



T.C.
Burdur Mehmet Akif Ersoy Üniversitesi
Eğitim Bilimleri Enstitüsü
Bilgisayar ve Öğretim Teknolojileri Eğitimi Anabilim Dalı
Bilgisayar ve Öğretim Teknolojileri Eğitimi Tezli Yüksek Lisans
Programı

**AKADEMİK ALANDA YAPAY ZEKÂ ARAÇLARI TARAFINDAN
ÜRETİLEN İÇERİKLERİN TESPİT EDİLMESİNE YÖNELİK BİR
UYGULAMA**

Sinem KOÇ
Yüksek Lisans Tezi

Tez Danışmanı
Doç. Dr. Onur SEVLİ

Burdur, 2025

T.C.
Burdur Mehmet Akif Ersoy Üniversitesi
Eğitim Bilimleri Enstitüsü
Bilgisayar ve Öğretim Teknolojileri Eğitimi Anabilim Dalı
Bilgisayar ve Öğretim Teknolojileri Eğitimi Tezli Yüksek Lisans
Programı

AKADEMİK ALANDA YAPAY ZEKÂ ARAÇLARI TARAFINDAN
ÜRETİLEN İÇERİKLERİN TESPİT EDİLMESİNE YÖNELİK BİR
UYGULAMA

Sinem KOÇ
Yüksek Lisans Tezi

Tez Danışmanı
Doç. Dr. Onur SEVLİ

Burdur, 2025



**MAKÜ EĞİTİM
BİLİMLERİ ENSTİTÜSÜ**

YÜKSEK LİSANS JÜRİ ONAY FORMU

M.A.K.Ü Eğitim Bilimleri Enstitüsü Yönetim Kurulu'nun 23.06.2025 tarih ve E.535212 sayılı kararıyla oluşturulan jüri tarafından 08.07.2025 tarihinde tez savunma sınavı yapılan Sinem KOÇ'un Akademik Alanda Yapay Zekâ Araçları Tarafından Üretilen İçeriklerin Tespit Edilmesine Yönelik Bir Uygulama konulu tez çalışması Bilgisayar ve Öğretim Teknolojileri Eğitimi Anabilim Dalında YÜKSEK LİSANS tezi olarak kabul edilmiştir.

JÜRİ

ÜYE (Tez Danışmanı)

: Doç. Dr. Onur SELVİ

ÜYE

: Doç. Dr. Ahmet Ali SÜZEN

ÜYE

: Dr. Öğr. Üyesi Vesile Gül BAŞER

ONAY

M.A.K.Ü Eğitim Bilimleri Enstitüsü Yönetim Kurulu'nun
...../...../..... tarih ve/..... sayılı kararı.

İMZA/MÜHÜR

BİLDİRİM

Tez yazma sürecinde bilimsel ve etik ilkelere uyduğumu, yararlandığım tüm kaynakları kaynak gösterme ilkelerine uygun olarak kaynakçada belirttiğimi ve bu bölümler dışındaki tüm ifadelerin şahsıma ait olduğunu taahhüt edip, tezimin kaynak göstermek koşuluyla aşağıda belirttiğim şekilde fotokopi ile çoğaltılmasına izin veriyorum.

[] Tezimin/Raporumun tamamı her yerden erişime açılabilir.

[] Tezim/Raporum sadece Mehmet Akif Ersoy Üniversitesi yerleşkelerinden erişime açılabilir.

[] Tezimin/Raporumun yıl süreyle erişime açılmasını istemiyorum. Bu sürenin sonunda uzatma için başvuruda bulunmadığım takdirde, tezimin/raporumun tamamı her yerden erişime açılabilir.

Sinem KOÇ

Tarih

İmza

TEŞEKKÜR

Bu araştırma sürecinde beni yönlendiren, karşılaştığım zorluklara bilgi ve tecrübesi ile yardımcı olan, desteğini esirgemeyen; sürecin her aşamasında akademik disiplini, yönlendirmeleri ve yapıcı geri bildirimleri ile araştırmamın daha nitelikli bir yapıya ulaşmasında önemli rol oynayan ve akademik rehberliğiyle yolumu aydınlatan değerli danışman hocam **Doç. Dr. Onur SEVLİ**'ye teşekkürlerimi sunarım.

Eğitim hayatımın her aşamasında koşulsuz sevgileri, sabırları, anlayışları ve maddi-manevi destekleriyle daima yanımda olan kıymetli annem **Dilek KOÇ** ve babam **Ramazan KOÇ**, bana duydukları inancı her fırsatta hissettirdikleri ve motivasyon kaynağım oldukları için sonsuz teşekkürlerimi, sevgi ve saygılarımı sunarım.

Ayrıca, yüksek lisans eğitimim süresince bilgi ve deneyimlerinden faydalandığım, akademik gelişimime katkı sunan ve her zaman desteklerini hissettiren Burdur Mehmet Akif Ersoy Üniversitesi Eğitim Bilimleri Enstitüsü Bilgisayar ve Öğretim Teknolojileri Eğitimi Anabilim Dalı'nın değerli hocalarımıza teşekkür ederim.

Akademik Alanda Yapay Zekâ Araçları Tarafından Üretilen İçeriklerin Tespit Edilmesine Yönelik Bir Uygulama

(Yüksek Lisans Tezi)

Sinem KOÇ

ÖZ

Yapay zekâ teknolojilerinin özellikle doğal dil işleme alanında gösterdiği hızlı gelişim, metin üretimi ve içerik düzenleme süreçlerinde köklü dönüşümlere yol açmıştır. Üretken dil modellerinin, özellikle akademik yazım süreçlerinde yaygın biçimde kullanılmaya başlanması, insan üretimi metinler ile yapay zekâ destekli içerikler arasındaki ayrımı zorlaştırmakta; bu durum, akademik etik, içerik güvenliği ve yayıncılık ilkeleri açısından çeşitli sorunları gündeme getirmektedir. Bu çalışmanın temel amacı, akademik metinlerde yapay zekâ tarafından oluşturulmuş içerikleri yüksek doğrulukla tespit edebilen sınıflandırma modeli geliştirmektir. Çalışma kapsamında 2022 öncesi Türkçe akademik makalenin özetleri seçilmiştir. Sağlık, mühendislik, eğitim bilimleri ve sosyal bilimler gibi çeşitli alanlarda, ULAKBİM, Scopus, Web of Science gibi dizinlerde taranan dergilerde yayımlanmış 2.000 adet Türkçe akademik makale kullanılarak özgün bir veri seti oluşturulmuştur. Veri setinde orijinal metinlerin yanında, orijinal makale verilerin üzerinden ChatGPT, Smodin ve Gemini gibi üretken yapay zekâ araçları kullanılarak üretilen özetler ve yine orijinal özetlerinin üretken yapay zekâ uygulamaları tarafından yeniden düzenlenmiş hallerinden (paraphrase) oluşan üç sınıflı (Gerçek Veri, Yapay Zekâ, Düzenlenmiş) halinde genişletilmiştir. Böylece toplamda 6.000 adet dengeli bir veri seti oluşturulmuştur. Sınıflandırma işlemlerinden derin sinir ağı temelli ve literatürde yaygın kullanılan XLM-Roberta, TurkishBERT ve Bert Base Multilingual modelleri kullanılmıştır. Model performanslarının değerlendirilmesinde doğruluk, kesinlik, duyarlılık ve F1 skoru gibi ölçütler esas alınmıştır. Söz konusu metrikler doğrultusunda analiz edilerek modellerin sınıflandırma başarımı karşılaştırılmıştır. Elde edilen sonuçlar değerlendirildiğinde, XLM-Roberta modelinin %97,22 Bert Base Multilingual modeli %96,40 ve TurkishBERT modeli ise %90,22 en yüksek doğruluk değerine ulaşmıştır. Elde edilen bulgular, geliştirilen sınıflandırıcının akademik metinlerdeki yapay zekâ içeriklerini tespit etmede etkili olduğunu göstermektedir. Kullanılan modellerin sağladığı yüksek performans hem akademik içeriklerin orijinalliğini kontrol etmek hem de özgün bilimsel üretkenliği artırmak açısından önemli bir adım olarak değerlendirilebilir. Bu çalışma, akademik etiğe katkı sağlamakta ve bilimsel çalışmalarda şeffaflık ve güvenilirliği artırmaya yardımcı olmaktadır.

Anahtar Kelimeler: Yapay Zekâ Üretimi İçerik Tespiti, Doğal Dil İşleme, Metin Analizi

Sayfa Sayısı : 72

Danışman : Doç. Dr. Onur SEVLİ

An Application for Detecting Content Produced by Artificial Intelligence Tools in Academic Field

(Master Thesis)

Sinem KOÇ

ABSTRACT

The rapid development of artificial intelligence technologies, particularly in the field of natural language processing, has led to fundamental changes in text production and content editing processes. The widespread use of generative language models, especially in academic writing processes, has made it difficult to distinguish between human-generated texts and AI-supported content, raising various issues in terms of academic ethics, content security, and publishing principles. The main objective of this study is to develop a classification model that can accurately detect AI-generated content in academic texts. For this study, abstracts of Turkish academic articles published before 2022 were selected. A unique dataset was created using 2,000 Turkish academic articles published in journals indexed in databases such as ULAKBİM, Scopus, and Web of Science, covering various fields such as health, engineering, educational sciences, and social sciences. In addition to the original texts, the dataset was expanded to include three classes (Real Data, Artificial Intelligence, and Edited): summaries generated using generative artificial intelligence tools such as ChatGPT, Smodin, and Gemini, and rephrasings of the original summaries by generative artificial intelligence applications. This resulted in a balanced dataset of 6,000 items. For the classification process, deep neural network-based models widely used in the literature, such as XLM-Roberta, TurkishBERT, and Bert Base Multilingual, were employed. The performance of the models was evaluated based on criteria such as accuracy, precision, sensitivity, and F1 score. The models' classification performance was compared by analyzing them according to these metrics. When the results were evaluated, the XLM-Roberta model achieved the highest accuracy value of 97.22%, the Bert Base Multilingual model achieved 96.40%, and the TurkishBERT model achieved 90.22%. The findings show that the developed classifier is effective in detecting artificial intelligence content in academic texts. The high performance provided by the models used can be considered an important step in both checking the originality of academic content and increasing original scientific productivity. This study contributes to academic ethics and helps to increase transparency and reliability in scientific studies.

Keywords: Artificial Intelligence Generation Content Detection, Natural Language Processing, Text Analytics

Page Number : 72

Supervisor : Assoc. Prof. Dr. Onur SEVLİ

İÇİNDEKİLER

BİLDİRİM	i
TEŞEKKÜR.....	ii
ÖZ	iii
ABSTRACT.....	iv
İÇİNDEKİLER	v
TABLolar LİSTESİ	viii
ŞEKİLLER LİSTESİ	ix
SİMGELER VE KISALTMALAR LİSTESİ.....	x
BÖLÜM I.....	1
GİRİŞ.....	1
1.1. Problem Durumu.....	2
1.2. Problem Cümlesi.....	3
1.3. Araştırmanın Amacı.....	4
1.4. Araştırmanın Önemi	5
1.5. Sınırlılıklar	6
1.6. Tanımlar	6
BÖLÜM II.....	7
KURAMSAL ÇERÇEVE VE İLGİLİ ARAŞTIRMALAR.....	7
2.1. Kuramsal Çerçeve.....	7
2.1.1. Doğal dil işleme (Natural Language Processing- NLP).....	7
2.1.2. Doğal dil işlemenin tarihçesi ve gelişimi.	7
2.1.3. Doğal dil yapısı.	9
2.1.3.1. Ses bilgisi.	9
2.1.3.2. Biçim bilgisi.	9
2.1.3.3. Söz dizim bilgisi.....	10
2.1.3.4. Anlam bilgisi.....	10
2.1.3.5. Kullanım bilgisi.....	11
2.1.4. Doğal dil üretimi (Natural Language Generation- NLG).....	11
2.1.5. Dil modellemesi.	11
2.1.6. Duygu analizi.	12

2.1.7. Makine çevirisi.	12
2.1.8. Konuşma tanıma ve sentezi.	13
2.1.9. Doğal dil işleme uygulama alanları.	13
2.1.9.1. Yazım yanlışlarının düzeltilmesi.	13
2.1.9.2. Büyük metin dosyalarından özet çıkarma.	14
2.1.9.3. Ana bilgi tespiti.	15
2.1.9.4. Talep anlamlandırma.	15
2.1.9.5. Arama motorları.	16
2.1.9.6. Konuşma metne dönüştürme ve metin seslendirme.	16
2.1.9.7. Doğal dil çevirisi.	16
2.1.9.8. Soru yanıtlama sistemleri.	17
2.1.9.9. Sohbet robotları (Chatbots).	17
2.1.9.10. Metin sınıflandırma (Text Classification).	17
2.1.9.11. Metin özetleme (Text Summarization).	18
2.1.10. BERT (Bidirectional Encoder Representations from Transformers). ...	19
2.1.10.1. XLM-RoBERTa.	20
2.1.10.2. TurkishBERT.	21
2.1.10.3. Bert base multilingual.	21
2.1.11. Üretken yapay zekâ.	22
2.1.12. Üretken yapay zekâ araçları.	23
2.1.13. ChatGPT.	23
2.1.14. Smodin.	23
2.1.15. Gemini.	24
2.2. Yurtiçi ve Yurtdışında Yapılmış İlgili Çalışmalar	24
BÖLÜM III	42
MATERYAL VE YÖNTEM	42
3.1. Araştırmanın Modeli.....	42
3.2. Veri Seti.....	43
3.3. Veri Çekme	44
3.4. Yazılım Geliştirme Ortamı.....	45
3.4.1. Veri ön işleme.....	46
3.4.2. XLM-RoBERTa modeli.	47
3.4.3. Bert Base Multilingual modeli.	48

3.4.4. TurkishBERT modeli.....	49
BÖLÜM IV	51
ARAŞTIRMA BULGULARI VE YORUMLAR	51
4.1. Modelin Sınıflandırma İşleyişi ve Tahmin Süreci	56
4.2. Model Bazlı Sınıflandırma Sonuçları	57
4.2.1. XLM-Roberta sınıflandırma sonuçları.	57
4.2.2. Bert Base Multilingual sınıflandırma sonuçları.	58
4.2.3. TurkishBERT sınıflandırma sonuçları.....	59
4.2.4. Bert Base Multilingual, XLM-Roberta ve Turkishbert modellerine ait sınıflandırma sonuçlarının karşılaştırılması.	60
BÖLÜM V	62
TARTIŞMA, SONUÇ VE ÖNERİLER	62
5.1. Tartışma.....	62
5.2. Sonuç	65
5.3. Öneriler	65
KAYNAKLAR.....	67

TABLÖLAR LİSTESİ

<u>Tablolar</u>	<u>Sayfa</u>
Tablo 1. Literatürdeki Çalışmaların Karşılaştırması	38
Tablo 2. Veri Kümesinin Etiketlenmiş Örnek Görüntüsü	46
Tablo 3. Modelde Kullanılan Hiper parametreler	51
Tablo 4. Modelin Sınıflandırma İşleyişi ve Tahmin Sonuçları.....	56
Tablo 5. XLM-Roberta Sonucu	58
Tablo 6. Bert Base Multilingual Sonucu	59
Tablo 7. TurkishBERT Sonucu.....	60
Tablo 8. Bert Base Multilingual, XLM-Roberta ve Turkishbert Sonuç Karşılaştırma Tablosu.	61
Tablo 9. Bulguların Literatürdeki Çalışmalar ile Karşılaştırılması	63

ŞEKİLLER LİSTESİ

<u>Şekiller</u>	<u>Sayfa</u>
Şekil 1. BERT model yapısı (Devlin vd., 2019).....	20
Şekil 2. Uygulamanın akış diyagramı	43
Şekil 3. Google Colab çalışma ortamı.....	45
Şekil 4. Sınıf dağılım grafiği.	47
Şekil 5. XLM-RoBERTa modeli	48
Şekil 6. Bert Base Multilingual modeli.	49
Şekil 7. TurkishBERT modeli.....	50
Şekil 8. XLM-RoBERTa modeli grafikleri.	52
Şekil 9. Bert Base Multilingual modeli grafikleri.	54
Şekil 10. TurkishBERT modeli grafikleri.....	55

SİMGELER VE KISALTMALAR LİSTESİ

BERT	: İki Yönlü Dil İşleme Modeli
Bert Base Multilingual	: Çok Dilli Temel BERT Modeli
CNN	: Görüntü İşleme İçin Sinir Ağı
CPU	: Merkezi İşlem Birimi
GPT	: Üretken Ön-Eğitilmiş Dönüştürücü
HMM	: Gizli Markov Modeli
LSTM	: Uzun-Kısa Süreli Bellek Ağı
NER	: Adlandırılmış Varlık Tanıma
NLG	: Doğal Dil Üretimi
NLP	: Doğal Dil İşleme

BÖLÜM I

GİRİŞ

Yapay zekâ teknolojilerinin özellikle Doğal Dil İşleme (Natural Language Processing -NLP) alanında kaydettiği hızlı gelişmeler, metin üretimi, içerik düzenleme ve bilgi yönetimi süreçlerinde köklü değişimlere yol açmaktadır. Üretken dil modelleri, insan benzeri metinler üretme becerisiyle akademik alanda da giderek yaygınlaşmakta, bu durum ise insan üretimi metinler ile yapay zekâ tarafından oluşturulan içerikler arasındaki sınırların giderek belirsizleşmesine neden olmaktadır. Akademik yazım süreçlerinde üretken yapay zekâ araçlarının kullanımının artması, içeriklerin özgünlüğü ve etik değerlere uygunluğu gibi temel kavramları tartışmalı hâle getirmiştir. ChatGPT, Gemini ve Smodin gibi büyük dil modeli uygulamalarının erişilebilirliği, bilimsel metinlerin hazırlanmasında kullanıcılar açısından önemli kolaylıklar sağlarken; aynı zamanda bu içeriklerin güvenilirliğini ve kaynağını tespit etmeyi güçleştirmektedir.

Yapay zekâ destekli metinlerin insanlar tarafından üretilen içeriklerden ayırt edilmesi, geleneksel denetim yöntemleriyle mümkün olamamakta; özellikle intihal tespit araçları, özgün ama yapay yollarla oluşturulmuş metinleri saptamakta yetersiz kalmaktadır. Bu durum, akademik etik ilkelerin ihlali, bilimsel güvenilirliğin zedelenmesi ve yayıncılık standartlarının tehdit altına girmesi gibi sorunları beraberinde getirmektedir. Üretken dil modelleri, yalnızca dilbilgisel açıdan değil, aynı zamanda içerik olarak da bağlamsal tutarlılığı yüksek metinler sunabilmektedir. Dolayısıyla, içeriklerin üretim kaynağının belirlenmesi, yalnızca yüzeysel dil özelliklerinde değil; derin anlamsal yapıya dayalı yöntemlerin geliştirilmesini gerektirmektedir.

Bu doğrultuda, akademik içeriklerin üretim kaynağını belirlemeye yönelik sınıflandırıcı sistemlerin geliştirilmesi önemli bir ihtiyaç hâline gelmiştir. Bu nedenle hem insan üretimi hem de yapay içerikleri barındıran veri setlerine dayalı analizler ve özellikle dilimiz Türkçe için de bir gereklilik haline gelmiştir. Bu doğrultuda Türkçe akademik metinlere odaklanılarak hazırlanan ve sağlık, mühendislik, sosyal bilimler ve eğitim bilimleri gibi çeşitli alanları kapsayan bir veri kümesi oluşturulmuştur. Veri

kümesi içerisinde insan üretimi orijinal içerikler “Gerçek Veri”, orijinal metinlerden yapay zekâ tarafından doğrudan özetlenenler “Yapay Zekâ” ve orijinal özetlerin yapay zekâ destekli düzenlemelerinden oluşan içerikler “Düzenlenmiş” olmak üzere üç sınıf türünde ve her sınıftan eşit sayılarda toplam 6.000 adet örnek yer almaktadır

Veri işleme ve sınıflandırma sürecinde Python programlama dilinin doğal dil işleme ve analiz kütüphanelerinden yararlanılmış; metin temizleme, normalizasyon, tokenizasyon gibi işlemlerin ardından transformer tabanlı modeller kullanılarak sınıflandırma gerçekleştirilmiştir. Literatürde yaygın kullanılan XLM-Roberta, TurkishBERT ve Bert Base Multilingual gibi güncel modeller test edilmiştir. Modellerin performansları doğruluk, kesinlik, duyarlılık, F1 skoru metrikleri açısından karşılaştırılmıştır.

Üretken dil modelleri tarafından sentezlenen içeriklerin tespitine yönelik sistemlerin geliştirilmesi, akademik yazının şeffaflık ve güvenilirlik temelli yapısını korumak açısından kritik öneme sahiptir. Üretken yapay zekâ araçlarının sağladığı kolaylıkların, etik dışı kullanım riskini beraberinde getirdiği göz önüne alındığında, insan üretimi olmayan içeriklerin tespiti için etkili yöntemlere olan ihtiyaç her geçen gün artmaktadır. Elde edilen bulgular, içerik tabanlı sınıflandırıcı sistemlerin yapay zekâ tarafından üretilmiş metinleri yüksek doğrulukla ayırt edebileceğini ve bu sayede akademik içeriklerin kalitesinin korunmasına katkı sağlayabileceğini göstermektedir. Bu bağlamda geliştirilen yöntemlerin üniversiteler, akademik dergiler ve araştırma kurumları tarafından desteklenmesi, bilimsel bilginin niteliği ve güvenilirliği açısından stratejik bir öneme sahiptir. Bu çalışmada geliştirilen sistem, yalnızca Türkçe akademik özetlerde değil, gelecekte farklı dil ve türdeki metinlerde de kullanılabilecek şekilde genelleştirilmeye uygundur.

1.1. Problem Durumu

Günümüzün hızla gelişen teknoloji dünyasında, yapay zekânın etin üretim yetenekleri önemli bir ivme kazanmıştır. Bu durum, insan üretimi içerikler ile yapay zekâ tarafından üretilen içerikler arasındaki sınırın giderek belirsizleşmesine neden olmaktadır. Akademik dünyanın temel prensipleri arasında yer alan güvenilirlik ve bilgi özgünlüğünün kalitesi bu belirsizlik nedeniyle tehdit altına girmektedir. Özellikle akademik alanda yapılan çalışmaların içeriklerinde yapay zekâ araçlarının kullanımı

hızla artmakta, ancak bu içeriklerin hangi oranda yapay zekâ tarafından üretildiğinin tespiti noktasında yetersizlikler bulunmaktadır. Yapay zekâ araçlarının akademik yazında kullanımı, bilgi üretiminin hızını ve erişilebilirliğini artırsa da bu araçların kontrolsüz ve etik dışı kullanımı ciddi sorunlara yol açabilmektedir. Bu bağlamda, akademik çalışmaların doğruluğunu, güvenilirliğini ve orijinalliğini korumak, akademik dünyanın sürdürülebilirliği açısından kritik bir öneme sahiptir. Bu nedenle, yapay zekâ tarafından üretilen içeriklerin tespit edilmesine yönelik yöntemlerin geliştirilmesi, akademik etik ve bilgi güvenliği açısından kaçınılmaz bir gereklilik olarak ortaya çıkmaktadır. Mevcut durumda, akademik yazında yapay zekâ ile üretilen içeriklerin tespitine yönelik belirgin ve yaygın olarak kabul görmüş yöntemler bulunmamaktadır. Bu eksiklik, bilgi kalitesinin düşmesine, akademik hırsızlık vakalarının artmasına ve akademik camiada güven kaybına neden olabilmektedir. Yapay zekâ ve insan üretimi içerikler arasındaki farkların net bir şekilde belirlenememesi, akademik dünya için önemli bir problem teşkil etmektedir. Dolayısıyla, bu alanda yapılacak çalışmalarda, yapay zekâ tarafından üretilen içeriklerin belirlenmesine yönelik güvenilir ve etkili araçların geliştirilmesi büyük bir ihtiyaç olarak öne çıkmaktadır. Özellikle dünya genelinde yaygın olarak kullanılan İngilizce dili için benzer çalışmalar varken, Türkçe dili bağlamında bu alandaki çalışmaların sınırlı sayıda olması, konunun Türkçe özelinde daha derinlemesine ele alınmasını ve yeni araştırmaların yapılmasını gerekli kılmaktadır.

Bu yöntemler, yapay zekâ ile insan üretimi içerikler arasındaki farkı belirginleştirmeli ve akademik etik standartlarına uygunluğu sağlanmalıdır. Yapay zekâ ile üretilen içeriklerin doğruluğunu ve güvenilirliğini belirleme konusunda şeffaf ve güvenilir bir mekanizmanın oluşturulması, akademik dünyanın sürdürülebilirliği açısından kritik bir öneme sahiptir. Akademik alanda yapay zekâ araçlarının artan kullanımıyla birlikte, içeriklerin yapay zekâ tarafından üretilmiş olup olmadığının belirlenmesinde yaşanan belirsizlikler, güvenilirlik ve bilgi kalitesi açısından ciddi bir problem oluşturmaktadır.

1.2. Problem Cümlesi

Yapay zekâ tarafından üretilen içerikler insan üretimi içeriklerden etkin bir şekilde ayırt edilebilir mi/tespit edilebilir mi?

1.3. Araştırmanın Amacı

Yapay zekânın metin üretimi yetenekleri nedeniyle, insan üretimi içerikler ile yapay zekâ üretimi içerikler arasındaki sınır belirsizleşmiştir. Bu durum, akademik dünyanın güvenilirliğini ve bilgi kalitesini tehlikeye atabilir. Türkçe gibi morfolojik açıdan karmaşık, biçimsel olarak zengin ve dil yapısı bakımından dünya genelinde daha az incelenmiş dillerde, yapay zekâ üretimi içeriklerin tespitine yönelik yöntemlerin sınırlı olması, bu tür içeriklerin ayırt edilmesini zorlaştırmakta ve bilgi kalitesine yönelik tehditleri artırmaktadır. Bu nedenle, çalışmanın Türkçe diline odaklanması, alandaki özgün katkıyı ve gerekliliği daha da pekiştirmektedir.

Yapılacak çalışmanın amacı, akademik alanda yapılan çalışmaların içeriklerinde kullanımı her geçen gün artan yapay zekâ araçları ile üretilen içeriklerin tespit edilmesine fayda sağlayacak bir model geliştirmektir. Çalışma, yapay zekâ tarafından üretilen içerikler ile insanlar tarafından üretilen içerikler arasındaki farkları belirlemeye yönelik yöntemler geliştirmeyi hedeflemektedir. Bu sayede, akademik alanda güvenilir ve etik içeriklerin korunmasına katkı sağlanması hedeflenmektedir. Bu katkı, akademik alanda sağlıklı bir bilgi akışını desteklemek adına kritik bir öneme sahiptir. Bunun için metin madenciliği, doğal dil işleme ve derin öğrenme gibi teknikleri kullanarak yapay zekâ tarafından üretilen içeriklerin özelliklerini analiz edilmesi hedeflenmektedir. Bu analizler, yapay zekâ ve insan üretimi içerikler arasındaki farkları belirlemeye yardımcı olacak ve insan üretimi olmayan içeriklerin tespit edilmesine yönelik bir temel sağlayacaktır.

Çalışmanın önemi yalnızca teknik bir model geliştirmekten ibaret değildir. Aynı zamanda bu model, akademik içeriklerin doğruluğunun korunması, etik ihlallerin önlenmesi ve sahte içeriklerin bilimsel yayınlarda yer bulmasının engellenmesi açısından önemli katkılar sunacaktır. Geliştirilecek model, üniversiteler, araştırma kurumları, editörler ve akademik dergiler tarafından bir ön denetim aracı olarak kullanılabilecek potansiyele sahiptir. Bunun yanı sıra, eğitim amacıyla da kullanılarak, akademik topluluğun yapay zekâ üretimi içeriklerin riskleri konusunda bilinçlenmesine katkı sağlama potansiyeline sahiptir.

1.4. Araştırmanın Önemi

Günümüzde üretken yapay zekâ sistemlerinin, özellikle doğal dil işleme temelli modellerin yaygınlaşması, akademik içerik üretiminde köklü bir paradigma değişikliğine yol açmıştır. ChatGPT, Gemini ve Smodin gibi büyük dil modeli uygulamalarının gelişmiş metin üretme becerileri sayesinde, kullanıcılar artık kısa süre içinde akademik nitelikte görünen içerikler üretebilmektedir. Bu durum, bilimsel yazının güvenilirliğini ve özgünlüğünü tehdit eder hâle gelmiştir. Zira yapay zekâ ile oluşturulan metinler, dilsel akıcılık ve anlamsal tutarlılık açısından yüksek kalite sunmakta; böylece insan üretimi içeriklerle görsel ve yapısal açıdan ayırt edilemez hâle gelmektedir. Bu bağlamda, yapay zekâ destekli içeriklerin tespit edilmesine yönelik etkili ve güvenilir araçların geliştirilmesi, bilimsel etik ve akademik orijinallik açısından kaçınılmaz bir gereklilik olarak öne çıkmaktadır.

Mevcut intihal tespit sistemleri, geleneksel anlamda benzerlik oranlarına dayalı olarak çalışmakta ve yapay zekâ tarafından sıfırdan oluşturulmuş, özgün görünümlü içerikleri tespit etmede yetersiz kalmaktadır. Bu nedenle, akademik kurumlar, hakemli dergiler ve araştırma kuruluşları açısından insan üretimi olmayan içerikleri belirleyebilecek daha gelişmiş, dilsel ve anlamsal düzeyde analiz yapabilen sistemlere ihtiyaç duyulmaktadır. Bu çalışmanın temel önemi de tam bu noktada ortaya çıkmaktadır. Geliştirilecek sınıflandırıcı sistem, üretim kaynağı bilinmeyen akademik metinlerin analiz edilmesini ve insan üretimi olup olmadıklarının belirlenmesini mümkün kılacak, böylece akademik şeffaflık ve bilimsel güvenilirliğe katkı sağlayacaktır.

Ayrıca bu çalışma, Türkçe metinler üzerinde yoğunlaşarak yerel bağlamda önemli bir boşluğu doldurmayı hedeflemektedir. Literatürde genellikle İngilizce içeriklere odaklanan çalışmaların aksine, bu araştırma Türkçe akademik içeriklerin yapay zekâ tarafından üretilip üretilmediğini saptamaya yönelik özgün bir yöntem geliştirmeyi amaçlamaktadır. Sağlık, eğitim, mühendislik ve sosyal bilimler gibi farklı disiplinlerden derlenen 6.000 adet akademik özetin analizine dayanan bu araştırma, alandaki sayılı kapsamlı Türkçe örneklerden biri olma özelliğini taşımaktadır.

Bu yönüyle çalışma, yalnızca teknik bir sınıflandırma modeli geliştirmekle kalmayacak; aynı zamanda etik içerik üretiminin desteklenmesine, akademik yozlaşmanın önlenmesine ve üniversitelerdeki yayın politikalarının güncellenmesine yönelik somut katkılar sunacaktır. Ayrıca geliştirilen sistemin, editörler, araştırmacılar

ve öğretim üyeleri tarafından bir ön denetim aracı olarak kullanılması mümkün olup, eğitim amaçlı olarak da akademik yazım sürecinde öğrencilere geri bildirim sağlayabilecek bir yapı olarak işlev görmesi beklenmektedir. Böylelikle, üretken yapay zekâ araçlarının oluşturduğu risklere karşı etkili bir kontrol mekanizması sunularak, akademik ekosistemin sürdürülebilirliği desteklenecektir.

1.5. Sınırlılıklar

Bu araştırmanın sınırlılıkları

1. Veri kümesi sağlık, mühendislik, sosyal bilimler ve eğitim bilimleri alanlarıyla sınırlandırılmıştır.
2. Kullanılan yapay zekâ araçları (ChatGPT, Gemini, Smodin) ve seçilen modeller ile elde edilen sonuçlarla sınırlandırılmıştır.
3. Araştırma yalnızca Türkçe dilinde yazılmış akademik çalışmalar üzerinde sınırlandırılmıştır.

1.6. Tanımlar

Gerçek Veri: İnsanlar tarafından üretilen orijinal akademik özetler.

Yapay Zekâ Verisi: Orijinal metinlerden yapay zekâ tarafından doğrudan üretilen özetler.

Düzenlenmiş Veri: İnsan üretimi özetlerin yapay zekâ araçları ile yeniden düzenlenmiş halleri.

Transformer Tabanlı Modeller: Doğal dil işleme alanında yaygın kullanılan, özellikle derin öğrenme tabanlı gelişmiş dil modelleri; XLM-Roberta, TurkishBERT ve BERT Base Multilingual gibi modelleri içerir.

BÖLÜM II

KURAMSAL ÇERÇEVE VE İLGİLİ ARAŞTIRMALAR

2.1. Kuramsal Çerçeve

2.1.1. Doğal dil işleme (Natural Language Processing- NLP). Doğal dil işleme insan diliyle bilgisayar sistemlerinin etkileşimini sağlayan ve bu dili anlamaya, analiz etmeye ve üretmeye yönelik yöntemleri kapsayan bir yapay zekâ alt alanıdır. Dilin yapısal ve anlamsal boyutlarını ele alan bu disiplin, metin ve konuşma verilerini işleyerek anlamlı bilgiler çıkarmayı amaçlar. İnsan dilinin yapısal karmaşıklığı, bağlamsal çeşitliliği ve anlamsal belirsizliği nedeniyle doğal dil işleme, disiplinler arası bir yaklaşım gerektirir; dilbilim, bilgisayar bilimi, istatistik ve makine öğrenimi gibi birçok alanın birleşiminden oluşur (Adalı, 2016). Metinlerin otomatik olarak analiz edilmesi, anlamlandırılması ve uygun çıktılar üretilmesi gibi görevleri üstlenen doğal dil işleme, dijital çağda veriyle etkileşimin temel araçlarından biri haline gelmiştir.

Doğal dil işleme süreçleri, genellikle dilin farklı analiz seviyelerine dayanır. Bu seviyeler arasında morfolojik analiz (kelimelerin kök ve ek yapılarını belirleme), sözdizimsel analiz (cümle içi yapısal ilişkilerin çözümü), anlamsal analiz (sözcük ve cümle düzeyinde anlamın çıkarımı) ve pragmatik analiz (bağlam temelli yorumlama) yer alır. Bu analiz türleri, metin verisinin daha derinlikli biçimde anlaşılmasını sağlar. Bu çok katmanlı yaklaşım, doğal dilin iç yapısını modelleyerek bilgi çıkarımı, duygu analizi, soru-cevap sistemleri ve otomatik özetleme gibi uygulamalarda yüksek doğrulukla sonuçlar elde edilmesine imkân tanır (Talan & Aktürk, 2021).

2.1.2. Doğal dil işlemenin tarihçesi ve gelişimi. Doğal dil işlemenin tarihsel gelişimine bakıldığında, ilk adımlar 1950’li yıllarda atılmıştır. Alan Turing’in 1950’de yayımladığı ve makinelerin düşünme yetisine sahip olup olamayacağını tartıştığı makalesiyle başlayan süreç, “Turing Testi” kavramını ortaya koymuştur. Bu test, bir makinenin insanla ayırt edilemeyecek düzeyde iletişim kurabilme yetisini değerlendirme amacı taşır (Turing, 1950). Bu dönemde dilin işlenmesine dair ilk uygulama alanı makine çevirisi olmuştur. Özellikle 1954’te gerçekleştirilen

Georgetown-IBM deneyi, İngilizceden Rusçaya yapılan otomatik çeviriyle alanda önemli bir ivme yaratmıştır. Basit sözcük dağarcığı ve sınırlı gramer kurallarına rağmen elde edilen başarılı sonuçlar, makine çevirisinin gelecek vaat eden bir alan olduğunu göstermiştir (Hutchins, 2004).

1960'lı yıllar, sembolik (kural tabanlı) yaklaşımların ön planda olduğu bir dönem olmuştur. Bu yöntemler, dilin kurallarla ifade edilmesini ve bilgisayar sistemlerinin bu kurallar doğrultusunda işlemler yapmasını öngörür. Dönemin en bilinen uygulamalarından biri, ELIZA adlı programdır. ELIZA, önceden tanımlı kalıplarla kullanıcılara yanıt vererek temel düzeyde diyaloglar kurabilmiş ve bilgisayarların dil etkileşimine dair potansiyelini göstermiştir. Ancak sembolik yöntemlerin esneklikten yoksun olması ve dildeki varyasyonları sınırlı şekilde modelleyebilmesi, alternatif yaklaşımların gerekliliğini doğurmuştur (Weizenbaum, 1966).

1980'ler ve 1990'lar, doğal dil işlemede istatistiksel yöntemlerin ön plana çıktığı bir dönüşüm sürecidir. Bu dönemde, büyük veri kümelerinden öğrenen sistemler geliştirilmeye başlanmış ve kurallara dayalı sistemlerin yerini olasılık temelli modeller almaya başlamıştır. Gizli Markov Modelleri (Hidden Markov Model- HMM), özellikle konuşma tanıma ve sözcük türü etiketleme gibi görevlerde yüksek başarı sağlamıştır (Rabiner, 1989). HMM, gözlemlenebilir veri dizileri üzerinden gizli durumların modellenmesini sağlayarak dilin istatistiksel özelliklerini yakalayabilmiştir. 1990'lı yıllarda ise açıklama veya etiket içeren (anotasyonlu-annotated) dil veri kümelerinin (lügat-corpus) yaygınlaşması, NLP sistemlerinin gelişiminde önemli rol oynamıştır (Marcus vd., 1993). Penn Treebank gibi kapsamlı kaynaklar, dil modellerinin eğitilmesi için referans veri sağlamıştır. Aynı dönemde IBM tarafından geliştirilen CANDIDE adlı istatistiksel makine çevirisi sistemi, dil çiftleri arasındaki olasılık ilişkilerini modelleyerek çeviri doğruluğunda büyük gelişme kaydetmiştir (Brown vd., 1990).

Sonuç olarak, doğal dil işleme alanı, sembolik yöntemlerden istatistiksel modellere ve günümüzde derin öğrenme temelli yaklaşımlara evrilmiş; insan diliyle bilgisayar sistemlerinin anlamlı biçimde iletişim kurabilmesini sağlayan temel bir yapay zekâ bileşeni haline gelmiştir. Bu gelişim, büyük veri kaynaklarının artması ve hesaplama gücünün ilerlemesiyle birlikte, doğal dilin dijital ortamlarda daha doğru ve etkili şekilde işlenmesini mümkün kılmıştır.

2.1.3. Doğal dil yapısı. Doğal dil yapısını beş aşamada incelemek mümkündür (Covington, 1994):

- Ses bilgisi (fonoloji)
- Biçim bilgisi (morfoloji)
- Söz dizim bilgisi (syntactic)
- Anlam bilgisi (semantik)
- Kullanım bilgisi (pragmatik)

2.1.3.1. Ses bilgisi. Ses bilgisi, harflerin oluşturduğu ses yapılarının incelenmesi ve bu yapıların dil içerisindeki işlevlerinin analiz edilmesiyle ilgilenen dilbilim dalıdır. Doğal dillerde anlamı ayırt eden en küçük ses birimine fon (phon) adı verilir. Sözlü dilin yazılı dile aktarılabilmesi için bu ses birimlerinin belirli sembollerle, yani harflerle eşleştirilmesi gereklidir. Bazı dillerde her fon belirli bir harfle doğrudan temsil edilirken, bazı dillerde bu birebir karşılık olmayabilir. Her dil, kendine özgü bir alfabe sistemine sahiptir ve bu sistemdeki her harfin belirli bir ses değeri bulunmaktadır. Ancak bu ses değerleri, aynı dilin farklı sözcüklerinde çeşitlilik gösterebileceği gibi, dilden dile de değişkenlik arz edebilir. Ses bilgisi, konuşma dilindeki seslerin duyumsanması, yapısal özelliklerinin çözülmesi ve bu seslerin dilsel bağlamlar içerisindeki işlevlerinin araştırılmasıyla ilgilenir.

2.1.3.2. Biçim bilgisi. Morfolojik çözümleme, sözcüklerin iç yapısının sistematik olarak incelendiği ve bu yapının dilbilgisel işlevleriyle ilişkilendirildiği dilsel bir süreçtir. Bu analizde, kelimelerin kök biçimleri ve taşıdığı ekler ayrıştırılarak sözcüğün türetilmiş ya da çekimlenmiş hali tespit edilir. Morfolojik analiz, özellikle Türkçe gibi eklemeli dillerde önemli bir yer tutmakta ve dilin anlamlandırılmasında kritik bir rol üstlenmektedir. Bu süreçte yalnızca biçim birimsel yapılar değil, aynı zamanda kelimenin ait olduğu sözcük türü, zaman, çoğulluk ya da kişi gibi dilbilgisel özellikler de belirlenmektedir. Morfolojik analiz, dilin yüzey yapısının ötesine geçerek köken bilgisi ve biçimbilimsel özelliklerle doğal dil işleme sistemlerine derinlik kazandırır. Sözcüklerin doğru çözümlemesi, daha üst düzey işlemler olan sözdizimsel

ya da anlamsal analizlerin doğruluğunu doğrudan etkiler. Otomatik morfolojik ayrıştırma araçları, bu süreci hızlandırmakta ve özellikle bilgi çıkarımı, makine çevirisi ve anlamsal analiz gibi uygulamalarda temel girdi olarak kullanılmaktadır.

2.1.3.3. Söz dizim bilgisi. Söz dizimsel analiz, cümle yapılarının dilbilgisel kurallara uygun biçimde incelenmesi ve bu yapılar içerisindeki sözcüklerin birbirleriyle olan ilişkilerinin belirlenmesidir. Bu analiz türü, cümle bileşenlerinin görevlerini (özne, nesne, yüklem vb.) ve bunlar arasındaki yapısal bağlantıları saptayarak, dilin biçimsel organizasyonunun çıkarılmasını sağlar. Sözdizimsel çözümleme, dilin anlamının doğru bir şekilde yorumlanabilmesi için vazgeçilmez bir aşamadır; çünkü anlam genellikle kelime dizilişi ve sözdizimsel yapı aracılığıyla ifade edilir. Bu süreçte bağımlılık gramerleri (dependency grammars) ve söz dizimi ağaçları (parse trees) gibi yapılar kullanılarak kelimeler arası ilişkiler modellenir. Sözdizimsel analiz, makine çevirisi, soru-cevap sistemleri ve metin sınıflandırma gibi pek çok doğal dil işleme uygulamasının altyapısını oluşturmaktadır. Bu bağlamda, otomatik sözdizimsel ayrıştırıcılar, bir metindeki gramatikal yapıların belirlenmesinde yüksek doğrulukla çalışan araçlardır.

2.1.3.4. Anlam bilgisi. Anlamsal analiz, dilin içerdiği anlamın biçimsel olarak çözülmesini ve metnin içeriksel boyutunun ortaya konulmasını amaçlayan bir süreçtir. Bu analizde temel olarak, sözcüklerin bireysel anlamlarının yanı sıra cümle ve metin düzeyindeki kavramsal bütünlük de değerlendirilir. Dilin doğasında bulunan çok anlamlılık (ambiguity), eş anlamlılık (synonymy) ve bağlama göre değişen anlam (polysemy) gibi zorluklar, anlamsal analiz sürecinde dikkate alınır. Bu nedenle, anlamsal analiz yalnızca sözcük düzeyinde kalmaz; aynı zamanda sözcükler arası anlamsal ilişkiler, bağlam temelli çözümlemeler ve anlam çıkarımı yöntemleri ile desteklenir. Anlam çözümleme sürecinde kullanılan teknikler arasında kelime ağı (WordNet), vektör uzay temsilleri, kelime gömme (word embedding) ve anlamsal benzerlik ölçümleri öne çıkar. Bu analiz, özellikle metin özetleme, bilgi çıkarımı ve anlam tabanlı arama motorları gibi alanlarda vazgeçilmez bir rol oynamaktadır.

2.1.3.5. Kullanım bilgisi. Pragmatik analiz, dilin yalnızca biçimsel ve anlamsal yönlerini değil, aynı zamanda kullanım bağlamını, konuşur niyetlerini ve iletişim ortamının etkilerini de kapsayan bir süreçtir. Bu çözümleme, dilin söylem bağlamında nasıl işlediğini ve dil kullanıcılarının sosyal rollerine göre anlamın nasıl şekillendiğini değerlendirir. Pragmatik çözümleme sürecinde konuşmacının amacı, dinleyicinin bilgi düzeyi, konuşma ortamı, örtük anlamlar ve ima gibi etmenler dikkate alınarak metnin bağlama özgü yorumu yapılır. Doğal dil işleme sistemlerinde pragmatik analiz, özellikle diyalog sistemleri, sanal asistanlar ve etkileşimli uygulamalarda kullanıcı deneyimini derinleştirmek adına kritik önemdedir. Bu bağlamda, soru-cevap etkileşimleri, emir-komut analizleri ve duygusal niyet belirleme gibi görevlerde pragmatik bilgi, sistemin doğru yanıt üretme becerisini artırmaktadır. Dilin bağlama duyarlı yorumlanması, yalnızca söylenenin değil, nasıl ve neden söylendiğinin de anlaşılmasını gerekli kılmaktadır. Bu nedenle, pragmatik analiz, dilin gerçek dünyadaki işlevselliğini modellemeye yönelik önemli bir bileşen olarak değerlendirilir.

2.1.4. Doğal dil üretimi (Natural Language Generation- NLG). Doğal dil üretimi bilgisayar sistemlerinin yapılandırılmış verileri anlamlı ve bağlama uygun doğal dil metinlerine dönüştürme sürecidir. Bu teknoloji, yapay zekâ tabanlı sistemlerin insan benzeri açıklamalar, özetler veya anlatılar oluşturmasını mümkün kılar. Doğal dil üretimi öncelikle verilerin analiz edilmesi, ardından bu analiz doğrultusunda dilbilgisel olarak tutarlı cümlelerin oluşturulması adımlarını içerir. Bu süreçte dilin sözdizimsel, anlamsal ve pragmatik boyutları dikkate alınır. Özellikle otomatik rapor oluşturma, haber metinleri yazma, finansal özetler üretme ve müşteri hizmetleri alanlarında geniş kullanım alanı bulmuştur. Günümüzde birçok medya kuruluşu ve kurumsal yapı, haber üretimi ve veri raporlamada doğal dil üretimi sistemlerini aktif olarak kullanmaktadır (Dong vd. 2022).

2.1.5. Dil modellemesi. Dil modellemesi, bir dildeki sözcüklerin ve cümlelerin olasılık dağılımlarını tahmin ederek, dilin yapısal ve anlamsal özelliklerini modellemeyi amaçlayan bir yaklaşımdır. Bu modelleme, doğal dil işleme sistemlerinde metinlerin otomatik analizinde, üretiminde ve çevirisinde temel rol

oyunar. Klasik n-gram modellerinden, son yıllarda yaygınlaşan Transformer tabanlı ileri düzey modellere (GPT, BERT gibi) kadar çok çeşitli dil modelleme teknikleri geliştirilmiştir. Bu modeller, büyük veri kümeleri üzerinde eğitilerek dilin istatistiksel örüntülerini öğrenmekte ve bağlamdan yola çıkarak anlamlı tahminlerde bulunabilmektedir. Dolayısıyla, dil modellemesi hem dilin içsel yapısını çözümleme hem de makine tarafından üretilecek dilin doğruluğunu ve bağlamsallığını sağlama açısından kritik öneme sahiptir.

2.1.6. Duygu analizi. Duygu analizi (sentiment analysis), yazılı ya da sözlü ifadelerden kullanıcının duygusal eğilimini (örneğin, olumlu, olumsuz ya da nötr) belirlemeyi hedefleyen doğal dil işleme alt alanıdır. Bu analiz türü, kullanıcıların düşünce ve duygularını yorumlamak için metin madenciliği, makine öğrenmesi ve istatistiksel yöntemlerle birleştirilmiş algoritmalar kullanır. Sosyal medya analizlerinden müşteri geri bildirimlerinin sınıflandırılmasına, kamuoyu yoklamalarından film/dizi değerlendirmelerine kadar birçok alanda duygu analizi teknolojisinden yararlanılmaktadır. Duyguların doğru bir şekilde sınıflandırılması, kurumların hedef kitle analizleri yapmasını, pazarlama stratejilerini yeniden yapılandırmasını ve kriz yönetiminde hızlı refleks göstermesini mümkün kılmaktadır.

2.1.7. Makine çevirisi. Makine çevirisi (machine translation), bir dildeki metinlerin başka bir dile otomatik olarak çevrilmesini sağlayan doğal dil işleme teknolojisidir. Bu süreç, kaynak dildeki metnin sözdizimsel ve anlamsal analizinin yapılması ve elde edilen içeriğin hedef dildeki karşılığıyla eşleştirilmesi adımlarını içerir. Geleneksel kurallara dayalı yöntemlerden, istatistiksel modellere (SMT) ve son olarak sinir ağı tabanlı çeviri sistemlerine (NMT) kadar önemli evrimler geçirmiştir. Günümüzde Google Translate, DeepL, Microsoft Translator gibi sistemler, yapay sinir ağları ve büyük dil modelleri sayesinde anlamlı ve bağlama uygun çeviriler sunabilmektedir. Makine çevirisi, dil engellerinin aşılmasında, uluslararası iletişimde ve çok dilli içerik üretiminde büyük kolaylık sağlamaktadır.

2.1.8. Konuşma tanıma ve sentezi. Konuşma tanıma (speech recognition), sesli girdilerin metin biçimine dönüştürülmesini sağlayan bir doğal dil işleme alt dalıdır. Bu teknoloji, ses dalgalarının dijital sinyallere çevrilmesi, bu sinyallerin sözcük ya da cümlelerle eşleştirilmesi ve en sonunda doğal dil metni hâline getirilmesi süreçlerini içerir. Siri, Alexa, Google Assistant gibi sanal asistanlarda yaygın olarak kullanılmaktadır.

Konuşma sentezi (speech synthesis) ise, yazılı metinlerin doğal ve akıcı biçimde sesli konuşmaya dönüştürülmesini sağlar. Metin okuma sistemleri (Text-to-Speech- TTS), eğitim, sağlık, otomasyon ve engellilere yönelik teknolojilerde önemli rol oynamaktadır. Her iki teknoloji de kullanıcıların makinelerle daha doğal ve sezgisel etkileşim kurmalarına olanak tanır.

2.1.9. Doğal dil işleme uygulama alanları. Doğal dil işleme teknolojileri, insan dilini anlamaya, analiz etmeye ve üretmeye yönelik geniş kapsamlı uygulamalarıyla dikkat çekmektedir. Günümüzde birçok sektörde dil temelli problemleri çözmeye yönelik sistemlerin tasarımı ve entegrasyonu sağlanmaktadır.

2.1.9.1. Yazım yanlışlarının düzeltilmesi. Doğal Dil İşleme teknolojileri, metinlerde yer alan yazım yanlışlarının otomatik olarak düzeltilmesi konusunda önemli bir rol oynamaktadır. Yazım hatalarının tespiti ve düzeltilmesi, özellikle büyük ölçekli metin verilerinin işlendiği uygulamalarda verimliliği artıran kritik bir adımdır. Bu süreç, metinlerin daha anlaşılır ve doğru olmasını sağlayarak hem kullanıcı deneyimini iyileştirir hem de bilgi işlem sistemlerinin doğruluğunu artırır (Adalı, 2016: 46). Yazım yanlışlarının düzeltilmesi için kullanılan yöntemler arasında sözlük tabanlı yaklaşımlar, dilbilgisel kurallara dayalı yöntemler ve istatistiksel veya derin öğrenme tabanlı modeller bulunmaktadır. Örneğin, sözlük tabanlı yöntemlerde, metindeki kelimeler mevcut bir kelime listesiyle karşılaştırılarak yazım hataları tespit edilir. Bu yöntem, eksik harf, fazla harf veya yanlış harf içeren kelimelerin belirlenmesi için kullanılır.

Dilbilgisel kurallara dayalı yöntemler ise daha karmaşık dil yapılarının anlaşılmasını gerektirir. Bu yaklaşımlar, cümlenin yapısı ve kelime türleri arasındaki ilişkileri

dikkate alarak doğru kelime formlarını belirler. Örneğin, Türkçede eklerin yanlış kullanımı veya fiil çekim hataları bu tür yöntemlerle analiz edilebilir. Son yıllarda derin öğrenme tekniklerinin gelişmesiyle birlikte, yazım düzeltme sistemlerinde dilin bağlamını daha iyi anlayabilen modeller yaygınlaşmıştır. Bu modeller, büyük veri kümelerinden öğrenerek kelime anlamlarını ve bağlam ilişkilerini daha doğru şekilde analiz edebilir. Bu sayede, sadece basit yazım hataları değil, aynı zamanda anlam kaymalarını da düzeltebilen daha güçlü sistemler geliştirilmiştir.

Yazım yanlışlarının düzeltilmesi, birçok farklı uygulama alanında kullanılır. Özellikle metin tabanlı arama motorları, sosyal medya platformları ve otomatik müşteri hizmetleri gibi geniş kullanıcı kitlesine hitap eden sistemlerde bu teknolojiler kritik öneme sahiptir. Ayrıca akademik yazılar, resmî belgeler ve haber metinleri gibi yüksek doğruluk gerektiren içeriklerde de yazım düzeltme araçları yaygın olarak kullanılmaktadır.

Bu bağlamda, doğal dil işleme teknolojilerinin gelişimi, doğru ve tutarlı metinlerin oluşturulmasını sağlamak için büyük önem taşımaktadır. Yazım hatalarının tespiti ve düzeltilmesi üzerine yapılan çalışmalar, dilin bilgisayarlar tarafından daha iyi anlaşılmasını sağlarken, aynı zamanda kullanıcıların metinlerini daha doğru ve etkili bir şekilde ifade etmelerine olanak tanımaktadır.

2.1.9.2. Büyük metin dosyalarından özet çıkarma. Büyük boyutlu metin dosyalarından özet çıkarma, doğal dil işleme alanında önemli bir araştırma konusudur. Özetleme, uzun metinlerin kısa ve öz bir şekilde ifade edilmesini sağlayarak bilgiye hızlı erişim imkânı sunar. Bu süreç, metindeki temel bilgilerin, ana düşüncelerin ve önemli noktaların tespit edilmesine dayanır. Özetleme sistemleri genellikle iki ana yaklaşımla geliştirilir: çıkarımsal (extractive) ve üreteçsel (abstractive) özetleme. Çıkarımsal özetleme yöntemleri, orijinal metindeki cümleleri doğrudan seçerek özet oluştururken, üreteçsel yöntemler daha karmaşık olup metnin anlamını analiz ederek özgün cümleler oluşturur.

Çıkarımsal yöntemler, dilbilgisi kuralları ve kelime frekansları gibi metin içi özelliklere dayanarak önemli cümleleri seçme eğilimindedir. Bu tür sistemler, özellikle haber özetleri, yasal belgeler ve bilimsel makaleler gibi yapılandırılmış metinlerde etkili olabilir. Üreteçsel yöntemler ise doğal dilin daha derin bir anlayışını gerektirir.

Bu yöntemler, makine öğrenme ve derin öğrenme tekniklerini kullanarak metnin bağlamını, anlamını ve dil yapısını dikkate alır. Bu sayede daha doğal ve akıcı özetler üretilir.

Bu alanda yapılan çalışmalar, metin özetleme sistemlerinin doğruluğunu artırmak için metinlerin anlamsal analizine odaklanmaktadır. Özellikle büyük veri kümeleri üzerinde eğitilen modeller, dilin bağlamsal ilişkilerini daha iyi kavrayarak özetleme kalitesini artırmaktadır. Günümüzde, haber siteleri, akademik veri tabanları ve sosyal medya platformları gibi birçok uygulama alanında özetleme teknolojileri yaygın olarak kullanılmaktadır.

Bu tür sistemlerin geliştirilmesi, bilgiye hızlı erişim sağlamak ve kullanıcıların bilgi yükünü azaltmak açısından büyük önem taşımaktadır. Metin özetleme, doğal dil işleme teknolojilerinin en zorlu ve en kritik uygulamalarından biri olarak kabul edilmektedir.

2.1.9.3. Ana bilgi tespiti. Büyük veri çağında, metinlerden ana bilgiyi doğru ve hızlı bir şekilde çıkarmak, doğal dil işleme teknolojilerinin temel hedeflerinden biridir. Metinlerde yer alan özne, fiil, yer ve zaman gibi temel bileşenlerin belirlenmesi, bilgiye dayalı karar destek sistemlerinin verimliliğini artırır. Bu süreç, metinlerin anlam yapısının doğru bir şekilde analiz edilmesini ve anahtar bilgilerin etkin bir şekilde çıkarılmasını gerektirir. Örneğin, haber metinlerinde olayların ne zaman ve nerede gerçekleştiği, bilimsel makalelerde ana hipotezler ve bulgular gibi kritik unsurlar bu yöntemlerle tespit edilebilir. Bu amaçla, ad öbeği çıkarma, özne-fiil ilişkisinin belirlenmesi ve anlamsal rol etiketleme gibi teknikler kullanılır. Doğal dil işleme sistemlerinde bu tür analizlerin yapılması, metin madenciliği, bilgi çıkarımı ve belge sınıflandırma gibi birçok uygulama için önemli bir adım teşkil etmektedir.

2.1.9.4. Talep anlamlandırma. Talep anlamlandırma, özellikle müşteri hizmetleri, sanal asistanlar ve otomatik çağrı merkezleri gibi alanlarda önemli bir yer tutmaktadır. Bu teknolojiler, kullanıcının yazılı veya sözlü olarak ilettiği talepleri anlamak ve uygun yanıtları sunmak amacıyla geliştirilir. Doğal dil işleme teknikleri, kullanıcıların ifadelerini analiz ederek bu taleplerin türünü ve içeriğini belirler.

Örneğin, "Siparişimi iptal etmek istiyorum" gibi bir cümledeki niyetin doğru anlaşılması, müşteri memnuniyetini artırmak için kritik öneme sahiptir. Bu bağlamda, duygu analizi, niyet sınıflandırması ve bağlam algılama gibi teknikler sıkça kullanılmaktadır. Ayrıca, bu tür sistemlerin daha etkin çalışabilmesi için metinlerin dil yapısı, bağlamı ve kullanıcı geçmişi gibi faktörlerin de dikkate alınması gerekmektedir.

2.1.9.5. Arama motorları. Arama motorları, kullanıcıların internet üzerinde bilgiye hızlı erişim sağlayabilmesi için gelişmiş doğal dil işleme teknikleri kullanır. Bu sistemler, arama sorgularını analiz ederek en alakalı sonuçları sunmak üzere tasarlanmıştır. Kullanıcıların girdikleri anahtar kelimeler, cümle yapıları ve bağlamsal ipuçları değerlendirilerek en uygun sonuçların seçilmesi sağlanır. Örneğin, "en iyi yapay zekâ kitapları" gibi bir arama sorgusu, içeriğinde "yapay zekâ" ve "kitap" terimlerinin sıkça geçtiği sayfalarla eşleştirilir. Bunun yanı sıra, semantik analiz, kelime anlamı çıkarma ve kullanıcı davranışı takibi gibi ileri seviye doğal dil işleme teknikleri de bu sistemlerde yaygın olarak kullanılmaktadır. Bu sayede, kullanıcıların arama deneyimi daha kişisel ve etkili hale getirilebilir.

2.1.9.6. Konuşma metne dönüştürme ve metin seslendirme. Konuşma metne dönüştürme (speech-to-text) ve metin seslendirme (text-to-speech) teknolojileri, sesli komut sistemleri, otomatik altyazı oluşturma ve engelli bireyler için geliştirilen iletişim araçlarında yaygın olarak kullanılmaktadır. Bu teknolojiler, konuşulan dilin doğru bir şekilde metne aktarılması veya metinlerin doğal seslerle okunmasını sağlar. Örneğin, telefonlarda kullanılan sesli asistanlar veya video platformlarındaki altyazı sistemleri bu teknolojilerin günlük kullanımına örnek olarak verilebilir. Konuşma analizi, fonetik modelleme ve sinirsel dil modelleri gibi tekniklerle desteklenen bu uygulamalar, dilin doğal akışını ve tonlamalarını daha iyi anlamak için sürekli geliştirilmektedir.

2.1.9.7. Doğal dil çevirisi. Doğal dil çevirisi, farklı diller arasındaki iletişimi kolaylaştırmak amacıyla geliştirilen önemli bir doğal dil işleme uygulamasıdır. Bu

teknolojiler, metinlerin veya konuşmaların hedef dile en doğru şekilde çevrilmesini sağlar. Geleneksel kural tabanlı yaklaşımlardan, günümüzde yaygın olarak kullanılan sinirsel makine çevirisi (Neural Machine Translation, NMT) yöntemlerine kadar birçok teknik bu alanda kullanılır. Özellikle bağlam bağımlı anlam analizi, dil çiftleri arasındaki semantik farklılıkların doğru aktarılmasında kritik öneme sahiptir. Bu teknolojiler, sosyal medya platformları, seyahat uygulamaları ve uluslararası iş iletişimi gibi birçok alanda geniş bir kullanım alanına sahiptir.

2.1.9.8. Soru yanıtlama sistemleri. Soru yanıtlama sistemleri, kullanıcıların belirli bir konu hakkında sordukları sorulara otomatik ve doğru yanıtlar verebilen doğal dil işleme teknolojileridir. Bu sistemler, bilgi tabanlarından, dökümanlardan veya büyük veri kümelerinden anlamlı yanıtlar üretmek için geliştirilmiştir. Örneğin, sanal asistanlar, chatbotlar ve müşteri destek uygulamaları bu tür sistemlerin yaygın kullanım alanlarıdır. Bu teknolojilerin başarısı, metinlerin bağlamsal analizi, soru türlerinin sınıflandırılması ve doğru bilgi çıkarımı gibi faktörlere bağlıdır. Gelişmiş soru yanıtlama sistemleri, kullanıcıların bilgiye hızlı ve doğru bir şekilde erişmesini sağlayarak verimliliği artırmaktadır.

2.1.9.9. Sohbet robotları (Chatbots). Sohbet robotları, kullanıcılarla doğal dil aracılığıyla etkileşim kurabilen otomatik yanıt sistemleridir. Kullanıcıların yazılı ya da sesli taleplerine anında geri dönüş sağlayarak müşteri hizmetleri, rehberlik, bilgi sağlama gibi görevlerde kullanılmaktadır. Chatbot sistemleri, doğal dil anlama (NLU), yanıt üretme (NLG) ve bağlamsal bellek gibi bileşenleri bir araya getirir. Facebook Messenger, WhatsApp, Slack gibi platformlarda kullanılan müşteri hizmeti botları, bu teknolojinin somut örnekleri arasında yer alır.

2.1.9.10. Metin sınıflandırma (Text Classification). Metin sınıflandırma, dijital belgelerin içeriklerine göre önceden belirlenmiş temalara, kategorilere veya etiketlere otomatik biçimde atanması sürecidir. Bu yöntem, özellikle doğal dil işleme ve makine öğrenimi algoritmalarının etkileşimli olarak kullanılmasıyla gerçekleştirilmektedir. Günümüzde çeşitli uygulama alanları arasında spam tespiti,

haber başlıklarının kategorilendirilmesi, sosyal medya içeriklerinin duygu temelli analizi ve müşteri geri bildirimlerinin değerlendirilmesi yer almaktadır. Örneğin, bir e-posta uygulaması, kullanıcıya ulaşan mesajları "önemli", "tanıtım", "spam" gibi gruplara ayırırken, bu sınıflandırmayı içerik analizine dayalı karar ağacı veya destek vektör makineleri gibi öğrenme algoritmalarını kullanarak yapar. Metin sınıflandırma süreci genellikle iki aşamada yürütülmektedir: ilk aşamada, metinlerdeki dilsel özellikler (örneğin anahtar kelimeler, sözcük sıklıkları, sözdizimsel yapı vb.) çıkarılır; ardından bu özellikler üzerinden makine öğrenmesi algoritmaları aracılığıyla sınıflandırma işlemi gerçekleştirilir. Bu bağlamda, metinler üzerinde yapılan kelime vektörleştirme (örneğin TF-IDF, Word2Vec gibi) teknikleri ile metin, sayısal biçime dönüştürülerek algoritmalara uygun hale getirilir. Son yıllarda, derin öğrenme tabanlı modeller (özellikle LSTM, BERT gibi dönüşüm tabanlı yapılar) metin sınıflandırma doğruluğunu önemli ölçüde artırmıştır. Bu teknolojiler, dilin anlamsal bağlamını daha derin düzeyde kavrayarak, metinleri yalnızca kelime düzeyinde değil, bağlam ve anlam düzeyinde de analiz edebilmektedir. Metin sınıflandırma, böylece bilgi yönetimi, içerik filtreleme, otomatik karar destek sistemleri ve dijital güvenlik gibi birçok alanda kritik işlevler üstlenmektedir.

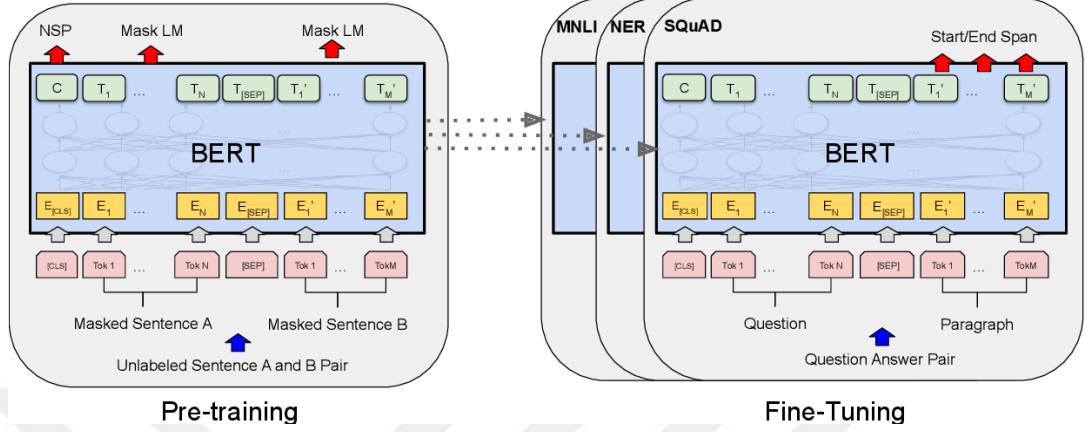
2.1.9.11. Metin özetleme (Text Summarization). Metin özetleme, hacimli yazılı içeriklerin temel bilgilerini veya ana mesajlarını koruyarak daha kısa ve anlamlı bir biçime indirgenmesi sürecidir. Bu süreç, özellikle bilgiye hızlı erişimin gerekli olduğu durumlarda kullanıcıya büyük kolaylık sağlamaktadır. Özetleme işlemi, genel olarak iki yöntemle gerçekleştirilir: çıkarımsal (extractive) ve türetici (abstractive) özetleme. Çıkarımsal özetleme, orijinal metindeki en anlamlı cümle veya paragrafların seçilmesi yoluyla özet üretirken; türetici özetleme ise kaynak metindeki bilgilerden hareketle yeni ifadeler oluşturarak daha doğal ve akıcı metinler oluşturur. Doğal dil işleme teknikleri bu noktada metnin yapısını, bağlamını ve anlamsal bütünlüğünü analiz ederek özetleme sürecini yönlendirir. Geleneksel yöntemlerde TF-IDF gibi istatistiksel modeller kullanılsa da günümüzde özellikle derin öğrenme mimarilerine dayalı seq2seq modelleri ve transformer yapıları (örneğin BART, T5 gibi) özetleme işlemlerinde yüksek başarı sağlamaktadır. Bu sistemler, uzun metinleri anlamlandırma, önemli bilgileri tespit etme ve dilsel bütünlüğü koruma noktasında

gelişmiş yetenekler sunmaktadır. Eğitim, haber, sağlık ve hukuk gibi disiplinlerde özetleme araçları; raporlar, araştırmalar veya belgeler üzerinde zaman tasarrufu sağlamanın yanı sıra, bilgi yoğunluğunun azaltılmasına da katkı sağlamaktadır. Özellikle eğitimde metin özetleyiciler, öğrencilerin içerikleri daha verimli şekilde kavramalarını kolaylaştırmakta ve bireysel öğrenme süreçlerini desteklemektedir. Özetleme teknolojileri, bilgiye erişimin hızlanması ve karar destek sistemlerinin etkinliğinin artırılması açısından stratejik öneme sahiptir.

2.1.10. BERT (Bidirectional Encoder Representations from Transformers). BERT (Bidirectional Encoder Representations from Transformers), Google tarafından 2018 yılında geliştirilen ve doğal dil işleme alanında devrim yaratan bir dil modelidir. Transformer mimarisine dayanan BERT, dilin bağlamsal anlamını iki yönlü olarak öğrenme yeteneğine sahiptir, yani hem soldan sağa hem de sağdan sola bağlamı aynı anda değerlendirir. Bu, modelin metinlerdeki kelimelerin anlamını daha derin ve doğru bir şekilde anlamasına olanak tanır. BERT'in eğitimi, büyük miktarda metin verisi üzerinde, iki temel görev olan Masked Language Modeling (MLM) ve Next Sentence Prediction (NSP) kullanılarak gerçekleştirilir. İki temel görev olan MLM ve NSP kullanılarak gerçekleştirilir. MLM, metindeki rastgele seçilen kelimelerin maskelenmesi ve modelin bu maskeleri doğru şekilde tahmin etmesi üzerine kurulu bir yaklaşımdır. Bu yöntem, modelin bağlamı daha iyi anlamasını sağlar. NSP ise modelin iki cümlelerin birbiriyle ilişkili olup olmadığını öğrenmesine olanak tanır, bu da daha uzun metinlerin anlaşılmasını kolaylaştırır (Devlin vd., 2019).

BERT, metin sınıflandırması, duygu analizi, soru-cevap sistemleri, adlandırılmış varlık tanıma NER ve makine çevirisi gibi birçok doğal dil işleme görevinde üstün performans göstermektedir. BERT 'in başarısı, modelin bağlamı derinlemesine anlaması ve dilin karmaşık yapısını yakalamasıyla ilişkilidir. Bu nedenle, BERT tabanlı modeller, çok dilli ve düşük kaynaklı dillerde bile yüksek doğruluk oranlarına ulaşabilmektedir. BERT, 2019 yılında Google araştırmacılarından Jacob Devlin ve ekibi tarafından geliştirilmiş olup, dilin bağlamsal yapısını çift yönlü olarak öğrenme kapasitesiyle öne çıkmaktadır (Devlin vd., 2019). Model, cümle içerisindeki sözcükleri hem önceki hem de sonraki bağlamla birlikte değerlendirerek, anlam ilişkilerini daha etkili biçimde yakalamaktadır. Bu yapı, dilin yapısal derinliğini

çözümlemede önemli bir avantaj sunar. Modelin ön eğitimi, BookCorpus ve İngilizce Wikipedia gibi kapsamlı metin koleksiyonları üzerinde gerçekleştirilmiştir. BERT dil modelinin çalışma yapısı Şekil 2’de verilmiştir.



Şekil 1. *BERT model yapısı* (Devlin vd., 2019)

2.1.10.1. XLM-RoBERTa. XLM-RoBERTa (Cross-lingual Language Model-RoBERTa), Facebook AI tarafından geliştirilmiş çok dilli bir dil modelidir. Transformer tabanlı bu model, BERT mimarisinin genişletilmiş ve optimize edilmiş bir versiyonudur. XLM-RoBERTa, farklı dillerdeki metinleri anlamak, işlemek ve birbirleriyle ilişkilendirmek amacıyla geliştirilmiş olup, 100’den fazla dili kapsayan geniş bir veri setiyle eğitilmiştir. Model, temel olarak BERT ‘in çift yönlü dikkat mekanizmasını kullanırken, veri setinin çeşitliliği ve büyüklüğü sayesinde diller arası anlam ilişkilerini daha başarılı bir şekilde öğrenir.

Bu modelin en önemli özelliği, metinlerin bağlamını iki yönlü olarak (hem soldan sağa hem de sağdan sola) değerlendirerek daha zengin ve doğru anlam temsilleri oluşturmasıdır. XLM-RoBERTa’nın eğitimi sırasında, MLM yöntemi kullanılır. Bu yöntemde, metinlerin belirli kısımları maskelenerek modelin bu boşlukları doğru şekilde doldurması beklenir. Bu, modelin bağlamı daha iyi anlamasına olanak tanır. XLM-RoBERTa, dil çevirisi, duygu analizi, metin sınıflandırması ve adlandırılmış varlık tanıma gibi çok çeşitli doğal dil işleme görevlerinde başarılı sonuçlar elde etmektedir. Aynı zamanda, dil kaynakları kısıtlı olan diller için de yüksek performans gösterir, çünkü diller arası transfer öğrenmeyi etkili bir şekilde kullanır.

Model, XLM ve RoBERTa'nın güçlü yönlerini birleştirerek çok dilli ortamda daha yüksek doğruluk oranlarına ulaşmayı hedefler. Bu sayede, düşük kaynaklı dillerde de yüksek kaliteli dil temsilleri sağlamak mümkün hale gelmiştir. XLM-RoBERTa, özellikle düşük kaynaklı dillerdeki performansı ile dikkat çekerken, aynı zamanda yüksek kaynaklı dillerde de rekabetçi sonuçlar elde etmektedir.

2.1.10.2. *TurkishBERT*. TurkishBERT, Türkçe metinler üzerinde özelleşmiş bir BERT modelidir. Türkçe, eklemeli ve karmaşık yapılı bir dil olduğu için, genel çok dilli modeller bu dilin özgün yapısına tam olarak uyum sağlayamayabilir. Bu nedenle TurkishBERT, Türkçe'nin morfolojik ve sentaktik özelliklerini daha iyi yakalayabilmek adına özel olarak eğitilmiştir. Bu model, Türkçedeki eklemeli dil yapısını ve kelime türetme kurallarını daha iyi anlamak amacıyla geniş ve çeşitli Türkçe veri kümeleri üzerinde önceden eğitilmiştir.

Modelin eğitim sürecinde, genellikle Wikipedia makaleleri, haber metinleri, forum yazıları ve sosyal medya içerikleri gibi çeşitli Türkçe kaynaklar kullanılmıştır. Bu geniş veri tabanı, TurkishBERT'in dilin farklı tonlarını, karmaşık dil yapısını ve bağlamsal anlamları daha iyi öğrenmesine olanak tanır. TurkishBERT'in eğitimi sırasında MLM ve NSP gibi BERT'in temel görevleri kullanılmıştır. Bu teknikler, modelin bağlamsal anlamları ve dil yapısını daha derinlemesine kavramasına yardımcı olur. TurkishBERT, duygu analizi, adlandırılmış varlık tanıma, metin sınıflandırması, soru-cevap sistemleri ve bilgi çıkarımı gibi birçok doğal dil işleme görevinde başarılı sonuçlar elde etmiştir. Aynı zamanda, Türkçenin zengin morfolojik yapısı nedeniyle kelime içi anlam ilişkilerini yakalamada da güçlü bir performans sergiler. Bu model, özellikle Türkçe veriyle çalışan araştırmacılar ve uygulama geliştiriciler için önemli bir araç olarak öne çıkmaktadır.

2.1.10.3. *Bert base multilingual*. Bert Base Multilingual, BERT modelinin çok dilli bir versiyonu olup, 100'den fazla dilde geniş kapsamlı metinlerle eğitilmiştir. Google tarafından geliştirilen bu model, çok dilli veri üzerinde eğitim alarak farklı dillerde benzer anlamları ve dilsel kalıpları öğrenme kapasitesine sahiptir. Bert Base Multilingual, özellikle düşük kaynaklı dillerde etkili bir performans sergileyebilmek için tasarlanmıştır ve diller arası transfer öğrenme yetenekleriyle dikkat çeker.

Modelin eğitiminde, çeşitli dillerdeki Wikipedia makaleleri ve diğer geniş çaplı veri kümeleri kullanılmıştır. Bu çeşitlilik, modelin farklı dil gruplarındaki ortak kalıpları öğrenmesine ve diller arası bağlamı daha iyi kavramasına olanak tanır. Bert Base Multilingual, MLM ve NSP görevleriyle eğitilmiş olup, metinleri bağlamsal olarak daha derin ve doğru şekilde anlamak üzere optimize edilmiştir.

Bert Base Multilingual, çok dilli metin sınıflandırması, adlandırılmış varlık tanıma, metin tamamlama ve duygu analizi gibi birçok doğal dil işleme görevinde yüksek başarı oranlarına ulaşmıştır. Özellikle düşük kaynaklı dillerde dil modelleri oluşturmak ve diller arası bilgi transferini geliştirmek amacıyla sıkça kullanılmaktadır. Modelin bu esnek yapısı, onu çok dilli projelerde ve küresel doğal dil işleme uygulamalarında önemli bir araç haline getirmiştir.

2.1.11. Üretken yapay zekâ. Üretken yapay zekâ, bilgisayar sistemlerinin insan benzeri yaratıcı süreçleri taklit ederek özgün içerikler üretebilmesini sağlayan bir yapay zekâ alanıdır. Bu teknoloji, metin, görüntü, ses ve video gibi farklı medya türlerinde yeni ve anlamlı veriler oluşturma kapasitesine sahiptir. Üretken yapay zekânın temelini, özellikle son yıllarda büyük veri kümelerinin kullanılabilirliğinin artması ve derin öğrenme tekniklerinin gelişmesi oluşturur. 2017 yılında Google tarafından geliştirilen Transformer mimarisi, bu alandaki ilerlemelerin önünü açmış; sonrasında GPT (Generative Pre-trained Transformer) ve BERT gibi modeller aracılığıyla üretken yapay zekâ sistemleri daha erişilebilir ve etkili hale gelmiştir.

Üretken yapay zekâ, yalnızca mevcut bilgiyi işlemenin ötesine geçerek, yeni fikirler ve yaratıcı içerikler geliştirebilme kapasitesiyle geleneksel yapay zekâdan ayrılmaktadır. Bu yönüyle dijital içerik üretiminden eğitim materyallerinin oluşturulmasına, yazılı metinlerin iyileştirilmesinden görsel tasarım ve müzik üretimine kadar geniş bir kullanım alanına sahiptir. Ayrıca kişiselleştirme, otomasyon ve üretkenliği artırma gibi avantajlar sunarak bireysel ve kurumsal düzeyde yaratıcı süreçleri yeniden şekillendirmektedir.

Günümüzde üretken yapay zekâ, özellikle doğal dil işleme temelli sistemlerde ön plana çıkmakta; insan benzeri dil üretimi, içerik oluşturma, soru-cevap sistemleri ve özetleme gibi görevlerde yaygın olarak kullanılmaktadır. Bu bağlamda, üretken yapay

zekâ teknolojileri, sadece teknik bir yenilik değil, aynı zamanda bilgi üretimi ve aktarımı süreçlerinde dönüşümsel bir rol üstlenmektedir.

2.1.12. Üretken yapay zekâ araçları. Üretken yapay zekâ araçları, içerik oluşturma, veri analizi ve yaratıcı süreçleri otomatikleştirmek amacıyla kullanılan ileri teknoloji çözümleridir. Bu araçlar, metin, görüntü, ses ve video gibi çok çeşitli medya formatlarında içerik üretebilir. Genellikle büyük dil modelleri (LLM'ler) ve derin öğrenme algoritmalarıyla desteklenirler. Üretken yapay zekâ kavramı, derin öğrenme alanındaki gelişmeler ve büyük veri kümelerinin artışıyla 2010'lardan itibaren hızla popülerlik kazanmıştır. Özellikle Transformer tabanlı modellerin (ör. GPT, BERT) geliştirilmesi, üretken yapay zekâ araçlarının performansını önemli ölçüde artırmıştır. Bu araçlar, dijital içerik üretiminde verimliliği artırmak, yaratıcı süreçleri hızlandırmak ve daha kişiselleştirilmiş içerikler oluşturmak için yaygın olarak kullanılır. Öne çıkan üretken yapay zekâ araçları arasında ChatGPT, Smodin ve Gemini bulunmaktadır.

2.1.13. ChatGPT. ChatGPT, OpenAI tarafından geliştirilen büyük bir dil modelidir. İnsan dilini anlamak ve doğal dilde yanıtlar üretmek amacıyla tasarlanmıştır. ChatGPT'nin kökeni, 2017 yılında Google tarafından geliştirilen ve doğal dil işleme alanında devrim yaratan Transformer mimarisine dayanmaktadır. İlk olarak 2018'de GPT adıyla tanıtılan bu model, geniş veri kümeleriyle eğitilmiş ve metin tabanlı etkileşimlerde yüksek doğruluk ve tutarlılık sağlar. 2020 yılında tanıtılan GPT-3 modeli, 175 milyar parametreye sahip olmasıyla dikkat çekmiştir ve dil modelleme alanında önemli bir adım olmuştur. ChatGPT, müşteri hizmetlerinden içerik üretimine, eğitimden kişisel asistanlara kadar geniş bir kullanım yelpazesine sahiptir.

2.1.14. Smodin. Smodin, metin yeniden yazma, özetleme, içerik üretme ve dil çevirisi gibi çeşitli dil işleme görevlerini yerine getiren bir yapay zekâ platformudur. Platformun temel amacı, kullanıcıların metinlerini hızlı ve doğru bir şekilde düzenlemesine olanak tanımadır. Smodin'in tarihçesi, yapay zekâ destekli içerik

üretme ihtiyacının artmasıyla şekillenmiştir. Özellikle çok dilli içerik üretimi ve metin optimizasyonu alanlarında önemli bir araç olarak öne çıkmıştır. Smodin, içerik üreticileri, akademisyenler ve öğrenciler tarafından sıkça tercih edilmekte olup, dil bariyerlerini aşma ve içerik kalitesini artırma noktasında etkili çözümler sunar.

2.1.15. Gemini. Google Gemini, Google tarafından geliştirilen ve doğal dil işleme teknolojilerine dayalı olarak çalışan bir üretken yapay zekâ aracıdır. Gemini'nin kökeni, Google'ın uzun süredir geliştirdiği BERT modeline dayanır. 2018'de BERT'in tanıtılması, doğal dil işleme alanında büyük bir dönüm noktası olmuştur. Gemini, bu teknolojiyi daha ileriye taşıyarak kullanıcıların yaratıcı içerik oluşturmalarını, fikirlerini geliştirmesini ve karmaşık metinleri daha anlaşılır hale getirmesini sağlar. Gelişmiş metin analitiği ve bağlam anlama yetenekleriyle dikkat çeken gemini, içerik üretimi ve metin analizi alanlarında güçlü bir çözümdür. Ayrıca, büyük veri kümeleriyle eğitilmiş olması, metinlerin anlamını ve bağlamını daha iyi kavramasına olanak tanır.

2.2. Yurtiçi ve Yurtdışında Yapılmış İlgili Çalışmalar

Munika vd. (2019) tarafından gerçekleştirilen çalışmada, duygu sınıflandırması için transformer tabanlı derin öğrenme mimarisi kullanılmıştır. Çoğu doğal dil işleme modelinin genellikle ikili duygu sınıflandırması üzerine yoğunlaşırken, bu çalışma daha ayrıntılı bir duygu sınıflandırma görevini ele almıştır. Deneysel sonuçlar, BERT modelinin diğer yaygın kullanılan modellere kıyasla daha üstün bir performans sergilediğini ortaya koymuştur.

Türkçe metinler üzerinde yapılan duygu analizi çalışmaları, literatürde henüz sınırlı sayıda yer almaktadır. Bu alandaki uygulamalarda genellikle iki temel yaklaşım öne çıkmaktadır: ilki çok dilli BERT modellerinin doğrudan kullanılması, diğeri ise Türkçe içeriklerin İngilizce 'ye çevrilerek BERT 'in özgün (İngilizce) versiyonuyla analiz edilmesidir. Acikalin vd. (2020), Türk filmleri ile otellere ait kullanıcı yorumlarından oluşan veri kümeleri üzerinde gerçekleştirdikleri uygulamalarda, BERT tabanlı modellerin duygu sınıflandırmasında oldukça başarılı sonuçlar verdiğini ortaya koymuştur.

Öztürk vd. (2020), tarafından gerçekleştirilen çalışmada Türkçede yaygın olarak karşılaşılan ‘de/da’ ekine ilişkin yazım hatalarını incelemişlerdir. Morfolojik yapının karmaşıklığı nedeniyle bu tür hataların otomatik sistemler tarafından tespit edilmesi oldukça güçtür; bu durum, mevcut yazım denetleyicilerin yeterince etkili sonuçlar verememesine yol açmaktadır. Bu bağlamda, bağlamsal bilgiye dayalı kelime yerleştirmeleri üretebilen BERT modelinin potansiyeli değerlendirilmiştir. Araştırmada, BERT’in bağlamı dikkate alarak sunduğu tahminlerin, bağlamsız modellere kıyasla belirgin biçimde daha başarılı olduğu sonucuna ulaşılmıştır.

Sevli ve Kemaloğlu (2021) tarafından gerçekleştirilen çalışmada, olağandışı olaylar hakkındaki tweet’lerin gerçek ve gerçek dışı olarak Google BERT modeli ile sınıflandırılması başlıklı çalışmalarında, deprem, kaza, olumsuz hava olayları gibi felaket durumları hakkında atılan 7613 adet gerçek veya gerçek dışı olarak etiketlenmiş tweet içeren veri seti Bert modeli kullanılarak sınıflanmıştır. Sinir ağı temelli bir model olan BERT ile sınıflanmıştır. %80 eğitim, %20 test seti olarak ayrılan veri kümesi üzerinde 10 epoch olarak gerçekleştirilen eğitim işlemi sonucunda %98,88 doğruluk başarıları elde edilmiştir.

Ozan ve Taşar (2021) tarafından gerçekleştirilen çalışmada, belirli bir alana ait kısa konuşma cümlelerinin otomatik olarak etiketlenmesini sağlayacak bir yöntem geliştirilmesi hedeflenmiştir. Bu amaçla, müşteri temsilcileri ile web sitesi ziyaretçileri arasında gerçekleşen diyaloglardan elde edilen yaklaşık 14 bin örnekten oluşan veri kümesi, on temel kategoriye ayrılarak elle etiketlenmiştir. Elde edilen veri, anlamlı ve işlevsel yanıtlar üretebilecek bir sohbet robotunun oluşturulması amacıyla, transformatör mimarisi temelli dil modelleriyle analiz edilmiştir. Yapılan karşılaştırmalar sonucunda, en yüksek sınıflandırma başarısının BERT modeli aracılığıyla elde edildiği rapor edilmiştir.

Ayan vd. (2021) tarafından gerçekleştirilen çalışmada, Türkiye'deki en büyük 25 şirketin web sitelerinden elde edilen veriler üzerinden anahtar kelimelerin çıkarılması süreci ele alınmıştır. Çalışma hem metodolojik hem de sonuçsal açıdan doğal dil işleme alanına yeni bir perspektif kazandırmayı hedeflemiştir. Yapılan analizler, şirketlerin kullandığı anahtar kelimelerin kümelenmesi yoluyla iş alanlarına dair önemli çıkarımlar sunmuş ve bu sayede sektörel eğilimler hakkında değerli bilgiler ortaya konmuştur.

Ozan ve Taşar (2021) tarafından gerçekleştirilen çalışmada, müşteri temsilcileri ile müşteriler arasındaki yazışmalar analiz edilerek, kısa konuşma cümlelerinin anlamsal özelliklerine dayalı bir inceleme yapılmıştır. Araştırmada, 46 farklı kategoriye ayrılmış bir veri seti kullanılarak, GPT-2 ve BERT modelleri ile sınıflandırma işlemi gerçekleştirilmiştir. Bu analiz, müşteri hizmetleri süreçlerinde kullanılan dilin daha iyi anlaşılmasını ve otomatik sınıflandırma tekniklerinin etkinliğini değerlendirmeyi amaçlamaktadır.

İşeri vd. (2021) tarafından gerçekleştirilen çalışmada, çağrı merkezi işlemlerini otomatikleştirmek amacıyla makine öğrenmesi ve doğal dil işleme tekniklerinden yararlanarak bir sohbet robotu geliştirilmesi hedeflenmiştir. Araştırmada kullanılan veri seti, sıkça sorulan soruların yanı sıra telefon görüşmeleri ve anlık mesajlaşma uygulamalarından elde edilen verileri içermektedir. Geliştirilen sohbet robotu, kullanıcıların firma veya ürünlerle ilgili sorularına hızlı ve etkili bir şekilde yanıt verebilmekte, kullanıcı girdileri doğrultusunda model, doğru niyet ve yanıtları üretmektedir. Yapılan testler sonucunda, sohbet robotunun yüksek doğruluk oranlarıyla soruları cevapladığı gözlemlenmiştir.

İnternetin ve sosyal medya platformlarının yaygınlaşması, bilginin üretilme ve dolaşıma girme biçimlerini köklü bir şekilde dönüştürmüştür. Bu dijital ortamda oluşan veri yığını anlamlandırmak amacıyla, makinelerin insan dilini anlamasını sağlayan yöntemler giderek daha fazla önem kazanmıştır. Özellikle son yıllarda, doğal dil işleme alanında önceden eğitilmiş modellerin sergilediği yüksek performans dikkat çekmektedir. Hark (2022) tarafından gerçekleştirilen çalışmada, sahte haberleri tespit etmek amacıyla kullanılan bu tür modeller, %91'e ulaşan doğruluk oranlarıyla etkili sonuçlar vermiştir.

Toğaçar vd. (2022) tarafından gerçekleştirilen çalışmada, yapay zekâ tabanlı doğal dil işleme yaklaşımını kullanarak internet ortamında yayınlanmış sahte haberlerin tespiti başlıklı çalışmalarında veri seti, 6335 haber başlığı ve içerikten oluşmaktadır. Bu haberlerin 3171'i gerçek haber niteliği taşırken; 3164'ü ise sahte haber niteliği taşımaktadır. Doğal dil işleme yöntemi ile birlikte LSTM modeli kullanıldı ve veri setinin eğitimi bu model sayesinde gerçekleştirilmiştir. Sonuç olarak, bu çalışmada eğitim verilerinden elde edilen genel doğruluk başarısı %99,83 iken ve test verilerinden %91,48 doğruluk elde edilmiştir.

İnternetin ve sosyal medya kullanımının giderek yaygınlaşması, dijital ortamda gerçekleştirilen zorbalık vakalarının da önemli ölçüde artmasına zemin hazırlamıştır. Bu artış, dünya çapında çok sayıda bireyin siber zorbalığa maruz kalmasına ve bunun sonucunda çeşitli psikolojik ve sosyal sorunların ortaya çıkmasına neden olmaktadır. Yazgılı ve Baykara (2022), sosyal medyadan derlenen 3000 örnek içeren Türkçe bir veri seti üzerinden yürüttükleri çalışmada, siber zorbalığın otomatik olarak tespit edilmesini amaçlamışlardır. Araştırmalarında, literatürde sıkça başvurulmuş sınıflandırma tekniklerine ek olarak, Bagging ve Boosting gibi topluluk öğrenme yöntemlerini karşılaştırmalı olarak değerlendirerek bu yöntemlerin performansını analiz etmişlerdir.

Bayram (2022) gerçekleştirdiği çalışmada, COVID-19 pandemisine ilişkin farklı boyutları ele aldığı çalışmada, sosyal medya içerikleri ile dijital haber arşivlerini karşılaştırmalı olarak analiz etmiştir. Sosyal medya verilerinin bireylerin duygusal ve tepkisel ifadelerini yansıttığı, haber kaynaklarının ise daha çok nesnel bilgi sunduğu gözlemlenmiştir. Araştırmada, pandeminin toplumsal etkileri çok yönlü olarak değerlendirilmiş; istatistiksel analizler yoluyla öne çıkan sorun alanları tespit edilmiştir. Ayrıca, makine öğrenmesi yöntemleri kullanılarak pandeminin yarattığı problemler derinlemesine incelenmiştir. Elde edilen bulgular, ekonomik sorunların en fazla dikkat çeken başlıklardan biri olduğunu ortaya koyarken; sözcüksel ağlar aracılığıyla bu konuya bağlı diğer temalar da haritalandırılmıştır. Çalışma, pandeminin etkilerini kapsamlı şekilde anlayabilmek adına doğal dil işleme ve makine öğrenmesine dayalı daha fazla araştırmanın gerekliliğini vurgulamaktadır.

Hark (2022) tarafından yapılan çalışmada, sahte haber tespiti için derin bağlamsal kelime gömülmeleri ve sinirsel ağların performans değerlendirildiği deneysel çalışmada, açık erişimli COVID-19 isimli sahte haber tespit veri seti (10700) kullanılmış, başarımları çeşitli performans metrikleri ile hesaplanmıştır. En yüksek başarımlar %91 ile LSTM tarafından rapor edilmiştir. Ön-eğitilmiş kelime gömülmelerinin farklı sinirsel ağlardan bağımsız olarak yüksek bir hassasiyetle sahte haberlerin tespitinde kullanılabileceğini gösteren umut verici sonuçlar sunulmuştur.

Özkan ve Kar (2022) yapılan çalışmada Google tarafından Ekim 2018’de geliştirilen BERT, makine öğrenimi ve doğal dil işleme alanlarında devrim niteliğinde bir model olarak kabul edilmiştir. Bu yapı, transformatör mimarisi temelinde geliştirilen ve dilin

bağlamını her iki yönden değerlendirebilen bir yaklaşıma sahiptir. Türkçenin yapısal özellikleri — eklemeli dil yapısı ve zengin biçim bilgisel yapısı — çoklu sınıflandırma sorunlarında önemli zorluklar yaratmaktadır. Bu kapsamda gerçekleştirilen çalışmada, son on yıla ait Türkçe bilimsel ve akademik içerikler içeren bir veri kümesiyle çok sınıflı bir metin sınıflandırma süreci yürütülmüştür. Hazır Türkçe BERT modeli, bu veriler üzerinde yeniden eğitilerek optimize edilmiştir. Uygulanan deneylerde modelin doğruluk oranı %96 seviyesine ulaşmış; böylece BERT mimarisinin Türkçe metinler üzerinde sınıflandırma görevlerinde etkili ve başarılı bir performans sunduğu gözlemlenmiştir. Bulgular, yöntemin dilsel karmaşıklığa karşı güçlü bir çözüm sunduğunu göstermektedir.

Yapay zekâ ve makine öğrenmesi, günümüzde öne çıkan ileri teknolojiler arasında yer almakta olup, bilgisayar sistemlerinin insan benzeri düşünme ve öğrenme yetkinlikleri kazanmasını hedeflemektedir. Bu teknolojiler, çeşitli sektörlerde verimlilik artışı sağlamak ve günlük yaşamı önemli ölçüde kolaylaştırmaktadır. Bununla birlikte, bu gelişmeler sadece olumlu sonuçlar doğurmamakta; etik, sosyal ve ekonomik açılardan bazı tartışmaları da beraberinde getirmektedir. Yapay zekânın uzun vadeli etkileri konusunda akademik camiada fikir birliği bulunmamakta, farklı bakış açıları ortaya konulmaktadır. Altıntop (2023), ChatGPT gibi yaygın olarak kullanılan yapay zekâ sistemleri üzerinden bu teknolojilerin sunduğu olanakları ve beraberinde getirdiği sorunsalları irdeleyerek çeşitli görüşleri analiz etmiştir.

Yapay zekâ teknolojilerinin, özellikle ChatGPT gibi dil tabanlı sistemlerin, iletişim süreçleri üzerindeki etkisi giderek daha fazla dikkat çekmekte ve alanda yeni tartışmaların önünü açmaktadır. Bu tür uygulamaların yaygınlaşmasıyla birlikte, iletişim biçimlerinde köklü dönüşümlerin yaşanması öngörülmektedir. (Koçyiğit ve Darı (2023), yapay zekâ ile iletişim arasındaki etkileşimi ChatGPT örneği üzerinden ele alarak, bu teknolojilerin iletişim alanına getirdiği yenilikleri, potansiyel faydaları ve ortaya çıkardığı sorunları değerlendirmiştir. Çalışmalarında, içinde bulunduğumuz dönemi "İnsanlaşan Dijitalleşme Çağı" olarak nitelendirerek, teknolojik gelişmelerin iletişim üzerindeki etkilerine farklı bakış açısı sunmayı hedeflemişlerdir.

Yiğit vd. (2023) tarafından yürüttüğü araştırmada, doğal dil işleme tabanlı bir yapay zekâ aracı olan ChatGPT'nin sağlık sektöründeki potansiyel kullanımları detaylı biçimde ele alınmıştır. Çalışmada, bu teknolojinin sağlık profesyonellerine yönelik

akademik yayın hazırlama, eğitim içerikleri oluşturma ve çeşitli sağlık hizmetlerinin sunumunda yardımcı olabileceği vurgulanmıştır. Ayrıca, bireye özgü tedavi süreçlerinin desteklenmesi, toplumun sağlık bilgilerine erişiminin artırılması ve sağlık okuryazarlığının geliştirilmesi gibi katkılar da öne çıkarılmıştır. Bununla birlikte, ChatGPT'nin kullanımında etik ve hukuki açıdan bazı önemli sorunlar gündeme gelmektedir. Özellikle veri gizliliği, hasta bilgilerinin korunması gibi konular dikkatle ele alınmalı; bu sürecin güvenli ve etkili bir şekilde yürütülebilmesi için teknoloji geliştiricileri ile sağlık hizmeti sağlayıcıları arasında güçlü bir iş birliği tesis edilmelidir. Araştırma, sistemin daha güvenilir ve verimli hale gelmesi için ileri düzey veri setlerine ve sürekli geliştirmelere ihtiyaç duyulduğuna da işaret etmektedir.

Sağlık hizmeti tüketicileri, hasta dostu belgelerde yer almayan bilgileri genellikle tıbbi literatürden edinmeye çalışmaktadırlar. Ancak, tıbbi makaleler genellikle karmaşık terimler içerdiği için bu kaynakları anlamak zorlayıcı olabilmektedir. Bu sorunu çözmek amacıyla, August vd. (2023) bir prototip sistem geliştirmiştir. Sistem, kullanıcıların araştırma makalelerini daha kolay bir şekilde okuyup anlamalarını sağlamış, bu kullanıcılar, sistemi kullanmayanlara kıyasla daha başarılı bir şekilde metinlere hâkim olabilmişlerdir. Çalışma, önerilen yaklaşımın tıbbi makalelerin anlaşılmasını kolaylaştırdığını ve okuyucularda güven duygusu oluşturduğunu göstermektedir.

Müller vd. (2023) tarafından gerçekleştirilen çalışmada, COVID-19 ile ilgili Twitter mesajlarından oluşan kapsamlı bir veri kümesi kullanılarak geliştirilen CT-BERT modeli tanıtılmıştır. Bu model, COVID-19'a özgü sosyal medya verilerini işlemek amacıyla özel olarak optimize edilmiştir ve sınıflandırma, soru cevaplama gibi farklı doğal dil işleme (NLP) görevlerinde etkinlik gösterebilmektedir. Araştırma, alana özgü önceden eğitilmiş dil modellerinin, pandemi döneminde sosyal medyadaki bilgi akışını anlamak ve analiz etmek için ne denli önemli olduğunu ortaya koymaktadır. Bu bağlamda, CT-BERT'in geliştirilmesi, COVID-19 konulu metinlerin daha doğru ve hızlı bir şekilde analiz edilmesine yönelik önemli bir katkı sağlamaktadır.

Amini vd. (2023) araştırmasında, nöropsikolojik testlerin ses kayıtlarından otomatik olarak metne dönüştürülmesine dayalı bir doğal dil işleme modeli geliştirilmiştir. Bu model, demansın farklı evrelerinin tespiti için önemli bir adım olarak değerlendirilmektedir. Araştırmada kullanılan yöntem, transkripsiyon sürecinin

doğruluğunu artırarak klinik değerlendirmelerde zaman ve maliyet tasarrufu sağlamayı amaçlamaktadır. Çalışmanın sonuçları, modelin %92'nin üzerinde bir performans sergilediğini ortaya koyarak, bu tür yenilikçi yaklaşımların demans taramalarında etkili bir çözüm sunduğunu ortaya koymaktadır.

Çetindağ vd. (2023) tarafından gerçekleştirilen araştırması, hukuki metinlerde adlandırılmış varlık tanıma (NER) teknolojilerinin kullanımına odaklanmaktadır. NER, hukuki belgelerde yer alan kişi, kurum, yer ve diğer önemli varlıkların doğru bir şekilde tanımlanması açısından kritik bir öneme sahiptir. Ancak, hukuk alanına özgü yüksek doğruluklu NER modellerinin sınırlı olması, bu alanda önemli bir boşluk yaratmaktadır. Bu eksikliği gidermek amacıyla, çalışma Türkçe hukuki metinlere özel olarak tasarlanmış yeni bir NER modeli geliştirmiştir. Model, farklı NER mimarileri ve kelime yerleştirme stratejilerinin kapsamlı bir karşılaştırmasını içermektedir. Sonuçlar, Türkçe gibi sondan eklemeli dillerde hukuki metinlerin otomatik işlenmesi için önemli bir adım olarak değerlendirilmektedir.

Jiang vd. (2023) tarafından gerçekleştirilen çalışmada, insan hareketini dilsel yapılarla entegre eden MotionGPT adlı yeni bir model geliştirilmiştir. Bu model, hareket verilerini metinlerle anlamlı bir şekilde birleştirerek hem hareket analizini hem de doğal dil modellemeyi tek bir çerçevede birleştirmektedir. MotionGPT, insan hareketlerini ayırık vektör kuantizasyonu yöntemiyle kodlayarak 3 boyutlu hareket verilerini belirteçler (tokenlar) şeklinde temsil eder. Bu yaklaşım, modelin hareket oluşturma, hareket açıklaması ekleme, hareket tahmini gibi çeşitli görevlerde üstün performans sergilemesini sağlamaktadır. Deney sonuçları, MotionGPT'nin farklı hareket görevlerinde yüksek doğruluk ve esneklik sunduğunu ortaya koymuştur.

Zong ve Krishnamachari (2023) tarafından gerçekleştirilen çalışmada, öğrencilerin matematik problemlerini çözme süreçlerinde güçlü dönüştürücü tabanlı modellerin, özellikle GPT-3'ün, nasıl kullanılabileceğini incelemektedir. Araştırma, matematiksel sözlü problemler için üç temel zorluk düzeyini ele almıştır: problem sınıflandırma, denklemlerin çıkarılması ve yeni problem oluşturma. Yapılan deneyler, GPT-3'ün bu üç aşamada da yüksek doğruluk oranları sergilediğini ortaya koymuştur. Çalışma, bu tür yapay zekâ tabanlı araçların, öğrencilerin matematiksel düşünme becerilerini geliştirmelerine yardımcı olma potansiyelini vurgulayarak, eğitim teknolojilerinde önemli bir adım olarak değerlendirilmektedir.

Duygu analizi, doğal dil işleme alanında giderek artan bir öneme sahip araştırma konularından biridir. Bu analiz türü, özellikle e-ticaret kullanıcı yorumları, siyasal eğilimlerin belirlenmesi ve kamu yönetimi gibi çeşitli alanlarda yaygın olarak kullanılmaktadır. Başarılı bir duygu analizi gerçekleştirebilmek için etkili metin temsil yöntemlerine ihtiyaç duyulmaktadır. Bu yöntemler genel olarak sözlük tabanlı ve makine öğrenimi tabanlı yaklaşımlar şeklinde iki grupta sınıflandırılmaktadır. Ancak her iki yöntemin de bazı kısıtları bulunmaktadır. Örneğin, Word2Vec gibi popüler yöntemler kelimelerin bağlam içerisindeki duygusal tonlarını yeterince yansıtamayabilir. Mutinda vd. (2023) tarafından geliştirilen LeBERT adlı model, bu eksiklikleri gidermeyi amaçlayan hibrit bir yaklaşım sunmaktadır. Modelde, duygu analizine yönelik sözlük tabanlı bilgi ve BERT tabanlı bağlamsal kelime temsilleri bir arada kullanılmakta, sınıflandırma işlemi ise bir evrişimsel sinir ağı (CNN) aracılığıyla gerçekleştirilmektedir. LeBERT modeli, çeşitli standart veri kümeleri üzerinde yapılan değerlendirmelerde mevcut yöntemlerin çoğuna kıyasla daha yüksek doğruluk oranları sergilemiştir.

Bird vd. (2023) tarafından gerçekleştirilen çalışmada insan-merkezli etkileşimleri temel alan bir transformatör tabanlı sohbet robotu mimarisine yönelik yeni bir eğitim çerçevesi geliştirmiştir. "Yapay zekâ destekli sohbet etkileşimi çerçevesi" olarak adlandırılan bu yapı, özellikle teknik bilgiye sahip olmayan kullanıcıların yapay zekâ sistemleriyle daha doğal ve erişilebilir bir biçimde etkileşim kurmasını amaçlamaktadır. Söz konusu model, insan komutlarını yalnızca dilsel açıdan değil, aynı zamanda sosyal bağlamda da anlamlandırma kapasitesi sunarak, daha gerçekçi ve anlamlı insan-makine iletişimini mümkün kılmaktadır. Elde edilen bulgular, geliştirilen sistemin yüksek performanslı bir etkileşim altyapısı sunduğunu ortaya koymaktadır.

Geleneksel haber öneri sistemleri, görev hedefi ile model yapısı arasındaki uyumsuzluklar nedeniyle çoğu zaman sınırlı başarı göstermektedir. Bu bağlamda, son dönemde dikkat çeken anlık öğrenme (prompt-based learning) yaklaşımı, doğal dil işleme alanında kayda değer ilerlemelere olanak tanımaktadır. Zhang ve Wang (2023) tarafından geliştirilen Prompt4NR isimli çerçeve, haber öneri sistemlerinde bu yeni öğrenme paradigmasının uygulanmasına yönelik özgün bir model sunmaktadır. Modelde, kullanıcının belirli bir habere tıklama eğilimini tahmin etmek amacıyla

maske tahminine dayalı bir görev tanımı yapılmış, ayrıca ayırık, sürekli ve hibrit biçimlerde çeşitli bilgi istemi (prompt) şablonları geliştirilmiştir. Bu şablonlardan elde edilen çıktıların bütüncül şekilde değerlendirilmesi için birleştirici bir bilgi istemi mekanizması tasarlanmıştır. MIND veri kümesi üzerinde gerçekleştirilen deneysel analizler, Prompt4NR yaklaşımının haber öneri performansında anlamlı bir iyileşme sağladığını ortaya koymaktadır.

Kanan vd. (2023), tarafından gerçekleştirilen çalışmada çevrimiçi işletmelere ait sosyal medya platformlarından elde edilen büyük ölçekli veri kümeleri aracılığıyla müşteri memnuniyetini belirlemeye yönelik bir araştırma gerçekleştirmiştir. Bu kapsamda, makine öğrenmesi ve derin öğrenme yöntemleri kullanılarak çeşitli analizler yapılmış; elde edilen sonuçlar, özellikle derin sinir ağlarının müşteri duyarlılığını doğru biçimde modelleme ve sınıflandırma konusundaki üstün performansını ortaya koymuştur. Çalışma, sosyal medya verilerinin işletmeler açısından stratejik değere sahip içgörülere dönüştürülmesinde yapay zekâ tabanlı yaklaşımların etkinliğini gözler önüne sermektedir.

Saleh vd. (2023), tarafından gerçekleştirilen çalışmada çevrimiçi ortamlarda yaygınlaşan nefret söyleminin otomatik olarak tespitine yönelik derin öğrenme temelli yaklaşımlar üzerine bir çalışma gerçekleştirmiştir. Bu kapsamda, LSTM tabanlı modellerin yanı sıra BERT gibi önceden eğitilmiş dil modelleri kullanılarak ikili sınıflandırma çerçevesinde analizler yürütülmüştür. Elde edilen bulgular, bu tür modellerin nefret söylemi tespitinde yüksek doğruluk oranlarına ulaştığını ve özellikle modelin eğitim verisinin niteliği ile boyutunun performans üzerinde belirleyici etkiler yarattığını göstermektedir. Ayrıca, alan odaklı olarak eğitilen modellerin genel amaçlı modellere kıyasla daha başarılı sonuçlar verdiği vurgulanmaktadır. Bu çalışma, sosyal medya üzerindeki zararlı içeriklerin tespitine yönelik etkili ve uygulanabilir yöntemler sunarak literatüre değerli bir katkı sağlamaktadır.

Salvagno vd. (2023) tarafından yapılan çalışmada, Yapay zekâ Chatbot'unun bilimsel yazımda kullanımını tartışmışlardır. Yapay zekâ sohbet robotu ve ChatGPT özellikle bilimsel yazımda yararlı araçlar olarak görüldüğü ve araştırmacılara ve bilim insanlarına materyali organize etmede, ilk taslağı oluşturmada ve/veya redaksiyonda yardımcı olduğu vurgulanmıştır. ChatGPT çalışması, insan yargısının yerine kullanılmaması ve çıktı, herhangi bir kritik karar alma veya uygulamada

kullanılmadan önce her zaman uzmanlar tarafından incelenmesi gerektiğinden bahsedilmiştir. Ayrıca, bu araçların kullanımıyla ilgili intihal ve yanlışlıklar riski ve ayrıca yazılımın ödeme yapması durumunda yüksek ve düşük gelirli ülkeler arasında erişilebilirliğinde potansiyel bir dengesizlik gibi çeşitli etik sorunlar ortaya çıkabileceğine değinilmiştir. Bu nedenle, araştırmacılar bilimsel yazımda chatbot kullanımının nasıl düzenleneceği konusunda yakın zamanda bir fikir birliğine ihtiyaç duyulacağı sonucuna varmışlardır.

Kim (2023) araştırmasında bilimsel makalelerde dil düzenleme için ChatGPT'yi kullanma başlıklı çalışmada, ChatGPT sadece dil düzenleme amacıyla kullanılmasının yanısıra, bilimsel makaleler hazırlamak için kullanılmasında bir sakınca olmadığını bildirmiştir. Ancak, ChatGPT tarafından üretilen herhangi bir yeni fikir, gerçek deneylerle ve sonuçları insanlar tarafından doğrulanmalıdır. Yapay zekâ tabanlı araçların yanlış bilgi üretme potansiyeline sahip olduğu için yazarların makalelerine eklemekten önce ChatGPT tarafından oluşturulan bilgileri dikkatlice gözden geçirmeleri ve doğrulamaları önemli olduğu çalışmada vurgulanmıştır.

Gao vd. (2023) tarafından yapılan çalışmada ChatGPT tarafından oluşturulan bilimsel özetleri, algılayıcılar ve gözleri kör edilmiş insan hakemlerle gerçek özetlerle karşılaştırma amaçlı yaptıkları çalışmada, ChatGPT gibi büyük dil modelleri, bu modelleri bilimsel yazımda kullanmanın doğruluğu ve bütünlüğü hakkında bilinmeyen bilgiler içeren, giderek daha gerçekçi metinler üretebileceğinden ve özet oluşturma, oluşturulan verileri içermesine rağmen okunabilir bilimsel özetler oluşturmak için güçlü bir araç olduğundan bahsetmişlerdir. Çalışmalarında, 50 bilimsel tıbbi makale için ChatGPT tarafından oluşturulan özetleri değerlendirmişlerdir. Oluşturulan özetler, bir intihal dedektörü web sitesi ve iThenticate aracılığıyla çalıştırıldığında orijinal özetlerden daha düşük puan aldığı sonucunu elde etmişlerdir. Orijinal ve genel özetlerin bir karışımı verildiğinde, özetlerin %68'inin ChatGPT tarafından oluşturulduğunu doğru bir şekilde elde edilmiştir. Bu çalışmada hem insanların hem de yapay zekâ çıktı algılayıcılarının ChatGPT tarafından oluşturulan özetlerin bir kısmını tanımlayabildiğini ancak ikisinin de mükemmel ayırmacı olmadığını tespit etmişlerdir.

Uzun (2023) tarafından yapılan çalışmada yapay zekânın ürettiği içerikleri tespit etmeyi amaçladığı çalışmada, yapay zekâ tarafından oluşturulan içeriği tespit etmek

için birkaç yöntem olsa da hiçbirinin kusursuz olmadığından bahsetmiş ve yapay zekâ tarafından üretilen içeriği doğru bir şekilde tanımlamak için tekniklerin ve araçların bir kombinasyonu gerekli olabileceğini söylemiştir. Stilometri, meta veri analizi ve CopyLeaks ve Turnitin gibi çevrimiçi araçlar, yapay zekâ tarafından üretilen içeriği tespit etmek için kullanılan etkili araçlardan ve tekniklerden bazılarıdır. Bununla birlikte, bu yöntemlerin sınırlamaları da kabul edilmeli yapay zekâ tarafından üretilen içeriği tespit etmeye yönelik etkili ve etik yöntemler geliştirmek için daha fazla araştırmaya ihtiyaç olduğu bu çalışmada vurgulanmıştır.

Sosyal medya platformları, bireylerin çeşitli konulara ilişkin düşüncelerini içeren büyük hacimli veri sağlamaktadır. Bu nedenle, pazar araştırmaları ve kamuoyu analizleri gibi alanlarda değerli bilgi kaynaklarıdır. Ancak sosyal medya kullanıcıları toplumun tamamını temsil etmediği için, bu verilerdeki önyargıyı azaltmak amacıyla demografik bilgiler dikkate alınarak ek işlemler uygulanmalıdır. Bu çalışmada, Türkçe Twitter kullanıcılarının gönderileri üzerinden yaş aralığı ve cinsiyet tahmini yapılmıştır. 1040 kullanıcıdan oluşan etiketli bir veri seti oluşturulmuş ve kelime, karakter, retweet, fastText ve BERT temelli beş yöntem geliştirilmiştir. Deneyler, kullanıcı içeriklerinin demografik tahminler için anlamlı göstergeler sunduğunu ortaya koymuştur (Görgün vd., 2023).

Mujahid vd. (2023) tarafından gerçekleştirilen çalışmada, ChatGPT'ye yönelik kullanıcı geri bildirimlerini analiz etmek amacıyla Arapça dilinde atılmış tweet'ler üzerinde otomatik duygu analizi gerçekleştirilmektedir. Özellikle Arapça içeriklerin analizinde sınırlı çalışma bulunması, bu araştırmayı önemli kılmaktadır. Twitter'dan SNscrape aracıyla 27.780 adet yapılandırılmamış tweet toplanmış ve TextBlob Arapça kütüphanesi kullanılarak bu tweet'ler olumlu, olumsuz ve nötr olarak etiketlenmiştir. Ham veriler, gereksiz karakterler, durdurma kelimeleri ve alakasız ögelerden arındırılarak işlenmiş, ardından analiz için hazır hâle getirilmiştir. Analiz sürecinde RoBERTa, XLNet ve DistilBERT gibi güçlü transformatör tabanlı modeller kullanılmış; bunun yanı sıra, RoBERTa ve BERT'in çıktılarını birleştiren özel bir hibrit model önerilmiştir. Bu model, ReLU aktivasyon katmanı ve çok sınıflı softmax sınıflandırıcı ile desteklenmiştir. Deneysel sonuçlar, önerilen hibrit modelin %96'dan fazla doğruluk ve özellikle olumsuz tweetlerde %100 hassasiyet sağladığını ortaya

koymuřtur. Çalışma, Arapça duygu analizinde hem yöntemsel katkı sunmakta hem de transformatör tabanlı hibrit modellerin etkinliğini ortaya koymaktadır.

Sosyal medyada her geçen gün artan bilgi yoğunluğu, söylentilerin ve sahte haberlerin zamanında tespit edilmesini hayati hale getirmiştir. Özellikle yapay zekâ tarafından oluşturulan içeriklerin yaygınlaşması, doğru bilgiyi ayırt etmeyi daha da zorlaştırmaktadır. Geleneksel yöntemlerle yapılan doğrulama süreçleri, bu kadar büyük bir veri akışı içinde yetersiz kalmakta ve oldukça zaman almaktadır. Bu nedenle, sahte içerikleri otomatik olarak tanıyabilen sistemlere olan ilgi giderek artmaktadır. Bu kapsamda yapılan bir çalışmada, doğal dil işleme alanında öne çıkan BERT ve RoBERTa modelleri, yapay zekâ kaynaklı haberleri tanıma konusunda oldukça başarılı sonuçlar vermiştir. Özellikle ince ayar yapılmış RoBERTa modeli, %98 doğruluk oranına ulaşarak dikkat çekici bir performans sergilemiştir. Elde edilen bulgular, bu tür derin öğrenme tabanlı modellerin, ChatGPT gibi yapay zekâ sistemleri tarafından üretilen yanıltıcı haberlerin ayıklanmasında etkili araçlar olabileceğini ortaya koymaktadır. Sonuç olarak, bu modellerin yanlış bilginin yayılmasını engellemede önemli katkılar sunabileceği anlaşılmaktadır (Wang vd., 2023).

Literatürde yükseköğretim haberlerinde önyargı tespiti sınırlı kalmış olup, bu boşluk DistilBert modeliyle doldurulmaya çalışılmıştır. Ayrıca açıklama görevlerinde GPT-3.5-turbo ve Google kullanılarak modellerin performansları karşılaştırılmıştır. Gerçek etiketler Amazon Mturk aracılığıyla elde edilmiş ve açıklamaların doğruluğu değerlendirilmiştir. Bulgular, DistilBert'in önyargı tespitinde orta düzeyde performans gösterdiğini, GPT-3.5-turbo ve Gemini açıklamalarının ise insan açıklamalarıyla karşılaştırıldığında düşük doğruluk sunduğunu ortaya koymaktadır. Bu sonuçlar, büyük dil modellerinin bu alandaki açıklamalarına temkinli yaklaşılması gerektiğini göstermektedir. Bin Shiha vd. (2023) tarafından gerçekleştirilen çalışma, araştırma odaklı bir İngiliz üniversitesinin haber makalelerinde önyargılı dilin tespitine odaklanmaktadır.

Yang vd. (2024) araştırmasında, konuşma sentezi alanında yenilikçi bir yaklaşıma odaklanmaktadır. Geleneksel metinden konuşma üretim sistemlerinin ötesine geçerek, doğal dilin farklı konuşma stillerini esnek bir şekilde kontrol edebilmek amacıyla InstructTTS adlı bir model geliştirilmiştir. Bu modelin eğitimi için, stil bilgilerini içeren geniş kapsamlı bir konuşma veri seti oluşturulmuştur. InstructTTS, kendi

kendine denetlenen öğrenme ve metrik temelli optimizasyon tekniklerinden yararlanarak üç aşamalı bir eğitim süreciyle geliştirilmiştir. Yapılan deneyler, modelin çeşitli konuşma tarzlarını başarıyla sentezleyebildiğini ve yüksek kalite standartlarını karşıladığını ortaya koymaktadır.

Hsu vd. (2024) tarafından gerçekleştirilen araştırması, doğal dil işleme teknolojilerinin ilaç geliştirme süreçlerinde sağladığı faydaları ele almaktadır. Bu çalışma, klinik farmakolojiye yönelik üç ana kullanım alanına odaklanmaktadır: farmakokinetik parametrelerin analizi, klinik soruların yanıtlanması ve düzenleyici belgelerin incelenmesi. Doğal dil işlemenin, büyük hacimli verileri kısa sürede işleyerek kritik bilgileri çıkarmada ve analiz etmede nasıl etkili olabileceği ortaya konmuştur. Araştırma, bu tür teknolojilerin ilaç geliştirme süreçlerini daha hızlı ve verimli hale getirme potansiyeline dikkat çekmekte ve klinik verilerin daha etkin kullanımını sağlamanın önemini vurgulamaktadır.

Tejaswini vd. (2024) tarafından gerçekleştirilen çalışması, sosyal medya platformlarında paylaşılan metinlerin doğal dil işleme yöntemleriyle analiz edilerek depresyon belirtilerinin erken aşamalarda tespit edilmesine odaklanmaktadır. Bu amaçla, araştırmacılar "Uzun Kısa Süreli Belleğe Sahip Hızlı Metin Evrişim Sinir Ağı" (Fast Text Convolutional Neural Network with Long Short-Term Memory) adı verilen yenilikçi bir hibrit derin öğrenme modeli geliştirmiştir. Bu model, metinlerin anlamını daha etkili bir şekilde temsil etmek için hızlı metin yerleştirme yöntemleri ile evrişim sinir ağlarının (CNN) genel özellik çıkarma gücünü ve LSTM ağlarının dizisel bilgi yakalama yeteneklerini birleştirmektedir. Gerçek dünya veri kümeleri üzerinde yapılan testler, modelin depresyon belirtilerini tespit etmede yüksek doğruluk sağladığını göstermiştir.

Daungsupawong ve Wiwanitkit (2024) araştırması, üretken dil modellerinin sağlık hizmetlerinde sunduğu geniş kullanım olanaklarını incelemektedir. Bu modeller, hasta bilgilendirme materyallerinin hazırlanması, tıp öğrencilerine sanal eğitim sağlanması, idari işlerin otomasyonu ve klinik araştırmaların desteklenmesi gibi birçok alanda önemli potansiyele sahiptir. Örneğin, hastalar için kişiselleştirilmiş rehberler oluşturabilir, eğitim süreçlerini zenginleştirebilir ve klinik veri analizlerini hızlandırabilir. Ancak, bu teknolojilerin kullanımı bazı zorlukları da beraberinde getirmektedir. Özellikle, güncel veriye erişim kısıtları, veri önyargılarının yansıtılması,

etik kaygılar ve otomasyonun aşırı kullanımına bağlı riskler dikkate alınmalıdır. Bu nedenle, modellerin doğru ve güvenilir şekilde kullanılabilmesi için insan denetimi büyük önem taşımaktadır. Gelecekte yapılacak çalışmaların, bu sınırlamaları aşarak daha güvenilir ve etkili çözümler geliştirmeye odaklanması gerektiği vurgulanmaktadır.

Deguchi vd. (2024) tarafından gerçekleştirilen araştırmasında, doğal dil kullanılarak oluşturulan yol tariflerinden topolojik haritalar üreten yeni bir yöntem geliştirilmiştir. Bu yaklaşım, metinsel ifadelerden çıkarılan mekânsal ilişkileri kullanarak robotların çevrelerini daha iyi kavramalarına ve karmaşık yol planlamaları yapmalarına olanak tanımaktadır. Çalışmada, bu yöntemle oluşturulan haritaların, robotların kullanıcı komutlarını daha hassas bir şekilde yorumlayarak yeni güzergahlar oluşturmada önemli avantajlar sunduğu gösterilmiştir. Bu, robotların çevresel uyum kabiliyetlerini artırarak daha otonom ve esnek hareket etmelerini mümkün kılmaktadır.

Hata raporları, kullanıcıların yazılımda karşılaştıkları sorunları geliştiricilere iletmeleri açısından temel bir iletişim aracı olup, yazılımın bakım ve iyileştirme süreçlerinde önemli katkı sağlar. Bununla birlikte, mevcut otomatik hata yeniden üretme sistemleri, çoğunlukla raporların yapısal bütünlüğüne ve içerik kalitesine bağımlı oldukları için istenilen düzeyde performans sunamamaktadır. Bu soruna çözüm olarak Feng ve Chen (2024), büyük dil modellerinin doğal dili anlama konusundaki yüksek yetkinliğinden faydalanarak, AdbGPT isimli yeni ve düşük maliyetli bir yaklaşım geliştirmiştir. Bu yöntem, herhangi bir ek eğitim süreci ya da sabit kodlama gereksinimi olmadan çalışmakta ve hata raporlarını geliştirici perspektifinden yeniden üretebilmek için LLM'lerin bilgi temelli akıl yürütme becerilerinden yararlanmaktadır. Yapılan deneysel çalışmalar, AdbGPT'nin hata raporlarının %81,3'ünü ortalama 253,6 saniye içinde başarılı şekilde yeniden oluşturabildiğini ortaya koymuştur. Bu sonuçlar, önerilen yöntemin yazılım hata yönetiminde pratik ve etkili bir çözüm sunduğunu göstermektedir.

Chen (2024) tarafından gerçekleştirilen çalışma, üç büyük dil modeli sohbet robotunun (ChatGPT-3.5, ChatGPT-4 ve Gemini) kendiliğinden konuşmalardan türetilen metinler aracılığıyla Alzheimer Demansı (AD) ve Bilişsel Olarak Normal (CN) bireyleri ayırt etme performansını incelemektedir. Sıfır atışlı öğrenme ve düşünce zinciri istemiyle yapılan değerlendirmelerde, Gemini AD tanısında en yüksek hatırlama (%89) ve F1

puanı (%71) elde etmiş, ancak CN'yi AD olarak yanlış tanıma eğilimi göstermiştir. GPT-4 ise CN tanısında en yüksek F1 puanını (%62) üretmiştir. Modeller, şans düzeyinin üzerinde performans sergilemiş ancak klinik uygulama için yeterli bulunmamıştır.

Xu vd. (2024) çalışmasında, hisse senedi yorumlarının konu tanımlamasını iyileştirmeye yönelik iki yenilikçi strateji sunmaktadır. Finansal metinlerin doğru sınıflandırılması, yatırımcıların bilinçli kararlar almasına katkı sağlamaktadır. Mevcut yöntemlerden BERT modeli, güçlü performansına rağmen uzun metinlerde sınırlı giriş kapasitesi nedeniyle bazı zorluklarla karşılaşmaktadır. Bu sorunu aşmak amacıyla ilk olarak, ChatGPT tarafından oluşturulan özetler BERT modeline giriş olarak verilmiştir. ChatGPT'nin üretken özetleme yeteneği sayesinde metinler hem daha kısa hale getirilmiş hem de anlam bütünlüğü korunarak modelin performansı artırılmıştır. İkinci olarak, BERT'in farklı katmanlarında çok katmanlı özellik ablasyonu uygulanarak, konu tanımlaması açısından en etkili temsillerin belirlenmesi hedeflenmiştir. Bu yöntem, modelin yalnızca yüzeysel değil, derin bağlamsal bilgileri de öğrenmesini sağlayarak doğruluğu artırmaktadır. Gerçekleştirilen deneyler, önerilen iki stratejinin birlikte kullanıldığında hisse senedi yorumlarının konu tanımlamasında kayda değer iyileşmeler sağladığını ortaya koymaktadır. Bulgular, finansal metin analizi için daha etkili ve verimli sınıflandırma modelleri geliştirilmesine katkı sunmaktadır. Literatürde yer alan çalışmaların karşılaştırmalı bir özeti, kullanılan veri kümeleri, yöntemler ve doğruluk oranları temelinde Tablo 1'de sunulmuştur.

Tablo 1.

Literatürdeki Çalışmaların Karşılaştırması

Yazar(lar) ve Yıl	Çalışma Konusu	Kullanılan Veri Kümesi	Kullanılan Yöntem(ler)	Doğruluk (%) / Performans
Munika vd. (2019)	Detaylı duygu sınıflandırması	Çok sınıflı duygu veri seti	Transformer, BERT	BERT en yüksek başarı
Acikalin vd. (2020)	Türkçe film ve otel yorumlarında duygu analizi	Türkçe yorum veri seti	Çok dilli BERT & çeviri destekli analiz	Başarılı sınıflandırma, oran verilmemiş
Öztürk vd. (2020)	'de/da' yazım hatalarının tespiti	Türkçe yazım hatası örnekleri	BERT, bağlamsal analiz	BERT bağlamsız modellere göre daha başarılı

Sevli & Kemaloğlu (2021)	Olağandışı olay tweet'lerinin sınıflandırılması	7613 tweet	Google BERT, sinir ağı	%98,88
Ozan & Taşar (2021)	Diyalogların kategorik etiketlenmesi	14.000 müşteri diyalogu	Transformatör mimarisi, BERT	%98 BERT
Ayan vd. (2021)	Şirket anahtar kelime analizi	Türkiye'nin en büyük 25 şirketi	NLP kümeleme	Nitel sonuçlar, oran verilmemiş
Ozan & Taşar (2021)	46 kategoriye ayrılmış müşteri-temsilci yazışmalarının sınıflandırılması	46 kategori içeren veri seti	GPT-2 ve BERT modelleri	BERT ile yüksek başarı elde edilmiş (kesin oran yok)
İşeri vd. (2021)	Sohbet robotu geliştirme	SSS, görüşmeler, mesajlar	NLP + ML tabanlı chatbot	Yüksek doğruluk gözlemlenmiş
Hark (2022)	Sahte haber tespiti	COVID-19 sahte haber veri seti (10.700 örnek)	LSTM, derin bağlamsal gömme	%91
Toğaçar vd. (2022)	Sahte haber tespiti	6335 haber (3171 gerçek, 3164 sahte)	NLP + LSTM	Eğitim: %99,83; Test: %91,48
Yazgılı & Baykara (2022)	Siber zorbalık tespiti	3000 sosyal medya örneği	Bagging & Boosting karşılaştırmalı analiz	Karşılaştırmalı başarılar (rakam yok)
Bayram (2022)	COVID-19 içerik analizi	Sosyal medya + dijital haber arşivi	ML yöntemleri, sözcüksel ağlar	Nicel başarı belirtilmemiş
Hark (2022)	COVID-19 sahte haber tespiti (başka veri seti)	Açık erişimli COVID-19 veri seti (~10.700 öge)	Derin bağlamsal gömülmeler + sinirsel ağlar (LSTM)	En yüksek: LSTM ile %91 başarı
Özkan & Kar (2022)	Türkçe bilimsel içeriklerde çok sınıflı metin sınıflandırma	Türkçe bilimsel/akademik içerik veri kümesi	Hazır Türkçe BERT yeniden eğitimi	%96 doğruluk
Altıntop (2023)	ChatGPT gibi AI sistemlerinin sunduğu olanaklar ve sorunlar	-	Literatür incelemesi	-
Koçyiğit & Darı (2023)	ChatGPT temelli yapay zekâ ve iletişim etkileşiminin değerlendirilmesi	-	İnceleme/tartışma çalışması	-
Yiğit vd. (2023)	ChatGPT'nin sağlık sektöründe kullanımı	-	Uygulama analizleri	Nitel katkı, doğruluk belirtilmedi
August vd. (2023)	Tıbbi makale anlama sistemi	Akademik makaleler	Özel prototip sistem	Başarı: Kullanıcı anlama başarısı arttı
Müller vd. (2023)	CT-BERT ile COVID-19 NLP görevleri	Twitter mesajları	CT-BERT modeli	Çok görevli başarı, oranlar değişken

Amini vd. (2023)	Demans tespiti	Nöropsikolojik ses kayıtları	Otomatik transkripsiyon NLP modeli	%92 üzeri doğruluk
Çetindağ vd. (2023)	Türkçe hukuki metinlerde NER	Hukuki belgeler	Yeni NER modeli	Türkçe için yüksek başarı
Jiang vd. (2023)	MotionGPT: Hareket + Dil entegrasyonu	3D hareket verileri	MotionGPT, vektör kuantizasyonu	Yüksek doğruluk ve esneklik
Zong & Krishnamachari (2023)	GPT-3 ile matematiksel problem çözme	Sözlü matematik problemleri	GPT-3	Tüm aşamalarda yüksek başarı
Mutinda vd. (2023)	LeBERT: Duygu analizi	Standart veri setleri	BERT + sözlük + CNN	Mevcut yöntemlere kıyasla daha iyi
Bird vd. (2023)	İnsan merkezli sohbet çerçevesi	Kullanıcı testleri	AI destekli sohbet sistemi	Gerçekçi insan-makine iletişimi
Zhang & Wang (2023)	Haber öneri sistemi	MIND veri kümesi	Prompt4NR, maske tahmini	Performansta iyileşme sağlandı
Kanan vd. (2023)	Müşteri memnuniyeti tahmini	Sosyal medya verisi	DNN	Derin öğrenme ile yüksek başarı
Saleh vd. (2023)	Nefret söylemi tespiti	Online veri	LSTM, BERT	Yüksek doğruluk, veri boyutu önemli
Salvagno vd. (2023)	Chatbot kullanımı ve etik sorunlar bilimsel yazımda	–	Literatür temelli analiz	-
Kim (2023)	ChatGPT'nin bilimsel yazımda dil düzenleme kullanımının incelenmesi	–	İnceleme çalışması	-
Gao vd. (2023)	ChatGPT özetlerinin gerçek özetlerle karşılaştırılması	50 tıbbi makale özeti	ChatGPT özetleri karşılaştırmalı değerlendirme	Özetlerin %68'i ChatGPT tarafından üretildi
Long (2023)	AI içerik üretiminin saptanması yöntemlerinin değerlendirilmesi	–	Stilometri, CopyLeaks, Turnitin vb.	Teknik kombinasyon gerektiği vurgusu
Görgün vd. (2023)	Demografik tahmini: yaş ve cinsiyet	1.040 Türkçe Twitter kullanıcısı metni	fastText, BERT, karakter/kelime özellikleri	İçerikler demografik tahmin için anlamlı gösterge
Mujahid vd. (2023)	Arapça tweet'lerde ChatGPT'ye yönelik kullanıcı duygu analizleri	27.780 Arapça tweet (TextBlob + RoBERTa/XLNet/DistilBERT hibrit modeli)	Hibrit model: RoBERTa + BERT + ReLU + softmax	%96'dan fazla doğruluk, olumsuz tweetlerde %100 hassasiyet
Wang vd. (2023)	ChatGPT üretilmiş sahte haber tanıma	–	RoBERTa; ince ayarlı	RoBERTa %98 doğruluk

Bin Shiha vd. (2023)	Haber makalelerinde önyargı ve açıklama kalitesi	Üniversite haber makaleleri	DistilBERT, GPT 3.5 turbo, Gemini	DistilBERT orta; GPT 3.5/Gemini düşük doğruluk
Yang vd. (2024)	InstructTTS ile konuşma stili kontrolü	Stil etiketli konuşma veri seti	InstructTTS model, üç aşamalı eğitim	Yüksek kalite ve stil kontrolü sağlanmış
Hsu vd. (2024)	NLP'nin ilaç geliştirmedeki klinik farmakoloji uygulamaları	Klinik farmakoloji veri kümeleri	NLP yöntemleri	-
Tejaswini vd. (2024)	Depresyon tespiti için hibrit model	Gerçek sosyal medya metin kümeleri	Fast Text CNN + LSTM modeli	Yüksek doğruluk sağlandığı rapor edilmiş
Daungsupawong & Wiwanitkit (2024)	Generatif dil modellerinin sağlık hizmetlerinde kullanımı	—	Literatür/uygulama incelemesi	İnsan denetimi vurgusu, potansiyel & risk analizi
Deguchi vd. (2024)	Metinden topolojik harita çıkarımı	Yol tariflerinden oluşturulan metin veri kümesi	NLP tabanlı yöntem	Robot sistemlerine avantaj sağladığı belirtilmiş
Feng & Chen (2024)	Yazılım hata raporlarının yeniden üretilmesi (AdbGPT)	Hata raporu metinleri	AdbGPT (LLM temelli)	%81,3 başarılı yeniden üretim, ~253,6 sn ortalama süre
Chen (2024)	Alzheimer/Demens ayrımı için ChatGPT 3.5/4/Gemini'nin değerlendirilmesi	Konuşmalardan türetilen metinler	Sıfır atış + düşünce zinciri istemi	Gemini: %89 recall, F1 %71; GPT 4: CN tanısında F1 %62
Xu vd. (2024)	Hisse senedi yorumlarında konu tanımlamasını iyileştirme	Finans yorumları veri kümesi	ChatGPT özet + BERT; katmanlı ablasyon stratejileri	-

BÖLÜM III

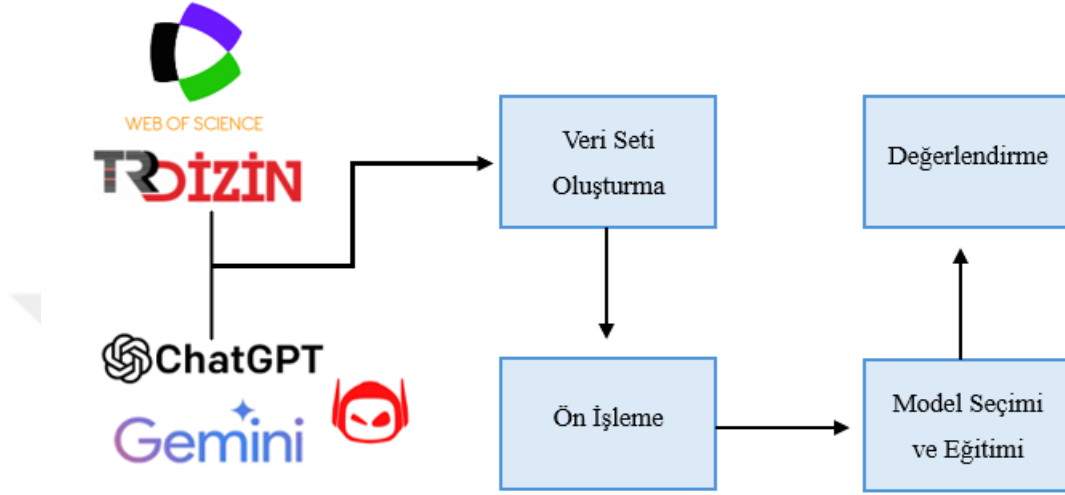
MATERYAL VE YÖNTEM

3.1. Araştırmanın Modeli

Bu çalışma, akademik alanda yapay zekâ tarafından üretilen içeriklerin, insan üretimi içeriklerden ayırt edilmesini amaçlayan derin öğrenme tabanlı bir tespit modelinin geliştirmesi hedeflenmiştir. Bu doğrultuda önce veri toplama veri toplama, etiketleme, ön işleme, ardından sınıflandırma için model eğitimi ve adımları uygulanmıştır. İlk aşamada aşamasında, ChatGPT gibi üretken yapay zekâ tabanlı dil modellerinin 2022 yılı başlarından itibaren yaygın biçimde kullanılmaya başlanması nedeniyle, bu tarihten sonra yayımlanan akademik içeriklerin üretiminde bu tür araçların kullanılmış olma ihtimali bulunmaktadır. Bu durum, insan üretimi içeriğin özgünlüğünü etkileyebileceğinden, çalışmada 2022 öncesi yayımlanmış akademik metinler dikkate alınmıştır. Bu kapsamda, Türkçe dilinde ve açık erişimli kaynaklardan (TR Dizin, Scopus, SCI, ESCI gibi güvenilir indekslerden) 2.000 adet özgün insan üretimi akademik makale özeti toplanmıştır. Bu özetler temel alınarak, her biri için üretken yapay zekâ araçları (ChatGPT, Gemini, Smodin) aracılığıyla 2.000 adet yapay zekâ tarafından üretilmiş özet oluşturulmuştur. Ayrıca, bu yapay özetlerin dil ve anlatım açısından insan müdahalesiyle gözden geçirilip yeniden yapılandırılmasıyla 2.000 adet yapay zekâ destekli, insan tarafından düzenlenmiş özet elde edilmiştir. Toplamda 6.000 adet metinden oluşan dengeli ve üç kategorili bir veri seti meydana getirilmiştir.

Bu içerikler, insan üretimi veri setini oluşturmak amacıyla sınıf dengesi gözetilerek seçilmiştir. Paralel olarak, aynı başlık ve anahtar kelimeler temel alınarak ChatGPT, Gemini ve Smodin gibi güncel üretken yapay zekâ araçları kullanılarak yapay içerikler üretilmiş ve ikinci bir veri seti oluşturulmuştur. Elde edilen metinler üzerinde ön işleme adımında Türkçe doğal dil işleme teknikleri (tokenizasyon, durak kelime kaldırma, kök/sonek ayrıştırma vb.) uygulanmıştır. Ön işlenen veriler, yapay ve insan üretimi metinler olarak etiketlenmiş ve sınıflandırma için uygun bir formata dönüştürülmüştür. Ardından, çeşitli derin öğrenme tabanlı metin sınıflandırma algoritmaları (XLM-RoBERTa, TurkishBERT, Bert Base Multilingual) kullanılarak model eğitimi gerçekleştirilmiştir. Modellerin hiper parametre optimizasyonları

yapılarak sınıflandırma sürecine geçilmiştir. Model değerlendirme sürecinde doğruluk, kesinlik, duyarlılık ve F1 skoru esas alınmış; geliştirilen modeller karşılaştırmalı olarak analiz edilmiştir. Uygulamanın aşamalarını belirten akış diyagramı Şekil 2’de gösterilmektedir.



Şekil 2. Uygulamanın akış diyagramı

Şekil 2’ de Yapay zekâ tarafından üretilen içeriklerin tespiti amacıyla geliştirilecek modelin genel işlem akışı verilmiştir. Öncelikle, Web of Science, TR Dizin, Scopus gibi akademik dizinlerden farklı disiplinlerde makale özetleri derlenmiştir. Ardından ChatGPT, Gemini ve Smodin gibi üretken yapay zekâ araçları kullanılarak elde edilen metinlerden yeni özetler oluşturulmuştur. Ardından yine bu yapay zekâ araçları kullanılarak orijinal özetler yeniden düzenlenmiştir (paraphrase). Bu şekilde orijinal, yapay zekâ üretimi ve yapay zekâ tarafından düzenlenmiş içerikleri barındıran üç sınıflı bir veri seti oluşturulmuştur. Ardından veriler üzerinde ön işleme adımları uygulanmış ve model eğitime geçilmiştir. Eğitilen modeller, çeşitli performans kriterlerine göre değerlendirilmiş ve en iyi sonuç veren model, nihai tespit aracı olarak seçilmiştir.

3.2. Veri Seti

Bu çalışmada kullanılan veri seti oluşturulurken 2022 öncesi SCI, SCI-E, ESCI, Scopus ve TR Dizin gibi tanınmış akademik dizinlerde yer alan açık erişimli Türkçe

bilimsel makalelerin özetleri alınmıştır. Toplam 2.000 adet orijinal yazılmış metin derlenmiştir. Çalışmanın odak noktası, bu içeriklerin yapay zekâ tarafından üretilen metinlerle karşılaştırılıp ve ayrıştırıla bilirliliğinin test edilmesidir. ChatGPT'nin yaygın kullanılmaya başladığı 2022 sonrasında akademik makalelere yapay zekâ üretimi içeriklerin de karışmış olma ihtimaline karşı 2022 öncesine odaklanılmıştır.

2022 öncesi seçilen 2.000 adet makalenin gerçek özetlerinin yanında bu makaleler ChatGPT, Smodin ve Gemini gibi yapay zekâ destekli metin üretme araçları kullanılarak yeniden özetlenmiş ve toplam 2.000 adet yapay zekâ üretimi özet metin elde edilmiştir. Daha sonra, özgün özetler üzerinde yapay zekâ araçları ile, yazım düzeni, cümle yapısı ve ifadesel düzenlemeler yapılarak, özgün metinlerin yeniden düzenlenmiş versiyonu oluşturulmuş ve bu metinler "Düzenlenmiş" etiketiyle tanımlanmıştır.

- **Gerçek Veri (İnsan Üretimi):** 2022 yılı öncesinde yayımlanmış, farklı akademik disiplinlerden (eğitim bilimleri, sağlık, sosyal bilimler ve mühendislik), kaynaklar (WOS, TR Dizin, Scopus) seçilen, özgün metinlerden oluşturulmuştur.
- **Yapay Zekâ (YZ Üretimi):** Gerçek içeriklerin dokümanları dosyası olarak yapay zekâ araçlarına yüklenmiş ve dil modelinin bu metinden yeni bir özet oluşturması istenmiştir.
- **Düzenlenmiş Veri:** Orijinal özetler yapay zekâ modelleri ile revize edilmiştir.

3.3. Veri Çekme

Bu çalışmada kullanılan veri seti, 2022 yılı öncesinde yayımlanmış ve ulusal ile uluslararası düzeyde tanınırlığı olan akademik dizinlerde (SCI, SCI-E, ESCI, Scopus ve TR Dizin) yer alan açık erişimli Türkçe bilimsel makalelerden elde edilmiştir. Veriler, tamamen manuel (elle) toplama yöntemiyle derlenmiştir. Bu süreçte, açık erişimli makaleler tercih edilmiştir.

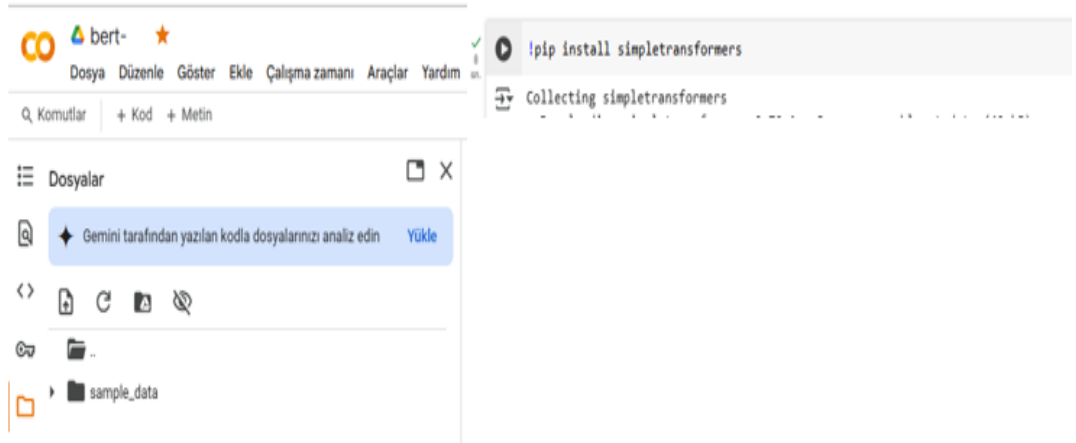
Veri toplama süreci, ilgili dizinlerin arama arayüzleri aracılığıyla gerçekleştirilmiş; eğitim bilimleri, sağlık, mühendislik, sosyal bilimler alanlarından mümkün olduğunca dengeli sayıda ve güncel yayınlar seçilmiştir. Seçilen içeriklerin özetleri, biçimsel bütünlük sağlanarak tablolaştırılmıştır. Her bir özetin karşısına etiketleri eklenmiştir.

Veri seti tüm platformlar için evrensel CSV (Comma Separated Values) dosya formatına dönüştürülmüştür.

3.4. Yazılım Geliştirme Ortamı

Bu çalışma kapsamında geliştirilen tüm makine öğrenmesi ve derin öğrenme modelleri, Google tarafından sunulan bulut tabanlı bir geliştirme ortamı olan Google Colab üzerinde yürütülmüştür. Google Colab, Python tabanlı kodların çevrim içi olarak yazılmasına, çalıştırılmasına ve görselleştirilmesine olanak tanımaktadır. En önemli avantajlarından biri, GPU (Graphics Processing Unit) ve TPU (Tensor Processing Unit) gibi yüksek hesaplama gücüne sahip donanımları kullanıcılara ücretsiz olarak sunmasıdır. Bu donanımsal destek, özellikle BERT gibi büyük dil modellerinin eğitimi ve test edilmesi gibi yoğun işlem gücü gerektiren işlemlerde kritik bir rol oynamaktadır.

Google Colab ayrıca, veri yükleme, analiz ve sonuçları görselleştirme gibi süreçleri destekleyen zengin bir eklenti altyapısına sahiptir. Kullanıcıların kodları doğrudan Google Drive üzerinde kaydedebilmesi, projelerin paylaşılmasını ve sürüm kontrolünü kolaylaştırmaktadır. Bu çalışma kapsamında Google Colab ortamı, kodların test edilmesi, hiper parametre ayarlarının optimize edilmesi ve model çıktılarının analiz edilmesinde yoğun olarak kullanılmıştır. Ayrıca, SimpleTransformers ve TensorFlow gibi kütüphanelerin kolaylıkla entegre edilmesi sayesinde, deneysel tasarım süreci hızlandırılmıştır. Colab çalışma ortamından örnek bir görsel Şekil 3'te gösterilmektedir.



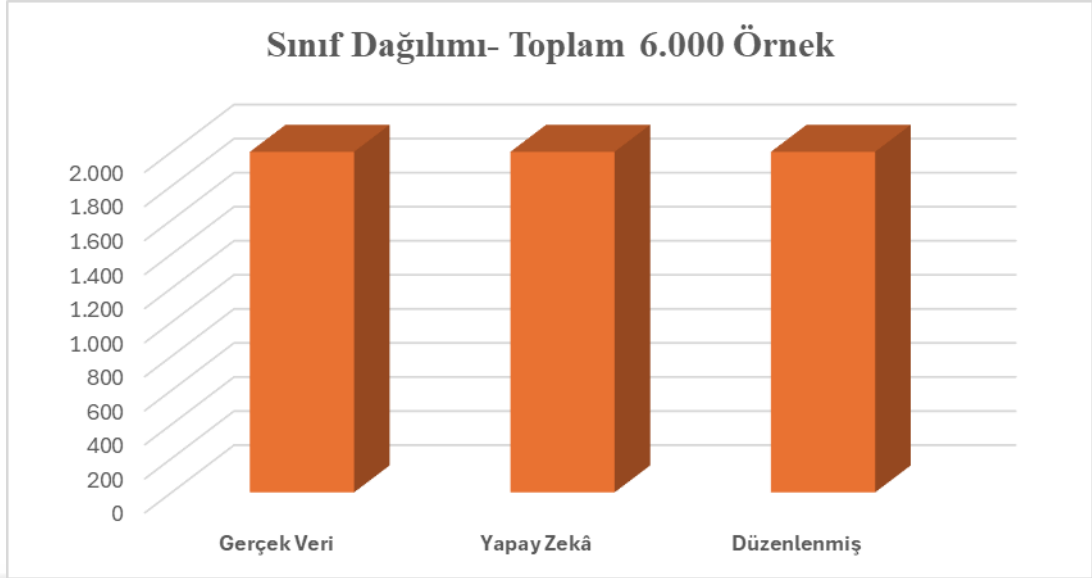
Şekil 3. Google Colab çalışma ortamı

3.4.1. Veri önışleme. Doğal dil işleme tabanlı sınıflandırma modellerinde doğru sonuçlar elde edebilmek için verinin uygun şekilde önışleme kritik öneme sahiptir. Bu çalışmada, toplanan metin verileri üzerinde bir dizi önışleme adımı uygulanmıştır. Bu adımlar, yazın dilinin analize uygun bir formata dönüştürölmesini sağlamaktadır. İlk olarak tüm metinler küçük harfe dönüştürölerek, standart sağılanmıştır. Ardından, tüm noktalama işaretları metinlerden kaldırılmış, anlamsız boşluk karakterleri atılmıştır. Her bir metne ait sınıf bilgisi, yani etiketi ("Gerçek Veri", "Yapay Zekâ", "Düzenlenmiş") sayısal formata dönüştürölerek 0, 1, 2 değeri ile etiketlenmiştir. Etiketlerin örnek metinlerle birlikte nasıl yapılandırıldığına ilişkin Tablo 2’de sunulmuştur. Bu tabloda, her bir sınıfa ait örnek bir metin ve karşılık gelen sayısal etiket değeri yer almaktadır. Veri seti içerisindeki örnek dağılımı Şekil 4’te sunulmuştur. Bu grafik, çalışmanın başlangıcında kullanılan veri setindeki "Gerçek Veri", "Yapay Zekâ", "Düzenlenmiş" dağılımını göstermektedir.

Tablo 2. Veri Kümesinin Etiketleri ve Örnek Görüntüsü

Veri Kümesinin Etiketlenmiş Örnek Görüntüsü

Örnek Metin	Etiket	Sayısal Karşılık
Bu çalışmada öğrenci başarıları üzerine yapılan analizler ele alınmıştır.	Gerçek Veri	0
Yapay zekâ, son yıllarda eğitimde etkin bir araç hâline gelmiştir.	Yapay Zekâ	1
Araştırmada öğrencilerin başarı düzeyleri üzerinde durulmuştur.	Düzenlenmiş	2



Şekil 4. Sınıf dağılım grafiği

Son olarak, eğitim ve test setlerine ayırma işlemi gerçekleştirilmiştir. Veri setinin %80'i eğitim, %20'si test için kullanılmıştır. Bu oran, literatürde sıkça tercih edilen bir bölme yöntemidir; modelin öğrenme sürecinde yeterli veriyle eğitilmesini sağlarken, genelleme performansının sağlıklı bir şekilde değerlendirilebilmesine de imkân tanımaktadır. Bölme işlemi veri seti karıştırılarak, rastgele ve her grupta sınıflardan eşit sayıda olacak şekilde gerçekleştirilmiştir.

Türkçe akademik içeriklerin yapay zekâ tarafından üretilip üretilmediğini tespit etmek amacıyla derin öğrenme tabanlı güncel dil modelleri kullanılmıştır. Bu amaçla, özellikle çok dilli ve Türkçe desteği olan TurkishBERT, RoBERTaBERT ve Bert Base Multilingual tercih edilmiştir. Bu modellerin seçilme nedenleri, geniş veri setleri üzerinde önceden eğitilerek farklı dillerde üstün performans göstermeleri ve Türkçe gibi zengin morfolojiye sahip dillerde yüksek doğruluk oranları sunmalarıdır.

3.4.2. XLM-RoBERTa modeli. XLM-RoBERTa, Facebook AI Research (FAIR) tarafından geliştirilen, çok dilli metinlerin anlamını modellemek amacıyla optimize edilmiş, Transformer tabanlı bir dil modelidir. Bu model, 100'den fazla dilde eğitim almış olup, geniş çaplı dil çeşitliliğine yönelik üstün genelleme yetenekleri nedeniyle tercih edilmiştir. XLM-RoBERTa, 2.5 TB boyutunda CommonCrawl verisinden türetilmiş büyük ölçekli veri setlerinde eğitilmiştir. Modelin eğitimi sırasında MLM stratejisi kullanılmıştır.

Bu strateji, cümle içerisindeki bazı kelimelerin maskeleyme işlemiyle gizlenmesi ve modelin bu kelimeleri doğru tahmin etmesi üzerine kuruludur. Bu yöntem, modelin bağlama dayalı anlam çıkarma yeteneğini geliştirmektedir. Özellikle Türkçe gibi zengin morfolojik yapıya sahip dillerde yüksek doğruluk oranları sunar. Şekil 5'te XLM-RoBERTa modeli gösterilmektedir.

```
from simpletransformers.classification import ClassificationModel
import torch

model = ClassificationModel(
    "xlmroberta",
    "xlm-roberta-base",
    num_labels=3,
    use_cuda=torch.cuda.is_available(),
    args={
        "reprocess_input_data": True,
        "overwrite_output_dir": True,
        "num_train_epochs": 3,
        "train_batch_size": 16,
        "max_seq_length": 512,
        "save_steps": -1,
        "save_model_every_epoch": False,
        "learning_rate": 4e-5,
    }
)
```

Şekil 5. XLM-RoBERTa modeli

Model, GPU desteği varsa CUDA kullanarak hızlandırılmış eğitim yapabilir. Bu modelin geniş veri setlerinde eğitilmiş olması, bağlama dayalı anlam çıkarma yeteneğini artırarak, Türkçe gibi eklemeli dillerde yüksek performans sunmasını sağlar. Ayrıca, maksimum dizin uzunluğu (max_seq_length) ve öğrenme oranı (learning_rate) gibi hiper parametreler modelin doğruluğunu doğrudan etkileyen kritik ayarlardır.

3.4.3. Bert Base Multilingual modeli. Bert Base Multilingual, BERT tabanlı çok dilli bir dil modelidir ve özellikle çoklu dil desteği sunmak üzere geliştirilmiştir. Bert-base-multilingual-cased sürümü, büyük çaplı çok dilli veri setleri kullanılarak eğitilmiştir ve çapraz dil analizi için optimize edilmiştir. Bu model, dil bağımsız temsil öğrenmesini destekler ve metin sınıflandırma, duygu analizi gibi görevlerde yüksek doğruluk oranları sunar. Şekil 6'da Bert Base Multilingual kurulumu gösterilmektedir.

```

model = ClassificationModel(
    "bert",
    "bert-base-multilingual-cased",
    num_labels=3,
    use_cuda=torch.cuda.is_available(),
    args={
        "reprocess_input_data": True,
        "overwrite_output_dir": True,
        "num_train_epochs": 3,
        "train_batch_size": 16,
        "max_seq_length": 512,
        "save_steps": -1,
        "save_model_every_epoch": False,
        "learning_rate": 4e-5,
    }
)

```

Şekil 6. *Bert Base Multilingual modeli*

Bu model, çok dilli veri setlerinde metin sınıflandırma için optimize edilmiştir. Geniş dil kapsama alanı sayesinde, farklı dillerde yüksek doğruluk oranları sunar. Modelin eğitimi sırasında kullanılan hiper parametreler, modelin genelleme kapasitesini doğrudan etkiler.

3.4.4. TurkishBERT modeli. TurkishBERT, Türkçe diline özgü olarak geliştirilmiş bir BERT modelidir. Türkçe'nin eklemeli morfolojik yapısı göz önünde bulundurularak eğitilmiştir ve daha karmaşık cümle yapılarını daha iyi anlamak üzere optimize edilmiştir. Model, Türkçe Wikipedia ve çeşitli Türkçe haber sitelerinden toplanan büyük hacimli veri setleri kullanılarak eğitilmiştir. Bu sayede, model Türkçe metinlerde anlam çıkarma, dilsel bağlamı kavrama ve semantik ilişkileri keşfetme konusunda yüksek performans göstermektedir. Şekil 7'de TurkishBERT kurulumu gösterilmektedir.

```

model = ClassificationModel(
    "bert",
    "dbmdz/bert-base-turkish-cased",
    num_labels=3,
    use_cuda=torch.cuda.is_available(),
    args={
        "reprocess_input_data": True,
        "overwrite_output_dir": True,
        "num_train_epochs": 3,
        "train_batch_size": 16,
        "max_seq_length": 512,
        "save_steps": -1,
        "save_model_every_epoch": False,
        "learning_rate": 4e-5,
    }
)

```

Şekil 7. *TurkishBERT modeli*

Bu model, Türkçe metinlerin anlamını daha doğru bir şekilde analiz edebilmek için optimize edilmiştir. Özellikle dilin özgün gramer yapısı ve eklemeli morfolojik özelliklerini dikkate alarak eğitilmiş olması, Türkçe içeriklerde yüksek doğruluk oranları sunmasını sağlar.

BÖLÜM IV

ARAŞTIRMA BULGULARI VE YORUMLAR

Bu çalışmada kullanılan her bir sınıflandırma modeli, yığın boyutu (batch size) 16 ve 3 epoch olacak şekilde eğitilmiştir. Eğitim sürecinde doğruluk (accuracy) ve kayıp (loss) değerleri dikkate alınarak hiper parametreler deneysel olarak belirlenmiştir. Modelin öğrenme eğrileri; eğitim ve doğrulama doğrulukları ile kayıp değerleri temelinde sistematik biçimde izlenmiş, her bir modelin öğrenme dinamiği bu metrikler üzerinden değerlendirilmiştir. Modelin eğitim sürecine doğrudan etki eden temel hiper parametreler ve bu parametrelerin seçim gerekçeleri Tablo 3'te sunulmuştur. Bu değerler hem literatürde önerilen aralıklar hem de ön denemelerde elde edilen sonuçlar doğrultusunda belirlenmiştir.

Hiper parametre optimizasyon sürecinde Grid Search, Random Search veya Bayesian Optimization gibi sistematik algoritmalar kullanılmamıştır. Bunun yerine, hiper parametreler literatürde yaygın olarak önerilen değer aralıkları esas alınarak ve sınırlı sayıda kontrollü deneme yapılarak manuel olarak belirlenmiştir. Örneğin, yığın boyutu için 8, 16 ve 32 gibi farklı değerler test edilmiş; 16 değeri eğitim süresi, istikrar ve model başarımı açısından en uygun sonucu verdiği için tercih edilmiştir. Benzer şekilde, epoch sayısının artırılması doğrulama kaybında yükselmeye neden olduğundan eğitim süreci 3 epoch ile sınırlandırılmıştır. Öğrenme oranı olarak ise Transformer mimarilerinde önerilen $4e-5$ değeri esas alınmıştır.

Sonuç olarak, bu çalışmada hiper parametreler algoritmik yöntemler yerine literatür destekli ön kabuller ve deneysel gözlemler doğrultusunda manuel olarak optimize edilmiştir. Bu yaklaşım hem bilimsel temellere dayanmaktadır hem de uygulama açısından pratiklik sağlamaktadır.

Tablo 3.

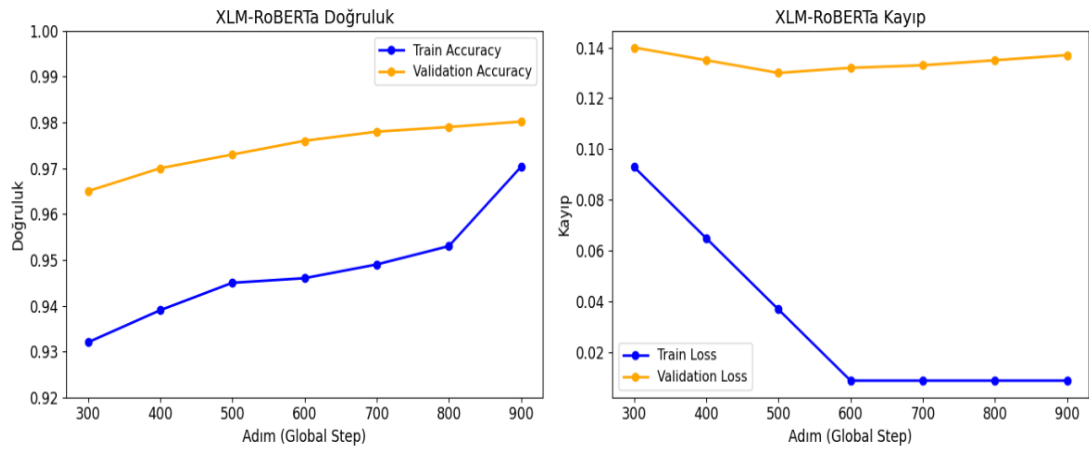
Modelde Kullanılan Hiper Parametreler

Parametre Adı	TurkishBERT	Bert Base Multilingual	XLNet
learning_rate	4e-5	4e-5	4e-5
train_batch_size	16	16	16
num_train_epochs	3	3	3

Parametre Adı	TurkishBERT	Bert Base Multilingual	XLM-RoBERTa
evaluate_during_training	True	True	True
evaluation_strategy	"steps"	"steps"	"steps"
save_eval_checkpoints	False	False	False
save_model_every_epoch	False	False	False
manual_seed	42	42	42
eval_metric	"accuracy"	"accuracy"	"accuracy"

Eğitim sürecinde modelin doğrulama performansını izlemek amacıyla evaluate_during_training parametresi etkinleştirilmiş ve değerlendirmelerin belirli adımlarda yapılabilmesi için evaluation_strategy değeri "steps" olarak tanımlanmıştır. Böylece, doğrulama verisi üzerindeki başarı düzenli olarak takip edilmiş ve aşırı öğrenme belirtileri erken aşamada fark edilebilmiştir. Modelin her adımda veya epoch sonunda kaydedilmesi ek kaynak tüketimine yol açtığı için save_eval_checkpoints ve save_model_every_epoch parametreleri False olarak ayarlanmıştır. Bu yaklaşım, kaynak kullanımını optimize ederek yalnızca nihai model çıktısına odaklanılmasına imkân tanımıştır. Deneylerin tekrarlanabilirliğini sağlamak amacıyla tüm çalışmalarda sabit bir tohum değeri (manual_seed = 42) kullanılmıştır. Sınıflandırma başarımının değerlendirilmesinde temel metrik olarak doğruluk oranı tercih edilmiştir.

XLM-RoBERTa modeline ait eğitim süreci performansı aşağıdaki şekil 8’de sunulmaktadır. Sol tarafta eğitim ve doğrulama doğrulukları, sağ tarafta ise eğitim ve doğrulama kayıpları karşılaştırmalı olarak verilmiştir. Bu grafikler aracılığıyla modelin öğrenme eğrisi ve genelleme kabiliyeti detaylı şekilde analiz edilmiştir.



Şekil 8. XLM-RoBERTa modeli grafikleri

Modelin eğitim doğruluğu, yaklaşık %91,50 seviyesinden başlayarak düzenli bir artış göstermiş ve 900. adıma ulaşıldığında %97,22 civarına ulaşmıştır. Doğrulama doğruluğu ise bu artışa paralel bir eğilim izlemekle birlikte, daha yavaş bir ivmeyle ilerlemiş ve yaklaşık %97,22 düzeyinde sabitlenmiştir. Bu durum, modelin eğitim verisine yüksek düzeyde uyum sağladığını ancak doğrulama verisi üzerinde benzer ölçüde bir başarı elde edemediğini göstermektedir.

Kayıp değerleri açısından değerlendirildiğinde, eğitim kaybının eğitim sürecinin ilk aşamalarında yatay bir seyir izlediği, yaklaşık 400. adımdan itibaren ise belirgin bir düşüş eğilimi gösterdiği ve nihayetinde sıfıra yakın değerlere ulaştığı gözlemlenmektedir. Doğrulama kaybı ise başlangıçta nispeten sabit bir düzeyde kalmış; ancak 600. adımdan sonra artış göstermeye başlamıştır. Bu durum, modelin eğitim verisine fazla uyum sağlamaya başladığını, yani aşırı öğrenme (overfitting) eğilimine girdiğini göstermektedir. Bu nedenle, modelin genelleme performansını koruyabilmek için eğitim süresinin dikkatli bir şekilde sınırlandırılması gerekmektedir.

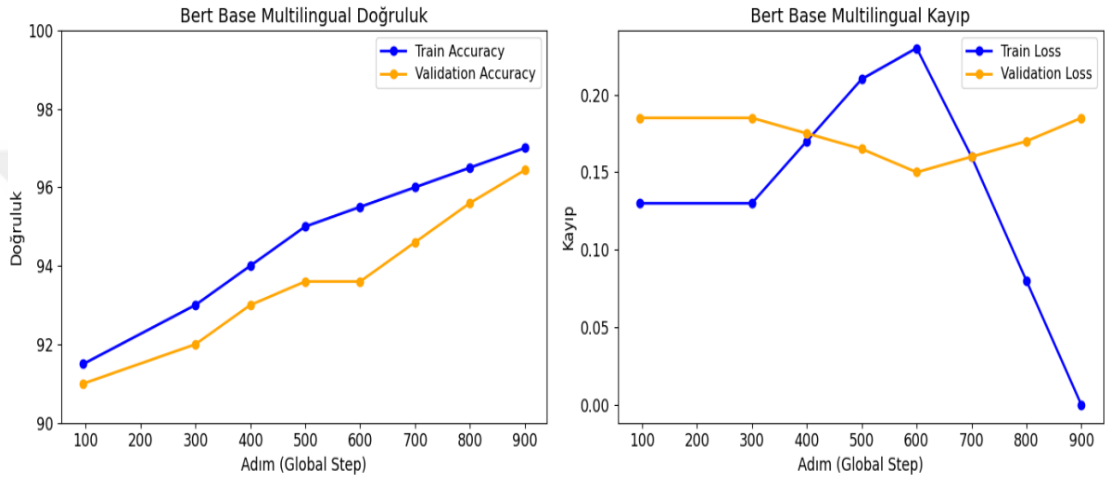
Şekil 8’de yer alan “Adım (Step)” ifadesi, modelin her bir mini yığın veri üzerinde gerçekleştirdiği ağırlık güncellemelerini temsil etmektedir. Bu bağlamda, adım ile epoch kavramları birbirinden farklıdır. Epoch, eğitim verisinin tamamının bir kez modele sunulması sürecidir. Buna karşılık, adım, her bir yığın modele verilmesiyle birlikte gerçekleştirilen tekil ağırlık güncellemesini ifade etmektedir.

Bu çalışmada kullanılan toplam veri sayısı 6.000’dir. Verilerin %80’i eğitim verisi olarak ayrılmış, geri kalan %20’si test verisi olarak kullanılmıştır. Eğitim sırasında yığın değeri 16 olarak belirlenmiştir. Bu doğrultuda, her bir epoch sürecinde model yaklaşık 300 adım gerçekleştirmiştir. Model, 3 epoch boyunca eğitildiğinden, eğitim süreci toplamda yaklaşık 900 adımda tamamlanmıştır. Grafiklerde gösterilen adım sayıları, bu toplam güncellemeleri yansıtmaktadır. Epoch = 3 değeri, bu toplam adım sayısının belirlenmesinde temel belirleyici olmuş; modelin ne kadar süre eğitileceği bu parametre aracılığıyla kontrol edilmiştir. Sonuç olarak, grafiklerdeki adımlar, modelin zaman içindeki öğrenme davranışını, doğruluk artışlarını ve olası aşırı öğrenme eğilimlerini ortaya koyma açısından kritik önem taşımaktadır.

XLM-RoBERTa modeli hem doğruluk hem de kayıp değerleri açısından en istikrarlı performansı sergileyen model olmuştur. Öğrenme eğrilerinin paralellliği, doğrulama başarımının yüksekliği ve düşük kayıp değerleri modelin dengeli bir biçimde

eğitildiğini ve genellenebilir çıktılar üretebildiğini göstermektedir. Bu yönüyle, çok dilli metin sınıflandırma görevleri için son derece uygun bir model profili sunmaktadır.

Şekil 9’da BERT Base Multilingual modelinin eğitim sürecine ilişkin doğruluk ve kayıp değerlerinin değişimi sunulmaktadır. Sol tarafta eğitim ve doğrulama doğrulukları, sağ tarafta ise eğitim ve doğrulama kayıpları yer almaktadır. Grafikler, modelin eğitim sürecinde nasıl bir öğrenme eğrisi izlediğini ve genelleme kapasitesini görselleştirmektedir.



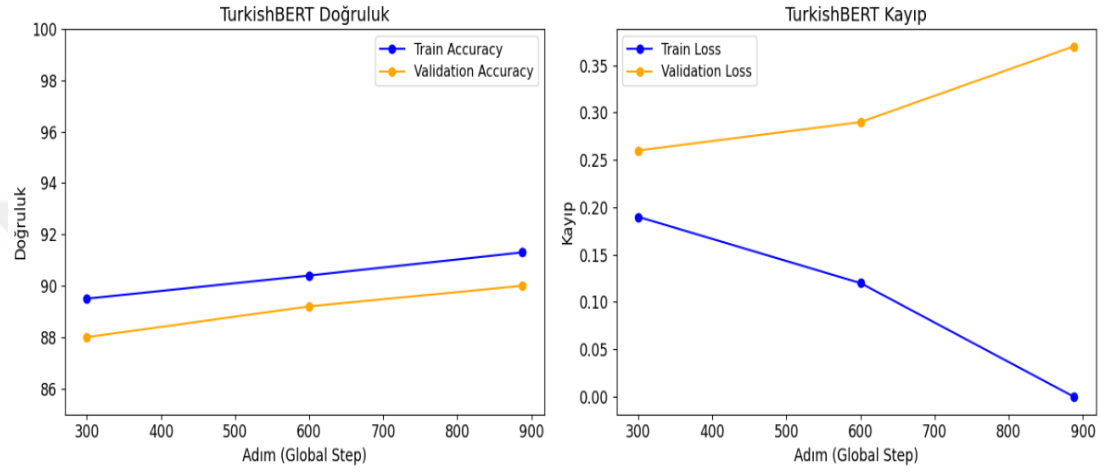
Şekil 9. Bert Base Multilingual modeli grafikleri

Şekil 9’da eğitim doğruluğunda gözlemlenen artış oldukça düzenlidir; yaklaşık %91,50 seviyesinde başlayan doğruluk, 900. adıma ulaşıldığında %96,40 seviyesine ulaşmıştır. Doğrulama doğruluğu da genel anlamda benzer bir artış eğilimi gösterse de eğitim doğruluğuna kıyasla daha yavaş artmış ve yaklaşık %96,40 seviyesinde dengelenmiştir. Bu durum, modelin eğitim verisine güçlü biçimde uyum sağladığını, ancak doğrulama verisinde tam olarak aynı başarıyı gösteremediğini ortaya koymaktadır.

Kayıp değerleri açısından modelin davranışı daha kritik bilgiler sunmaktadır. Eğitim kaybı başlangıçta yatay bir seyir izlemekte, 400. adımdan sonra belirgin biçimde azalarak sıfıra yaklaşmaktadır. Ancak doğrulama kaybı aynı şekilde azalmamış, 600. adımdan sonra artış göstermeye başlamıştır. Bu eğilim, modelin belirli bir noktadan sonra eğitim verisine fazla uyum sağladığını, yani aşırı öğrenme eğilimi gösterdiğini işaret etmektedir. Genel olarak model, kısa sürede yüksek doğruluk elde etme

kapasitesine sahip olsa da doğrulama kaybındaki yükselme, eğitim süresinin dikkatle sınırlandırılması gerektiğini göstermektedir.

TurkishBERT modeline ait eğitim süreci performansı aşağıdaki Şekil 10'da sunulmaktadır. Eğitim ve doğrulama doğrulukları ile kayıp değerlerinin birlikte incelenmesi, modelin öğrenme davranışını ve genelleme kapasitesini kapsamlı biçimde değerlendirme olanağı sağlamaktadır.



Şekil 10. *TurkishBERT modeli grafikleri*

Şekil 10'da eğitim doğruluğu, 300. adımdan başlayarak 900. adıma kadar düzenli ancak sınırlı bir artış göstermiştir. Başlangıçta %90,22 seviyelerinde olan doğruluk, eğitim sonunda %91,20 seviyelerine ulaşmıştır. Doğrulama doğruluğu ise benzer şekilde %88,33 civarından başlayarak %90,22 seviyesine ulaşmıştır. Bu değerler, modelin öğrenme sürecinin sınırlı düzeyde etkili olduğunu, ancak diğer modellere kıyasla daha düşük doğruluk seviyelerine ulaşıldığını göstermektedir.

Kayıp grafikleri bu durumu daha net biçimde ortaya koymaktadır. Eğitim kaybı her adımdan sonra düzenli biçimde azalırken, doğrulama kaybı tersine bir eğilim göstermektedir. Başlangıçta yaklaşık 0.26 olan doğrulama kaybı, eğitim ilerledikçe sürekli artarak 0.37 seviyesine çıkmıştır. Bu durum, modelin eğitim verisine gittikçe daha fazla uyum sağladığını, ancak doğrulama verisinde genelleme başarısının zayıfladığını ve öğrenme sorunu yaşandığını göstermektedir.

TurkishBERT modeli, sınırlı doğruluk artışına karşın doğrulama kaybındaki yükseliş nedeniyle, eğitim sürecinde dengeli bir performans sergileyememiştir.

4.1. Modelin Sınıflandırma İşleyişi ve Tahmin Süreci

Uygulanan sınıflandırma modeli, test aşamasında sunulan akademik metinleri analiz ederek her bir metni ait olduğu kategoriye atamaktadır. Model, sınıflandırma işlemini üç ana sınıf üzerinden gerçekleştirmektedir: Gerçek Veri, Düzenlenmiş ve Yapay Zekâ Verisi. Bu süreçte, modelin temel hedefi; dilbilgisel yapılar, sözdizimi, kelime tercihleri ve metin bütünlüğü gibi dilsel ve yapısal unsurları dikkate alarak, içeriklerin üretim kaynaklarına ilişkin ayrımları tespit etmektir.

Model, test veri kümesinde yer alan her metin için hem tahmin edilen sınıf etiketini hem de metnin gerçek etiketini dikkate almaktadır. Bu sayede modelin başarımı, doğru sınıflandırma yapma oranı üzerinden değerlendirilebilmektedir. Özellikle kelime tercihi, cümle yapısı, bağlaç kullanımı ve metin içi tutarlılık gibi özellikler, modelin karar verme sürecinde belirleyici olmaktadır.

Aşağıda yer alan Tablo 4'te sınıflandırma modelinin test verileri üzerinde yaptığı tahmin sonuçlarını özetlemektedir. Görselde her bir metne ait üç temel bilgi sunulmaktadır: metnin gerçek etiketi (hedef sınıf), model tarafından yapılan tahmin ve ilgili metin içeriği. Böylece her bir örnek üzerinden modelin sınıflandırma doğruluğu açık biçimde gözlemlenebilmektedir. Tahmin edilen sınıfın gerçek etiket ile örtüşüp örtüşmediği doğrudan karşılaştırılarak, modelin başarılı ya da hatalı sınıflandırmaları kolaylıkla ayırt edilebilmektedir. Modelin test verisi üzerinde gerçekleştirdiği sınıflandırma süreci ve örnek tahmin çıktıları.

Tablo 4.

Modelin Sınıflandırma İşleyişi ve Tahmin Sonuçları

Gerçek Etiket	Model Tahmini	Metin
Düzenlenmiş	Düzenlenmiş	Bu çalışma, Türkçe öğretmenlerinin sosyal medya kullanımlarını incelemektedir.
Düzenlenmiş	Düzenlenmiş	Bu araştırma, okul kütüphaneleri ve sınıf kitaplıklarının öğrenci başarısı üzerindeki etkisini araştırmaktadır.
Gerçek Veri	Gerçek Veri	Bu araştırma, ilkökul 4. sınıfların artırılmış gerçeklik uygulamalarıyla öğrenme süreçlerine etkisini analiz etmektedir.
Yapay Zekâ	Yapay Zekâ	Bu araştırma, ortaokul şirketlerinin yapay zekâ tabanlı muhasebe çözümlerine yaklaşımını ele almaktadır.

Yapay Zekâ	Yapay Zekâ	Okul öncesi dönemdeki çocuklarda değer davranışlarının gelişimi yapay zekâ destekli oyunlarla incelenmiştir.
Gerçek Veri	Gerçek Veri	Alternatif tedarikçiler arasında en uygun olanların seçimi yapılmıştır.
Yapay Zekâ	Yapay Zekâ	Tarihi merkezler, kentlerin fiziksel ve sosyokültürel yapılarıyla birlikte ele alınmıştır.
Düzenlenmiş	Düzenlenmiş	Otomobillerde kullanılan ön salıncak kolu amortisör sistemleri üzerine yapılmış bir analizdir.
Düzenlenmiş	Düzenlenmiş	Kuantum ağ yapısı mimari bakımdan klasik bilgisayar sistemlerinden ayrılmaktadır.
Gerçek Veri	Gerçek Veri	İnme, fiziksel, mental ve algısal problemlere yol açmakta ve bu durum rehabilitasyon sürecini etkilemektedir.

4.2. Model Bazlı Sınıflandırma Sonuçları

Bu çalışmada, üretken yapay zekâ araçları tarafından oluşturulmuş akademik içeriklerin tespiti amacıyla geliştirilen sınıflandırma sistemi kapsamında, derlenen Türkçe akademik özetlerden oluşan veri seti üzerinde çeşitli derin öğrenme tabanlı modeller uygulanmıştır. Oluşturulan veri seti; insan üretimi içerikler “Gerçek Veri”, doğrudan yapay zekâ ile üretilmiş metinler “Yapay Zekâ” ve yapay zekâ destekli düzenlemelerden geçmiş içerikler “Düzenlenmiş” olmak üzere üç sınıfa ayrılmıştır.

Veri seti, model eğitim ve değerlendirme sürecinde kullanılmak üzere eğitim %80 ve test %20 olmak üzere ikiye ayrılmıştır. Sınıflandırma görevinde, doğal dil işleme alanında yüksek başarı göstermeleriyle bilinen üç farklı önceden eğitilmiş model kullanılmıştır: XLM-Roberta, TurkishBERT ve Base Bert Base Multilingual. Çalışmada kullanılan tüm model bazlı sınıflandırma sonuçları numaralar halinde yazılmıştır.

4.2.1. XLM-Roberta sınıflandırma sonuçları. Bu çalışmada XLM-Roberta modeli kullanılarak gerçekleştirilen sınıflandırma deneylerinde, modelin farklı sınıflar üzerindeki performansı doğruluk, kesinlik, duyarlılık ve F1 skoru metrikleriyle değerlendirilmiştir. Üç sınıf için elde edilen sonuçlar, modelin yapay zekâ tarafından oluşturulan içerikleri, düzenlenmiş metinleri ve gerçek insan üretimi içerikleri başarıyla ayırt ettiğini göstermektedir.

Gerçek veri sınıfında model, doğruluk %97,22 kesinlik %95,21 ve duyarlılık %97,30 oranları ile %96,26 F1 skoru elde etmiş; bu da modelin gerçek içerikleri yüksek doğrulukla tanımlayabildiğini ortaya koymaktadır. Düzenlenmiş içerikler için ise doğruluk %96,20 kesinlik %95,23 duyarlılık %95,23 ve %95,36 F1 skoru gözlemlenmiştir. Bu sonuçlar, modelin yapay zekâ destekli düzenleme süreçlerinden geçmiş metinleri ayırt etmede de oldukça başarılı olduğunu göstermektedir. En yüksek başarı ise yapay zekâ sınıfında elde edilmiştir; model, bu sınıfa ait içerikleri doğruluk %100, kesinlik %100, duyarlılık %100 ve %100 F1 skoru ile kusursuz bir şekilde sınıflandırmıştır.

Genel olarak XLM-Roberta modelinin genel doğruluk değeri %97,22 olarak hesaplanmıştır. Bu yüksek performans, modelin akademik metinlerde yapay zekâ kaynaklı içeriklerin tespitinde güvenilir bir araç olduğunu göstermektedir. Sınıf bazlı sonuçlar, modelin hem özgün hem de yapay içeriklerin ayırt edilmesinde dengeli ve tutarlı performans sergilediğini ortaya koymaktadır. Sınıflandırmalar sonucu elde edilen bulgular Tablo 5’de verilmiştir.

Tablo 5. XLM-Roberta Sonucu

XLM-Roberta Sonucu

Sınıf	Doğruluk	Kesinlik	Duyarlılık	F1 Skoru
Gerçek Veri	97,22	95,21	97,30	96,26
Düzenlenmiş	96,20	95,23	95,23	95,36
Yapay Zekâ	100	100	100	100
Genel Doğruluk				97,22

4.2.2. Bert Base Multilingual sınıflandırma sonuçları. Bert Base Multilingual modeli ile yapılan sınıflandırma, modelin özellikle yapay zekâ sınıfında doğruluk %100, kesinlik %100, duyarlılık ve F1 skoru ile en yüksek performansı sergilediği gözlemlenmiştir. Gerçek veri sınıfında model doğruluk %94,20 kesinlik %94,40 duyarlılık %97,20 ile %96,44 F1 skoru elde ederken, düzenlenmiş veri sınıfında ise doğruluk %92,30 kesinlik %97,21 duyarlılık %93,20 ve %95,25 F1 skoru hesaplanmıştır. Bu sonuçlar, modelin genel olarak yüksek ayırt ediciliğe sahip olduğunu ve sınıflar arasında dengeli bir performans sergilediğini göstermektedir.

Özellikle “Gerçek Veri” sınıfında yüksek duyarlılık, modelin gerçek yazıları büyük oranda doğru şekilde tespit ettiğini ortaya koymaktadır. Modelin genel doğruluk değeri %96,40 olarak ölçülmüştür. Bu durum, Bert Base Multilingual çok dilli ortamlarda bile istikrarlı ve güçlü bir sınıflandırma kapasitesine sahip olduğunu göstermektedir. Sınıflandırmalar sonucu elde edilen bulgular Tablo 6’da verilmiştir.

Tablo 6.

Bert Base Multilingual Sonucu

Sınıf	Doğruluk	Kesinlik	Duyarlılık	F1 Skoru
Gerçek Veri	94,20	94,40	97,20	96,44
Düzenlenmiş	92,30	97,21	93,20	95,25
Yapay Zekâ	100	100	100	100
Genel Doğruluk				96,40

4.2.3. TurkishBERT sınıflandırma sonuçları. TurkishBERT modeli, yalnızca Türkçe dilinde eğitilmiş bir dil modeli olarak, bu çalışmada dilin yapısal özelliklerine özgü derinlikli analiz yeteneği ile dikkat çekmektedir. Ancak çok dilli modellere kıyasla genel başarı oranı nispeten düşüktür. Gerçek veri sınıfında doğruluk %91,20 ile, kesinlik %83,10 duyarlılık %92,31 ile %86,21 F1 skoru, düzenlenmiş sınıfında doğruluk %81,22 kesinlik %90,21 duyarlılık %79,60 ile %82,42 F1 skoru elde edilmiştir. Yapay zekâ sınıfında ise model, diğer modellerde olduğu gibi %100 başarı göstermiştir. Modelin genel doğruluk %90,22 olarak hesaplanmıştır. Bu sonuçlar, TurkishBERT’in özellikle düzenlenmiş" veri sınıfında duyarlılık açısından zorlandığını göstermektedir. Düşük duyarlılık, modelin bu tür verileri ayırt ederken bazı örnekleri gözden kaçırdığını ortaya koymaktadır. Bununla birlikte, TurkishBERT’in Türkçeye özgü dil yapısını daha iyi anlayabilmesi, gerçek veri sınıfında yüksek duyarlılık ile kendini göstermektedir. Sınıflandırmalar sonucu elde edilen bulgular Tablo 7’de verilmiştir.

Tablo 7.

TurkishBERT Sonucu

Sınıf	Doğruluk	Kesinlik	Duyarlılık	F1 Skoru
Gerçek Veri	91,20	83,10	92,31	86,21
Düzenlenmiş	81,22	90,21	79,60	82,42
Yapay Zekâ	100	100	100	100
Genel Doğruluk				90,22

4.2.4. Bert Base Multilingual, XLM-Roberta ve Turkishbert modellerine ait sınıflandırma sonuçlarının karşılaştırılması. Bu çalışmada karşılaştırılan üç farklı dil modeli BERT Base Multilingual, XLM-RoBERTa ve TurkishBERT üzerinden yapılan sınıflandırma performansları, özellikle genel doğruluk ölçütüyle değerlendirilmiştir. Genel doğruluk, modelin tüm sınıflardaki örnekleri doğru tahmin etme oranını yansıtır ve bu açıdan modellerin genel sınıflandırma başarısı hakkında güçlü bir gösterge sunar. XLM-RoBERTa modeli, %97,22 genel doğruluk oranı ile en yüksek başarıyı göstermiştir. Bu durum, modelin hem çok dilli bağlamları etkili biçimde öğrenebildiğini hem de sınıflar arasındaki ayrımı yüksek isabetle gerçekleştirebildiğini ortaya koymaktadır. BERT Base Multilingual modeli ise %96,40 genel doğruluk ile yine oldukça güçlü bir performans sergilemiştir. Her iki model de yapay zekâ sınıfında %100 doğruluk ile kusursuz sonuçlar üretmiş ve sahte içeriklerin tanımlanmasında oldukça güvenilir olduklarını kanıtlamıştır.

Gerçek veri ve düzenlenmiş sınıflarında ise her iki model arasında küçük farklılıklar gözlemlense de sonuçlar genel olarak dengelidir. XLM-RoBERTa, bu sınıflarda da yüksek duyarlılık ve kesinlik değerleriyle dikkat çekmektedir.

Türkçe dilinde eğitilmiş olan TurkishBERT modeli, %90,22 genel doğruluk oranı ile diğer iki çok dilli modele kıyasla daha düşük performans sergilemiştir. Bu durum, özellikle “Düzenlenmiş” sınıfta düşük duyarlılık ile ilişkilendirilebilir. TurkishBERT’in Türkçe’nin morfolojik yapısına olan duyarlılığı, belirli avantajlar sağlasa da bağlamsal çeşitliliği yakalama konusunda sınırlı kaldığı gözlemlenmiştir.

Sonuç olarak, çok dilli büyük modeller özellikle XLM-RoBERTa ve BERT Base Multilingual metin türleri arasında ayırım yapmada daha yüksek doğruluk sunmakta ve

sahte içerik tespiti gibi uygulamalarda daha uygun alternatifler olarak öne çıkmaktadır. Tablo 8’te BERT Base Multilingual, XLM-RoBERTa ve TurkishBERT modellerine ait sınıflandırma sonuçlarının karşılaştırmalı tablosu yer almaktadır.

Tablo 8.

Bert Base Multilingual, XLM-Roberta ve Turkishbert Sonuç Karşılaştırma Tablosu.

Model	Sınıf	Doğruluk	Kesinlik	Duyarlılık	F1 Skoru
XLM-RoBERTa	Gerçek Veri	97,22	95,21	97,30	96,26
	Düzenlenmiş	96,20	95,23	95,23	95,36
	Yapay Zekâ	100	100	100	100
	Genel Doğruluk				97,22
Bert Base Multilingual	Gerçek Veri	94,20	94,40	97,20	96,44
	Düzenlenmiş	92,30	97,21	93,20	95,25
	Yapay Zekâ	100	100	100	100
	Genel Doğruluk				96,40
TurkishBERT	Gerçek Veri	91,20	83,10	92,31	86,21
	Düzenlenmiş	81,22	90,21	79,60	82,42
	Yapay Zekâ	100	100	100	100
	Genel Doğruluk				90,22

BÖLÜM V

TARTIŞMA, SONUÇ VE ÖNERİLER

5.1. Tartışma

Bu tez çalışması, üretken yapay zekâ tarafından oluşturulan akademik metinlerin içerik tabanlı analizle tespit edilmesine odaklanmıştır; bu amaçla transformer tabanlı üç ayrı model BERT Base Multilingual, XLM-RoBERTa ve TurkishBERT kullanılarak yüksek hassasiyetli bir sınıflandırma sistemi geliştirilmiştir.

Özellikle XLM-RoBERTa modelinin %97,22 genel doğruluk, BERT Base Multilingual modelinin ise yaklaşık %96,40 doğruluk oranı elde etmesi, çok dilli ve geniş bağlamsal öğrenme kapasitesine sahip modellerin, akademik metinlerde yapay zekâ üretimi ile insan üretimi içerikler arasındaki ayrımı yüksek başarıyla gerçekleştirebildiğini ortaya koymaktadır. Bu doğruluk oranları, her iki modelin sınıflar arasında dengeli, istikrarlı ve güvenilir bir şekilde sınıflandırma yaptığına işaret etmektedir.

Buna karşılık yalnızca Türkçe verilerle önceden eğitilmiş olan TurkishBERT modeli, %90,22 genel doğruluk oranı ile daha sınırlı bir performans ortaya koymuştur. Bu sonuç, özellikle modelin düzenlenmiş sınıfında daha düşük başarı göstermesiyle ilişkilidir. Ayrıca bu durum, Türkçenin eklemeli yapısı ve morfolojik zenginliği göz önüne alındığında, tek bir dile odaklanan modellerin bağlamsal çeşitliliği yeterince yakalayamaması nedeniyle dezavantajlı olabileceğine işaret etmektedir.

Model performansları karşılaştırıldığında, her üç modelin yapay zekâ etiketi altındaki içerikleri %100 doğrulukla sınıflandırabilmesi, bu içeriklerin belirgin biçimsel ve anlamsal imzalar taşıdığını; ancak düzenlenmiş sınıfında gözlenen başarı düşüşü, insan müdahalesinin yapay zekâ çıktısını daha doğal kılarak tespit zorluğunu artırdığını ortaya koymaktadır. Bu bulgu, gelecekte daha karmaşık, semantik tabanlı ve dilsel nüansları hesaba katan ek yöntemlerin geliştirilmesi gerekliliğini vurgulamaktadır.

Çalışmanın özgün yönlerinden biri, araştırmacı tarafından tamamen yenilikçi biçimde oluşturulmuş 6.000 adet Türkçe akademik özet veri setinin kullanılmasıdır. Buna

karşılık, bu tez çalışması bağlamsal yoğunluğu yüksek akademik özetler üzerinde üçlü sınıflandırma sunarak metodolojik derinliği artırmıştır. Tablo 9’da mevcut çalışmada elde edilen sonuçları, farklı yöntem ve veri kümeleriyle literatürde bildirilen performanslarla yan yana getirmektedir. Özellikle Bert Base Multilingual ve XLM-RoBERTa’nın, ikili sınıflandırmada kullanılan modellere kıyasla üç sınıflı yapıda dahi benzer veya daha yüksek başarı sergilemesi, çok dilli bağlamsal modellerin genel geçerliğini desteklemektedir.

Tablo 9.

Bulguların Literatürdeki Çalışmalar ile Karşılaştırılması

Referanslar	Yöntem / Model	Veri Seti / İçerik	Doğruluk / Performans Oranı
Ayan vd. (2021)	NLP + Anahtar kelime çıkarımı	Şirket web siteleri	Doğruluk belirtilmemiş
Ozan & Taşar (2021)	BERT	14.000 kısa diyalog	%98 BERT
Sevli & Kemaloğlu (2021)	Google BERT	7613 tweet (afet içerikli)	%98,88
Hark (2022)	Derin bağlamsal gömü + LSTM	COVID-19 sahte haber veri seti (10.700 örnek)	%91
Toğaçar vd. (2022)	NLP + LSTM	6335 haber (3171 gerçek, 3164 sahte)	Eğitim: %99,83 / Test: %91,48
Özkan & Kar (2022)	Türkçe BERT	Türkçe bilimsel metinler	%96
Bayram (2022)	Makine öğrenmesi + NLP	Pandemi içerikleri	İstatistiksel analiz, doğruluk oranı yok
Yazgılı & Baykara (2022)	Bagging, Boosting, klasik yöntemler	3000 Türkçe sosyal medya gönderisi (siber zorbalık)	Karşılaştırmalı performans raporu
Altıntop (2023)	ChatGPT analizi	Yorumlayıcı/teorik çalışma	Performans verisi yok
Koçyiğit & Darı (2023)	ChatGPT üzerinden sosyolojik çözümleme	Kavramsal tartışma	Performans verisi yok
Yiğit vd. (2023)	ChatGPT sağlık uygulamaları	Akademik metin yazımı, sağlık içeriği	Teorik katkı
Görgün vd. (2023)	BERT, FastText vb.	1040 Twitter kullanıcısı	Anlamli göstergeler / doğruluk belirtilmemiş
Gao vd. (2023)	ChatGPT	50 tıbbi makale özeti	%68 ChatGPT özetleri doğru tespit edildi
Kim (2023)	ChatGPT + İnsan doğrulaması	Bilimsel makale yazımı	Doğruluk gerektiği vurgulanmış
Amini vd. (2023)	NLP + Otomatik transkripsiyon	Nöropsikolojik test ses kayıtları	%92+ doğruluk
Feng & Chen (2024)	AdbGPT	Yazılım hata raporları	%81,3 başarı / Ortalama süre: 253.6s

Referanslar	Yöntem / Model	Veri Seti / İçerik	Doğruluk / Performans Oranı
Daungsupawong & Wiwanitkit (2024)	LLM'ler genel analiz	Sağlık hizmetleri	Etik analiz, doğruluk oranı yok
Bu Çalışma	Bert Base Multilingual, XLM-RoBERTa, TurkishBERT	Akademik içeriklerde yapay zekâ üretimi tespiti	Bert: %96,40XLM-RoBERTa: %97,22TurkishBERT: %90,22

Literatürde yapılan çeşitli çalışmalar, farklı doğal dil işleme ve makine öğrenmesi teknikleriyle metin sınıflandırma süreçlerinde yüksek başarı oranlarına ulaşıldığını göstermektedir. Örneğin, Sevlı ve Kemaloğlu (2021), afet temalı tweet'lerin sınıflandırılmasında Google BERT modelini kullanarak %98,88 doğruluk oranı bildirmiştir. Toğaçar vd. (2022) ise sahte haber verisi üzerinde LSTM tabanlı sistem ile test aşamasında %91,48 doğruluk elde etmiştir. Özkan ve Kar (2022), Türkçe bilimsel metinler üzerinde %96 başarı sağlayarak, BERT tabanlı modellerin akademik içerik analizinde etkili olduğunu ortaya koymuştur. Öte yandan, Gao vd. (2023), ChatGPT tarafından oluşturulan tıbbi özetlerin yalnızca %68'inin doğru şekilde tespit edilebildiğini vurgulayarak üretken yapay zekâ içeriklerinin tespitinde zorluklara dikkat çekmiştir. Bu çalışma ise üç sınıflı yapısıyla çok daha karmaşık bir yapay içerik analizine olanak tanımakta ve bu yönüyle literatüre anlamlı bir katkı sunmaktadır.

Sonuç olarak, yürütölen araştırma yalnızca teknik olarak güçlü sınıflayıcılar geliştirmekle kalmamış; aynı zamanda akademik yayıncılıkta yapay zekâ destekli içeriklerin denetiminde somut bir araç sunmuştur. Geliştirilen sistem, üniversiteler, hakemli dergiler ve araştırma kurumları tarafından bir ön inceleme aracı olarak kullanılmaya uygundur. Bununla birlikte, farklı diller, üretken yapay zekâ modelleri ve tam metinler üzerinde ek çalışmalar yapılması; sınıflandırma performansını daha da artırmak ve sistemi genelledebilmek adına önerilmektedir. Bu tez çalışması, yapay zekâ üretimi ile insan üretimi akademik metinleri ayırt edebilen bir sınıflandırma modeli geliştirmeyi hedeflemektedir. Ancak çalışmanın kapsamı ve uygulama alanı bazı sınırlılıklarla çevrilidir. Öncelikle, kullanılan veri kümesi tamamen Türkçe akademik özetlerden oluşmaktadır. Bu durum, geliştirilen modelin diğeri dil yapılarına ve tam metinlere genellenebilirliğini kısıtlamaktadır.

İkinci olarak, yapay metinler inceleme süreci sadece belirli üretken yapay zekâ araçları (ChatGPT, Gemini, Smodin) aracılığıyla üretilmiştir. Bu nedenle, farklı modellerden

gelen içeriklerde modelin performansı test edilmemiştir. Bu araştırmanın sınırları dahilinde sınıflandırıcı modellerin eğitim verileri belirli disiplinlerle (sağlık, mühendislik, eğitim, sosyal bilimler) sınırlı tutulmuştur. Bu nedenle, elde edilen bulguların genellenebilirliği belirli ölçüde sınırlıdır. Daha kapsamlı ve disiplinden bağımsız sonuçlara ulaşılabilmesi için, gelecekte gerçekleştirilecek çalışmalarda tam metin analizlerine dayalı, farklı alanları kapsayan daha geniş veri kümeleriyle modelin performansının değerlendirilmesi önerilmektedir.

5.2. Sonuç

Bu tez çalışmasında, akademik metinlerde yapay zekâ katkılı içeriklerin tespitine yönelik bir sınıflandırma sistemi geliştirilmiş ve bu sistem, Türkçe dilinde hazırlanmış özgün bir veri seti üzerinde test edilmiştir. Günümüzde üretken dil modellerinin özellikle ChatGPT, Smodin ve Gemini gibi yaygınlaşmasıyla birlikte, insan üretimi ile yapay zekâ destekli içerikler arasındaki farkın belirlenmesi giderek zorlaşmakta; bu durum, akademik etik, içerik güvenilirliği ve yayın politikaları açısından ciddi riskler doğurmaktadır. Çalışmanın temel amacı, bu belirsizliği ortadan kaldıracak yüksek doğruluklu bir sınıflandırma modeli oluşturmaktır.

Bu kapsamda geliştirilen veri seti, 6.000 adet Türkçe akademik özetten oluşmakta; gerçek, yapay ve düzenlenmiş içerikler olmak üzere üç sınıfa ayrılmıştır. Sınıflandırma sürecinde XLM-RoBERTa, Bert Base Multilingual ve TurkishBERT olmak üzere üç farklı transformer tabanlı model kullanılmıştır. Elde edilen sonuçlara göre, XLM-RoBERTa modeli %97,22 genel doğruluk ile en yüksek başarıyı sağlamıştır. Bert Base Multilingual modeli %96,40 ve TurkishBERT modeli ise genel olarak %90,22 oranında başarıya ulaşmış; ancak düzenlenmiş içeriklerde görece düşük performans göstermiştir.

5.3. Öneriler

Çalışmanın en önemli avantajlarından biri, tamamen araştırmacı tarafından oluşturulan özgün ve çok sınıflı Türkçe veri setine dayanmasıdır. Sadece Türkçe içeriklere odaklanması sebebiyle, çalışmanın çok dilli veri kümeleriyle desteklenerek farklı dil yapıları ve kültürel bağlamlarda uygulanabilir hale getirilmesi önerilmektedir. Ayrıca,

geliştirilen sınıflandırma modelinin üniversitelerde intihal denetim süreçlerine entegre edilmesi; etik değerlendirme kurulları tarafından kullanılacak modüllerle desteklenmesi ve öğrencilere yönelik eğitim amaçlı geri bildirim sistemlerine dönüştürülmesi önemli bir potansiyel taşımaktadır. Bu bağlamda, çalışma yalnızca teknik bir katkı sunmakla kalmayıp, akademik içerik denetimi, etik ilkelerin güçlendirilmesi ve akademik dürüstlüğün tesis edilmesi açısından da anlamlı bir katkı sağlamaktadır. Geliştirilen sistemin, üniversiteler, akademik yayın kurulları ve araştırma merkezleri tarafından bir ön denetim aracı olarak kullanılması mümkündür. İlerleyen dönemlerde farklı diller, metin türleri ve model mimarileri ile bu sistemin genellenebilirliğinin artırılması önerilmektedir.



KAYNAKLAR

- Acikalin, U. U., Geminiak, B., & Kutlu, M. (2020). Turkish sentiment analysis using bert. *2020 28th signal processing and communications applications conference (SIU)*, 1-4.
- Adalı, E. (2016). Doğal dil işleme. *Türkiye Bilişim Vakfı Bilgisayar Bilimleri ve Mühendisliği Dergisi*, 5(2), Article 2.
- Altıntop, M. (2023). Yapay zekâ/akıllı öğrenme teknolojileriyle akademik metin yazma: Chatgpt örneği. *Süleyman Demirel Üniversitesi Sosyal Bilimler Enstitüsü Dergisi*, 46, Article 46.
- Amini, S., Hao, B., Zhang, L., Song, M., Gupta, A., Karjadi, C., Kolachalama, V. B., Au, R., & Paschalidis, I. C. (2023). Automated detection of mild cognitive impairment and dementia from voice recordings: A natural language processing approach. *Alzheimer's & Dementia*, 19(3), 946-955.
- August, T., Wang, L. L., Bragg, J., Hearst, M. A., Head, A., & Lo, K. (2023). Paper plain: Making medical research papers approachable to healthcare consumers with natural language processing. *ACM Transactions on Computer-Human Interaction*, 30(5), 1-38.
- Ayan, E. T., Arslan, R., Zengin, M. S., Duru, H. A., Salman, S., & Geminiak, B. (2021). Turkish keyphrase extraction from web pages with BERT. *2021 29th Signal Processing and Communications Applications Conference (SIU)*, 1-4. <https://doi.org/10.1109/SIU53274.2021.9477842>
- Bayram, U. (2022). Revealing the reflections of the pandemic by investigating COVID-19 Related news articles using machine learning and network analysis. *Bilişim Teknolojileri Dergisi*, 15(2), Article 2. <https://doi.org/10.17671/gazibtd.949599>
- Bin Shiha, R., Atwell, E., & Abbas, N. (2023). Detecting Bias in university news articles: A comparative study using BERT, GPT-3.5 and Google Gemini annotations. *International Conference on Innovative Techniques and Applications of Artificial Intelligence*, 487-492.

- Bird, J. J., Ekárt, A., & Faria, D. R. (2023). Chatbot Interaction with Artificial Intelligence: Human data augmentation with T5 and language transformer ensemble for text classification. *Journal of Ambient Intelligence and Humanized Computing*, 14(4), 3129-3144.
- Brown, P. F., Cocke, J., Della Pietra, S. A., Della Pietra, V. J., Jelinek, F., Lafferty, J., Mercer, R. L., & Roossin, P. S. (1990). A statistical approach to machine translation. *Computational linguistics*, 16(2), 79-85.
- Covington, MA, Grosz, BJ, & Pereira, FC (1994). *Prolog programcıları için doğal dil işleme* . Upper Saddle River: Prentice hall.
- Chen, J.-M. (2024). Performance Assessment of ChatGPT vs Gemini in Detecting Alzheimer's Dementia. *arXiv preprint arXiv:2402.01751*.
- Çetindağ, C., Yazıcıoğlu, B., & Koç, A. (2023). Named-entity recognition in Turkish legal texts. *Natural Language Engineering*, 29(3), 615-642.
- Daungsupawong, H., & Wiwanitkit, V. (2024). Natural Language Processing in vascular surgery. *EJVES Vascular Forum*, 61, 2.
- Deguchi, H., Shibata, K., & Taguchi, S. (2024). Language to Map: Topological map generation from natural language path instructions. *2024 IEEE International Conference on Robotics and Automation (ICRA)*, 9556-9562.
- Dong, C., Li, Y., Gong, H., Chen, M., Li, J., Shen, Y. ve Yang, M. (2022). Doğal dil üretiminin bir araştırması. *ACM Hesaplama Araştırmaları*, 55 (8), 1-38.
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). *BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding* (No. arXiv:1810.04805). arXiv. <https://doi.org/10.48550/arXiv.1810.04805>
- Feng, S., & Chen, C. (2024). Prompting is all you need: Automated android bug replay with large language models. *Proceedings of the 46th IEEE/ACM International Conference on Software Engineering*, 1-13.
- Gao, C. A., Howard, F. M., Markov, N. S., Dyer, E. C., Ramesh, S., Luo, Y., & Pearson, A. T. (2023). Comparing scientific abstracts generated by ChatGPT to real abstracts with detectors and blinded human reviewers. *NPJ digital medicine*, 6(1), 75.

- Görgün, M. K., Demirok, G. B., & Kutlu, M. (2023). Türkçe sosyal medya mesajlarından kullanıcıların yaş ve cinsiyetini tahmin etme. *Niğde Ömer Halisdemir Üniversitesi Mühendislik Bilimleri Dergisi*, 12(2), 325-333.
- Hark, C. (2022). Sahte haber tespiti için derin bağlamsal kelime gömülmeleri ve sinirsel ağların performans değerlendirilmesi. *Fırat Üniversitesi Mühendislik Bilimleri Dergisi*, 34(2), 733-742.
- Hsu, J. C., Wu, M., Kim, C., Vora, B., Lien, Y. T., Jindal, A., Yoshida, K., Kawakatsu, S., Gore, J., & Jin, J. Y. (2024). Applications of advanced natural language processing for clinical pharmacology. *Clinical Pharmacology & Therapeutics*, 115(4), 786-794.
- Hutchins, W. J. (2004). The Georgetown-IBM experiment demonstrated in January 1954. *Conference of the Association for Machine Translation in the Americas*, 102-114.
- İşeri, İ., Aydın, Ö., & Tutuk, K. (2021). Müşteri hizmetleri yönetiminde yapay zekâ temelli chatbot geliştirilmesi. *Avrupa Bilim ve Teknoloji Dergisi*, 29, 358-365.
- Jiang, B., Chen, X., Liu, W., Yu, J., Yu, G., & Chen, T. (2023). Motiongpt: Human motion as a foreign language. *Advances in Neural Information Processing Systems*, 36, 20067-20079.
- Kanan, T., Mughaid, A., Al-Shalabi, R., Al-Ayyoub, M., Elbes, M., & Sadaqa, O. (2023). Business intelligence using deep learning techniques for social media contents. *Cluster Computing*, 26(2), 1285-1296.
- Kim, S.-G. (2023). Using ChatGPT for language editing in scientific articles. *Maxillofacial plastic and reconstructive surgery*, 45(1), 13.
- Koçyiğit, A., & Darı, A. B. (2023). Yapay zekâ iletişimde Chatgpt: İnsanlaşan dijitalleşmenin geleceği. *Stratejik ve Sosyal Araştırmalar Dergisi*, 7(2), Article 2. <https://doi.org/10.30692/sisad.1311336>
- Marcus, M., Santorini, B., & Marcinkiewicz, M. A. (1993). Building a large annotated corpus of English: The Penn Treebank. *Computational linguistics*, 19(2), 313-330.

- Mujahid, M., Kanwal, K., Rustam, F., Aljedaani, W., & Ashraf, I. (2023). Arabic ChatGPT tweets classification using RoBERTa and BERT ensemble model. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 22(8), 1-23.
- Munikaar, M., Shakya, S., & Shrestha, A. (2019). Fine-grained Sentiment Classification using BERT. *2019 Artificial Intelligence for Transforming Business and Society (AITB)*, 1, 1-5. <https://doi.org/10.1109/AITB48515.2019.8947435>
- Mutinda, J., Mwangi, W., & Okeyo, G. (2023). Sentiment analysis of text reviews using lexicon-enhanced bert embedding (LeBERT) model with convolutional neural network. *Applied Sciences*, 13(3), 1445.
- Müller, M., Salathé, M., & Kummervold, P. E. (2023). Covid-twitter-bert: A natural language processing model to analyse covid-19 content on twitter. *Frontiers in artificial intelligence*, 6, 1023281.
- Ozan, Ş., & Taşar, D. E. (2021). Auto-tagging of short conversational sentences using natural language processing methods. *2021 29th Signal Processing and Communications Applications Conference (SIU)*, 1-4. <https://doi.org/10.1109/SIU53274.2021.9477994>
- 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, H., Değirmenci, A., Güngör, O., & Uskudarlı, S. (2020). The role of contextual word embeddings in correcting the ‘de/da’ clitic errors in Turkish. *2020 28th Signal Processing and Communications Applications Conference (SIU)*, 1-4. <https://doi.org/10.1109/SIU49456.2020.9302477>
- Rabiner, L. R. (1989). A tutorial on hidden Markov models and selected applications in speech recognition. *Proceedings of the IEEE*, 77(2), 257-286.
- Saleh, H., Alhothali, A., & Moria, K. (2023). Detection of hate speech using bert and hate speech word embedding with deep model. *Applied Artificial Intelligence*, 37(1), 2166719.

- Salvagno, M., Taccone, F. S., & Gerli, A. G. (2023). Artificial intelligence hallucinations. *Critical Care*, 27(1), 180. <https://doi.org/10.1186/s13054-023-04473-y>
- Sevli, O., & Kemaloğlu, N. (2021). Olağandışı olaylar hakkındaki tweet'lerin gerçek ve gerçek dışı olarak Google BERT modeli ile sınıflandırılması. *Veri Bilimi*, 4(1), Article 1.
- Talan, T., & Aktürk, C. (2021). *Bilgisayar Bilimlerinde Teorik Ve Uygulamalı Araştırmalar*. Efe Akademi Yayınları.
- Tejaswini, V., Sathya Babu, K., & Sahoo, B. (2024). Depression detection from social media text analysis using natural language processing techniques and hybrid deep learning model. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 23(1), 1-20.
- Toğaçar, M., Eşidir, K. A., & Ergen, B. (2022). Yapay zekâ tabanlı doğal dil işleme yaklaşımını kullanarak internet ortamında yayınlanmış sahte haberlerin tespiti. *Journal of Intelligent Systems: Theory and Applications*, 5(1), Article 1. <https://doi.org/10.38016/jista.950713>
- Turing, A. M. (1950). *Mind*, 59(236), 433-460.
- Uzun, L. (2023). ChatGPT and academic integrity concerns: Detecting artificial intelligence generated content. *Language Education and Technology*, 3(1).
- Wang, Z., Cheng, J., Cui, C., & Yu, C. (2023). Implementing BERT and fine-tuned RobertA to detect AI generated news by ChatGPT. *arXiv preprint arXiv:2306.07401*.
- Weizenbaum, J. (1966). ELIZA—a computer program for the study of natural language communication between man and machine. *Communications of the ACM*, 9(1), 36-45.
- Xu, M., Chen, H., Deng, F., Liu, Y., Zhao, F., & Zhang, Y. (2024). Practical applications of ChatGPT and BERT: Long text summarization and model feature ablation study for stock comment topic identification. *Proceedings of the 2024 Guangdong-Hong Kong-Macao Greater Bay Area International Conference on Digital Economy and Artificial Intelligence*, 787-793.

- Yang, D., Liu, S., Huang, R., Weng, C., & Meng, H. (2024). Instructtts: Modelling expressive tts in discrete latent space with natural language style prompt. *IEEE/ACM Transactions on Audio, Speech, and Language Processing*.
- Yazgılı, E., & Baykara, M. (2022). Türkçe metinlerde makine öğrenmesi yöntemleri ile siber zorbalık tespiti. *Gümüşhane Üniversitesi Fen Bilimleri Dergisi*, 12(2), 443-453.
- Yiğit, S., Berşe, S., & Dirgar, E. (2023). Yapay zekâ destekli dil işleme teknolojisi olan ChatGPT'nin sağlık hizmetlerinde kullanımı. *Eurasian Journal of Health Technology Assessment*, 7(1), Article 1. <https://doi.org/10.52148/ehta.1302.000>
- Zhang, Z., & Wang, B. (2023). Prompt learning for news recommendation. *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval*, 227-237.
- Zong, M., & Krishnamachari, B. (2023). Solving math word problems concerning systems of equations with gpt-3. *Proceedings of the AAAI conference on artificial intelligence*, 37(13), 15972-15979.

