-past-2sg think-evid
 “Ali thought that you were going to Istanbul.”
- b. ?*Sen-i_i Ali [_{CP} t_i İstanbul-a gid-iyor-du-n] san-mış.
 you-acc Ali İstanbul-dat go-prog-past-2sg think-evid
 “Ali thought that you were going to Istanbul.”
- c. ?Ali [_{CP} t_i İstanbul-a gid-iyor-du-n] san-mış sen-i_i.
 Ali İstanbul-dat go-prog-past-2sg think-evid you-acc
 “Ali thought that you were going to Istanbul.”

Comparing the examples above, we can once again observe the effect of moving the extraction to a position after the verb on grammaticality. Therefore, when we apply the same test to two structures, one accepted as a phase and one not accepted as a phase, we see that we can obtain similar results depending on the position of the extraction. This leaves us with two options. First, using Bošković's (2014) extraction hypothesis as a phase diagnostic would create an unhealthy framework for discussion; second, CPs in Turkish cannot be phases, which would be an undesired outcome.

The extraction does not seem to work well from the perspective of vPs:

- (40) a. Ali kitab-ı oku-du.
 Ali book-acc read-past
 “Ali read a book.”
- b. Kitab-ı Ali oku-du.
 book-acc Ali read-past
 “Ali read a book.”
- c. Ali kitap oku-du.
 Ali book read-past
 “Ali read a book.”

- d. Kitap Ali oku-du.
 book Ali read-past
 “Ali read a book.”

In terms of transitive *vP* structures, the assignment of accusative case does not have any effect on the extraction process. Note that Turkish is a language with flexible word order (Göksel & Kerslake, 2005; Kornfilt, 1997), which means that all the tests we have conducted so far may have been permitted due to PF-driven displacement. Therefore, we are once again faced with two options. The first is that applying the extraction diagnostic may not be possible in Turkish because the free word order in Turkish generally allows words to change positions. The other option is that if we accept Bošković’s (2014) extraction claim as a phase diagnostic notwithstanding these independent factors, we are forced to conclude that all phrases and syntactic structures in Turkish are phases, which would be another unwelcome result.

Bošković’s (2014: 9) final claim regarding extraction is that there are two different PP structures in Turkish. PP is a phase if it allows extraction; otherwise, it is not a phase:

- (41) a. *Biz [_{NP} Pelin-in arkadaş-ı]_i
 we Pelin-gen friend-poss.3sg
 dün [_{PP} *t_i* için] para topladı-k.
 yesterday for money collect-past-1pl
 “We collected money yesterday for Pelin’s friend.”
 b. Ben araba-nın_i dün [_{PP} *t_i* önünde] dur-du-m.
 I car-gen yesterday in front of stop-past-1sg
 “I stood in front of the car yesterday.”

The locative morpheme -DA attaches to the nominal paradigm, as mentioned earlier. Thus, the alleged PP in example (41b) is not a PP but a DP. Accordingly, as in example (41a), extraction cannot be made from genuine bare PPs in Turkish. This again leads us to accept that PPs in Turkish are not phases.

Considering that Bošković’s (2014) extraction hypothesis cannot be employed as a phase diagnostic in Turkish due to the scrambling property of Turkish and the fact that extraction cannot be performed in PP structures, we should deduce either that there is no concept of phase in Turkish or that Bošković’s (2014) extraction hypothesis is not a phase diagnostic.

2.4 Ellipsis as a phase diagnostic

Another tool used as a phase diagnostic in the literature is ellipsis (*see* Bošković, 2014; Rouveret, 2012; Gengel, 2007; Bošković, 2012). İnce (2012: 248) argues that, contrary to the common belief, ellipsis can occur in languages like Turkish, even though they do not use *wh*-movement. Since the sluicing is known to be TP deletion, in languages like English,

the wh-word moves to CP, and the TP is deleted. However, Turkish is a language that does not obligatorily perform wh-movement (Arslan, 1999):

- (42) a. Hasan biri-yile konuş-uyor.
 Hasan someone-with talk-prog
 ama Hasan kim-le konuş-uyor bil-mi-yor-um.
 but Hasan who-with talk-prog know-neg-prog-1sg
 ‘‘Hasan is talking with someone, but Hasan doesn't know with
 whom Hasan is talking.’’
- b. Hasan biri-yile konuş-uyor.
 Hasan someone-with talk-prog
 ama kim-le bil-mi-yor-um.
 but whom-with know-neg-prog-1sg
 ‘‘Hasan is talking to someone, but I don't know with whom.’’

(İnce, 2012: 250)

In example (42a), ‘‘Hasan’’ and ‘‘konuşuyor’’ (talk-prog) have been elided in example (42b). However, since Turkish does not undergo wh-movement, after the ellipsis operation, the entire clause ‘‘Hasan kimle konuşuyor’’ (‘‘Hasan is talking with whom’’) should have been deleted. At this point, İnce (2012: 252) proposed that Turkish actually does perform wh-movement, but due to the weak wh-question feature in the C head, the lower copy is pronounced instead. In other words, through the remnant movement phenomenon we mentioned earlier, ‘‘kim-le’’ (whom-with) first moves to the specifier position of CP, leaving a copy in its original position. However, because the wh-question feature at the specifier of CP is weak, it's not the wh-word in the specifier of CP that is pronounced, but the lower copy of the wh-word. İnce (2012) suggested in this regard that movement in Turkish occurs not at the LF but at the PF level. On the other hand, there are also views suggesting that this structure in Turkish is a cleft construction, and that ellipsis does not occur in languages that do not perform wh-movement:

- (43) *The case feature of the elided wh-word must be the same as that of the main clause.*

(Merchant, 1999: 48)

That is, in accordance with the constraint above, the case feature in the ellipsis structures we construct must be identical to the case feature in the antecedent clause. Now, see these examples:

- (44) a. Ahmet birin-i döv-müş
 Ahmet someone-acc beat-evid
 ama kim-i bil-mi-yor-um.
 but who-acc know-neg-prog-1sg
 ‘‘Ahmet has beaten someone, but I don't know whom.’’

- b. Ahmet birin-e kitap ver-miş-ti
 Ahmet someone-dat book give-evid-past
 ama kim-e-ydi hatırla-mı-yor-um.
 but who-dat-past remember-neg-prog-1sg
 "Ahmet had given a book to someone, but I don't remember to whom it was."

(İnce, 2012: 255)

As we examine the above examples, in example (44a), the accusative case marked on "kim-i" (who-ACC) is identical to "biri-ni" (someone-ACC) in the main clause. Changing the case on "kim" (who) yields ungrammaticality. In example (44b), the case marked on "kim" is the same as "biri" ("someone") in the main clause. Thus, it is insufficient for us to observe the difference between stripping and ellipsis in terms of case. However, we can see this difference when using the copula "ol-" (to be). In Turkish, while the copula "ol-" (to be) can be used in stripping constructions, using this copula in ellipsis constructions leads to ungrammaticality:

- (45) a. Ahmet-in borç ver-diğ-i-nin
 Ahmet-gen debt give-vnom-poss.3sg-gen
 Hasan ol-duğ-un-u bil-iyor-um.
 Hasan be-vnom-acc-poss.3sg know-prog.1sg
 "I know that it is Hasan to whom Ahmet lent money."
 b. *Araba-yı onar-dı ama
 car-acc repair-past but
 nasıl ol-duğ-un-u bil-mi-yor-um.
 how be-vnom-poss.3sg-acc know-neg-prog.1sg
 "He repaired the car, but I don't know how it happened."

(İnce, 2012: 11)

In example (45a) the copular verb "ol-" (be) can be used in a stripping sentence, in example (45b), which involves an ellipsis construction, the copular verb "ol-" (to be) results in ungrammaticality.

In English VP-ellipsis, the elided structure must be structurally identical to the main clause. First, observe examples involving objects that have received the accusative case and those that have not, in relation to the ellipsis operation:

- (46) a. Ali kitap oku-du, Ayşe de ~~kitap oku-du~~.
 Ali book read-past Ayşe too book read-past
 "Ahmet read a book, and Ayşe read too."
 b. Ali kitab-ı oku-du, Ayşe de ~~kitab-ı oku-du~~.
 Ali book-acc read-past Ayşe too book-acc read-past
 "Ahmet read a book, and Ayşe read too."

The unit "kitab" (book) in (46a) that has received the accusative case is in the VP domain. In contrast, the unit "kitab-ı" (the book-acc) in (46b) that has received the accusative case is in the specifier position of the *v*P. Bošković (2014) has stated that only the complements of phase heads (i.e., spell-out domains) or entire phases can be elided. In the numbered examples (46) we analyzed above, Bošković (2014) argues that the V head is a phase, and that the fact that the object in its complement and the entire VP can be elided demonstrates the phasehood of the VP. Note that, in example (46a), the entire VP (= spell-out domain) is elided; whereas in example (46b), the entire *v*P (= phase) is elided. In fact, the ellipsis here suggests the phasehood of the *v*P, not the VP. If we accept the view that the VP is a phase and that only the complement of the V head and the VP are elided, we cannot explain the elision of "kitabı" (the book), which has been moved to the specifier of the *v*P in example (46b).

Another issue at this point is that the tense marker on the verb "oku-du" (read-past) is also elided. We have previously noted that in Turkish, verbs must raise to the T head to acquire tense and agreement features. Unlike in English, where tense can be carried by an auxiliary, in Turkish, it is directly marked on the main verb. Consequently, if the verb "oku-du" (read-past) displays tense marking, it indicates that the verb has raised to T. Therefore, if the part of the structure undergoing ellipsis also includes the tense morphology from T, then the domain of ellipsis must encompass T. This implies that the construction in (46) cannot be mere VP-ellipsis. Instead, a larger syntactic domain (including T) is elided. Thus, examples like those in (46) do not constitute a straightforward VP-ellipsis case.

In terms of DPs see this example:

- (47) [DP [AGRP Ben-im ev-im] sat-ıl-dı
 I-gen house-poss.1sg sell-pass-past
 ama [DP [AGRP Ayşe-nin ~~ev-i~~] sat-ıl-ma-dı.
 but Ayşe-gen house-poss.3sg sell-pass-neg-past
 "My house was sold, but Ayşe's house was not sold."

In (47) we are dealing with a structure that poses a problem for ellipsis. First, within the DP, the fact that the element "ev-i" (the house-ACC), which has raised to the head of the AGRP, can be elided while "Ayşe'nin" (Ayşe-GEN) remains unelided does not clearly indicate which exact domain has been elided. If we also consider the possibility that the element "Ayşe'nin" (Ayşe-GEN) has moved out of the AGRP to another position, then extracting and eliding the genitive-case possessor units in the DPs becomes impossible:

- (48) *[DP [AGRP Ben-im ev-im] sat-ıl-dı
 I-gen house-poss.1sg sell-pass-past
 ama [DP [AGRP ~~Ayşe-nin~~ ~~ev-i~~] sat-ıl-ma-dı.
 but Ayşe-gen house-poss.3sg sell-pass-neg-past
 "My house was sold, but Ayşe's house was not sold."

As we can see from the example above, in example (48), “Ayşe-nin” (Ayşe-GEN) was not elided because it moves to a higher position, but rather that the elements in the specifier positions of the AGRP are not identical. Accordingly, it is evident that independent factors, such as the use of different words within the phrase, also influence the ellipsis operation.

Now see apply the same test to sentential DPs as well:

- (49) a. [DP Ali-nin [CP kereviz ye-diğ-i]] yalan ama
 Ali-gen celery eat-vnom-poss.3sg lie but
 [DP Ayşe-nin [CP ~~kereviz ye-diğ-i~~]] gerçek.
 Ayşe-gen celery eat-vnom-poss.3sg real
 “It is a lie that Ali ate celery, but it is true that Ayşe ate celery.”
- b. [DP Ali-nin [CP kereviz ye-diğ-i]] duy-ul-muş
 Ali-gen celery eat-vnom-poss.3sg hear-pass-evid
 ama [DP Ali'nin [CP ~~kereviz ye-diğ-i~~]]
 but Ali-gen celery eat-vnom-poss.3sg
 hiç gör-ül-me-miş.
 ever see-pass-neg-evid
 “Ali’s having eaten celery has been heard, but Ali’s having eaten celery has never been seen.”
- (Özgen, 2018: 13)
- c. [DP Ali-nin [CP kereviz ye-diğ-i]] duyulmuş
 Ali-gen celery eat-vnom-poss.3sg hear-pass-evid
 ama [DP Ali'nin [CP ~~kereviz in~~]] turşu-sun-u
 but Ali-gen celery-gen pickle-poss.3sg-acc
 kur-duğ-u hiç gör-ül-me-miş.
 set-vnom-poss.3sg ever see-pass-neg-evid
 “It has been heard that Ali ate celery, but it has never been seen that Ali pickled the celery.”

In the sentential DP types above, we encounter a similar situation. If at this stage we accept the view that in certain cases an entire DP can be elided while in other cases only its complement can be elided, then, considering that the elided portions differ depending on the specifier element even though they appear in the same structure, we run into a problematic discussion in terms of phasehood. Further, if we observe the example in (49c), we see that the elided portion differs from its antecedent, suggesting that independent factors may influence the ellipsis process. Consequently, using ellipsis to test the phasehood of DPs does not provide us with a reliable framework because ellipsis may fail to occur for reasons unrelated to phasehood. As in example (49), it may be unclear precisely which domain is elided.

CPs, first, in order to see whether there is a difference in terms of ellipsis between the ECM structure, which is not accepted as a phase, and the CP, which is accepted as a phase, we need to apply this test to ECM structures as well:

- (50) a. *Murat [_{CP} ben-i öl-dü] san-ıyor,
 Murat I-acc die-past think-prog
 Özgün-se [_{CP} ben-i ~~öl-dü~~] san-mı-yor.
 Özgün-while I-acc die-past think-neg-prog
 “Murat thinks I am dead while Özgün does not think I am dead.”
- b. *Murat [_{CP} ben öl-dü-m] san-ıyor,
 Murat I die-past.1sg think-prog
 Özgün-se [_{CP} ben ~~öl-dü-m~~] san-mı-yor.
 Özgün-while I die-past.1sg think-neg-prog
 “Murat thinks I am dead while Özgün does not think I am dead.”

With regard to both of the above examples, we observe that in the ECM construction in (50a) and in the root clause structure in (50b), the sentences become ungrammatical if the verb is deleted. However, in both the ECM construction and the root clause, the entire CP can be elided:

- (51) a. Murat [_{CP} ben-i öl-dü] san-ıyor,
 Murat I-acc die-past think-prog
 Özgün-se [_{CP} ~~ben-i öl-dü~~] san-mı-yor.
 Özgün-while I-acc die-past think-neg-prog
 “Murat thinks I am dead while Özgün does not think I am dead.”
- b. Murat [_{CP} ben öl-dü-m] san-ıyor,
 Murat I die-past.1sg think-prog
 Özgün-se [_{CP} ~~ben öl-dü-m~~] san-mı-yor.
 Özgün-while I die-past.1sg think-neg-prog
 “Murat thinks I am dead while Özgün does not think I am dead.”

From example (51a), we can see that the entire CP can be elided in the ECM construction, which is not accepted as a phase. At this stage, we would expect the root clause structure in (51b), which is accepted as a phase, to present a contrasting scenario; however, as we can observe from the examples, the entire CP can be elided in both constructions. Therefore, using ellipsis to test the phasehood of the CP does not provide a sound framework for discussion.

2.5 Sentential stress rule as a phase diagnostic

In the literature, the sentential stress rule is another structure recognized as a phase diagnostic. Legate (2003), following Bresnan

(1972), argues that the Sentential Stress Rule (SSR) must be used in its phonetic form to demonstrate the phasehood of the VP, functioning as a phase diagnostic. Similarly, Kahnemuyipour (2004) states that the SSR constitutes evidence for the presence of multiple spell-out domains within phases. Focusing first on Legate (2003:511), she observes that in English, the primary stress is assigned to the unit containing the final stress within VPs:

(52) Mary fixed the bike¹.

(Legate, 2003: 511)

In the example above, the points marked with ¹ indicate the stress.

Another study that uses the SSR as a phase diagnostic is found in Kahnemuyipour (2004):

(53) *Sentential Stress Rule*

Sentential stress is assigned to the highest element in the spell-out domain within a phase.

Regarding Turkish, Tarhan (2006:20) has noted that the same situation applies in Turkish as well:

(54) Ali [hızlı kitap oku-du].
Ali quickly book read-past
“Ali read a book quickly.”

(Tarhan, 2006: 20)

First, in the example above, we need to determine whether the stress on “hızlı” (quickly) arises from its position in the sentence or if it is related to the fact that “hızlı” is functioning as an adverb:

(55) Ali kitab-ı hızlı oku-du.
Ali book-acc quickly read-past
“Ali read a book quickly.”

When we change the position of the adverb “hızlı” the stress it carries remains unchanged. Therefore, the first conclusion we can draw is that the stress on the adverb is not affected by its position.

Now, observe the focus on unergative and unaccusative sentences:

(56) a. A: Sabah ne ol-du?
morning what happen-past
“What happened in the morning?”

B: Ali koş-tu.
Ali run-past
“Ali ran.”

b. A: Haber-ler-i duy-du-n mu?
 news-3pl-acc hear-past-2sg Q
 “Have you heard the news?”

B: Hayır. Ne ol-muş?
 No what happen-evid
 “No. What happened?”

A: **Bir gemi** bat-mış.
 A ship sink-evid
 “A ship has sunk.”

(Tarhan, 2006: 50-51)

As for the unergative example in (56a), Tarhan (2006: 50) states that the stress falls on the verb “koştu” (“ran”). According to Tarhan (2006), this is because in an unergative structure, the subject remains in vP, making “koş-” (run) the highest element. However, since there is agreement between “Ali” and “koş-tu” (run-past) and “Ali” is in the nominative case, “Ali” must have moved up to the specifier of T. Under these conditions, the highest element in the spell-out domain should not be “koştu,” but rather “Ali.” Despite that, we observe that “koş-tu” (run-past) bears stress, even though it is not actually the highest element in the domain. Therefore, according to this test, unergative verbs commonly accepted in the literature as phases cannot be phases after all. By contrast, in the unaccusative example (56b), which is not accepted as a phase, “bir gemi” (a ship) is the highest element in the spell-out domain. We see that stress falls on “bir gemi,” showing once again that the Sentential Stress Rule can yield similar results in structures not considered phases and can even give contrary results in structures that are considered phases.

3. Concluding Remarks

Up to this point, we have examined whether the structures viewed as phase diagnostics can be explained without resorting to the concept of phasehood, and whether these structures are affected by factors independent of phasehood. Based on the data we obtained and the comparisons we made, we found that these diagnostics do not work in Turkish. Let us look at this situation through the table below:

Table 1. The applicability of phase diagnostics in Turkish

PHASE DIAGNOSTICS	TURKISH
Chomsky (2000) Agreement determines phase heads.	×
Gallego (2010) Uninterpretable features determine phase heads.	×
Miyagawa (2011) Case determines phase heads.	×
Den Dikken (2003) Phases can be tested via the WH-movement structure.	Not applicable
Bošković (2014) Phases can be tested via deletion.	×
Bošković (2014) Phases can be tested via extraction and stranding.	×
Matushansky (2005) & Legate (2003) Phases can be tested via stranding.	×
Legate (2003) Phases can be tested via sentential stress rules.	×
Den Dikken (2006) Phases can be tested via WH-fronting.	×

From the table above, we can see that none of the structures proposed in the literature for testing the concept of phase are applicable to Turkish. This observation leads us to several different conclusions. The contradictory results that emerged from our analyses point to three different possibilities:

1. If the testing methods we have examined are considered valid phase diagnostics, then it becomes impossible to identify any phrase as a phase in Turkish or in other languages with similar typological features. In other words, these purported phase diagnostics may not actually test for phasehood in a general sense.
2. If these phase diagnostics are accepted, then no phrase in Turkish or in any language with similar typological characteristics should be considered a phase.
3. The phenomenon called 'phase' is, in fact, a misconception.

When supported by crosslinguistic data, it will become clearer which of these options is more plausible. Likewise, further research on

phasehood could offer different perspectives on the universality of phase theory. If, as a result of such research, phasehood proves to be illusory, then a more comprehensive theory that yields consistent findings despite varying typological features of languages could be proposed instead of phase theory.

References

- Arslan, Z. C. (1999). *Approaches to wh-structures in Turkish* (Doctoral dissertation). Boğaziçi University.
- Bošković, Ž. (2012). On NPs and clauses. In G. Grewendorf & T. E. Zimmermann (Eds.), *Discourse and grammar: From sentence types to lexical categories* (pp. 179–245). Berlin, Germany: De Gruyter.
- Bošković, Ž. (2014). Now I'm a phase, now I'm not a phase: On the variability of phases with extraction and ellipsis. *Linguistic Inquiry*, 45(1), 27–89. https://doi.org/10.1162/LING_a_00148
- Chomsky, N. (2000). Minimalist inquiries: The framework. In R. Martin, D. Michaels, & J. Uriagereka (Eds.), *Step by step: Essays on minimalist syntax in honor of Howard Lasnik* (pp. 89–155). Cambridge, MA: MIT Press.
- Chomsky, N. (2001). Derivation by phase. In M. Kenstowicz (Ed.), *Ken Hale: A life in language* (pp. 1–52). Cambridge, MA: MIT Press.
- Citko, B. (2014). *Phase theory: An introduction* (Research Surveys in Linguistics). Cambridge, England: Cambridge University Press.
- Den Dikken, M. (2006). *A reappraisal of vP being phasal: A reply to Legate*.
- Drummond, A., Hornstein, N., & Lasnik, H. (2010). A puzzle about P-stranding and a possible solution. *Linguistic Inquiry*, 41(4), 689–692.
- Gallego, Á. J. (2010). *Phase theory*. Amsterdam, Netherlands: John Benjamins Publishing Company.
- Gengel, K. (2007). *Focus and ellipsis: A generative analysis of pseudogapping and other elliptical structures* (Doctoral dissertation, University of Stuttgart).
- Göksel, A., & Kerslake, C. (2005). *Turkish: A comprehensive grammar*. London, England/New York: Routledge.
- Huang, J. T. C. (1982). Move wh in a language without wh movement. In H. Lasnik & H. Safir (Eds.), *An annotated syntax reader* (pp. 151–169). Oxford, England: Wiley-Blackwell.
- Ince, A. (2012). Sluicing in Turkish. In J. A. van Craenenbroeck & T. Temmerman (Eds.), *Sluicing: Cross-linguistic perspectives* (pp. 248–269). Oxford, England: Oxford University Press.
- Jo, J., & Palaz, B. (2019a). Licensing pseudo-incorporation in Turkish. In M. Baird & J. Pesetsky (Eds.), *Proceedings of the 49th meeting of the North East Linguistic Society* (pp. 155–164). Amherst, MA: GLSA.
- Kahnemuyipour, A., & Massam, D. (2004). Deriving the order of heads and adjuncts: The case of Niuean DPs. *ZAS Papers in Linguistics*, 34, 135–147.
- Kornfilt, J. (1997). *Turkish*. London, England/New York: Routledge.

- Legate, J. A. (2003). Some interface properties of the phase. *Linguistic Inquiry*, 34(3), 506–515.
- Matushansky, O., McGinnis, M., & Richards, N. (2005). Going through a phase. In S. Epstein & T. D. Seely (Eds.), *Perspectives on phases* (pp. 157–183). Cambridge, MA: MIT Working Papers in Linguistics.
- Miyagawa, S. (2011). Genitive subjects in Altaic and specification of phase. *Lingua*, 1217, 1265–1282.
- Moore, J. (1998). Turkish copy-raising and A-chain locality. *Natural Language & Linguistic Theory*, 16, 149–189.
- Merchant, J. (1999). *The syntax of silence: Shuicing, islands, and identity in ellipsis* (Doctoral dissertation, University of California, Santa Cruz).
- Ouali, H. (2008). On C-to-T phi-feature transfer: The nature of agreement and anti-agreement in Berber. In R. D'Alessandro, G. H. Hrafnbjargarson, & S. Fischer (Eds.), *Agreement restrictions* (pp. 159–180). Berlin, Germany: Mouton de Gruyter.
- Ouhalla, J. (1993). Subject-extraction, negation and the anti-agreement effect. *Natural Language and Linguistic Theory*, 11, 477–518.
- Özgen, M., & Aydın, Ö. (2016). What type of defective feature do exceptionally case-marked clauses of Turkish bear? *Open Journal of Modern Linguistics*, 6(4), 302–316.
- Özgen, M. (2018). Phasehood of DPs in Turkish: An implication for non-simultaneity. *Dokuz Eylül University Journal of Faculty of Letters*, 5(2), 1–31.
- Öztürk, B., Taylan, E. E., & Zimmer, K. (2015). Possessive-free genitives in Turkish. In D. Zeyrek, Ç. S. Şimşek, U. Ataş, & J. Rehbein (Eds.), *Ankara papers in Turkish and Turkic linguistics* (pp. 189–204). Wiesbaden, Germany: Harrassowitz Verlag.
- Rouveret, A. (2012). VP ellipsis, phases and the syntax of morphology. *Natural Language & Linguistic Theory*, 30, 897–963.
- Stepanov, A. (2007). The end of CED? Minimalism and extraction domains. *Syntax*, 10(1), 80–126.
- Şener, S. (2006). Strandability of a p is about its extendability [Handout from the Workshop on Antilocality]. Cambridge, MA: Harvard University.
- Şener, S. (2008). Non-canonical case licensing is canonical: Accusative subjects of CPs in Turkish [Manuscript, University of Connecticut]. Storrs, CT.
- Ulutaş, S. (2009). Feature inheritance and subject case in Turkish. In *Essays on Turkish linguistics: Proceedings of the 14th International Conference on Turkish Linguistics* (pp. 141–151). Wiesbaden, Germany: Harrassowitz Verlag.
- Üntak Tarhan, F. A. (2006). *Topics in syntax-phonology interface in Turkish: Sentential stress and phases* (Master's thesis, Boğaziçi University). Istanbul, Turkey.