



TÜRKÇE DOĞAL DİL İŞLEME TEMELLİ ÇALIŞMALARIN TEORİK DEĞERLENDİRMESİ: YÖNTEMSEL ZORLUKLAR VE GELECEK PERSPEKTİFLERİ

THEORETICAL EVALUATION OF TURKISH NATURAL LANGUAGE PROCESSING BASED STUDIES: METHODOLOGICAL CHALLENGES AND FUTURE PERSPECTIVES

Zülfü ALANOĞLU¹

<https://doi.org/10.55071/ticaretfbid.1677269>

Sorumlu Yazar
(Corresponding Author)
zalanoglu@gmail.com

Geliş Tarihi
(Received)
16.04.2025

Revizyon Tarihi
(Revised)
11.08.2025

Kabul Tarihi
(Accepted)
02.09.2025

Öz

Bu çalışma, son beş yılda Türkçe doğal dil işleme alanında gerçekleştirilen gelişmeleri, karşılaşılan metodolojik zorlukları ve geleceğe yönelik araştırma perspektiflerini kapsamlı bir şekilde ele almıştır. Türkçenin eklemeli dil yapısı ve morfolojik zenginliği, NLP alanında dilin yapısal karmaşıklığına uygun özgün yöntemlerin geliştirilmesini gerektirmektedir. Çalışmada, metin sınıflandırma, duygu analizi, soru-cevap sistemleri ve kelime gömme modelleri gibi yaygın NLP uygulamaları değerlendirilmektedir. Özellikle BERT ve GPT gibi transformer tabanlı modellerin Türkçe üzerindeki performansı ve uyarılma çalışmaları detaylandırılmıştır. Türkçe gibi düşük kaynaklı dillerde veri yetersizliğinin NLP modellerinin başarısını kısıtladığı belirtilmiş ve bu sorunun çözümüne yönelik olarak açık kaynak veri kümeleri ile veri artırma tekniklerinin sağladığı katkılar tartışılmıştır. Türkçe için geliştirilen BERTürk, BioBERTürk ve benzeri transformer tabanlı modellerin başarılı sonuçlar vermesine rağmen makine çevirisi, isim tanıma ve metin üretme gibi alanlarda daha fazla çalışmaya ihtiyaç duyulduğu belirtilmiştir. Çalışma, literatürdeki boşluklara işaret ederek Türkçeye özgü veri kaynaklarının ve NLP yöntemlerinin geliştirilmesinin, diğer eklemeli diller için de yol gösterici olabileceğini vurgulamaktadır. Sonuç olarak, bu derleme, Türkçe NLP alanında karşılaşılan mevcut zorlukları ve gelişmeleri ortaya koymakta; düşük kaynaklı dillerde etkin NLP çözümleri üretmeye yönelik öneriler sunmakta ve gelecekte yapılacak araştırmalar için kapsamlı bir yön belirlemektedir.

Anahtar Kelimeler: Düşük kaynaklı diller, morfolojik analiz, BertTURK, transformer, veri kıtlığı ve veri artırma.

Abstract

This study comprehensively addresses the developments in the field of Turkish natural language processing over the past five years, the methodological challenges encountered, and future research perspectives. The agglutinative structure and morphological richness of Turkish require the development of unique methods suitable for the structural complexity of the language in the NLP field. The study evaluates common NLP applications such as text classification, sentiment analysis, question-answer systems, and word embedding models. In particular, the performance of transformer-based models like BERT and GPT on Turkish and their adaptation studies are detailed. It is noted that data scarcity in low-resource languages like Turkish limits the success of NLP models, and discussions are provided on the contributions of open-source datasets and data augmentation techniques to address this problem. Despite the successful results of transformer-based models developed for Turkish, such as BERTürk and BioBERTürk, it is stated that further research is needed in areas such as machine translation, named entity recognition, and text generation. The study emphasizes that addressing the gaps in the literature and developing Turkish-specific data resources and NLP methods could also be informative for other agglutinative languages. In conclusion, this review highlights the current challenges and advancements encountered in the field of Turkish NLP and offers suggestions for producing effective NLP solutions in low-resource languages.

Keywords: Low resource languages, morphological analysis, BertTURK, transformer, data scarcity and data augmentation.

¹Hatay Mustafa Kemal Üniversitesi, Antakya Meslek Yüksekokulu, Bilgisayar Teknolojileri Bölümü, Hatay, Türkiye.
zalanoglu@gmail.com.

1. GİRİŞ

Doğal dil işleme (Natural Language Processing-NLP), insan dilini anlamlandırma, analiz etme ve işleme üzerine odaklanan bir disiplin olarak, yapay zekâ teknolojilerinin gelişmesiyle önemli bir araştırma alanı hâline gelmiştir. Dil tabanlı uygulamalardaki bu ilerlemeler, özellikle makine öğrenmesi ve derin öğrenme yöntemlerinin yaygınlaşması ile hız kazanmıştır. NLP çalışmaları; metin sınıflandırma, duygu analizi, soru cevap (Question Answering - QA) sistemleri ve dil modelleri gibi birçok alanda kritik çözümler sunmaktadır (Acikalin ve ark., 2020). Örneğin, chatbot'lar ve dijital asistanlar, müşteri deneyimini dönüştürürken, sosyal medya içeriklerinin duygu analizi, halkın eğilimlerini anlama konusunda stratejik veri sağlamaktadır (Doğan ve ark., 2024; Kuruca ve ark., 2022; Xu ve ark., 2022).

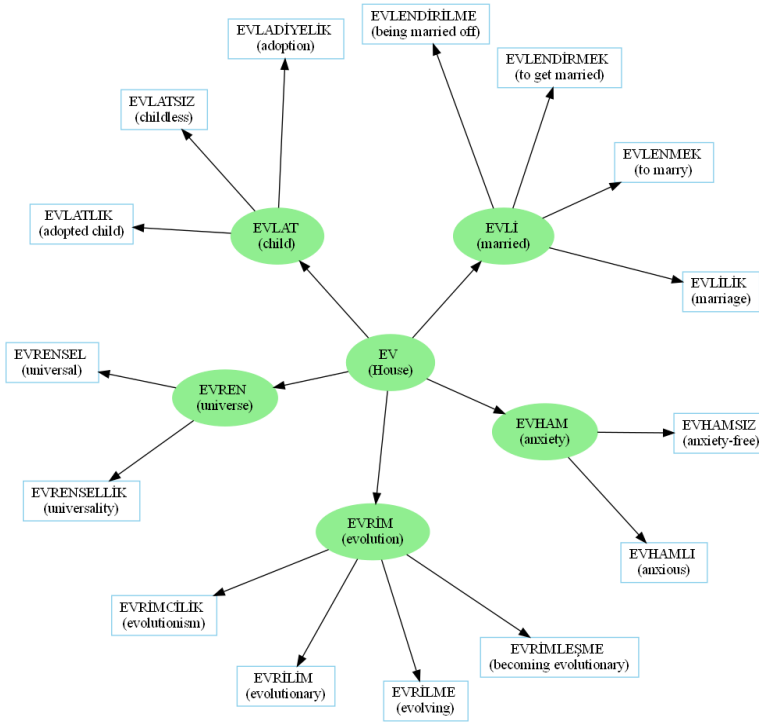
NLP teknolojilerindeki atılımlar, Transformer mimarisi ile yeni bir dönemin kapısını aralamıştır. BERT (Bidirectional Encoder Representations from Transformers) ve GPT (Generative Pre-trained Transformer) gibi modeller, dilin bağlamsal anlamını öğrenme yetenekleriyle yüksek performans sergilemiştir (Katar ve ark., 2023). Ancak, bu teknolojiler genellikle İngilizce ve Çince gibi yüksek kaynaklı dillere odaklanmıştır (Yazar ve ark., 2023). Düşük kaynaklı dillerde aynı başarıya ulaşmak için daha fazla araştırmaya ihtiyaç duyulmaktadır.

Türkçe, eklemeli yapısı ve zengin morfolojik özellikleri nedeniyle NLP modelleri için önemli zorluklar sunmaktadır. Örneğin, bir kelimenin birçok türevi farklı eklerle oluşturulabilir ve bu durum, metin işleme algoritmalarının doğruluğunu olumsuz etkileyebilir. Literatürde BERTurk ve BioBERTurk gibi Türkçe'ye özgü modeller geliştirilmiş olsa da özellikle veri eksikliği, bağlamsal anlamları yakalama ve model performansını optimize etme konularında çalışmalar yetersiz kalmaktadır.

Bu çalışmada, son beş yılda (2019-2024) Türkçe NLP alanında gerçekleştirilen araştırmalar kapsamlı bir şekilde incelenmiştir. Çalışmamız, metin sınıflandırma, duygu analizi, makine çevirisi ve QA sistemleri gibi yaygın NLP uygulamalarının performansını değerlendirerek, literatürdeki boşlukları ortaya koymayı amaçlamaktadır. Aynı zamanda, düşük kaynaklı diller için geliştirilen veri artırma ve model uyarlama tekniklerinin etkisini analiz etmekte ve gelecekteki araştırmalar için öneriler sunmaktadır.

1.1. Konunun Önemi ve Motivasyonu

Türkçe, Dünya genelinde yaklaşık 221 milyon kişi tarafından konuşulmakta olup, Türkiye başta olmak üzere çeşitli ülkelerde yaygın olarak kullanılmaktadır (Wikipedia, 2024). Bununla birlikte, Türkçe dili için NLP çalışmaları hâlâ gelişim aşamasındadır ve birçok zorluk içermektedir. Türkçe'nin eklemeli yapısı, kelimelere eklenen eklerin farklı anlamlar yaratabilmesi nedeniyle NLP modellerinde önemli zorluklara yol açmaktadır (Çavuşoğlu ve ark., 2020). Şekil 1'de görüldüğü gibi, bir fiil köküne eklenen farklı çekim ekleri ile aynı kelimenin birçok türevi ortaya çıkmakta ve bu durum, metin işleme algoritmalarının performansını olumsuz yönde etkileyebilmektedir.



Şekil 1. Türkçede Kök ve Eklerle Kelime Türetimi

Buna ek olarak, Türkçe için veri setlerinin yetersizliği ve etiketli içeriklerin eksikliği, dil modellerinin eğitilmesinde önemli bir kısıt oluşturmaktadır (Yucalar, 2023). NLP çalışmalarında kullanılan transformer tabanlı modeller, büyük dil modellerine dayandığından, yeterli büyüklükte veri setlerine ihtiyaç duymaktadır. Ancak Türkçe için erişilebilir veri kümeleri, İngilizce’ye kıyasla çok daha sınırlıdır (Şahin, 2024). Bu durum, modellerin yeterince iyi performans göstermesini engellemekte ve araştırmacıları veri artırma (data augmentation) gibi yöntemlere başvurmaya zorlamaktadır (Koksal ve ark., 2020).

Son dönemde, özellikle BERT, GPT ve benzeri transformer tabanlı modellerin Türkçe ‘ye uyarlanmasıyla önemli gelişmeler sağlanmıştır. Türkçe ‘ye özgü BERT türevleri (örn. BERTurk, BioBERTurk, DistilBERT) geliştirilerek çeşitli NLP görevlerinde başarı elde edilmiştir (Sanh, 2019; Schweter, 2020; Türkmen ve ark., 2023). Bununla birlikte, hukuki ve biyomedikal alan gibi belirli disiplinlerde daha özgün ve yerleştirilmiş modellere ihtiyaç duyulmaktadır (Öztürk ve ark., 2023).

Bu gelişmeler ışığında, NLP ‘nin Türkçe için önemi hem akademik hem de endüstriyel anlamda giderek artmaktadır. Türkçe dil işleme modelleri, otomatik metin analizleri, duygu analizi, chatbot sistemleri ve QA uygulamaları gibi birçok pratik uygulamaya olanak sağlamaktadır. Özellikle sosyal medya içeriklerinin analiz edilmesi (Ballı ve ark., 2022; Demirci ve ark., 2019; Kemaloglu ve ark., 2021; Köksal ve Özgür, 2021; Najafi ve Varol, 2024), haber metinlerinin sınıflandırılması (Alzoubi ve ark., 2023;

Aydoğan & Karci, 2020; Çelikten & Bulut, 2021; Kuyumcu ve ark., 2019) ve hukuki metinlerin işlenmesi (Akça, 2023; Mumcuoğlu ve ark., 2021; Sert ve ark., 2022) gibi alanlarda Türkçe NLP çalışmaları kritik bir rol oynamaktadır (Girgin ve ark., 2024).

Motivasyon açısından bakıldığında, Türkçe NLP'nin geliştirilmesi, yalnızca akademik literatürdeki eksiklikleri gidermekle kalmayıp, Türkçe içeriklerin daha verimli işlenmesini ve anlamlandırılmasını da sağlayacaktır. Ayrıca, dil işleme modellerinin düşük kaynaklı diller üzerindeki uygulanabilirliğini artırarak, diğer eklemeli diller için de faydalı olabilecek yöntemlerin geliştirilmesine olanak tanımaktadır (Özçift ve ark., 2021). Gelecekte, Türkçe NLP'nin daha büyük dil modellerine entegre edilmesi ve dil çeşitliliğini destekleyecek çalışmaların artırılması hedeflenmektedir (Acikgoz ve ark., 2024).

1.2. Türkçe NLP Alanında Güncel Sorunlar ve Araştırma Boşlukları

Türkçe NLP alanı, dilin eklemeli yapısı ve morfolojik zenginliği nedeniyle çeşitli zorluklar barındırmaktadır. Eklemeli dillerde her kelime kökü birçok ek alarak farklı anlamlar kazanabilmekte, bu da hem sözcük tabanlı modelleme yöntemlerinin hem de transformer tabanlı modellerin performansını sınırlandırmaktadır (Bağcı & Amasyalı, 2021). Özellikle, geleneksel makine öğrenimi algoritmaları ve morfolojik analiz tabanlı yöntemlerin bu yapısal karmaşıklığı işleyebilme kapasitesi sınırlıdır. Derin öğrenme modelleri ise yüksek kaliteli ve büyük hacimli veri kümelerine ihtiyaç duyar, ancak Türkçe gibi düşük kaynaklı diller için veri setleri kısıtlıdır (Altınok, 2023).

Bir diğer önemli zorluk, Türkçe NLP araştırmalarında veri kümesi eksikliği ve farklı veri türlerine sınırlı erişimdir. Haber metinleri, sosyal medya içerikleri ve hukuki belgeler gibi farklı içerik türlerine uygun, etiketlenmiş veri kümelerinin eksikliği, model performansını olumsuz yönde etkilemektedir (Türk ve ark., 2022). Çeşitli araştırmalar, Türkçe metinlerde duygu analizi ve metin sınıflandırma alanlarında BERT tabanlı modellerin kullanılabilirliğini göstermiş olsa da, hala bazı uygulama alanlarında performans sorunları devam etmektedir (Acikalin ve ark., 2020). Özellikle, isim tanıma (named entity recognition) (Aras ve ark., 2021; Çarık ve Yeniterzi, 2022; Kaya ve Tantug, 2022; Küçük & Can, 2019; Ozelik & Toraman, 2022) ve QA sistemleri (Ayverdi ve ark., 2020; Soygazi ve ark., 2021; Tohma ve ark., 2023) gibi görevlerde hata oranları yüksek seviyelerde seyretmektedir.

Bunlara ek olarak, dil modellerinin Türkçe gibi morfolojik açıdan zengin dillere uyarlanması sürecinde yerelleştirme sorunları yaşanmaktadır (Dündar ve ark., 2020; Uymaz & Metin, 2023a). Transformer tabanlı modellerin, İngilizce gibi yüksek kaynaklı dillerde gösterdiği başarı, Türkçe için aynı düzeyde elde edilememektedir. Bu durum, dil modellerinin eğitimi sırasında karşılaşılan veri kıtlığı ve bağlamsal anlamların yakalanamamasıyla ilişkilidir (Aytan & Sakar, 2022). Ayrıca, Türkçe'de kelimelerin eklerle türetilme şekli nedeniyle doğru anlamın modellenmesi, daha ileri seviye morfolojik ayrıştırma gerektirmektedir (Şahin, 2024).

Bunların yanı sıra, sosyal medya gibi dinamik veri kaynaklarının hızlı değişimi, veri kalitesini ve model genelleştirilebilirliğini olumsuz etkileyen faktörler arasında yer alır. Sosyal medyada kullanılan argo ve günlük konuşma dili, modellerin bağlamı doğru bir şekilde anlamasını zorlaştırmaktadır. Türkçe için etkili NLP modellerinin

geliştirilebilmesi, yalnızca veri miktarını artırmakla değil, aynı zamanda veri çeşitliliği ve kalitesini de iyileştirmekle mümkündür (Koksal ve ark., 2020).

1.3. Literatür Taramasının Amacı ve Kapsamı

Bu çalışmanın amacı, 2019-2024 yılları arasında yayımlanan Türkçe NLP literatürünü kapsamlı bir şekilde incelemek ve mevcut gelişmeleri, yöntemleri ve karşılaşılan sorunları derinlemesine ele almaktır. Türkçe için geliştirilmiş metin sınıflandırma, duygu analizi, isim tanıma, makine çevirisi ve QA sistemleri gibi uygulamalarda kullanılan yöntemlerin performans değerlendirmeleri yapılmıştır.

Bu derleme çalışması, Türkçe NLP’de karşılaşılan sorunların yanı sıra bu sorunları aşmak için önerilen çözüm yollarını da ele almaktadır. Açık kaynak veri kümelerinin önemi, veri artırma tekniklerinin etkisi ve düşük kaynaklı diller için geliştirilen metodolojiler tartışılmıştır. Aynı zamanda, NLP araştırmalarında morfolojik analizlerin rolü ve Türkçe’nin kendine özgü yapısının NLP modellerine nasıl yansıtılabileceği konusunda yeni yaklaşımlar incelenmiş ve detaylı analiz edilmiştir.

Çalışmanın bir diğer hedefi, literatürdeki boşlukları tespit ederek gelecekteki araştırmalar için yol gösterici olmaktır. NLP modellerinin Türkçe ’ye daha etkin şekilde uyarlanabilmesi için, yalnızca dil modelleri değil, aynı zamanda dil politikaları ve veri toplama stratejileri de gözden geçirilmelidir (Acikgoz ve ark., 2024). Hukuki metinler, tıp alanı ve sosyal medya içerikleri gibi farklı disiplinlerde Türkçe NLP uygulamalarının yaygınlaştırılması, dil işleme alanında sürdürülebilir çözümler geliştirilmesine katkı sağlayacaktır (Öztürk ve ark., 2023).

Bu kapsamda, çalışmanın alt amaçları şunlardır:

- Transformer tabanlı modellerin (BERTurk, GPT gibi) Türkçe metin işleme görevlerindeki başarı oranlarını analiz etmek.
- Sosyal medya, haber ve hukuki metinler gibi farklı veri türleri üzerinde mevcut NLP modellerinin karşılaştırmalı performans değerlendirmesini yapmak.
- Veri kıtlığını aşmak için kullanılan veri artırma tekniklerini tartışmak ve bu tekniklerin NLP modellerine etkisini incelemek.
- Türkçe'nin eklemeli yapısının NLP modellerine etkisini değerlendirerek morfolojik analiz ve model optimizasyonu için öneriler sunmak.
- Gelecekte yapılacak çalışmalara ışık tutacak şekilde, Türkçe NLP'nin düşük kaynaklı diller üzerindeki uygulanabilirliğini analiz etmek.

Bu literatür taraması, Türkçe NLP’nin hem akademik hem de pratik uygulamalarda nasıl geliştiğini anlamak için kapsamlı bir çerçeve sunmaktadır. Derin öğrenme ve transformer tabanlı modellerin başarısını değerlendiren çalışma, aynı zamanda geleneksel yöntemlerle yeni nesil yaklaşımları karşılaştırarak araştırmacılar için kapsamlı bir referans kaynağı olmayı amaçlamaktadır. Çalışma sonunda, literatürdeki boşluklara yönelik öneriler sunulacak ve Türkçe NLP’nin gelecekteki yönelimlerine dair çıkarımlar yapılacaktır.

2. METODOLOJİ

Bu derleme çalışmasında, 2019-2024 yılları arasında yayımlanmış Türkçe NLP konulu araştırmalar sistematik ve kapsamlı bir şekilde incelenmiştir. Bu dönem, transformer tabanlı modellerin NLP’de yaygınlaştığı ve BERT, GPT gibi modellerin farklı diller için başarıyla adapte edilmeye başlandığı kritik bir dönemi temsil etmektedir (Okur ve Sertbaş, 2021). Çalışmanın amacı doğrultusunda, literatürde Türkçe için geliştirilen metin sınıflandırma, duygu analizi, makine çevirisi, QA sistemleri ve diğer NLP uygulamalarını kapsayan yayınlar seçilmiştir. Literatür taramasının kapsamını genişletmek ve konuya dair en güncel çalışmaları dahil edebilmek için hem akademik dergi makaleleri hem de uluslararası konferans bildirileri kullanılmıştır. Aşağıda, kaynak seçimi sürecinin adımları ve kriterleri detaylı bir şekilde açıklanmaktadır.

2.1. Veri Tabanları ve Tarama Süreci

Araştırmada kullanılan yayınlar, prestijli akademik veri tabanları üzerinden erişilmiştir. Bu veri tabanları arasında Google Scholar, IEEE Xplore, Scopus, Web of Science ve ACM Digital Library gibi akademik platformlar yer almaktadır. Tarama sırasında “Turkish NLP,” “BERT in Turkish,” “Turkish sentiment analysis,” “Transformer models for Turkish,” “low-resource languages NLP,” gibi anahtar kelimeler kullanılmıştır. Arama filtreleri, yalnızca 2019-2024 yılları arasında yayımlanmış ve Türkçe NLP uygulamalarıyla doğrudan ilişkili çalışmaları içerecek şekilde ayarlanmıştır.

Seçilen yayınların yalnızca başlık ve özet bölümleri değil, tam metinleri de incelenerek çalışmanın kapsamına uygunlukları değerlendirilmiştir. NLP alanında hızla gelişen literatüre paralel olarak, literatür taramasına en güncel ve yüksek atıf potansiyeline sahip çalışmaların dahil edilmesine özen gösterilmiştir. Buna ek olarak, açık erişimli veri kümeleri sunan ve Türkçe dil modellerinin eğitiminde kullanılan yöntemleri ele alan araştırmalara özel önem verilmiştir.

2.2. Dahil Etme ve Hariç Tutma Kriterleri

Çalışmanın kapsamına uygun kaynakları seçmek için belirli dahil etme ve hariç tutma kriterleri uygulanmıştır. Dahil etme kriterleri aşağıdaki maddelerde sıralanmıştır.

- Kapsam: Çalışmanın doğrudan Türkçe diline yönelik NLP uygulamaları ve metodolojileri içermesi.
- Yöntem: Transformer tabanlı modeller derin öğrenme, geleneksel makine öğrenimi yöntemleri veya morfolojik analiz tabanlı yaklaşımlar kullanılması.
- Disiplin: Metin sınıflandırma, duygu analizi, makine çevirisi, isim tanıma ve QA sistemleri gibi uygulama alanlarını kapsaması.
- Yayın Türü: Hakemli dergi makaleleri, konferans bildirileri gibi akademik nitelik taşıyan yayınlar.

Hariç tutma kriterleri de aşağıdaki maddelerde belirtilmiştir.

- Dil: İngilizce veya başka dillerde NLP çalışmaları içerip Türkçe'ye özgü bulgular sunmayan yayınlar.
- Zaman: 2019 öncesinde yayımlanan ve güncelliğini yitirmiş çalışmalar.
- Teknoloji: Güncel NLP yaklaşımlarından uzak olan ve modern dil modellerini ele almayan yöntemler.
- Tekrar: Aynı veri setini veya modeli farklı bağlamlarda ele alan, içerik bakımından benzer çalışmaları içeren yayınlar.

2.3. Literatür Taramasının Sınırlılıkları ve Kısıtlamaları

Bu çalışmanın kapsamı, yalnızca 2019-2024 yılları arasında yayımlanan ve doğrudan Türkçe NLP uygulamalarını ele alan çalışmaları içermektedir. Ancak, Türkçe NLP'ye özgü veri setlerinin sınırlılığı nedeniyle bazı çalışmaların tekrarlı veri setlerine dayandığı görülmüştür. Ayrıca, uluslararası yayınlarda Türkçe'ye yönelik çalışmaların sayısının sınırlı olması, bazı alanlarda literatür taramasının genişletilmesini engellemiştir. Bununla birlikte, derin öğrenme ve transformer tabanlı modeller üzerine yapılan güncel araştırmalar, çalışmanın kapsamına dahil edilerek genel çerçevenin geniş tutulması sağlanmıştır.

Veri Seti Sınırlılığı: Türkçe NLP çalışmaları, İngilizce gibi yüksek kaynaklı dillere kıyasla sınırlı veri setlerine dayanmakta ve araştırmacılar veri artırma (data augmentation) gibi yöntemlere başvurmaktadır. Bu durum, bazı modellerin genelleştirme kapasitesini sınırlayabilir ve karşılaştırma yapmayı zorlaştırabilir (Türk ve ark., 2022).

Dil Yapısının Karmaşıklığı: Türkçe'nin eklemeli yapısı, özellikle kelimelerin çok sayıda ek alması nedeniyle morfolojik analizlerin karmaşıklığını artırmaktadır. Transformer tabanlı modeller, bu karmaşıklığı tam olarak işleyemeyebilir ve kullanılan veri setlerinin boyutu ve çeşitliliği, dil modellerinin bağlamı doğru anlamasını kısıtlayabilir (Bağcı & Amasyalı, 2021).

Düşük Kaynaklı Dil Problemi: Türkçe, NLP çalışmalarında düşük kaynaklı bir dil olarak kabul edilmektedir. Mevcut çalışmaların büyük bir kısmı İngilizce dil modellerine dayandığından, bu modellerin doğrudan Türkçe'ye uygulanmasında performans düşüşleri yaşanabilir (Girgin ve ark., 2024). Bu durum, literatürde karşılaştırılabilir sonuçların elde edilmesini zorlaştırmakta ve daha fazla yerleştirilmiş model geliştirilmesi gerekliliğini ortaya koymaktadır (Acikgoz ve ark., 2024).

Güncel Araştırmaların Yetersizliği: Her ne kadar çalışma 2019-2024 yılları arasındaki yayınları kapsasa da bazı NLP alanlarında hâlâ literatür eksiklikleri bulunmaktadır.

Çok Dilli Modellerle Sınırlı Karşılaştırma: Çalışma, Türkçe dil modelleri üzerine odaklandığı için çok dilli modellerin performansı ile sınırlı karşılaştırmalar yapılmıştır. Ancak, çok dilli BERT gibi modellerin Türkçe üzerindeki performansı, özellikle düşük kaynaklı dil sorunlarını ele alırken dikkate alınmıştır (Okur & Sertbaş, 2021).

3. KAPSAMLI LİTERATÜR ANALİZİ

3.1. Transformer Tabanlı Modellerin Türkçe NLP Uygulamaları

Transformer tabanlı modeller, son yıllarda NLP alanında büyük bir atılım yapmış ve birçok dilde yüksek performans göstermiştir. Transformer mimarisi, sekans-veri tabanlı modellerin kısıtlamalarını aşmak amacıyla geliştirilmiştir (Song ve ark., 2021) ve self-attention (Yao ve ark., 2021) mekanizması sayesinde dilin bağlamsal anlamını çift yönlü olarak öğrenebilme kabiliyetine sahiptir. Bu yenilik, NLP 'deki birçok uygulamada performans iyileştirmeleri sağlamıştır. Özellikle BERT ve GPT gibi modeller, dil modelleri arasında yeni bir standart oluşturmuş ve NLP 'nin pratik uygulamalarda daha yaygın kullanılmasını mümkün kılmıştır. Türkçe gibi eklemeli ve morfolojik olarak zengin diller için bu modellerin uyarlanması, NLP çalışmalarının başarısını artırmak amacıyla önemli bir alan haline gelmiştir. NLP alanında önde gelen transformer tabanlı modellerin özellikleri ve odak alanları Tablo 1 'de gösterilmiştir.

Tablo 1. Önde Gelen Transformer Tabanlı NLP Modellerinin Özellikleri ve Odak Alanları

Model	Geliştirici	Odak Alanı	Özellik
BERT	Google	Anlama, sınıflandırma	Bağlam duyarlı, çift yönlü
GPT	OpenAI	Dil üretimi	Tek yönlü, dil oluşturma
T5	Google	Çok amaçlı NLP	Her görev için metin girişi-çıkışı
RoBERTa	Facebook AI	Optimize edilmiş BERT	Uzun eğitim, daha iyi performans
XLNet	Google & CMU	Karma dil modeli	Hem oto-regresif hem oto-encoding
DistilBERT	Hugging Face	Hız ve hafızlık	Parametre azaltma
Electra	Google	Verimli eğitim	Hatalı token tahmini

Transformer tabanlı modellerin, özellikle Türkçe gibi morfolojik açıdan zengin dillerde NLP'de etkili bir çözüm sunduğu görülmektedir. BERT ve türevleri, Türkçe için yüksek doğruluk oranları sağlamış ve farklı görevlerde başarılı sonuçlar elde etmiştir (Yıldırım ve ark., 2021).

BERT modeli, bağlamı hem soldan hem de sağdan öğrenen çift yönlü bir dil modeli olarak NLP çalışmalarında çığır açmıştır. BERT'in başarısını takiben, GPT gibi ileri seviye transformer modelleri de geliştirildi ve metin üretim, özetleme, QA gibi birçok görevde kullanılmaya başlandı. Ancak, bu modellerin eğitildiği veri kümelerinin büyük çoğunluğu İngilizce gibi yüksek kaynaklı dillere odaklanmıştır. Türkçe gibi düşük kaynaklı ve morfolojik olarak karmaşık dillerde aynı performansı sağlamak için yerleştirilmiş versiyonların geliştirilmesi zorunlu hale gelmiştir (Girgin ve ark., 2024). Transformer tabanlı dil modelleri geleneksel yöntemlere kıyasla daha etkili olmaktadır. Örneğin, cosmosGPT modeli, yalnızca Türkçe verilerle eğitilmiş monolingual bir model olarak diğer çok dilli modellere kıyasla daha iyi sonuçlar elde

etmiştir (Kesgin ve ark., 2024). TURNA modeli, Türkçe'nin doğal dil anlama ve üretme görevlerinde monolingual modellerin üstün performansını göstermiştir. Bu model, özellikle mBERT ve XLM-R gibi çok dilli modellerden daha iyi sonuçlar elde etmiş ve Türkçe için hem anlama hem de üretim görevlerinde yüksek doğruluğa ulaşmıştır (Uludoğan ve ark., 2024).

BERT modellerinin yanı sıra, BERT2D gibi çift boyutlu pozisyonel gömme sistemine sahip yeni varyantlar da Türkçe NLP görevlerinde performans artışı sağlamıştır. Kaya ve Tantuğ'un çalışmasında, Türkçe'nin serbest sözdizimine uygun olarak geliştirilen BERT2D modelinin, Adlandırılmış Varlık Tanıma (Named Entity Recognition - NER) ve QA sistemleri gibi görevlerde yüksek doğruluk sunduğu görülmüştür (Yığıt Bekir Kaya ve A. Cüneyd Tantuğ, 2024). Bu kapsamda, RoBERTa-Turk ve ConvBERT gibi diğer transformer modellerinin Türkçe duygu analizi ve metin sınıflandırmada da etkili olduğu gösterilmiştir (Aytan & Sakar, 2022).

Küçük ölçekli Türkçe BERT modellerinin geliştirilmesi, düşük kaynaklı diller için önemli bir adım olarak değerlendirilmektedir. Kesgin ve arkadaşlarının (2023) geliştirdiği küçük boyutlu Türkçe BERT modelleri, daha hızlı çalışmasına olanak tanırken sıfır- atış sınıflandırma gibi görevlerde de etkili sonuçlar vermektedir. BioBERTurk modeli, Türkçe biyomedikal metinleri işlemek için özel olarak eğitilmiş bir transformer modelidir. Bu model, tıbbi belgeler, klinik raporlar ve biyomedikal araştırmalar gibi özel metinlerin anlamlandırılmasını hedeflemektedir. BioBERTurk'un performansı, genel dil modellerine göre tıbbi veriler üzerinde önemli ölçüde daha yüksektir (Türkmen ark., 2023). Yasal metinlerin işlenmesi, hukuk dilinin karmaşıklığı ve uzun cümle yapıları nedeniyle NLP uygulamaları için özel bir alan oluşturmaktadır. (Öztürk ve ark., 2023) 'de yapılan çalışma, Türk hukuk sistemi için öncelikli yasal kararların otomatik olarak tespit edilmesini hedefleyen bir transformer tabanlı modeli tanıtmaktadır.

Automated Essay Scoring (AES) sistemlerinde de Türkçe'ye özgü dil modellerinin etkin kullanımı önem kazanmıştır. Firoozi ve arkadaşları (2023), AES sistemlerinin Türkçe için etkili bir şekilde kullanılabilmesi adına BERT ve LaBSE gibi modellerin kullanıldığı çalışmaların umut verici sonuçlar sunduğunu belirtmiştir.

Türkçe NLP çalışmalarında BERT tabanlı modellerin başarısı dikkat çekici olsa da, her modelin performansı farklı görevlerde değişiklik göstermektedir. BERTurk ve mBERT (Li ve ark., 2022) gibi modeller, metin sınıflandırma ve duygu analizi gibi görevlerde karşılaştırılmıştır. Yapılan çalışmalar, BERTurk'un, Türkçe verilerle eğitildiği için mBERT'e kıyasla daha iyi performans gösterdiğini ortaya koymuştur (Okur & Sertbaş, 2021). Ancak, çok dilli bir model olan mBERT'in, dil çeşitliliği gerektiren görevlerde daha geniş bir kapsama sahip olduğu gözlemlenmiştir.

Transformer tabanlı modellerin Türkçe NLP uygulamalarına adaptasyonu, dilin eklemeli yapısının getirdiği zorluklara rağmen başarılı sonuçlar vermiştir. BERTurk ve BioBERTurk gibi yerelleştirilmiş modeller hem genel hem de özel alanlardaki metin işleme görevlerinde yüksek performans göstermiştir. Bununla birlikte, Türkçe 'de mBERT ve GPT gibi çok dilli modellerin kullanımının sınırlı kaldığı ve yerel modellerin geliştirilmesine daha fazla ihtiyaç duyulduğu görülmektedir. Bu çalışmalar, Türkçe NLP 'nin gelişimini desteklemekle kalmayıp, düşük kaynaklı diller için de örnek

teşkil etmektedir. Transformer tabanlı modellerin NLP görevlerindeki performansı Tablo 2 'de gösterilmiştir.

Tablo 2. Transformer Tabanlı Modellerin NLP Görevlerindeki Performans Analizi

Kaynak	Model	Geliştirilen Model	F1-Score	Doğruluk (%)	Kesinlik	Geri Çağırma
(Oztürk ve ark., 2023)	BERT-Base	BERTurk-Legal	0,428	43,17	0,425	-
(Firoozi ve ark., 2023)	BERT	mBERT	0,821	82,3		-
(Firoozi ve ark., 2023)	BERT	BERT-multilingual	0,969	97,1	0,968	0,972
(Akça ve ark., 2022)	BERT	Fine-tuned BERT	0,87	95	0,86	0,88
(Çelikten & Bulut, 2021)	BERT	BERTurk	0,93	-	0,94	0,93
(Uludoğan ve ark., 2024)	TURNA	TURNA-Encoder	0,956	99,46	0,952	0,96

Transformer tabanlı modellerin performansı değerlendirildiğinde, özellikle BERT tabanlı modellerin farklı versiyonlarının üstün başarı sağladığı görülmektedir. Genel olarak, BERTurk ve mBERT gibi modeller, bağlamsal anlamayı güçlendiren yapıları sayesinde yüksek doğruluk oranlarına ulaşmaktadır. TURNA gibi Türkçeye özgü geliştirilmiş modeller ise doğruluk ve güvenilirlik açısından diğer modellerle kıyaslandığında oldukça başarılı sonuçlar elde etmiştir. Buna karşın, belirli alanlara uyarlanmış bazı BERT modelleri, daha karmaşık dil yapıları içeren görevlerde performans açısından sınırlamalar göstermektedir. Transformer tabanlı modellerin bu başarısı, Türkçe gibi morfolojik olarak zengin dillerde NLP uygulamalarını daha etkili hale getirmektedir.

3.2. Türkçe Metin Sınıflandırma

Metin sınıflandırma NLP alanında yaygın olarak kullanılan tekniklerden biridir. Türkçe gibi eklemeli dillerde bu görevlerin yerine getirilmesi, dilin morfolojik yapısı ve veri kaynaklarının kısıtlı olması nedeniyle önemli zorluklar içermektedir.

Metin sınıflandırma, temel olarak metinleri önceden belirlenmiş kategorilere ayırmayı amaçlar. Metin sınıflandırma alanında, özellikle fastText gibi kelime gömme tekniklerinin kullanımı öne çıkmıştır. Kuyumcu ve arkadaşları (2019), fastText kullanılarak Türkçe verilerin ön işlem gerektirmeden sınıflandırılması üzerine bir model geliştirilmiş ve başarılı sonuçlar elde edilmişlerdir. FastText, özellikle bağlam içindeki kelimeleri dikkate alarak daha güçlü temsil özellikleri sunduğundan, BOW (Bag of Words) gibi geleneksel yöntemlerin ötesine geçmiştir. Çalışmalar, BOW'un sınırlı bir performans sağladığını, ancak fastText'in özellikle üç sınıflı duygu analizinde daha iyi sonuçlar verdiğini göstermektedir (Aksu & Karaman, 2020).

Aktif öğrenme yöntemleri de metin sınıflandırma alanında dikkate değer bir alternatif sunmaktadır. Sapcı ve arkadaşları (2022) çalışmasında, etiketlenmemiş veri örneklerinden anlamlı ve bilgi içeren örneklerin seçilerek daha küçük ama etkili veri kümeleriyle yüksek performanslı sınıflandırma modellerinin eğitildiği gösterilmiştir.

Köksal & Yılmaz (2022) çalışmalarında, Türkçe haber veri kümelerinde makine öğrenmesi ve önceden eğitilmiş dil modelleri kullanılarak yüksek başarı oranlarına ulaşılmış ve bu metodolojinin haber sınıflandırma performansını iyileştirdiği

gösterilmiştir. Türkçe haberlerin otomatik olarak sınıflandırılması, özellikle haberlerin içerdiği çoklu anlam ve uzun cümle yapıları nedeniyle zordur. Transformer tabanlı BERT modelleri, bağlamı anlamada yüksek performans gösterse de haber içeriklerinin dilsel karmaşıklığı performansı sınırlandırabilir (Okur & Sertbaş, 2021). Alzoubi ve arkadaşları (2023), farklı makine öğrenmesi algoritmalarının (SVM, Naive Bayes, LSTM, Random Forest, Logistic Regression) Türkçe metin sınıflandırma üzerindeki performansını karşılaştırmış ve LSTM algoritmasının doğruluk açısından en yüksek başarıyı sağladığını rapor etmiştir. Makine öğrenimi yöntemlerinden Naive Bayes, SVM ve rastgele ormanlar (random forests) gibi algoritmalar, hızlı ve etkin sınıflandırma yapabilmekte; ancak çoklu eklemeli yapıya sahip Türkçe gibi dillerde performansları sınırlı kalmaktadır. Ayrıca, bu yöntemler öznetelik mühendisliği gerektirdiği için daha fazla manuel işlem gerektirir (Çöltekin ve ark., 2023).

Metin sınıflandırma, özellikle hukuk ve haber gibi alanlarda, Türkçe metinleri doğru kategorilere ayırmak için geniş bir uygulama alanı bulmuştur. Akça ve arkadaşları (2022) Türkçe hukuk metinlerini sınıflandırma çalışmasında, transformers tabanlı modellerin diğer makine öğrenmesi algoritmalarından üstün olduğu gösterilmiştir. Bu çalışmada BERT gibi önceden eğitilmiş modellerin hukuk alanına uyarlanmasıyla yüksek başarı oranları elde edilmiştir. Benzer şekilde, Özkan & Kar (2022), BERT tabanlı modellerin Türkçe akademik metinlerin çoklu sınıflandırılmasında %96 doğruluk sağladığını ortaya koymuşlardır. Nöral metin sınıflandırıcılarda kelime segmentasyonu seçimleri, özellikle Türkçe gibi eklemeli dillerde sınıflandırma performansını önemli ölçüde etkileyebilir. Al Nahas ve arkadaşları (2020), Türkçe metin sınıflandırmasında farklı segmentasyon yöntemlerinin etkisini deneysel olarak değerlendirmiştir. Tablo 3 'de metin sınıflandırma modellerinin NLP görevlerindeki performansı gösterilmiştir.

Tablo 3. Metin Sınıflandırma Modellerinin NLP Görevlerindeki Performans Analizi

Kaynak	Model	Geliştirilen Model	F1-Score	Doğruluk (%)	Kesinlik	Geri Çağırma
(Türkmen ve ark., 2023)	BERTurk	BioBERTurkcon (+trR)	0,93	-	0,93	0,93
(Kuyumcu ve ark., 2019)	fastText	-	-	93,5	-	-
(Özçift ve ark., 2021)	BERT	BERT-multilingual	0,98	-	0,98	0,98
(Uymaz & Metin, 2023a)	BERT	EEM2 BERT	0,64	64,3	-	-
(Aytan & Sakar, 2022)	BERT	BERTurk	-	94	-	-
(Zorapaci, 2023)	TF-IDF ve IG	DPC	-	99,69	-	-
(Şapıcı ve ark., 2020)	Lojistik Regresyon	-	-	97	-	-
(Suncak & Aktaş, 2022)	LSTM- CNN	C-LSTM	0,87	88,5	0,87	0,86
(Aydoğan & Karci, 2020)	GRU	-	0,81	85	0,84	0,8
(Alzoubi ve ark., 2023)	SVM	-	0,77	78	0,77	0,78
(Kontuk & Turan, 2020)	SVM	-	0,64	70	0,81	0,70
(Aydoğan & Kocaman, 2023)	Naive Bayes	Naive Bayes (N-Gram)	0,82	81,9	-	-
(Özkan & Kar, 2022)	BERT	-	0,97	95	1,00	-

BERT tabanlı modeller, özellikle BERTurk ve BERT-multilingual gibi versiyonlarıyla yüksek doğruluk ve F1-score değerleri elde etmiş, en yüksek doğruluğu sağlayan modellerden biri olarak öne çıkmıştır. TF-IDF ve IG gibi geleneksel yöntemler ise %99,69 gibi çok yüksek doğruluk oranlarına ulaşarak dikkat çekmiştir. fastText ve

lojistik regresyon modelleri de güçlü performans sergilerken, Convolutional Neural Networks (CNN) ve Long Short-Term Memory (LSTM) hibrit modeli LSTM-CNN ve Gated Recurrent Unit (GRU) gibi derin öğrenme yaklaşımları istikrarlı sonuçlar sunmaktadır. Öte yandan, SVM ve Naive Bayes gibi klasik makine öğrenmesi algoritmaları nispeten daha düşük doğruluk oranlarına ulaşmıştır. Bu çeşitlilik, metin sınıflandırma performansının model seçimi ve veri özelliklerine bağlı olarak önemli ölçüde değişebileceğini göstermektedir.

3.3. Türkçe Duygu Analizi Çalışmaları

Duygu analizi, metinlerdeki duygusal ifadeleri belirleyerek pozitif, negatif veya nötr gibi kategorilere ayırmayı amaçlayan doğal dil işleme tekniklerinin kullanıldığı bir süreçtir. Duygu analizi çalışmalarında genellikle TF-IDF ve n-gram tabanlı featurizasyon yöntemleri Türkçe duygu tahmini için kullanılmıştır. Özellikle e-ticaret alanındaki yorumlar üzerinde yapılan çalışmalarda, TF-IDF ve Doc2Vec gibi yöntemlerin duygu analizinde kullanılabilirliği araştırılmıştır (Cavusoglu ve ark., 2020). Ayrıca, Türkçe'nin duygu analizi çalışmaları için kullanılan veri setleri genellikle sınırlıdır, bu nedenle çeşitli yöntemlerle elde edilen zenginleştirilmiş vektör temsillerinin duygu tespiti performansını artırdığı gözlemlenmiştir (Uymaz & Metin, 2023a).

Duygu analizi modellerinde Transformer tabanlı yaklaşımlar da araştırılmıştır. Transformer modelleri, Türkçe gibi eklemeli dillerde başarılı sonuçlar sunarak daha özgün duygu içeriklerinin tespit edilmesine yardımcı olmaktadır. Uymaz ve Metin (2023b) çalışmalarında, BERT gibi transformer tabanlı gömme modelleri kullanılarak duygu tespiti performansının arttığı gösterilmiştir. Bu çalışmada, Türkçe'nin eklemeli yapısının duygu tespitine etkisi ele alınarak, duygusal zenginleştirme yöntemlerinin performansı üzerinde durulmuştur.

Kirelli & Arslankaya (2020) çalışmalarında, küresel ısınma üzerine Twitter paylaşımları analiz edilerek, insanların çevresel sorunlara duyarlılığı ölçülmüştür. Bu çalışmada, SVM, Naive Bayes ve K-NN gibi sınıflandırıcıların performansı karşılaştırılmış ve duygu analizinde başarılı sonuçlar elde edilmiştir.

Duygu analizinde verimliliği artırmak için metin çoğaltma teknikleri de uygulanmıştır. Örneğin, Onan & Balbal (2024) çalışmalarında, Türkçe duygu sınıflandırmasını iyileştirmek amacıyla görev-özel ve evrensel dönüşümler kullanmışlardır. Bu yöntem, veri kümesinin boyutunu artırarak sınıflandırma doğruluğunu önemli ölçüde iyileştirmiştir. Ayrıca, Altınel & Şahin (2023), çok anlamlı kelimelerin farklı bağlamlarda farklı duygu skorlarına sahip olabileceğini göstererek, polisemik kelimelerin duygu analizine etkisini araştırmıştır. Bu kapsamda, polisemik kelimelerin anlam ayrıştırmasını dikkate alan özel bir modül geliştirilmiştir.

Köksal & Özgür (2021) çalışmalarında, Twitter üzerinden Türkçe duygu analizi için oluşturulan BounTi veri kümesini kullanarak BERTurk gibi Türkçe diline özgü transformer modellerinin performansını incelemiştir. Ayrıca Ahmetoğlu & Daş (2020), Türkçe film ve otel yorumları üzerinde yapılan çalışmada, Word2Vec ve GRU modellerinin duygu analizinde başarılı sonuçlar sunduğu gözlemlenmiştir. Türkçe otel yorumları üzerine gerçekleştirilen çalışma, Word2Vec gibi kelime gömme

yöntemlerinin büyük kelime havuzlarıyla eğitildiğinde daha doğru duygu sınıflandırmaları sağladığını göstermiştir. Bunun yanı sıra, Acikalin ve arkadaşları (2020), Türkçe metinlerin İngilizceye çevrilmesi yoluyla BERT modelinin kullanıldığı çalışmalarında, bu yöntemle yüksek duyarlılıklı sonuçlar elde etmişlerdir. Tokcaer (2021), Türkçe duygu analizi çalışmalarının sınırlı olduğunu belirtmekte ve literatürde açık kaynaklı kütüphaneler ile veri tabanlarının eksikliğine dikkat çekmektedir. Bu tür eksiklikler, Türkçe metinlerde duygu analizi yapmak isteyen araştırmacılar için önemli bir engel teşkil etmektedir.

Türkçe tweet'lerle yapılan duygu analizi çalışmaları, özellikle derin öğrenme modellerinin doğruluk oranını artırdığını göstermektedir. Ancak geleneksel makine öğrenimi algoritmaları, verinin özelliklerine bağlı olarak daha hızlı sonuçlar verebilmektedir. Örneğin, Naïve Bayes ve destek vektör makineleri (SVM) gibi yöntemler, küçük veri setlerinde tatmin edici sonuçlar sağlarken, daha büyük veri setlerinde derin öğrenme modelleri öne çıkmaktadır (Koksall ve ark., 2020).

Duygu analizinde kullanılan modeller, temel olarak makine öğrenimi ve derin öğrenme yaklaşımlarına ayrılmaktadır. Makine öğrenimi yöntemleri, özellikle küçük veri kümeleri üzerinde etkili olsa da derin öğrenme modelleri daha karmaşık ve büyük veri setlerinde üstün performans göstermektedir.

Duygu tespitinde kullanılan kelime gömme yöntemleri, metinlerin özniteliklerinin doğru şekilde temsil edilmesinde kritik bir rol oynamaktadır. Türkçe metinlerde kullanılan başlıca gömme yöntemleri arasında Word2Vec, GloVe ve FastText yer alır. Word2Vec ve GloVe kelimelerin vektörlerle ifade edilmesini sağlar ve metinler arasındaki semantik ilişkileri yakalamada etkilidir. Ancak, Türkçe'nin eklemeli yapısı nedeniyle kelime gömme performansında bazı sınırlamalar ortaya çıkmaktadır (Bağcı & Amasyalı, 2021). FastText kelime parçacıklarını kullanarak kelimeleri daha doğru temsil eder ve Türkçe gibi morfolojik olarak zengin diller için daha uygun bir seçenektir (Türk ve ark., 2022). Transformer tabanlı modeller ise doğrudan bağlamsal gömme sağlayarak, öznitelik seçimine duyulan ihtiyacı azaltmaktadır. Bu, model performansını artırmakta ve manuel öznitelik mühendisliğini en aza indirmektedir (Girgin ve ark., 2024).

Sonuç olarak Türkçe metinlerin sınıflandırılması ve duygu analizi, sosyal medya, haber ve kullanıcı yorumu metinlerinde büyük önem taşımaktadır. Sosyal medya verilerinin düzensiz yapısı, haber metinlerinin dilsel karmaşıklığı ve yorumların kısa olması, NLP modellerinin performansını etkileyen başlıca faktörlerdir. Makine öğrenimi yöntemleri, küçük veri kümelerinde etkili olurken, derin öğrenme tabanlı yaklaşımlar, daha büyük ve karmaşık veri kümelerinde üstünlük sağlamaktadır. Özellikle kelime gömme yöntemleri, metinlerin öznitelik çıkarımında önemli bir rol oynarken, transformer tabanlı modeller bağlamsal anlamları yakalamada öne çıkmaktadır. Ancak, Türkçe'nin dil yapısından kaynaklanan zorluklar, veri setlerinin çeşitliliği ve kalitesini artırmaya yönelik çalışmaları gerekli kılmaktadır.

Türkçe duygu analizi çalışmalarında kullanılan farklı makine öğrenmesi ve derin öğrenme modellerinin performansları, çeşitli ölçüm kriterlerine göre karşılaştırılabilir. Tablo 4 'de, yaygın olarak kullanılan modellerin F1-score, doğruluk, kesinlik ve recall değerleri sunulmuştur.

Türkçe duygu analizinde kullanılan farklı makine öğrenmesi ve derin öğrenme modelleri, çeşitli başarı oranları göstermektedir. BERT tabanlı modeller, özellikle Multilingual BERT ve BerTurk gibi versiyonlarıyla, yüksek doğruluk ve F1-score değerleriyle öne çıkmaktadır. ConvBERT, DistilBERT ve BERT-base gibi modeller de güçlü performans sergilerken; alternatif olarak GRU, CNN ve ensemble sınıflandırıcılar dikkate değer doğruluk oranlarına ulaşmıştır. SVM ve LSTM gibi geleneksel makine öğrenmesi modelleri ise daha tutarlı fakat nispeten düşük başarı oranları sunmaktadır. Bu modellerin performansı, kullanılan veri setine, model yapılandırılmasına ve uygulama alanlarına göre değişiklik göstermekte olup, BERT ve GRU tabanlı modeller genellikle daha yüksek doğruluk sağlamaktadır.

Tablo 4. Duygu Analizi Modellerinin NLP Görevlerindeki Performans Analizi

Kaynak	Model	Geliştirilen Model	F1-Score	Doğruluk (%)	Kesinlik	Geri Çağırma
(Cavusoglu ve ark., 2020)	TF-IDF	n-gram tabanlı TF-IDF	0,924	-	0,88	-
(Özçift ve ark., 2021)	BERT	BERT-multilingual	0,99	-	0,99	0,99
(Aytan & Sakar, 2022)	BERT	ConvBERT	0,929	-	-	-
(Ballı ve ark., 2022)	SVM			87,47	0,86	0,87
(Çam & Özgür, 2023)	BERT	BERT-base-offensive-2	0,584		0,64	0,57
(Tunali, 2022)	BERT	DistilBERT	0,78	Çok Yüksek	-	-
(Girgin & Şahin, 2023)	-	GDELT	0,926	91,56	-	-
(Yılmaz & Toraman, 2021)	BERT	BerTurk	0,812	91,2	0,91	-
(Dündar ve ark., 2020)	BERT	-	0,899	90,28	0,91	-
(Safaya ve ark., 2022)	ERSAT Z	-	0,89		0,88	0,86
(Ba Alawi & Bozkurt, 2024)	CNN	-	0,837	88,11	0,85	0,82
(Kirelli & Arslankaya, 2020)	K-NN	-	0,746	74,63	0,748	0,75
(Tohma ve ark., 2023)	Linear SVM	-	0,852	84,77	0,86	0,84
(Kemaloğlu ve ark., 2021)	LSTM	-	0,84	84,46	0,84	0,84
(Ahmetoğlu & Daş, 2020)	GRU	Kelime Vektör Modeli	0,94	94	0,93	0,93

3.4. Sosyal Medya İçeriklerinin NLP ile Analizi ve Yanıltıcı Bilgi Tespiti

Sosyal medya platformları, bilgi yayılımının hızlı gerçekleştiği ve kullanıcıların duygu, düşünce ve eğilimlerini serbestçe ifade ettiği dijital alanlar olarak günümüzde büyük önem taşımaktadır. Bununla birlikte, sosyal medya platformlarında yayılan yanıltıcı bilgiler, yanlış bilgilendirme ve dezenformasyon gibi sorunlar, kamuoyunu olumsuz etkileyebilmektedir. Bu nedenle, NLP teknikleri ile sosyal medya içeriklerinin analiz edilmesi hem duygu ve tema analizi hem de yanıltıcı bilgilerin tespiti için kritik öneme sahiptir (Freiling, 2019; Zovikoğlu & Çetin, 2024).

Küçük ve Can (2019) çalışmalarında, Türkçe tweet veriseti, NER ve tutum bilgileriyle etiketlenerek, sosyal medya içeriklerinde bu iki tekniğin ilişkisinin incelenmesine olanak sağlamışlardır. Araştırmada, NER'in tutum analizi performansını artırdığı gözlemlenmiştir. Ayrıca Çarık ve Yeniterzi (2022) çalışmalarında, 5000 Türkçe tweeti, yüksek anlaşma oranıyla NER amacıyla etiketlenmiştir. Bu çalışmada, özellikle kişi, organizasyon, lokasyon, zaman, para, ürün gibi farklı adlandırılmış varlık türlerinin tespiti yapılmış ve bu türlerin, sosyal medya içeriklerinde doğru bilgiye ulaşma sürecini hızlandırdığı vurgulanmıştır.

COVID-19 pandemisi sürecinde, yanıltıcı bilgi tespiti için çeşitli NLP modelleri geliştirilmiştir. Toraman ve arkadaşları (2022) tarafından geliştirilen ARC-NLP modeli, COVID-19 bağlamında yanıltıcı tweetlerin gerçek yaşam olaylarıyla çelişip çelişmediğini analiz ederek zararlı içerik tespiti yapmıştır. Bozuyulu & Özçift (2022) tarafından geliştirilen COVID-19 temalı sahte haber tespit modeli ise BerTURK gibi Türkçe odaklı derin öğrenme tabanlı dil modelleri kullanarak yüksek doğruluk oranlarına ulaşmıştır. Zovikoğlu & Çetin (2024) ise, yanıltıcı bilgi tespitinde lojistik regresyon ve cümle tabanlı transformer modelleri kullanarak, yanıltıcı içeriğin başlığının haberin içeriğinden daha belirleyici bir faktör olduğunu bulmuşlardır.

Türkçe sosyal medya içeriklerinde siber zorbalık ve nefret söylemi tespiti, sosyal medya analizinin bir diğer önemli araştırma alanıdır. Balcıoğlu (2024), Türkçe tweetlerde siber zorbalığı tespit etmek için SVM ve Rastgele Orman algoritmalarını kullanmış ve Zemberek-NLP gibi Türkçe odaklı NLP araçlarıyla model performansını artırmıştır. Aynı doğrultuda, Çam & Özgür (2023), BERT ve ChatGPT tabanlı modellerin Türkçe nefret söylemi tespitindeki performansını karşılaştırmış, ChatGPT'nin de bu konuda umut verici sonuçlar sunduğunu göstermiştir.

Sosyal medya içeriklerindeki yazım hataları, kısaltmalar ve yanlış yazımlar gibi gürültüler, NLP modellerinin performansını olumsuz etkileyebilmektedir. Bu soruna yönelik olarak Demir & Topcu (2022), grafik tabanlı bir normalizasyon yaklaşımı geliştirerek, Türkçe metinlerin anlamlı hale getirilmesi sürecinde büyük bir iyileşme sağlandığını ortaya koymuşlardır. FastText gibi modellerin, sosyal medya içeriklerinde yeni veya hatalı yazımlar için sağladığı esneklik, duygu analizinde öne çıkan bir özellik olmuştur (Joulin ve ark., 2016).

Tablo 5. Sosyal Medya İçeriklerinin NLP Analizi ve Yanıltıcı Bilgi Tespiti ile ilgili Modellerinin NLP Görevlerindeki Performans Analizi

Kaynak	Model	Geliştirilen Model	F1-Score	Doğruluk (%)	Kesinlik	Geri Çağırma
(Suncak & Aktaş, 2022)	Logistic Regression	-	-	93,81	-	-
(Toraman ve ark., 2022)	BERTurk	ARC-NLP-contr	0,600	79,4	-	-
(Balcıoğlu, 2024)	SVM	-	0,91	95,9	0,9081	0,912
(Koru & Uluyol, 2024)	BERT+ CNN	BERTurk + CNN	0,94	94	-	-
(Demir & Topcu, 2022)	-	Grafik Tabanlı Normalizatör	0,665	-	0,678	0,653
(Demirci ve ark., 2019)	Random Forest	-	0,81	81,74	-	-
(Ballı ve ark., 2022)	SVM	-	-	87,47	-	-
(Carik & Yeniterzi, 2021)	BERT	BERTurk Ensemble	0,7342	-	0,6666	-
(Ozdemir & Yeniterzi, 2020)	BERTurk	Ensemble	0,816	88,3	-	-

Sosyal medya içeriklerinin NLP ile analizi, farklı görevlerde çeşitli zorluklar ve avantajlar sunmaktadır. Yanıltıcı bilgi tespiti ve doğrulama çalışmaları, BERT ve mBERT gibi transformer tabanlı modellerin bağlamı anlamadaki başarısını göstermektedir. Ancak, çok dilli modeller, Türkçe gibi dillerde performans kaybı yaşayabilirken, Türkçe 'ye özel geliştirilen modeller daha iyi sonuçlar vermektedir. Sosyal medya içeriklerinin analizi, yanıltıcı bilgi tespiti ve duygu analizi gibi görevlerde, doğru ve çeşitli veri kümelerinin kullanılmasının önemini ortaya

koymaktadır. Transformer tabanlı modeller, sosyal medya analizlerinde geleneksel yöntemlere kıyasla üstünlük sağlasa da veri kalitesi ve çeşitliliği, model performansının kritik belirleyicileri olarak öne çıkmaktadır. Tablo 5 ‘deki çalışmalar, sosyal medya içeriklerinde yanıltıcı bilgi tespiti amacıyla farklı makine öğrenmesi ve derin öğrenme modellerinin performanslarını karşılaştırmaktadır.

Logistic Regression ve SVM gibi klasik yöntemler iyi sonuçlar elde ederken, özellikle BERTurk temelli modeller daha yüksek doğruluk ve güvenilirlik sağlamıştır. Grafik tabanlı normalizatör göreceli olarak daha düşük performans sergilemiştir. Bu analizler, yanıltıcı bilgi tespitinde BERT tabanlı modellerin etkili olduğunu göstermektedir.

3.5. Adlandırılmış Varlık Tanıma ve Soru-Cevap Sistemleri

NER, metin içerisindeki kişi, yer, organizasyon gibi özel adların tanımlanması ve kategorize edilmesi sürecidir. Türkçe dilinde NER, morfolojik açıdan zengin ve eklemeli yapısı nedeniyle diğer dillere göre daha karmaşık bir yapıya sahiptir (Aras ve ark., 2021; Arslan & Eryiğit, 2023).

Türkçe'deki NER çalışmalarında BiLSTM ve Transformer tabanlı modellerin etkili olduğu görülmüştür. Örneğin, BiLSTM modelleri, karakter, kelime ve morfolojik özellikleri birleştirerek başarılı sonuçlar elde etmektedir. Bununla birlikte, son dönemde BERTurk gibi Türkçe'ye özel transformer modelleri kullanılarak daha yüksek doğruluk oranlarına ulaşılmıştır. Aras ve arkadaşları (2021), BiLSTM-CRF tabanlı bir modelle Türkçe NER alanında yeni bir başarı seviyesine ulaşmıştır. Transformer tabanlı BERT modelleri, bağlamsal bilgi yakalama yetenekleri ile NER görevlerinde öne çıkmaktadır. BERTurk gibi Türkçe için optimize edilmiş modeller, standart NER görevlerinde yüksek doğruluk oranlarına ulaşmıştır. mBERT gibi çok dilli modellerin performansı, farklı dil yapılarını işleyebilme yeteneği sayesinde genel olarak iyidir; ancak Türkçe'ye özgü modellerin, bu dil için daha iyi sonuçlar sağladığı gözlemlenmiştir (Okur & Sertbaş, 2021). Türkçe hukuki metinlerde NER üzerine yapılan çalışmalar, yasal dokümanlardaki karmaşık yapılar nedeniyle özel bir önem taşır. Çetindağ ve arkadaşları (2023) çalışmalarında, BiLSTM ve CRF tabanlı bir model kullanarak, Türkçe hukuk metinlerinde %92,27 F1 skoruna ulaşmışlardır. Tablo 6 ‘da NER modellerinin NLP görevlerindeki performansları detaylı bir şekilde analiz edilmiştir.

Tablo 6. NER Modellerinin NLP Görevlerindeki Performans Analizi

Kaynak	Model	Geliştirilen Model	F1-Score	Doğruluk (%)	Kesinlik	Geri Çağırma
(Altınok, 2023)	spaCy	tr_core_news_trf	0,91	90	0,92	-
(Aras ve ark., 2021)	BERT	BERTurk + CRF	0,96	99,42	0,956	0,963
(Küçük & Can, 2019)	SVM	-	0,84	-	0,85	0,83
(Kaya & Tantug, 2022)	BERT	ITUTurkBERT	0,937	-	-	-
(Çetindağ ve ark., 2023)	CNN-LSTM-CRF	-	0,966	-	-	-

QA sistemleri, NLP alanında, metin veya dokümanlar üzerinden kullanıcının sorduğu sorulara en uygun cevapları bulmayı amaçlamaktadır. Özellikle Türkçe gibi morfolojik olarak zengin dillerde bu tür sistemlerin geliştirilmesi, veri setlerinin yetersizliği ve

dilin yapısal farklılıkları nedeniyle zorluk teşkil etmektedir (Gemirter & Goularas, 2021; Soygazi ve ark., 2021).

THQuAD, Türkçe dilinde tarihi metinler üzerinde geliştirilmiş bir okuma anlama veri setidir. Bu veri seti, Osmanlı tarihi ve İslam bilimleri gibi belirli konularda sorulara yanıt bulmak için tasarlanmıştır (Soygazi ve ark., 2021). Bir diğer çalışma, Türk bankacılık sektörüne yönelik bir QA sistemi geliştirmiştir. BERT tabanlı bu sistem, sorulara en uygun yanıtları bulmak için geniş veri setleri üzerinde ince ayar (fine-tuning) yapılmıştır ve Türkçe gibi karmaşık dillerde başarılı sonuçlar elde edilmiştir (Gemirter & Goularas, 2021).

Türkçe QA sistemlerinde duygu analizi kullanılarak insan-robot etkileşimini daha gerçekçi hale getirmek mümkün olmuştur. Tohma ve arkadaşları (2023) çalışmalarında, insan-robot etkileşiminde duygusal tepkiyi artırmak için duygu analizi entegre etmişlerdir. Bir başka önemli araştırma ise, zaman ve nesne temelli Türkçe QA sistemleri üzerine odaklanmıştır. Bu sistemlerde, tarihsel olaylar veya nesnelerle ilgili sorulara yanıt bulmak için zaman damgaları ve nesne tanımlayıcıları kullanılmıştır (Ayverdi ve ark., 2020).

Türkçe dilinde NER ve QA sistemleri üzerine yapılan çalışmalar, dilin morfolojik zenginliği ve veri seti eksiklikleri gibi zorluklara rağmen, özellikle derin öğrenme ve transformer tabanlı modellerle önemli ilerlemeler kaydetmiştir. QA görevlerinde kullanılan BERT tabanlı yaklaşımlar, sorunun bağlamını doğru şekilde analiz ederek cevap üretir. Geleneksel CRF (Panchendrarajan ve Amaresan, 2018) ve HMM tabanlı NER modelleri (Morwal ve ark., 2012), genellikle özellik mühendisliği gerektirir ve Türkçe gibi karmaşık morfolojik yapılara sahip dillerde düşük performans gösterir. Buna karşılık, BERTurk modeli, kelimelerin bağlam içindeki anlamını kavrayarak çok daha yüksek doğruluk oranlarına ulaşmıştır (Türk ve ark., 2022).

NER ve QA sistemleri, Türkçe NLP çalışmaları için kritik öneme sahiptir. BERT tabanlı modellerin bu görevlerdeki başarısı, bağlamsal bilgi yakalama yeteneği sayesinde diğer modellere göre üstünlük sağlamaktadır. İsim tanıma görevlerinde BERTurk, çok dilli modellere kıyasla daha iyi sonuçlar verirken, THQuAD veri kümesi ile yapılan çalışmalar, derin öğrenme tabanlı modellerin Türkçe QA sistemlerinde etkili olduğunu ortaya koymaktadır. Ancak, Türkçe için daha fazla veri kümesine ve yerleştirilmiş modellere ihtiyaç duyulmaktadır.

Bu çalışmalar, yalnızca Türkçe NLP'nin gelişimine katkı sağlamakla kalmayıp, aynı zamanda düşük kaynaklı diğer diller için de yol gösterici niteliktedir. Gelecekte, daha fazla veri kümesi ve model geliştirilmesiyle Türkçe NER ve QA sistemlerinin daha da iyileştirilmesi beklenmektedir.

Tablo 7'deki sonuçlar, Türkçe dilinde QA sistemlerinin performansını ölçmek için kullanılan çeşitli modellerin karşılaştırmalı analizini sunmaktadır.

Tablo 7. Soru Cevap Sistemleri ile İlgili Modellerin NLP Görevlerindeki Performans Analizi

Kaynak	Model	Geliştirilen Model	F1-Score	EM (%)
(Gemirter & Goularas, 2021)	BERT	Derin öğrenme tabanlı QA sistemi	79,01	54,09
(Alecakir ve ark., 2022)	LSTM	BiLSTM-CRF	67,39	-
(Soygazi ve ark., 2021)	ELECTRA	-	81,55	63,08

BERT tabanlı model, bağlamsal anlamlandırma konusunda oldukça başarılı olurken, BiLSTM-CRF modeli nispeten daha düşük performans göstermiştir. ELECTRA, bağlamsal ilişkileri öğrenme kapasitesiyle dikkat çekerken, BERT de kapsamlı dil anlama yeteneği sunmaktadır. Genel olarak, modern derin öğrenme modellerinin Türkçe NLP görevlerinde geleneksel yöntemlere göre daha üstün performans sergilediği görülmektedir.

Transformer tabanlı BERT modelleri, özellikle BERTurk ve ITUTurkBERT gibi versiyonlarla, yüksek doğruluk ve güvenilirlik oranlarına ulaşarak NER görevlerinde öne çıkmıştır. Bu modeller, bağlamsal dil ilişkilerini öğrenme yetenekleri sayesinde geleneksel yöntemlere göre üstün performans sergilemektedir. Bununla birlikte, CNN-LSTM-CRF gibi derin öğrenme kombinasyonları da dikkat çekici performans göstermiştir. Daha geleneksel bir yöntem olan SVM ise, modern derin öğrenme yaklaşımlarına kıyasla daha sınırlı bir performans sergilemektedir.

3.6. Düşük Kaynaklı Diller için Model Geliştirme ve Veri Artırma Teknikleri

Düşük kaynaklı diller, sınırlı dil kaynaklarına ve veri setlerine sahip olmaları nedeniyle NLP alanında zorluklarla karşılaşmaktadır. Türkçe gibi diller, zengin morfolojik yapıları ve özgün dil bilgisel kuralları nedeniyle özellikle zordur. Son yıllarda, düşük kaynaklı dillerde model geliştirme ve veri artırma teknikleri için çeşitli yöntemler araştırılmıştır.

Düşük kaynaklı diller için büyük dil modellerinin (Large Language Models - LLMs) adaptasyonu, mevcut kaynakların yetersizliği nedeniyle zorluk arz etmektedir. Örneğin, Türkçe üzerinde yapılan çalışmalarda, İngilizce'den önceden eğitilmiş modellerin Türkçe'ye adaptasyonu başarı ile gerçekleştirilmiştir (Nasution & Onan, 2024). Acikgöz ve arkadaşları (2024) çalışmalarında, Türkçe dil modeli geliştirmede veri artırma ve sürekli ön-eğitim stratejilerini detaylandırmıştır. Bu çalışmada, İngilizce modellerden Türkçe'ye bilgi transferi sağlanmış ve yeni Türkçe veri setleri oluşturularak modellerin performansı artırılmıştır.

Türkçe gibi düşük kaynaklı dillerde etiketlenmiş veri eksikliği, veri artırma yöntemlerini önemli kılmaktadır. Veri artırma teknikleri, veri setlerinin çeşitliliğini artırarak bu sorunun aşılmasına katkıda bulunmaktadır (Feng ve ark., 2021). Çeşitli yöntemler arasında çevrim içi veri üretimi (back-translation), paraphrasing ve rastgele kelime ekleme veya silme gibi teknikler yer almaktadır (Çöltekin ve ark., 2023). Back-translation, bir metnin başka bir dile çevrilip yeniden Türkçe'ye çevrilmesi ile yeni veri setleri oluşturulmasını sağlamaktadır (Avşaroğlu & Karadağ, 2019). Paraphrasing teknikleri ise aynı anlamı farklı ifadelerle yeniden yazma sürecini içermektedir (Karaoğlu ve ark., 2019).

İngilizce'den Türkçe'ye makine çevirisi kullanılarak büyük veri setlerinin Türkçe'ye kazandırılması ile ilgili çalışmalar yapılmıştır (Budur ve ark., 2020). Çalışmalarda, otomatik çeviri yoluyla elde edilen veri setlerinin, insan çevirisi ile karşılaştırıldığında yüksek kaliteye sahip olduğu gösterilmiştir. Türkçe biyomedikal alanında yapılan çalışmalarda, domain-spesifik ön-eğitim stratejilerinin model performansını artırdığı gözlemlenmiştir. BioBERTurk ve TurkRadBERT gibi modeller, Türkçe biyomedikal metinler kullanılarak sürekli ön-eğitim stratejisi ile geliştirilmiştir (Türkmen ve ark., 2023). Bu modellerde, genel dil modellerinin Türkçe biyomedikal veri setleri ile yeniden eğitilmesi, sınırlı veri ile daha yüksek doğruluk oranları elde edilmesine olanak sağlamıştır.

Türkçe gibi düşük kaynaklı dillerde, manuel veri etiketleme maliyetli ve zaman alıcıdır. Ancak, son yıllarda LLMs kullanılarak otomatik etiketleme yöntemleri geliştirilmektedir. Nasution & Onan (2024), insan ve model tabanlı etiketleme tekniklerini karşılaştırmış ve LLMs'lerin özellikle duygu analizi gibi görevlerde rekabetçi sonuçlar elde ettiğini bildirmişlerdir.

Tarihsel dillerde veri eksikliği daha büyük bir sorun teşkil etmektedir. Osmanlı Türkçesi için bağımlılık anotasyonu yapılmış ve BERT tabanlı bir model kullanılarak etiketleme süreci hızlandırılmıştır (Özateş ve ark., 2024).

Düşük kaynaklı diller, büyük veri setleri ve etiketlenmiş içeriklerin yetersizliği zorluklarını aşmak için geliştirilen yöntemlerden biri, transfer öğrenimi (transfer learning) (Zhuang ve ark., 2020) ve çok dilli dil modellerinin (multilingual models) (Srinivasan ve ark., 2021) kullanımudur. Transfer öğrenimi, başka dillerde önceden eğitilmiş modellerin düşük kaynaklı bir dile uyarlanmasını sağlar. Özellikle, çok dilli modeller olan mBERT (Conneau, 2019), XLM-R (Barbieri ve ark., 2021) ve mT5 (Xue, 2020) gibi modeller, birden fazla dilde çalışabilme yeteneği sunarmakta ve düşük kaynaklı dillerde veri eksikliğini kısmen telafi etmektedir (Acikgoz ve ark., 2024; Nezhad & Agrawal, 2023). Ancak, çok dilli modellerin başarısı, dilin yapısına göre değişiklik göstermektedir (Muller ve ark., 2023). Özellikle Türkçe'nin eklemeli yapısı nedeniyle çok dilli modellerin bazı durumlarda anlam kaybına uğradığı ve doğruluk oranlarının düştüğü gözlemlenmiştir (Acikalin ve ark., 2020; Karayigit ve ark., 2022).

Tablo 8. Düşük Kaynaklı Diller için Geliştirilen Veri Arttırma Modellerinin NLP Görevlerindeki Performans Analizi

Kaynak	Model	Geliştirilen Model	Doğruluk (%)	F1 Score	Kesinlik	Geri Çağırma
(Acikgoz ve ark., 2024)	GPT2	HamzaMistral	39,85		-	-
(Kaya&Tantuğ, 2024)	BERT	BERT2D	94,20	0,94	-	-
(Kesgin ve ark., 2023)	BERT	-	91,85		-	-
(Türkmen ve ark., 2023)	BERT	TurkRadBERT	-	0,96	0,97	0,95
(Karagöz ve ark., 2024)	BERT	BERTurk + PT	-	0,89	0,82	0,90

Sonuç olarak, düşük kaynaklı diller için model geliştirme süreci hem çok dilli hem de yerelleştirilmiş modellerin güçlü yönlerinden yararlanmayı gerektirir. Türkçe için geliştirilen NLP uygulamalarının, veri artırma ve transfer öğrenimi gibi yöntemlerle daha da iyileştirilmesi mümkündür. Gelecekte, Türkçe gibi dillerde NLP çalışmalarını destekleyecek daha fazla veri kümesi ve yerel modelin geliştirilmesi beklenmektedir.

Tablo 8’de, düşük kaynaklı dillerde gerçekleştirilen model geliştirme çalışmaları, çeşitli önceden eğitilmiş dil modelleri kullanılarak farklı metriklerle değerlendirilmiştir.

GPT-2 tabanlı HamzaMistral modeli, doğruluk açısından sınırlı bir performans sergilemiştir. Bunun aksine, BERT tabanlı modeller, özellikle BERT2D ve TurkRadBERT, yüksek doğruluk ve F1 skoru ile öne çıkmıştır. Özellikle TurkRadBERT modeli, geri çağırma ve kesinlik metriklerinde etkileyici sonuçlar elde etmiştir. Ayrıca, BERTurk modeline ek olarak ön eğitim (Pre-Training-PT) ile geliştirilen modelin F1 skoru bir miktar daha düşük olsa da, kesinlik açısından güçlü bir performans sergilemiştir. Bu durum, dil modellerinin yerleştirilmiş verilerle eğitildiğinde performans artışı sağladığını göstermektedir.

3.7. Türkçe İçin Morfolojik Analiz ve Eklemeli Dillerin Zorlukları

Türkçe, eklemeli (agglutinative) dil yapısına sahip olup, bu durum morfolojik analiz süreçlerini diğer dillere kıyasla daha karmaşık hale getirmektedir (Tohma & Kutlu, 2020). Eklemeli dillerde, kelimeler köklere eklenen çeşitli ekler aracılığıyla türetilir, bu da morfolojik çözümleme süreçlerinde zorluklara yol açar (Güler & Tantug, 2020). Türkçe gibi dillerde morfolojik yapı, bir kelimenin anlamını ve dil bilgisel özelliklerini kök ve eklerle belirler. Bu nedenle, NLP ve ilgili alanlardaki araştırmalarda Türkçe'nin morfolojik analizine dair çeşitli yaklaşımlar geliştirilmiştir.

Eklemeli dillerde, kelimeler köklerine ekler eklenerek türetilir. Örneğin, Türkçede bir fiil köküne eklenen ekler, kelimenin anlamını ve dil bilgisel fonksiyonunu önemli ölçüde değiştirebilir. Bu yapı, kelime köklerinin ve eklerin doğru bir şekilde ayrıştırılmasını gerektirir. Ancak bu ayrıştırma süreci, eklemeli dillerdeki eklerin çok çeşitli anlam ve fonksiyonları nedeniyle karmaşıktır (Suncak & Aktaş, 2021) (Boltayevich ve ark., 2023). Özellikle, sözcük türü etiketi belirleme (part-of-speech - POS) etiketleme ve kök bulma (stemming) gibi görevler, eklerin farklı bağlamlarda farklı anlamlar taşıması nedeniyle karmaşık hale gelir. Türkçe gibi eklemeli POS etiketleme ve kelime ayrıştırma görevleri, dilin yapısal özelliklerinden dolayı zorluklar içerir (Boltayevich ve ark., 2023; Bölücü & Can, 2019). Aynı kelime köküne eklenen farklı eklerin, kelimenin cümle içindeki görevini değiştirmesi, POS etiketlemenin doğruluğunu olumsuz etkileyebilir (Adalı & Adamov, 2016). POS etiketleme sürecinde karşılaşılan diğer bir zorluk, Türkçe ’deki kelimelerin cümle içinde farklı görevler üstlenebilmesidir (Bozuyula, 2024).

FastText gibi alt birim (subword) tabanlı modeller, kelime içindeki ekleri ve kökleri daha iyi analiz edebildiğinden Türkçe için daha uygun bir çözüm sunmaktadır (Kilimci & Akyokuş, 2019). Bu model, kelimenin parçalarını ayrı ayrı işleyerek, nadir kelimelerin bile anlamını doğru bir şekilde temsil edebilir. Ancak, alt birim temelli gömme tekniklerinin avantajlarına rağmen, bu modellerin doğruluğu ve etkinliği, kullanılan veri setinin çeşitliliğine bağlıdır (Saritaş ve ark., 2024).

Son yıllarda derin öğrenme teknikleri, Türkçe gibi eklemeli dillerin morfolojik analizinde önemli bir ilerleme kaydetmiştir. Özellikle, CNN ve LSTM gibi modeller, geleneksel kural tabanlı yaklaşımlara göre daha yüksek doğruluk oranları elde etmiştir (Suncak ve Aktaş, 2022).

Türkçe'nin eklemeli yapısı, kelime köklerinin ve eklerin doğru bir şekilde ayrıştırılmasını gerektirir. Bu süreç, tokenizasyon ve kelime segmentasyonu gibi tekniklerle gerçekleştirilir. Ancak, Türkçe gibi dillerde kelime segmentasyonu, yüksek sayıda kelime formu ve eğitim sırasında görmediği ve bu nedenle tanımadığı kelimeler (out-of-vocabulary - OOV) nedeniyle zorlu bir süreçtir (Al Nahas ve ark., 2020; Yiğit Bekir Kaya ve A. Cüneyd Tantuğ, 2024). Çeşitli segmentasyon teknikleri, kelime köklerine eklenen ekleri ayırarak NLP modellerinin performansını artırmayı hedeflemektedir. Byte Pair Encoding (BPE) ve WordPiece gibi algoritmalar, kelimeleri alt birimlere ayırarak kelime haznesini küçültür ve OOV oranını azaltır. Bu yaklaşımlar, özellikle sınırlı dil kaynaklarına sahip Türkçe gibi diller için önemlidir.

Kelime gömme teknikleri, kelimelerin anlamlarını vektör temelli bir biçimde temsil ederek NLP modellerinin performansını artırır. Ancak, eklemeli dillerde kelime gömme işlemi bazı zorluklar yaratır (Baykara & Güngör, 2022; Uymaz & Metin, 2023b). Word2Vec ve GloVe gibi geleneksel gömme yöntemleri, kelimeyi bir bütün olarak ele alırken, Türkçe gibi eklemeli dillerde aynı köke sahip fakat farklı eklerle türetilmiş kelimeleri doğru bir şekilde temsil edemeyebilir. Bu durum, kelimeler arasındaki semantik ilişkilerin yakalanmasını zorlaştırır (Tulu, 2022).

Türkçe'nin anlamsal zenginliğini ortaya koymak amacıyla FrameNet gibi projeler, Türkçe için bir çerçeve semantiği oluşturmayı hedeflemektedir (Marsan ve ark., 2021). FrameNet, kelime ve kavramlar arasındaki ilişkileri tanımlayarak Türkçe dil işleme çalışmalarına önemli katkılar sağlamaktadır.

Türkçe gibi eklemeli dillerde morfolojik analiz, NLP süreçlerinde önemli zorluklar sunmaktadır. Ancak, derin öğrenme tabanlı yaklaşımlar ve çeşitli tokenizasyon yöntemleri, bu zorlukların üstesinden gelmek için umut verici çözümler sunmaktadır. Özellikle Türkçe gibi eklemeli dillerde, kelime kökleri ve ekler arasındaki ilişkilerin doğru bir şekilde modellenmesi, dil işleme görevlerinin başarısını artırabilir. Türkçe için daha kapsamlı veri kümelerinin oluşturulması ve yerleştirilmiş modellerin geliştirilmesi, NLP çalışmalarının başarısını artıracaktır. Tablo 9'da morfolojik analiz modellerinin NLP görevlerindeki performansı detaylı bir şekilde analiz edilmiştir.

Tablo 9. Morfolojik Analiz Modellerinin NLP Görevlerindeki Performans Analizi

Kaynak	Geliştirilen Model	Model	Doğruluk (%)	F1-Score
(Aytan & Şakar, 2023)	LSTM	Bi-LSTM	95	0,87
(Dönmez & Adalı, 2015)	-	M5	68	0,74

Bi-LSTM (Bidirectional Long Short-Term Memory) modeli, yüksek doğruluk ve güçlü bir F1-Skoru ile en başarılı model olarak öne çıkmıştır. Buna karşın, M5 modeli, diğer yaklaşımlara göre daha düşük doğruluk ve F1-Skoru sergileyerek performans açısından geride kalmıştır. Bu bulgular, derin öğrenme tabanlı yaklaşımların, özellikle karmaşık dil yapısına sahip Türkçe gibi dillerde, geleneksel yöntemlere göre üstün olduğunu göstermektedir.

3.8. Disiplinler Arası Uygulamalar için NLP Modelleri

NLP modelleri, son yıllarda farklı disiplinlerde oldukça çeşitli ve yenilikçi uygulamalarla kullanılmıştır. Bu alandaki çalışmalar, NLP'nin disiplinler arası etkisini ortaya koyarak hem sosyal bilimler hem de mühendislik alanlarında yeni çözümler sunmuştur.

Hukuki metinlerin analizi ve yargı kararlarının tahmin edilmesi, NLP'nin hukuki alanda kullanıldığı önemli örneklerden biridir. Türkiye'de yüksek mahkemelerde yargı kararlarının tahmini üzerine yapılan çalışmalarda, ileri beslemeli sinir ağları (Feedforward Neural Networks- FFNN) ve destek vektör makineleri (SVM) gibi makine öğrenimi yöntemleri kullanılarak büyük başarıya ulaşılmıştır (Mumcuoğlu ve ark., 2021; Aydemir, 2023). Ayrıca, Türk Anayasa Mahkemesi'nin kararlarının NLP ve yapay zekâ teknikleri kullanılarak analiz edildiği çalışmalar, davaların sonuçlarını yüksek doğruluk oranları ile tahmin edebilmiştir (Sert ve ark., 2022). Transformer tabanlı modellerin, özellikle uzun ve karmaşık cümle yapılarında bağlamı anlamadaki başarısı, bu tür metinlerin işlenmesinde önemli avantajlar sunar. Geleneksel makine öğrenimi yöntemleri ile karşılaştırıldığında, transformer tabanlı modellerin doğruluk oranlarının yüksek olduğu ve daha karmaşık görevlerde bile üstün performans sergilediği gözlemlenmiştir (Öztürk, 2023).

Tarih ve edebiyat alanında da NLP'nin güçlü bir etki alanı bulunmaktadır. Özellikle, Osmanlı Türkçesi gibi tarihi dillerde NLP uygulamaları, metinlerin dijital edisyonlarının oluşturulması ve dil analizi gibi çalışmalarla tarihsel araştırmaları desteklemiştir. Örneğin, Evliya Çelebi'nin Seyahatnamesi üzerinde yapılan bir çalışmada, GPT modelleri ve BERT tabanlı duygu analizi yöntemleri kullanılarak metinlerdeki semantik ve tematik örüntüler analiz edilmiştir (Aladağ, 2023). Özellikle Osmanlı Türkçesinin zorlu yapısını analiz etmek ve modern NLP modellerine entegre edebilmek, dil araştırmaları için yeni kapılar açmıştır (Karagöz ve ark., 2024).

Biyomedikal alan, geniş kapsamlı ve terminolojik açıdan yoğun metinlerin hızlı ve doğru bir şekilde işlenmesini gerektirmektedir (Hazal Türkmen ve ark., 2022). BioBERTTurk, biyomedikal metinler için geliştirilen Türkçe dil modellerinden biridir ve tıbbi metinlerin anlamlandırılması için özel olarak optimize edilmiştir (Türkmen ve ark., 2023). Genel dil modelleri, tıbbi terimleri anlamada sınırlı kalırken, BioBERTTurk'un tıbbi metinlere uyarlanması sayesinde daha yüksek doğruluk oranları elde edilmektedir (H. Türkmen ve ark., 2022).

Eğitim alanında, öğrenci destek hizmetlerinde GPT-3 gibi ileri seviye NLP modelleri kullanılarak, öğrenci danışmanlığı ve destek sistemlerinde otomasyon sağlanmıştır (Firat, 2020). Bu tür uygulamalar, öğrenci memnuniyetini ve destek hizmetlerinin verimliliğini artırmak açısından önemli avantajlar sunmaktadır.

Robotik ve otomasyon sistemlerinde de NLP modelleri kullanılmaktadır. Örneğin, Türkçe dilinde verilen komutlarla mobil robotların formasyon kontrolü gerçekleştirilmiş ve bu alandaki ilk çalışmalar arasında yer almıştır (Aram ve ark., 2024). Türkçe'nin eklemeli yapısı nedeniyle dilin çözümlemesi zor olmasına rağmen, bu tür projeler Türkçe NLP 'nin potansiyelini göstermektedir.

Disiplinler arası NLP uygulamaları hem sosyal bilimler hem de mühendislik gibi teknik alanlarda büyük bir potansiyel sunmaktadır. Hukuk, tarih, eğitim ve robotik gibi farklı disiplinlerde NLP modellerinin kullanımı, sadece teorik katkılar sağlamakla kalmamış, aynı zamanda pratik uygulamaların da kapısını aralamıştır. Disiplinler arası NLP uygulamalarında karşılaşılan en büyük zorluk, her iki alanın da farklı terminolojik yapılar ve dil modelleri gerektirmesidir. Transformer tabanlı modeller, farklı bağlamları işleme konusunda büyük bir potansiyele sahip olsa da, her iki alan için ayrı ayrı optimize edilmiş modeller geliştirilmesi gerekmektedir. Gelecekte, disiplinler arası NLP çalışmalarının daha da yaygınlaşması, bu disiplinlerde verimliliği artıracak ve hataları en aza indirecektir. Tablo 10'da NLP görevlerinde disiplinler arası uygulamalar için geliştirilen modellerin analizi yapılmıştır.

Tablo 10. Disiplinler Arası Uygulamalar için Geliştirilen Modellerinin NLP Görevlerindeki Performans Analizi

Kaynak	Model	F1-Score	Doğruluk (%)	Kesinlik	Geri Çağırma	Kappa
(Aydemir, 2023)	DecisionTree Classifier	0,8889	88,9	0,9013	0,9	-
(Aras ve ark., 2022)	PCA	0,85	91,8	0,82	0,88	-
(Mumcuoğlu ve ark., 2021)	LSTM + Dikkat	0,67	91,8	-	-	-
(Nazaretsky ve ark., 2023)	DistilBERTTurk	0,89	89	-	-	0,78
(Sert ve ark., 2022)	MLP	0,987	97,8	0,976	1,0	-

Geleneksel yöntemlerden DecisionTreeClassifier ve boyut indirgeme tekniği (Principal Component Analysis – PCA) ile elde edilen modellerin belirli ölçütlerde yüksek doğruluk ve kesinlik sunduğu görülmüştür. Bununla birlikte, derin öğrenme tabanlı yaklaşımlar arasında LSTM modeli, dikkat mekanizmasıyla birlikte optimize edilerek tatmin edici bir doğruluk oranına ulaşmıştır. Özellikle Türkçe dil işleme özelinde kullanılan DistilBERTTurk modeli, güçlü bir sınıflandırma performansı ve yüksek tutarlılık sergilemiştir. Son olarak, çok katmanlı algılayıcı (MLP) modelinin, neredeyse kusursuz doğruluk ve geri çağırma değerleri ile en yüksek performansı gösterdiği tespit edilmiştir.

3.9. Açık Kaynak Platformlar ve Veri Setlerinin Türkçe NLP'ye Katkısı

Türkçe NLP alanı, diğer popüler dillerle kıyaslandığında, sınırlı kaynaklara sahip olmasından dolayı uzun süre geri planda kalmış bir alan olmuştur. Ancak son yıllarda, açık kaynak platformlar ve veri setlerinin ortaya çıkması, bu alanda önemli bir ivme kazandırmıştır. Örneğin, VNLP adlı araç kiti, Türkçe dilinde kapsamlı bir çözüm sunarak, cümle bölme, metin normalizasyonu, duygu analizi ve NER gibi görevleri gerçekleştirebilmektedir (Turker ve ark., 2024). Benzer şekilde, TurkishDelightNLP, Türkçe için geliştirilmiş kapsamlı bir sinirsel NLP araç kiti olarak, kök bulma, morfolojik ayrıştırma ve bağımlılık çözümlemesi gibi birçok NLP görevini kapsar. Araç kitinin açık kaynak olması, araştırmacıların ve geliştiricilerin Türkçe metin işleme süreçlerini hızlandırmasına olanak tanımaktadır (Alecakir ve ark., 2022).

Açık kaynak platformlar, araştırmacıların veri kümelerine ve dil modellerine erişimlerini kolaylaştırarak, yenilikçi NLP projelerinin önünü açar. TULAP (Turkish Natural Language Processing) (Uskudarlı ve ark., 2023), İTÜ NLP grubu (Eryiğit, 2014)

ve Mukayese (Safaya ve ark., 2022) Türkçe NLP alanında araştırma yapmak isteyenler için kapsamlı bir araç ve veri kaynağı sunmaktadır. TULAP, araştırmacılara Türkçe NLP için hazırlanmış açık kaynaklı modeller, eğitim veri setleri ve örnek projeler sunar. Bu platformlar, veri eksikliklerinin üstesinden gelmek için önemli bir çözümdür ve Türkçe'nin NLP çalışmaları daha geniş kullanımını sağlar.

Veri kümeleri ve benchmark setleri, NLP modellerinin performansını değerlendirmek ve karşılaştırmak için gereklidir. TRSAv1, Türkçe metin işleme görevleri için geliştirilen önemli bir veri kümesidir (Aydoğan & Kocaman, 2023). TRSAv1, metin sınıflandırma, duygu analizi, NER ve makine çevirisi gibi birçok NLP görevine uygun veri içerir (Touheed ve ark., 2024). Bu veri kümesinin temel avantajı hem çeşitli konuları kapsaması hem de model performanslarını objektif bir şekilde ölçmeye olanak tanınmasıdır.

Açık kaynak projeler, NLP araştırmalarına hız kazandırmak ve bilimsel ilerlemeyi teşvik etmek için önemli bir araçtır (Çöltekin, 2014). Açık kaynaklı veri setleri ve modeller, araştırmacıların aynı altyapıyı kullanarak deneyler yapmalarına olanak tanır, böylece çalışmaların tekrarlanabilirliğini artırır. Türkçe NLP'de açık kaynak projeler, veri eksikliği sorununu çözmenin yanı sıra, araştırma maliyetlerini düşürmekte ve daha fazla araştırmacının bu alana katkıda bulunmasını sağlamaktadır (Alecakir ve ark., 2022). Açık kaynak projelerin bir diğer önemli katkısı, yeni nesil modellerin geliştirilmesi için geniş bir zemin sunmasıdır. Ayrıca, açık kaynak projeler, akademik ve endüstriyel paydaşlar arasında iş birliğini teşvik ederek, Türkçe NLP araştırmalarının daha hızlı ve etkili bir şekilde ilerlemesine olanak tanımaktadır (Yıldız ve ark., 2019).

Sonuç olarak, Türkçe NLP'nin gelişimi için daha fazla açık kaynak platformunun kurulması ve yeni veri setlerinin sunulması, dil modellerinin performansını artıracak ve araştırma ekosistemini güçlendirecektir. Açık kaynak projelerin teşviki hem akademi hem de endüstri açısından yeni fırsatlar yaratacaktır.

4. TARTIŞMA

4.1. Literatürdeki Boşluklar ve Eksikler

Türkçe NLP alanında son yıllarda kaydedilen ilerlemelere rağmen, literatürde bazı önemli boşluklar ve eksiklikler devam etmektedir. Bu eksiklikler, hem kullanılan yöntemlerin dilin yapısına tam uyum sağlayamamasından hem de Türkçe'nin düşük kaynaklı bir dil olarak sınırlı veri ve model desteğine sahip olmasından kaynaklanmaktadır.

Türkçe, NLP çalışmalarında düşük kaynaklı bir dil olarak kabul edilmekte ve dil modellerinin gelişimi için yeterli veri seti ve açık kaynak proje bulunmamaktadır. Çok dilli modellerin (mBERT, XLM-R) Türkçe üzerinde kullanılması kısmen çözüm sağlasa da bu modellerin eklemeli yapıya sahip Türkçe'nin bağlamlarını tam olarak kavrayamadığı gözlemlenmiştir. Özellikle yerleştirilmiş modellerin eksikliği, Türkçe NLP araştırmalarının ilerlemesini yavaşlatmaktadır. BERTurk gibi bazı modeller geliştirilmeye başlanmış olsa da dilin yapısına özgü daha fazla modelin geliştirilmesine ihtiyaç vardır.

Türkçe için geliştirilen veri setlerinin ve benchmark küme sayısının azlığı, dil modellerinin performansının değerlendirilebilmesini zorlaştırmaktadır. Örneğin, TRSAv1 benchmark seti, Türkçe NLP’de önemli bir boşluğu doldurmuş olsa da daha fazla sayıda çeşitlendirilmiş veri setine ihtiyaç duyulmaktadır. BioBERTurk gibi yerleştirilmiş modellerin geliştirilmesine rağmen, tıp alanında kullanılan veri setlerinin sınırlı olması model başarısını olumsuz etkilemektedir.

Türkçe'nin eklemeli yapısı, NLP modellerinin performansını doğrudan etkileyen bir faktördür. Mevcut kelime gömme teknikleri (Word2Vec, GloVe) eklemeli diller için yeterli olmamakta, FastText gibi alt birim tabanlı modeller bile sınırlı başarı elde edebilmektedir. Bağlamsal bilgi yakalama yeteneğine sahip transformer tabanlı modellerin (BERT, BERTurk) geliştirilmesiyle bazı ilerlemeler sağlanmıştır. Ancak eklemeli diller için daha kapsamlı ve optimize edilmiş modellerin geliştirilmesine ihtiyaç duyulmaktadır. POS etiketleme ve kelime ayrıştırma süreçlerinde de benzer zorluklar mevcuttur, bu nedenle daha fazla morfolojik analiz çalışmasına ihtiyaç vardır. Literatürdeki boşluklar ve eksiklikler, Türkçe NLP araştırmalarının daha geniş ve etkin bir şekilde ilerlemesi için çözüm bekleyen alanlar sunmaktadır. Özellikle yerleştirilmiş modellerin sayısının artırılması, daha çeşitli ve kapsamlı veri setlerinin geliştirilmesi ve disiplinler arası NLP çalışmalarına öncelik verilmesi gerekmektedir. Eklemeli dil yapısının getirdiği zorlukların aşılması için yeni morfolojik analiz tekniklerine ve POS etiketleme modellerine ihtiyaç duyulmaktadır. Hukuk ve tıp gibi alanlarda NLP’nin daha yaygın kullanımını desteklemek için daha geniş veri kümeleri ve özelleştirilmiş modeller geliştirilmelidir. Sosyal medya içeriklerinin analizi için yeni veri setleri ve yanıltıcı bilgi tespiti için optimize edilmiş modeller, literatürdeki eksikliklerin giderilmesine katkı sağlayacaktır.

4.2. Konunun İlerlemesine Engel Teşkil Eden Unsurlar

Türkçe NLP alanında kaydedilen ilerlemelere rağmen, araştırma ekosisteminin tam potansiyeline ulaşmasını engelleyen bazı önemli unsurlar mevcuttur. Düşük kaynaklı diller için yaygın olan bu sınırlamalar, veri eksikliği, metodolojik sınırlılıklar, yerleştirilmiş modellerin yetersizliği ve disiplinler arası iş birliğinin eksikliği gibi çeşitli başlıklarda toplanabilir.

Türkçe NLP araştırmalarının en büyük engellerinden biri, yeterli büyüklükte ve çeşitli veri setlerinin bulunmamasıdır. İngilizce gibi yüksek kaynaklı diller için geniş veri kümeleri (GLUE Benchmark(Nangia ve Bowman, 2019), SQuAD(Rajpurkar ve ark., 2018), IMDb Movie Reviews (Haque ve ark., 2019), OpenAI GPT-3 Datasets(Ryu ve Nakajima, 2022) vb.) mevcutken, Türkçe ‘deki veri setleri sınırlıdır. Sosyal medya verilerinin yanı sıra farklı alanlarda daha fazla etiketlenmiş veriye ihtiyaç duyulmaktadır.

Türkçe için yerleştirilmiş NLP modellerinin sayısının yetersizliği, araştırmaların ilerlemesini engelleyen bir diğer önemli faktördür. BERTurk ve BioBERTurk gibi bazı başarılı modeller geliştirilmiş olsa da bunların kapsamı sınırlıdır ve farklı görevlerde kullanılabilecek daha fazla yerleştirilmiş modellere ihtiyaç duyulmaktadır.

Araştırma ekosisteminin gelişmesini destekleyen önemli unsurlardan biri olan açık kaynak projelerin sayısı Türkçe NLP için yetersizdir. TULAP gibi bazı girişimler veri

paylaşımını teşvik etmeye çalışsa da araştırmacıların iş birliği yapabileceği ve geniş veri kümelerine erişebileceği daha fazla platforma ihtiyaç vardır. Akademik kurumlar ve özel sektör arasındaki veri paylaşımı ve iş birliği eksikliği, inovasyonu kısıtlayan bir diğer engeldir.

Türkçe'nin eklemeli yapısı ve morfolojik zenginliği, dil modelleri için ciddi zorluklar yaratmaktadır. Kelime köklerine eklenen çok sayıda ek ve kelimelerin bağlamlarına göre değişen anlamları, POS etiketleme ve kelime ayrıştırma görevlerini zorlaştırmaktadır. Mevcut kelime gömme teknikleri, bu karmaşık dil yapısına tam olarak uyum sağlayamamakta, bu da modellerin doğruluk oranlarını olumsuz yönde etkilemektedir.

Sosyal medya platformlarındaki verilerin sürekli değişmesi ve argo dil kullanımı, duygu analizi ve yanıltıcı bilgi tespiti gibi görevlerin başarısını sınırlandırmaktadır. TurkishTweets veri kümesi, sosyal medya içeriklerinin doğruluğunu değerlendirmek için bir başlangıç sunmuş olsa, daha fazla tematik veri setine ihtiyaç duyulmaktadır. Yanıltıcı bilgi tespiti için kullanılan modellerin, sürekli güncellenen veri akışına uyum sağlayabilmesi gerekmektedir.

Bu engellerin aşılması, daha geniş veri kümelerinin oluşturulması, açık kaynak projelerin teşvik edilmesi ve dil yapısına uygun modellerin geliştirilmesi ile mümkün olacaktır. Ayrıca, hukuk ve tıp gibi disiplinlerde daha yaygın NLP uygulamaları, bu alanlardaki süreçleri hızlandırabilir ve araştırma ekosistemine katkı sağlayabilir. Sosyal medya içeriklerinin analizi ve yanıltıcı bilgi tespiti alanlarında yapılacak çalışmalar hem toplumsal fayda sağlayacak hem de NLP'nin kapsamını genişletecektir.

4.3. Gelecekte Yapılabilecek Çalışmalar

Türkçe NLP alanında son yıllarda önemli gelişmeler kaydedilmiş olsa da, literatürdeki eksikliklerin giderilmesi ve daha kapsamlı çözümlerin geliştirilmesi için birçok fırsat ve çalışma alanı mevcuttur. Mevcut modellerin performansını artırmak, veri eksikliklerini gidermek ve disiplinler arası kullanım alanlarını genişletmek, gelecekteki çalışmaların temel hedefleri arasında yer almaktadır.

Çok dilli modellerin Türkçe üzerindeki sınırlı başarısı, dilin eklemeli yapısından kaynaklanan zorlukları tam olarak işleyememesine bağlıdır. Gelecekte, Türkçe için optimize edilmiş daha fazla yerelleştirilmiş modelin geliştirilmesi gerekmektedir. Bu modellerin eğitiminde geniş ve çeşitli veri setleri kullanılarak, farklı alanlara özgü ince ayar süreçlerinin uygulanması önerilmektedir. Ayrıca, GPT benzeri LLMs'lerin Türkçe için yerelleştirilmesi, metin üretimi ve özetleme gibi görevlerde yenilikçi çözümler sunabilir.

Veri eksikliği, Türkçe NLP'deki en büyük sorunlardan biridir. Bu sorunun aşılması için veri artırma yöntemlerinin geliştirilmesi ve yaygınlaştırılması önem arz etmektedir. Çevrim içi veri üretimi, paraphrasing ve yapay zekâ destekli veri üretim teknikleri ile mevcut veri kümeleri genişletilebilir. Özellikle yapay zekâ modelleri kullanılarak, eksik veri setleri otomatik olarak oluşturulabilir ve veri kalitesi artırılabilir. Gelecekte, Türkçe için özel veri kümelerinin çeşitlendirilmesi, sosyal medya verileri, hukuk metinleri, tıp ve biyomedikal içerikler gibi farklı veri kaynaklarını kapsayacak şekilde

genişletilmelidir. Tıp ve hukuk gibi disiplinler için geliştirilen NLP modelleri, hata oranlarını azaltarak karar destek sistemlerini güçlendirebilir.

Sosyal medya içerikleri, bilgi yayılımının hızlı ve yoğun olduğu platformlardır. Bu nedenle, duygu analizi ve yanıltıcı bilgi tespiti için kullanılan modellerin performansının artırılması gerekmektedir. Yanıltıcı bilgi tespiti için geliştirilmiş modellerin başarısı artırılarak, sosyal medya platformlarındaki dezenformasyonun önüne geçilebilir.

Gelecekte, daha fazla açık kaynak platform ve proje geliştirilmesi, araştırmacılar arasında iş birliğini artırarak Türkçe NLP çalışmalarını hızlandırabilir. Açık kaynak projeler ayrıca, endüstri ile akademi arasında köprü kurarak, inovasyonu teşvik edecek ve yeni araştırmaların önünü açacaktır.

Gelecekte, yalnızca metin tabanlı değil, çok modelli (multimodal) NLP çalışmalarının da geliştirilmesi beklenmektedir. Metin, görsel ve ses verilerini birlikte işleyebilen modeller, Türkçe NLP'ye yeni bir boyut kazandırabilir. Bu tür çalışmalar, özellikle otomatik altyazı üretimi, sesli asistanlar ve çok dilli çeviri sistemlerinde yenilikçi çözümler sunabilir.

Gelecekteki Türkçe NLP çalışmaları, veri setlerinin çeşitlendirilmesi, yerelleştirilmiş modellerin geliştirilmesi ve disiplinler arası uygulamaların yaygınlaştırılması gibi hedeflere odaklanmalıdır. Özellikle, morfolojik analiz ve eklemeli dil zorluklarının üstesinden gelmek için yeni tekniklerin geliştirilmesi gerekmektedir. Açık kaynak projelerin teşviği ve yapay zekâ destekli veri üretimi, Türkçe NLP çalışmalarını daha ileriye taşıyacaktır. Sosyal medya analizleri, yanıltıcı bilgi tespiti ve çok modelli çalışmalar, gelecekte Türkçe NLP 'nin yeni uygulama alanları arasında yer alacaktır. Tüm bu çalışmalar, Türkçe'nin dijitalleşmesine ve NLP 'de daha etkin kullanılmasına katkı sağlayacaktır.

5. SONUÇ

Bu çalışma, Türkçe doğal dil işleme (NLP) alanında son yıllarda gerçekleştirilen ilerlemeleri ve karşılaşılan zorlukları kapsamlı bir şekilde incelemiştir. Türkçenin eklemeli yapısı ve zengin morfolojisi, NLP çalışmalarında hem potansiyel fırsatlar hem de metodolojik zorluklar yaratmaktadır. Özellikle veri kıtlığı ve dil modellerinin yerelleştirilmesi konularında yaşanan zorluklar, Türkçe NLP modellerinin performansını sınırlandırmakta ve bu durum daha geniş ve çeşitli veri kümelerinin oluşturulması gerekliliğini ortaya koymaktadır.

Çalışmada, BERTurk ve BioBERTurk gibi Türkçe'ye özgü transformer tabanlı modellerin çeşitli NLP görevlerinde başarılı sonuçlar verdiği gözlemlenmiştir. Ancak, mevcut araştırmaların sonuçları, Türkçe gibi eklemeli dillere yönelik daha fazla optimize edilmiş modellere ve veri artırma stratejilerine ihtiyaç olduğunu göstermektedir. Duygu analizi, metin sınıflandırma, isim tanıma gibi temel NLP görevlerinin yanı sıra hukuki ve biyomedikal alan gibi spesifik disiplinlerde yerelleştirilmiş modellerin geliştirilmesi, Türkçe NLP alanında önemli bir ilerleme sağlayacaktır.

Bu derleme, Türkçe doğal dil işlemenin mevcut durumunu ve literatürdeki boşlukları belirleyerek gelecekteki çalışmalara yol göstermeyi amaçlamaktadır. Türkçe için geliştirilen metodolojilerin, düşük kaynaklı diğer dillere uyarlanabilmesi ve bu alandaki çeşitliliğin artırılması yönündeki çalışmalar hem akademik araştırmalar hem de endüstriyel uygulamalar için büyük önem taşımaktadır. Bu bağlamda, Türkçe'nin özgün yapısına daha duyarlı NLP modellerinin geliştirilmesi gerek akademik gerekse pratik uygulamalarda verimliliği ve etkinliği artıracak potansiyele sahiptir.

Araştırma ve Yayın Etiği Beyanı

Yapılan çalışmada araştırma ve yayın etiğine uyulmuştur.

KAYNAKÇA

- Acıkalın, U. U., Bardak, B., & Kutlu, M. (2020). Turkish sentiment analysis using BERT. *2020 28th Signal Processing and Communications Applications Conference (Siu)*. <https://doi.org/10.1109/siu49456.2020.9302492>
- Acıkgöz, E. C., Erdogan, M., & Yuret, D. (2024). Bridging the Bosphorus: Advancing turkish large language models through strategies for low-resource language adaptation and benchmarking. *arXiv preprint arXiv:2405.04685*.
- Adalı, E., & Adamov, A. Z. (2016). Sentiment analysis for agglutinative languages. 2016 IEEE 10th International Conference on Application of Information and Communication Technologies (AICT), Baku, Azerbaijan, 2016, pp. 1-3, doi: 10.1109/ICAICT.2016.7991659.
- Ahmetoğlu, H., & Daş, R. (2020). Türkçe otel yorumlarıyla eğitilen kelime vektörü modellerinin duygu analizi ile incelenmesi. *Süleyman Demirel Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, 24(2), 455-463.
- Akça, O. (2023). Natural language processings in legal domain: Classification of turkish legal texts. [Yüksek Lisans Tezi] Marmara Üniversitesi.
- Akça, O., Bayrak, G., Issifu, A. M., & Ganiz, M. C. (2022). Traditional machine learning and deep learning-based text classification for turkish law documents using transformers and domain adaptation, " 2022 International Conference on INnovations in Intelligent SysTems and Applications (INISTA), Biarritz, France, 2022, pp. 1-6, doi: 10.1109/INISTA55318.2022.9894051.
- Aksu, M. Ç., & Karaman, E. (2020). FastText ve kelime çantası kelime temsil yöntemlerinin turistik mekanlar için yapılan türkçe incelemeler kullanılarak karşılaştırılması. *Avrupa Bilim ve Teknoloji Dergisi*(20), 311- 320.
- Al Nahas, A., Kulunk, A., Gozutok, B., Kalkan, S. C., & Erdinc, H. Y. (2020). how to segment turkish words for neural text classification. 2020 International Conference on INnovations in Intelligent SysTems and Applications (INISTA),

- Aladağ, F. (2023). Osmanlı çalışmalarında GPT'nin potansiyeli: Evliya Çelebi Seyahatnamesinin NLP ve metin madenciliği ile uygulamalı analizi ve TEI yöntemiyle dijital edisyonu. I. Evliya Çelebi Sempozyumu. İstanbul.
- Alecakir, H., Bölücü, N., & Can, B. (2022). TurkishDelightNLP: A neural Turkish NLP toolkit, Proceedings of the 2022 Conference of the North American
- Altinok, D. (2023). A diverse set of freely available linguistic resources for Turkish. Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Acl 2023): Long Papers, Vol 1, 13739-13750.
- Alzoubi, Y. I., Topcu, A. E., & Erkaya, A. E. (2023). Machine learning-based text classification comparison: Turkish language context. *Applied Sciences*, 13(16), 9428.
- Aram, K., Erdemir, G., & Can, B. (2024). Formation control of multiple autonomous mobile robots using Turkish natural language processing. *Applied Sciences*, 14(9), 3722.
- Aras, A. C., Öztürk, C. E., & Koç, A. (2022). Feedforward neural network based case prediction in Turkish higher courts.
- Aras, G., Makaroğlu, D., Demir, S., & Cakir, A. (2021). An evaluation of recent neural sequence tagging models in Turkish named entity recognition. *Expert Systems with Applications*, 182, 115049.
- Arslan, T. P., & Eryiğit, G. (2023). Incorporating dropped pronouns into coreference resolution: the case for Turkish.
- Avşaroğlu, M., & Karadağ, A. B. (2019). "Foreign language creation" and "textless back translation": A case study on Turkish translations of jason goodwin's ottoman-themed works written in English. *Advances in Language and Literary Studies*, 10(5), 107-119.
- Aydemir, E. (2023, 2023). Estimation of Turkish constitutional court decisions in terms of admissibility with NLP.
- Aydoğan, M., & Karci, A. (2020). Improving the accuracy using pre-trained word embeddings on deep neural networks for Turkish text classification. *Physica A: Statistical Mechanics and its Applications*, 541, 123288.
- Aydoğan, M., & Kocaman, V. (2023). TRSAv1: a new benchmark dataset for classifying user reviews on Turkish e-commerce websites. *Journal of Information Science*, 49(6), 1711-1725.
- Aytan, B., & Sakar, C. O. (2022, 2022). Comparison of transformer-based models trained in Turkish and different languages on Turkish natural language processing problems.

- Aytan, B., & Şakar, C. O. (2023). Deep learning-based Turkish spelling error detection with a multi-class false positive reduction model. *Turkish Journal of Electrical Engineering and Computer Sciences*, 31(3), 581-595.
- Ayverdi, S., Öncevarlık, A., Uçar, M., & Adalı, E. (2020, 2020). Time and object based question and answering system for Turkish.
- Ba Alawi, A., & Bozkurt, F. (2024). Performance analysis of embedding methods for deep learning-based Turkish sentiment analysis Models. *Arabian Journal for Science and Engineering*, 1-23.
- Bağcı, A., & Amasyali, M. F. (2021, 2021). Comparison of Turkish paraphrase generation models.
- Balcıoğlu, Y. S. (2024). Detecting Turkish cyberbullying tweets using machine learning. *Düzce Üniversitesi Bilim ve Teknoloji Dergisi*, 12(3), 1410-1428.
- Balli, C., Guzel, M. S., Bostanci, E., & Mishra, A. (2022). Sentimental analysis of Twitter users from Turkish content with natural language processing. *Computational Intelligence and Neuroscience*, 2022(1), 2455160.
- Barbieri, F., Anke, L. E., & Camacho-Collados, J. (2021). XLM-T: Multilingual language models in Twitter for sentiment analysis and beyond. arXiv preprint arXiv:2104.12250.
- Baykara, B., & Güngör, T. (2022). Abstractive text summarization and new large-scale datasets for agglutinative languages Turkish and Hungarian. *Language Resources and Evaluation*, 56(3), 973-1007.
- Boltayevich, E. B., Adalı, E., Mirdjonovna, K. S., Xolmo'Minovna, A. O., Yuldashevna X. Z., & Uktamboy O'g'li, X. N. (2023). The problem of pos tagging and stemming for agglutinative languages (Turkish, Uyghur, Uzbek Languages). 2023 8th International Conference on Computer Science and Engineering (UBMK)
- Bozuyla, M. (2024). Sentiment analysis of Turkish drug reviews with bidirectional encoder representations from transformers. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 23(1), 1-17.
- Bozuyla, M., & Özçift, A. (2022). Developing a fake news identification model with advanced deep languagetransformers for Turkish COVID-19 misinformation data. *Turkish Journal of Electrical Engineering and Computer Sciences*, 30(3), 908-926.
- Bölücü, N., & Can, B. (2019). Unsupervised joint PoS tagging and stemming for agglutinative languages. *ACM Transactions on Asian and Low-Resource Language Information Processing (TALLIP)*, 18(3), 1-21.
- Budur, E., Özçelik, R., Güngör, T., & Potts, C. (2020). Data and representation for Turkish natural language inference. arXiv preprint arXiv:2004.14963.

- Carik, B., & Yeniterzi, R. (2021, 2021). SU-NLP at CheckThat! 2021: Check-Worthiness of Turkish Tweets.
- Cavusoglu, I., Pielka, M., & Sifa, R. (2020). Adapting established text representations for predicting review sentiment in Turkish. 2020 Ieee 7th International Conference on Data Science and Advanced Analytics (Dsa 2020), 755-756. <https://doi.org/10.1109/Dsa49011.2020.00100>
- Conneau, A. (2019). Unsupervised cross-lingual representation learning at scale. arXiv preprint arXiv:1911.02116.
- Çam, N. B., & Özgür, A. (2023). Evaluation of chatgpt and bert-based models for turkish hate speech detection. 2023 8th International Conference on Computer Science and Engineering (UBMK) 229-233
- Çarık, B., & Yeniterzi, R. (2022). A Twitter Corpus for named entity recognition in Turkish. Proceedings of the Thirteenth Language Resources and Evaluation Conference(LREC),4546-4551
- Çelikten, A., & Bulut, H. (2021). Turkish medical text classification using BERT. 2021 29th Signal Processing and Communications Applications Conference (SIU) ,1-4
- Çetindağ, C., Yazıcıoğlu, B., & Koç, A. (2023). Named-entity recognition in Turkish legal texts. *Natural Language Engineering*, 29(3), 615-642.
- Çöltekin, Ç. (2014). A set of open source tools for Turkish natural language processing. Proceedings of the Thirteenth Language Resources and Evaluation Conference(LREC) 1079-1086.
- Çöltekin, Ç., Dogruöz, A. S., & Çetinoglu, Ö. (2023). Resources for Turkish natural language processing. *Language Resources and Evaluation*, 57(1), 449-488. <https://doi.org/10.1007/s10579-022-09605-4>
- Demir, S., & Topcu, B. (2022). Graph-based Turkish text normalization and its impact on noisy text processing. *Engineering Science and Technology, an International Journal*, 35, 101192.
- Demirci, G. M., Keskin, Ş. R., & Doğan, G. (2019). Sentiment analysis in Turkish with deep learning. 2019 IEEE International Conference on Big Data (Big Data), 2215-2221.
- Doğan, B., Balcioglu, Y. S., & Elçi, M. (2024). Multidimensional sentiment analysis method on social media data: comparison of emotions during and after the COVID-19 pandemic. *Kybernetes*.
- Dönmez, İ., & Adalı, E. (2015). Türkçe tümce çözümlemede vektör yaklaşımı. *Afyon Kocatepe Üniversitesi Fen Ve Mühendislik Bilimleri Dergisi*, 15(3), 1-11.

- Dündar, E. B., Kiliç, O. F., Cekiç, T., Manav, Y., & Deniz, O. (2020, 2020). large scale intent detection in Turkish short sentences with contextual word embeddings. In Proceedings of the 12th International Joint Conference on Knowledge Discovery, Knowledge Engineering and Knowledge Management (IC3K 2020)
- EmreOztürk, C. (2023). Retrieving turkish prior legal cases with deep learning. [Doktora Tezi]. Bilkent Üniversitesi.
- Eryiğit, G. (2014). ITU Turkish NLP web service, Proceedings of the Demonstrations at the 14th Conference of the European Chapter of the Association for Computational Linguistics. 2014. p. 1-4.
- Feng, S. Y., Gangal, V., Wei, J., Chandar, S., Vosoughi, S., Mitamura, T., & Hovy, E. (2021). A survey of data augmentation approaches for NLP. arXiv preprint arXiv:2105.03075.
- Firat, M. (2020). Öğrenci destek servislerinde doğal dil işleme: GPT-3 örneği. International Conference of Strategic Research in Social Science and Education. 2020. p. 532-536.
- Firoozi, T., Bulut, O., & Gierl, M. (2023). Language models in automated essay scoring: Insights for the Turkish language. *International Journal of Assessment Tools in Education*, 10(Special Issue), 149-163.
- Freiling, I. (2019). Detecting misinformation in online social networks: A think-aloud study on user strategies. *Scm Studies in Communication and Media*, 8(4), 471-496. <https://doi.org/10.5771/2192-4007-2019-4-471>
- Gemirter, C. B., & Goularas, D. (2021). A Turkish question answering system based on deep learning neural networks [Derin Öğrenme Sinir Ağlarına Dayalı Türkçe Soru Cevaplama Sistemi]. *Journal of Intelligent Systems: Theory and Applications*, 4(2), 65-75. <https://doi.org/10.38016/jista.815823>
- Girgin, A. B. A., Gümüşçekiççi, G., & Birdemir, N. C. (2024). Turkish sentiment analysis: A comprehensive review. *Sigma Journal of Engineering and Natural Sciences-Sigma Muhendislik ve Fen Bilimleri Dergisi*, 42(4), 1292-1314. <https://doi.org/10.14744/sigma.2024.00033>
- Girgin, A. B. A., & Şahin, S. (2023). Improving the performance of sentiment analysis by ensemble hybrid learning algorithm with nlp and cascaded feature extraction. *International Journal of Advances in Engineering and Pure Sciences*, 35(1), 125-141.
- Güler, G., & Tantuğ, A. C. (2020). Comparison of Turkish word representations trained on different morphological forms. arXiv preprint arXiv:2002.05417.

- Haque, M. R., Lima, S. A., & Mishu, S. Z. (2019). Performance analysis of different neural networks for sentiment analysis on IMDb movie reviews. 3rd International conference on electrical, computer & telecommunication engineering (ICECTE). IEEE, 2019. p. 161-164.
- Joulin, A., Grave, E., Bojanowski, P., & Mikolov, T. (2016). Bag of tricks for efficient text classification. arXiv preprint arXiv:1607.01759.
- Karagöz, F., Doğan, B., & Özateş, Ş. (2024). Towards a clean text corpus for Ottoman Turkish. Proceedings of the First Workshop on Natural Language Processing for Turkic Languages (SIGTURK 2024). 2024. p. 62-70.
- Karaoğlu, B., Yorgancıoğlu, H. E., Kışla, T., & Metin, S. K. (2019, 2019). The Impact of sentence embeddings in Turkish paraphrase detection. 2019 27th Signal Processing and Communications Applications Conference (SIU) (pp. 1-4).
- Karayığit, H., Akdaglı, A., & Aci, Ç. İ. (2022). Homophobic and hate speech detection using multilingual-bert model on turkish social media. Information Technology and Control, 51(2), 356-375.
- Katar, O., Özkan, D., Yıldırım, Ö., & Acharya, U. R. (2023). Evaluation of GPT-3 AI language model in research paper writing. *Turkish Journal of Science and Technology*, 18(2), 311-318.
- Kaya, Y. B., & Tantug, A. C. (2022,). Finding the optimal vocabulary size for Turkish named entity recognition. ALTNLP. 2022. p. 99-106.
- Kaya, Y. B., & Tantug, A. C. (2024). BERT2D: Two Dimensional positional embeddings for efficient Turkish NLP. IEEE Access.
- Kaya, Y. B., & Tantug, A. C. (2024). Effect of tokenization granularity for Turkish large language models. Intelligent Systems with Applications, 21, 200335.
- Kemaloğlu, N., Küçüksille, E., & Özgünsür, M. (2021). Turkish sentiment analysis on social media. Sakarya University Journal of Science, 25(3), 629-638.
- Kesgin, H. T., Yuce, M. K., & Amasyali, M. F. (2023). Developing and evaluating tiny to medium-sized turkish bert models. arXiv preprint arXiv:2307.14134.
- Kesgin, H. T., Yuce, M. K., Dogan, E., Uzun, M. E., Uz, A., Seyrek, H. E., Zeer, A., & Amasyali, M. F. (2024). Introducing cosmosGPT: Monolingual training for Turkish language models. arXiv preprint arXiv:2404.17336.
- Kilimci, Z. H., & Akyokuş, S. (2019, 2019). The evaluation of word embedding models and deep learning algorithms for Turkish text classification.
- Kirelli, Y., & Arslankaya, S. (2020). Sentiment analysis of shared tweets on global warming on twitter with data mining methods: a case study on Turkish language. Computational Intelligence and Neuroscience, 2020(1), 1904172.

- Koksal, A. T., Bozal, O., Yürekli, E., & Gezici, G. (2020). TurkishTweets: A benchmark dataset for Turkish text correction. In Findings of the Association for Computational Linguistics: EMNLP 2020.
- Kontuk, R., & Turan, M. (2020). NLP kullanılarak haberlerin yaş gruplarına göre sınıflandırılması. *Gazi University Journal of Science Part C: Design and Technology*, 8(2), 372-382.
- Koru, G. K., & Uluyol, Ç. (2024). Detection of Turkish fake news from tweets with bert models. *IEEE Access*.
- Köksal, A., & Özgür, A. (2021). Twitter dataset and evaluation of transformers for Turkish sentiment analysis. 29th Signal Processing and Communications Applications Conference (SIU). IEEE, 2021. p. 1-4.
- Köksal, Ö., & Yılmaz, E. H. (2022). Improving automated Turkish text classification with learning-based algorithms. *Concurrency and Computation: Practice and Experience*, 34(11), e6874.
- Kuruca, Y., Üstüner, M., & Şimşek, I. (2022). Dijital pazarlamada yapay zekâ kullanımı: Sohbet robotu (Chatbot). *Medya ve Kültür*, 2(1), 88-113.
- Kuyumcu, B., Aksakalli, C., & Delil, S. (2019). An automated new approach in fast text classification (fastText) A case study for Turkish text classification without pre-processing. *Proceedings of the 2019 3rd international conference on natural language processing and information retrieval*. 2019. p. 1-4.
- Küçük, D., & Can, F. (2019). A tweet dataset annotated for named entity recognition and stance detection. *arXiv preprint arXiv:1901.04787*.
- Li, G., Wang, Z., Zhao, M., Song, Y., & Lan, L. (2022). Sentiment analysis of political posts on Hong Kong local forums using fine-tuned mBERT. 2022 IEEE International Conference on Big Data (Big Data), 2022 IEEE International Conference on Big Data (Big Data). IEEE, 2022. p. 6763-6765.
- Marsan, B., Kara, N., Özçelik, M., Arıcan, B. N., Cesur, N., Kuzgun, A., Sanıyar, E., Kuyrukçu, O., & Yıldız, O. T. (2021). Building the turkish framenet. South African Centre for Digital Language Resources (SADiLaR) Potchefstroom, South Africa, 118.
- Morwal, S., Jahan, N., & Chopra, D. (2012). Named entity recognition using hidden Markov model (HMM). *International Journal on Natural Language Computing (IJNLC)* Vol, 1.
- Muller, B., Gupta, D., Fauconnier, J.-P., Patwardhan, S., Vandyke, D., & Agarwal, S. (2023). Languages you know influence those you learn: Impact of language characteristics on multi-lingual text-to-text transfer. *Transfer Learning for Natural Language Processing Workshop*. PMLR, 2023. p. 88-102.

- Mumcuoğlu, E., Öztürk, C. E., Ozaktas, H. M., & Koç, A. (2021). Natural language processing in law: Prediction of outcomes in the higher courts of Turkey. *Information Processing & Management*, 58(5), 102684.
- Najafi, A., & Varol, O. (2024). Turkishbertweet: Fast and reliable large language model for social media analysis. *Expert Systems with Applications*, 255, 124737.
- Nangia, N., & Bowman, S. R. (2019). Human vs. muppet: A conservative estimate of human performance on the GLUE benchmark. arXiv preprint arXiv:1905.10425.
- Nasution, A. H., & Onan, A. (2024). ChatGPT label: Comparing the quality of human-generated and LLM-generated annotations in low-resource language NLP tasks. IEEE Access.
- Nazaretsky, T., Yolcu, H. H., Ariely, M., & Alexandron, G. (2023). Towards automated assessment of scientific explanations in Turkish using language transfer. Proceedings of the 16th International Conference on Educational Data Mining. 2023. p. 453-457.
- Nezhad, S. B., & Agrawal, A. (2023). mBBC: Exploring the multilingual maze. arXiv preprint arXiv:2310.05404.
- Okur, H. I., & Sertbaş, A. (2021). Pretrained neural models for turkish text classification. 2021 6th International Conference on Computer Science and Engineering (UBMK). IEEE, 2021. p. 174-179.
- Onan, A., & Balbal, K. F. (2024). Improving Turkish text sentiment classification through task-specific and universal transformations: an ensemble data augmentation approach. IEEE Access.
- Ozcelik, O., & Toraman, C. (2022). Named entity recognition in Turkish: A comparative study with detailed error analysis. *Information Processing & Management*, 59(6), 103065.
- Ozdemir, A., & Yeniterzi, R. (2020). Su-nlp at semeval-2020 task 12: Offensive language identification in turkish tweets. Proceedings of the Fourteenth Workshop on Semantic Evaluation. 2020. p. 2171-2176.
- Özateş, Ş. B., Tıraş, T. E., Genç, E. E., & Taşdemir, E. F. B. (2024). Dependency annotation of Ottoman Turkish with multilingual BERT. arXiv preprint arXiv:2402.14743.
- Özçift, A., Akarsu, K., Yumuk, F., & Söylemez, C. (2021). Advancing natural language processing (NLP) applications of morphologically rich languages with bidirectional encoder representations from transformers (BERT): an empirical case study for Turkish. *Automatika*, 62(2), 226-238. <https://doi.org/10.1080/00051144.2021.1922150>

- Özkan, M., & Kar, G. (2022). Türkçe dilinde yazılan bilimsel metinlerin derin öğrenme tekniği uygulanarak çoklu sınıflandırılması. *Mühendislik Bilimleri ve Tasarım Dergisi*, 10(2), 504-519.
- Öztürk, C. E., Özçelik, S. B., & Koç, A. (2023). A Transformer-based prior legal case retrieval method. 2023 31st Signal Processing and Communications ApplicationsCon.,Siu <https://doi.org/10.1109/Siu59756.2023.10223938>
- Panchendrarajan, R., & Amaresan, A. (2018, 2018). Bidirectional LSTM-CRF for named entity recognition.
- Rajpurkar, P., Jia, R., & Liang, P. (2018). Know what you don't know: Unanswerable questions for SQuAD. arXiv preprint arXiv:1806.03822.
- Ryu, M., & Nakajima, K. (2022). Analysis and mitigation of dataset artifacts in OpenAI GPT-3. In.
- Safaya, A., Kurtuluş, E., Göktoğan, A., & Yuret, D. (2022). Mukayese: Turkish NLP strikes back. arXiv preprint arXiv:2203.01215.
- Sanh, V. (2019). DistilBERT, A Distilled Version of BERT: Smaller, faster, cheaper and lighter. arXiv preprint arXiv:1910.01108.
- Sarıtaş, K., Öz, C. A., & Güngör, T. (2024). A comprehensive analysis of static word embeddings for Turkish. *Expert Systems with Applications*, 252, 124123.
- Schweter, S. (2020). Berturk-bert models for Turkish, April 2020. URL <https://doi.org/10.5281/zenodo.3770924>.
- Sert, M. F., Yıldırım, E., & Haşlak, İ. (2022). Using artificial intelligence to predict decisions of the Turkish constitutional court. *Social Science Computer Review*, 40(6), 1416-1435.
- Song, B., Li, Z., Lin, X., Wang, J., Wang, T., & Fu, X. (2021). Pretraining model for biological sequence data. *Briefings in functional genomics*, 20(3), 181-195.
- Soygazi, F., Çiftçi, O., Kök, U., & Cengiz, S. (2021). THQuAD: Turkish historic question answering dataset for reading comprehension. 2021 6th international conference on computer science and engineering (UBMK). IEEE, 2021. p. 215-220.
- Srinivasan, A., Sitaram, S., Ganu, T., Dandapat, S., Bali, K., & Choudhury, M. (2021). Predicting the performance of multilingual nlp models. arXiv preprint arXiv:2110.08875.
- Suncak, A., & Aktaş, Ö. (2021). A novel approach for detecting defective expressions in Turkish. *Journal of Artificial Intelligence and Data Science*, 1(1), 35-40.

- Suncak, A., & Aktaş, Ö. (2022). Detecting Defective Expressions in Turkish Sentences Using a Hybrid Deep Learning Method. *Dokuz Eylül Üniversitesi Mühendislik Fakültesi Fen ve Mühendislik Dergisi*, 24(72), 825-834.
- Şahin, G. G. (2024). Bridging the Gap Between Wikipedians and Scientists with Terminology-Aware Translation: A Case Study in Turkish. *Wikimedia Research Fund* 2024
- Şapcı, A. O. B., Taştan, Ö., & Yeniterzi, R. (2020). Active Learning for Turkish Text Classification. 28th Signal Processing and Communications Applications Conference (SIU). IEEE, 2020. p. 1-4.
- Tohma, K., & Kutlu, Y. (2020). Challenges Encountered in Turkish Natural Language Processing Studies. *Natural and Engineering Sciences*, 5(3), 204-211.
- Tohma, K., Okur, H. I., Kutlu, Y., & Sertbas, A. (2023). Sentiment Analysis in Turkish Question Answering Systems: An Application of Human-Robot Interaction. *IEEE Access*.
- Tokcaer, S. (2021). Türkçe metinlerde duygu analizi. *Yaşar Üniversitesi E-Dergisi*, 16(63), 1514-1534.
- Toraman, C., Ozcelik, O., Sahinuç, F., & Sahin, U. (2022, 2022). ARC-NLP at CheckThat!-2022: Contradiction for Harmful Tweet Detection. *CLEF (Working Notes)*. 2022. p. 722-739.
- Touheed, M., Zubair, U., Sabir, D., Hassan, A., Butt, M. F. U., Riaz, F., Abdul, W., & Ayub, R. (2024). Applications of Pruning Methods in Natural Language Processing. *IEEE Access*.
- Tulu, C. N. (2022). Experimental comparison of pre-trained word embedding vectors of Word2Vec, glove, FastText for word level semantic text similarity measurement in turkish. *Advances in Science and Technology. Research Journal*, 16(4).
- Tunali, V. (2022). Improved prioritization of software development demands in Turkish with deep learning-based NLP. *IEEE Access*, 10, 40249-40263.
- Turker, M., Ari, M. E., & Han, A. (2024). VNLP: Turkish NLP Package. *arXiv preprint arXiv:2403.01309*.
- Türk, U., Atmaca, F., Özates, S. B., Berk, G., Bedir, S. T., Köksal, A., Basaran, B. Ö., Güngör, T., & Özgür, A. (2022). Resources for Turkish dependency parsing: introducing the BOUN Treebank and the BoAT annotation tool. *Language Resources and Evaluation*, 56(1), 259-307. <https://doi.org/10.1007/s10579-021-09558-0>

- Türkmen, H., Dikenelli, O., Eraslan, C., Çallı, M. C., & Özbek, S. S. (2023). BioBERTurk: Exploring Turkish Biomedical Language Model Development Strategies in Low-Resource Setting. *Journal of Healthcare Informatics Research*, 7(4), 433-446.
- Türkmen, H., Dikenelli, O., Eraslan, C., Çallı, M. C., & Ozbek, S. S. (2022). Developing Pretrained Language Models for Turkish Biomedical Domain. 2022 IEEE 10th International Conference on Healthcare Informatics (ICHI), IEEE, 2022. 597-598.
- Türkmen, H., Dikenelli, O., Eraslan, C., Çallı, M. C., & Özbek, S. S. (2023). Harnessing the power of BERT in the Turkish clinical domain: pretraining approaches for limited data scenarios. *arXiv preprint arXiv:2305.03788*.
- Uludoğan, G., Balal, Z. Y., Akkurt, F., Türker, M., Güngör, O., & Üsküdarlı, S. (2024). Turna: A turkish encoder-decoder language model for enhanced understanding and generation. *arXiv preprint arXiv:2401.14373*.
- Uskudarli, S., Şen, M., Akkurt, F., Gürbüz, M., Güngör, O., Özgür, A., & Güngör, T. (2023). TULAP-An Accessible and Sustainable Platform for Turkish Natural Language Processing Resources. *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics: System Demonstrations*. 2023. p. 219-227.
- Uymaz, H. A., & Metin, S. K. (2023a). Emotion-enriched word embeddings for Turkish. *Expert Systems with Applications*, 225, 120011.
- Uymaz, H. A., & Metin, S. K. (2023b). Enriching Transformer-Based Embeddings for Emotion Identification in an Agglutinative Language: Turkish. *It Professional*, 25(4), 67-73.
- Wikipedia. (2024). Türk dilleri. Wikipedia. Retrieved 18.10.2024 from https://tr.wikipedia.org/wiki/T%C3%BCrk_dilleri
- Xu, Q. A., Chang, V., & Jayne, C. (2022). A systematic review of social media-based sentiment analysis: Emerging trends and challenges. *Decision Analytics Journal*, 3, 100073.
- Xue, L. (2020). mt5: A massively multilingual pre-trained text-to-text transformer. *arXiv preprint arXiv:2010.11934*.
- Yao, S., Peng, B., Papadimitriou, C., & Narasimhan, K. (2021). Self-attention networks can process bounded hierarchical languages. *arXiv preprint arXiv:2105.11115*.
- Yazar, B. K., Şahin, D. Ö., & Kiliç, E. (2023). Low-resource neural machine translation: A systematic literature review. *IEEE Access*, 11, 131775-131813.

- Yıldırım, A., Cetiner, M., Öksüz, C., & Onay, B. (2021). A search tool in Turkish using contextual vectors. 2021 29th Signal Processing and Communications Applications Conference (SIU). IEEE, 2021. p. 1-4.
- Yıldız, O. T., Avar, B., & Ercan, G. (2019). An open, extendible, and fast Turkish morphological analyzer. Proceedings of Recent Advances in Natural Language Processing, pages 1364–1372
- Yilmaz, E. H., & Toraman, C. (2021). Intent classification based on deep learning language model in turkish dialog systems. 2021 29th Signal Processing and Communications Applications Conference (SIU) p 1-4
- Yucalar, F. (2023). Developing an advanced software requirements classification model using bert: An empirical evaluation study on newly generated turkish data. *Applied Sciences*, 13(20), 11127.
- Zhuang, F., Qi, Z., Duan, K., Xi, D., Zhu, Y., Zhu, H., Xiong, H., & He, Q. (2020). A comprehensive survey on transfer learning. Proceedings of the IEEE, 109(1), 43-76.
- Zorarpaci, E. (2023). A Turkish Text Classification Based Feature Selection and Density Peaks Clustering. In 2023 31st Signal Processing and Communications Applications Conference (SIU) (pp. 1-4). IEEE.
- Zovikoğlu, M., & Çetin, U. (2024). Detecting misinformation on social networks with NLP. *Transactions on Computer Science and Applications*, 1(1), 11-16.