



Atf (cite). RESULOĞLU, Erkut. (2026) Üretken Büyük Dil Modellerinde Türkçe Ses Bilgisi Terimlerinin Temsili, *Disiplinler Arası Dil Araştırmaları Dergisi*, II. Uluslararası Dil Araştırmaları Sempozyumu Özel Sayısı (Türkçe Sesbilgisi Terminolojisi), 2, 543-567. doi: <https://doi.org/10.48147/dada.1990460>

Üretken Büyük Dil Modellerinde Türkçe Ses Bilgisi Terimlerinin Temsili¹

Representation of Turkish Phonological Terms in Generative Large Language Models

Erkut RESULOĞLU²

Özet

Çağdaş araştırmalar, büyük dil modellerinin özellikle doğal dil işleme ve dil öğretimi alanlarındaki kullanımına, dilsel görevlerdeki performansına ve teknolojik gelişmelerin model başarıları üzerindeki etkisine yoğunlaşmaktadır. Dilin bilgisayarlarca işlenebilmesi için seslerin tanımlanması, modellenmesi ve sayısallaştırılması gerekmektedir. Bu çalışmanın amacı, büyük dil modellerinin ses bilgisi terimlerini tanım düzeyinde çözümleme becerilerini ortaya koymaktır. Dil modellerinin ünlüleri, ünsüzleri ve ses olaylarını adlandırmaları, kıyaslama veri kümeleri aracılığıyla değerlendirilmiştir. Araştırmada kullanılan üç farklı kıyaslama veri kümesi açık uçlu ve çoktan seçmeli soru türlerini barındırmaktadır. Ses bilgisi terimleri, Ollama çalışma ortamında bulut tabanlı büyük dil modelleri üzerinden yanıtlar aracılığıyla elde edilmiştir. Bulgular, dil modellerinin terminolojik açıdan büyük çeşitlilik gösterdiğini, tanımlarda ayrıntıya ve alt terimlendirmelere yönelme eğilimi sergilediğini göstermektedir. Sonuçlar, büyük dil modellerinin dili çözümleme becerisine sahip olduğunu ancak terminolojik standartlaşma konusunda hem model bazında hem de modeller arasında ortaklık sağlayamadığını göstermektedir.

Anahtar Sözcükler: Büyük dil modelleri (BDM), üretken yapay zekalar, Türkçe ses bilgisi, metin tabanlı modeller, terminoloji.

Abstract

Contemporary research focuses on the use of LLMs, particularly in fields of natural language processing and language teaching, their performance in linguistic tasks and the impact of technological advancements on model success. For language to be processed by computers, it is necessary to identify, model and digitize sounds. The purpose of this study is to demonstrate the ability of LLMs to analyze phonological terms at the definition level. The naming of vowels, consonants and phonological processes by LLMs was evaluated through benchmark datasets. Three different benchmark datasets used in the study include open-ended and multiple-choice question types. Phonological terms were obtained through responses from cloud-based LLMs in the Ollama working environment. The findings indicate that LLMs show great diversity in terms of terminology, exhibit a tendency toward detail and sub-categorization of terms in definitions. The results show that LLMs possess the ability to analyze language, but they fail to achieve terminological standardization, both on a model-by-model basis and across models.

Key Words: Large language models (LLMs), generative artificial intelligence, Turkish phonology, text-based models, terminology.

Makale Türü: Araştırma

Paper Type: Research

¹ Bu makale, Eskişehir Osmangazi Üniversitesi, Sözlükbilimi Uygulama ve Araştırma Merkezi tarafından 11-12 Haziran 2026 tarihinde düzenlenen II. Uluslararası Dil Araştırmaları Sempozyumu'nda sunulan bildirinin gözden geçirilmiş hâlidir.

² Bilim Uzmanı, Bağımsız Araştırmacı, erkutresul@gmail.com, ORCID No: (0000-0002-2391-9525)

1. Giriş

Dil, bir iletişim aracı olmanın ötesinde, yapay zekâ teknolojilerindeki gelişmelere bağlı olarak bilgisayarlarca işlenmekte ve modellenmektedir. Dilin sosyokültürel yapısı ve sürekli değişen doğası, dijital mecralardaki kullanım biçimlerine de yansımaktadır. Bu bağlamda, dilin farklı unsurları, nitelikli kaynakların dijital ortama aktarılmasıyla yapay zekâ çalışmalarının odağı hâline gelmiştir.

Bu noktada, dil biliminin alt disiplinleri temel alınarak dilin sayısallaştırılması çabası ortaya çıkmıştır. Dilin sayısal olarak işlenebilmesi için de doğal dil işleme (NLP) modelleri geliştirilmiştir. Doğal dil işleme modelleri, makinelerin insan dilini anlaması, sayısallaştırması ve üretmesi üzerine odaklanan yapay zekânın bir alt dalı olarak karşımıza çıkmaktadır (Sriharsha vd., 2024: 507). İnsan dilinin sosyolengüistik yapısı, anlamlandırılması ve dildeki kalıpların modellenmesi gibi çeşitli zorluklarla başa çıkmayı amaçlayan doğal dil işleme modelleri, metin madenciliğinden ses tanımlamaya kadar bünyesinde çok çeşitli uygulamalar barındırmaktadır (Dong, 2024: 4).

Doğal dil işleme modelleri, bu becerileri uygun yapılarda işlemek ve dilin çeşitli yapılarını öğrenmek için farklı yöntemler ve yaklaşımlar kullanmaktadır. Veri kaynaklarına geleneksel istatistiksel yöntemlerle yaklaşmanın yanı sıra derin öğrenme ve sinir ağları gibi daha modern teknikleri de merkezinde barındırmaktadır. Doğal dil işleme modelleri, bu noktada dilin anlam bilimsel yapısını, dil bilgisi kurallarını ve bağlamsal ilişkilerini sayısallaştırarak metinlerdeki anlamı sayısal temsillere dönüştürmeyi amaçlamaktadır (Cai, 2024: 62-64). Başlangıçta kural tabanlı sistemler olarak geliştirilen doğal dil işleme modelleri, özellikle büyük dil modellerinin ortaya çıkmasıyla birlikte yerini veri odaklı yaklaşım, örüntü tanıma ve tahmine dayalı öğrenmeye bırakmıştır. Böylece dil modellerinin uygulama alanları genişlemiş ve günümüzde birçok sektörde yaygın olarak kullanılmaya başlanmıştır. Bu uygulama alanlarından bazılarını metin sınıflandırma ve duygu analizi, makine çevirisi ve sohbet uygulamaları oluşturmaktadır (Bowden, 2023: 194-203).

Çalışmanın ana unsuru olan büyük dil modelleri, derin öğrenme tekniklerine dayalı olarak büyük miktarda metin verisi üzerinde eğitilmiş yapay sinir ağları olarak karşımıza çıkmaktadır. Dil modelleri, büyük ölçekli metin verileri üzerinden eğitilerek insan dilinin gramer yapısını, anlamsal bağlarını ve bağlamsal özelliklerini öğrenmeyi amaçlamaktadır. *Metin üretimi, soru yanıtama, çeviri, özetleme ve metin sınıflandırma* gibi çeşitli görevleri yüksek doğrulukla gerçekleştirmek için geliştirilmişlerdir (Ren, 2024: 55-58; Song vd., 2024: 11). Büyük dil modelleri, özellikle *Transformer* mimarisi üzerine inşa edilmiştir. Transformer mimarisi, kodlayıcı (*encoder*) ve kod çözücü (*decoder*) katmanlardan oluşan sinir ağı yapısı olarak “öz-dikkat” (*self-attention*) yaklaşımıyla kelimeler arasındaki bağlamsal ilişkileri ağırlıklandırmaktadır (Vaswani vd., 2017: 1-3). Yapısı gereği milyarlarca, hatta trilyonlarca değişken içeren karmaşık bir sinir ağından oluşmaktadır. Bu dil modellerine, *OpenAI* tarafından geliştirilen *GPT* modelleri ve *Google* tarafından geliştirilen *Gemini* örnek olarak gösterilebilir (Matarazzo ve Torlone, 2025: 6-10).

Dil modelleri, dilin yapısı üzerindeki üretkenlik yeteneğini yalnızca modelin mimari yapısına veya eğitim verilerine dayandırmamaktadır. Aynı zamanda dil modelleri, çevrim içi kullanıcı girdileriyle sağlanan etiketsiz ve çoğu durumda nitelsiz veri yüküne de maruz kalmaktadır. Bu durum, dil modellerinin kararsız çalışmasına, yüksek halüsinasyon değerlerine ve akademik düzlemden uzaklaşmasına yol açmaktadır. Dil modellerindeki yapısal özelliklerin yanı sıra dil çeşitliliği, uzmanlaşmış kaynak eksikliği ve yöntem sorunlarının da önemli kısıtlamalara yol açtığı görülmektedir. Yapay zekâ tabanlı teknolojilerin, özellikle büyük dil modellerinin, uluslararası geçerliliğe sahip dil olarak İngilizce kaynaklar üzerinden eğitilmesi, Türkçe gibi farklı dillerdeki yeterliliklere ilişkin becerilerin başka bir eksene kaymasına yol açmaktadır. Türkçe; sonradan eklemeli dil yapısı, zengin ek sistemi, esnek kelime dizilimi ve özgün dil bilgisi kurallarıyla doğal dil işleme çalışmalarında ayırıcı nitelikler barındırmaktadır (Durrant, 2013: 30-32). Bu durum, Türkçeyi hem öğrenme hem de dili işleme süreçlerinde diğer dillerden ayırmakta ve dil modelleri için farklı güçlükler ortaya çıkarmaktadır. Bu sebeple, Türkçe gibi morfolojik açıdan zengin ve bağlamsal olarak karmaşık dillerin özgün nitelikleri dikkate alınarak uygun yaklaşımların geliştirilmesi gerekmektedir.

Bu çalışma, büyük dil modellerinin ses bilgisi terimlerini tanım düzeyinde anlama ve doğru biçimde kullanma performansını incelemeyi amaçlamaktadır. Bu kapsamda, büyük dil modelleri, söz konusu performansı değerlendirmeye yönelik kıyaslama veri kümeleriyle sınanmıştır. Çalışmada ayrıca büyük dil modellerinde görülen çok terimliliğe ilişkin bulgular yorumlanmıştır.

Araştırma, büyük dil modellerinin Türkçe ses bilgisi terimlerine ilişkin performanslarını kıyaslama veri kümeleri aracılığıyla inceleyerek alanyazındaki boşluğa dikkat çekmektedir.

2. İlgili Araştırmalar

Ses bilgisi, dilin temel yapısını oluşturan sesleri ve bu seslere ilişkin değişimleri ele alması bakımından alanyazında önemli bir yer tutmaktadır. Çağdaş teknolojiler, bu kaynakların anlamlandırılması ve dilin yapısının çözümlenmesi bakımından önem taşımaktadır. Bu doğrultuda, dil bilgisinin alt disiplinleri dâhil olmak üzere tüm kurallı yapılarının, terim bilgisinin ve ilgili uygulama alanlarının sayısallaştırılarak veri olarak işlenmesi gerekmektedir. Bu sebeple, Türkçe terimlerde adlandırma ve terim birliğinin sağlanması, hem eğitim öğretim açısından hem de dil modellerinin terimleri doğru biçimde işlemesi bakımından kritik rol oynamaktadır. Her yazılı kaynakla birlikte Türkçedeki çok terimlilik artmakta; bu durum, dil modellerinde terim ve tanım karmaşasına neden olabilmektedir. Akademik çevrelerde çok terimliliğin sorun oluşturduğuna dair görüşler bulunmakla birlikte, bu durumun tarihsel kökeni Tanzimat döneminde terimlerin Türkçeleştirilmesi çalışmalarına kadar uzanmaktadır. Boz'a göre, alanyazındaki çok terimliliğin sebepleri; Osmanlı bakiyesinden süregelen eski terimlerin güncellenmesi, Batılı dillerden yapılan tercüme, dil bilgisi kavramlarına nokta atışı karşılıklar bulma arayışı, anlaşılama veya yanlış anlaşılma, tercih edilen

terimlerdeki ses değışiklikleri, çok anlamlılıktan kurtulma çabası ve kişisel tercihlerdir (Boz, 2018: 1593-1600). Bu bağlamda, kavramların birden çok terimle karşılanması (*çok terimlilik*), yalnızca alanyazında değil, büyük dil modellerinin terim kullanımında da bir sorun olarak karşımıza çıkmaktadır.

Ses bilgisindeki çok terimliliği, Karaağaç'ı referans noktası olarak açıklamak gerekirse, alanyazında şu şekilde karşımıza çıkmaktadır: *ünsüzler* için “*sesdeş, sessiz, konson ve konsonant*”; *ünlüler* için ise “*sesli ve vokal*” terimleri kullanılmaktadır (Banguoğlu, 1995: 34-40; Ergin, 2013: 39-43). *Ötümlü ve ötümsüz ünsüz* terimleri için “*tonlu, tonsuz, sedalı, sedasız, sert ve yumuşak*” gibi adlandırmalar kullanılmaktadır (Vural ve Böler, 2011: 73-74; Ergin, 2013: 46; Hengirmen, 2007: 65-66). *Patlamalı ünsüzler* için “*patlayıcı ünsüz, süreksiz ünsüz, patlamalı sesdeş, kapanma ünsüzü ve patlayıcı konson*” gibi adlandırmalara rastlanmaktadır (Eker, 2006: 274; Hengirmen, 2007: 64; Banguoğlu, 1995: 41; Güneş, 1999: 47). *Sürtünmeli ünsüzler* için ise “*sızmalı sesdeş, daralmalı sesdeş, sürekli ve sızıcı ünsüz*” gibi adlandırmalara yer verilmiştir (Banguoğlu, 1995: 41-42; Eker, 2006: 274; Hengirmen, 2007: 65). *Kalın ve ince ünlüler* için “*ön ve art (arka) ünlü, kalın ve ince vokal*” terimleri kullanılmaktadır (Eker, 2006: 266; Ergin, 2013: 40). Ses olaylarının adlandırılmasında benzer durum devam etmektedir. *Aykırılma* terimi için “*benzeşmezlik, başkalaşma ve disimilasyon*” gibi terimler kullanılmaktadır (Hengirmen, 2007: 92; Eker, 2006: 310). Ayrıca alanyazında *göçüşme* terimini karşılayan “*yer değıştirme ve metatez*” terimleri de bulunmaktadır (Güneş, 1999: 54; Eker, 2006: 308; Ergin, 2013: 52). Özellikle MEB düzeyinde *büyük ünlü uyumu ve küçük ünlü uyumu* olarak öğretilen yaygın terim kullanımına ek olarak, “*dudak ve damak uyumu, artlık-önlük uyumu, kalınlık-incelik uyumu ve düzlük-yuvarlaklık uyumu*” gibi farklı adlandırmalara da rastlanmaktadır (Eker, 2006: 283-285; Güneş, 1999: 61; Ergin, 2013: 70-72; Demir ve Yılmaz, 2017: 57-60). Bu örnekler, alanyazındaki çok terimlilik durumlarının yalnızca bir bölümünü oluşturmaktadır (Karaağaç, 2013: 95-151).

Büyük dil modelleri, çok terimlilik ortamında, öğrenme kaynaklarının toplumsal veya akademik çevrelere dayanıp dayanmadığına bakılmaksızın, terim kullanımında tutarsızlıklar gösterebilmektedir. Bu nedenle, büyük dil modellerinin anlam kayması, yanlış bağlamda ilerleme veya yanıltıcı bilgiler üretme durumlarının değerlendirilmesi gerekmektedir. Bu amaçla, dil modelleri üzerinde terim ve tanım düzeyinde doğruluk kıyaslamaları yapılmaktadır.

Alanyazında, Türkçe büyük dil modellerini değerlendirmek amacıyla geliştirilen “*CETVEL*” kıyaslama veri kümesinin, dilsel ve kültürel özgünlüğü içeren kapsamlı bir değerlendirme yapısına sahip olduğu görülmektedir. Dil modellerinin kıyaslanmasında “*TR-MMLU*” ise başka bir değerlendirme aracı olarak karşımıza çıkmaktadır. Görev çeşitliliği ve kültürel bağlam noktasındaki sınırlılıkları ortadan kaldıran “*CETVEL*”, atasözleri ve bilmece gibi özel yapıları barındırarak değerlendirme alanını geniş

tutmaktadır. Bu doğrultuda yapılan kıyaslamalarda, “LLaMa 3” dil modeli ailesi diğer modellere kıyasla daha yüksek başarı göstermiştir. Buna karşın, Türkçe odaklı geliştirilen dil modellerinin genel olarak çok dilli modellere göre daha düşük başarı sergilediği görülmektedir. Yalnızca Türkçe veri kümeleriyle eğitilen modellerin ise çeviri, kapsamlı tanım yapıları ve genel muhakeme görevlerinde yetersiz kaldıkları ifade edilmiştir. Alanyazındaki bu araştırmalar, Türkçe büyük dil modellerinin geliştirilmesinde çok dilli ve geniş kapsamlı veri bloklarının kullanımının büyük önem taşıdığını göstermektedir. Doğal dil işleme teknolojileri alanında, dil bilimsel ayırt ediciliğe sahip modellerin geliştirilmesinde ana unsurlardan birinin de kıyaslama veri kümeleriyle değerlendirme olduğu görülmektedir (Er vd., 2026: 1-9; Bayram vd., 2024: 1-6).

3. Yöntem

Çalışma, büyük dil modellerinin farklı kıyaslama veri kümelerine verdikleri yanıtların karşılaştırmalı olarak incelenmesine dayanmaktadır. Büyük dil modellerinin *çoktan seçmeli ses bilgisi testine* verdikleri yanıtlar doğruluk oranları üzerinden değerlendirilmiştir. *Eşdeğer seçenekler arası seçim ve tanımdan terim çıkarım testlerine* verdikleri yanıtlar ise karşılaştırmalı olarak incelenerek yorumlanmıştır. Bu doğrultuda, çalışmada nicel ve nitel değerlendirmeler birlikte kullanılmıştır. Çalışmanın verisi, Türkçenin ses bilgisi terminolojisi kapsamında derlenen kaynaklardan yararlanılarak oluşturulan tanımlar ve bu tanımlara büyük dil modelleri tarafından verilen yanıtlardan oluşmaktadır. Bu veriler JSON biçiminde yapılandırılarak saklanmıştır.

Çalışmada kullanılan büyük dil modelleri (*Qwen3 Next, Cogito v2.1, Gemini 3.5 Flash, Gemma 4, MiniMax M2, MiniMax M2.5, Devstral 2, Nemotron 3 Super, GPT-OSS, ChatGPT 5.5, GLM-4.6 ve GLM-4.7*), 18 Mayıs 2026 tarihinde erişilebilen ve Türkçe dil desteği sunan modeller arasından seçilmiştir.

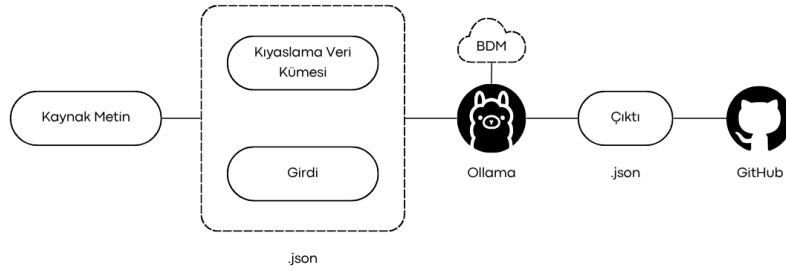
Buna göre, “*Türkçe dil desteği*”nin belirlenmesinde üç ölçüt dikkate alınmıştır:

1. Dil modelinin Türkçe girdileri işleyebilmesi ve Türkçe yanıt üretebilmesi.
2. Geliştiricisi tarafından Türkçenin desteklenen diller arasında belirtilmesi veya Türkçe metinlerle çalışabildiğinin ifade edilmesi.
3. Dil modelinin Türkçe desteğinin yayımlanmış araştırmalar veya kıyaslama çalışmalarıyla doğrulanabilmesi.

Belirlenen ölçütleri karşılayan büyük dil modelleri arasından farklı geliştiricilere (*Google, OpenAI, MiniMax, Zhipu AI vb.*) ait modeller seçilmiştir. Bu seçim, veri toplama sürecine başlanmadan önce gerçekleştirilmiştir. Araştırmada kullanılan büyük dil modelleri *akıl yürütme (thinking)* modunda çalıştırılmıştır. *ChatGPT 5.5 ve Gemini 3.5 Flash*, kişisel kullanım için sunulan web arayüzleri üzerinden; diğer dil modelleri ise Ollama API’si aracılığıyla *bulut (cloud)* sürümleri kullanılarak çalıştırılmıştır.

Çalışmanın ön test (*duman testi*) aşamasında, Türkçeye özgü geliştirilen büyük dil modellerinin (*Kumru 2B* ve *Kanarya 2B*) veri kümesinin kompakt yapısını (*JSON biçiminde yapılandırılmış*) doğru biçimde çözümleyemedikleri tespit edilmiştir. Soru, seçenek ve cevap etiketleri arasındaki yapısal ilişkileri doğru biçimde kuramadıkları için çalışmanın dışında tutulmuşlardır.

Görsel 1. Sistem Erişimi Akış Diyagramı



Büyük dil modellerine erişim, Ollama çalışma ortamı üzerinden (*ChatGPT 5.5* ve *Gemini 3.5 Flash hariç*) sağlanmıştır. Soru ve seçeneklerin büyük dil modellerine aktarılmasını ve işlenmesini kolaylaştırmak amacıyla veri yapısı JSON biçiminde tasarlanmıştır. *OpenAI* ve *Google* tarafından geliştirilen dil modelleri ise kendi web arayüzleri üzerinden test edilmiştir. Çalışmada kullanılan giriş yönergeleri, geçmiş sohbet bağlamının etkisini önlemek amacıyla yeni kullanıcı bilgileri üzerinden uygulanmıştır. Seçilen dil modellerinin bağlam kapasitesi 128 bin (*GPT-OSS*) ile 256 bin (*Qwen3 Next*, *Gemma 4*, *Devstral 2* ve *Nemotron 3 Super*) belirteç arasında değişiklik göstermektedir. Ayrıca modellerin değişken (parametre) sayıları 31 milyar (*Gemma 4*) ile 671 milyar (*Cogito v2.1*) arasında değişmektedir. Bağlam kapasitesi ve değişken sayılarındaki farklılıklar, örneklem dengesi gözetilerek dikkate alınmıştır.

Çalışmanın temelini, büyük dil modellerine yöneltilen soruları içeren kıyaslama veri kümeleri oluşturmaktadır. Bu bağlamda, büyük dil modellerinin Türkçe ses bilgisi terimlerine yönelik performanslarını değerlendirmek amacıyla, iki çoktan seçmeli (*çoktan seçmeli ses bilgisi ve eşdeğer seçenekler arası seçim testleri*) ve bir açık uçlu (*tanımdan terim çıkarım testi*) olmak üzere üç kıyaslama veri kümesi oluşturulmuştur. Bu veri kümeleri, dil modellerinin ses bilgisi terimlerine yönelik performanslarını referanslı ve referanssız soru biçimleri üzerinden değerlendirmek üzere tasarlanmıştır.

Bu kapsamda, dil modellerinin seçenekler arasından doğru yanıtı seçme (veya eşdeğer seçenekler arasından terim tercih etme) ve verilen tanıma uygun yanıtı üretme biçimlerinin incelenmesi amaçlanmıştır.

Çoktan seçmeli ses bilgisi testi, Anadolu Üniversitesi Açıköğretim Fakültesi tarafından “*Türkçe Ses Bilgisi*” dersi kapsamında gerçekleştirilen final sınavlarından derlenmiştir. Kıyaslama veri kümesinin oluşturulmasında 2012–2025 yılları arasındaki sınav soruları kullanılmıştır. Her eğitim yılından rastgele beş soru seçilmiş, tekrarlanan sorular çıkarılarak toplam elli dört benzersiz çoktan seçmeli soru elde edilmiştir. Her soru beş seçenek içermekte ve tek doğru cevap etiketi ile sunulmaktadır. Seçilen sorular ünlü ve ünsüz seslerin sınıflandırılması, ses olayları ve ünlü uyumları gibi çeşitli konuları kapsamaktadır.

Eşdeğer seçenekler arası seçim testi, soru yapılarının oluşturulmasında *Güncel Türkçe Sözlük* (TDK, t.y.) tanımları ile kaynakçada yer verilen ses bilgisi çalışmalarından yararlanılarak hazırlanmıştır. Söz konusu çalışmalar, alanyazında sıklıkla başvuru kaynağı olarak kullanılan kaynaklar arasından seçilmiştir. Kıyaslama veri kümesi elli yedi çoktan seçmeli sorudan oluşmakta, her soru beş seçenek içermekte ve tek doğru cevap bulunmamaktadır. Aynı tanıma karşılık gelen terimler seçenek olarak verilmiştir. Dil modellerinin yanıtları puanlanmamış, tercih örüntüleri karşılaştırmalı tablolar hâlinde sunulmuştur.

Tanımdan terim çıkarım testi, terim tanımları ve ses olaylarını kapsayan yetmiş dört boşluk doldurma sorusundan oluşmaktadır. Ünlülerin ve ünsüzlerin oluşumu, ses uyumları ve ses olaylarını (*ünlü düşmesi, ünsüz türemesi, göçüşme vb.*) tanımlayan otuz beş soru ile seslerin (8 ünlü ve 21 ünsüz) boğumlanma özelliklerine göre sınıflandırılmasını isteyen otuz dokuz soru içermektedir.

Çalışmada, *çoktan seçmeli ses bilgisi testi* için nicel analiz yöntemi kullanılmıştır. Dil modellerinin yanıtları doğru seçeneklerle karşılaştırılarak doğru/yanlış biçiminde ikili (*binary*) olarak kodlanmıştır. Ayrıca soru yapısını bozan veya bağlam dışı metinler içeren yanıtlar halüsinasyon olarak etiketlenmiştir. Bu yöntemle modellerin doğruluk, hata ve halüsinasyon oranları, ilgili yanıt sayılarının toplam soru sayısına bölünmesiyle hesaplanmıştır. Örneklem boyutunun küçük olması dikkate alınarak doğruluk oranlarına ilişkin %95 güven aralıkları *Wilson Score yöntemi* kullanılarak belirlenmiştir. İkili olarak kodlanan yanıt değişkeninin varyansı, *Bernoulli dağılımı varsayımına göre* her modelin doğruluk oranı (p) ve doğru olmayan yanıt oranı ($q = 1 - p$) dikkate alınarak $p \times q$ formülüyle hesaplanmıştır. Dil modelleri arasındaki performans farklılıklarını değerlendirmek amacıyla *Gemini 3.5 Flash* referans olarak belirlenmiştir. Diğer modellerin 54 sorudaki doğru/yanlış sonuçları, *Gemini 3.5 Flash*’ın aynı sorulardaki sonuçlarıyla eşleştirilmiş ve *iki yönlü exact McNemar testi* uygulanmıştır. Dil modelleri arasındaki ikili karşılaştırmalar sonucunda elde edilen p -

değerleri Bonferroni yöntemiyle düzeltilmiş ve anlamlılık düzeyi $\alpha = 0,05$ olarak belirlenmiştir. Hata ve halüsinasyon olarak sınıflandırılan yanıtlar, McNemar analizinde ve varyans hesaplamalarında “doğru olmayan yanıt” kategorisi altında birleştirilmiştir. Eşdeğer seçenekler arası seçim testi için dil modellerinin yanıtları incelenerek çok terimlilik bulguları tablolar hâlinde sunulmuş ve yorumlanmıştır. Tanımdan terim çıkarım testi için ise önceki bulgulara ek olarak, modellerin açık uçlu soru yapılarındaki farklılaşma durumları karşılaştırmalı olarak verilmiştir.

Ayrıca büyük dil modellerinden Ollama çalışma ortamında elde edilen yanıtlar, GitHub deposu üzerinden açık kaynak olarak sunulmaktadır.³

4. Bulgular ve Yorumlar

Bu bölümde, seçili büyük dil modellerinin çıktılarında elde edilen bulgulara yer verilmiştir.

4.1. Dil Modellerinin Türkçe Ses Bilgisi Sorularındaki Başarı Düzeyleri

Dil modelleri, Anadolu Üniversitesi Açıköğretim Fakültesinin lisans düzeyindeki sorularına verdikleri yanıtlar açısından incelendiğinde, temel ses bilgisi terimleri, sınıflandırmalar ve boğumlanma özelliklerine ilişkin sorularda anlam kayması yaşamadıkları (*istisnalar dışında*) ve başarılı oldukları görülmüştür. Özellikle dil modellerinin *ünlülerin niteliksel sınıflandırmaları* ile *çift dudak ünsüzü*, *yarı ünlü* ve *ötümlülük* gibi tanımları eşleştirirken daha az hata yaptıkları görülmüştür. Bununla birlikte, *Qwen3 Next* ve *GPT-OSS* dil modellerinin “a” ünlüsünü “ön” veya “o” ünlüsünü “dar” olarak tanımlaması, ayrıca *akıcı* ve *ötümlü* özellikleri ayırt etmeden yanıt üretmesi, bazı durumlarda terim bilgisi açısından kararsızlık yaşadıklarını ortaya koymaktadır.

Aynı zamanda, dil modellerinin Açıköğretim Fakültesi sorularında şu konularda belirgin zayıf noktaları olduğu görülmektedir:

Gerileyici benzeşme, *nöbetleşme* ve *ön damaksıllaşma* terimlerine ilişkin sorular ile özellikle “*geniz n*” sesine dair sorularda başarısız olmuşlardır. Aynı aileye ait dil modellerinin bazı sorularda birbiriyle çelişen yanıtlar ürettiği görülmüştür. Örneğin, *GLM 4.6* dil modeli *göçüşme* ile ilgili soruyu doğru yanıtlarken *GLM 4.7* aynı soruyu yanlış yanıtlamıştır. *MiniMax M2* dil modeli *geniz n* sesi ile ilgili soruyu doğru yanıtlarken *MiniMax M2.5* aynı soruyu yanlış yanıtlamıştır.

³ Bu çalışmada, dil modelleri tarafından üretilen çıktılar, araştırma verilerinin erişilebilirliğini sağlamak amacıyla açık kaynak olarak GitHub üzerinden paylaşılmaktadır. Erişim bağlantısı: <https://github.com/erkutresul/trphonologicalterms>

Tablo 1. Büyük Dil Modellerinin Açıköğretim Fakültesi Sorularındaki Başarısı

<i>Büyük Dil Modelleri</i>	<i>Doğruluk (%)</i>	<i>Hata (%)</i>	<i>Halü. (%)</i>	<i>%95 GA</i>	<i>Varyans</i>	<i>p Değeri</i>
Gemini 3.5 Flash	96,30	3,70	-	87,46 – 98,98	0,0357	Ref.
Gemma 4	83,33	16,67	-	71,26 – 90,98	0,1389	0,4297
ChatGPT 5.5	72,22	27,78	-	59,11 – 82,38	0,2006	0,0107
Cogito v2.1	72,22	25,93	1,85	59,11 – 82,38	0,2006	0,0027
MiniMax M2	72,22	27,78	-	59,11 – 82,38	0,2006	0,0027
Nemotron 3 Super	70,37	29,63	-	57,17 – 80,86	0,2085	0,0057
GLM 4.7	66,67	31,48	1,85	53,36 – 77,76	0,2222	0,0016
GPT OSS	64,81	33,33	1,85	51,48 – 76,18	0,2281	< 0,001
Devstral 2	64,81	35,19	-	51,48 – 76,18	0,2281	0,0024
MiniMax M2.5	62,96	37,04	-	49,63 – 74,58	0,2332	< 0,001
GLM 4.6	59,26	40,74	-	45,97 – 71,32	0,2414	< 0,001
Qwen3 Next	53,70	46,30	-	40,61 – 66,31	0,2486	< 0,001

Tablo 1'e göre, Açıköğretim Fakültesine ait Türkçe ses bilgisi sorularından oluşan kıyaslama veri kümesiyle yapılan değerlendirmede, en yüksek doğruluk oranına sahip dil modelinin *Gemini 3.5 Flash* (%96,3) olduğu görülmektedir. En düşük doğruluk oranına sahip dil modeli ise *Qwen3 Next* (%53,7) olmuştur. *Gemini 3.5 Flash* dil modeli, 54 sorunun 52'sini doğru yanıtlayarak en yüksek başarıyı gösterirken, diğer modellerin doğru yanıtladığı soru sayıları 30 ile 45 arasında değişmektedir. Doğruluk oranları açısından değerlendirildiğinde, *Gemini 3.5 Flash*, *Gemma 4*, *ChatGPT 5.5*, *Cogito v2.1*, *MiniMax M2* ve *Nemotron 3 Super* dil modellerinin %70 ve üzerinde doğruluk oranlarına ulaştığı görülmektedir. Buna karşılık özellikle *GLM 4.6* ve *Qwen3 Next* modellerinin doğruluk oranları %60'ın altında kalmıştır.

Doğruluk oranlarına ilişkin %95 güven aralıkları incelendiğinde, *Gemini 3.5 Flash* için güven aralığının %87 ile %99, *Qwen3 Next* için ise %42 ile %68 arasında olduğu

görülmektedir. Hata ve varyans değerlerindeki farklılıklar, model performanslarının birbirinden farklılaştığını göstermektedir.

Bununla birlikte, güven aralıkları, hata oranları ve varyans değerleri tek başına dil modelleri arasındaki farklılıkların istatistiksel olarak anlamlı olup olmadığına ilişkin yeterli kanıt sunmamaktadır. Bu nedenle, *Gemini 3.5 Flash* referans model olarak belirlenmiş ve aynı elli dört soruya verilen yanıtlar eşleştirilerek modeller arasındaki farklılıkların değerlendirilmesi amacıyla *iki yönlü exact McNemar testi* uygulanmıştır. *McNemar* analizinde, hata ve halüsinasyon değerleri birleştirilmiştir. *Bonferroni düzeltmesi* sonrasında, *Gemini 3.5 Flash* ile *Gemma 4* arasındaki farklılığın istatistiksel olarak anlamlı olmadığı ($p = 0,430$), buna karşılık *Gemini 3.5 Flash* ile diğer dil modelleri arasındaki farklılıkların istatistiksel olarak anlamlı olduğu ($p < 0,05$) belirlenmiştir.

Buna göre, *Gemini 3.5 Flash* dil modeli en yüksek doğruluk oranına sahip olmasına rağmen *Gemma 4* karşısında istatistiksel olarak anlamlı bir üstünlük göstermemiş; diğer dil modelleri karşısında ise bu üstünlük istatistiksel olarak anlamlı bulunmuştur. Bu durum, dil modellerinin Türkçe ses bilgisi konularındaki başarı düzeylerinin birbirinden farklılaşabileceğini göstermiştir.

4.2. Dil Modellerinin Ünlüleri Sınıflandırma Biçimleri

Tablo 2. Büyük Dil Modellerinin Ünlüleri Sınıflandırması

<i>Büyük Dil Modelleri</i>	<i>Açıklık Derecesi</i>	<i>Dil Konumu</i>	<i>Dudak Durumu</i>
Gemini 3.5 Flash	geniş/dar	kalın/ince	düz/yuvarlak
Gemma 4	geniş/dar	kalın/ince	düz/yuvarlak
ChatGPT 5.5	açık/dar	kalın/ince	düz/yuvarlak
Cogito v2.1	geniş/dar	kalın/ince	düz/yuvarlak
MiniMax M2	geniş/dar	art/ön	düz/yuvarlak
Nemotron 3 Super	açık/kapalı	kalın/ön	düz/yuvarlak
GLM 4.7	geniş/dar	kalın/ince	düz/yuvarlak
GPT OSS	açık/dar	kalın/ince	düz/yuvarlak
Devstral 2	geniş/dar	kalın/ince	düz/yuvarlak

MiniMax M2.5	geniş/dar	art/ön damak	düz/yuvarlak
GLM 4.6	geniş/dar	kalın/ince	düz/yuvarlak
Qwen3 Next	açık/dar	kalın/ince	düz/yuvarlak

Tablo 2’de büyük dil modellerinin ünlüleri sınıflandırma biçimleri karşılaştırmalı olarak incelenmiştir. Bu doğrultuda, dil modellerinin Türkçe ses bilgisi terimlerini doğru biçimde tanımlayabildiği görülmüştür. Özellikle “*dar ünlü*”, “*kalın ünlü*”, “*ince ünlü*”, “*düz ünlü*” ve “*yuvarlak ünlü*” terimlerinde belirgin birer ortaklık bulunmaktadır (*ChatGPT 5.5*, *Cogito v2.1*, *Gemini 3.5 Flash*, *Gemma 4*, *GLM-4.6*, *GLM-4.7 vb.*). Bununla birlikte, dil modelleri arasında terim tercihleri bakımından bazı farklılıklar göze çarpmaktadır. Bu dil modelleri (*MiniMax M2*, *MiniMax M2.5*, *Nemotron 3 Super vb.*) ise farklılaşarak “*art ünlü*”, “*ön ünlü*”, “*açık ünlü*” ve “*kapalı ünlü*” gibi alternatif terimleri tercih etmişlerdir. Ayrıca dil modelleri bu konu başlığı altında sorulan ünlü uyumu tanımlarına çoğunlukla “*büyük ünlü uyumu*” ve “*küçük ünlü uyumu*” olarak cevap vermişlerdir. Buna rağmen *MiniMax M2* dil modeli, ünlü uyumlarında farklı kavramlara yönelerek “*dudak uyumu*” ve “*damak uyumu*” gibi terimleri tercih etmiştir. Ayrıca *Gemini 3.5 Flash* dil modelinin, terim tercihlerinde alternatif terimleri (*kalınlık-incelik uyumu*, *düzlük-yuvarlaklık uyumu*, *vokal*, *konsonant*, *ön ünlü*, *art ünlü vb.*) parantez içinde verdiği görülmüştür. Bulgular, özellikle dil modellerinin açık uçlu soru yapılarında alternatif terimlere (*açık veya kapalı ünlü*, *ön veya arka ünlü vb.*) yöneldiğini göstermiştir. Bu durum özellikle *Nemotron 3 Super*, *GPT-OSS* ve *Qwen3 Next* gibi dil modellerinde görülmüştür.

4.3. Dil Modellerinin Ünsüzleri Sınıflandırma Biçimleri

Tablo 3. Büyük Dil Modellerinin Ünsüzleri Sınıflandırması

<i>Büyük Dil Modelleri</i>	<i>Ses Tellerine Göre</i>	<i>Boğumlanma Biçimine Göre</i>		
Gemini 3.5 Flash	ötümlü/ötümsüz	süreksiz	sızıcı	akıcı
Gemma 4	ötümlü/ötümsüz	patlayıcı	sürekli	akıcı
ChatGPT 5.5	ötümlü/ötümsüz	süreksiz	sürekli	akıcı
Cogito v2.1	ötümlü/ötümsüz	patlayıcı	sürtünmeli	akıcı
MiniMax M2	ötümlü/ötümsüz	patlayıcı	sürekli	akıcı
Nemotron 3 Super	sedalı/tonsuz	patlayıcı	sürekli	akıcı
GLM 4.7	ötümlü/ötümsüz	süreksiz	sürekli	akıcı

GPT OSS	tonlu/ötümsüz	kapanma	sürtünmeli	sürtünmesiz
Devstral 2	ötümlü/ötümsüz	süreksiz	sürtünmeli	akıcı
MiniMax M2.5	ötümlü/ötümsüz	patlayıcı	sürtünmeli	akıcı
GLM 4.6	ötümlü/ötümsüz	patlayıcı	sürtünmeli	akıcı
Qwen3 Next	tonlu/tonsuz	patlayıcı	sürtünmeli	akıcı

Tablo 3’te büyük dil modellerinin Türkçedeki ünsüz sınıflandırmalarına yönelik verdikleri yanıtlar karşılaştırmalı olarak incelenmiştir. Elde edilen bulgular, dil modellerinin büyük çoğunluğunun ses bilgisi terimlerini doğru biçimde tanımlayabildiğini göstermiştir. Özellikle “ötümlü ve ötümsüz ünsüz” ile “akıcı ünsüz” gibi terim tanımlarına yönelik, aralarında yüksek düzeyde örtüşme gösteren yanıtlar üretmişlerdir. Bu durum, dil modellerinin terim tanımları alanında birbiriyle kesişen veri yapılarına sahip olduklarını ortaya koymaktadır.

Bununla birlikte, büyük dil modelleri arasında terim tercihlerinde kısmi farklılıklar dikkat çekmektedir. Dil modelleri, terim kullanımında “sürekli ve süreksiz ünsüz”, “sürtünmeli ve sürtünmesiz ünsüz” ile “tonlu ve tonsuz ünsüz” gibi farklı terimlere yönelmiştir.

Bu farklılık, modellerin eğitildiği veri kaynaklarının çeşitliliğinden ve terminolojik standartlaşmanın doğal dil işleme sistemlerinde henüz tam olarak sağlanamamış olmasıyla açıklanabilir. Özellikle “patlayıcı ünsüz” ile “süreksiz ünsüz” kullanımındaki ayrım, akademik kaynaklar ile MEB’in öğretim odaklı dil bilgisi yaklaşımı arasındaki kaynak farklılaşmasını açık biçimde göstermektedir. Öyle ki bazı modellerin daha ayrıntılı tanımlamalara ulaşan yanıtlar üretme eğiliminde olduğu görülmektedir. *ChatGPT 5.5*, *Devstral 2* ve *GLM-4.7* gibi modellerin “süreksiz ünsüz” ve “sürtünmeli ünsüz” gibi terimleri kullanmayı tercih ettiği görülmüştür. Ayrıca *Qwen3 Next*, *GPT-OSS* ve *Nemotron 3 Super* gibi dil modellerinin, terim kullanımında diğer modellerden farklılaşma eğilimi gösterdiği belirlenmiştir.

Bu konu başlığı altında sorulan “m, n” seslerini, beş dil modeli dışında hepsi “nazal ünsüz” olarak tanımlamıştır. *Cogito v2.1*, *Gemini 3.5 Flash* ve *GLM-4.7* modelleri “geniz ünsüzü”; *MiniMax* modelleri ise “genizsil sesler” terimini seçmiştir. Bütün dil modelleri “y” sesi için yaptıkları tanımlamalarda “yarı ünlü” teriminde ortaklaşırken, iki dil modeli dışında (*Qwen3 Next* ve *Nemotron 3 Super* dil modelleri “bağlayıcı ünsüz” terimini tercih etmiştir.) tüm dil modelleri “y” sesi için aynı zamanda “kaynaştırma ünsüzü” tanımı yapmıştır.

Dil modelleri, yukarıdaki tanımlamalardan farklı olarak, açık uçlu soru yapılarında farklı terim kullanımları sergilemekte ve bağlam dışı metinler üretebilmektedir. Açık uçlu sorularda “sesli ve sessiz ünsüz” ile “tonlu ve tonsuz ünsüz” terimlerine yönelen modeller, aynı kavramı model bazında farklı adlandırmalarla ifade edebilmektedir.

4.3.1. Dil Modellerinin Ünsüzleri Boğumlanma Yerine Göre Adlandırma Biçimleri

Tablo 4. Büyük Dil Modellerinin Ünsüzleri Boğumlanma Noktasına Göre Sınıflandırması

<i>Büyük Dil Modelleri</i>	<i>Boğumlanma Noktası</i>	<i>Ünsüzler</i>
Gemini 3.5 Flash	Çift dudak ünsüzü	b, m, p
	Diş-dudak ünsüzü	f, v
	Dil ucu-diş ünsüzü	d, n, s, t, z
	Diş eti ünsüzü	l, r
	Diş eti-damak ünsüzü	c, ç, j, ş
	Ön damak ünsüzü	g, k, y
	Yumuşak damak ünsüzü	ğ
	Gırtlak ünsüzü	h
Gemma 4	Dudak ünsüzü	b, p
	Diş-dudak ünsüzü	f, v
	Dudak-geniz ünsüzü	m
	Dil ucu-diş ünsüzü	s, t, z
	Diş eti ünsüzü	d, n, l, r
	Diş eti-damak ünsüzü	c, ç, j, ş
	Ön damak ünsüzü	g, k, y
	Art damak ünsüzü	ğ
	Gırtlak ünsüzü	h

Erkut RESULOĞLU

ChatGPT 5.5	Çift dudak ünsüzü	b, p
	Diş-dudak ünsüzü	f, v
	Dudak-geniz ünsüzü	m
	Diş eti ünsüzü	l, r, z
	Diş eti ardı ünsüzü	ç, ş
	Dil ucu-diş ünsüzü	d, n, t
	Asıl diş ünsüzü	s
	Diş eti-damak ünsüzü	c, j
	Ön damak ünsüzü	g, y
	Sert damak ünsüzü	k
	Yumuşak damak ünsüzü	ğ
Cogito v2.1	Gırtlak ünsüzü	h
	Çift dudak ünsüzü	b, m, p
	Diş-dudak ünsüzü	f, v
	Diş eti ünsüzü	d, l, n, r, s, t, z
	Diş eti-damak ünsüzü	c, ç, j, ş
	Art damak ünsüzü	ğ
	Ön damak ünsüzü	y
MiniMax M2	Gırtlak ünsüzü	h
	Dudak ünsüzü	b, m, p
	Diş-dudak ünsüzü	f, v
	Diş eti ünsüzü	d, n, l, r, s, t, z
	Damak ünsüzü	g
	Diş-damak ünsüzü	c, ç, j, ş
	Ön damak ünsüzü	k, y
	Yumuşak damak ünsüzü	ğ
	Gırtlak ünsüzü	h

Nemotron 3 Super	Çift dudak ünsüzü	b, m, p
	Diş-dudak ünsüzü	f, v
	Diş eti ünsüzü	d, n, l, r, s, t, z
	Diş eti ardı ünsüzü	c, ç, j
	Diş eti-damak ünsüzü	ş
	Damak ünsüzü	g, k
	Yumuşak damak ünsüzü	ğ
	Sert damak ünsüzü	y
	Gırtlak ünsüzü	h
GLM 4.7	Çift dudak ünsüzü	b, m, p
	Diş-dudak ünsüzü	f, v
	Diş ünsüzü	t
	Diş eti ünsüzü	d, l, n, r, s, z
	Diş eti ardı ünsüzü	j, ş
	Damak ünsüzü	g, k
	Diş eti-damak ünsüzü	c, ç
	Ön damak ünsüzü	y
	Yumuşak damak ünsüzü	ğ
	Gırtlak ünsüzü	h

Erkut RESULOĞLU

GPT OSS	Çift dudak ünsüzü	b, m, p
	Diş-dudak ünsüzü	f, v
	Diş ünsüzü	d, n
	Diş eti ünsüzü	l
	Diş eti ardı ünsüzü	ç, ş
	Asıl diş ünsüzü	s, z
	Dil ucu-diş ünsüzü	t
	Damak ünsüzü	g, k, r
	Diş-damak ünsüzü	c
	Diş eti-damak ünsüzü	j
	Ön damak ünsüzü	y
	Yumuşak damak ünsüzü	ğ
Devstral 2	Gırtlak ünsüzü	h
	Çift dudak ünsüzü	b, p
	Diş-dudak ünsüzü	f, v
	Dudak-geniz ünsüzü	m
	Diş eti ünsüzü	d, n, l, r, s, t, z
	Diş-damak ünsüzü	c, ç, j, ş
	Art damak ünsüzü	ğ
	Sert damak ünsüzü	k, y, g
	Gırtlak ünsüzü	h

MiniMax M2.5	Dudak ünsüzü	b, m, p
	Diř-dudak ünsüzü	f, v
	Diř ünsüzü	t
	Diř eti ünsüzü	d, j, n, l, r, s, ř, z
	Diř eti ardı ünsüzü	ç
	Diř eti-damak ünsüzü	c
	Ön damak ünsüzü	k, y
	Art damak ünsüzü	ğ
	Sert damak ünsüzü	g
	Gırtlak ünsüzü	h
GLM 4.6	Dudak ünsüzü	b, m, p
	Diř-dudak ünsüzü	f, v
	Diř ünsüzü	n, t
	Diř eti ünsüzü	d, l, r, s
	Diř eti ardı ünsüzü	ç, ř
	Damak ünsüzü	g, k
	Diř eti-damak ünsüzü	c, j
	Arka damak ünsüzü	ğ
	Ön damak ünsüzü	y
	Gırtlak ünsüzü	h

Qwen3 Next	Dudak ünsüzü	b, m, p
	Diş-dudak ünsüzü	f, v
	Damak ünsüzü	g, k
	Diş eti-damak ünsüzü	c, ç, j
	Art damak ünsüzü	ğ
	Ön damak ünsüzü	y
	Dil ucu-diş ünsüzü	d, n, s, t, z
	Diş eti ünsüzü	l, r
	Diş eti ardı ünsüzü	ş
	Gırtlak ünsüzü	h

Tablo 4’te büyük dil modellerinin Türkçedeki ünsüzleri boğumlanma noktalarına göre sınıflandırma biçimleri karşılaştırmalı olarak verilmiştir. Elde edilen bulgular, dil modellerinin sınıflandırmalarda doğruluk ve tutarlılık (*dudak ünsüzleri (b, p, m) ve diş-dudak ünsüzleri (f, v) gibi*) sergileme eğilimi gösterdiği yönündedir. Özellikle dil modelleri “*diş-dudak ünsüzü (f, v sesleri için)*” ve “*gırtlak ünsüzü (h sesi için)*” gibi sınıflandırmalarda ortaklık göstermektedir. Bunun yanı sıra, dil modelleri arasında “*diş*” ve “*damak*” sesleri için yapılan sınıflandırmalarda farklılıklar bulunmaktadır. Bununla birlikte, dil modellerinde terim bilgisi açısından belirgin standartlaşma eksikliği tespit edilmiştir. Özellikle dil modellerinin “*damak*”, “*ön damak*”, “*art damak*”, “*arka damak*”, “*sert damak*” ve “*yumuşak damak*” gibi farklı kavramlara yöneldiği görülmüştür.

Tablo 4’e göre, *Qwen3-Next*, *GPT-OSS* ve *GLM* modellerinin daha ayrıntılı sınıflandırma kullandıkları görülmektedir. Bu dil modelleri, boğumlanma noktalarını alt sınıflara ayırmaları bakımından dikkat çekmektedir. Buna karşılık, *Cogito v2.1* dil modelinde görüldüğü üzere, sınıflandırmalar bazı modellerde ayrıntı düzeyinde değildir. Bu durum, dil modelleri arasında yalnızca terim karşılama doğruluğu açısından değil, aynı zamanda terminolojik temsil düzeyi bakımından da farklılıklar bulunduğunu göstermektedir. Ayrıca bazı ünsüzlerin, dil modelleri tarafından farklı alt kategorilere yerleştirilerek birbirinden uzaklaştığı görülmektedir. Özellikle “*ş*”, “*k*”, “*g*”, “*j*” ve “*y*” seslerindeki çeşitlilik, hem dil modellerinin terim temsillerindeki farklılaşma eğilimini hem de model bazı girdilerin yorumlanabilir doğasını göstermektedir.

Dil modellerinin, yukarıdaki tanımlamalar dışında kalan açık uçlu sorulara verdikleri yanıtlarda, sınama sayısı arttıkça seslerin tanımlamalarında kavram kayması yaşadıkları görülmektedir. *GPT-OSS*, *ChatGPT 5.5* ve *GLM* modelleri, bağlamdan uzaklaşma eğilimi göstermektedir. Bu noktada, terim tercihlerinde tutarlılığı kısmen sürdüren iki model (*Cogito 2.1* ve *Devstral 2*) olmuştur. *ChatGPT 5.5* dil modeli, “*c, ç, l, z*” seslerini “*diş eti ünsüzü*” terimi altına çekmiştir. *Gemini 3.5 Flash* dil modeli, “*d, t*” seslerini “*diş ünsüzü*” ve “*g, k*” seslerini

“art damak ünsüzü” terimi altında vermiştir. *Gemma 4* dil modeli ise “g, k” seslerini “art damak ünsüzü” terimi altına çekmiştir. Ayrıca tüm dil modellerinde terim kullanımına ilişkin farklılıklar gözlemlenmiş; *GPT-OSS* ve *MiniMax M2.5* dil modellerinde ise bu farklılıkların daha belirgin olduğu ve bağlam dışı metinlerin de üretildiği görülmüştür.

4.4. Dil Modellerinin Ses Olaylarını Adlandırma Biçimleri

Çalışmada elde edilen bulgular, büyük dil modellerinin özellikle “ünlü daralması”, “ünlü genişlemesi”, “ünlü düzleşmesi”, “ünlü yuvarlaklaşması”, “ünlü ve ünsüz düşmesi”, “ünsüz türemesi”, “ünsüz ikizleşmesi” ve “ünsüz benzeşmesi” adlandırmalarında ses olaylarını yüksek düzeyde örtüşme göstererek işlediğini göstermektedir. Bu durum, dil modellerinin ses olaylarını ve örüntüleri belirli bir doğruluk düzeyinde öğrendiğini ortaya koymaktadır.

Bununla birlikte, dil modellerinin terim tercihlerinde, modeller arasında terminolojik tercih ve kavramsal yaklaşım bakımından belirgin farklılıklar olduğu açıkça görülmektedir. Büyük dil modellerinin “ünlü kaynaşması” yerine “ünlü birleşmesi” (*Gemini 3.5 Flash* ve *Nemotron 3 Super*) veya “büzişme” (*GLM* dil modelleri), “ünlü uyumu” yerine “ahenk kaidesi” (*Nemotron 3 Super*), “ünsüz yumuşaması” yerine “ötümlüleşme” (*MiniMax* modelleri) veya “sedahlama” (*Nemotron 3 Super*) gibi terimleri seçmesi, dil modellerinin farklı dil bilimsel yaklaşımlardan beslendiğini göstermektedir. Dil modellerinden bazıları ses olaylarında terim adlandırmalarını “ünlü incilmesi ve ünlü kalınlaşması” olarak tanımlarken, bazı modeller “öndamaksullaşma ve artdamaksullaşma” olarak tanımlamıştır.

Dil modelleri, TDK’nin *Güncel Türkçe Sözlük* (TDK, t.y.) tanımlarından yola çıkarak oluşturulan ve eşdeğer terimler içeren sorulara ise farklı yanıtlar üretmiştir. *Benzeşmezlik* terimi için büyük dil modellerinden ikisi “aykırılama”, diğerleri ise “disimilasyon” terimini tercih etmiştir. *Göçüşme* terimini dil modellerinden dördü “göçüşme” olarak tanımlarken sekizi “metatez” terimini tercih etmiştir. *Hece tekleşmesi* terimi için büyük dil modellerinden ikisi “hece düşmesi”, diğerleri ise “haploloji” terimini tercih etmişlerdir. *Ünlü türemesi* terimini büyük dil modellerinden beş tanesi “epentez” terimiyle tanımlarken yedi tanesi “ünlü türemesi” terimini korumuştur. Dil modellerinin yaklaşımı, bazı durumlarda terim tercihlerini ödünçlemelerden yana kullandıklarını göstermiştir. Dil modellerinin terim yaklaşımları değerlendirildiğinde, akademik anlamda terim karmaşası yaşadıkları görülmüştür.

Dil modelleri, açık uçlu soru yapılarını içeren sorgulamalarda, *benzeşmezlik* terimi için “ayırışma, uzaklaşma ve ünsüz değişimi” gibi adlandırmalarda bulunmuştur. *Göçüşme* terim tanımına “metatesis ve yer değiştirme” tanımlamaları yapmıştır. *Ünsüz türemesi* terimi için bazı modeller “ünsüz eklenmesi” olarak adlandırmada bulunmuştur. *Hece tekleşmesi* terimi için “haploloji ve hece yutumu” birer kez tercih edilirken diğer modeller “hece düşmesi” terimini seçmişlerdir. *Ünlü kaynaşması* tanımı için iki model “ünlü aşınması”, yedi model “ünlü birleşmesi”, birer model ise “hece birleşmesi ve kaynaşma” adlandırmaları yapmıştır. *Ünlü ve ünsüz benzeşmeleri* terimleri yerine “ünlü ve ünsüz uyumu” terimleri tercih edilmiştir.

Ayrıca açık uçlu soru yapılarındaki terim tanımlamalarında genel ifadeler tercih edilmiştir. Dil modelleri, terimleri *ünlü* ve *ünsüz* olarak alt sınıflandırmalara ayırmadan “*kalınlaşma, yumuşama, incelme ve düzleşme*” gibi terimlerle tanımlamıştır. Ek olarak, dil modelleri ses olaylarını sınıflandırırken ağırlıklı olarak bağlam dışı metinler üretmişlerdir.

5. Sonuç

Büyük dil modelleri, Türkçe ses bilgisi terminolojisi bağlamında yüksek düzeyde öğrenme çeşitliliği göstermektedir. Genel amaçlı dil modellerinin (*ChatGPT 5.5* ve *Gemini 3.5 Flash*), özel amaçlı dil modellerine kıyasla hem tanım öngörme hem de terim uygunluk bakımından tutarlılık gösterme alışkanlığı sergilediği görülmüştür. Buna karşılık dil modelleri, referanssız terim tercihlerinde merkezden uzaklaşma, halüsinasyon, kendine özgü terimlendirme ve anlam kayması gibi etmenler göstermektedir.

Dil modelleri, özellikle çok terimlilik göstermeleri açısından, tanımlarda ayrıntıya ve alt terimlendirmelere yönelme eğilimi göstermektedir.

Dil modellerinin öğrenme kaynakları erişilebilir olmamakla birlikte, ürettikleri yanıtların standartlaşmış veya MEB kaynaklı terimlere yönelme eğilimi gösterdiği değerlendirilmektedir.

Araştırmada, dil modellerinin akademik sınavlarda gösterdiği başarı dikkat çekmektedir. Bu durum, dil modellerinin Anadolu Üniversitesi Açıköğretim Fakültesinin geçme notu şartlarını sağlayabileceğini göstermektedir.

Dil modelleri, ünlü ve ünsüzleri başarılı biçimde sınıflandırmaktadır. Boğumlanma noktası bakımından ünsüzleri sınıflandırmada, büyük dil modelleri çok çeşitli ve ayrıntılı tanımlamalar yapmaktadır. Standart bir tanımlama biçimi gözlemlenmemiştir. Dil modellerinin örneklerden yola çıkarak çok terimlilik sergiledikleri görülmektedir.

Ses olaylarının sınıflandırmasında, büyük dil modelleri ödünclemelerden yararlanarak terim tanımlı yapmaktadır. Ses olayları bağlamında dil modelleri terimlerde kısmen ortaklaşma sergilemektedir.

Sonuç olarak, büyük dil modellerinin çözümleme becerisine sahip olduğu fakat terminolojik standartlaşma açısından hem model düzeyinde hem de modeller arasında terminolojik ortaklığı sağlayamadığı görülmektedir. Birçok etmeni eş zamanlı olarak içerisinde barındıran model yapısı, terimleri öğrenme kaynaklarındaki örnek karşılaştırmasına dayandırarak seçtiğini düşündürmektedir. Aynı zamanda, yaygın kullanım dil modellerinin eğitim öğretim ile ölçme ve değerlendirme süreçlerinde yardımcı araç olarak kullanılabileceğini ancak akademik hassasiyetten uzak olduklarını göstermektedir. Büyük dil modellerinin çıktılarının uzman doğrulaması olmaksızın doğrudan kullanılması uygun görünmemektedir.

Son olarak, büyük dil modellerinin ortaklaşan Türkçe ses bilgisi terimleri doğrultusunda iyileştirilmesi gerekmektedir.

6. Sınırlılıklar ve Öneriler

Bu araştırma, *Qwen3 Next*, *Cogito v2.1*, *Gemini 3.5 Flash*, *Gemma 4*, *MiniMax M2*, *MiniMax M2.5*, *Devstral 2*, *Nemotron 3 Super*, *GPT-OSS*, *ChatGPT 5.5*, *GLM-4.6* ve *GLM-4.7* dil modelleri ile sınırlıdır. Dil modellerinin kıyaslanması, sadece bulut teknolojileri ve web tabanlı (*sadece ChatGPT 5.5 ve Gemini 3.5 Flash modelleri*) yapılar kullanılmıştır.

Kıyaslama veri kümeleri, açık uçlu ve çoktan seçmeli olmak üzere toplamda yüz seksen beş (185) tanım sorusu barındırmaktadır. *Çoktan seçmeli ses bilgisi testi*, Anadolu Üniversitesi Açıköğretim Fakültesinin 2012–2025 yılları arasındaki ses bilgisi sorularından rastgele seçilen elli dört (54) sorudan oluşmaktadır. Soru sayısının sınırlı olması ve soruların belirli alt kategoriler gözetilmeksizin rastgele seçilmiş olması, kategorik değerlendirme yapılmasına olanak tanımamaktadır. Bu hususta testin değerlendirme kapsamı sınırlı kalmaktadır.

Bulgular bölümünde dil modellerine ait yanıtlar, *eşdeğer seçenekler arası seçim testi* esas alınarak ikinci, üçüncü ve dördüncü tablolarda sunulmuştur. Bu durum, dil modellerinin açık uçlu tanım sorularının yapısını bozmaları, boşlukları eksik veya yanlış tanımlarla doldurmaları ve bağlam dışı metinler üretmelerinden kaynaklanmaktadır. Bu sebeple çalışmada, açık uçlu tanım soruları bağımsız olarak değerlendirilip karşılaştırılamamıştır. Dil modellerinin tanımlara verdiği yanıtlar *eşdeğer seçenekler arası seçim testi* esas alınarak model bazında karşılaştırmalı biçimde yorumlanmıştır.

Çalışmadaki tanım sorularında dil modellerinin bağlamdan uzaklaşmalarının engellenmesi amacıyla, kıyaslama veri kümesindeki yanıt seçenekleri beş eşdeğer terimle sınırlandırılmıştır. Beş eşdeğer terim seçilirken, TDK'nin Güncel Türkçe Sözlük tanımlarında ve kaynakçadaki çalışmalarda yer almaları koşulu aranmıştır. Bu husus, çalışmanın ana bulgularının kapsamını sınırlandırmıştır.

Gelecekteki araştırmalarda, terim düzeyinde kapsamlı veri kümelerinin oluşturulması, eşdeğer terim sayısının sınırlandırılmaması ve farklı soru türlerinin kullanılması önerilmektedir. Soruların çeşitli kaynaklardan eşit sayıda ve sistematik biçimde seçilmesi, belirli kategoriler doğrultusunda sınıflandırılması da yararlı olacaktır. Ayrıca dil modellerinin bağlam dışı metin üretimleri, eksik veya yanlış terim tercihleri ve tanım sorularının yapısını bozma durumlarının ayrı ayrı hata kategorileri altında incelenmesi uygun olacaktır. Çok terimlilik çerçevesinin belirlenmesi ve kapsam dışında tutulan unsurların oluşturabileceği yanlışlığın değerlendirilmesi önerilmektedir. Bu doğrultuda, üretken dil modellerinin ses bilgisi becerilerine ilişkin performanslarının daha kapsamlı ve nesnel olarak değerlendirilmesi mümkün olacaktır.

Araştırma ve Yayın Etiği Beyanı

“Bu makale için etik kurul izni alınmasına gerek yoktur. Araştırma ve yayın etiğine uygun hareket edilmiştir.”

Yazarların Makaleye Olan Katkıları

“Makale tek yazarlıdır.”

Destek Beyanı

“Araştırma herhangi bir kurum veya kuruluş tarafından desteklenmemiştir.”

Çıkar Beyanı

“Makale tek yazarlıdır. Herhangi bir çıkar çatışması yoktur.”

Kaynaklar

- Aksan, D. (2020). *Her Yönüyle Dil: Ana Çizgileriyle Dilbilim*, Ankara: Türk Dil Kurumu Yayınları.
- Alibaba Cloud. (2025). Qwen3 Next (80B Cloud sürümü) [Büyük dil modeli]. *Ollama Model Kayıt Defteri*. Ollama yerel sunucu uygulaması üzerinden 22–23 Mayıs 2026 tarihlerinde <https://ollama.com/zerocopia/qwen3-next> adresinden erişildi.
- Banguoğlu, T. (1959). *Türk Grameri Ses Bilgisi*, Ankara: Türk Tarih Kurumu Basımevi.
- Banguoğlu, T. (1995). *Türkçenin Grameri*, Ankara: Türk Dil Kurumu Yayınları.
- Baran, B. (2023). Türkiye Türkçesinde Ünlülerle İlgili Ses Olaylarının Öğretimi Üzerine. *Sosyal Bilimler Araştırma Dergisi*, 12(6), 758-765.
- Baran, B. (2023). Türkiye Türkçesinde Ünsüzlerle İlgili Ses Olaylarının Öğretimi Üzerine. *Karadeniz Uluslararası Bilimsel Dergi*, 58, 134-142.
- Bayram, M. A., Fincan, A. A., Gümüş, A. S., Diri, B., Yıldırım, S. ve Aytaş, Ö. (2024). Setting standards in Turkish NLP: TR-MMLU for large language model evaluation. *arXiv preprint arXiv:2501.00593*.
- Bowden, M. (2023). A Review of Textual and Voice Processing Algorithms in the Field of Natural Language Processing. *Journal of Computing and Natural Science*, 194-203.
- Boz, E. (2001). Ünsüz Düşme ve Kaybolmalarında Terim ve Tasnif Sorunu. *Türk Dili*, 600, 856-864.
- Boz, E. (2018). Türk Dil Bilgisi Terminolojisinde Çok Terimlilik (Çok Adlılık) Sorunu. *Uluslararası Türkçe Edebiyat Kültür Eğitim Dergisi*, 7(3), 1592-1603.
- Cai, H. (2024). Research on the intersection of natural language processing and deep learning. *Applied and Computational Engineering*, 42(1), 61-66.
- Coşkun, M. V. (2010). Standart Türkçede Ses Olaylarının Sebep-Sonuç İlişkisi Çerçevesinde Yeniden Sınıflandırılması. *Bilig*, 55, 41-64.

- Deep Cogito. (2025). Cogito v2.1 (671B Cloud sürümü) [Büyük dil modeli]. *Ollama Model Kayıt Defteri*. Ollama yerel sunucu uygulaması üzerinden 22–23 Mayıs 2026 tarihlerinde <https://ollama.com/zerocopia/cogito-2.1> adresinden erişildi.
- Demir, N. ve Yılmaz, E. (2017). *Türkçe Ses Bilgisi*, Eskişehir: Anadolu Üniversitesi Yayınları.
- Dogan, E., Uzun, M. E., Uz, A., Seyrek, H. E., Zeer, A., Sevi, E., Kesgin, H. T., Yüce, M. K. ve Amasyali, M. F. (2024). Türkçe Dil Modellerinin Performans Karşılaştırması Performance Comparison of Turkish Language Models. *arXiv preprint arXiv:2404.17010*.
- Dong, J. (2024). Natural Language Processing Pretraining Language Model for Computer Intelligent Recognition Technology. *ACM Transactions on Asian and Low Resource Language Information Processing*, 23(8), 1-12.
- Durrant, P. (2013). Formulaicity in an agglutinating language: The case of Turkish. *Corpus Linguistics and Linguistic Theory*, 9(1), 1-38.
- Dursunoğlu, H. (2017). Adlandırma, Terimler ve Sınıflandırma Açısından Türkiye Türkçesinde Ünsüzlere Bir Bakış. *Türk Dili Arařtırmaları Yıllığı - Belleten*, 65(1), 27-59.
- Ediskun, H. (1999). *Türk Dil Bilgisi: Sesbilgisi, Biçimbilgisi, Cümlebilgisi*, İstanbul: Remzi Kitabevi.
- Efendioğlu, S. ve İşcan, A. (2010). Türkçe Ses Bilgisi Öğretiminde Ses Olaylarının Sınıflandırılması. *Atatürk Üniversitesi Türkiyat Arařtırmaları Enstitüsü Dergisi*, 17(43), 183-200.
- Eker, S. (2006). *Çağdaş Türk Dili*, Ankara: Grafiker Yayınları.
- Er, Y. A., Kesen, İ., Şahin, G. G. ve Erdem, A. (2026). Cetvel: A unified benchmark for evaluating language understanding, generation and cultural capacity of LLMs for Turkish. *Proceedings of the 19th Conference of the European Chapter of the Association for Computational Linguistics*, 1, 1052–1085.
- Ergin, M. (2013). *Türk Dil Bilgisi*, İstanbul: Bayrak Yayınları.
- Ersöz, S. (2018). Türkçede Göçüşme: Tanım ve Tasnif. *Uluslararası Türkçe Edebiyat Kültür Eğitim Dergisi*, 7(4), 2189-2203.
- Gencan, T. N. (2001). *Dilbilgisi*, Ankara: Türk Dil Kurumu Yayınları.
- Google DeepMind. (2026). Gemma 4 (Cloud sürümü) [Büyük dil modeli]. *Ollama Kütüphanesi*. Ollama yerel sunucu uygulaması üzerinden 22–23 Mayıs 2026 tarihlerinde <https://ollama.com/library/gemma4> adresinden erişildi.
- Google. (2026). *Gemini 3.5 Flash* (Web sürümü) [Büyük dil modeli]. Google Gemini web sitesinden 23 Mayıs 2026 tarihinde <https://gemini.google.com/app> adresinden erişildi.

- Güneş, S. (1999). *Türk Dili Bilgisi*, İzmir: Dokuz Eylül Üniversitesi Matbaası.
- Hengirmen, M. (2007). *Türkçe Dilbilgisi*, Ankara: Engin Yayınları.
- Karaağaç, G. (2013). *Türkçenin Dil Bilgisi*, Ankara: Akçağ Yayınları.
- Karademir, F. (2012). Türkçe Ses Bilgisindeki Terim ve Tasnif Kargaşası Üzerine. *VI. Uluslararası Türk Dili Kurultayı Bildirileri*, 20-25 Ekim 2008, Ankara: Türk Dil Kurumu Yayınları, 2545-2563.
- Kolaç, E. (2025). Yapay Zekâ Tabanlı Dil Modellerinde Anlamlandırma Sorunu: Geleneksel Sözlük Tanımlarının Yeniden Değerlendirilmesi. *Disiplinler Arası Dil Araştırmaları Dergisi*, 1, 307-358.
- Korkmaz, Z. (2021). Türkiye Türkçesinin Ses Bilgisi Üzerine Notlar-1. *Türk Dili*, 70(829), 4-15.
- Korkmaz, Z. (2021). Türkiye Türkçesinin Ses Bilgisi Üzerine Notlar-2. *Türk Dili*, 70(830), 4-13.
- Matarazzo, A. ve Torlone, R. (2025). A survey on large language models with some insights on their capabilities and limitations. *arXiv preprint arXiv:2501.04040*.
- MiniMax. (2026). MiniMax M2 (Cloud sürümü) [Büyük dil modeli]. *Ollama Model Kayıt Defteri*. Ollama yerel sunucu uygulaması üzerinden 22-23 Mayıs 2026 tarihlerinde <https://ollama.com/zerocopia/minimax-m2> adresinden erişildi.
- MiniMax. (2026). MiniMax M2.5 (Cloud sürümü) [Büyük dil modeli]. *Ollama Model Kayıt Defteri*. Ollama yerel sunucu uygulaması üzerinden 22-23 Mayıs 2026 tarihlerinde <https://ollama.com/zerocopia/minimax-m2.5> adresinden erişildi.
- Mistral AI. (2026). Devstral 2 (123B Cloud sürümü) [Büyük dil modeli]. *Ollama Model Kayıt Defteri*. Ollama yerel sunucu uygulaması üzerinden 22-23 Mayıs 2026 tarihlerinde <https://ollama.com/zerocopia/devstral-2> adresinden erişildi.
- NVIDIA. (2026). Nemotron 3 Super (Cloud sürümü) [Büyük dil modeli]. *Ollama Kütüphanesi*. Ollama yerel sunucu uygulaması üzerinden 22-23 Mayıs 2026 tarihlerinde <https://ollama.com/library/nemotron-3-super> adresinden erişildi.
- Oğuzkan, A. (2005). *Örneklerle Türkçe ve Kompozisyon Bilgileri*, İstanbul: İnkılap Yayınları.
- Ollama. (2026). *Ollama* [Bilgisayar yazılımı]. Açık ağırlıklı büyük dil modellerini yerel sunucuda çalıştırmak ve bu modellere erişim sağlamak amacıyla kullanılmıştır. Ollama web sitesinden 21 Mayıs 2026 tarihinde <https://ollama.com/download> adresinden erişildi.
- OpenAI. (2025). GPT OSS (120B Cloud sürümü) [Büyük dil modeli]. *Ollama Kütüphanesi*. Ollama yerel sunucu uygulaması üzerinden 22-23 Mayıs 2026 tarihlerinde <https://ollama.com/library/gpt-oss> adresinden erişildi.

- OpenAI. (2026). *ChatGPT 5.5* (Web sürümü) [Büyük dil modeli]. ChatGPT web sitesinden 23 Mayıs 2026 tarihinde <https://chatgpt.com/> adresinden erişildi.
- Ren, M. (2024). Advancements and Applications of Large Language Models in Natural Language Processing: A Comprehensive Review. *Applied and Computational Engineering*, 97(1), 55-63.
- Song, X., Hsieh, W.-C., Bi, Z., Jiang, C., Liu, M., Wang, T., Jing, B., Yang, J., Song, J., Chen, K., Chen, S., Li, M., Niu, Q., Liu, J., Peng, B., Zhang, S., Xu, J., Pan, X., Wang, J. ve Wang, Y. (2024). Large Language Models in Nonprofit and Public Good Sectors. *Center for Open Science*. <https://doi.org/10.31219/osf.io/gtfj2>
- Sriharsha, A. V., Jahnavi, M. R., Kousik, D. S., Hemanth, V., Hari, M. G. ve Vasili, P. P. (2024). Language Detection using Natural Language Processing. *Proceedings of the International Conference on Computational Innovations and Emerging Trends*, Atlantis Press International BV., 112, 507-517.
- TÜBİTAK. (2025). *Destek Süreçlerinde Üretken Yapay Zekânın (ÜYZ) Sorumlu ve Güvenilir Kullanımı Rehberi*. TÜBİTAK web sitesinden 20 Mayıs 2026 tarihinde <https://tubitak.gov.tr/> adresinden erişildi.
- Türk Dil Kurumu. (t.y.). *Güncel Türkçe Sözlük*. Türk Dil Kurumu web sitesinden 12 Mayıs 2026 tarihinde <https://sozluk.gov.tr/> adresinden erişildi.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł. ve Polosukhin, I. (2017). Attention is all you need. *Advances in neural information processing systems*, 30.
- Vural, H. ve Böler, T. (2011). *Ses ve Şekil Bilgisi*, İstanbul: Kesit Yayınları.
- Yavuz, S. ve Gürlek, M. (2012). İlköğretim-Ortaöğretim ve Yükseköğretimde Türk Dili Alanındaki Terim Farklılıkları I-Ses Bilgisi. *Turkish Studies - International Periodical for the Languages, Literature and History of Turkish or Turkic*, 7(4), 3235–3252.
- Zhipu AI. (2026). GLM 4.6 (Cloud sürümü) [Büyük dil modeli]. *Ollama Kütüphanesi*. Ollama yerel sunucu uygulaması üzerinden 22–23 Mayıs 2026 tarihlerinde <https://ollama.com/library/glm-4.6> adresinden erişildi.
- Zhipu AI. (2026). GLM 4.7 (Cloud sürümü) [Büyük dil modeli]. *Ollama Kütüphanesi*. Ollama yerel sunucu uygulaması üzerinden 22–23 Mayıs 2026 tarihlerinde <https://ollama.com/library/glm-4.7> adresinden erişildi.