



T.C.
KONYA TEKNİK ÜNİVERSİTESİ
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ



TÜRKÇE E-TİCARET YORUMLARININ ÇOK
ETİKETLİ ANALİZİ İÇİN DERİN ÖĞRENME
MODELLERİNİN UYGULANMASI

Abdulkadir ŞEN

YÜKSEK LİSANS TEZİ

Bilgisayar Mühendisliği Anabilim Dalı

Ocak-2025
KONYA
Her Hakkı Saklıdır

TEZ KABUL VE ONAYI

Abdulkadir ŞEN tarafından hazırlanan “TÜRKÇE E-TİCARET YORUMLARININ ÇOK ETİKETLİ ANALİZİ İÇİN DERİN ÖĞRENME MODELLERİNİN UYGULANMASI” adlı tez çalışması 08/01/2025 tarihinde aşağıdaki jüri tarafından oy birliği ile Konya Teknik Üniversitesi Lisansüstü Eğitim Enstitüsü Bilgisayar Mühendisliği Anabilim Dalı’nda YÜKSEK LİSANS TEZİ olarak kabul edilmiştir.

Jüri Üyeleri

İmza

Başkan

Doç. Dr. Emine BAŞ
(Konya Teknik Üniversitesi)

.....

Danışman

Dr. Öğr. Üyesi Fatma Zehra SOLAK
(Konya Teknik Üniversitesi)

.....

Üye

Dr. Öğr. Üyesi Ayşe Merve ACILAR
(Necmettin Erbakan Üniversitesi)

.....

Yukarıdaki sonucu onaylarım.

Prof. Dr. Mevlüt UYAN
Enstitü Müdürü

TEZ BİLDİRİMİ

Bu tezdeki bütün bilgilerin etik davranış ve akademik kurallar çerçevesinde elde edildiğini ve tez yazım kurallarına uygun olarak hazırlanan bu çalışmada bana ait olmayan her türlü ifade ve bilginin kaynağına eksiksiz atıf yapıldığını bildiririm.

DECLARATION PAGE

I hereby declare that all information in this document has been obtained and presented in accordance with academic rules and ethical conduct. I also declare that, as required by these rules and conduct, I have fully cited and referenced all material and results that are not original to this work.

Abdulkadir ŞEN
Tarih: 08.01.2025

ÖZET

YÜKSEK LİSANS TEZİ

TÜRKÇE E-TİCARET YORUMLARININ ÇOK ETİKETLİ ANALİZİ İÇİN DERİN ÖĞRENME MODELLERİNİN UYGULANMASI

Abdulkadir ŞEN

**Konya Teknik Üniversitesi
Lisansüstü Eğitim Enstitüsü
Bilgisayar Mühendisliği Anabilim Dalı**

Danışman: Dr. Öğr. Üyesi Fatma Zehra SOLAK

2025, 106 Sayfa

Jüri

Dr. Öğr. Üyesi Fatma Zehra SOLAK

Doç. Dr. Emine BAŞ

Dr. Öğr. Üyesi Ayşe Merve ACILAR

E-ticaret, dijitalleşen dünyanın hızla büyüyen sektörlerinden biri olarak öne çıkmaktadır. Tüketicilerin çevrimiçi alışveriş kararlarını etkileyen en önemli unsurlardan biri, diğer kullanıcıların deneyimlerini paylaştığı müşteri yorumlarıdır. Ancak, bu yorumların manuel olarak analiz edilmesi, büyük veri setleri için oldukça zaman alıcı ve maliyetli bir süreçtir. Bu çalışma, e-ticaret platformlarından elde edilen müşteri yorumlarının yapay zekâ tabanlı sınıflandırma yöntemleri kullanılarak hızlı ve doğru bir şekilde işlenmesini amaçlamaktadır. Türkiye'nin önde gelen e-ticaret platformlarından toplanan müşteri yorumları, "Fiyat Performans", "İade Edilen", "Paketleme İyi, Hızlı Teslimat", "Kalitesiz, Kusurlu, Kötü Paketleme" ve "Tavsiye Edilen, Kaliteli Ürünler" gibi çok etiketli sınıflandırma kategorilerine ayrılmıştır. Bu kategoriler, yorumların yalnızca içeriklerini değil, aynı zamanda müşterilerin olumlu ya da olumsuz duygu durumlarını da yansıtarak çalışmanın duygu analizi alanına önemli katkılar sunmasını sağlamıştır. Veri saklama ve yönetimi için Microsoft SQL Server kullanılmış, Selenium Web Driver ve Html Agility Pack araçlarıyla veri kazıma işlemleri gerçekleştirilmiştir. Yorumlar doğal dil işleme süreçlerinde TF-IDF temsiliyle işlenmiş ve LSTM, CNN, GRU, RNN, BiLSTM ve BERT gibi yapay sinir ağları tabanlı derin öğrenme modelleri ve Lojistik Regresyon, Rastgele Orman gibi geleneksel makine öğrenmesi algoritmaları ile sınıflandırılmıştır. Performans değerlendirmelerinde doğruluk, duyarlılık, kesinlik, F1-Skor, log loss ve konfüzyon matrisi gibi metrikler kullanılmıştır. Özellikle BERT modeli, bağlamsal anlam çıkarma yetenekleriyle büyük veri setlerinde doğru sınıflandırma ve anlam çıkarma süreçlerinde diğer yöntemlere göre üstün bir performans sergilemiştir. BERT'in, transformer tabanlı yapısıyla dilin bağlamını çift yönlü analiz etme yeteneği, metinlerin anlamını derinlemesine kavrayarak daha doğru ve anlamlı sınıflandırmalar yapmasına olanak sağlamıştır. Bu çalışmada kullanılan yapay zekâ modelleri, müşteri yorumlarını sınıflandırmada yüksek doğruluk oranlarına ulaşmış ve geliştirilen yapay zekâ destekli sistemin, tüketicilerin bilinçli kararlar almasına rehberlik ettiği ve işletmelerin pazarlama ile ürün stratejilerini iyileştirmelerine katkı sağladığı ortaya konulmuştur. BERT'in üstün performansı, özellikle çok etiketli sınıflandırma görevlerinde, dilin bağlamını doğru anlamadaki başarısı ile dikkat çekmiştir. Bu kapsamda, çalışma e-ticaret sektöründe müşteri yorumlarının analizine yönelik yenilikçi ve etkili bir çözüm sunmaktadır.

Anahtar Kelimeler: E-Ticaret, Derin Öğrenme, Duygu Analizi, Veri Madenciliği, Web Kazıma

ABSTRACT

MS THESIS

APPLYING DEEP LEARNING MODELS FOR MULTI-LABEL ANALYSIS OF TURKISH E-COMMERCE COMMENTS

Abdulkadir ŞEN

**Konya Technical University
Institute of Graduate Studies
Department of Computer Engineering**

Advisor: Asst. Prof. Dr. Fatma Zehra SOLAK

2025, 106 Pages

Jury

**Asst. Prof. Dr. Fatma Zehra SOLAK
Assoc. Prof. Dr. Emine BAŞ
Asst. Prof. Dr. Ayşe Merve ACILAR**

E-commerce stands out as one of the rapidly growing sectors in the increasingly digital world. One of the most important factors influencing consumers' online shopping decisions is customer reviews, where users share their experiences. However, manually analyzing these reviews is a time-consuming and costly process for large datasets. This study aims to process customer reviews obtained from e-commerce platforms quickly and accurately using artificial intelligence-based classification methods. Customer reviews collected from leading e-commerce platforms in Turkey have been categorized into multi-label classification categories such as "Price-Performance", "Returned", "Good Packaging, Fast Delivery", "Low Quality, Defective, Bad Packaging", and "Recommended, Quality Products". These categories have contributed significantly to the field of sentiment analysis by reflecting not only the content of the reviews but also the positive or negative emotional states of the customers. Microsoft SQL Server was used for data storage and management, and data scraping was carried out using Selenium Web Driver and Html Agility Pack tools. The reviews were processed with TF-IDF representation in natural language processing tasks and classified using deep learning models based on artificial neural networks, such as LSTM, CNN, GRU, RNN, BiLSTM, and BERT, as well as traditional machine learning algorithms like Logistic Regression and Random Forest. Performance evaluation metrics such as accuracy, sensitivity, precision, F1-score, log loss, and confusion matrix were used. Notably, the BERT model outperformed other methods in large datasets, demonstrating superior performance in classification and context extraction due to its contextual understanding capabilities. BERT's transformer-based architecture, with its ability to analyze language context bidirectionally, allowed for a deeper understanding of text meaning and more accurate and meaningful classifications. The artificial intelligence models used in this study achieved high accuracy rates in classifying customer reviews, and it was demonstrated that the developed AI-supported system guided consumers in making informed decisions and contributed to improving marketing and product strategies for businesses. BERT's superior performance, particularly in multi-label classification tasks, stands out for its success in understanding language context accurately. In this context, the study provides an innovative and effective solution for analyzing customer reviews in the e-commerce sector.

Keywords: E-Commerce, Deep Learning, Sentiment Analysis, Data Mining, Web Scraping

ÖNSÖZ

Tez konusunun belirlenmesi ve yürütülmesindeki katkıları, çalışma sürecinde göstermiş olduğu ilgi, destek ve anlayışı ile değerli hocam Dr. Öğr. Üyesi Fatma Zehra Solak’a,

Akademik çalışmaya teşvik ederek, eğitimim boyunca beni destekleyen aileme, çalışmamın her aşamasında en büyük destekçim olan eşim Hüsna Çetin Şen’e sonsuz teşekkürlerimi sunarım.

Abdulkadir ŞEN
KONYA-2025



İÇİNDEKİLER

ÖZET	iv
ABSTRACT.....	v
ÖNSÖZ	vi
İÇİNDEKİLER	vii
SİMGELER VE KISALTMALAR	ix
1. GİRİŞ	1
1.1. Tezin Amacı ve Önemi	2
2. KAYNAK ARAŞTIRMASI	5
3. MATERYAL VE YÖNTEM.....	10
3.1. Veri Saklama ve Yönetimi.....	10
3.2. Web Kazıma	10
3.3. Web Sitesi Araçları.....	11
3.4. Doğal Dil İşleme	12
3.4.1. Veri Ön İşleme.....	13
3.4.2. Kelime Temsili	13
3.5. Makine Öğrenmesi.....	16
3.5.1. CNN (Evrişimsel Sinir Ağları - Convolutional Neural Network)	18
3.5.2. RNN (Tekrarlayan sinir ağları - Recurrent Neural Network).....	19
3.5.3. LSTM (Uzun Kısa Süreli Bellek - Long Short-Term Memory).....	20
3.5.4. BiLSTM (Çift Yönlü Uzun Kısa Süreli Bellek - Bidirectional LSTM)	21
3.5.5. GRU (Kapı Özyinelemeli Geçitler - Gated Recurrent Unit)	22
3.5.6. BERT (Dönüştürücülerden Çift Yönlü Kodlayıcı Temsiller - Bidirectional Encoder Representations from Transformers)	23
3.5.7. Lojistik Regresyon.....	24
3.5.8. Rastgele Orman.....	26
3.7. Performans Gösterme Metrikleri	27
3.8. Çalışma Ortamı ve Paketler	32
3.8.1. Python Kütüphaneleri ve Kullanımları	33
4. ARAŞTIRMA SONUÇLARI VE TARTIŞMA.....	35
4.1. Web Kazıma Projesi, Verilerin Toplanması ve Saklanması.....	35
4.2. Verilerin Etiketlenmesi ve Veri Seti Oluşturulması	39
4.3. Veri Setinin Dengelenmesi	46
4.4. Keşifsel Veri Analizi	46
4.5. Veri Ön İşleme.....	50
4.6. Veri Modelleme	51

4.7. Yapay Zeka Modelleri	52
4.7.1. CNN	54
4.7.2. RNN	57
4.7.3. LSTM	60
4.7.4. BiLSTM	64
4.7.5. GRU	68
4.7.6. BERT	71
4.7.7. Lojistik Regresyon	76
4.7.8. Rastgele Orman	80
4.8. E-Ticaret Yorum Analiz Uygulaması	83
5. SONUÇLAR VE ÖNERİLER	90
KAYNAKLAR	93



SİMGELER VE KISALTMALAR

Simgeler

E	: Doğal logaritmanın tabanı olan Euler sabiti (≈ 2.71)
FP	: Toplam pozitif tahminler
FN	: Yanlış negatif tahminler
f(t,d)	: Belirli bir terim (t), belirli bir belgede (d) kaç kez yer alan sayıdır
Log	: Doğal logaritma işlemi
N	: Toplam veri sayısı
y_i	: Gerçek sınıf etiketleri
$\hat{y}_i$	: Modelin <i>i</i> -inci örnek için tahmin ettiği olasılık
sig(t)	: Lojistik regresyon sonucunda elde edilen olasılık değeri
t	: Bağımsız değişkenin belirli bir değeri
TP	: Doğru pozitif tahminler
TN	: Doğru negatif tahminler

Kısaltmalar

BERT	: Dönüştürücülerden Çift Yönlü Kodlayıcı Temsiller (Bidirectional Encoder Representations from Transformers)
BiLSTM	: Çift Yönlü Uzun Kısa Süreli Bellek (Bidirectional Long Short-Term Memory)
CNN	: Evrişimsel Sinir Ağı (Convolutional Neural Network)
GPU	: Grafik İşlem Birimi (Graphic Processing Unit)
GRU	: Kapı Özyinelemeli Geçitler (Gated Recurrent Unit)
HTML	: Hipermetin İşaretleme Dili (HyperText Markup Language)
LSTM	: Uzun Kısa Süreli Bellek (Long Short-Term Memory)
MVC	: Model-Görünüm-Denetleyici (Model-View-Controller)
MSSQL	: Microsoft SQL Server
NLP	: Doğal Dil İşleme (Natural Language Processing)
NLTK	: Doğal Dil Araç Takımı (Natural Language Toolkit)
RDBMS	: İlişkisel Veritabanı Yönetim Sistemi (Relational Database Management System)
RNN	: Özyinelemeli Sinir Ağı (Recurrent Neural Network)
SQL	: Yapılandırılmış Sorgu Dili (Structured Query Language)
TPU	: Tensör İşlem Birimi (Tensor Processing Unit)
T-SQL	: İşlemsel-SQL (Transact-SQL)
TVT	: Eğitim - Doğrulama - Test (Train-Validation-Test)
TF-IDF	: Terim Frekansı - Ters Doküman Frekansı (Term Frequency- Inverse Document Frequency)
URL	: Tekdüzen Kaynak Bulucu (Uniform Resource Locator)

1. GİRİŞ

Günümüzde, e-ticaret sektörü hızla dijitalleşen dünyada merkezi bir rol oynamaktadır. İnternetin sunduğu erişim ve kolaylık, tüketicilerin alışveriş tercihlerini büyük ölçüde şekillendirmektedir. Bu bağlamda, online alışveriş yapan müşterilerin kararlarını etkileyen en önemli unsurlardan biri, ürün ve hizmetlere dair diğer tüketicilerin deneyimleridir. Bu deneyimlerin en belirgin göstergesi ise e-ticaret sitelerinde yer alan müşteri yorumlarıdır. Milyonlarca kullanıcının farklı platformlarda paylaştığı bu yorumlar, alıcıların ürün ve marka tercihlerinde kritik bir rol oynamaktadır. Ancak, bu devasa veri setlerini elle sınıflandırmak ve analiz etmek zaman alıcı ve maliyetli bir süreçtir. Bu noktada, derin öğrenme tekniklerinin e-ticaret sitelerinden toplanan müşteri yorumlarını otomatik olarak sınıflandırma ve etiketleme konusunda sunduğu potansiyel büyük bir çözüm sunmaktadır.

Bu tez çalışması, e-ticaret müşteri yorumlarının yapay sinir ağları tabanlı makine öğrenmesi algoritmaları olan derin öğrenme modelleriyle analiz edilerek, önceden tanımlanmış çok sayıda etiket arasından en uygun olanına atanması üzerine odaklanmaktadır. Çalışmada, e-ticaret sitelerinden toplanan müşteri yorumlarının içerikleri, otomatik etiketleme yöntemleri kullanılarak analiz edilecek ve her bir yorum, tanımlı etiketlerden birine atanacaktır. Bu kapsamda, derin öğrenme tekniklerinden yararlanılarak müşteri geri bildirimlerinin sınıflandırılması, yorumların anlamının daha iyi kavranması ve kullanıcı deneyiminin geliştirilmesi hedeflenmektedir.

Tez çalışmasında, müşteri yorumlarının analiz ve sınıflandırma süreçlerinde hem yapay sinir ağları tabanlı makine öğrenmesi yani derin öğrenme modellerinden hem de geleneksel makine öğrenmesi yöntemlerinden yararlanılmaktadır. Geleneksel makine öğrenmesi yöntemleri, sınırlı özelliklere dayanarak sınıflandırma yaparken, derin öğrenme modelleri özellikle doğal dil işleme (NLP) alanında daha gelişmiş ve güçlü sonuçlar sunmaktadır. Derin öğrenme, katmanlı yapılarıyla veri setlerini daha derinlemesine analiz edebilir, bağlamı daha doğru anlamak için gelişmiş algoritmalar kullanabilir. Özellikle Dönüştürücülerden Çift Yönlü Kodlayıcı Temsiller (Bidirectional Encoder Representations from Transformers- BERT) modeli, bağlamı çift yönlü olarak analiz ederek dilin anlamını daha doğru kavrayabilmektedir. BERT, metinlerin tüm bağlamını anlamak için dilin her yönünü dikkate alarak, sadece bir kelimenin ya da cümlelerin değil, tüm metnin anlamını çıkarabilir. Bu özellik, e-ticaret yorumlarının doğru

bir şekilde sınıflandırılmasını sağlayarak, daha doğru etiketler ve sonuçlar elde edilmesine olanak tanır.

Çalışma kapsamında, müşteri yorumlarının analiz ve sınıflandırma süreçlerinde derin öğrenme modelleri olarak Uzun Kısa Süreli Bellek (Long Short-Term Memory - LSTM), Evrişimsel Sinir Ağı (Convolutional Neural Network - CNN), Kapı Özyinelemeli Geçitler (Gated Recurrent Unit - GRU), Çift Yönlü Uzun Kısa Süreli Bellek (Bidirectional Long Short-Term Memory - BiLSTM) ve Özyinelemeli Sinir Ağı (Recurrent Neural Network - RNN) modelleri kullanılmıştır. Bu modeller, büyük veri setlerinin işlenmesi konusunda yüksek verimlilik sağlarken, özellikle doğal dil işleme alanında daha anlamlı ve doğrusal sonuçlar elde edilmesine katkı sağlar. Ancak, bu tezde kullanılan asıl güçlü model BERT'tir. BERT, dilin bağlamını çift yönlü şekilde analiz etme yeteneği sayesinde diğer derin öğrenme tekniklerinden daha üstün sonuçlar vermektedir. Bu, e-ticaret müşteri yorumlarının daha doğru ve etkili bir şekilde etiketlenmesi ve işletmelere değerli bilgiler sunulması anlamına gelmektedir.

Bunlara ek olarak, tez çalışmasında, geleneksel makine öğrenmesi algoritmalarına da yer verilmiştir. Lojistik Regresyon ve Rastgele Orman gibi yöntemler, daha temel fakat güçlü analizler yaparak doğrulama ve karşılaştırma amacıyla kullanılmaktadır. Ancak derin öğrenme teknikleri, geleneksel yöntemlerin ötesine geçerek büyük veri setlerini daha hassas bir şekilde analiz edebilmekte, anlamlı bilgilere dönüştürmektedir. Bu çalışma, müşteri yorumlarının derin öğrenme yöntemleri ile daha hızlı ve daha hassas bir şekilde sınıflandırılmasını sağlayarak, işletmelere veriye dayalı karar destek sistemleri oluşturmada katkı sunmayı hedeflemektedir.

1.1. Tezin Amacı ve Önemi

E-ticaret sektörü, dijitalleşmenin sunduğu olanaklar sayesinde çağımızın en hızlı büyüyen ticaret yöntemlerinden biri haline gelmiştir. Bu süreçte, tüketicilerin satın alma kararlarını etkileyen faktörlerin başında, ürün ve hizmetlerle ilgili müşteri deneyimlerine dayanan yorumlar gelmektedir. Milyonlarca kullanıcı tarafından oluşturulan bu büyük veri havuzu, işletmeler için önemli fırsatlar sunarken aynı zamanda verimli bir şekilde işlenmesi gereken bir zorluk teşkil etmektedir.

Bu tezin önemi, müşteri yorumlarının derin öğrenme modelleri kullanılarak çok kategoriden birine otomatik sınıflandırılmasına odaklanmasından kaynaklanmaktadır. Geleneksel yöntemlerle yapılması güç olan bu sınıflandırma işlemi, işletmelerin müşteri

memnuniyetini artırmak, ürünlerini geliştirmek ve hedef kitleye daha etkili bir şekilde ulaşmak için ihtiyaç duyduğu iç görüyü sağlamayı amaçlamaktadır. Aynı zamanda, kullanıcıların ihtiyaçlarına uygun ürünleri seçme sürecini kolaylaştırarak, tüketici tarafında da bir katma değer yaratmaktadır.

Bu çalışmanın temel amacı, e-ticaret platformlarından toplanan Türkçe müşteri yorumlarını derin öğrenme modelleri kullanarak otomatik olarak çoklu etiket analizi yapmak ve her bir yorumu belirli kategorilerden birine atamaktır. Yorumlar, "Fiyat Performans", "İade Edilen", "Paketleme İyi, Hızlı Teslimat", "Kalitesiz, Kusurlu Ürün, Kötü Paketleme" ve "Tavsiye Edilen, Kaliteli Ürünler" gibi etiketlerden birine atanarak sınıflandırılacaktır. Bu etiketleme süreci, yorumların içeriği üzerinden dolaylı olarak müşterilerin duygu durumlarının analizini de sağlamaktadır.

Bu bağlamda, Türkiye'nin önde gelen bir e-ticaret platformundan elde edilen müşteri yorumları üzerinde çalışılmıştır. Tez, hem teknik hem de stratejik bir çözüm sunmayı hedeflemekte, e-ticaret sektöründe veri analitiği ve yapay zeka uygulamalarına yenilikçi bir bakış açısı getirmektedir.

Çalışmanın başlıca hedefleri şunlardır:

- i. **Veri Toplama:** E-ticaret sitesinden müşteri yorumlarının web kazıma yöntemleriyle elde edilmesi.
- ii. **Manuel Etiketleme:** Yorumların web ortamında yayınlanarak gerçek kullanıcılar tarafından önceden belirlenmiş kategorilerden biri ile etiketlenmesinin sağlanması.
- iii. **Yapay Zeka Modelleri İçin Temsil:** Yorumları, derin öğrenme ve geleneksel makine öğrenmesi modelleri için vektörel olarak temsil etmek.
- iv. **Otomatik Sınıflandırma:** Derin öğrenme ve geleneksel makine öğrenmesi modelleri kullanarak müşteri yorumlarını belirli kategorilere otomatik olarak sınıflandırmak.
- v. **Çoklu Etiket Analizi:** Yorumların farklı kategorilere sınıflandırılmasıyla, dolaylı olarak müşteri yorumlarının genel fikrini ve duygu durumlarını analiz etmek. Her bir etiket, yorumun duygusal yönünü anlamaya yönelik bir işaret teşkil etmektedir.
- vi. **Model Performansının Değerlendirilmesi:** Derin öğrenme modellerinin (CNN, LSTM, BiLSTM, GRU, RNN) ve transformer tabanlı modellerin (BERT) müşteri yorumlarını daha doğru ve etkili bir şekilde etiketleme süreçlerindeki etkinliğini değerlendirmek.

- vii. ***Sonuçların Uygulanabilirliği:*** Otomatik sınıflandırma sonuçlarının, işletmeler için müşteri memnuniyeti ve ürün geliştirme süreçlerine sağladığı faydalar ile tüketicilere karar verme süreçlerindeki etkilerinin analiz edilmesi.

Sonuç olarak, bu tez, hem işletmelerin hem de tüketicilerin ihtiyaçlarına çözüm sunan bir araştırma olmayı hedeflemektedir. Derin öğrenme modellerinin sunduğu yenilikçi bakış açısıyla, e-ticaret platformlarının mevcut sorunlarını çözmeye katkıda bulunmaktadır.



2. KAYNAK ARAŞTIRMASI

Yapılan araştırmada, önceden gerçekleştirilen çalışmalar titizlikle gözden geçirilmiştir. Bu kapsamlı inceleme sürecinde, kullanılan çeşitli yöntemler, uygulanan algoritmalar ve elde edilen sonuçlar detaylı bir şekilde ele alınmıştır. Bu çalışmanın seçimi, kullanılan metodoloji, tezin hedefleri ve benimsenen tekniklerin daha geniş bir bakış açısıyla anlaşılmasını sağlamak amacıyla özenle gerçekleştirilmiştir.

Karakuş ve diğ. (2018), çalışmalarında, Türkçe film incelemelerine odaklanan bir veri setini kullanarak duygu analizi gerçekleştirmişlerdir. Bu analizlerde, CNN, LSTM ve bu iki yöntemin kombinasyonlarını kullanmışlardır. Ayrıca, Word2Vec tekniğini veri setlerine uygulayarak elde ettikleri sonuçlar üzerinden kelime gömme kullanımının derin öğrenme modelleri üzerindeki etkilerini detaylı bir şekilde ele almışlardır. Elde edilen bulgular, Türkçe film yorumlarına yönelik duygu analizinde farklı yöntemlerin ve tekniklerin etkili olabileceğini göstermektedir.

Rumelli ve diğ. (2019), e-ticaret ürün yorumları üzerinden bir duygu analizi modeli geliştirmişlerdir. Çalışmada, hepsiburada.com e-ticaret sitesinden toplam 272 bin adet ürün yorumu çekilerek bir veri seti oluşturulmuştur. Bu veri setinde, 1 ve 2 yıldız ile değerlendirilen yorumlar olumsuz, 3 yıldız olanlar nötr, 4 ve 5 yıldız alan yorumlar ise olumlu olarak etiketlenmiştir. Son aşamada, 13 bin olumlu ve 13 bin olumsuz yorum seçilerek çalışmada kullanılmak üzere bir veri seti hazırlanmıştır. Çalışmada, Naive Bayes, Rastgele Orman, Destek Vektör Makinesi ve K-En Yakın Komşuluk algoritmaları kullanılarak analizler gerçekleştirilmiştir. Yapılan çalışma sonucunda, ortalama %73 doğruluk oranına ulaşılmıştır. Bu bulgular, e-ticaret ürün yorumları üzerinden duygu analizi yapmak için çeşitli algoritmaların etkili olabileceğini göstermektedir.

Mengutayci ve Temurtas (2021), Türkçe otel yorumlarını yapay sinir ağıları kullanarak sınıflandırma üzerine odaklanan bir çalışma gerçekleştirmiştir. Bu araştırmada, çevrimiçi otel rezervasyonu yapılan bir platformda bulunan müşteri yorumları üzerinde detaylı bir inceleme yapmışlardır. Yapılan sınıflandırma, yorumların 4-5 yıldız alması durumunda olumlu, 1-2 yıldız alması durumunda ise olumsuz olarak etiketlenmesi üzerine odaklanmıştır, böylece iki sınıflı bir çalışma ortaya konmuştur. Veri seti üzerinde gerekli ön işlemleri gerçekleştirdikten sonra, yorumları TF-IDF modeli ile bir matrise dönüştürmüş ve her terimin sınıflandırmadaki önem derecesini belirlemişlerdir. Veri setini iki farklı yöntemle eğitim ve test verilerine ayırarak, bu veri

setlerinin başarısını değerlendirmişlerdir. Yapılan testler, çeşitli yapay sinir ağı mimarileri üzerinden gerçekleştirilmiştir.

Mayda ve Korkmaz (2018), müşteri yorumlarının sınıflandırılması için sözlük tabanlı bir yöntem kullanarak ortak bir çalışma gerçekleştirmişlerdir. Bu çalışmada, yazarlar tarafından özel olarak oluşturulan bir sıfat sözlüğü kullanılmıştır. Kitap satışları yapılan bir e-ticaret platformundan topladıkları okuyucu yorumları üzerinden bir veri seti hazırlanmış ve sıfat sözlüğü kullanılarak olumlu, olumsuz ve nötr sınıflar için bir duygu analizi çalışması yapılmıştır. Ayrıca, Türkçe diline özgü kurallar geliştirilerek çalışmanın başarısının artırılması amaçlanmıştır. Yapılan çalışmanın sonuçlarına göre, yaklaşık %62'lik bir doğruluk oranına ulaşılmıştır. Bu bulgular, sözlük tabanlı yaklaşımın müşteri yorumlarının duygu analizi için etkili bir yöntem olabileceğini göstermektedir.

Gezici ve Yanıkoğlu (2018), Türkçe film incelemelerinin duygu analizini yapmak amacıyla makine öğrenmesi algoritmalarını kullanmışlardır. Bu çalışmada, denetimli öğrenme ve sözlük tabanlı yaklaşımları birleştirerek özgün bir yöntem geliştirmişlerdir. Beyazperde adlı film sitesinden önceden oluşturulan bir veri seti üzerinde çalışmışlar ve bu veri setindeki incelemeleri, 4 ve 5 yıldız alanları olumlu, 1 ve 2 yıldız alanları ise olumsuz olarak etiketleyerek ikili sınıflandırma gerçekleştirmişlerdir. Naive Bayes ve Destek Vektör Makinesi algoritmalarının sonuçlarını karşılaştırmışlar ve bu karşılaştırma sürecinde sözlük tabanlı yaklaşım ile denetimli öğrenme yaklaşımlarının sonuçlarını da dikkate almışlardır. Yapılan analizler sonucunda, Naive Bayes algoritmasının %75 doğruluk oranına ulaştığı belirlenmiştir. Bu bulgular, denetimli öğrenme ve sözlük tabanlı yöntemlerin birleştirilmesinin Türkçe film incelemelerinin duygu analizi için etkili bir strateji olabileceğini göstermektedir.

Pervan ve Keleş (2017), Türkçe yorumlarda duygu analizi yapmak amacıyla Rastgele Orman algoritmasını kullanmışlardır. Çalışmalarında, popüler bir e-ticaret platformundan elektronik kategorisi altında bulunan ürünlere yapılan müşteri yorumlarından oluşan bir veri seti hazırlamışlardır. Etiketleme işlemini yıldız puanına göre otomatik olarak gerçekleştirerek, yorumları olumlu ve olumsuz olarak sınıflandırmışlardır. Rastgele Orman algoritması kullanılarak elde ettikleri sonuçlar, yaklaşık %85 doğruluk oranı ile başarılı bir duygu analizi gerçekleştirdiklerini göstermiştir. Bu çalışma, Rastgele Orman algoritmasının Türkçe yorumlardaki duygu analizi üzerinde etkili bir araç olabileceğini vurgulamaktadır.

Demircan ve diğ. (2021), Türkçe metinlerde duygu analizi yöntemlerini inceledikleri bir çalışmada, e-ticaret ürün yorumlarını veri seti olarak kullanmışlardır.

Yorumlardaki puanlardan faydalanarak veri etiketleme işlemi gerçekleştirmişlerdir. Bu çalışma kapsamında, Destek Vektör Makinesi, Rastgele Orman, Karar Ağacı, Lojistik Regresyon ve K - En Yakın Komşu algoritmalarını kullanarak modeller geliştirmiş ve bu modellerin testlerini gerçekleştirmişlerdir. Yapılan testlerin sonuçlarına göre, Destek Vektör Makinesi ve Rastgele Orman modellerinin daha iyi performans gösterdiği gözlemlenmiştir. Bu bulgular, Türkçe metinlerde duygu analizi için özellikle bu iki algoritmanın etkili olabileceğini işaret etmektedir.

Toçoğlu ve diğ. (2019), Türkçe tweetlerden duygu analizi üzerine bir çalışma gerçekleştirmişlerdir. Çalışma kapsamında oluşturdukları veri setini sözlük tabanlı bir yaklaşım kullanarak her tweet için altı duygu kategorisi için otomatik açıklamalar eklemiştir. Bu çalışmada, yapay sinir ağı, CNN, LSTM ve RNN mimarilerini inceleyerek bu mimariler üzerinden çeşitli testler gerçekleştirmişlerdir. Yapılan testlerin sonuçlarına göre, sözlük tabanlı yaklaşım ile otomatik eklenen açıklamaların etkilerini gözlemlemiştir. Karşılaştırma sonucunda, CNN mimarisinin %74'lük doğruluk puanı ile en yüksek performansı gösterdiği belirlenmiştir. Bu sonuçlar, Türkçe tweetler üzerinde duygu analizi için CNN mimarisinin etkili bir seçenek olabileceğini göstermektedir.

Çataltaş ve diğ. (2020), e-ticaret olumsuz ürün yorumlarından çıkarılan ürün kusurlarını belirlemek amacıyla bir metin analizi yöntemi önermişler ve bu kusurlardan yapılandırılmış bir özet oluşturmayı hedeflemiştir. Çalışma kapsamında, Amazon.com üzerinden ayakkabı kategorisine ait 10 bin olumsuz yorum çekerek bir veri seti oluşturmuşlardır. Yapılan ön işlemler sonucunda veri seti 6190 yoruma düşmüş ve çalışmayı bu verilerle gerçekleştirmişlerdir. Ürün kusurlarını belirlemek için DBSCAN algoritması kullanılarak yorumlar kümelerine ayrılmıştır. Her bir kusur, daha önceden tanımlanmış Konuşma Parçası (Part-of-Speech) kalıplarını arayarak kendi görüş ifadelerini bulmak için kullanılmıştır. Bu bulunan kusurlar ve onların görüş ifadeleri kullanılarak yapılandırılmış bir özet oluşturulmuştur. Bu çalışma, e-ticaret olumsuz yorumlarından çıkarılan kusurların analizi ve özetlenmesi konusunda etkili bir metin analizi yöntemi sunmaktadır.

Çabuk ve diğ. (2023), Ürün yorumları, web kazıma yöntemiyle toplanarak manuel bir sınıflandırma sürecinden geçirilmiştir, bu süreçte yorumlar olumlu, olumsuz ve nötr olarak sınıflandırılmıştır. Nötr değerler çıkarılarak, iki sınıflı bir veri seti oluşturulmuştur. İkinci aşamada, Uzun Kısa Süreli Bellek Ağları (LSTM) ile otomatik bir veri seti oluşturulmuştur. Daha sonra, manuel ve otomatik olarak oluşturulan veri setleri üzerinde, RNN, LSTM, GRU ve CNN gibi derin öğrenme algoritmaları kullanılarak duygu analizi

testleri yapılmıştır. Bu testler sonucunda, otomatik sonuçların manuel sınıflandırmalara göre yüksek bir başarı elde ettiği gözlemlenmiştir. Ardından, otomatik etiketlenmiş veri seti ile Rastgele Orman, Lojistik Regresyon, Naive Bayes, Destek Vektör Makineleri gibi temel geleneksel makine öğrenmesi algoritmaları kullanılarak testler yapılmıştır. Bu testlerin sonucunda, derin öğrenme mimarilerinin geleneksel makine öğrenmesi algoritmalarından daha iyi sonuçlar verdiği gözlemlenmiştir.

Deniz ve diğ. (2023), doktora tezinde, “fiyat performans”, “hızlı kargo”, “iyi paketlenme”, “uygun fiyat” ve “güzel ürün” gibi etiketlerle benzer bir çalışma gerçekleştirmiştir. Veri seti elde edildikten sonra, sırasıyla klasik yöntemler ve derin öğrenme yöntemleri analizlerde kullanılmıştır. Çok etiketli analiz için uygulanan klasik sınıflandırma yöntemlerinden İkili İlgililik (Binary Relevance) ve Biri Geriye Karşı (One vs Rest) algoritmalarının doğru parametre seçimiyle yeterli sonuçlar verdiği, ancak derin öğrenme tekniklerinin gerisinde kaldığı gözlemlenmiştir. Kelime gömme yaklaşımlarında Word2Vec ve GloVe modelleri, TF-IDF yöntemine kıyasla daha düşük performans göstermiştir. Bu durum, kelime gömme modellerinin oluşturulmasında kullanılan veri kaynaklarının yetersizliği ile açıklanmıştır. Aynı şekilde, özellik çıkarımı amacıyla kullanılan BERT modeli de tatmin edici sonuçlar verememiştir. Bunun temel nedeni, BERT modelinin ince ayar (fine-tuning) gerektirmesidir. Derin öğrenme yöntemlerinde ise LSTM ve BiLSTM modelleri ile önemli başarılar elde edilmiştir. Elektronik veri setinde LSTM %85 doğrulama doğruluğuna ulaşırken, BiLSTM modeli %87 doğruluk oranı ile başarılı sonuçlar elde etmiştir. Bu da BiLSTM modelinin doğruluk ve stabilite açısından en iyi sonucu verdiğini göstermektedir. Bu bulgular, derin öğrenme tekniklerinin klasik yöntemlere göre üstünlüğünü ve BERT modelinin etkili kullanılabilmesi için ince ayar gerekliliğini ortaya koymaktadır.

Polat ve diğ. (2021), bir e-ticaret sitesinde farklı kategorilerden ürünlere ait Türkçe yorumları içeren bir veri seti kullanılmıştır. Her bir yorumun duygu durumunun olumlu veya olumsuz olarak belirlenmesi hedeflenmiştir. Bu amaçla, son yıllarda duygu analizinde sıklıkla kullanılan CNN, LSTM ve GRU derin öğrenme yöntemleri ile modeller geliştirilmiştir. Ayrıca, “FastText” tarafından Atlamalı Gramer (Skip-gram) ve Sürekli Kelime Torbası (Continuous Bag of Words) yöntemleri ile oluşturulan önceden eğitilmiş kelime vektörlerinden faydalanılarak bu yöntemlerin modellerin performansına etkisi karşılaştırılmıştır. Yapılan deneysel sonuçlar, Türkçe metinlerin duygu analizi konusunda GRU yöntemi ile oluşturulan modellerin diğer yöntemlere kıyasla daha başarılı sonuçlar verdiğini göstermiştir.

Yönet ve diğ. (2023), çalışmalarında, farklı kategorilerde yazılmış şiirlerin otomatik olarak sınıflandırılması üzerine odaklanmıştır. Çalışmada, metin sınıflandırması için altı farklı makine öğrenmesi algoritmaları (Destek Vektör Makineleri, Çok Katmanlı Algılayıcı, Naive Bayes, K-En Yakın Komşu, Karar Ağacı ve Rastgele Orman) ve doğal dil işleme teknikleri kullanılmıştır. Veri seti, web kazıma yöntemiyle elde edilmiş ve ön işleme aşamasında Zemberek Kütüphanesi'nden faydalanılmıştır. Metinlerin sayısal formata dönüştürülmesi için TF-IDF yöntemi tercih edilmiştir. Eğitim ve test veri setleri %80 eğitim, %20 test olarak ayrılmış ve sınıflandırma işlemi gerçekleştirilmiştir. Denemeler sonucunda en yüksek doğruluk oranı %89 ile Rastgele Orman algoritmasından elde edilmiştir. Ayrıca hiperparametre analizi için “RandomizedSearchCV” ve “GridSearchCV” yöntemleri kullanılarak model performansı optimize edilmiştir. Çalışma, şiirlerin kategorik sınıflandırılmasında makine öğrenimi algoritmalarının etkinliğini ortaya koymaktadır.

Akdeniz ve diğ. (2022), çalışmalarında, Twitter'da uzaktan eğitim konulu Türkçe tweetler üzerine duygu analizi gerçekleştirmiştir. Çalışmada, veri seti zemberek kütüphanesi kullanılarak ön işleme tabi tutulmuş ve metinlerin sayısallaştırılması için BoW, TF-IDF ve Word2Vec yöntemleri uygulanmıştır. Ayrıca N-gram yaklaşımı ile farklı kelime kombinasyonları kullanılarak sınıflandırma performansına etkileri değerlendirilmiştir. Sınıflandırma işlemi için Lojistik Regresyon, Stokastik Gradyan İnişi, Destek Vektör Makineleri, Rastgele Orman ve Naive Bayes algoritmaları karşılaştırılmıştır. En yüksek doğruluk oranı %88 ile Lojistik Regresyon algoritması ile elde edilmiştir. Bunun yanında manuel etiketlenmiş veri seti ile hazır modellerin (TextBlob, Vader, Bert) performansları kıyaslanmış ve manuel etiketli verilerin üstünlüğü gözlemlenmiştir. Analiz sonucunda, uzaktan eğitimle ilgili paylaşımların ağırlıklı olarak olumsuz duygular içerdiği tespit edilmiştir. Bu çalışma, Türkçe duygu analizi ve makine öğrenmesi algoritmalarının sınıflandırma performansındaki etkilerini ortaya koymaktadır.

3. MATERYAL VE YÖNTEM

Gerçekleştirilen tez çalışmasının amacına yönelik olarak, Türkiye'nin önde gelen bir e-ticaret sitesinden elde edilen müşteri yorumlarının toplanmasından işlenmesine, analiz edilmesinden performans değerlendirmesine kadar olan süreç aşağıdaki başlıklar çerçevesinde yürütülmüştür.

3.1. Veri Saklama ve Yönetimi

Çalışmada kullanılmak üzere e-ticaret sitesinden elde edilen veriler gerek etiketleme işleminden önce gerek kullanıcılar tarafından manuel olarak gerçekleştirilen etiketleme işleminden sonra, saklanma ve işleme süreçlerinde Microsoft SQL Server kullanılarak yönetilmiştir. Bu sayede, veri düzenleme ve erişim işlemleri kolaylaşmıştır.

Microsoft SQL Server (MSSQL), Microsoft tarafından geliştirilen ve geniş kullanım alanı sunan bir ilişkisel veritabanı yönetim sistemidir (Relational Database Management System- RDBMS). Büyük veri kümeleri ve karmaşık işlemler için yüksek performans, veri güvenliği ve ölçeklenebilirlik sağlar.

Veritabanı işlemlerinde kullanılan Transact-SQL (T-SQL), standart SQL'e ek işlevler kazandırarak veri yönetimini ve uygulama geliştirme süreçlerini optimize eder. T-SQL'in esnekliği, karmaşık sorgular ve özelleştirilmiş çözümler geliştirilmesine olanak tanır (Yeşilyurt, 2024).

Sonuç olarak, Microsoft SQL Server, güçlü özellikleri ve esnek yapısıyla sadece büyük ölçekli kurumsal uygulamalar için değil, aynı zamanda küçük ve orta ölçekli projeler için de tercih edilen bir veritabanı yönetim sistemidir. T-SQL dilinin sunduğu imkânlar, veritabanı işlemleri ve yönetimini daha verimli ve esnek hale getirmekte, kullanıcıların ihtiyaçlarına göre özelleştirilmiş çözümler geliştirmelerine olanak tanımaktadır.

3.2. Web Kazıma

E-ticaret platformundan müşteri yorumlarının toplanabilmesi için web kazıma işlemi gerçekleştirilmiştir.

Web kazıma, belirli bir web sitesinden veya birden fazla kaynaktan veri çıkarma işlemi olarak tanımlanır ve genellikle yazılım araçlarıyla gerçekleştirilir. Bu yöntem,

çevrimiçi ortamda büyük veri kümelerinden belirli veri setlerini hızla ve verimli bir şekilde toplamak amacıyla kullanılmaktadır.

Kullanım alanı olarak, büyük veri analizi ve veri madenciliği gibi alanlarda önemli bir yöntem olarak öne çıkar. Bot adı verilen otomatik yazılımlar, web sayfalarında gezinerek kullanıcı ihtiyaçlarına yönelik verileri toplar ve analiz için hazır hâle getirir. Web kazıma, doğru bir şekilde uygulandığında, veri analistlerine anlamlı ve güvenilir bilgiler sunarak karar alma süreçlerine katkı sağlar (Divya, Rastogi, Yadav ve Perwej, 2022).

Bu çalışmada, Selenium¹ ve HtmlAgilityPack² kütüphaneleri kullanılmıştır. Selenium, dinamik içeriklere erişim sağlayarak web sayfalarında gezinme, form doldurma ve buton tıklama gibi işlemleri otomatikleştirir. Bu özellikler, özellikle JavaScript ile yüklenen içeriklerin toplanmasını mümkün kılar. Selenium'un sunduğu esneklik, statik ve dinamik içeriklere erişim açısından önemli avantajlar sağlamaktadır. HtmlAgilityPack ise web sayfalarının HTML kaynak kodlarını analiz ederek verilerin ayrıştırılmasını sağlar. Karmaşık HTML yapılarında bile hedeflenen verilere ulaşmayı mümkün kılan bu kütüphane, veri çıkarma işlemlerinde doğruluk ve hız kazandırır. Her iki kütüphane, web kazıma sürecini daha verimli ve güvenilir hâle getirerek, veri analizi ve bilgi işleme süreçlerinde geniş bir kullanım alanı sunmaktadır.

3.3. Web Sitesi Araçları

Web kazıma yöntemiyle elde edilen müşteri yorumları, web ortamında yayınlanarak gerçek kullanıcılar tarafından önceden belirlenmiş kategorilere uygun şekilde etiketlenmiştir. Bu süreçte, web tabanlı araçlar (ASP.NET MVC, Entity Framework ve HTML/CSS/JavaScript) kullanılarak kullanıcıların etiketleme işlemini kolaylaştıracak bir sistem oluşturulmuştur. Böylece, otomatik sınıflandırma işlemlerinde kullanılacak derin öğrenme ve geleneksel makine öğrenmesi algoritmalarının doğru ve güvenilir verilerle eğitilmesi sağlanmıştır.

ASP.NET MVC, Microsoft tarafından modern ve modüler web uygulamaları geliştirmek amacıyla sunulan güçlü bir çerçevedir. Model-Görünüm-Denetleyici (Model-

¹ <https://www.selenium.dev/documentation/>
Selenium, web tarayıcılarının otomasyonu için araçlar ve kütüphaneler sağlayan bir açık kaynaklı projedir. Projenin temelinde, birden fazla tarayıcıda çalışabilen talimatlar yazmayı mümkün kılan WebDriver arayüzü bulunur.

² <https://html-agility-pack.net/documentation>

View-Controller – MVC) tasarım deseni, uygulamayı üç ana katmana ayırarak, modülerlik, bakım kolaylığı ve ölçeklenebilirlik sağlar. Bu yapı, büyük projelerde daha düzenli ve yönetilebilir bir mimari sunar.

- **Veri Katmanı (Model):** Model, uygulamanın veri katmanını temsil eder ve veritabanı işlemleri ile iş mantığını içerir. Veri saklama, işleme ve doğrulama gibi süreçleri yönetir. Entity Framework gibi araçlarla entegre çalışarak veritabanına kolay erişim sağlar ve diğer katmanlardan bağımsız olarak geliştirilebilir.
- **Görünüm Katmanı (View):** View, kullanıcılara görsel arayüz sunar ve HTML, CSS, JavaScript gibi teknolojilerle desteklenir. Model'den gelen verileri görselleştirir ve dinamik içeriklerle kullanıcı dostu bir deneyim sağlar. Bu katman, kullanıcıyla etkileşim için temel bir yapı sunar.
- **Denetleyici Katmanı (Controller):** Controller, Model ve View arasında köprü görevi görür. Kullanıcı isteklerini işler, iş mantığını uygular ve sonuçları View'a aktarır. Dinamik veri akışını yöneterek uygulamanın tüm işleyişini düzenler.

ASP.NET MVC, bu üç bileşeni uyum içinde çalıştırarak sürdürülebilir, test edilebilir ve esnek web uygulamaları geliştirmeye olanak tanır. Bu yapı, geliştiricilere daha verimli ve ölçeklenebilir bir geliştirme süreci sunar (Turan, 2018).

3.4. Doğal Dil İşleme

Doğal dil işleme (NLP) teknikleri sayesinde, toplanan ve kullanıcılar tarafından etiketlenen verilerin ön işleme ve kelime temsili süreçleri titizlikle gerçekleştirilmiştir. Bu adımlar, derin öğrenme ve geleneksel makine öğrenmesi modellerinin daha doğru ve güvenilir bir şekilde eğitilmesi için gerekli olan temiz ve organize bir veri seti oluşturmayı amaçlamaktadır.

Teknolojinin hızlı gelişimi, insan dilinin bilgisayarlara aktarılması ve analiz edilmesi ihtiyacını doğurmuştur. Bu süreç, bilgisayar bilimi ve dil biliminin kesiştiği Doğal Dil İşleme (Natural Language Processing - NLP) alanı ile ele alınmaktadır. NLP, bilgisayarların insan dilini anlayabilmesi, işleyebilmesi ve doğru anlamlar çıkarabilmesi için geliştirilmiş bir tekniktir (Çoban, 2016; Gordin, 2015; Yıldırım, 2018). Temel amacı, dilin yapısını kavrayarak bilgisayarlara anlamlandırma yeteneği kazandırmaktır.

Doğal Dil İşleme, sosyal medya gönderileri, haberler ve e-postalar gibi geniş bir veri yelpazesinde kullanılır (Korkusuz, 2018). Duygu analizi, metin özetleme, otomatik çeviri ve soru-cevaplama gibi uygulamalar, NLP'nin öne çıkan kullanım alanlarıdır

(Yelmen, 2016). Örneğin, duygu analizi bir metnin olumlu ya da olumsuz içeriğini belirlerken, soru-cevaplama, doğru yanıtları otomatik olarak çıkarır.

NLP, büyük metin verilerinin işlenmesinde önemli kolaylıklar sağlar. Soru-cevaplama, bilgi çıkarımı ve duygu analizi gibi yöntemler, sosyal medya analizinden müşteri hizmetlerine kadar geniş bir alanda kullanılır. Bu yetenekleriyle NLP, modern teknolojiye kritik bir rol oynamaktadır (Kızılırmak, 2020).

3.4.1. Veri Ön İşleme

Metin tabanlı veri setlerinin analiz için uygun hale getirilmesi, analizlerin doğruluğu açısından kritik öneme sahiptir. Bu süreçte, yazım hatalarının düzeltilmesi, gereksiz içeriklerin temizlenmesi, kısaltmaların açılması ve kelimelerin anlamını koruyacak şekilde işlenmesi gereklidir. Genellikle yapılandırılmamış formda olan metin verileri, analizden önce yapılandırılmış hale getirilerek daha güvenilir ve tutarlı sonuçlar elde edilmesi sağlanır (Kuzucu, 2015; Çınar, 2020).

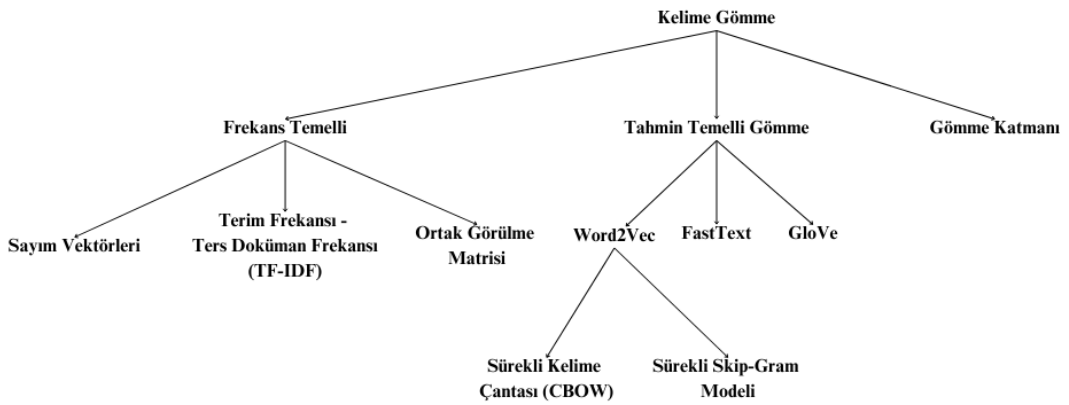
Gürültü Temizleme aşamasında, "gürültü" olarak adlandırılan eksik, hatalı ve tutarsız veriler analiz süreçlerini olumsuz etkiler. Gürültü Temizleme aşaması, bu verilerin temizlenmesi, eksik değerlere uygun tahminler yapılmasını ya da hatalı verilerin düzeltilmesini içerir (Özkan, 2008). Örneğin, eksik veriler yerine sabit bir değer atanabilir veya diğer verilerden tahmin edilerek tamamlanabilir.

Sonuç olarak, veri setinin ön işleme ve gürültü temizleme süreçleri, analizlerin doğruluğunu ve modelin güvenilirliğini artırır. Bu adımlar, büyük veri analizi ve makine öğrenmesi gibi alanlarda model performansını doğrudan etkileyen temel unsurlardır.

3.4.2. Kelime Temsili

NLP çalışmaları, temel olarak metinlerin anlamlandırılması ve işlenmesini hedefler. Ancak bilgisayar sistemleri, doğal dildeki metinleri doğrudan anlayamaz. İnsanların dilsel bağlamı ve anlamı kavrama yeteneği, makineler için mevcut değildir. Bu nedenle, metinlerin işlenebilir hale gelmesi için sayısal temsillere dönüştürülmesi gereklidir. Sayısal temsiller, metin madenciliği, duygu analizi ve anlamsal analiz gibi NLP uygulamalarının gerçekleştirilmesinde vazgeçilmez bir adım olarak işlev görür (Sar, 2021).

Metin işleme ve dil modelleme çalışmalarında, kelimelerin doğru şekilde temsil edilmesi büyük önem taşır. Bilgisayar sistemlerinin dili anlaması, kelimelerin ya da ifadelerin sayısal bir formatta temsil edilmesine dayanır. Bu bağlamda, kelime gömme (word embedding) yöntemleri, metinleri sayısal verilere dönüştürmek için etkili bir yaklaşım sunar. Kelime gömme yöntemleri, kelimeleri yüksek boyutlu vektörler olarak temsil eder ve böylece kelimelerin bağlam içindeki anlamlarını daha doğru bir şekilde modellemeyi mümkün kılar (Aydoğan ve Karcı, 2019). Şekil 3.1. kelime gömme yöntemlerinin şematik gösterimini sunmaktadır.



Şekil 3.1. Kelime Gömme Yöntemleri Şematik Gösterim (Akdeniz, 2022)

Kelime gömme yöntemleri, genellikle iki ana yaklaşıma dayanır: frekans temelli ve tahmin temelli yöntemler. Frekans temelli yöntemler, kelimelerin sıklığını ve kelimeler arasındaki ilişkileri analiz ederek metinlerin yapısal özelliklerini ortaya koyar. Tahmin temelli yöntemler ise kelimenin bağlamını ve semantik ilişkilerini dikkate alarak daha anlamlı vektörler üretir. Örneğin, Word2Vec ve GloVe gibi algoritmalar, kelimelerin çevresindeki diğer kelimelerle olan ilişkisini anlamak için bağlam odaklı bir yaklaşım kullanır.

Kelime gömme yöntemlerinin, metin tabanlı analizlerde derinlemesine ve başarılı sonuçlar elde edilmesine olanak tanıdığı söylenebilir. Bu yöntemler, yalnızca kelimelerin bireysel anlamlarını çözümlemekle kalmaz, aynı zamanda kelimeler arasındaki ilişkileri de modelleyerek dilin daha zengin bir temsiline olanak sağlar. Bu durum, metin madenciliği, duygu analizi ve diğer NLP uygulamalarında daha doğru ve kapsamlı sonuçlar elde edilmesini mümkün kılar (Akdeniz, 2022).

Bu tez çalışmasında yukarıda bahsedilen yöntemlerden TF-IDF kullanılmıştır. Metin madenciliği ve NLP alanlarında, bir terimin bir metindeki önemini belirlemek için kullanılan istatistiksel bir yöntemdir. Bu yöntem, kelimelerin metnin genel anlamına katkısını değerlendirmeyi amaçlar ve Terim Sıklığı (Term Frequency - TF) ile Ters Doküman Sıklığı (Inverse Document Frequency - IDF) olmak üzere iki bileşene dayanır. TF-IDF, bir terimin hem içinde bulunduğu metindeki hem de genel veri kümesindeki önemini ölçerek, metin analizlerinde etkili bir araç sunar.

TF, Denklem 3.1 de görüldüğü üzere bir terimin metindeki ham frekansını baz alır, yani bir terimin bir metinde kaç kez geçtiğini belirler. Bu değer, terimin metin için önemine dair bir ilk gösterge sunar. Ancak yalnızca TF değeri, terimin metne anlamlı bir katkı yapıp yapmadığını tek başına açıklayamaz, çünkü sık kullanılan terimler her zaman anlam açısından önemli olmamaktadır (Taşkiran, 2021).

$$\text{Ham Frekans} = \frac{\text{Terimin belgedeki tekrar sayısı}}{\text{Belgedeki toplam kelime sayısı}} = f_{t,d} \quad (3.1)$$

Bu formülde;

f(t,d) : Belirli bir terim (*t*) bir belge (*d*) içinde kaç kez geçiyorsa, bu sayıyı ifade eder. Yani f(t,d) , terimin o belge içindeki tekrar sayısıdır.

IDF, bir terimin tüm belgelerdeki yaygınlık durumunu ifade eder. IDF, terimin diğer metinlerde ne kadar nadir bulunduğunu değerlendirir. Yüksek bir IDF değeri, terimin yalnızca belirli metinlerde yer aldığını ve bu nedenle ilgili metin için daha anlamlı olabileceğini gösterir. Buna karşılık, düşük bir IDF değeri, terimin tüm belgelerde sıkça kullanıldığını ve metin için özgünlüğünün düşük olduğunu belirtir. Bu nedenle, hem yüksek TF hem de yüksek IDF değerine sahip bir terim, metnin anlamı ve içeriği açısından daha önemli kabul edilir (Denklem 3.2.) (Christian ve diğ., 2016).

$$IDF = \log \left(\frac{\text{Veri setindeki toplam belge sayısı}}{\text{Terimin geçtiği belge sayısı}} \right) \quad (3.2)$$

TF-IDF, bu iki değerın çarpımıyla her terimin önemini daha doğru bir şekilde hesaplar. Yüksek TF-IDF değeri, terimin ilgili metinde sık kullanıldığını ve diğer belgelerde nadiren yer aldığını, dolayısıyla metnin özgünlüğüne önemli katkıda

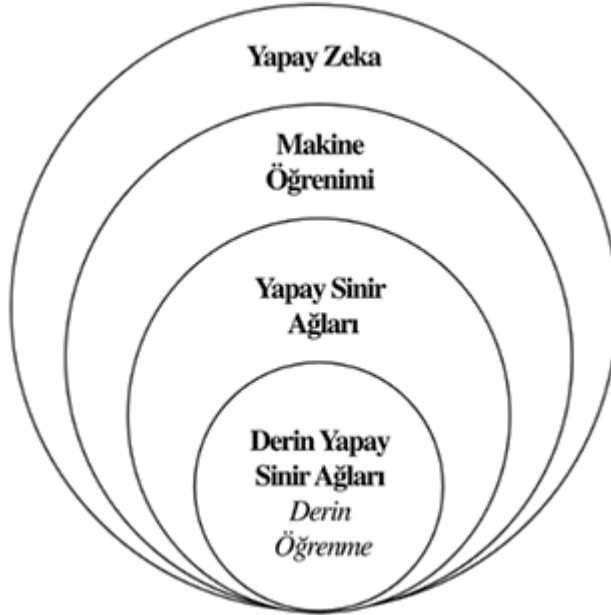
bulunduğunu gösterir. Bu yöntem, özellikle belge sıralama, metin sınıflandırma ve bilgi çıkarımı gibi uygulamalarda etkili bir araçtır.

Bilgi erişimi ve arama motorları için TF-IDF, kullanıcılara en alakalı ve özgün içerikleri sunabilmek adına temel bir ölçü aracı olarak kullanılır. Ayrıca, doğal dil işleme süreçlerinde metinler arası benzerliklerin ölçülmesi, anahtar kelime çıkarımı ve metin madenciliği gibi işlemlerde de sıklıkla tercih edilir. Özetle, metinleri sayısal bir forma dönüştürerek modelin işleyebileceği bir giriş oluşturmak ve bu sayede metindeki her kelimenin önemini vektörel olarak ifade etmektedir.

3.5. Makine Öğrenmesi

1950 yılında Alan Turing, "Makineler düşünebilir mi?" sorusunu sorarak bu konuyu "Bilgisayar Makineleri ve Zekâ" başlıklı makalesinde tartışmıştır. İlk makine öğrenmesi uygulaması, Arthur Samuel'in geliştirdiği ve öğrenme yeteneğiyle dama oynayan bir program olmuştur (Doğançay, 2023). 1959 yılında Samuel, makine öğrenmesini, bilgisayarların doğrudan programlanmadan öğrenme yeteneği sergilediği bir alan olarak tanımlamıştır (Bilekyiğit, 2022).

Yapay zeka hiyerarşisi Şekil 3.2'de verilmiştir;



Şekil 3.2. Yapay Zekâ ve Alt Kolları (Bingöl, 2020)

Makine öğrenmesi, yapay zekânın bir alt dalı olarak, veri setlerinin işlenmesi, dönüştürülmesi ve birleştirilmesi gibi ön işleme adımlarını içerir. Bu süreçten sonra,

çeşitli algoritmalar kullanılarak veri üzerinde matematiksel ve istatistiksel işlemler gerçekleştirilir. Bu işlemler sonucunda tahmin yapma, kümeleme, sınıflandırma ve takviyeli öğrenme gibi farklı görevler yerine getirilir (Erdaş, 2022).

NLP ile hazır hale getirilen müşteri yorumlarının sınıflandırılması amacıyla CNN, RNN, LSTM, BiLSTM, GRU, BERT, Lojistik Regresyon ve Rastgele Orman makine öğrenmesi algoritmaları uygulanmıştır. Bu yöntemler, hızlı ve etkili sonuçlar elde edilmesini sağlamıştır.

Sınıflandırma işlemleri yapay sinir ağı temelli derin öğrenme modelleri ve geleneksel makine öğrenmesi yöntemleriyle gerçekleştirilmiştir. CNN, RNN, LSTM, GRU, BiLSTM ve BERT modelleri kullanılarak müşteri yorumlarının daha derinlemesine analizi yapılmış ve karmaşık veri yapılarında yüksek doğruluk oranlarına ulaşılmıştır. Ek olarak Lojistik Regresyon ve Rastgele Orman geleneksel makine öğrenmesi algoritmaları tercih edilmiştir. Bu yöntemler, özellikle büyük ölçekli veri setlerinde ideal performans sergileyerek, yorumların anlamını daha detaylı şekilde kavrama ve kategorilere ayırma süreçlerine katkı sağlamıştır.

Derin öğrenme, en basit ifadeyle, insan beyninin işlevlerini taklit ederek bilgisayarların öğrenme ve karar verme süreçlerini gerçekleştirebilmesini sağlayan bir yapay zekâ yöntemidir. Makine öğreniminin bir alt dalı olarak kabul edilen derin öğrenme, yapay zekânın yaygınlaşmasıyla birlikte önemli bir çalışma alanı bulmuş ve bilgisayarların, insan beynine benzer şekilde davranabileceğini ortaya koymuştur.

İnsan beyni, öğrenmeyi milyonlarca nöronun etkileşimiyle gerçekleştirir. Derin öğrenme algoritmaları, bu biyolojik süreçten ilham alarak yapay sinir ağlarıyla verileri işler. Yapay sinir ağları, matematiksel işlemler yapan yapay nöronlardan oluşur ve genelde üç ana katmandan meydana gelir:

Girdi Katmanı: Verileri alarak işlemin ilk adımını oluşturur.

Gizli Katman: Verileri işleyip dönüştürür; nöron sayısı arttıkça model daha karmaşık ve güçlü olur.

Çıktı Katmanı: İşlenen veriden sonuç üretir ve sunar.

Derin öğrenme, bu katmanlar aracılığıyla veriyi anlamlı bilgilere dönüştürür ve öğrenme sürecini yönlendirir (Aslanpençesi, 2024).

Derin öğrenme, etiketlenmemiş büyük veri setlerinden sonuçlar elde etmekte önemli avantajlar sunsa da, bu avantajlarla birlikte bazı zorluklar da getirmektedir. Verimli bir derin öğrenme modeli oluşturabilmek için, yüksek kaliteli ve büyük miktarda veri setine ihtiyaç duyulmaktadır. Ayrıca, bu verileri işleyebilecek güçlü bir işlem

altyapısının sağlanması gereklidir. Bu faktörler, derin öğrenmenin hem potansiyelini hem de sınırlarını belirleyen temel unsurlardır.

3.5.1. CNN (Evrişimsel Sinir Ağları - Convolutional Neural Network)

CNN, özellikle görüntü işleme ve doğal dil işleme gibi alanlarda kullanılan, derin öğrenmeye dayalı güçlü bir algoritmadır. CNN'ler, giriş verilerindeki özellikleri çıkarmak ve belirli bir amacı gerçekleştirmek için tasarlanmış katmanlardan oluşur. Giriş verisi genellikle bir görüntü ya da metin şeklinde olur ve bu veri CNN mimarisindeki çeşitli katmanlardan geçirilerek işlenir.

CNN mimarisi genellikle dört temel katmandan oluşur: Konvolüsyonel Katman, Aktivasyon Katmanı, Havuzlama Katmanı ve Tam Bağlantı (Dense) Katmanı. İlk olarak, Konvolüsyonel Katman, giriş verisindeki özelliklerin çıkarılmasını sağlayarak desenler ve yerel ilişkilerin anlaşılmasına olanak tanır. Çekirdekler (kernels) kullanılarak veri üzerinde kaydırma işlemleri yapılır ve bu sayede önemli özellikler çıkarılarak öğrenme sürecine katkı sağlanır. Ardından, Aktivasyon Katmanı, çıkarılan özelliklere doğrusal olmayan dönüşümler uygulayarak modelin daha karmaşık ilişkileri öğrenebilmesini sağlar. Aktivasyon fonksiyonu olarak genellikle ReLU (Rectified Linear Unit) tercih edilir.

Çıkarılan bu özellikler genellikle büyük boyutlu olduğundan, işlem maliyetini azaltmak ve veri boyutunu küçültmek için Havuzlama (Pooling) Katmanı devreye girer. Havuzlama işlemi, maksimum havuzlama veya ortalama havuzlama yöntemleriyle gerçekleştirilir. Maksimum havuzlama, belirli bir bölgede en yüksek değeri seçerken; ortalama havuzlama, belirlenen bölgedeki değerlerin ortalamasını alır. Bu süreç, verinin önemli özelliklerini koruyarak modele daha sade ve etkili bir veri sunar (Aslanpençesi, 2024).

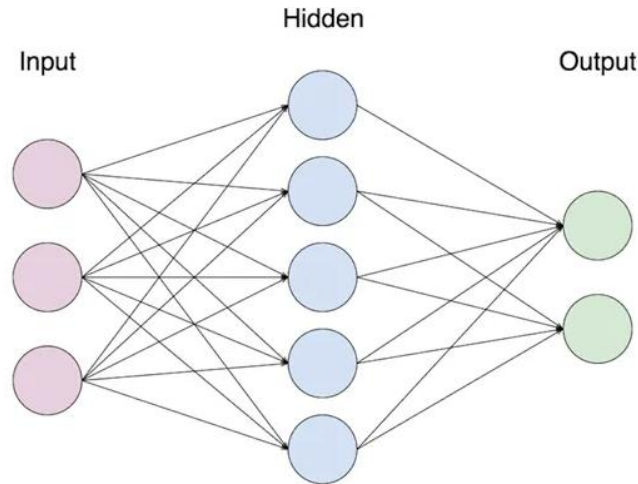
Son aşamada, konvolüsyonel, aktivasyon ve havuzlama katmanlarından geçen veriler, Tam Bağlantı (Dense) Katmanına iletilir. Bu katmanda veriler, matris halinden vektör formuna dönüştürülerek işlenir. Dense katmanında her düğüm, diğer tüm düğümlere bağlıdır ve bu bağlantılar, modelin tahmin ya da sınıflandırma yapmasını sağlayan ağırlıklara sahiptir. Ayrıca, bazı CNN mimarilerinde Normalizasyon Katmanları (örneğin Batch Normalization) bulunarak öğrenme sürecinin hızlanması ve kararlılığın artırılması sağlanır. Böylece model, giriş verisinden elde edilen özellikleri kullanarak istenilen sonuçları üretir.

3.5.2. RNN (Tekrarlayan sinir ağıları - Recurrent Neural Network)

RNN, sıralı verileri veya zaman serilerini işleyebilen bir yapay sinir ağı türüdür. Bu ağlar, özellikle zamanla değişen veri dizilerini analiz etmekte ve işlemekte oldukça etkili olmaktadır. RNN'ler, konuşma tanıma, dil çevirisi, doğal dil işleme, görüntü etiketleme ve video analizi gibi sıralı veya zaman duyarlı problemlere yönelik çözüm sunmaktadır. Bu ağların en yaygın kullanım alanları arasında Siri, sesli arama sistemleri ve Google Çeviri gibi popüler uygulamalar yer almaktadır. RNN'ler, sıralı verilerdeki bağımlılıkları öğrenme ve önceki girdileri dikkate alarak tahminlerde bulunma kabiliyetleriyle öne çıkmaktadır.

Diğer sinir ağı türlerine, örneğin CNN ve ileri yönlü (feed forward) sinir ağlarına benzer şekilde, RNN'ler de eğitim verileri aracılığıyla öğrenme gerçekleştirmektedir. Ancak, RNN'lerin ayırt edici özelliği, sahip oldukları "hafıza" kapasitesidir. Bu hafıza, ağın önceki girişlerden aldığı bilgileri mevcut giriş ve çıkışları etkileyebilmesi için saklamasına olanak tanımaktadır. Geleneksel derin öğrenme modelleri, girdilerin ve çıktının birbirinden bağımsız olduğunu varsaymaktadır. Buna karşın, RNN'lerde her bir çıktının önceden gelen verilerle ilişkisi bulunmakta, bu da sıralı verilerdeki uzun vadeli bağımlılıkları anlamayı mümkün kılmaktadır. Bu özellik, RNN'leri sıralı ve zaman serisi verilerini işlemek için son derece etkili hale getirmektedir (Chentouf, 2024).

Şekil 3.3'te basit bir RNN örneği sunulmakta olup, bu ağın çalışma prensibini daha iyi anlamak için görselleştirilmiştir. Bu şekilde, RNN'lerin zaman adımlarında bilgi akışını nasıl işlediği ve her bir adımda nasıl çıktılar ürettiği daha net bir şekilde gözlemlenebilmektedir.



Şekil 3.3. RNN Mimarisi (Chentouf, 2024)

Sonuç olarak, RNN'ler, sıralı ve zaman serisi verilerindeki bağımlılıkları öğrenebilme yeteneğiyle, çeşitli uygulama alanlarında önemli bir rol oynamaktadır. Tıbbi görüntü analizinden dil işleme ve sesli arama sistemlerine kadar pek çok alanda verimli sonuçlar elde edilebilmektedir. Ancak, RNN'lerin potansiyelinden tam anlamıyla faydalanabilmek için, ileri düzey modeller ve algoritmaların geliştirilmesi gerekmektedir.

3.5.3. LSTM (Uzun Kısa Süreli Bellek - Long Short-Term Memory)

LSTM, geleneksel RNN modelinin uzun vadeli bağımlılıkları daha verimli bir şekilde öğrenebilmesine olanak tanımak amacıyla geliştirilmiş bir yapay sinir ağı mimarisidir. LSTM, özellikle sıralı verilerin işlenmesinde etkin olarak kullanılmakta ve RNN'lerin bellek problemi olan kaybolan gradyan (vanishing gradient) sorununu çözme yeteneği ile öne çıkmaktadır. LSTM'nin temel bileşeni olan hücreler, geçmiş bilgi ve yeni veri arasındaki etkileşimi düzenlemek için üç ana kapıdan yararlanmaktadır: unutma kapısı (forget gate), giriş kapısı (input gate) ve çıkış kapısı (output gate). Bu kapılar, hücrelerin iç durumlarını güncelleyerek, uzun süreli bağımlılıkların korunmasına olanak tanımaktadır (Gers, 2000).

LSTM'nin çalışma prensibi, geçmişteki bilgilerin gerektiği şekilde unutulmasını ve yeni bilgilerin hücreye eklenmesini kontrol eden kapı mekanizmaları üzerine inşa edilmiştir. Unutma kapısı, hücrenin geçmişteki bilgileri ne kadar koruyacağına karar verirken, giriş kapısı yeni bilgilerin hücreye ne kadar dahil edileceğini belirlemektedir. Çıkış kapısı ise hücrenin hangi bilgileri dışa aktaracağına karar vermektedir. Bu süreç, zaman serileri ve sıralı verilerdeki uzun vadeli bağımlılıkların öğrenilmesine olanak sağlamaktadır (Greff, 2017). LSTM'nin en belirgin avantajı, daha karmaşık veri setlerinde, özellikle uzun vadeli ilişkilerde bilgi kaybını önleyerek daha sağlam bir öğrenme süreci sunmasıdır.

LSTM'nin başarısı, büyük ölçüde geri yayılım (backpropagation) algoritmasının etkili bir şekilde uygulanması ve gradyan silinmesi sorunlarının aşılabilmesinden kaynaklanmaktadır. LSTM, hata gradyanlarını zamanın ilerleyen adımlarında geri iletebilme yeteneği sayesinde, uzun vadeli ilişkileri daha iyi öğrenebilmektedir. Bu özellik, modelin geçmiş verilerden gelen bilgileri etkin bir şekilde kullanabilmesini sağlamaktadır. LSTM'nin sıralı verilerle etkileşimde bulunabilmesi, özellikle doğal dil işleme (NLP), zaman serisi tahmini, konuşma tanıma ve metin sınıflandırma gibi alanlarda büyük avantajlar sağlamaktadır. Bu alanlarda, LSTM'nin uzun vadeli

bağımlılıkları yakalayabilme becerisi, modelin yüksek doğrulukla çalışmasına yardımcı olmaktadır (Graves, 2012).

Sonuç olarak, LSTM, sıralı verilerdeki uzun vadeli bağımlılıkları öğrenme yeteneği sayesinde, birçok gerçek dünya uygulamasında başarıyla kullanılmaktadır. LSTM'nin uzun vadeli bilgi akışını etkin bir şekilde yöneten yapısı, doğal dil işleme gibi alanlarda çok değerli bir araç haline gelmesine olanak tanımaktadır. Ancak, eğitim ve hesaplama sürelerinin uzun olabilmesi, modelin verimli bir şekilde uygulanabilmesi için uygun hiperparametre ayarlamaları ve yeterli eğitim verisi gereksinimini doğurmaktadır. Ayrıca, overfitting gibi sorunların önlenmesi için düzenleme yöntemlerinin (regularization) kullanılması önemli görülmektedir.

3.5.4. BiLSTM (Çift Yönlü Uzun Kısa Süreli Bellek - Bidirectional LSTM)

BiLSTM, geleneksel LSTM'nin geliştirilmiş bir versiyonudur ve veriyi hem geçmişten geleceğe hem de gelecekte geçmişe işleyerek bağlamı daha geniş bir perspektifle anlamayı sağlar. BiLSTM, iki bağımsız LSTM birimini birleştirerek ileri ve geri yönde işlem yapar. Bu çift yönlü işlem, modelin geçmişten gelen ve gelecekteki bilgileri aynı anda değerlendirmesine olanak tanır. Böylece, özellikle doğal dil işleme (NLP) gibi sıralı veri uygulamalarında daha derinlemesine bir bağlam anlayışı geliştirilir.

BiLSTM'in temel çalışma prensibi, her bir zaman dilimi için iki ayrı gizli durum oluşturmaktır: biri ileri yönde (geçmişten geleceğe), diğeri ise geri yönde (gelecekte geçmişe) bilgiyi temsil eder. Bu özellik, modelin yalnızca önceki kelimelere değil, aynı zamanda sonraki kelimelere de dayanarak anlamlı sonuçlar üretmesini sağlar. Örneğin, bir cümledeki kelimelerin anlamı, hem öncesindeki hem de sonrasındaki kelimelerle ilişkilidir. BiLSTM bu ilişkileri göz önünde bulundurarak metinlerin anlamını daha doğru bir şekilde analiz eder.

BiLSTM'in uygulama alanları oldukça geniştir. Dil modelleme, metin sınıflandırma, makine çevirisi ve ses tanıma gibi birçok alanda başarıyla kullanılmaktadır. Örneğin, bir cümledeki duygu analizini ele alalım: Geleneksel LSTM yalnızca önceki kelimelere odaklanırken, BiLSTM hem cümlenin başlangıcındaki hem de sonundaki kelimeleri dikkate alarak daha kapsamlı bir analiz yapabilir. Bu özellik, BiLSTM'nin bağlam odaklı uygulamalarda etkili bir çözüm sunmasını sağlar.

BiLSTM'in bir diğer avantajı, kısa ve uzun vadeli bağımlılıkları öğrenme kapasitesine sahip olmasıdır. Metinlerde kelimelerin sırası değişse bile, model bağlamı

doğru şekilde anlayarak uygun sonuçlar üretebilir. Ayrıca, ileri ve geri yönlü bilgilerin birleştirilmesi, modelin anlam çıkarımında doğruluğunu artırır. Örneğin, gelecekteki bir kelimeyi tahmin ederken, BiLSTM hem önceki hem de sonraki kelimeleri dikkate alarak daha isabetli tahminler yapabilir (Emirali, 2022).

Sonuç olarak, BiLSTM, geleneksel LSTM'nin daha kapsamlı ve güçlü bir versiyonudur. Çift yönlü işlem yapabilme özelliği sayesinde, dilin karmaşıklıklarını ve bağlamın zenginliğini daha iyi analiz edebilir. Bu özellikler, BiLSTM'yi metin analizi, duygu analizi, makine çevirisi ve zaman serisi verileriyle ilgili pek çok uygulamada vazgeçilmez bir araç haline getirmiştir.

3.5.5. GRU (Kapı Özyinelemeli Geçitler - Gated Recurrent Unit)

GRU, RNN'nin daha verimli ve basitleştirilmiş bir versiyonudur. Geleneksel RNN'lerdeki kaybolan gradyan problemini aşmak ve uzun vadeli bağımlılıkları öğrenmek için geliştirilmiştir. GRU, daha az parametre ve işlem gerektirerek hızlı ve etkili öğrenme sunar.

GRU'nun temelini oluşturan iki ana mekanizma, Güncelleme Kapısı (Update Gate) ve Sıfırlama Kapısıdır (Reset Gate). Güncelleme kapısı, önceki bilgilerin ne kadarının korunacağını belirlerken, sıfırlama kapısı mevcut durumu hesaplarken geçmiş verilerin ne kadarını dikkate alacağını kontrol eder. Bu yapı, ağırlı bilgi akışını verimli bir şekilde yönetmesine olanak tanır (Dey ve Salem, 2017).

GRU, geleneksel RNN'ler ve LSTM'lere kıyasla daha basit bir yapıya sahiptir, ancak uzun vadeli bağımlılıkları öğrenmede oldukça etkilidir. Bu özellikleri, doğal dil işleme, konuşma tanıma, makine çevirisi ve video analizi gibi alanlarda GRU'yu ideal bir seçenek haline getirir. Daha düşük hesaplama maliyetleri ve daha hızlı öğrenme süreçleriyle, büyük veri projelerinde ve hızlı model eğitimi gerektiren uygulamalarda sıklıkla tercih edilir (Öztürk ve diğ., 2022).

Sonuç olarak, GRU, kapı mekanizmaları sayesinde bilgi akışını etkili bir şekilde yöneterek sıralı veri işleme alanında önemli bir yere sahiptir. Bu basit ama güçlü yapısı, modeli, uzun vadeli bağımlılıkların kritik olduğu uygulamalarda öne çıkarır.

3.5.6. BERT (Dönüştürücülerden Çift Yönlü Kodlayıcı Temsiller - Bidirectional Encoder Representations from Transformers)

BERT, NLP alanında bir dönüm noktası olarak, 2019 yılında Google Research ekibinden Jacob Devlin ve çalışma arkadaşları tarafından geliştirilmiştir. Bu model, metinleri çift yönlü olarak, yani hem soldan sağa hem de sağdan sola okuyarak kelime ilişkilerini derinlemesine analiz etme yeteneğine sahiptir (Devlin ve diğ., 2019). Çift yönlü analiz yeteneği, metin bağlamını daha gerçekçi bir şekilde öğrenmesini sağlar ve bu özellik, BERT'i NLP alanında ön plana çıkaran en önemli faktörlerden biridir.

BERT, özellikle BookCorpus ve Wikipedia gibi geniş ve zengin veri kümesi kaynakları kullanılarak eğitilmiştir. Bu büyük veri setleri, modelin geniş bir dil yelpazesinde kelime ilişkilerini anlamasına olanak tanır. Modelin geliştirilme sürecinde, hem kelime seviyesinde hem de cümle seviyesinde dil ilişkilerinin öğrenilmesine odaklanılmıştır. Bu sayede BERT, dil bağlamını daha iyi kavrayarak karmaşık metin yapılarını anlayabilme yeteneği kazanmıştır.

BERT'in temel yapısı model, kelimelerin çevresindeki bağlamı dikkate alarak onları anlamlandırma üzerine kuruludur. Bu yaklaşım, yalnızca bireysel kelime anlamlarını öğrenmekle kalmayıp, aynı zamanda metnin genel yapısını ve anlamını da modelin kavrayabilmesine olanak tanır.

3.5.6.1. BERT' in Temel Eğitim Yöntemleri

BERT modeli, metin bağlamını anlamada üstün başarı göstermesini iki temel eğitim stratejisine borçludur: Maskeleye Temelli Dil Modellemesi (Masked Language Modeling) ve Sonraki Cümle Tahmini (Next Sentence Prediction).

Maskeleye Temelli Dil Modellemesi: Bu yöntem, metindeki bazı kelimelerin maskelenmesi ve modelin bu kelimeleri bağlama dayalı olarak tahmin etmesi prensibine dayanır. Örneğin, "Ali okula ____" gibi bir ifadede, BERT çevredeki kelimelerden yola çıkarak boşluğu doğru bir şekilde doldurmayı öğrenir. Bu strateji, modelin kelimeler arasındaki bağlam ilişkilerini derinlemesine anlamasını sağlar.

Sonraki Cümle Tahmini: Modelin cümleler arasındaki ilişkileri kavramasını hedefler. İki cümleden birinin diğerinin devamı olup olmadığını tahmin etmeye yönelik bir yöntemdir. Örneğin, "Ali okula gitti. Sınıfta öğretmenini gördü." cümlelerini analiz

ederek, model bu iki ifadenin bağlantılı olup olmadığını belirler. Bu strateji, metin düzeyindeki anlam bütünlüğünü anlamada önemli bir rol oynar.

Bu iki yöntem bir araya gelerek, BERT'in hem kelime hem de cümle seviyesinde bağlam bilgisini etkili bir şekilde öğrenmesini sağlamaktadır (Uçar, 2020).

3.5.6.2. BERT 'in Çift Yönlü Yapısı ve Yenilikleri

BERT, doğal dil işleme alanında bir dönüm noktasıdır ve çift yönlü koldayıcı (encoder) yapısı sayesinde bağlam analizi konusunda önemli bir yenilik sunar. Geleneksel NLP modellerinin yalnızca tek yönlü çalışması, bağlamsal anlamları yakalamada sınırlı kalırken, BERT hem önceki hem de sonraki kelimelerden gelen bağlam bilgisini eşzamanlı olarak değerlendirerek bu sınırlamayı aşmıştır. Bu çift yönlü analiz, kelimeler arasındaki ilişkileri daha derinlemesine anlamayı ve metin anlam bütünlüğünü daha iyi yakalamayı sağlar.

BERT'in temelinde, Dönüştürücü (Transformer) mimarisi bulunur. Paralel hesaplama kapasitesiyle öne çıkan Transformer, büyük veri setleri üzerinde hızlı ve etkili eğitim olanağı sunar. BERT, bu mimarinin encoder kısmını kullanarak dil ilişkilerini hiyerarşik bir yapıda modelleyip karmaşık metin yapılarını öğrenebilir. Bu sayede model, hem işlem hızını artırır hem de bağlamsal zenginliği başarıyla yakalar (Erkan, 2023).

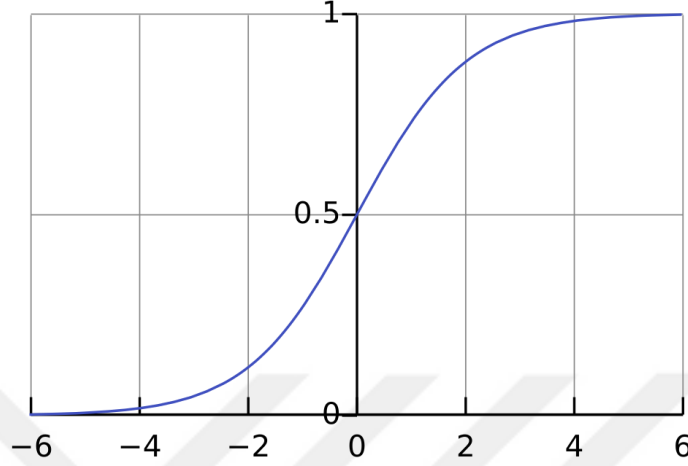
3.5.6.3. BERT' in Uygulama Alanları

BERT, NLP alanında çeşitli uygulamalarda yüksek başarı sağlayan bir modeldir. Metin sınıflandırmada (örneğin, spam tespiti veya duygu analizi), varlık adı tanımada (kişiler, yerler, organizasyonlar), soru yanıtlama modellerinde, metin tamamlama ve özetleme işlemlerinde etkili bir şekilde kullanılır. Ayrıca, soru yanıtlama görevlerinde bağlamdan doğru yanıtlar üretme, metin tamamlama ve özetleme işlemlerinde anlamlı sonuçlar sunar. BERT, çift yönlü bağlam analizi ve Transformer mimarisiyle NLP alanında güçlü ve çok yönlü bir araçtır.

3.5.7. Lojistik Regresyon

Lojistik Regresyon, geleneksel makine öğrenmesi algoritmalarından biri olup, bağımlı ve bağımsız değişkenler arasındaki ilişkiyi modelleyerek olayların olasılığını

tahmin etmekte ve sınıflandırma işleminde kullanılabilir. Bu yöntem, özellikle iki sınıf arasında karar verme işlemlerinde yaygın olarak kullanılmaktadır. Lojistik Regresyonun temelinde, Şekil 3.4'te grafiksel olarak gösterilen Sigmoid fonksiyonu yer almaktadır.



Şekil 3.4. Sigmoid Fonksiyonu (Han ve Moraga, 1995)

Sigmoid fonksiyonu, bir değişkenin değerlerini doğrusal bir şekilde 0 ile 1 arasında sıkıştıran bir matematiksel yapıdır. Bu yapı, Denklem 3.3'te belirtilen formülle ifade edilmiştir. Bu fonksiyon sayesinde, lojistik regresyon çıktıları her zaman 0 ile 1 arasında bir değer alır. Şekil 3.3'de görüleceği üzere, bağımsız değişkeni temsil eden t değeri artı sonsuza yaklaştığında, bağımlı değişkenin çıktısı 1'e yaklaşırken; t değeri eksi sonsuza yaklaştığında ise bağımlı değişken 0'a yaklaşmaktadır.

Model, kullanıcı tarafından belirlenen bir eşik değerine (genellikle 0.5) göre sonuçları sınıflandırır. Eğer hesaplanan lojistik regresyon değeri bu eşik değerinin altında kalırsa, veri 0 olarak değerlendirilir ve ilgili sınıfa ait olmadığı kabul edilir. Aksi durumda, yani değer eşik seviyesini aşarsa, veri mevcut sınıfa ait olarak kabul edilir.

$$\text{sig}(t) = \frac{1}{1+e^{-t}} \quad (3.3)$$

Denklemden kullanılan ifadeler şu şekilde açıklanabilir:

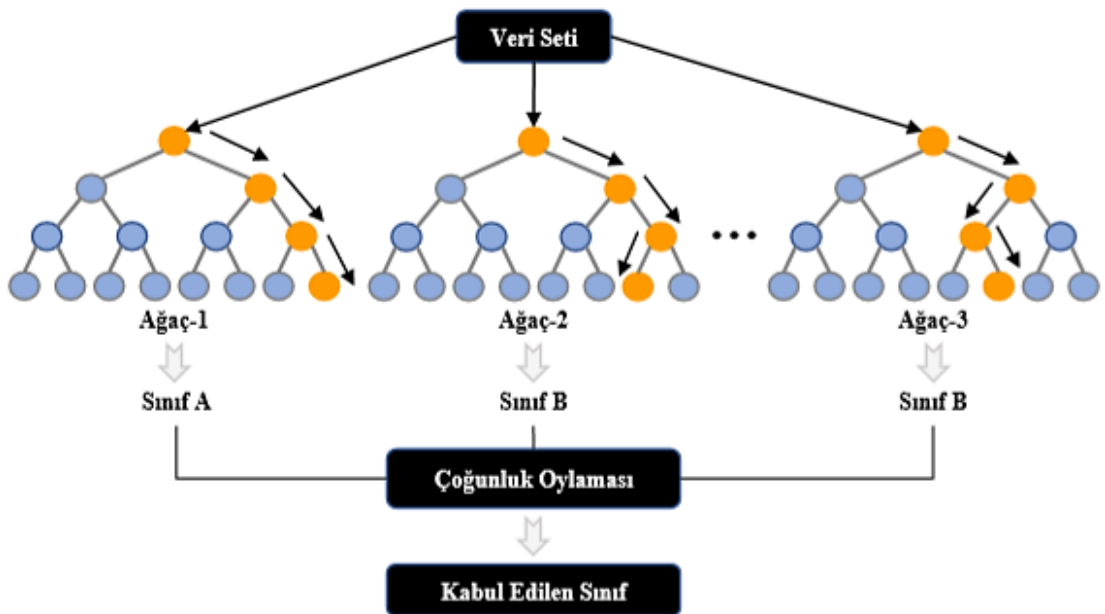
- $\text{sig}(t)$: Lojistik regresyon sonucunda elde edilen olasılık değeri,
- t : Bağımsız değişken için regresyon modelinin tahmin ettiği değer,
- e : Doğal logaritmanın tabanı olan Euler sabiti (≈ 2.71).

Bu yöntem, özellikle doğruluk ve hassasiyetin kritik olduğu sınıflandırma problemlerinde güçlü bir araç olarak kullanılmaktadır (Han ve Moraga, 1995).

3.5.8. Rastgele Orman

Rastgele Orman (Random Forest), sınıflandırma ve regresyon problemlerinin çözümünde kullanılan güçlü ve esnek bir denetimli geleneksel makine öğrenmesi algoritmasıdır. Algoritmanın temel çalışma prensibi, birden fazla Karar Ağacının oluşturulması ve bu ağaçların çıktılarının birleştirilmesine dayanır. Karar Ağaçları, tek bir ağaç üzerinden karar verirken, Rastgele Orman algoritması birden fazla ağacı bağımsız olarak oluşturur.

Rastgele Orman algoritmasının en dikkat çekici özelliklerinden biri, her bir ağacın eğitiminde kullanılan alt veri kümelerinin bagging yöntemiyle ana veri setinden rastgele oluşturulması ve düğüm bölünme süreçlerinde özniteliklerin rastgele seçilmesidir. Bu rastgelelik, modeli aşırı öğrenmeye (overfitting) karşı dirençli hale getirirken, aynı zamanda genelleme yeteneğini artırır. Şekil 3.5'te görüldüğü gibi, algoritma, birbirinden bağımsız çalışan birçok Karar Ağacının çıktısını birleştirerek son tahmin veya sınıflandırmayı gerçekleştirir. Sınıflandırma problemlerinde bu işlem, çoğunluk oyu prensibine göre yapılırken, regresyon problemlerinde ortalama alınarak sonuç belirlenir. Bu yapı, Rastgele Orman algoritmasına hem hızlı çalışma yeteneği hem de yüksek doğruluk oranı kazandırır.



Şekil 3.5. Rastgele Orman Algoritması Çalışma Prensibi (Bonissone ve diğ., 2010)

Algoritmanın etkin bir şekilde çalışabilmesi için kullanıcı tarafından iki ana parametrenin belirlenmesi gerekmektedir:

Karar Ağaçlarının Sayısı: Rastgele Ormanda oluşturulacak ağaç sayısı, modelin genelleme gücünü ve doğruluğunu doğrudan etkiler. Daha fazla ağaç, genellikle daha iyi performans anlamına gelse de, hesaplama maliyetlerini artırabilir.

Her Bir Düğümde Kullanılacak Öznitelik Sayısı: Ağaçların dallanma ve öğrenme süreçlerinde kullanılacak özniteliklerin sayısı, algoritmanın rastgelelik düzeyini belirler. Daha az öznitelik kullanılması, ağaçların daha fazla çeşitlilik göstermesine olanak tanır ve genelleme yeteneğini artırır.

Rastgele Orman, sınıflandırma problemlerinde yüksek doğruluk oranı sunmasının yanı sıra regresyon analizlerinde de güvenilir sonuçlar üretir. Geniş veri setlerinde bile hızlı ve verimli bir şekilde çalışabilir. Aynı zamanda, eksik veri ve öznitelik seçimi gibi problemlere karşı dayanıklı bir yöntem olarak öne çıkar. Bu esneklik ve performans, Rastgele Orman algoritmasını günümüzde makine öğrenmesi projelerinde sıklıkla tercih edilen bir yöntem haline getirmiştir (Akar ve Güngör, 2012).

3.7. Performans Gösterme Metrikleri

Sınıflandırma algoritmalarının performansını değerlendirmek amacıyla doğruluk, duyarlılık, kesinlik ve F1-Skor gibi metrikler sıklıkla kullanılmaktadır. Bu ölçütler, bir modelin tahminlerinin gerçek sonuçlarla ne derece örtüştüğünü analiz etme olanağı sunar. Performans değerlendirmesi sürecinde temel araçlardan biri olan konfüzyon (karmaşıklık matrisi), bir sınıflandırma modelinin doğru ve yanlış tahminlerini pozitif ve negatif sınıflar bazında özetleyen bir yapı olarak kullanılmaktadır (Albayrak, 2018).

Konfüzyon matrisi (Confusion Matrix), sınıflandırma algoritmalarının performansını değerlendirmek için kullanılan temel araçlardan biridir. Bu matris, modelin tahmin ettiği sonuçlar ile gerçek sonuçları bir tablo formatında karşılaştırarak modelin doğruluk oranını ve diğer metrikleri hesaplamaya olanak tanır. Konfüzyon matrisinin temel amacı, bir sınıflandırma algoritmasının performansını görselleştirerek doğru ve yanlış tahminlerin pozitif ve negatif sınıflar özelinde değerlendirilmesini sağlamaktır.

Bu yapı, sınıflandırma performansının hem genel başarı oranı hem de her bir sınıf için detaylı analiz yapılması açısından değerli bir araç olarak öne çıkmaktadır. Şekil 3.6'ta konfüzyon matrisi şeması gösterilmektedir.

		Gerçek Sınıf	
		Pozitif	Negatif
Tahmin Edilen Sınıf	Pozitif	TP	FP
	Negatif	FN	TN

Şekil 3.6. Konfüzyon Matrisi (Yönet, 2023)

Konfüzyon matrisinde 4 temel bileşen vardır.

- Gerçek Pozitif (True Positive - TP): Modelin pozitif bir durumu doğru bir şekilde pozitif olarak tahmin etmesidir.
- Gerçek Negatif (True Negative - TN): Modelin negatif bir durumu doğru bir şekilde negatif olarak tahmin etmesidir.
- Yanlış Pozitif (False Positive - FP): Modelin negatif bir durumu yanlış bir şekilde pozitif olarak tahmin etmesidir.
- Yanlış Negatif (False Negative - FN): Modelin pozitif bir durumu yanlış bir şekilde negatif olarak tahmin etmesidir.

Gerçek Pozitifler (TP + FN) ve Gerçek Negatifler (TN + FP) toplamı, sınıflandırmanın genel doğruluk düzeyini değerlendirmek için kullanılabilir. TP, FP, FN ve TN değerlerinin toplamı tüm gözlemlerin toplamını verir (Küçük, 2023) (Denklem 3.4.).

$$Toplam = TP + FP + FN + TN \quad (3.4)$$

Doğruluk, İngilizce literatürde *Accuracy* olarak adlandırılmakta olup, bir sınıflandırma modelinin genel başarı oranını ifade etmektedir. Bu ölçüt, modelin doğru tahminlerinin (doğru pozitifler [TP] ve doğru negatifler [TN]) toplam tahminlere oranı olarak hesaplanır. Doğruluk metriği, özellikle pozitif ve negatif sınıfların dengeli olduğu veri kümelerinde model performansını değerlendirmek için yaygın bir şekilde kullanılmaktadır. Metrik değeri 1 değerine yakın oldukça başarılı kabul edilir. Doğruluk metriği Denklem 3.5'te görüldüğü gibi hesaplanabilir (Patro ve Patra, 2014).

$$Doğruluk = \frac{TP+TN}{TP+TN+FP+FN} \quad (3.5)$$

Duyarlılık, İngilizce literatürde Recall veya Sensitivity olarak bilinmekte olup, bir sınıflandırma modelinin pozitif sınıfa ait örnekleri doğru bir şekilde tahmin etme oranını ifade etmektedir. Metrik değeri 1 değerine yakın oldukça başarılı kabul edilir. Bu metrik, doğru pozitif tahminlerin (TP), doğru pozitifler ile yanlış negatif tahminlerin (TP + FN) toplamına oranı ile hesaplanır formül gösterimi Denklem 3.6'da gösterilmiştir.

$$Duyarlılık = \frac{TP}{TP+FN} \quad (3.6)$$

Duyarlılık, özellikle pozitif sınıfların eksik tahmin edilmesinin ciddi maliyetlere yol açabileceği durumlarda, model performansının değerlendirilmesinde kritik bir ölçüt olarak öne çıkmaktadır. Düşük duyarlılık, modelin pozitif sınıfa ait örnekleri yeterince iyi tanıyamadığını veya bazı pozitif örnekleri atladığını göstermektedir. Bu durum, modelin eksik tahmin oranını azaltmak için iyileştirilmesi gerektiğine işaret etmektedir. Metrik değeri 1 değerine yakın oldukça başarılı kabul edilir.

Kesinlik, İngilizce literatürde *Precision* olarak adlandırılmakta olup, bir modelin pozitif sınıfa ait olarak yaptığı tahminlerin ne kadarının doğru olduğunu ifade eden bir ölçüttür. Metrik değeri 1 değerine yakın oldukça başarılı kabul edilir. Kesinlik, doğru pozitif tahminlerin (TP), modelin toplam pozitif tahminleri (TP + FP) içindeki oranı olarak hesaplanır Denklem 3.7'deki formül ile ifade edilir.

$$Kesinlik = \frac{TP}{TP+FP} \quad (3.7)$$

Bu metrik, bir sınıflandırma modelinin pozitif sınıfa yönelik seçiciliğini değerlendirmek amacıyla kullanılmaktadır. Kesinlik metriği, özellikle yanlış pozitif tahminlerin maliyetinin yüksek olduğu durumlarda kritik bir öneme sahiptir. Bu bağlamda, yüksek bir kesinlik metriği, modelin yalnızca gerçekten pozitif olan örnekleri tahmin ettiğini ve gereksiz pozitif tahminlerde bulunmadığını gösterir. Ancak, yalnızca kesinlik metriği ile modelin genel performansını değerlendirmek yetersiz olabilir. Bu nedenle, kesinlik genellikle duyarlılık gibi diğer ölçütlerle birlikte ele alınarak daha kapsamlı bir değerlendirme yapılmaktadır.

F1 Skoru, İngilizce literatürde F1-Score olarak ifade edilmekte olup, bir modelin Kesinlik ve Duyarlılık değerlerinin harmonik ortalaması olarak tanımlanmaktadır. Bu metrik, özellikle dengesiz veri kümelerinde model performansını değerlendirmek için yaygın olarak kullanılmaktadır. Metrik değeri 1 değerine yakın oldukça başarılı kabul edilir. F1 Skoru, hem yanlış pozitif (FP) hem de yanlış negatif (FN) tahminleri dikkate alarak dengeli bir performans ölçütü sunar ve denklem 3.8'deki formül ile ifade edilir.

$$F1\ Skoru = 2 \times \frac{Kesinlik \times Duyarlılık}{Kesinlik + Duyarlılık} \quad (3.8)$$

Bu metrik, modelin genel etkinliğini yansıtarak, kesinlik ile duyarlılık arasında bir dengeyi ifade etmektedir.

F1 Skoru, modelin hem güçlü hem de zayıf yönlerini analiz etmede etkili bir araçtır. Özellikle sınıflar arasında dengesizliklerin olduğu veri setlerinde, yalnızca kesinlik veya duyarlılık metriği üzerinden yapılan değerlendirmelerin eksikliklerini tamamlar. Bu nedenle, F1 Skoru, model performansını daha adil bir şekilde yansıtmak için sıklıkla tercih edilmektedir.

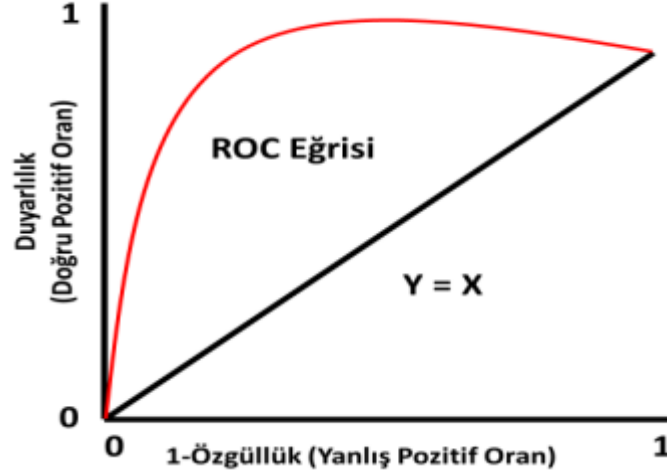
Yukarıda açıklanan metrikler, iki sınıflı sınıflandırmalar için tasarlanmıştır. Ancak, çoklu sınıflı problemler için bu metrikler sınıf bazında hesaplanarak Makro Ortalama ve Ağırlıklı Ortalama yöntemleriyle geliştirilebilir:

- **Sınıf Bazında Hesaplama:** Her sınıf ayrı ayrı pozitif sınıf olarak ele alınır ve TP, FP, FN, TN değerleri bu bağlamda hesaplanır.
- **Makro Ortalama:** Tüm sınıfların metriği eşit ağırlıkla ortalaması alınır.
- **Ağırlıklı Ortalama:** Sınıfların örnek büyüklüklerine göre ağırlıklandırılarak genel bir metrik değeri hesaplanır.

Bu yöntemlerle, çok sınıflı problemler için modelin genel performansı daha adil ve detaylı bir şekilde değerlendirilir.

ROC eğrisi (Receiver Operating Characteristic Curve), bir sınıflandırma modelinin ayırt etme gücünü ve performansını değerlendirmek için kullanılan önemli bir yöntemdir. Özellikle metin sınıflandırma ve etiketleme gibi NLP uygulamalarında, modelin karar sınırlarını belirlemek ve doğruluğunu izlemek için geniş bir kullanım alanına sahip bir tekniktir. ROC eğrisinde, y ekseninde doğru pozitif oranı (True Positive Rate - TPR, yani duyarlılık) yer almakta, x ekseninde ise yanlış pozitif oranı (False Positive Rate - FPR, ya da 1-özgüllük) gösterilmektedir. Bu ilişki, modelin başarısını

yansıtarak, güçlü bir modelin ROC eğrisinin sol üst köşeye daha yakın bir konumda olmasını sağlamaktadır. Bu da modelin yüksek duyarlılık ve özgüllük sağladığını gösterir (Tomak ve Yüksel, 2009; Çelik, 2011) (Şekil 3.7).



Şekil 3.7. Roc Eğrisi Formül Yapısı (Bozdoğan, 2024)

ROC eğrisi, modelin farklı eşik değerlerinde duyarlılık ve yanlış pozitiflik oranları arasındaki ilişkiyi görselleştirir. Bu bağlamda, ROC eğrisinin başlıca kullanım alanları aşağıdaki şekilde sıralanabilir:

- Modelin ayırt etme gücünün değerlendirilmesi,
- Farklı modellerin etkinliğinin karşılaştırılması,
- Uygun tahmin sınırlarının belirlenmesi,
- NLP uygulamalarında etiketleme doğruluğunun izlenmesi,
- Farklı model parametrelerinin performans karşılaştırmalarının yapılması.

ROC eğrisinin altında kalan alan (AUC), modelin doğruluğunu ölçen önemli bir metriktir. AUC değeri, doğru pozitif sonuçların yanlış pozitif sonuçlara oranını temsil eder. Modelin genel başarı düzeyini özetleyen önemli bir metrik olarak öne çıkmaktadır. AUC'nin değeri 0,50 ile 1,00 arasında değişir; değerin yüksek olması, modelin daha güçlü bir ayırım yapabildiğini gösterir (Faraggi ve Reiser, 2002).

Log Loss, sınıflandırma modellerinin probabilistik tahminlerini değerlendirmek için kullanılan bir hata metriğidir. Bu metrik, tahmin edilen olasılıkların gerçek sınıflarla ne kadar uyumlu olduğunu ölçer. Metrik değeri 0 değerine yakın oldukça başarılı kabul edilir. Log Loss değeri ne kadar düşükse, modelin tahmin doğruluğu o kadar yüksek olmaktadır. Özellikle dengesiz veri kümelerinde ve çok sınıflı problemler üzerinde etkili

bir performans değerlendirme aracı olarak kullanılmaktadır. Log Loss, denklem 3.9'da görülen formülle hesaplanmaktadır.

$$LogLoss = - \frac{1}{N} \sum_{i=1}^N [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)] \quad (3.9)$$

Denklemdaki;

N , Toplam veri sayısını;

y_i , Gerçek sınıf etiketlerini;

$\hat{y}_i$, Modelin i -inci örnek için tahmin ettiği olasılık

$\log$: Doğal logaritma işlemidir.

Log Loss, yanlış tahminlerdeki yüksek maliyeti cezalandırdığı için, modelin sınıflar arası tahmin performansını daha hassas bir şekilde ölçer.

3.8. Çalışma Ortamı ve Paketler

Son yıllarda, makine öğrenimi ve veri bilimi alanındaki gelişmeler, bu alanlarda kullanılan programlama dillerine olan ilgiyi artırmıştır. Günümüzde özellikle R, Python ve Java gibi genel amaçlı programlama dilleri, veri analitiği ve makine öğrenimi projelerinde sıklıkla tercih edilmektedir. Bu çalışmada ise dinamik yapısı, hızlı performansı, yerleşik kütüphaneleri ve gelişmiş veri madenciliği kabiliyetleri nedeniyle programlama dili olarak Python seçilmiştir. Python, ilk olarak 1991 yılında Guido van Rossum tarafından geliştirilmiş ve genel amaçlı bir programlama dili olarak tanıtılmıştır. Kullanıcı dostu yapısı, öğrenme kolaylığı ve geniş kütüphane desteği ile Python, günümüzde dünya genelinde geniş bir kullanıcı kitlesine hitap eden popüler bir dil haline gelmiştir (Malkoç, 2012).

Python kodlaması, hem yerel geliştirme ortamlarında (örneğin, PyCharm, Spyder, PyDev gibi IDE'ler) hem de bulut tabanlı platformlarda (örneğin, Jupyter Notebook, Google Colab gibi) gerçekleştirilebilmektedir. Bu çalışmada, geliştirme ortamı olarak Google Colaboratory (Colab) tercih edilmiştir. Colab, önceden yapılandırılmış ve Grafik İşlem Birimi (Graphic Processing Unit - GPU) ile Tensör İşlem Birimi (Tensor Processing Unit - TPU) kullanımını destekleyen bir bulut hizmetidir. Kullanıcılara ortam kurulumu gerektirmeden hızlı bir başlangıç yapma imkânı sunan Colab, özellikle makine

öğrenimi ve derin öğrenme projelerinde yaygın olarak kullanılmaktadır. Ayrıca, ücretsiz bir hizmet olması, projelerin Google Drive veya GitHub gibi platformlarla kolayca paylaşılabildiğini sağlaması, Colab'ı cazip bir seçenek haline getirmektedir.

Bunun yanı sıra, çalışmada çeşitli yazılım geliştirme ve otomasyon teknolojileri de kullanılmıştır. Windows Forms uygulamaları, masaüstü yazılımlarının geliştirilmesinde yaygın olarak tercih edilmiştir. Bu platformda, kullanıcı arayüzleri oluşturulmuş ve veri etkileşimi sağlanmıştır. ASP.NET MVC ile ise web tabanlı uygulamalar geliştirilmiş, kullanıcıların etkileşimli bir deneyim yaşaması sağlanmıştır. Web kazıma işlemleri için HTMLAgilityPack ve Selenium gibi araçlar kullanılmıştır. HTMLAgilityPack, HTML dokümanlarını analiz etmek ve verileri çıkarmak için kullanılırken, Selenium, web sayfalarında etkileşimli testler yapmak ve veri toplamak için tercih edilmiştir.

3.8.1. Python Kütüphaneleri ve Kullanımları

Python, esnek yapısı ve açık kaynaklı doğası ile birçok kütüphane sunarak geliştiricilere geniş bir çözüm yelpazesi sağlamaktadır. Veri madenciliği, makine öğrenimi ve metin analizi gibi disiplinlerde sıklıkla kullanılan bazı Python kütüphaneleri aşağıda detaylandırılmaktadır:

NLTK (Doğal Dil Araç Takımı - Natural Language Toolkit):

NLTK, NLP ve metin madenciliği alanında en çok tercih edilen Python kütüphanelerinden biridir. Bu kütüphane, kök bulma, kelime ayrıştırma, durak kelimelerin (stop words) çıkarılması gibi temel işlemleri kolaylıkla gerçekleştirmek için zengin bir paket içeriği sunmaktadır. Ayrıca, duygu analizi, sınıflandırma, otomatik özetleme ve anlamsal çıkarım gibi daha karmaşık doğal dil işleme görevlerini desteklemektedir (Mehta, 2021).

NumPy: NumPy³, bilimsel hesaplama için temel bir Python kütüphanesidir. Matris işlemleri, çok boyutlu dizi nesneleri üzerinde sıralama, matematiksel ve mantıksal işlemler gibi işlevleri hızlı ve verimli bir şekilde gerçekleştirme imkânı sağlamaktadır.

³ <https://numpy.org/doc/stable/>

NumPy, dizi tabanlı matematiksel işlemlerin performansını artırırken, büyük veri setleriyle çalışmayı da kolaylaştırmaktadır.

Pandas: Pandas⁴, veri analizi ve yönetimi için güçlü bir araç sunmaktadır. Etiketlenmiş veya ilişkisel verilerle çalışmaya olanak tanıyan bu kütüphane, veri çerçevelerini (dataframe) oluşturma, dönüştürme, eksik verileri işleme gibi görevleri hızlı bir şekilde gerçekleştirebilir. Pandas, veri temizleme ve ön işleme süreçlerinde sıklıkla kullanılan bir araçtır.

Matplotlib ve Seaborn: Matplotlib⁵, veri görselleştirme için temel bir Python kütüphanesidir. İki boyutlu grafik ve diyagramlar oluşturmak için sıklıkla kullanılan bu kütüphane, nesne yönelimli API desteği sayesinde uygulamalara kolayca entegre edilebilmektedir (Custer, 2020). Seaborn⁶ ise Matplotlib üzerine inşa edilmiş bir kütüphane olup, istatistiksel verilerin görselleştirilmesi için daha geniş bir araç seti sunmaktadır. Seaborn, veri dağılımlarını özetleyerek ve analiz ederek makine öğrenimi projelerinde anlamlı görselleştirmeler oluşturmayı kolaylaştırmaktadır.

SciKit-Learn: SciKit-Learn⁷, makine öğrenimi algoritmalarını uygulamak için yaygın olarak kullanılan bir Python kütüphanesidir. Doğrusal regresyon, karar ağaçları ve rastgele ormanlar gibi popüler algoritmaları destekleyen bu kütüphane, veri işleme, çapraz doğrulama, öznitelik mühendisliği ve model değerlendirme gibi süreçlerde etkili bir çözüm sunmaktadır (Yüceoğlu, 2017).

⁴ <https://pandas.pydata.org/pandas-docs/stable/>

⁵ <https://matplotlib.org/stable/users/index.html>

⁶ <https://seaborn.pydata.org/tutorial.html>

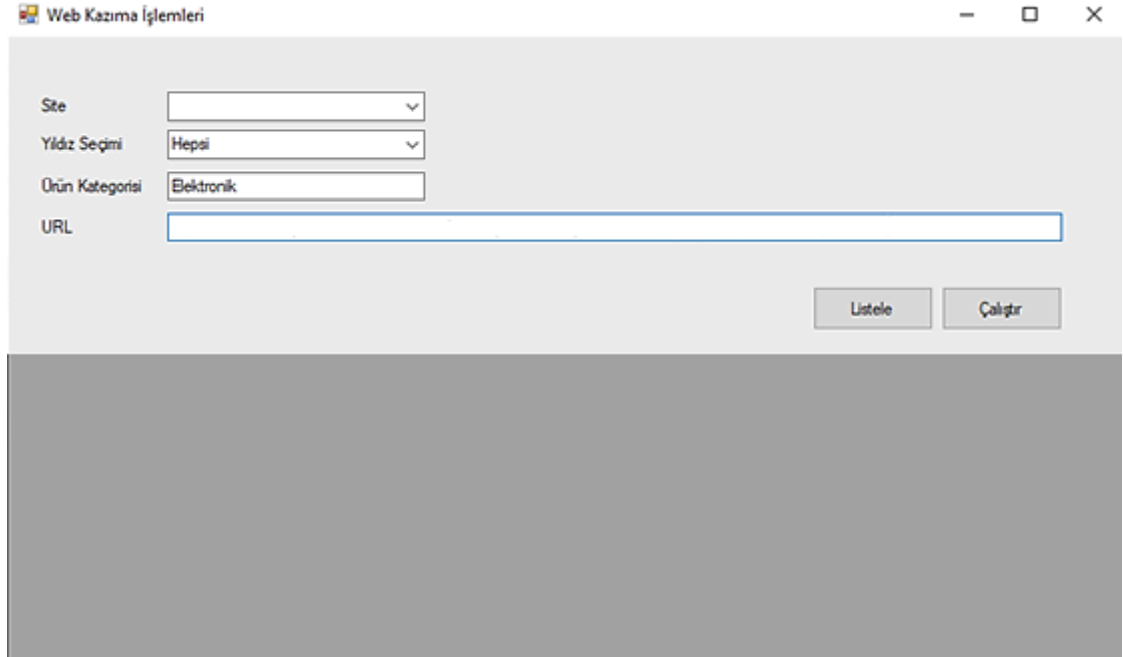
⁷ https://scikit-learn.org/stable/user_guide.html

4. ARAŞTIRMA SONUÇLARI VE TARTIŞMA

4.1. Web Kazıma Projesi, Verilerin Toplanması ve Saklanması

Gerçekleştirilen tez çalışması kapsamında, e-ticaret sitesinden veri kazıma işlemlerini gerçekleştirmek için Şekil 4.1’de görüldüğü gibi bir Windows Forms tabanlı uygulama geliştirilmiştir. Bu uygulama, Giyim, Elektronik, Ayakkabı ve Ev Eşyası kategorilerinde kullanıcı yorumlarını kazıyarak, elde edilen verileri bir veritabanına kaydetmiştir. Toplamda 120.000 veri kazınmış, ürünlerin marka, fiyat, yıldız değerlendirmesi ve yorum içerikleri gibi özellikleri işlenmiştir.

Uygulamanın geliştirilmesinde C# programlama dili ile birlikte Windows Forms Framework kullanılmıştır.



Şekil 4.1. Veri Kazıma Uygulaması için Windows Forms uygulaması

Bu aşamada iki temel modül geliştirilmiştir:

i. *Web Kazıma Modülü*

E-ticaret web sitesinden ürün bilgileri ve kullanıcı yorumlarını toplamak amacıyla, dinamik içeriklerin işlenmesine olanak tanıyan bir web kazıma modülü geliştirilmiştir. Bu süreçte HtmlAgilityPack ve Selenium WebDriver kütüphaneleri bir arada kullanılmıştır. Selenium, dinamik olarak yüklenen JavaScript içeriklerine erişim

sağlarken, HtmlAgilityPack bu içerikleri analiz ederek yapılandırılmış verilere dönüştürmüştür.

Web kazıma süreci, öncelikle ürünlerin yıldız, kategori ve URL bilgilerine göre hedeflenmesiyle başlamıştır. Selenium, sayfaların tam olarak render edilmesini sağlayarak dinamik HTML içeriklerini almış, bu içerikler daha sonra HtmlAgilityPack kullanılarak ayrıştırılmıştır. İlgili veriler (örneğin ürün adları, fiyatlar, yıldız değerlendirmeleri ve ürünlerin detay sayfalarına geçiş yapılarak her bir ürünün yorum içerikleri) XPath ifadeleri ve CSS Selector yöntemleriyle tespit edilerek çekilmiştir.

Kazınan veriler, Product_ID gibi birincil anahtarlar aracılığıyla veritabanında ilişkilendirilmiş ve organize edilmiştir. Bu yöntem, her bir ürünün doğru yorumlarla eşleştirilmesini sağlamıştır. Dinamik içerikleri analiz etmek ve yapılandırılmış verilere dönüştürmek için bu iki kütüphanenin birlikte kullanımı, süreçte verimlilik ve doğruluk açısından önemli bir avantaj sunmuştur.

ii. Veritabanı Entegrasyonu Modülü

Elde edilen veriler, SQL Server üzerinde oluşturulan ilişkisel tablolara kaydedilmiştir. Ürünlere ait temel bilgiler, ürün adı, marka, fiyat, yıldız değerlendirmesi ve ürün görseli gibi alanları içeren [TblProducts] tablosunda saklanmıştır (Şekil 4.2). Bu Tablo yapısı, ürünlerin benzersiz bir şekilde tanımlanmasına ve çeşitli meta verilerin kaydedilmesine olanak tanımaktadır. Ürünlere ait kullanıcı yorumları ise kullanıcı adı, yorum tarihi, içerik ve satıcı bilgisi gibi ayrıntıları barındıran [TblComments] tablosunda organize edilmiştir (Şekil 4.3). [TblTickets] tablosu oluşturularak, kullanıcıların yorumları etiketlemesi için belirlenen kategorik veya tematik bilgiler kaydedilmiştir (Şekil 4.4). Ayrıca, hangi e-ticaret web sitesinden kazıma işleminin gerçekleştirildiğine dair bilgiler, [TblSite] tablosunda tutulmuştur (Şekil 4.5). Bu yapılar, verilerin sistematik bir şekilde saklanmasını ve sorgulanmasını sağlamaktadır. Bunlara ek olarak, [TblUsers] tablosu, sistemin kullanıcılarını tanımlamak ve kimlik bilgilerini güvenli bir şekilde saklamak amacıyla tasarlanmıştır (Şekil 4.6). Bu tablo, sistemin kullanıcı yönetimi için temel bir yapı sağlar ve kullanıcıların yorum etiketleme süreçlerindeki hareketlerini takip etmeye olanak tanır. Son olarak da oluşturulan [TblUserComments] tablosu (Şekil 4.7) ile de kullanıcıların etiketleme sürecindeki işlemleri kayıt edilmiştir. Şekil 4.8'de Veritabanı ER Diagram verilmiştir.

[illegible]

	Ticket_ID	Ticket_Name	Site_ID
1	1	Fiyat Performans	1
2	2	İade Edilen	1
3	3	Paketleme İyi,Hızlı Teslimat	1
4	4	Kalitesiz,Kusurlu, Kötü Paketleme	1
5	5	Tavsiye Edilen, Kaliteli Ürünler	1
6	6	Diğer	1

	Site_ID	Site_Name
1	1	Trendyol

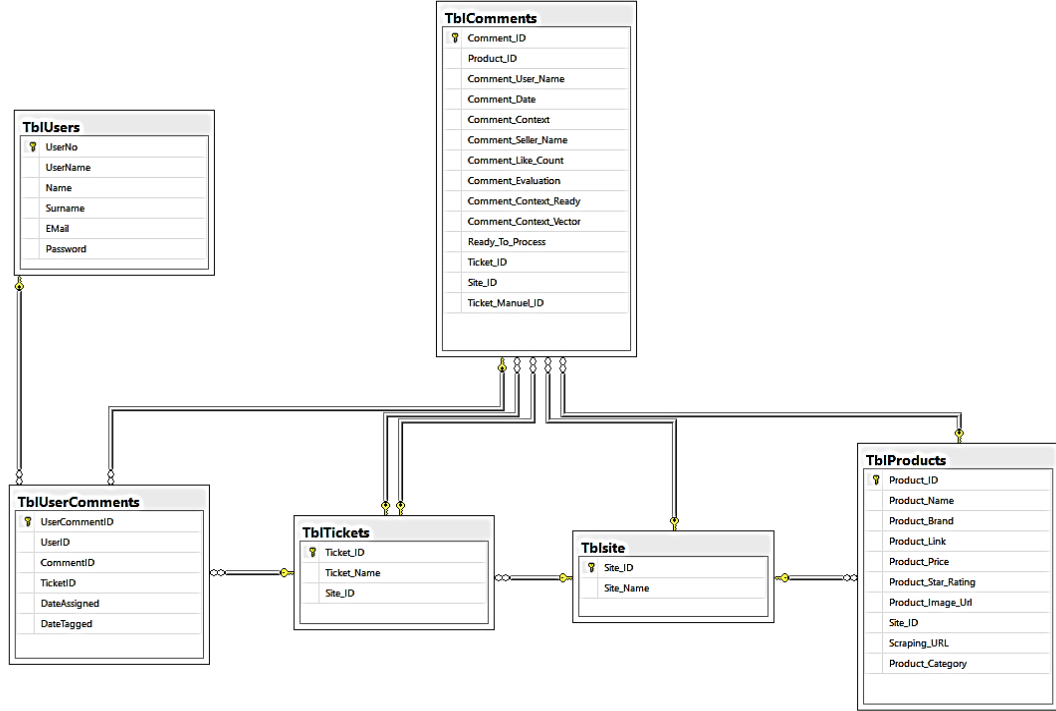
Şekil 4.5. Tblsite tablo yapısı

	UserNo	UserName	Name	Surname	Email	Password
1	1	t****	t****	t****	t***@test.com	*****
2	2	a****	a****	a****	a****@test.com	*****
3	3	H*****	H****	Ç*****	h*****@gmail.com	*****
4	4	f****	F*****	S****	f*****@ktun.edu.tr	*****
5	5	a****	A****	S****	a*****@ktun.edu.tr	*****
6	6	i****	i****	u*	i*****@gmail.com	*****
7	9	d****	d****	d****	d*****@gmail.com	*****
8	11	A****	A*****	C*****	F*****@ktun.edu.tr	*****
9	13	S****	S*****	Ç****	f*****@ktun.edu.tr	*****
10	14	E*****	E***	C****	f*****@ktun.edu.tr	*****
11	15	d*****	D****	U*****	u*****@gmail.com	*****
12	16	e*****	e***	a****	f*****@ktun.edu.tr	*****
13	17	T****	H****	A*****	h*****@gmail.com	*****
14	18	A*****	A****	B*****	f*****@ktun.edu.tr	*****
15	19	B****	B****	K***	f*****@ktun.edu.tr	*****
16	20	G****	G*****	K*****	f*****@ktun.edu.tr	*****
17	21	A*****	A*****	D****	f*****@ktun.edu.tr	*****
18	22	E****	M*****	K**	f*****@ktun.edu.tr	*****
19	23	A*****	A***	N****	f*****@ktun.edu.tr	*****
20	24	A****	A*****	A**	f*****@ktun.edu.tr	*****
21	25	e*****	E****	G*****	f*****@ktun.edu.tr	*****
22	26	Ö****	Ö*****	U*****	o*****@gmail.com	*****

Şekil 4.6. TblUsers tablo yapısı

	UserCommentID	UserID	CommentID	TicketID	DateAssigned
1	1	3	21527	1	2024-10-02 15:52:08.227
2	2	3	32732	4	2024-10-02 15:52:08.257
3	3	3	37948	3	2024-10-02 15:52:08.257
4	4	3	45574	1	2024-10-02 15:52:08.257
5	5	3	63721	4	2024-10-02 15:52:08.257
6	6	3	67038	1	2024-10-02 15:52:08.257
7	7	3	68517	1	2024-10-02 15:52:08.257
8	8	3	70363	2	2024-10-02 15:52:08.257
9	9	3	76388	3	2024-10-02 15:52:08.257
10	10	3	78482	3	2024-10-02 15:52:08.257
11	11	2	15730	4	2024-10-02 16:00:54.123
12	12	3	3633	2	2024-10-02 16:03:49.627
13	13	3	16608	2	2024-10-02 16:03:49.627
14	14	3	18960	2	2024-10-02 16:03:49.627
15	15	3	23664	2	2024-10-02 16:03:49.627
16	16	3	29050	2	2024-10-02 16:03:49.627
17	17	3	31148	4	2024-10-02 16:03:49.627
18	18	3	40775	2	2024-10-02 16:03:49.627
19	19	3	49402	4	2024-10-02 16:03:49.627
20	20	3	67078	4	2024-10-02 16:03:49.627
21	21	1	3099	4	2024-10-02 16:15:41.953
22	22	1	28991	5	2024-10-02 16:18:27.427

Şekil 4.7. TblUserComments tablo yapısı



Şekil 4.8. Veritabanı ER Diagram

4.2. Verilerin Etiketlenmesi ve Veri Seti Oluşturulması

Bu tez çalışmasında, yorum etiketleme işlemlerinin kullanıcılar tarafından kolaylıkla gerçekleştirilebilmesi amacıyla ASP.NET MVC mimarisi kullanılarak bir web tabanlı sistem geliştirilmiştir. Model-View-Controller prensibine dayalı bu yapı, katmanlı mimari sayesinde hem sürdürülebilir bir kod yapısı hem de esnek bir geliştirme ortamı sunmuştur. Kullanıcı arayüzünün daha estetik ve interaktif hale getirilmesi için Bootstrap ve JQuery kullanılmış, bu sayede yorum listeleme ve etiketleme işlemlerinde dinamik içeriklerin hızlı bir şekilde yüklenmesi sağlanmıştır.

Geliştirilen bu sistem, kullanıcıların sisteme üye olarak giriş yapmalarını ve belirlenmiş kategorilere uygun etiketleri seçerek yorumları etiketlemelerini sağlamaktadır. Yorumların doğru ve güvenilir bir şekilde etiketlenmesi için sistemde Entity Framework kullanılmış, bu Nesne-İlişkisel Eşleme (Object Relational Mapping - ORM) aracı sayesinde SQL sorgularına ihtiyaç duyulmadan veritabanı işlemleri gerçekleştirilmiştir. Veritabanı çözümü olarak SQL Server tercih edilmiş; büyük veri işleme kapasitesi ve güvenliği sayesinde sistemin arka planındaki veri yönetimi etkin bir şekilde sağlanmıştır.

Kullanıcılar, sisteme giriş yaptıktan sonra rastgele atanmış yorumlara erişim sağlayarak bu yorumlara ilgili etiketleri eklemiştir. Bu süreç, yorumların güvenliğini ve özgünlüğünü sağlamak adına dikkatle tasarlanmış olup, kullanıcılar yalnızca kendi katkıda bulundukları yorumları görebilmektedir. Böylece, veri bütünlüğü korunmuş ve sistem, kullanıcı dostu bir arayüzle desteklenerek işlemlerin hızlı ve etkili bir şekilde gerçekleştirilmesine olanak tanımıştır.

Proje, www.yorumetiketle.com.tr adresinde yayınlanarak kullanıma sunulmuştur. Sisteme toplam 70 kullanıcı kayıt olmuş ve bu kullanıcılar tarafından önceden belirlenmiş kategorilere uyan anahtar kelimeleri içeren 47 bin yorum etiketlenmiştir. Yorumların her biri, en az bir kez etiketlendikten sonra aynı yorumun farklı kullanıcılar tarafından birden fazla kez etiketlenmesi sağlanmış ve toplamda 65 bin etiketleme işlemi gerçekleştirilmiştir. Bu kapsamda, sistem hem veri çeşitliliği hem de etiketleme doğruluğu açısından önemli bir başarı elde etmiştir.

Projenin veri tabanı altyapısı, çalışmanın önceki bölümlerinde (bkz. Bölüm 4.1 - Veritabanı Entegrasyonu Modülü) açıklanan web kazıma sürecinde elde edilen yorumların saklandığı yapı ile paralellik göstermektedir. Bu yapı, hem kullanıcıların hem de yorumların düzenli bir şekilde yönetilmesine olanak sağlamıştır.

Geliştirilen Web Tabanlı Sistemin detayları aşağıdaki alt başlıklarda verilmiştir.

4.2.1. Kayıt Olma Süreci

Çalışma kapsamında, kullanıcıların sisteme en kısa ve hızlı şekilde üye olarak giriş yapmalarını sağlamak temel hedeflerden biri olmuştur. Kayıt olma süreci, kullanıcı deneyimini öncelikli tutarak basit ve erişilebilir bir yapıda tasarlanmıştır. Bu süreçte, kullanıcılar yalnızca temel bilgilerle (örneğin, e-posta adresi, şifre ve kullanıcı adı) sisteme kayıt olabilmektedir. Böylece, karmaşık işlemlerden kaçınılarak kullanıcıların motivasyonu artırılmaktadır. Kayıt sırasında girilen bilgilerin güvenliği ve gizliliği, şifreleme algoritmaları kullanılarak sağlanmaktadır. Şekil 4.9 kayıt ekranını göstermektedir.

Şekil 4.9. Kayıt olma ekranı

4.2.2. Giriş Ekranı

Web sitesinde, kullanıcıların giriş yapmadan herhangi bir bilgiye erişmesi sistematik olarak engellenmiştir. Bu tasarım, hem güvenlik hem de sistem işleyişi açısından büyük önem taşımaktadır. Kullanıcıların sisteme giriş yapmaları, yapılan tüm etiketleme işlemlerinin izlenebilirliğini sağlamaktadır. Her bir etiketleme işlemi, hangi kullanıcı tarafından ve hangi tarihte gerçekleştirildiği bilgisiyle kaydedilmektedir. Bu yaklaşım, hem veri güvenilirliğini artırmakta hem de kullanıcı bazlı detaylı analizlerin yapılabilmesine olanak tanımaktadır. Ayrıca, giriş ekranı tasarımında kullanıcı dostu bir arayüz tercih edilerek, giriş işlemlerinin hızlı ve sorunsuz bir şekilde gerçekleştirilmesi sağlanmaktadır. Şekil 4.10 giriş ekranını göstermektedir.

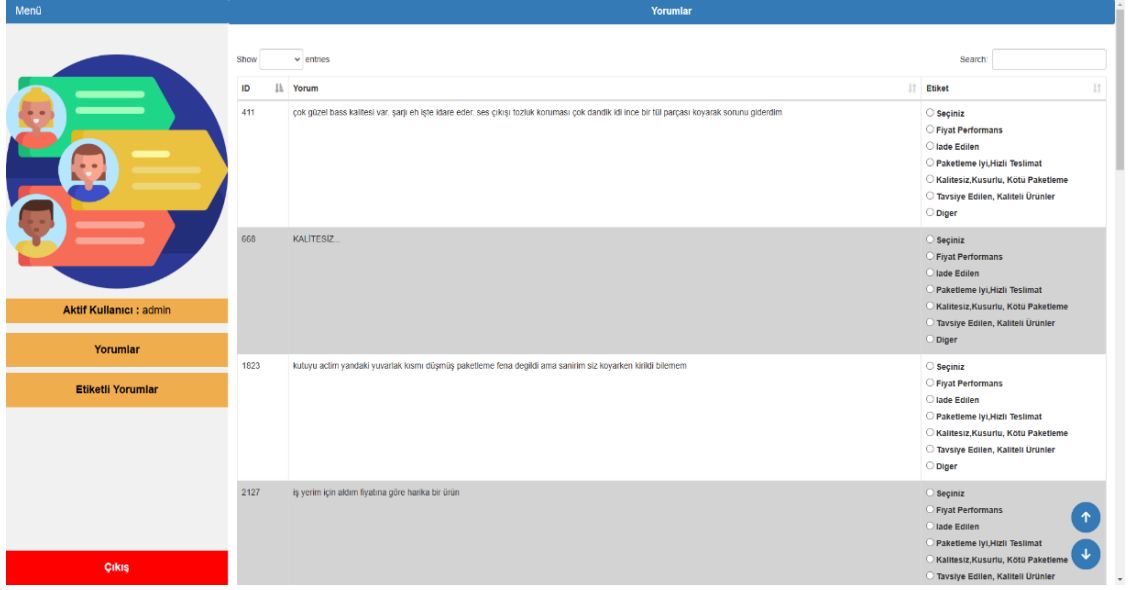
Şekil 4.10. Giriş ekranı

4.2.3. Etiketleme İşlemleri

Sisteme giriş yapan kullanıcılar, veritabanında bulunan rastgele yorumlarla karşılaşmaktadır. Bu yorumlar, sistemin algoritması tarafından rastgele bir şekilde atandığı için her kullanıcının eşit ve tarafsız bir şekilde yorumlara erişmesi sağlanmaktadır. Kullanıcılara sunulan etiketleme seçenekleri, yorumların kategorilere ayrılmasına olanak tanımaktadır. Belirlenen etiket kategorileri şu şekilde sıralanmaktadır:

- Fiyat Performans
- İade Edilen
- Paketleme İyi, Hızlı Teslimat
- Kalitesiz, Kusurlu, Kötü Paketleme
- Tavsiye Edilen, Kaliteli Ürünler
- Diğer

Her bir yorum, her kullanıcı için yalnızca bir kez ve yalnızca bir etiket seçilecek şekilde tasarlanmıştır. Bu tasarım, yorumların doğru bir şekilde kategorize edilmesini sağlamak ve birden fazla etikete izin verilmediği için veri tutarlılığını artırmaktadır. Kullanıcıların yorumlara kolay erişebilmesi ve hızlı bir şekilde etiketleme yapabilmesi amacıyla sade ve etkili bir kullanıcı arayüzü oluşturulmuştur. Şekil 4.11 etiketleme ekranını göstermektedir.



Şekil 4.11. Etiketleme ekranı

4.2.4. Etiketli Yorumlar

Sisteme giriş yapan kullanıcılar, kendi etiketledikleri yorumlara "Etiketli Yorumlar" sekmesi aracılığıyla erişebilmektedir. Bu sekmede, kullanıcıların etiketlediği yorumlar, eklenen etiketler ve etiketleme işleminin zamanı detaylı bir şekilde görüntülenmektedir. Bu tasarım sayesinde, kullanıcıların kendi katkılarını takip edebilmelerini sağlanmakta ve veri yönetiminde kullanıcı bazlı şeffaflık ile izlenebilirlik sağlanmaktadır.

Etiketli Yorumlar sekmesi, hem bireysel hem de toplu veri analizleri için faydalı bir yapı sunmaktadır. Kullanıcıların yalnızca kendi etiketleme işlemlerine erişim sağlayabilmesi, veri güvenliği açısından kritik bir öneme sahiptir. Şekil 4.12 kullanıcıların etiketlediği yorumları gördüğü etiketli yorumlar ekranını göstermektedir.

Menü		Etiketli Yorumlar		
 Aktif Kullanıcı : admin Yorumlar Etiketli Yorumlar Çıkış	Show 10 entries		Search:	
	ID	Yorum	Etiket	Tarih
	1894	ürünüm kusurlu geldi	Fiyat Performans	18.10.2024 22:17:42
	7150	paketlenme çok iyi sağlam geldi hızlı kargolandı.Celli beye teşekkür ederiz hediye olarak çok iyi bir seçenek tavsiye ederim	Fiyat Performans	18.10.2024 22:17:42
	8659	Fiyatına göre bir ürün tam numaranızı alın	İade Edilen	18.10.2024 22:17:42
	8736	sağ taraf çok iyiyken sol taraf çok sıkıyor ve acıttıyor, çok beğenmişim ama galiba İade	Paketleme İyi/Hızlı Teslimat	18.10.2024 22:17:42
	12756	Ürün güzel 39 giyiyorum ama yorumlarda bir numara büyük alın diyorodu ben de 40 aldım tam oktu rahat yürürken rahatsız eden herhangi bir şey yok	Tavsiye Edilen, Kaliteli Ürünler	18.10.2024 22:17:42
	15730	Daha önce alışveriş yaptığım memnun olduğum bir mağaza ama bu kez çok kötü.. 40 numara olan siparişimin bir tane 40 ölçen 39 geldi, ilk defa başıma böyle bir şey geliyor. Ayakkabıya gelsek, taraklı ayakkabı uzak durması gereken bir ürün ben de	Kalitesiz,Kusurlu, Kötü Paketleme	2.10.2024 16:00:54
	19238	güzel kargolamıştı	Tavsiye Edilen, Kaliteli Ürünler	18.10.2024 22:17:42
	19259	ten rengi istemişim ama başka renk geldi İade	Kalitesiz,Kusurlu, Kötü Paketleme	18.10.2024 22:17:42
	25908	Kalıba geniş biraz İade	Paketleme İyi/Hızlı Teslimat	18.10.2024 22:17:42
	27898	oversize fakat boyu kısa tam istediğim görüntüye ulaşamadım üzerinde İade	Kalitesiz,Kusurlu, Kötü Paketleme	18.10.2024 22:17:42
Showing 1 to 10 of 74 entries		Previous 1 2 3 4 5 ... 8 Next		

Şekil 4.12. Etiketli yorumlar ekranı

4.2.5. Etiketleme Süreci

Etiketleme işlemlerine başlamadan önce, sistemdeki yorumlar SQL tabanlı otomatik etiketleme algoritması kullanılarak belirli kurallar çerçevesinde önceden etiketlenmiştir. Bu süreçte, yorumların içerdiği kelimeler, ifade kalıpları ve puanlama dereceleri gibi kriterler temel alınmıştır. Örneğin:

Yorumun içeriğinde "Fiyat Performans", "F/P", "Fiyatına göre iyi" gibi ifadelerin bulunması ve yorum puanının 4 veya 5 olması durumunda, bu yorum "Fiyat Performans" kategorisi altında etiketlenmiştir.

Yorumun içeriğinde "İade" kelimesinin bulunması ve yorum puanının 1, 2 veya 3 olması durumunda, bu yorum "İade Edilen" kategorisine atanmıştır.

Bu süreçte toplamda 47 bin yorum otomatik olarak etiketlenmiş ve bu dağılımlar Çizelge 4.1'de detaylı bir şekilde sunulmuştur. Gerçek kullanıcılar tarafından gerçekleştirilen etiketleme işlemlerinden sonraki dağılımlar ise Çizelge 4.2'de gösterilmektedir. Bu iki dağılım karşılaştırıldığında, otomatik etiketleme sonuçlarının gerçek kullanıcı etiketleriyle örtüşme oranlarının analiz edilmesi mümkün hale gelmiştir.

Çizelge 4.1. Yorumların Otomatik Etiketlemeye Göre Dağılımları

Etiket	Adet
Fiyat Performans	8745
İade Edilen	9487
Paketleme İyi, Hızlı Teslimat	10085
Kalitesiz, Kusurlu, Kötü Paketleme	8325
Tavsiye Edilen, Kaliteli Ürünler	10358

Çizelge 4.2. Yorumların Kullanıcı Etiketlemelerine Göre Dağılımları

Etiket	Adet
Fiyat Performans	5240
İade Edilen	5506
Paketleme İyi, Hızlı Teslimat	6004
Kalitesiz, Kusurlu, Kötü Paketleme	12394
Tavsiye Edilen, Kaliteli Ürünler	14863
Diğer	2993

4.2.6. Kategorilerin Güvenilirliği ve Sonuçların Değerlendirilmesi

Web sitesi projesi sonucunda, otomatik olarak yapılan etiketleme işlemlerinin, gerçek kullanıcılar tarafından yapılan etiketler arasında kategori dağılımlarının değiştiği gözlemlenmiştir. Bu değişim, gerçek kullanıcı etiketlerinin, yorumların anlamlandırılması ve sınıflandırılması açısından daha güvenilir bir veri seti oluşturduğunu ortaya koymuştur. Kullanıcıların katkılarıyla elde edilen veri seti, hem yapay zeka hem de büyük veri analitiği çalışmaları için yüksek kaliteli bir temel oluşturmuştur. Bu süreçte, kullanıcı geri bildirimleri de dikkate alınarak sistem sürekli olarak geliştirilmiştir.

Sonuç olarak, sistemde gerçek kullanıcılar tarafından gerçekleştirilen etiketleme işlemleriyle daha tutarlı ve güvenilir bir veri seti oluşturulmuş ve yorumların analizi için sağlam bir temel sağlanmıştır. Bu yapı, yalnızca mevcut çalışma için değil, gelecekte benzer projeler için de değerli bir referans niteliği taşımaktadır.

4.3. Veri Setinin Dengelenmesi

Veri işleme sürecine, SQL ortamında etiketlenmiş kullanıcı yorumlarının CSV formatında dışa aktarılmasıyla başlanmıştır. Bu işlem, yorumların analize uygun bir veri yapısına dönüştürülmesini sağlamıştır. Ardından, veri setindeki sınıflar incelenmiş ve her bir sınıfın yorum sayısı analiz edilmiştir. 7.000'in altında yoruma sahip olan sınıflar belirlenerek, bu sınıfların temsil gücünü artırmak ve model performansını iyileştirmek amacıyla metin çoğaltma (text augmentation) yöntemi uygulanmıştır.

Metin çoğaltma sürecinde, Wordnet tabanlı bir araç kullanılmıştır. Belirlenen az örnekli sınıfların yorumları, anlamlarını koruyacak şekilde çoğaltılmış ve her bir sınıf için yeni yorumlar üretilmiştir. Bu işlem sonucunda elde edilen yeni yorumlar, veri setine entegre edilerek güncellenmiş bir veri çerçevesi oluşturulmuştur. Bu adım, hem sınıf çeşitliliğini artırmış hem de küçük sınıfların model öğrenim sürecinde daha etkili bir şekilde temsil edilmesini sağlamıştır.

Daha sonra, veri setindeki sınıf dengesizliğini gidermek için bir eşik değeri 12.500 tanımlanmıştır. Bu eşik, her sınıfın maksimum örnek sayısını belirleyerek büyük sınıfların veri setindeki baskınlığını azaltmayı amaçlamıştır. Sınıf örnek sayıları eşik değerini aşan durumlarda, rastgele örnekleme yöntemi kullanılarak sınıf boyutları eşitlenmiştir. Bu süreçte her bir sınıf dengelenmiş ve tüm sınıfların birleştirilmesiyle yeni bir veri seti oluşturulmuştur. Bu adım, modelin öğrenme sürecinde herhangi bir sınıfın diğerlerine kıyasla daha fazla etkili olmasını engellemiştir.

Son olarak, metin çoğaltma ve sınıf dengeleme adımları tamamlandıktan sonra orijinal veri seti yedeklenmiş ve genişletilmiş, dengeli güncel veri seti üzerinde çalışılmaya devam edilmiştir. Tüm bu işlemler, veri setini daha tutarlı ve analiz için uygun hale getirmiştir. Ayrıca, modelin performansını artırarak daha doğru ve güvenilir sonuçlar elde edilmesi için kritik bir rol oynamıştır.

4.4. Keşifsel Veri Analizi

Güncellenmiş veri setine göre, etiketlerin dağılımı belirli kategorilere göre sıralanmış ve her kategorinin frekansları hesaplanmıştır. Bu analiz, veri setinin genel yapısının anlaşılmasına katkı sağlamakta ve özellikle etiketler arasında dengesizlik olduğu durumlarda önemli bilgiler sunmaktadır. Çizelge 4.3'te, her etiketin kaç kez tekrar ettiği ve bu etiketlerin sıralı frekansları gösterilmektedir. Elde edilen veriler, hangi

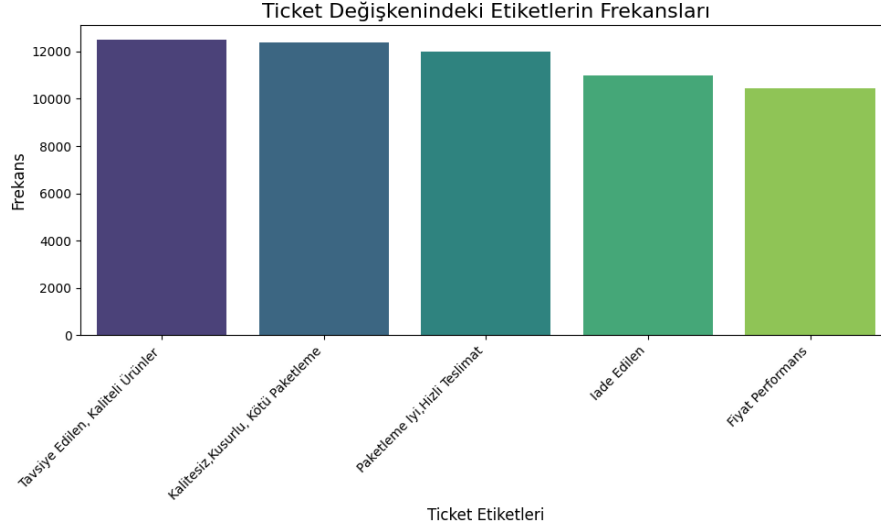
kategorilerin daha yaygın olduğunu ve hangi etiketlerin daha az sıklıkla kullanıldığını anlamak açısından değerli bir kaynak teşkil etmektedir.

Çizelge 4.3. Güncellenmiş Veri Seti Etiketlerine Göre Dağılımları

Etiket	Adet
Fiyat Performans	10462
İade Edilen	10990
Paketleme İyi, Hızlı Teslimat	12004
Kalitesiz, Kusurlu, Kötü Paketleme	12385
Tavsiye Edilen, Kaliteli Ürünler	12500

Güncellenmiş veri setine dayalı olarak oluşturulan Şekil 4.13'deki çubuk grafik, *ticket* değişkenindeki etiketlerin frekanslarını görsel olarak sunmaktadır. Grafik, her bir etiketin veri setinde ne sıklıkla yer aldığını açık bir şekilde göstermekte ve bu dağılımın görsel olarak daha kolay anlaşılmasına katkıda bulunmaktadır. Etiketler, sıralı bir şekilde düzenlenmiş olup, her birinin karşısındaki çubuklar ilgili frekans değerlerini temsil etmektedir.

Bu görselleştirme, veri setinin yapısına ilişkin derinlemesine bilgi sunmakta ve özellikle etiketler arasındaki yoğunluk farklılıklarını net bir biçimde gözler önüne sermektedir. Bu tür görsel analizler, veri bilimi ve istatistiksel modelleme süreçlerinde kritik bir öneme sahiptir; verilerin hangi sınıflara daha fazla yoğunlaştığını veya dengesizliklerin mevcut olup olmadığını belirleme açısından değerli bir araçtır. Çubuk grafik, verilerin görsel olarak temsil edilmesini sağlamakla birlikte, araştırmacılara hangi etiketlerin daha yaygın, hangilerinin ise daha az temsil edildiği konusunda ipuçları sunarak veri setinin analizini ve yorumlanmasını kolaylaştırmaktadır.



Şekil 4.13. Güncellenmiş veri setine dayalı olarak oluşturulan çubuk grafik

Güncellenmiş veri setine dayalı olarak gerçekleştirilen kelime frekansı analizi, yorumlarda en sık kullanılan kelimeleri ve bu kelimelerin sıklıklarını ortaya koymaktadır. Çizelge 4.4’de her bir kelimenin tekrar sayıları belirlenerek veri setindeki en yaygın kelimeler sıralanmıştır. Bu analiz, metin verisinin genel içeriği hakkında ayrıntılı bilgi sunmakta ve belirli temaların, kelimelerin ya da konuların ne ölçüde yaygın olduğunu tespit etmeye olanak tanımaktadır.

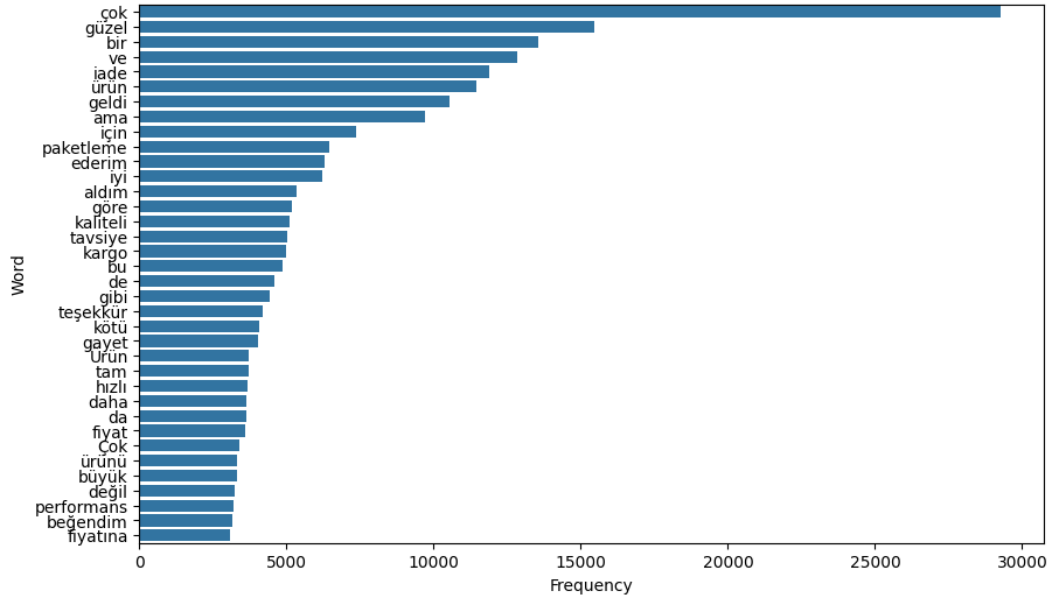
Kelime frekansları, yorumların genel yapısını anlamada ve veri setindeki anahtar kelimeleri öne çıkarmada etkili bir yöntemdir. Çizelgede yer alan sonuçlar, hangi kelimelerin daha fazla vurgulandığını ve öne çıktığını açık bir şekilde göstermektedir. Bu analiz, metin analizi ve doğal dil işleme (NLP) çalışmaları kapsamında kritik bir öneme sahip olup, özellikle veri odaklı araştırmalarda değerli bir içgörü sağlamaktadır.

Çizelge 4.4. En Sık Kullanılan Kelimeler

Kelime	Adet
Çok	29286
Güzel	15471
Bir	13568
ve	12857
iade	11900
ürün	11461
geldi	10543
ama	9702
İçin	7351
paketleme	6460

Kelime frekansı analizi sonucunda, belirli bir eşik değerin üzerinde yer alan kelimeler filtrelenmiş ve böylece daha odaklı bir görselleştirme elde edilmiştir. Bu tür bir filtreleme işlemi, düşük frekansa sahip kelimeleri dışlayarak veri setindeki yaygın ve anlamlı kelimelerin öne çıkarılmasına olanak tanımıştır. Sadece belirli bir sıklığın üzerindeki kelimelerin değerlendirilmesi, veri setinin daha temsil edici ve analitik açıdan zengin unsurlarına odaklanmayı sağlamıştır.

Şekil 4.14’de görülen çubuk grafiği, yüksek frekansa sahip kelimelerin görsel bir temsilini sunmaktadır. Her bir kelimenin sıklığı, grafik üzerinde net bir şekilde gösterilerek, veri setinin temel yapısına dair çıkarımlar yapılmasını kolaylaştırmıştır. Özellikle sık kullanılan kelimelerin ön plana çıkması, metnin ana temaları ve odak noktaları hakkında araştırmacılara önemli ipuçları sunmaktadır. Bu tür görselleştirmeler, metin madenciliği ve NLP çalışmalarında, verinin derinlemesine analizine yönelik kritik bir araç olarak değerlendirilmektedir.



Şekil 4.14. Yüksek frekansa sahip kelimelerin çubuk grafiği

Yorumların birleştirilmesi ve ardından kelime bulutu (WordCloud) görselleştirmesi, veri setindeki metinlerin temel temalarını ve öne çıkan kelimeleri vurgulamak için etkili bir yöntem olarak kullanılmaktadır. Bu süreçte, tüm yorumlar tek bir metin halinde birleştirilerek, veri setindeki tüm kelimeler bütüncül bir biçimde analiz edilmeye uygun hale getirilmiştir. Elde edilen birleşik metin üzerinden oluşturulan kelime bulutu, en sık kullanılan kelimeleri daha büyük ve belirgin boyutlarda sunarken, nadir kullanılan kelimeleri daha küçük boyutlarda göstermektedir. Bu yaklaşım, yorumlar

- **Noktalama İşaretlerinin Kaldırılması:** Noktalama işaretleri, metinsel analizde sıklıkla gereksiz veri olarak değerlendirilir. Dolayısıyla, bir başka düzenli ifade kullanılarak noktalama işaretleri metinden çıkarılır.
- **Durak Kelimelerin (Stop Words) Çıkarılması:** Türkçe dilinde sıkça kullanılan ve metnin genel anlamını etkilemeyen durak kelimeler ("ve", "ile", "de" gibi) bir stop words sözlüğü yardımıyla metinden temizlenir. Bu işlem, yalnızca anlamlı kelimelerin analiz edilmesini sağlayarak işlem yükünü azaltır.
- **Kelimelerin Köklerine İndirgenmesi (Lemmatization):** Son adımda, kelimeler bir lemmatizer (örneğin, `WordNetLemmatizer`) yardımıyla köklerine indirgenir. Bu işlem, farklı çekimlere sahip kelimeleri tek bir kök formda birleştirir ve metnin daha homojen bir yapıya kavuşmasını sağlar.

Bu süreçlerin her biri, ham metin verisini daha yapılandırılmış ve anlamlı bir hale dönüştürmeyi hedeflemektedir. Özellikle lemmatization gibi adımlar, kelimeler arasındaki ilişkiyi güçlendirerek, analiz modellerinin doğruluğunu artırır. Bu tür ön işlemler, metin madenciliği ve NLP projelerinde temel bir gereklilik olarak değerlendirilir.

4.6. Veri Modelleme

Bu aşamada, metin verisinin sayısal verilere dönüştürülmesi ve ardından eğitim ve test setlerine ayrılması işlemleri gerçekleştirilmiştir. İlk olarak, TF-IDF vektörleştirici yöntemi kullanılmakta ve metin verisindeki kelimelerin önem düzeyleri hesaplanmaktadır. Metin verilerinin bilgisayarlar tarafından işlenebilmesi için sayısal verilere dönüştürülmesi gerekmektedir; çünkü makine öğrenmesi algoritmaları yalnızca sayısal veriler üzerinde işlem yapabilmektedir. Bu dönüşüm, her bir kelimenin bir vektör (sayısal bir dizi) ile temsil edilmesini sağlamak ve böylece metnin anlamı ile yapısının sayısal biçimde analiz edilmesine olanak tanımaktadır.

TF-IDF yöntemi, bir kelimenin belirli bir belgede ne kadar önemli olduğunu belirlemek için iki temel ölçüt kullanmaktadır: kelimenin o belge içindeki geçiş sıklığı (Term Frequency, TF) ve kelimenin tüm belge kümesindeki nadirliği (Inverse Document Frequency, IDF). Bu yöntem, sıkça geçen ancak genelde anlam taşımayan kelimelere (örneğin "ve", "bu") düşük ağırlık atamakta, belirli belgelerde sık kullanılan ancak nadir

kelimelere ise yüksek ağırlık vermektedir. Böylece, modelin yalnızca anlamlı ve ayırt edici özelliklere odaklanması sağlanmaktadır.

Devamında; sayısal temsiller (X) ile hedef değişken olan etiketler (Y) arasındaki veriler ayrılmıştır. `train_test_split()` fonksiyonu kullanılarak veri kümesi, eğitim (%80) ve test (%20) setlerine bölünmüştür. Bu işlem, modelin eğitim sürecinde kullanılan verilerin, modelin performansını objektif bir biçimde değerlendirebilmek amacıyla test verisi ile karşılaştırılmasına olanak tanımaktadır. Eğitim verisi, modelin parametrelerini öğrenmesi için kullanılırken, test verisi modelin genelleme yeteneğini ölçmek için ayrılmıştır. Veri ayırma işlemi, aşırı öğrenme (overfitting) problemini engellemeye yardımcı olur ve modelin doğruluğunu daha güvenilir bir şekilde değerlendirmeyi mümkün kılar.

Bu adımlar, metin verisinin daha etkin bir biçimde analiz edilmesi için zorunlu olan işlemlerdir. Özellikle NLP ve makine öğrenmesi alanlarında, verinin sayısal temsillere dönüştürülmesi ve doğru bir şekilde bölünmesi, modelin başarısını doğrudan etkileyen temel faktörlerdendir. Metin verisi sayısal verilere dönüştürülerek, bilgisayarların dilin yapısal özelliklerini anlaması ve işlem yapması sağlanırken, TF-IDF yöntemi metinlerin içeriklerini anlamada ve anahtar kelimelerin öne çıkarılmasında etkin bir rol oynamaktadır. Eğitim ve test setlerine ayırma işlemi ise modelin doğruluğunu ve güvenilirliğini artırmakta, böylece modelin genel geçer sonuçlar üretmesi sağlanmaktadır.

4.7. Yapay Zeka Modelleri

Bu başlık altında, derin öğrenme (CNN, RNN, LSTM, BiLSTM, GRU, BERT) ve geleneksel makine öğrenmesi (Lojistik Regresyon, Rastgele Orman) modelleriyle elde edilen performans sonuçları ayrıntılı bir şekilde değerlendirilecektir. Her bir model, tezin amacına göre eğitilmiş ve çeşitli değerlendirme metrikleri kullanılarak analiz edilmiştir. Bu metrikler arasında doğruluk, kesinlik, duyarlılık, F1 Skoru, Log Loss ve ROC AUC yer almaktadır. Model eğitimi sırasında, doğruluk ve diğer istatistiksel ölçütler gibi temel metrikler üzerinden modelin başarısı ölçülürken, ROC eğrisi gibi grafiksel değerlendirmeler, modelin sınıflandırma başarısını ve hatalarını görselleştirmektedir. Ayrıca, her bir modelin Konfüzyon Matrisi hesaplanarak, sınıflandırma hatalarının detaylı bir şekilde analiz edilmesi sağlanmıştır.

Her model için kullanılan veri seti, eğitim ve test alt kümelerine ayrılmıştır. Veri ayrımı, %80 eğitim ve %20 test oranında gerçekleştirilmiş olup, bu dağılım sonucunda

elde edilen performans sonuçları detaylı bir şekilde analiz edilmiştir. Modellerin başarısını artırmak amacıyla, kullanılan yapay zeka modellerinin hiperparametreleri manuel hiperparametre ayarlaması yöntemiyle optimize edilmiştir. Bu süreçte, her bir model için farklı hiperparametre kombinasyonları test edilmiş ve en iyi sonuçları veren ayarlar seçilmiştir.

İlk olarak, dengelenmemiş ve dengelenmiş veri setlerine ilişkin genel bir değerlendirme yapmak amacıyla tüm modellerin sınıflandırma doğrulukları incelenmiştir. Çizelge 4.5, dengelenmemiş ve dengelenmiş veri setleri için test sonuçları sunmaktadır.

Çizelge 4.5. Veri Seti Dengeleme Test Sonuçları

Veri Seti	CNN	RNN	LSTM	BiLSTM	GRU	BERT	Lojistik Regresyon	Rastgele Orman
Dengelenmemiş Veri Seti Doğruluk	0.81	0.80	0.81	0.81	0.75	0.81	0.79	0.79
Dengelenmiş Veri Seti Doğruluk	0.85	0.84	0.86	0.85	0.78	0.89	0.85	0.84

Veri seti dengeleme işlemleri uygulandıktan sonra, modellerin başarı oranlarında anlamlı bir iyileşme gözlenmiştir. Bu süreçte, verilerin çeşitlendirilmesi sayesinde modellerin eğitim sürecinde daha genelleştirilebilir örneklerle çalışması sağlanmış ve doğruluk oranlarında belirgin bir artış elde edilmiştir. Elde edilen bu sonuçlar doğrultusunda, her model dengelenmiş veri seti kullanılarak detaylı metriklerle analiz edilmiştir.

Modeller için grafikler ve sonuçlar doğrultusunda, etiket numaralarının hangi sınıfa karşılık geldiği aşağıdaki gibi tanımlanmıştır. Bu sıralama, hem ROC eğrileri, hem de sınıf bazlı metrikler ve diğer modellerde kullanılan “LabelEncoder” ile tutarlı bir şekilde alfabetik sıralamaya dayanmaktadır:

- **0:** Fiyat Performans
- **1:** İade Edilen
- **2:** Kalitesiz, Kusurlu, Kötü Paketleme
- **3:** Paketleme İyi, Hızlı Teslimat
- **4:** Tavsiye Edilen, Kaliteli Ürünler

Bu numaralandırma, tüm metrik hesaplamalarında, görselleştirmelerde ve analizlerde standart olarak korunmuştur ve her sınıf, kendine özgü metriklerle bu sıralamaya göre değerlendirilmiştir.

4.7.1. CNN

Tez çalışmasında, metin sınıflandırma problemi için sadeleştirilmiş bir derin öğrenme modeli tasarlanmıştır. Giriş verileri doğrudan bir konvolüsyonel (Conv1D) katman ile işlenmeye başlanmıştır. İlk katmanda, 32 filtre ve 5x5 boyutunda kernel kullanılarak metinler üzerinde örüntüler öğrenilmiştir. Daha sonra MaxPooling katmanı ile boyutlar küçültülmüş ve önemli özellikler öne çıkarılmıştır.

Veriler, konvolüsyonel katmanlar sonrası düzleştirilmiş (Flatten) ve yoğun katmanlar (Dense) ile sınıflandırma işlemi gerçekleştirilmiştir. Modelin çıkış katmanı, çok sınıflı sınıflandırma için softmax aktivasyon fonksiyonu ile yapılandırılmıştır. Eğitim süreci, `sparse_categorical_crossentropy` kayıp fonksiyonu ve 'Adam' optimizasyon algoritması ile gerçekleştirilmiştir. Şekil 4.16'da CNN modeline ait pseudo kod, Çizelge 4.6'da ise bu model kurulurken kullanılan hiperparametre bilgileri verilmiştir.

Çizelge 4.7'de belirtilen modelin kullanılması sonucunda elde edilen test sonuçları gösterilmektedir. Çizelge 4.7 incelendiğinde, modelin ortalama doğruluk değeri olan 0.85 ile genel olarak başarılı bir performans sergilediği görülmektedir. F1 skoru 0.85, modelin hem kesinlik hem de duyarlılık açısından başarılı olduğunu ve sınıflandırma performansının genel olarak tatmin edici olduğunu göstermektedir. ROC AUC (0.97) değeri, modelin sınıflar arasındaki ayrımı çok iyi bir şekilde yapabildiğini ve doğru pozitifleri başarıyla sınıflandırabildiğini ortaya koymaktadır. Bunun yanı sıra, Log Loss değeri 0.43 ile modelin tahmin ettiği olasılıkların gerçek sınıflarla büyük ölçüde uyumlu olduğunu göstermektedir. Şekil 4.17 bu sonuçların elde edildiği konfüzyon matrisini vermektedir.

```

1  Baslat:
2      Veriyi özellik ve hedef değişken olarak ayır:
3          Özellik (X_raw): cleaned_comments sütunu
4          Hedef değişken (y): ticket sütunu
5      Veriyi eğitim ve test setlerine ayır:
6          Eğitim seti : %80 Test seti: %20
7  Ön İşleme:
8      Verileri aynı uzunlukta yapabilmek için doldurma işlemi uygula:
9          Maksimum uzunluk (maxlen): 200 kelime
10 CNN modeli tanımla:
11     Gömülü katman (Embedding):
12         Giriş boyutu (input_dim): 10.000
13         Çıkış boyutu (output_dim): 32
14         Giriş uzunluğu (input_length): 200
15     Konvolüsyonel katman (Conv1D):
16         Filtre sayısı: 32
17         Kernel boyutu: 5
18         Aktivasyon: relu
19     MaxPooling katmanı:
20         Havuzlama boyutu (pool_size): 2
21     Düzleştirme katmanı (Flatten)
22     Yoğun katman (Dense):
23         Birim sayısı: 128
24         Aktivasyon: relu
25     Çıkış katmanı (Dense):
26         Sınıf sayısı: len(np.unique(y_train))
27         Aktivasyon: softmax
28     Modeli derle:
29         Kayıp fonksiyonu: sparse_categorical_crossentropy
30         Optimizasyon algoritması: adam
31 Modeli eğitim seti ile eğit:
32     Epoch sayısı: 5
33     Batch boyutu: 64
34 Değerlendirme:
35     Test verisi üzerinde tahminler yap:
36         Sınıf tahminleri: argmax(axis=1)
37         Olasılık tahminleri: predict_proba(X_test_pad)
38     Model performans metriklerini hesapla:
39         Doğruluk (accuracy), Kesinlik (precision), Duyarlılık (recall), F1 skoru, ROC AUC, Log Loss
40     Sınıf bazlı metrikler:
41         Precision, Recall ve F1 Score
42 Görselleştirme:
43     Her bir sınıf için ROC eğrisi ve AUC değerlerini hesapla
44     ROC eğrilerini görselleştir
45     Konfüzyon matrisini hesapla ve görselleştir

```

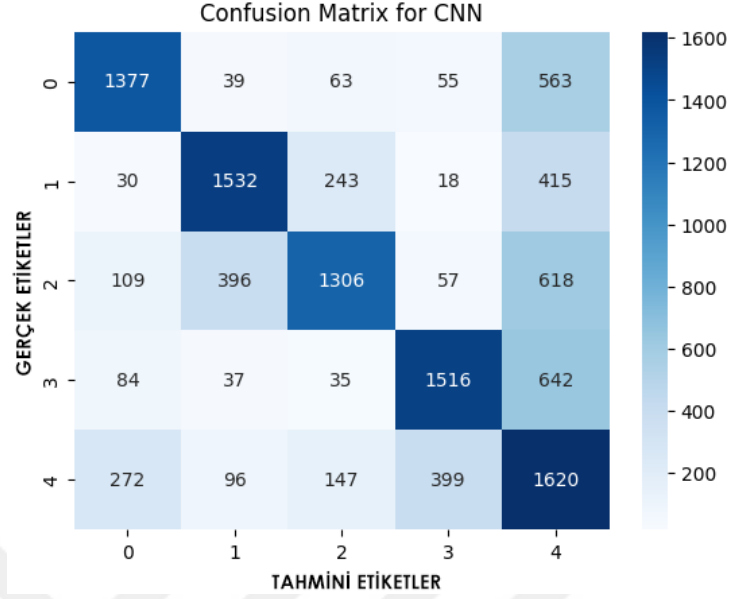
Şekil 4.16. CNN Modeli Pseudo Kod

Çizelge 4.6. CNN Hiperparametre Bilgileri

Parametre Adı	Değer
Giriş Boyutu	10.000
Çıkış Boyutu	32
Giriş Uzunluğu	200
Filtre Sayısı	32
Kernel Boyutu	5x5
Havuzlama Boyutu	2
Yoğun Katman Birim Sayısı	128
Yoğun Katman Aktivasyon	Relu
Çıkış Katmanı Aktivasyon	Softmax
Epoch Sayısı	5
Batch Boyutu	64

Çizelge 4.7. CNN Test Sonuçları

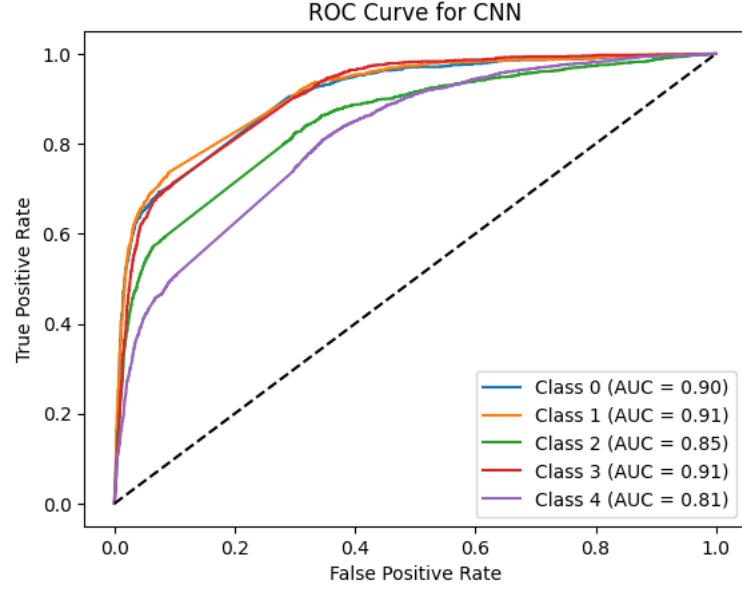
Sınıf	Kesinlik	Duyarlılık	F1 Skoru	Doğruluk	ROC AUC	Log Loss
Fiyat Performans	0.87	0.86	0.86			
İade Edilen	0.89	0.92	0.91			
Kalitesiz, Kusurlu, Kötü Paketleme	0.85	0.87	0.86			
Paketleme İyi, Hızlı Teslimat	0.81	0.86	0.83			
Tavsiye Edilen, Kaliteli Ürünler	0.81	0.73	0.77			
Ortalama Genel Sonuç	0.85	0.85	0.85	0.85	0.97	0.47



Şekil 4.17. CNN modeli ile elde edilen konfüzyon matrisi

Hem Çizelge 4.7, hem de Şekil 4.17'ye göre sınıf bazlı test sonuçları, modelin genel olarak iyi bir performans sergilediğini ancak bazı sınıflarda daha fazla iyileştirme yapılması gerektiğini göstermektedir. "Kalitesiz, Kusurlu, Kötü Paketleme" sınıfında kesinlik ve duyarlılık değerlerinin dengeli olduğu görülmekte, bu da modelin ilgili sınıfı tanımlamada genel olarak başarılı olduğunu göstermektedir. Ancak, sınıfın ayırımında bazı karışıklıkların yaşandığı da gözlemlenmiştir. "Paketleme İyi, Hızlı Teslimat" sınıfında ise kesinlik değeri biraz daha düşük olsa da, modelin duyarlılığı oldukça yüksek ve bu da sınıfın doğru tespit edilmesinde etkili olduğunu göstermektedir. Son olarak, "Tavsiye Edilen, Kaliteli Ürünler" sınıfında, kesinlik ve F1 skoru daha düşük olup, modelin bu sınıfı tanıma konusunda bazı zorluklar yaşadığı ve hatalı tahminler yaptığı gözlemlenmektedir. Konfüzyon matrisine bakıldığında, özellikle "Tavsiye Edilen, Kaliteli Ürünler" sınıfı ile ilgili hataların fazla olduğu ve modelin bu sınıfa ait örnekleri diğer sınıflarla karıştırdığı görülmektedir. Genel olarak modelin bazı sınıflarda başarılı olduğunu, ancak özellikle daha az temsil edilen sınıflarda hata oranının yüksek olduğunu söylemek mümkündür.

Şekil 4.18'de sunulan CNN modeline ait ROC eğrileri, sınıflar arasındaki ayırımın doğru bir şekilde görselleştirilmesine olanak tanımaktadır.



Şekil 4.18. CNN modeli ile elde edilen ROC eğrileri

4.7.2. RNN

RNN modelinin tasarımında, giriş verilerini işlemek için 128 birimli ve %20 dropout oranına sahip bir SimpleRNN katmanı kullanılmıştır. Dropout, öğrenme sürecinde overfitting riskini azaltmak amacıyla eklenmiştir. SimpleRNN katmanının ardından, çok sınıflı sınıflandırma görevleri için uygun bir şekilde yapılandırılmış, softmax aktivasyon fonksiyonuna sahip bir çıkış katmanı eklenmiştir. Modelin eğitimi, toplam 5 epoch boyunca ve her bir iterasyonda 64 veri örneğinden oluşan batch boyutlarıyla gerçekleştirilmiştir. Eğitim sırasında kayıp fonksiyonu olarak 'sparse_categorical_crossentropy', optimizasyon algoritması olarak ise 'Adam' tercih edilmiştir. Eğitim süreci boyunca, modelin doğruluk metriği ve doğrulama verisi üzerindeki performansı izlenmiştir. Şekil 4.19'da RNN modeline ait pseudo kod verilmiştir.

```

1  Başlat:
2  |   Özellik ve hedef değişkeni ayır:
3  |   |   Özellikler (X_raw): cleaned_comments
4  |   |   Hedef değişken (y): ticket
5  |   Veriyi eğitim ve test setlerine ayır:
6  |   |   Eğitim seti : %80 Test seti: %20
7  Ön İşleme:
8  |   Verileri aynı uzunlukta yapmak için doldurma işlemi uygula:
9  |   |   Maksimum uzunluk (maxlen): 200 kelime
10 Simple RNN modeli tanımla:
11 |   Gömülü katman (Embedding):
12 |   |   Giriş boyutu (input_dim): 10.000
13 |   |   Çıkış boyutu (output_dim): 128
14 |   |   Giriş uzunluğu (input_length): 200
15 |   Simple RNN katmanı:
16 |   |   Birim sayısı (units): 128
17 |   |   Dropout: 0.2
18 |   |   Tekrarlı dropout: 0.2
19 |   Dense katmanı:
20 |   |   Birim sayısı (units): len(np.unique(y_train))
21 |   |   Aktivasyon fonksiyonu: softmax
22 Modeli derle:
23 |   Kayıp fonksiyonu (loss): 'sparse_categorical_crossentropy'
24 |   Optimizasyon algoritması (optimizer): 'adam'
25 Modeli eğitim verisiyle eğit:
26 |   Epoch sayısı: 5
27 |   Batch boyutu: 64
28 Tahminler:
29 |   Test verisi üzerinde tahmin yap: y_pred_rnn = model.predict(X_test_pad)
30 |   Olasılıkları elde et: y_pred_rnn_proba = model.predict(X_test_pad)
31 Performans metrikleri:
32 |   Doğruluk (Accuracy), Kesinlik (Precision), Duyarlılık (Recall), F1 Skoru, ROC AUC, Log Loss
33 Sınıf Bazlı Analiz:
34 |   Precision, Recall ve F1 Skoru.
35 Görselleştirme:
36 |   Sınıf başına ROC eğrisi ve AUC.
37 |   Eğrileri çiz ve rastgele tahmin çizgisi ekle.
38 |   Konfüzyon matrisini hesapla ve görselleştir.

```

Şekil 4.19. RNN Modeli Pseudo Kod

Çizelge 4.8’de RNN modeline ait kullanılan hiperparametre bilgileri verilmiştir.

Çizelge 4.8. RNN Hiperparametre Bilgileri

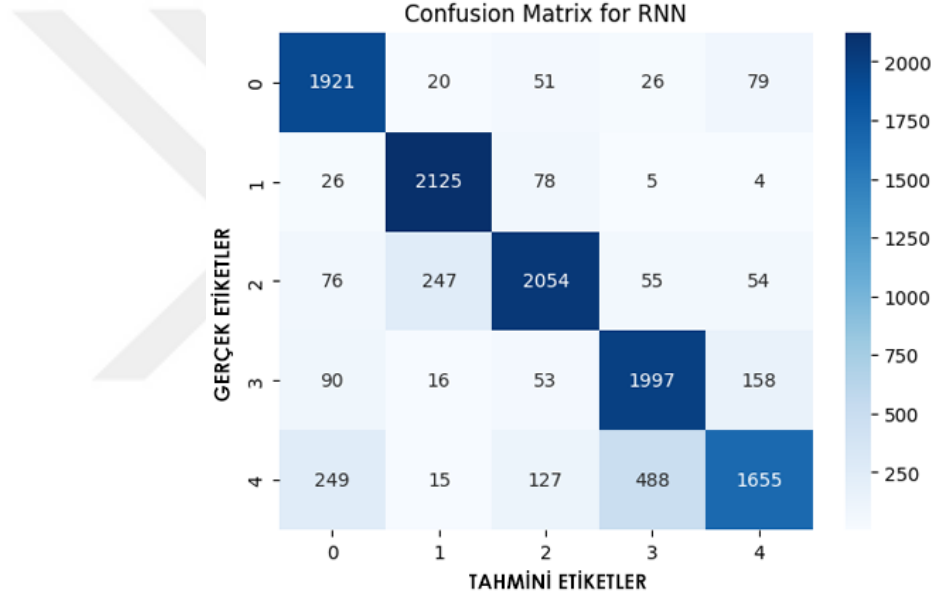
Parametre Adı	Değer
Giriş Boyutu	10.000
Çıkış Boyutu	128
Giriş Uzunluğu	200
Simple RNN Birim Sayısı	128
Dropout	0.2
Tekrarlı Dropout	0.2
Dense Katman Aktivasyon	Softmax
Epoch Sayısı	5
Batch Boyutu	64

Çizelge 4.9’da RNN modeli ile elde edilen test sonuçları gösterilmektedir.

Çizelge 4.9. RNN Test Sonuçları

Sınıf	Kesinlik	Duyarlılık	F1 Skoru	Doğruluk	ROC AUC	Log Loss
Fiyat Performans	0.81	0.92	0.86			
İade Edilen	0.88	0.95	0.91			
Kalitesiz, Kusurlu, Kötü Paketleme	0.87	0.83	0.85			
Paketleme İyi, Hızlı Teslimat	0.78	0.86	0.82			
Tavsiye Edilen, Kaliteli Ürünler	0.85	0.65	0.74			
Ortalama Genel Sonuç	0.84	0.84	0.84	0.84	0.96	0.50

Şekil 4.20’deki konfüzyon matrisi ise modelin hangi sınıflarda hata yaptığını ortaya koymuş ve modelin güçlü performansını detaylandırmıştır.



Şekil 4.20. RNN modeli ile elde edilen konfüzyon matrisi

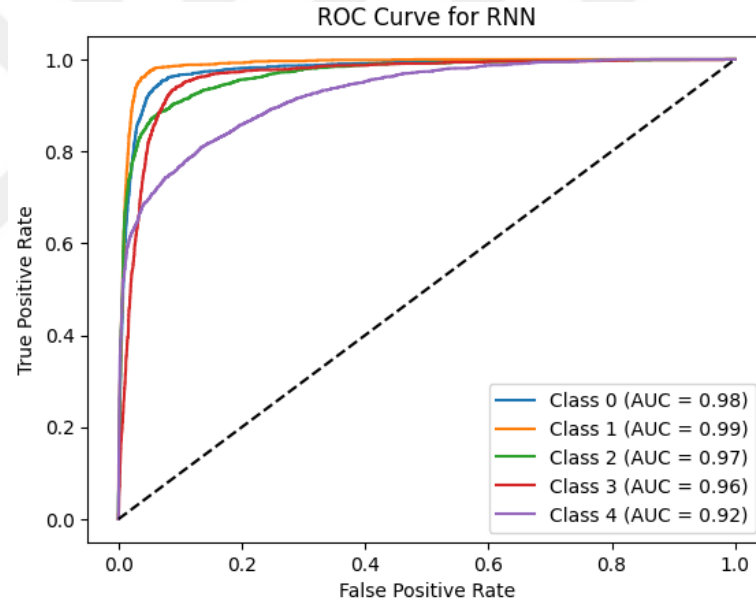
Çizelge 4.9’deki sonuçlar, modelin yüksek doğruluk ve güvenilirlik sağladığını, pozitif sınıfları doğru tahmin etmede başarılı olduğunu göstermektedir. Ayrıca, modelin ayırım gücü oldukça yüksek olup, tahminlerinin olasılık değerleri de doğru sınıflarla uyumludur.

RNN modelinin sınıf bazlı performans sonuçları ve konfüzyon matrisi, modelin genel anlamda iyi bir sınıflandırma performansı sergilediğini ancak bazı sınıflarda karışıklıklar yaşadığını göstermektedir. “İade Edilen” ve “Fiyat Performans” sınıfları, yüksek kesinlik, duyarlılık ve F1 skoru değerleri ile modelin en başarılı olduğu sınıflar olarak öne çıkmaktadır. Bu sınıflarda, modelin hem doğru pozitif oranı yüksek hem de yanlış pozitif oranı düşük kalmıştır. “Kalitesiz, Kusurlu, Kötü Paketleme” ve “Paketleme

İyi, Hızlı Teslimat” sınıflarında, duyarlılık ve kesinlik değerleri dengeli olmakla birlikte, sınıflar arasında bazı karışıklıkların yaşandığı gözlemlenmektedir. Özellikle “Tavsiye Edilen, Kaliteli Ürünler” sınıfında modelin duyarlılık skoru daha düşük kalmış ve bu durum, bu sınıfa ait örneklerin diğer sınıflar ile karıştırıldığını göstermektedir.

Şekil 4.20’de verilen konfüzyon matrisi incelendiğinde, özellikle “Tavsiye Edilen, Kaliteli Ürünler” sınıfına ait örneklerin diğer sınıflarla daha fazla karıştığı ve hatalı tahminlerin yoğunlaştığı görülmektedir. Ayrıca, diğer sınıflarda da küçük oranlarda karışıklıklar yaşanmış olsa da model, doğru tahminlerde genellikle tutarlı olmuştur. Bu durum, modelin bazı sınıflar arasındaki semantik yakınlıklar veya veri dengesizliklerinden etkilenmiş olabileceğini göstermektedir.

Modelin Şekil 4.21’deki ROC eğrileri, sınıflar arasındaki doğru ayırımın görselleştirilmesine olanak tanımaktadır.



Şekil 4.21. RNN modeli ile elde edilen ROC eğrileri

4.7.3. LSTM

Modelin ilk katmanı, verilerden anlamlı ilişkiler öğrenebilmek amacıyla 128 birimli bir LSTM katmanı olarak yapılandırılmıştır. LSTM katmanında kullanılan %20 dropout oranı, modelin öğrenme sürecinde aşırı öğrenmeyi (overfitting) önlemeye yardımcı olmaktadır. Bu katman, metinlerin sıralı yapısını ve bağlamını dikkate alarak daha derin özelliklerin öğrenilmesini sağlar.

Modelin son katmanı ise çok sınıflı sınıflandırma problemleri için softmax aktivasyon fonksiyonu ile tasarlanmıştır. Bu sayede, her bir giriş örneğinin belirli bir sınıfa ait olma olasılığı tahmin edilmektedir. Model, toplam 5 epoch boyunca eğitilmiş ve her bir iterasyonda 64 veri örneği işlenmiştir. Eğitim sırasında kayıp fonksiyonu olarak, çok sınıflı sınıflandırmalar için uygun bir seçim olan 'sparse_categorical_crossentropy' kullanılmıştır. Ayrıca, optimizasyon işlemleri için 'adam' algoritması tercih edilmiştir. Bu algoritma, model parametrelerini daha hızlı ve etkili bir şekilde güncelleyerek eğitim sürecini iyileştirmektedir.

Eğitim sürecinde modelin doğruluk metriği ve doğrulama verisi üzerindeki başarıyı düzenli olarak izlenmiştir. Bu süreç sonunda modelin, metin verileri üzerinden etkili sınıflandırmalar yapabildiği gözlemlenmiştir. Şekil 4.22’de LSTM modeline ait pseudo kod verilmiştir. Çizelge 4.10’da LSTM modeline ait kullanılan hiperparametre bilgileri verilmiştir.

```

1  Başlat:
2  | Veriyi özellik ve hedef değişken olarak ayır:
3  |   Özellik (X_raw): cleaned_comments sütunu
4  |   Hedef değişken (y): ticket sütunu
5  | Veriyi eğitim ve test setlerine ayır:
6  |   Eğitim seti : %80 Test seti: %20
7  On İşleme:
8  | Verileri aynı uzunlukta yapabilmek için doldurma işlemi uygula:
9  |   Maksimum uzunluk (maxlen): 200 kelime
10 Model:
11 | LSTM modeli tanımla:
12 |   Gömülü katman (Embedding):
13 |     Giriş boyutu (input_dim): 10.000
14 |     Çıkış boyutu (output_dim): 128
15 |     Giriş uzunluğu (input_length): 200
16 |   LSTM katmanı:
17 |     Birim sayısı (units): 128
18 |     Dropout: 0.2
19 |     Tekrarlayan dropout: 0.2
20 |   Çıkış katmanı (Dense):
21 |     Sınıf sayısı: len(np.unique(y_train))
22 |   Aktivasyon: softmax
23 Modeli derle:
24 | Kayıp fonksiyonu: sparse_categorical_crossentropy
25 | Optimizasyon algoritması: adam
26 Eğitim:
27 | Modeli eğitim seti ile eğit:
28 |   Epoch sayısı: 5
29 |   Batch boyutu: 64
30 Değerlendirme:
31 | Test verisi üzerinde tahminler yap:
32 |   Sınıf tahminleri: argmax(axis=1)
33 |   Olasılık tahminleri: predict_proba(X_test_pad)
34 | Model performans metriklerini hesapla:
35 |   Doğruluk (accuracy)
36 |   Kesinlik (precision)
37 |   Duyarlılık (recall)
38 |   F1 skoru
39 |   ROC AUC
40 |   Log Loss
41 Sınıf bazlı metrikler:
42 |   Precision, Recall ve F1 Score
43 Görselleştirme:
44 |   ROC eğrisini çiz:
45 |     Her bir sınıf için ROC eğrisi ve AUC değerlerini hesapla
46 |     ROC eğrilerini görselleştir
47 |   Konfüzyon matrisini hesapla ve görselleştir

```

Şekil 4.22. LSTM Modeli Pseudo Kod

Çizelge 4.10. LSTM Hiperparametre Bilgileri

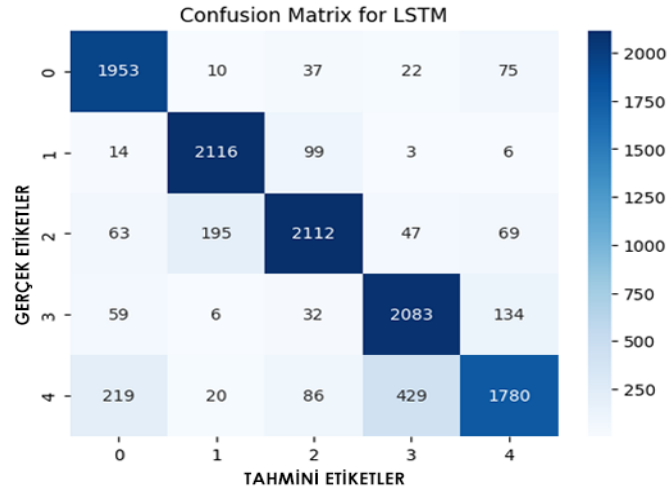
Parametre Adı	Değer
Giriş Boyutu	10.000
Çıkış Boyutu	128
Giriş Uzunluğu	200
LSTM Katmanı Birim Sayısı	128
Dropout	0.2
Tekrarlı Dropout	0.2
Dense Katman Aktivasyon	Softmax
Epoch Sayısı	5
Batch Boyutu	64

Çizelge 4.11’de LSTM modeli ile elde edilen test sonuçları gösterilmektedir.

Çizelge 4.11. LSTM Test Sonuçları

Sınıf	Kesinlik	Duyarlılık	F1 Skoru	Doğruluk	ROC AUC	Log Loss
Fiyat Performans	0.85	0.93	0.89			
İade Edilen	0.90	0.95	0.92			
Kalitesiz, Kusurlu, Kötü Paketleme	0.89	0.85	0.87			
Paketleme İyi, Hızlı Teslimat	0.81	0.90	0.85			
Tavsiye Edilen, Kaliteli Ürünler	0.86	0.70	0.77			
Ortalama Genel Sonuç	0.86	0.87	0.86	0.86	0.97	0.43

Modelin tahminlerinin doğruluğunu daha ayrıntılı bir şekilde incelemek için Şekil 4.23’deki konfüzyon matrisi kullanılmıştır.



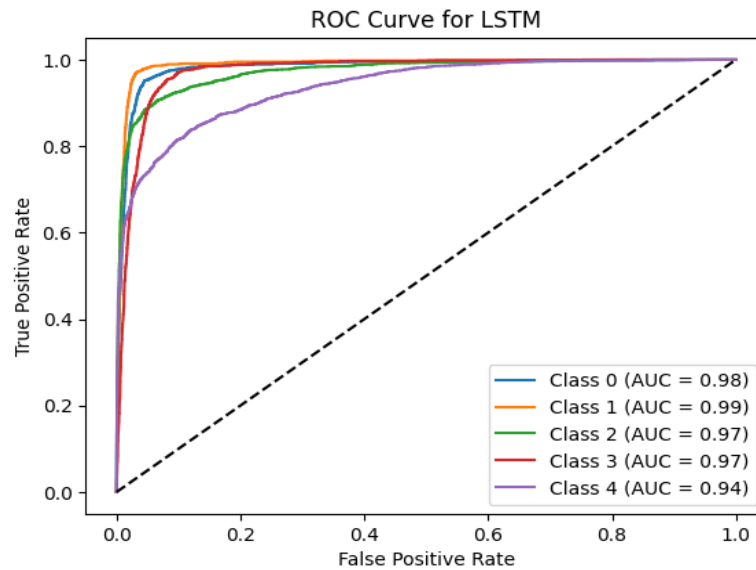
Şekil 4.23. LSTM modeli ile elde edilen konfüzyon matrisi

LSTM modelinin performans sonuçları ve konfüzyon matrisi incelendiğinde, modelin belirli sınıflarda yüksek performans gösterdiği, ancak diğer sınıflarda iyileştirme alanlarının bulunduğu gözlemlenmiştir. Özellikle "İade Edilen" ve "Fiyat Performans" sınıflarında duyarlılık ve F1 skorlarının oldukça yüksek olduğu, modelin bu sınıfları doğru şekilde ayırt etme kapasitesinin güçlü olduğunu göstermektedir. Benzer şekilde, "Kalitesiz, Kusurlu, Kötü Paketleme" sınıfında da dengeli bir kesinlik ve duyarlılık sergilenmiş, bu durum F1 skoruna olumlu yansımıştır.

Bununla birlikte, "Tavsiye Edilen, Kaliteli Ürünler" sınıfında duyarlılığın diğer sınıflara kıyasla daha düşük olması, modelin bu sınıfı tanımda zorlandığını ifade etmektedir. Bu durum, söz konusu sınıfın veri dağılımında dengesizlik veya diğer sınıflarla karışıklık yaşanmasından kaynaklanıyor olabilir. Konfüzyon matrisi incelendiğinde de bu sınıfta hatalı sınıflandırma oranlarının görece yüksek olduğu dikkat çekmektedir. Özellikle bu sınıfın, "Paketleme İyi, Hızlı Teslimat" sınıfı ile karışma eğilimi göstermesi, modelin bu iki sınıf arasındaki ayrımı yapmada güçlük yaşadığını göstermektedir.

Modelin metin verisinden anlamlı ilişkiler öğrenmekteki başarısı, embedding ve LSTM katmanlarının birlikte çalışmasıyla metin sırasındaki bağlamları etkili bir şekilde yakalamasına dayanmaktadır. Eğitim sürecindeki veri işleme adımları ve kullanılan teknikler, modelin performansını artırarak genel bir başarı sağlamıştır. Bu sonuçlar, LSTM'nin metin sınıflandırma problemlerinde etkili bir yöntem olduğunu bir kez daha ortaya koymaktadır.

Şekil 4.24 LSTM modeli sonucunda elde edilen ROC eğrilerini göstermektedir.



Şekil 4.24. LSTM modeli ile elde edilen ROC eğrileri

ROC eğrileri modelin her sınıf için oldukça güçlü bir ayırt edicilik yeteneğine sahip olduğunu göstermektedir. Özellikle ilk dört sınıf için ROC AUC değerlerinin 0.97'nin üzerinde olması, modelin bu sınıflarda pozitif ve negatif örnekleri yüksek bir doğrulukla ayırt edebildiğini ortaya koymaktadır. Ancak son sınıf için ROC AUC değerinin 0.94'e düşmesi, diğer sınıflara kıyasla bu sınıfta ayırt edicilik performansının bir miktar azaldığını işaret etmektedir. Bu durum, veri dengesizliği, sınıf örnekleri arasındaki benzerlikler veya modelin bu sınıfa özgü özellikleri öğrenmekte zorlanması gibi faktörlerden kaynaklanabilir. Genel olarak, tüm sınıflar için yüksek ROC AUC değerleri, modelin sınıflandırma performansının güvenilir ve başarılı olduğunu göstermektedir.

4.7.4. BiLSTM

Metin sınıflandırma problemini çözmek amacıyla BiLSTM tabanlı bir model geliştirilmiştir. Bu model, sıralı verilerin anlamını daha etkili bir şekilde yakalamak için her iki yönde bilgi taşıyan LSTM hücrelerinden yararlanır. Bu yapı, yalnızca ileri yöndeki bilgiyi dikkate alan modellerden farklı olarak, metinlerin bağlamını her iki yönden öğrenerek daha zengin özellikler çıkarılmasını sağlar. Modelin ana bileşeni, 128 birimli ve %20 dropout oranına sahip bir BiLSTM katmanıdır. Bu dropout oranı, modelin öğrenme sürecinde aşırı öğrenmeyi (overfitting) önlemeyi hedefler. BiLSTM katmanı, sıralı verinin hem ileri hem de geri yönlerini dikkate alarak verinin bağlamsal ilişkilerini etkili bir şekilde öğrenmektedir.

Çıkış katmanı, çok sınıflı sınıflandırma problemlerini çözmek için softmax aktivasyon fonksiyonuyla yapılandırılmıştır. Bu katman, her bir giriş örneğinin belirli bir sınıfa ait olma olasılığını tahmin etmektedir. Modelin eğitim sürecinde, toplam 5 epoch boyunca ve her iterasyonda 64 veri örneği işlenmiştir. Eğitim sırasında kayıp fonksiyonu olarak, çok sınıflı sınıflandırmalar için uygun olan 'sparse_categorical_crossentropy', optimizasyon işlemleri için ise 'Adam' algoritması tercih edilmiştir.

Eğitim süresince, modelin doğruluk metriği ve doğrulama verisi üzerindeki performansı düzenli olarak izlenmiştir. Modelin eğitimi tamamlandıktan sonra, test verisi üzerinde sınıflandırma sonuçları ve olasılıkları hesaplanmıştır. Şekil 4.25'de BiLSTM modeline ait pseudo kod görülmektedir.

```

1  Başlat:
2  | Özellik ve hedef değişkeni ayır:
3  |   Özellikler (X_raw): cleaned_comments
4  |   Hedef değişken (y): ticket
5  | Veriyi eğitim ve test setlerine ayır:
6  |   Eğitim seti : %80 Test seti: %20
7  Ön İşleme:
8  | Verileri aynı uzunlukta yapmak için doldurma işlemi uygula:
9  |   Maksimum uzunluk (maxlen): 200 kelime
10 BiLSTM modeli tanımla:
11 |   Gömülü katman (Embedding):
12 |     Giriş boyutu (input_dim): 10.000
13 |     Çıkış boyutu (output_dim): 128
14 |     Giriş uzunluğu (input_length): 200
15 |   Bidirectional LSTM katmanı:
16 |     Birim sayısı (units): 128
17 |     Dropout: 0.2
18 |     Tekrarlı dropout: 0.2
19 |   Çıkış katmanı (Dense):
20 |     Sınıf sayısı: len(np.unique(y_train))
21 |     Aktivasyon: softmax
22 |   Modeli derle:
23 |     Kayıp fonksiyonu: sparse_categorical_crossentropy
24 |     Optimizasyon: adam
25 Modeli eğit:
26 |   Epoch sayısı: 5
27 |   Batch boyutu: 64
28 Değerlendirme:
29 |   Test seti tahminleri:
30 |     Sınıf tahminleri: argmax(axis=1)
31 |     Olasılık tahminleri: predict_proba(X_test_pad)
32 |   Performans metrikleri:
33 |     Doğruluk (Accuracy), Kesinlik (Precision), Duyarlılık (Recall), F1 Skoru, ROC AUC, Log Loss
34 Sınıf Bazlı Analiz:
35 |   Precision, Recall ve F1 Skor.
36 Görselleştirme:
37 |   Sınıf başına ROC eğrisi ve AUC.
38 |   Eğrileri çiz ve rastgele tahmin çizgisi ekle.
39 |   Konfüzyon matrisini hesapla ve görselleştir

```

Şekil 4.25. BiLSTM Modeli Pseudo Kod

Çizelge 4.12’de BiLSTM modeline ait kullanılan hiperparametre bilgileri verilmiştir.

Çizelge 4.12. BiLSTM Hiperparametre Bilgileri

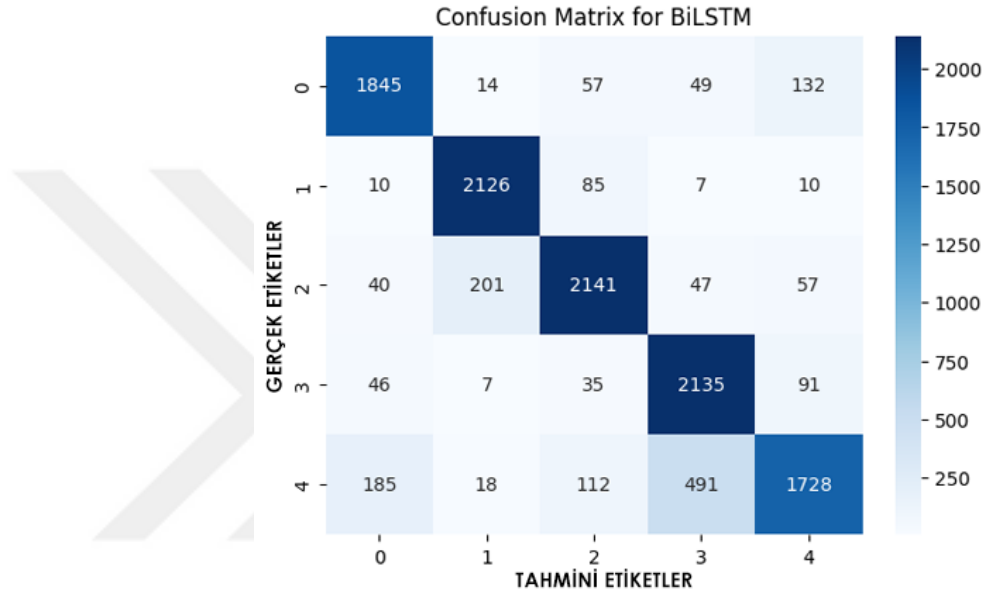
Parametre Adı	Değer
Giriş Boyutu	10.000
Çıkış Boyutu	128
Giriş Uzunluğu	200
BiLSTM Katmanı Birim Sayısı	128
Dropout	0.2
Tekrarlı Dropout	0.2
Dense Katman Aktivasyon	Softmax
Epoch Sayısı	5
Batch Boyutu	64

Çizelge 4.13’de BiLSTM modeli ile elde edilen test sonuçları gösterilmektedir.

Şekil 4.26 ise bu sınıflandırmaya ait konfüzyon matrisini içermektedir.

Çizelge 4.13. BiLSTM Test Sonuçları

Sınıf	Kesinlik	Duyarlılık	F1 Skoru	Doğruluk	ROC AUC	Log Loss
Fiyat Performans	0.87	0.88	0.87			
İade Edilen	0.90	0.95	0.92			
Kalitesiz, Kusurlu, Kötü Paketleme	0.88	0.86	0.87			
Paketleme İyi, Hızlı Teslimat	0.78	0.92	0.85			
Tavsiye Edilen, Kaliteli Ürünler	0.86	0.68	0.76			
Ortalama Genel Sonuç	0.86	0.86	0.85	0.85	0.97	0.44



Şekil 4.26. BiLSTM modeli ile elde edilen konfüzyon matrisi

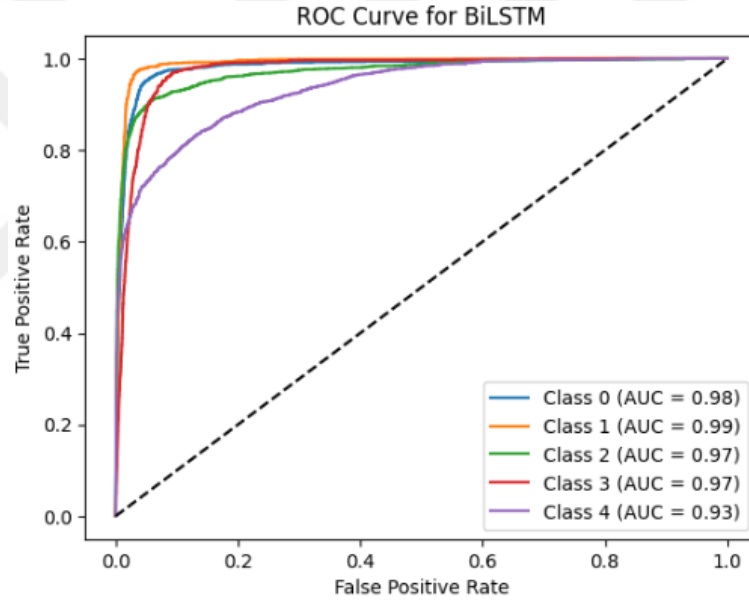
BiLSTM modelinin test sonuçları incelendiğinde, genel olarak modelin yüksek doğruluk sağladığı ve verimli bir performans sergilediği gözlemlenmektedir. "İade Edilen" sınıfı için en yüksek kesinlik ve duyarlılık değerleri elde edilmiştir, bu da modelin bu sınıfı doğru şekilde tanımladığını göstermektedir. "Fiyat Performans" ve "Kalitesiz, Kusurlu, Kötü Paketleme" gibi sınıflarda ise modelin başarı oranı yine yüksek olmakla birlikte, duyarlılık ve F1 skoru arasında küçük bir fark gözlemlenmektedir. Bu durum, modelin bu sınıflar için daha fazla yanlış pozitif sınıflandırma yapmış olabileceğini, ancak yine de genel başarıyı koruduğunu göstermektedir.

"Paketleme İyi, Hızlı Teslimat" sınıfı için, kesinlik ve F1 skoru daha düşük olmasına rağmen duyarlılık oldukça yüksek çıkmıştır. Bu da modelin bu sınıfı tespit etme konusunda güçlü olduğunu, ancak yanlış pozitif sınıflandırmalarda biraz daha dikkatli olması gerektiğini göstermektedir. "Tavsiye Edilen, Kaliteli Ürünler" sınıfı, daha düşük

kesinlik ve F1 skoru ile sonuçlanmıştır. Duyarlılık açısından ise orta seviyelerde kalmıştır, bu da modelin bu sınıfı doğru sınıflandırmada zorluk yaşadığını ve hatalı tahminlerde bulunduğunu gösterir.

Şekil 4.26'da verilen, konfüzyon matrisine bakıldığında, modelin doğru sınıflandırmalarda genellikle yüksek performans gösterdiği, ancak bazı sınıflar için yanlış sınıflandırmaların olduğu anlaşılmaktadır. Özellikle "Tavsiye Edilen, Kaliteli Ürünler" ve "Paketleme İyi, Hızlı Teslimat" sınıflarında hatalı sınıflandırmaların daha fazla olduğu görülmektedir. Genel olarak, modelin doğruluğu yüksek olmasına rağmen, daha zor sınıflar için iyileştirme yapılabilir.

Şekil 4.27'de BiLSTM modeli sonucunda elde edilen ROC eğrilerini göstermektedir.



Şekil 4.27. BiLSTM modeli ile elde edilen ROC eğrileri

BiLSTM modelinin sınıf bazlı ROC eğrisi sonuçları, modelin sınıflar arasında ayırt edici performansının genelde yüksek olduğunu göstermektedir. İlk dört sınıfta ROC AUC değerleri oldukça yüksek, yani modelin pozitif ve negatif örnekleri ayırt etme kapasitesi güçlüdür. Ancak, son sınıfta 0.93'lük bir değer, diğer sınıflara kıyasla biraz daha düşük olsa da hâlâ iyi bir performansa işaret etmektedir. Genel olarak, ROC eğrisi sonuçları modelin sınıflar arasında tutarlı ve başarılı bir performans sergilediğini desteklemektedir.

4.7.5. GRU

Metin sınıflandırma görevini gerçekleştirmek amacıyla GRU tabanlı bir model geliştirilmiştir. GRU, uzun dizilerle çalışırken hesaplama verimliliği sağlayan ve bellek sorunlarını azaltan bir RNN türüdür. Bu model, ardışık verilerdeki bağımlılıkları öğrenmek için 128 birimli bir GRU katmanı ile yapılandırılmıştır. GRU katmanında kullanılan %20 dropout ve %20 recurrent dropout oranları, aşırı öğrenmeyi (overfitting) önlemek ve modelin genelleme kapasitesini artırmak için uygulanmıştır.

Modelin çıkış katmanı, çok sınıflı sınıflandırma problemlerini çözmek için softmax aktivasyon fonksiyonuyla yapılandırılmıştır. Bu katman, her bir sınıf için olasılık tahminleri üretir ve giriş verisinin hangi sınıfa ait olduğunu belirler.

Modelin eğitiminde toplam 5 epoch boyunca çalışılmış ve her bir iterasyonda 64 veri örneği işlenmiştir. Eğitim sırasında kayıp fonksiyonu olarak 'sparse_categorical_crossentropy', optimizasyon algoritması olarak ise 'Adam' tercih edilmiştir. Adam algoritması, adaptif öğrenme oranı ile parametre güncellemelerini hızlı ve verimli bir şekilde gerçekleştirerek modelin performansını artırmıştır. Eğitim tamamlandıktan sonra model, test verisi üzerinde değerlendirilmiş ve tahmin edilen sınıflar ile olasılıklar elde edilmiştir. Şekil 4.28'de GRU modeline ait pseudo kod verilmiştir.

```

1  Başlat:
2  Veriyi özellik ve hedef değişken olarak ayır:
3  Özellik (X_raw): cleaned_comments sütunu
4  Hedef değişken (y): ticket sütunu
5  Veriyi eğitim ve test setlerine ayır:
6  Eğitim seti : %80 Test seti: %20
7  Ön işleme:
8  Verileri aynı uzunlukta yapmak için doldurma işlemi uygula:
9  Maksimum uzunluk (maxlen): 200 kelime
10 GRU modeli tanımla:
11 Gömülü katman (Embedding):
12 Giriş boyutu (input_dim): 10.000
13 Çıkış boyutu (output_dim): 128
14 Giriş uzunluğu (input_length): 200
15 GRU katmanı:
16 Birim sayısı (units): 128
17 Dropout: 0.2
18 Tekrarlı dropout (recurrent_dropout): 0.2
19 Çıkış katmanı (Dense):
20 Sınıf sayısı: len(np.unique(y_train))
21 Aktivasyon: softmax
22 Modeli derle:
23 Kayıp fonksiyonu: sparse_categorical_crossentropy
24 Optimizasyon algoritması: adam
25 Modeli eğitim seti ile eğit:
26 Epoch sayısı: 5
27 Batch boyutu: 64
28 Değerlendirme:
29 Test verisi üzerinde tahminler yap:
30 Sınıf tahminleri: argmax(axis=1)
31 Olasılık tahminleri: predict_proba(X_test_pad)
32 Model performans metriklerini hesapla:
33 Doğruluk (accuracy), Kesinlik (precision), Duyarlılık (recall), F1 skoru, ROC AUC, Log Loss
34 Sınıf bazlı metrikler:
35 Precision, Recall ve F1 Score
36 Görselleştirme:
37 Her bir sınıf için ROC eğrisi ve AUC değerlerini hesapla
38 ROC eğrilerini görselleştir
39 Konfüzyon matrisini hesapla ve görselleştir

```

Şekil 4.28. GRU Modeli Pseudo Kod

Çizelge 4.14’de GRU modeline ait kullanılan hiperparametre bilgileri verilmiştir.

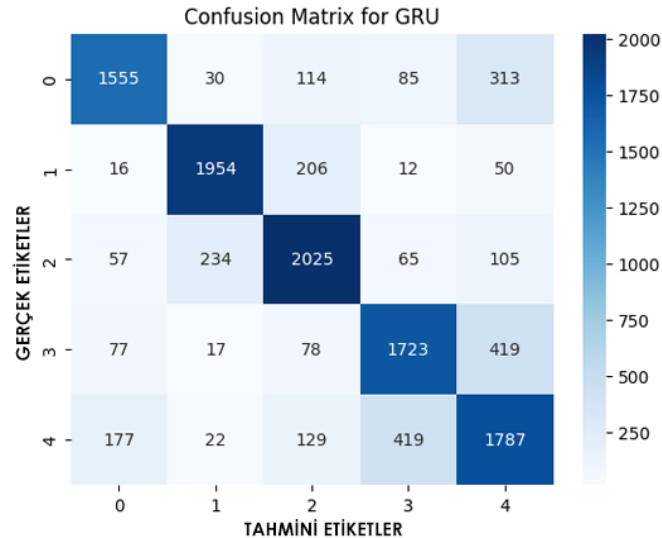
Çizelge 4.14. GRU Hiperparametre Bilgileri

Parametre Adı	Değer
Giriş Boyutu	10.000
Çıkış Boyutu	128
Giriş Uzunluğu	200
GRU Katmanı Birim Sayısı	128
Dropout	0.2
Tekrarlı Dropout	0.2
Dense Katman Aktivasyon	Softmax
Epoch Sayısı	5
Batch Boyutu	64

Çizelge 4.15’de GRU modeli ile elde edilen test sonuçları gösterilmektedir. Modelin sınıflandırma performansı Şekil 4.29’daki Konfüzyon Matrisi ile daha ayrıntılı bir şekilde analiz edilmiştir.

Çizelge 4.15. GRU Test Sonuçları

Sınıf	Kesinlik	Duyarlılık	F1 Skoru	Doğruluk	ROC AUC	Log Loss
Fiyat Performans	0.83	0.74	0.78			
İade Edilen	0.87	0.87	0.87			
Kalitesiz, Kusurlu, Kötü Paketleme	0.79	0.81	0.80			
Paketleme İyi, Hızlı Teslimat	0.75	0.74	0.75			
Tavsiye Edilen, Kaliteli Ürünler	0.67	0.71	0.69			
Ortalama Genel Sonuç	0.78	0.78	0.78	0.78	0.95	0.62



Şekil 4.29. GRU modeli ile elde edilen konfüzyon matrisi

Çizelge 4.15’te verilen GRU modelinin performansına göre, “İade Edilen” sınıfında yüksek doğruluk, kesinlik ve duyarlılık değerleri elde edilmiştir. Bu da modelin bu sınıfı doğru sınıflandırmada başarılı olduğunu gösterir. Diğer sınıflar ise daha düşük performans sergilemiş, özellikle “Fiyat Performans ve Tavsiye Edilen”, “Kaliteli Ürünler” sınıflarında modelin doğru sınıflandırma oranı düşmüştür.

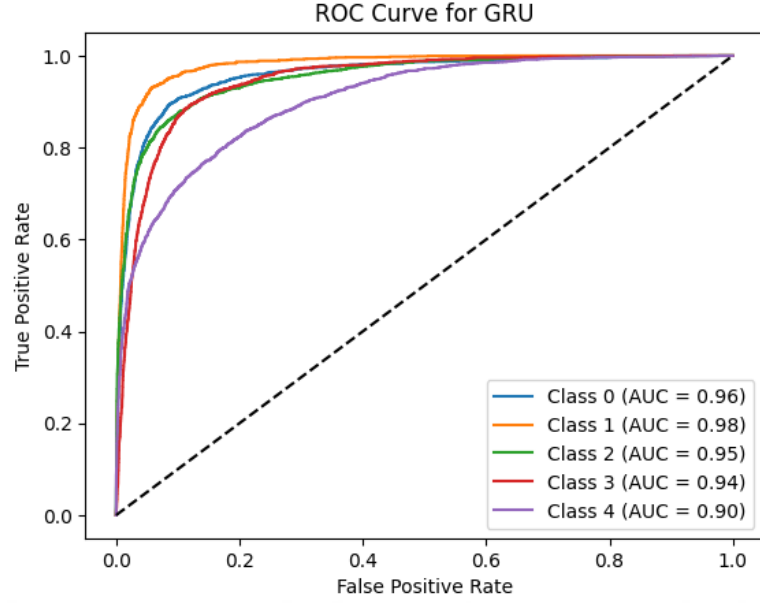
“Kalitesiz, Kusurlu, Kötü Paketleme” ve “Paketleme İyi, Hızlı Teslimat” sınıflarında ise modelin performansı ortalamanın üzerinde olsa da, her iki sınıfta da bazı hatalar mevcuttur. Özellikle “Paketleme İyi, Hızlı Teslimat” sınıfı, modelin doğru tahminler yapma konusunda zorluk yaşadığı bir kategori olarak dikkat çekiyor.

Şekil 4.29’da görülen konfüzyon matrisi de bu durumu desteklemekte olup, modelin “Tavsiye Edilen, Kaliteli Ürünler” sınıfında en fazla hata yaptığı, bu sınıfın örneklerinin genellikle “Fiyat Performans” ve “Kalitesiz, Kusurlu, Kötü Paketleme” sınıflarına yanlış sınıflandırıldığı gözlemlenmiştir. Bu, modelin bazı sınıflar arasında karışıklık yaşadığını ve iyileştirme yapılması gerektiğini göstermektedir.

GRU tabanlı modelin test sonuçlarına bakıldığında, genel olarak metin sınıflandırma görevinde oldukça sağlam bir performans sergilendiği görülmektedir. Modelin doğruluk, kesinlik, duyarlılık ve F1 skoru gibi temel metriklerde %78 civarında bir başarı elde ettiği, bu değerlerin birbirine yakın olması, modelin dengeli bir performans sergilediğini ve doğru sınıflandırmalar yaptığına işaret etmektedir. ROC AUC değeri 0.95 ile oldukça yüksek bir performans gösterdiği için modelin, özellikle sınıf ayrımında başarılı olduğunu söylemek mümkündür. Bu, modelin tüm sınıflar arasındaki ayrımı güçlü bir şekilde yapabildiğini ve yanlış sınıflandırma olasılığının düşük olduğunu gösterir.

Ancak, Log Loss değeri 0.62 ile biraz daha yüksek kalmıştır, bu da modelin sınıf tahminlerinin güven seviyesinin belirli bir noktada daha düşük olduğunu gösterebilir. Modelin doğruluğu yüksek olsa da, bazı sınıflarda daha fazla belirsizlik bulunabilir. GRU katmanının, LSTM ve BiLSTM gibi diğer sıralı veri modellerine kıyasla daha az hesaplama gerektirmesi, bu modelin daha hızlı eğitim ve test süresi sağlamasını mümkün kılmaktadır.

Şekil 4.30’da GRU modeli sonucunda elde edilen ROC eğrileri gösterilmektedir.



Şekil 4.30. GRU modeli ile elde edilen ROC eğrileri

Şekil 4.30'daki ROC eğrileri, modelin çoğu sınıf için yüksek doğrulukla ayırım yapabildiğini göstermektedir. 0.96, 0.98, 0.95, 0.94 ve 0.90 arasındaki değerler, modelin genellikle iyi performans gösterdiğini ancak son sınıfta biraz daha zorluk yaşadığını işaret ediyor. Yine de genel olarak, modelin doğru sınıflandırmalar yapma yeteneği oldukça güçlüdür.

4.7.6. BERT

BERT modeli, metin sınıflandırma ve NLP görevlerinde geniş bir uygulama yelpazesi sunan güçlü bir derin öğrenme modelidir. Gerçekleştirilen bu tez çalışmasında, Türkçe metin verisi üzerinde eğitilmiş bir BERT modeli kullanılarak çok sınıflı bir sınıflandırma problemi çözülmüştür. Bu süreç, verilerin hazırlanmasından modelin eğitilmesine, performans değerlendirmelerinin yapılmasına kadar birçok adımı içermektedir. Aşağıda, BERT modelinin eğitim süreci ve kullanılan parametrelerin etkileri açıklanacaktır.

Veri kümesi, metin yorumlarını içeren bir DataFrame'den alınmıştır. Bu metinler, `cleaned_comments` sütununda yer almakta ve her bir metin için bir `ticket` etiketi bulunmaktadır. İlk olarak, etiketler sayısal değerlere dönüştürülmüştür; bu işlem, `LabelEncoder` sınıfı kullanılarak yapılmıştır. Eğitim ve test veri kümeleri, `train_test_split` fonksiyonu ile rastgele bir şekilde ayrılmıştır. Verilerin %80'i eğitim, %20'si ise test olarak modelin doğruluğunu değerlendirmek için ayrılmıştır.

BERT modelinin eğitilmesinde kullanılan temel yapı, `dbmdz/bert-base-turkish-128k-uncased` adlı Türkçe dil modelidir. Bu model, Türkçe metinler üzerinde eğitilmiş olup, metinlerdeki kelimeleri anlamlı vektörlere dönüştürmek için kullanılan bir `BertTokenizer` aracılığıyla verileri işler.

- **Tokenization:** Her bir metin, BERT tokenizörü kullanılarak tokenlara ayrılmıştır. Bu işlemde, her metnin maksimum uzunluğu 256 token ile sınırlıdır. Tokenizer, metni `input_ids` ve `attention_masks` gibi iki parametreye dönüştürür. `input_ids`, metni BERT modelinin anlayabileceği sayısal verilere dönüştürürken, `attention_mask`, modelin hangi kelimelere odaklanması gerektiğini belirtir.
- **Model Yapılandırması:** Model, `BertForSequenceClassification` sınıfı ile oluşturulmuştur. Bu modelin son katmanı, sınıflandırma amacıyla kullanılan softmax aktivasyon fonksiyonu ile yapılandırılmıştır. Modelin çıktısı, her bir sınıf için olasılık değerlerini verir.

BERT modelinin eğitimi, `AdamW` optimizasyon algoritması ve `get_linear_schedule_with_warmup` scheduler kullanılarak gerçekleştirilmiştir. Modelin öğrenme oranı (learning rate) $1e-5$ olarak belirlenmiş ve `eps` parametresi $1e-8$ olarak ayarlanmıştır. Bu optimizasyon parametreleri, modelin verileri daha etkili bir şekilde öğrenmesini sağlamak için dikkatle seçilmiştir.

- **Epochs:** Model 10 epoch boyunca eğitilmiştir. Her epoch'ta eğitim verileri, küçük mini-batch'ler halinde modelin üzerinden geçirilmiştir. Her adımda, modelin tahmin hatası (loss) hesaplanmış ve geriye yayılım (backpropagation) ile modelin parametreleri güncellenmiştir. Eğitim süreci boyunca kayıp fonksiyonu, modelin başarısını izlemek için kullanılmıştır.
- **Ağırlık Güncelleme:** Modelin parametrelerini güncellemek için `optimizer.step()` fonksiyonu kullanılmış ve öğrenme oranı planlaması (`scheduler.step()`) ile aşamalı bir iyileştirme sağlanmıştır.

Şekil 4.31'de BERT modeline ait pseudo kod verilmiştir.

```

1  Başlat:
2      Veri kümesi yükle: df
3      Veriyi kopyala: data = df.copy()
4      Genel bilgi al: data.info()
5      GPU kontrolü yap: device_name = tf.test.gpu_device_name()
6  Veri Hazırlığı:
7      Etiket sütunu sayısal değerlere dönüştür: data['ticket'] = LabelEncoder().fit_transform(data['ticket'])
8      Veriyi eğitim ve test setlerine ayır:
9          Eğitim seti : %80 Test seti: %20
10     Eğitim ve test metinlerini hazırla:
11         training_texts = training.cleaned_comments.values
12         test_texts = test.cleaned_comments.values
13  Tokenizasyon:
14     BERT tokenizer yükle: tokenizer = BertTokenizer.from_pretrained('dbmdz/bert-base-turkish-128k-uncased')
15     Eğitim verisini tokenize et:
16         input_ids ve attention_masks oluştur
17         encoded_dict = tokenizer.encode_plus()
18         input_ids.append(encoded_dict['input_ids'])
19         attention_masks.append(encoded_dict['attention_mask'])
20  Model:
21     Modeli yükle: BertForSequenceClassification.from_pretrained
22     Num labels: number_of_categories
23     Çıktılar: output_attentions=False, output_hidden_states=False
24     Modeli GPU'ya yükle: model.cuda()
25  Optimizasyon:
26     Eğitim için optimizer ve scheduler oluştur:
27         optimizer = AdamW(model.parameters())
28         total_steps hesapla
29         scheduler = get_linear_schedule_with_warmup()
30  Eğitim döngüsüne başla:
31     Epoch sayısı: epochs = 10
32     training_stats listeye ekle
33  Doğrulama:
34     Test verisi üzerinde tahmin yap: predictions, true_labels listeleri oluştur
35     Her batch için logits ve etiketleri kaydet
36  Performans metrikleri:
37     Doğruluk (Accuracy), Kesinlik (Precision), Duyarlılık (Recall), F1 Skoru, ROC AUC, Log Loss
38  Sınıf Bazlı Analiz:
39     Precision, Recall ve F1 Skor.
40  Görselleştirme:
41     Sınıf başına ROC eğrisi ve AUC.
42     Eğrileri çiz ve rastgele tahmin çizgisi ekle.
43     Konfüzyon matrisini hesapla ve görselleştir.

```

Şekil 4.31. BERT Modeli Pseudo Kod

Çizelge 4.16’da BERT modeline ait kullanılan hiperparametre bilgileri verilmiştir.

Çizelge 4.16. BERT Hiperparametre Bilgileri

Parametre Adı	Değer
Giriş Uzunluğu	256
Modelin öğrenme oranı	1e-5
Tolerans Değeri (Eps)	1e-8
Dropout	0.2
Epoch Sayısı	10
Batch Boyutu	32

Eğitim tamamlandıktan sonra, test verileri ile modelin başarımlı değerlendirilmiştir. Test verileri de aynı şekilde tokenize edilmiş ve modelin çıktıları elde edilmiştir. Model, her bir test örneği için sınıflandırma tahminleri yapmış ve bu tahminler,

performans metrikleri ile değerlendirilmiştir. Çizelge 4.17’de BERT modeli ile elde edilen test sonuçları gösterilmektedir.

Çizelge 4.17. BERT Test Sonuçları

Sınıf	Kesinlik	Duyarlılık	F1 Skoru	Doğruluk	ROC AUC	Log Loss
Fiyat Performans	0.88	0.94	0.91			
İade Edilen	0.92	0.98	0.95			
Kalitesiz, Kusurlu, Kötü Paketleme	0.94	0.87	0.91			
Paketleme İyi, Hızlı Teslimat	0.82	0.92	0.87			
Tavsiye Edilen, Kaliteli Ürünler	0.87	0.73	0.79			
Ortalama Genel Sonuç	0.89	0.89	0.89	0.89	0.98	0.41

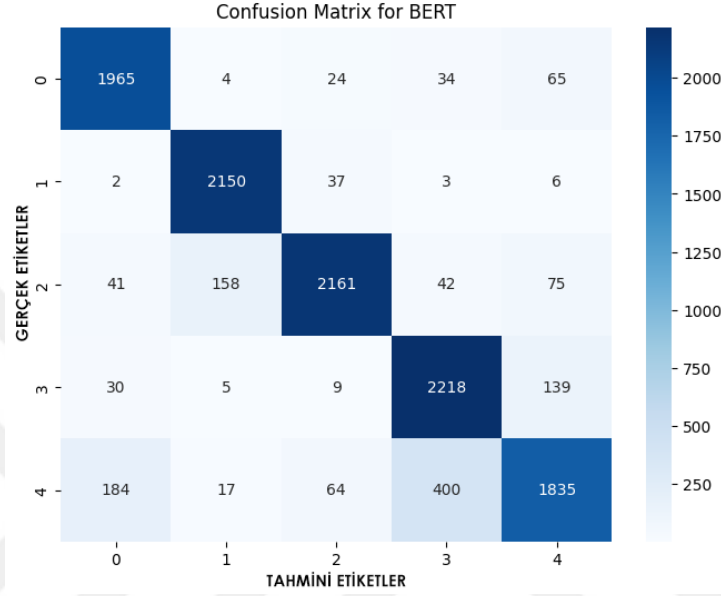
Çizelge 4.17’ye göre söylenebilir ki, BERT modelinin Türkçe metin sınıflandırma görevindeki performansı oldukça başarılıdır. Doğruluk, Kesinlik, Duyarlılık ve F1 Skoru gibi metriklerin hepsi 0.89 gibi yüksek değerlere ulaşmıştır, bu da modelin genel sınıflandırma yeteneğinin güçlü olduğunu göstermektedir. Bu sonuçlar, modelin metinleri doğru bir şekilde sınıflandırmada etkin olduğunu ve çeşitli sınıflar arasında iyi bir denge sağladığını belirtmektedir.

0.98 olan ROC AUC skoru, modelin pozitif ve negatif sınıflar arasında mükemmel bir ayırım gücüne sahip olduğunu ortaya koymaktadır. Bu, modelin doğru sınıf tahminleri yapma konusunda ne kadar başarılı olduğunu gösteren önemli bir metriktir. Log Loss değeri ise 0.41 ile kabul edilebilir bir seviyede olup, modelin kayıp fonksiyonunun düşük olduğunu ve tahminlerinin genellikle doğru olduğunu belirtmektedir.

Modelin sınıflandırma performansı Şekil 4.32’deki *Konfüzyon Matrisi* ile daha ayrıntılı bir şekilde analiz edilmiştir.

BERT modeli, metin sınıflandırma görevinde genel olarak başarılı sonuçlar elde etmiştir. "Fiyat Performans" ve "İade Edilen" sınıfları yüksek doğruluk ve F1 skoru sergilerken, "Kalitesiz, Kusurlu, Kötü Paketleme" sınıfı da iyi performans göstermiştir. "Paketleme İyi, Hızlı Teslimat" sınıfında doğruluk diğer sınıflara göre biraz daha düşük olmakla birlikte makul bir performans sergilemiştir. Ancak "Tavsiye Edilen, Kaliteli Ürünler" sınıfında daha fazla yanlış sınıflandırma görülmüştür ve yaklaşık 400 örneğin yanlış sınıflandırıldığı gözlemlenmiştir. Bu durum, modelin bu sınıfta diğerlerine kıyasla daha düşük bir genelleme performansı sergilediğini göstermektedir.

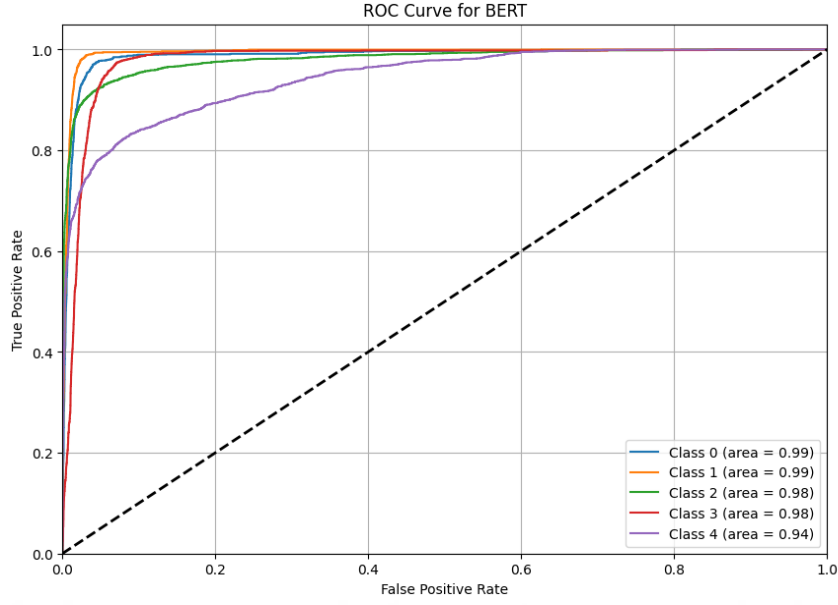
Konfüzyon matrisine göre model, çoğu sınıfı doğru bir şekilde sınıflandırmıştır ve genel olarak yüksek doğruluk ve F1 skoru elde etmiştir. Bununla birlikte, "Tavsiye Edilen, Kaliteli Ürünler" ve "Paketleme İyi, Hızlı Teslimat" sınıflarında iyileştirmeye ihtiyaç duyulmaktadır. Gelecekte, veri setinin boyutunun artırılması veya daha dengeli bir dağılım sağlanması gibi yöntemlerle modelin doğruluğu daha da artırılabilir.



Şekil 4.32. BERT modeli ile elde edilen konfüzyon matrisi

Sonuçlar, BERT modelinin Türkçe metin verisi üzerinde etkili bir şekilde çalıştığını ve çok sınıflı sınıflandırma problemlerinde yüksek doğruluk ve güvenilirlik sağladığını ortaya koymaktadır. Modelin eğitimi, optimizasyon stratejileri ve yapılandırılması, sonuçların kalitesini doğrudan etkileyen önemli faktörlerdir.

Modelin performansının görselleştirilmesi için Şekil 4.33'deki ROC eğrisinde her bir sınıf için AUC değeri hesaplanarak sınıfların ne kadar doğru ayırt edilebildiği görselleştirilmiştir. BERT modelinin ROC AUC değerleri oldukça etkileyici bir performansı yansıtmaktadır. İlk dört ROC AUC değeri (0.99, 0.99, 0.98, 0.98) modelin her bir sınıf için mükemmel bir ayırım gücüne sahip olduğunu gösteriyor. Bu, modelin her sınıf arasında doğru sınıflandırma yapma konusunda son derece başarılı olduğunu ve pozitif ile negatif sınıflar arasındaki farkları çok iyi ayırt ettiğini ifade eder. Son sınıf için ise ROC AUC değeri 0.94, hala oldukça yüksek bir değerdir ve modelin bu sınıf için de sağlam bir performans sergilediğini gösterir. Bu çalışma, Türkçe metinlerin sınıflandırılmasında BERT modelinin etkinliğini gösteren bir örnek olup, doğal dil işleme (NLP) alanındaki araştırmalar için değerli bir referans kaynağı sağlayabilir.



Şekil 4.33. BERT modeli ile elde edilen ROC eğrileri

4.7.7. Lojistik Regresyon

Geleneksel makine öğrenmesi yöntemlerinden olan Lojistik Regresyonun eğitim aşamasında, modelin performansını değerlendirebilmek ve aşamalı iyileştirme sağlayabilmek adına mini-batch yöntemi kullanılmıştır. Bu süreçte, model her iterasyonda bir epoch boyunca `lbfgs` optimizasyon algoritması ile eğitilmiş, toplamda 20 iterasyona kadar ilerlenmiştir. Modelin eğitiminde sıcak başlangıç yöntemi (`warm_start=True`) kullanılarak önceki iterasyonlarda öğrenilen ağırlıkların korunması sağlanmış ve bu sayede eğitim süreci daha verimli hale getirilmiştir. Eğitim tamamlandıktan sonra model, test verisi üzerinde tahminler yapmış ve her sınıf için olasılık tahminleri hesaplanmıştır.

Şekil 4.34’de Lojistik Regresyon modeline ait pseudo kod verilmiştir.

Çizelge 4.18’de ise Lojistik Regresyon modeline ait kullanılan hiperparametre bilgileri görülmektedir.

```

1  Başlat:
2      Veriyi eğitim ve test setlerine ayır:
3      |   Eğitim seti : %80 Test seti: %20
4  Model:
5      Logistic Regression modeli tanımla:
6      |   Maksimum iterasyon: 20
7      |   Sıcak başlangıç (warm_start): Aktif
8      |   Çözümleyici (solver): 'lbfgs'
9  Eğitim:
10     Mini-batch eğitim döngüsü:
11     |   20 iterasyon boyunca:
12     |   |   Maksimum iterasyon sayısını güncelle
13     |   |   Modeli eğitim verisi ile eğit
14  Değerlendirme:
15     Test seti üzerinde tahmin yap:
16     |   Sınıf tahminlerini al
17     |   Sınıf olasılıklarını al
18     Performans metriklerini hesapla:
19     |   Doğruluk (Accuracy)
20     |   Kesinlik (Precision)
21     |   Duyarlılık (Recall)
22     |   F1 skoru
23     |   ROC AUC skoru
24     |   Log loss
25  Görselleştirme:
26     ROC AUC eğrilerini çiz:
27     |   Her sınıf için ROC eğrisi ve AUC değerlerini hesapla
28     Konfüzyon matrisini görselleştir
29  Sonuçlar:
30     Performans metriklerini tablo halinde görüntüle
31     Sınıf bazlı precision, recall ve F1 skorlarını yazdır

```

Şekil 4.34. Lojistik Regresyon Modeli Pseudo Kod

Çizelge 4.18. Lojistik Regresyon Hiperparametre Bilgileri

Parametre Adı	Değer
Maksimum İterasyon	20
Sıcak Başlangıç	Aktif
Çözümleyici (solver)	Lbfgs
Hedef Değişken Sınıf Adeti	5
Sınıf Yapısı Algılama (multiclass)	Auto
Rastgele Durum (random state)	42

Çizelge 4.19’da Lojistik Regresyon modeli ile elde edilen test sonuçları gösterilmiştir. Bu sonuçlara göre söylenebilir ki, modelin performansı çok sayıda metrik aracılığıyla kapsamlı bir şekilde değerlendirilmiş ve sonuçlar genel anlamda oldukça başarılı bulunmuştur. Doğruluk skoru 0.85 olarak belirlenmiş ve bu değer, modelin genel sınıflandırma doğruluğunun yüksek olduğunu göstermiştir. Bu, veri setindeki örneklerin büyük bir kısmının doğru bir şekilde tahmin edildiğine işaret etmektedir. Kesinlik (0.85) ve duyarlılık (0.85) değerleri dengeli ve yüksek seviyede kalmıştır; bu da modelin doğru

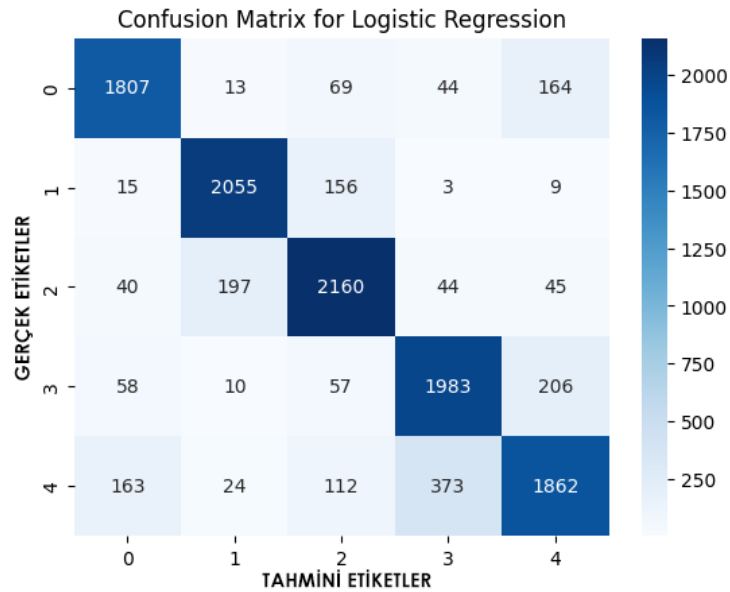
pozitif tahminlerde bulunma ve gerçek pozitif örnekleri tanıma yeteneğini tutarlı bir şekilde gösterdiğini ortaya koymaktadır.

F1 skoru ise yine 0.85 seviyesindedir ve kesinlik ile duyarlılık arasındaki dengeyi sağlayarak modelin sınıflandırma hatalarını minimize ettiğini kanıtlamaktadır. Bunun yanı sıra, ROC AUC skoru 0.97 gibi oldukça yüksek bir değer elde etmiştir. Bu skor, modelin sınıflar arasında ayırım yapma kapasitesinin güçlü olduğunu ve doğru sınıflandırma yapma ihtimalinin oldukça yüksek olduğunu ortaya koymaktadır. Log Loss değerinin 0.47 gibi makul bir seviyede kalması, modelin tahminlerinin genel olarak güvenilir olduğunu göstermektedir; ancak bu değer daha da düşürülmesi, tahminlerin doğruluğunu artırma potansiyeline sahiptir.

Çizelge 4.19. Lojistik Regresyon Test Sonuçları

Sınıf	Kesinlik	Duyarlılık	F1 Skoru	Doğruluk	ROC AUC	Log Loss
Fiyat Performans	0.87	0.86	0.86			
İade Edilen	0.89	0.92	0.91			
Kalitesiz, Kusurlu, Kötü Paketleme	0.85	0.87	0.86			
Paketleme İyi, Hızlı Teslimat	0.81	0.86	0.83			
Tavsiye Edilen, Kaliteli Ürünler	0.81	0.73	0.77			
Ortalama Genel Sonuç	0.85	0.85	0.85	0.85	0.97	0.47

Şekil 4.35 bu model ile elde edilen konfüzyon matrisini göstermektedir.



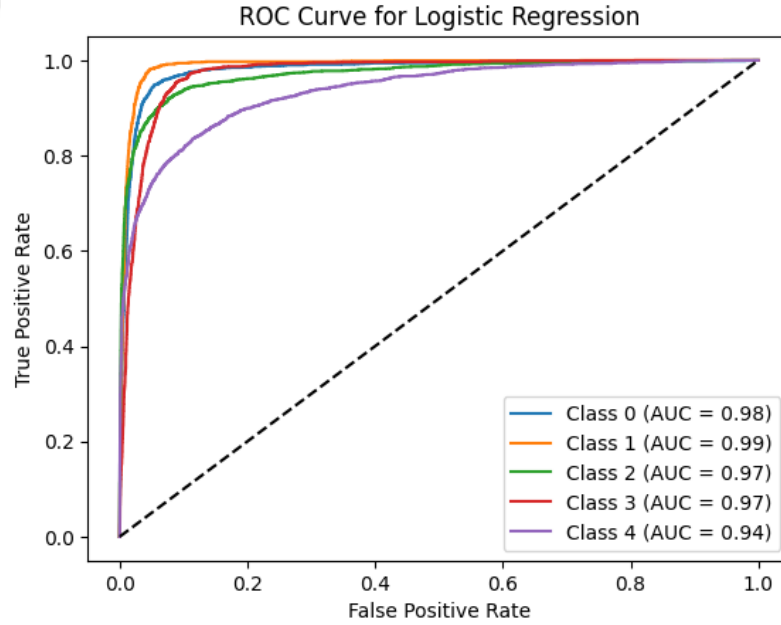
Şekil 4.35. Lojistik Regresyon modeli ile elde edilen konfüzyon matrisi

Çizelge 4.19 ve Şekil 4.35 göstermektedir ki, en başarılı performans “İade Edilen” sınıfında görülmüş olup, modelin kesinlik ve duyarlılık değerleri sırasıyla %89 ve %92 olarak belirlenmiştir. “Fiyat Performans” ve “Kalitesiz, Kusurlu, Kötü Paketleme” sınıfları da oldukça yüksek başarıya sahiptir (Kesinlik: 0.87, Duyarlılık: 0.86 ve Kesinlik: 0.85, Duyarlılık: 0.87).

“Paketleme İyi, Hızlı Teslimat” sınıfında biraz daha düşük performans gözlemlenmiş (Kesinlik: 0.81, Duyarlılık: 0.86), ancak model yine de doğru pozitif tahminlerde başarılı olmuştur. “Tavsiye Edilen, Kaliteli Ürünler” sınıfında ise kesinlik ve duyarlılık değerleri daha düşük kalmıştır (Kesinlik: 0.81, Duyarlılık: 0.73), bu da modelin bu sınıfı doğru tahmin etmekte zorlandığını göstermektedir.

Konfüzyon matrisine dayanarak yapılan değerlendirmeye göre söylenebilir ki, “İade Edilen” ve “Kalitesiz, Kusurlu, Kötü Paketleme” yüksek doğruluk gösterirken, model genel olarak iyi performans sergilemekle birlikte, bazı sınıflar arasında karışıklıkların giderilmesi ve özellikle “Tavsiye Edilen, Kaliteli Ürünler” sınıfı için iyileştirme yapılması gerekmektedir.

Şekil 4.36 Lojistik Regresyon modeli sonucunda elde edilen Roc Eğrilerini göstermektedir.



Şekil 4.36. Lojistik Regresyon modeli ile elde edilen ROC eğrileri

Modelin AUC değerleri, tüm sınıflar için oldukça yüksek olup 0.94 ile 0.99 arasında değişmektedir. Bu da modelin sınıflar arasında güçlü bir ayırım yapma yeteneğine sahip olduğunu ve pozitif sınıfları doğru şekilde ayırt etmekte başarılı olduğunu göstermektedir. Ancak, Sınıf 4 için AUC değeri diğerlerine göre biraz daha

düşük (0.94) kalmıştır. Buradan bu sınıf için diğerlerine kıyasla daha fazla karışıklık yaşandığı anlaşılabilmektedir. Diğer taraftan, karar eşikleri ile ilgili olarak, eğriler modelin farklı eşik değerlerinde nasıl performans gösterdiğini ortaya koymaktadır. Eğrilerin üst sol köşeye yakın olması, modelin düşük hata oranıyla çalıştığını ve yüksek doğruluk sağladığını ifade eder.

4.7.8. Rastgele Orman

Rastgele Orman modeli, sınıflandırma problemini çözmek amacıyla topluluk öğrenme (ensemble learning) yaklaşımıyla oluşturulan bir geleneksel makine öğrenmesi modeli olup, çalışmada eğitim verisi üzerinde eğitilmiştir. Bu model, birden fazla karar ağacından oluşarak karmaşık veri örüntülerini öğrenme ve sınıflar arasındaki ayrımı güçlendirme yeteneği sunmaktadır.

Modelin eğitim sürecinde, performansı optimize etmek amacıyla aşamalı bir eğitim yaklaşımı benimsenmiş ve modelin hiperparametrelerinden biri olan ağaç sayısı ($n_estimators$) kademeli olarak artırılmıştır. Eğitim sırasında, toplamda 50 ağaç kullanılmış ve her adımda ağaç sayısı 10 birim artırılmıştır. Eğitim tamamlandıktan sonra, model test verisi üzerinde tahminlerde bulunmuş, bu analizler modelin sınıflar arasındaki ayırım yeteneğini ve genel başarımını detaylı bir şekilde ortaya koymuştur.

Şekil 4.37’de Rastgele Orman modeline ait pseudo kod verilmiştir.

Çizelge 4.20’de Rastgele Orman modeline ait kullanılan hiperparametre bilgileri verilmiştir.

Çizelge 4.21’de Rastgele Orman modeli ile elde edilen test sonuçları gösterilmektedir. Şekil 4.38 ise Rastgele Orman modeli sonucunda elde edilen konfüzyon matrisini göstermektedir.

```

1  Başlat:
2  |   Veriyi eğitim ve test setlerine ayır:
3  |   Eğitim seti : %80 Test seti: %20
4  Model:
5  |   Random Forest modeli tanımla:
6  |   Ağaç sayısı (n_estimators): 10
7  |   Sıcak başlangıç (warm_start): Aktif
8  |   Rastgele durum (random_state): 42
9  Eğitim:
10 |   Aşamalı eğitim döngüsü:
11 |   |   Ağaç sayısını (n_estimators) 10'dan 50'ye kadar 10 adımlarla artır:
12 |   |   |   Modeli eğitim verisiyle eğit
13 |   |   |   Eğitim log loss değerlerini hesapla ve kaydet
14 |   |   |   Her iterasyon için log loss değerlerini yazdır
15 Değerlendirme:
16 |   Test seti üzerinde tahmin yap:
17 |   |   Sınıf tahminlerini al
18 |   |   Sınıf olasılıklarını al
19 |   Performans metriklerini hesapla:
20 |   |   Doğruluk (Accuracy)
21 |   |   Kesinlik (Precision)
22 |   |   Geri çağırma (Recall)
23 |   |   F1 skoru
24 |   |   ROC AUC skoru
25 |   |   Log loss
26 Görselleştirme:
27 |   ROC AUC eğrilerini çiz:
28 |   |   Her sınıf için ROC eğrisi ve AUC değerlerini hesapla
29 |   |   Konfüzyon matrisini görselleştir (heatmap)
30 Sonuçlar:
31 |   Performans metriklerini tablo halinde görüntüle
32 |   Sınıf bazlı precision, recall ve F1 skorlarını yazdır
33

```

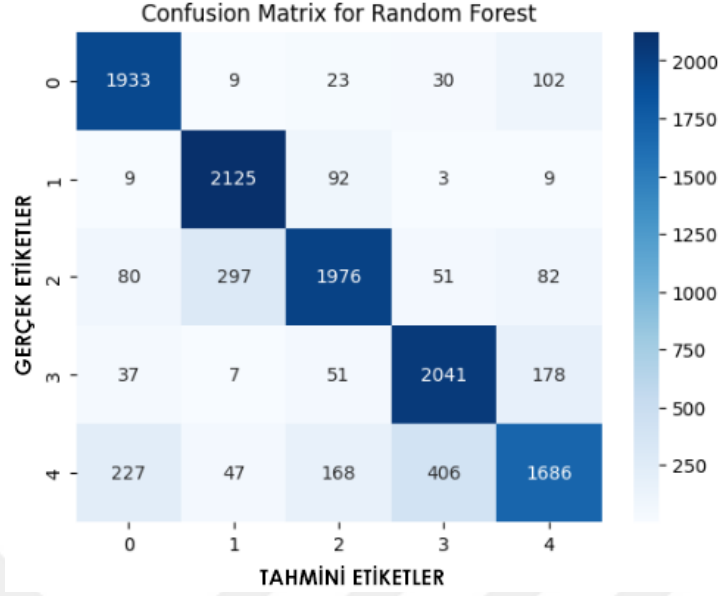
Şekil 4.37. Rastgele Orman Modeli Pseudo Kod

Çizelge 4.20. Rastgele Orman Hiperparametre Bilgileri

Parametre Adı	Değer
Maksimum Ağaç Sayısı	50
Sıcak Başlangıç	Aktif
Rastgele Durum (random state)	42

Çizelge 4.21. Rastgele Orman Test Sonuçları

Sınıf	Kesinlik	Duyarlılık	F1 Skoru	Doğruluk	ROC AUC	Log Loss
Fiyat Performans	0.85	0.92	0.88			
İade Edilen	0.86	0.95	0.90			
Kalitesiz, Kusurlu, Kötü Paketleme	0.86	0.79	0.82			
Paketleme İyi, Hızlı Teslimat	0.81	0.88	0.84			
Tavsiye Edilen, Kaliteli Ürünler	0.82	0.67	0.73			
Ortalama Genel Sonuç	0.84	0.84	0.84	0.84	0.96	0.53

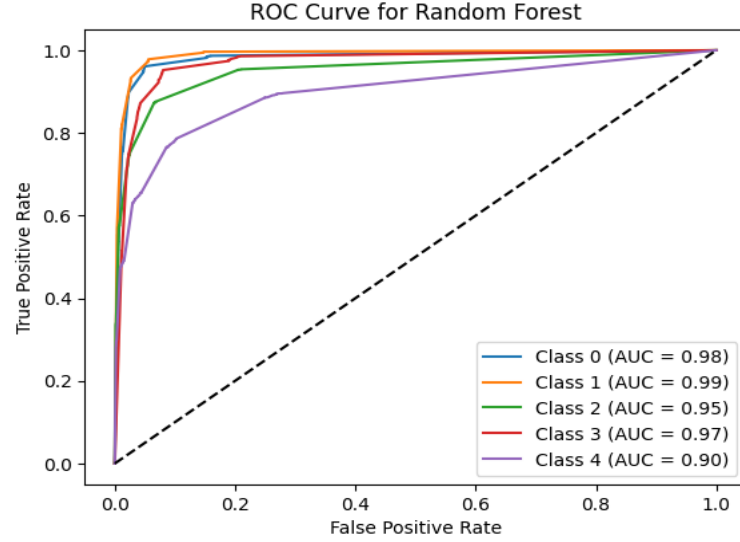


Şekil 4.38. Rastgele Orman modeli ile elde edilen konfüzyon matrisi

Çizelge 4.21’de verilen genel metriklere bakıldığında, doğruluk (0.84), kesinlik (0.84), duyarlılık (0.84) ve F1 skoru (0.84) değerleri modelin hem pozitif hem de negatif sınıfları iyi ayırt ettiğini ortaya koymaktadır. Yüksek bir ROC AUC değeri (0.96), modelin sınıflar arasında ayırma yeteneğinin oldukça güçlü olduğunu, düşük bir Log Loss değeri (0.53) ise modelin olasılık tahminlerinde genel olarak güvenilir olduğunu göstermektedir.

Rastgele Orman modeli, “İade Edilen” ve “Fiyat Performans” sınıflarında yüksek kesinlik (0.86 ve 0.85) ve duyarlılık (0.95 ve 0.92) değerleriyle başarılı sonuçlar elde etmiştir. “Kalitesiz, Kusurlu, Kötü Paketleme” sınıfında ise yüksek kesinlik (0.86) ile doğru pozitif tahminler yapılırken, duyarlılık (0.79) biraz daha düşük kalmıştır, bu da modelin bazı gerçek pozitif örnekleri kaçırdığını göstermektedir. “Paketleme İyi, Hızlı Teslimat” sınıfı, yüksek duyarlılık (0.88) ve F1 skoru (0.84) ile doğru sınıflandırma sağlar, ancak kesinlik (0.81) daha düşük kalmaktadır. “Tavsiye Edilen, Kaliteli Ürünler” sınıfı ise daha düşük duyarlılık (0.67) ve F1 skoru (0.73) değerlerine sahiptir, bu da modelin bu sınıfı doğru tanımakta zorlandığını göstermektedir. Konfüzyon matrisinden, özellikle “Tavsiye Edilen, Kaliteli Ürünler” ve “Fiyat Performans” sınıflarında karışıklıklar olduğu ve bazı sınıfların yanlış sınıflandırıldığı görülmektedir. Bu durum, sınıflar arasındaki benzerliklerin ve veri dengesizliğinin etkisiyle açıklanabilir.

Şekil 4.39 Rastgele Orman modeli sonucunda elde edilen Roc Eğrilerini göstermektedir.



Şekil 4.39. Rastgele Orman modeli ile elde edilen ROC eğrileri

Rastgele Orman modelinin AUC değerleri, tüm sınıflar için genellikle yüksek olup, bu da modelin sınıflar arasında güçlü bir ayırım yapma yeteneğine sahip olduğunu ve pozitif sınıfları doğru şekilde ayırt etmekte başarılı olduğunu göstermektedir. Ancak, Sınıf 4 için AUC değeri biraz daha düşük (0.90) kalmıştır, bu da bu sınıf için diğerlerine kıyasla daha fazla karışıklık yaşandığını göstermektedir.

4.8. E-Ticaret Yorum Analiz Uygulaması

Tez çalışması kapsamında, Türkçe metinler üzerinde en yüksek performansı sergileyen BERT modeli, e-ticaret platformlarından toplanan kullanıcı yorumlarının otomatik olarak sınıflandırılmasına yönelik optimize edilmiş ve bir e-ticaret yorum analiz uygulaması olarak geliştirilmiştir. Uygulama, e-ticaret sitelerine kullanıcılar tarafından girilen yorumları analiz ederek belirlenen kategorilere atamakta ve duygusal sınıflandırmalarını gerçekleştirmektedir. Yorumlar, "Fiyat Performans", "İade Edilen", "Kalitesiz, Kusurlu, Kötü Paketleme", "Paketleme İyi, Hızlı Teslimat" ve "Tavsiye Edilen, Kaliteli Ürünler" olmak üzere beş ana kategoriye ayrılmaktadır.

BERT modelinin güçlü metin sınıflandırma yetenekleri, uygulamanın yorumların anlamını doğru bir şekilde yakalayıp uygun kategorilere atamasını sağlamış ve böylece kullanıcılar ve işletmeler için değerli içgörüler sunan bir çözüm ortaya koyulmuştur. Modelin sağladığı yüksek doğruluk, kesinlik ve duyarlılık oranları sayesinde uygulama, e-ticaret firmalarının müşteri memnuniyetini artırıcı stratejiler geliştirmelerine ve kullanıcıların alışveriş deneyimlerini daha bilinçli bir şekilde değerlendirmelerine olanak

tanınmaktadır. Ayrıca, kullanıcı dostu bir arayüz ile geliştirilen uygulama, yorumların kategori ve duygu sınıflandırma sonuçlarını hızlı ve kolay bir şekilde görüntüleme imkânı sunmaktadır.

Uygulama, yalnızca akademik bir başarı olarak değil, aynı zamanda ticari uygulamalara uygun bir metin analizi çözümü olarak da değerlendirilmektedir. Bireysel kullanıcılar ve işletmeler için e-ticaret yorumlarının analizini kolaylaştırarak, ürün geliştirme, müşteri memnuniyeti ve pazar stratejilerinde önemli katkılar sağlamaktadır.

Streamlit platformu kullanılarak geliştirilen etkileşimli arayüz, kullanıcıların yorum analizi işlemlerini kolayca gerçekleştirmesine olanak tanımaktadır. Kullanıcılar, analiz edilmesini istedikleri yorumu arayüzde yer alan bir metin alanına girerek süreci başlatabilmektedir. Ardından, uygulamanın tahmin işlemini başlatan bir buton yardımıyla model devreye girmekte ve analiz sonucunda tahmin edilen kategori ve duygusal sınıflandırma kullanıcıya açık ve anlaşılır bir şekilde sunulmaktadır. Bu yaklaşım, uygulamanın hem basit hem de erişilebilir bir deneyim sunmasını sağlamaktadır.

Aşağıda yer alan görsellerde (Şekil 4.40, Şekil 4.41, Şekil 4.42, Şekil 4.43, Şekil 4.44 ve Şekil 4.45) uygulamanın farklı kullanıcı yorumlarına ilişkin analiz süreçlerini ve sonuçlarını örnekleyen arayüz tasarımı gösterilmektedir. Bu görseller, uygulamanın işlevselliğini ve kullanıcı dostu yapısını somut bir şekilde gözler önüne sermektedir.

E-Ticaret Yorum Analiz Uygulaması

Bu uygulama, e-ticaret sitelerinden alınan kullanıcı yorumlarını analiz eder ve hangi kategoriye ait olduğunu tahmin eder.

Yorumunuzu girin:

Çok kaliteli sayılmaz fakat fiyatına göre tatmin ediyor

Kategori Tahmin Et

✓ Tahmin Edilen Kategori: Fiyat Performans

🔍 Duygusal Sınıflandırma: Olumlu 👍

Şekil 4.40. Kullanıcı Arayüzü Örnek-1

Şekil 4.40'a göre kullanıcı tarafından girilen yorum, "Kategori Tahmin Et" butonuna basılarak analiz edilmiştir. Tasarlanan uygulama, yorumu "Fiyat Performans" kategorisine atamış ve duygusal sınıflandırmasını "Olumlu" olarak belirlemiştir.

Bu sınıflandırma, yorumun genel tonunu doğru bir şekilde yansıtmaktadır; çünkü kullanıcı, kaliteye dair eleştiriler içerse de ürünün fiyat-performans dengesinden memnuniyet duyduğunu ifade etmektedir. Başka bir deyişle, kullanıcı genel olarak "tatmin edici" ve olumlu bir sonuca ulaştığını belirtmektedir.

Yorumun "Fiyat Performans" kategorisine atanması da son derece yerinde bir değerlendirmedir. Kullanıcının yorumu, ürünün kalitesi ile fiyatı arasındaki dengeye odaklanmakta ve bu ilişkiyi yorumlamaktadır. Dolayısıyla, bu kategoriye yapılan atama, yorumun içeriğiyle tam olarak örtüşmektedir.

Sonuç olarak, hem duygusal sınıflandırma hem de kategori ataması, yorumun ana mesajını doğru bir şekilde yansıtmaktadır.



Şekil 4.41. Kullanıcı Arayüzü Örnek-2

Şekil 4.41, e-ticaret yorumlarını analiz eden uygulamanın kullanıcı arayüzüne ait bir başka örneğini göstermektedir. Kullanıcı tarafından girilen "Ürün yanlış geldi, satıcı firmayla görüştüm, iade ettim" yorumu analiz edilmiş ve "Kategori Tahmin Et" butonuna basılarak sınıflandırılmıştır.

Uygulama, yorumu "İade Edilen" kategorisine atamış ve duygusal sınıflandırmasını "Olumsuz" olarak belirlemiştir. Bu atama, yorumun içeriğiyle

uyumludur, çünkü kullanıcı doğrudan iade sürecini anlatmaktadır ve ürünün yanlış gelmesi nedeniyle bir memnuniyetsizlik yaşadığını ifade etmektedir.

E-Ticaret Yorum Analiz Uygulaması

Bu uygulama, e-ticaret sitelerinden alınan kullanıcı yorumlarını analiz eder ve hangi kategoriye ait olduğunu tahmin eder.

Yorumunuzu girin:

Ürün kırık geldi, bu kargo firmasından ne zaman ürün gelse hasarlı geliyor zaten, kargo firmanızı değiştirmenizi tavsiye ederim

Kategori Tahmin Et

✓ Tahmin Edilen Kategori: Kalitesiz, Kusurlu, Kötü Paketleme

🔍 Duygusal Sınıflandırma: Olumsuz 😞

Şekil 4.42. Kullanıcı Arayüzü Örnek-3

Şekil 4.42'ye göre uygulama, kullanıcı yorumunu "Kalitesiz, Kusurlu, Kötü Paketleme" kategorisine atamış ve duygusal sınıflandırmasını "Olumsuz" olarak belirlemiştir. Bu sınıflandırma, yorumun içeriğiyle tamamen uyumludur. Kullanıcı, ürünün kırık ve hasarlı bir şekilde teslim edildiğini, ayrıca kargo firmasının sürekli olarak sorunlu teslimatlar yaptığını ifade etmektedir. Bu durum, ürünle ilgili olumsuz bir deneyimi ve memnuniyetsizliği açıkça yansıtmaktadır.

E-Ticaret Yorum Analiz Uygulaması

Bu uygulama, e-ticaret sitelerinden alınan kullanıcı yorumlarını analiz eder ve hangi kategoriye ait olduğunu tahmin eder.

Yorumunuzu girin:

Elime çok hızlı ulaştı satıcıya çok teşekkür ederim, kutuyu balon naylonla iyice sarmış

Kategori Tahmin Et

✓ Tahmin Edilen Kategori: Paketleme İyi, Hızlı Teslimat

🔍 Duygusal Sınıflandırma: Olumlu 😊

Şekil 4.43. Kullanıcı Arayüzü Örnek-4

Şekil 4.43’de verilen görselde, uygulama, kullanıcı yorumunu "Paketleme İyi, Hızlı Teslimat" kategorisine atamış ve duygusal sınıflandırmasını "Olumlu" olarak belirlemiştir. Bu sınıflandırma, yorumun içeriğiyle tamamen uyumludur. Kullanıcı, ürünün hızlı bir şekilde eline ulaştığını ve paketlemenin özenli olduğunu vurgulayarak memnuniyetini dile getirmiştir. Hem kategori hem de duygusal sınıflandırma, yorumun genel tonunu ve içeriğini doğru bir şekilde yansıtmaktadır.

Şekil 4.44. Kullanıcı Arayüzü Örnek-5

Şekil 4.44’deki örnekte, uygulama, girilen kullanıcı yorumunu "Tavsiye Edilen, Kaliteli Ürünler" kategorisine atamış ve duygusal sınıflandırmasını "Olumlu" olarak belirlemiştir. Yorumun içeriği, ürünün daha önce de kullanıldığını ve benzer ürünlerle karşılaştırıldığında en kaliteli seçenek olarak öne çıktığını ifade etmektedir. Bu, kullanıcının üründen memnuniyetini açıkça dile getirdiğini ve ürünün kaliteli olduğunu vurguladığını göstermektedir.

Ayrıca, yorum, kullanıcının ürünü daha önce deneyimlemiş olması nedeniyle tavsiye edilebilir bir ürün olduğunu ima etmektedir. Bu bağlamda, sistemin yaptığı "Tavsiye Edilen, Kaliteli Ürünler" kategorisine atama, yorumun içeriğiyle tamamen uyumludur. Aynı şekilde, duygusal sınıflandırmanın "Olumlu" olarak değerlendirilmesi de yerindedir, çünkü yorum genel anlamda pozitif bir kullanıcı deneyimi sunmaktadır. Uygulama, hem ürünün niteliğini hem de kullanıcının yorum tonunu doğru bir şekilde analiz etmiştir.



Şekil 4.45. Kullanıcı Arayüzü Örnek-6

Şekil 4.45’de görülen örnekte ise, geliştirilen uygulama, kullanıcı yorumunu "Kalitesiz, Kusurlu, Kötü Paketleme" kategorisine atamış ve duygusal sınıflandırmasını "Olumsuz" olarak belirlemiştir. Yorumun içeriği incelendiğinde, kullanıcının teslimatın hızlı olmasından memnuniyet duyduğu ancak ürünün kalitesinden şikayet ettiği anlaşılmaktadır. Özellikle, ürünün ince ve dayanıksız bir yapıya sahip olması, kullanıcı tarafından memnuniyetsizlik kaynağı olarak belirtilmiştir.

Uygulamanın, bu yorumu "Kalitesiz, Kusurlu, Kötü Paketleme" kategorisine ataması yerindedir; çünkü yorumun ana vurgusu, ürünün kalitesizliği üzerinedir. Ayrıca, duygusal sınıflandırmanın "Olumsuz" olarak belirlenmesi de doğru bir değerlendirmedir, çünkü kullanıcı, genel olarak ürünle ilgili olumsuz bir deneyimini dile getirmiştir. Geliştirilen uygulama, yorumun hem kategorisini hem de duygusal tonunu başarılı bir şekilde analiz etmiştir.

Tez çalışmasında geliştirilen uygulamanın örnek yorumlar üzerindeki performansı incelendiğinde, yorumların hem kategoriler hem de duygusal tonlar açısından doğru bir şekilde sınıflandırıldığı gözlemlenmiştir. Bu başarı, uygulamanın yorumların anlamını doğru bir şekilde kavrayarak belirlenen kategorilere isabetli atamalar yapabildiğini göstermektedir.

Bu tez çalışması, e-ticaret platformlarında kullanıcı yorumlarının otomatik olarak analiz edilmesi amacıyla etkili bir çözüm sunmaktadır. Özellikle müşteri memnuniyetini artırmak ve ürün iyileştirme süreçlerine katkı sağlamak için bu tür bir sınıflandırma sistemi büyük önem taşımaktadır. Gelecekte, modelin daha geniş veri kümeleriyle

yeniden eğitilerek performansı artırılabilir ve daha fazla kategori eklenerek kullanım alanı genişletilebilir. Uygulama, hem akademik hem de ticari alanlarda metin sınıflandırma sistemlerinin bir örneği olarak değerlendirilebilir ve bu tür sistemlerin geliştirilmesi, veri analizi ve doğal dil işleme alanlarında önemli bir yer tutmaktadır.



5. SONUÇLAR VE ÖNERİLER

Gerçekleştirilen tez çalışmasında, Türkçe e-ticaret müşteri yorumlarının çok etiketli analizi amacıyla kullanılan derin öğrenme ve geleneksel makine öğrenmesi modellerinin performansları kapsamlı bir şekilde incelenmiş ve bu süreçte elde edilen bulgular oldukça değerli sonuçlar ortaya koymuştur.

Çizelge 5.1’de, çalışmada oluşturulan tüm modellerin performansları karşılaştırmalı olarak gösterilmiştir;

Çizelge 5.1. Çalışmada oluşturulan tüm modellerin performans metrikleri

Model	Doğruluk	Kesinlik	Duyarlılık	F1 Skor	ROC AUC	Log Loss
CNN	0.85	0.86	0.86	0.85	0.97	0.43
RNN	0.84	0.84	0.84	0.84	0.96	0.50
LSTM	0.86	0.86	0.87	0.86	0.97	0.43
BiLSTM	0.85	0.86	0.86	0.85	0.97	0.44
GRU	0.78	0.78	0.78	0.78	0.95	0.62
BERT	0.89	0.89	0.89	0.89	0.98	0.41
Lojistik Regresyon	0.85	0.85	0.85	0.85	0.97	0.47
Rastgele Orman	0.84	0.84	0.84	0.84	0.96	0.53

Çalışmada öne çıkan model BERT, bağlamsal anlam çıkarma yetenekleriyle tüm performans metriklerinde üstünlük sağlamıştır. BERT modelinin başarısı, transformer tabanlı yaklaşımların bağlam içeren metinlerin analizinde ne denli etkili olabileceğini gözler önüne sermektedir. Diğer derin öğrenme modelleri, özellikle LSTM, BiLSTM ve CNN gibi modeller de yüksek performans sergileyerek doğruluk, kesinlik ve duyarlılık gibi metriklerde tatmin edici sonuçlar elde etmiştir. Ancak BERT, özellikle bağlamı çift yönlü olarak analiz etme kabiliyeti, dilin anlamını daha derinlemesine kavrayarak diğer modellere göre daha doğru ve anlamlı sınıflandırmalar yapmasına olanak sağlamıştır.

Bu durum, derin öğrenme modellerinin karmaşık veri yapılarında geleneksel makine öğrenmesi algoritmalarına kıyasla daha etkili sonuçlar üretebildiğini ve özellikle metin tabanlı verilerde BERT gibi transformer tabanlı modellerin ne kadar güçlü olabileceğini bir kez daha göstermektedir. BERT, hem dilin içyapısını hem de bağlamını anlamadaki başarısı, çok etiketli sınıflandırma gibi karmaşık görevlerde büyük bir avantaj sunmaktadır. Ayrıca, elde edilen bulgular, doğru model seçiminin veri analitiği süreçlerinde ne kadar kritik olduğunu ve her modelin farklı veri setleri ve görevler için değişen avantajlar sunduğunu ortaya koymaktadır. Bu nedenle, model seçimi yapılırken

yalnızca doğruluk gibi temel metriklerin değil, aynı zamanda problem gereksinimlerinin ve sınıflandırma hedeflerinin de dikkate alınması gerektiği anlaşılmaktadır. Özellikle çok etiketten oluşan sınıflandırma problemi gibi karmaşık görevlerde, BERT bağlamsal bilgiyi işleyebilme gücü, etiket sayısı, veri dengesizliği ve bağlamsal faktörler gibi unsurları göz önünde bulundurulduğunda, oldukça belirleyici bir etkiye sahiptir.

Çalışmada kullanılan çok etiketli sınıflandırma yaklaşımı, e-ticaret müşteri yorumlarının farklı kategorilere ayrılarak daha zengin bir anlamlandırma yapılmasını sağlamış ve yorumların hem içerik hem de duygu boyutunda daha derinlemesine analiz edilebilmesine olanak tanımıştır. Bu, yalnızca bir yorumun basit bir şekilde olumlu ya da olumsuz olarak sınıflandırılmasının ötesine geçerek, yorumların taşıdığı daha geniş bağlamsal bilgilerin ortaya çıkarılmasına imkân tanımıştır. Bu analizlerin, işletmelerin pazarlama stratejileri, ürün geliştirme süreçleri ve müşteri memnuniyeti yönetimi gibi alanlarda somut katkılar sağlayabileceği öngörülmektedir.

Tez çalışmasının sonuçları, hem akademik hem de uygulamalı çalışmalar açısından geniş bir potansiyel sunmaktadır. Akademik olarak, Türkçe metin analizi alanında önemli bir katkı sağlanmış ve özellikle çok etiketli sınıflandırma gibi karmaşık görevlerde yapay zekâ tabanlı yöntemlerin başarısı ortaya konulmuştur. Çalışma, Türkçe gibi morfolojik açıdan zengin bir dilde, e-ticaret yorumlarının bağlam odaklı sınıflandırılmasıyla literatüre katkı sağlamış ve bu alanda gelecekte yapılacak çalışmalara güçlü bir temel oluşturmuştur. Uygulama bağlamında ise, geliştirilmiş modellerin e-ticaret platformlarına entegrasyonu ile yalnızca sınıflandırma süreçlerinin hızlandırılması değil, aynı zamanda müşteri geri bildirimlerinden elde edilen iç görülerle daha verimli pazarlama stratejileri oluşturulabileceği ve kullanıcı deneyiminin daha kişiselleştirilmiş hale getirilebileceği anlaşılmaktadır. Özellikle, kullanıcı yorumlarının doğru ve anlamlı bir şekilde etiketlenmesi, işletmelerin büyük ölçekli metin verilerinden değerli bilgiler çıkarmasını kolaylaştırmaktadır.

Veri setinin dengesizliği, metin çoğaltma ve rastgele örnekleme yöntemleriyle büyük ölçüde giderilmiş, böylece modellerin performansındaki tutarsızlıklar önemli ölçüde azaltılmıştır. Bununla birlikte, nadir bulunan sınıfların diğer sınıflarla karışma oranının yüksek olması, sınıf bazlı ayrımların iyileştirilmesi gerektiğini ortaya koymaktadır. Veri seti ve temsil açısından, daha fazla kategori eklenerek yorumların çeşitliliğinin artırılabilmesi vurgulanmıştır. Bu çeşitlilik, özellikle az temsil edilen sınıfların performansını iyileştirebilir ve modellerin genelleme kapasitesini artırabilir. Ayrıca, veri setlerinin manuel olarak etiketlenmiş, daha çok sayıda örnekle genişletilmesi,

modellerin hem genelleme kabiliyetini hem de hassasiyetini artırabilir. Hibrit model yaklaşımlarının, örneğin CNN ve LSTM gibi yapıların birleştirilerek yeni çözümler üretilmesi, çoklu etiketli sınıflandırma alanında daha önemli bir başarı sağlayabilir. Bunun yanı sıra, e-ticaret yorumlarının analizinde veri çeşitliliğinin artırılması, örneğin farklı platformlardan gelen yorumları içeren daha geniş bir veri tabanı oluşturulması, modellerin çeşitli veri kaynaklarına uyumunu güçlendirebilir.

Gelecekteki çalışmalar için birkaç önemli öneri geliştirilmiştir. Öncelikle, transformer tabanlı modellerin, örneğin Türkçe BERT gibi önceden eğitilmiş ve optimize edilmiş versiyonlarının, özellikle daha dengeli ve çeşitlendirilmiş veri setleriyle yeniden eğitilmesi oldukça faydalı olabilir. Veri setlerinin manuel olarak etiketlenmiş, daha çok sayıda örnekle genişletilmesi, modellerin hem genelleme kabiliyetini hem de hassasiyetini artırabilir. Ayrıca, hibrit model yaklaşımlarının, örneğin CNN ve LSTM gibi yapıların birleştirilerek yeni çözümler üretilmesi, çoklu etiketli sınıflandırma alanında daha önemli bir başarı sağlayabilir. Bunun yanı sıra, e-ticaret yorumlarının analizinde veri çeşitliliğinin artırılması, örneğin farklı platformlardan gelen yorumları içeren daha geniş bir veri tabanı oluşturulması, modellerin çeşitli veri kaynaklarına uyumunu güçlendirebilir.

Bu bulgular ve öneriler ışığında, çalışmanın e-ticaret sektöründe müşteri yorumlarının analizine yenilikçi ve etkili çözümler sunduğu, hem akademik hem de pratik alanlarda önemli katkılar sağladığı söylenebilir.

KAYNAKLAR

- Akar, Ö. ve Güngör, O., (2012), Rastgele Orman Algoritması Kullanılarak Çok Bantlı Görüntülerin Sınıflandırılması, *Jeodezi ve Jeoinformasyon Dergisi*, Sayı: 106, 139-146.
- Akdeniz, A. C., (2022), Makine Öğrenmesi Yöntemleri Kullanılarak Uzaktan Eğitim Konulu Türkçe Tweetlerin Duygu Analizi, Yüksek Lisans Tezi, *Konya Teknik Üniversitesi Lisansüstü Eğitim Enstitüsü, Bilgisayar Mühendisliği Anabilim Dalı*, Konya.
- Albayrak, A., (2018), Duygu Analizinde Farklı Vektör Temsil Yöntemleri ve Sınıflayıcıların Karşılaştırılması, Yüksek Lisans Tezi, *Sivas Cumhuriyet Üniversitesi Sosyal Bilimler Enstitüsü*, Sivas, 13-43.
- Aslanpençesi, Z., (2024), Derin Öğrenme Yöntem ve Modelleri Kullanılarak Sosyal Mühendislik Saldırılarının Tespitinde Yeni Yaklaşımların Geliştirilmesi, Yüksek Lisans Tezi, *Fırat Üniversitesi Fen Bilimleri Enstitüsü, Yazılım Mühendisliği Anabilim Dalı*, Elazığ.
- Aydoğan, M. ve Karcı, A., (2019), Kelime Temsil Yöntemleri ile Kelime Benzerliklerinin İncelenmesi, *Çukurova Üniversitesi Mühendislik-Mimarlık Fakültesi Dergisi*, Cilt: 34, Sayı: 2, 181-196.
- Bilekyiğit, S., & Eldem, A. (2022). Kalp Yetmezliği Riskinin Makine Öğrenmesi Yöntemleri İle Analiz Edilmesi, Yüksek Lisans Tezi. *Karamanoğlu Mehmetbey Üniversitesi, Fen Bilimleri Enstitüsü, Mühendislik Bilimleri Ana Bilim Dalı*.
- Bingöl, K. (2020). Depreme Dayanıklı Mimari Tasarım Aşamasında Derin Öğrenme Ve Görüntü Sınıflama Yöntemi İle Burulma Düzensizliği Tespiti, Yüksek Lisans Tezi. *Çankaya Üniversitesi, Fen Bilimleri Enstitüsü, Mimarlık Ana Bilim Dalı*.
- Bonissone, P., Cadenas, J. M., Garrido, M. C., & Díaz-Valladares, R. A. (2010). A fuzzy random forest. *International Journal of Approximate Reasoning*, 51(7), 729-747.
- Bozdoğan, A. (2024). Lokal İleri Meme Kanseri Cerrahi Tedavi Öncesi Uygulanan Kemoterapi Cevabını Etkileyen Faktörlerin Lojistik Regresyon Ve ROC Analizi İle Değerlendirilmesi. *İstanbul Ticaret Üniversitesi, Fen Bilimleri Enstitüsü, İstatistik Ana Bilim Dalı*.
- Chentouf, I. (2024). Diagnosis of Alzheimer's disease with deep learning: A hybrid 3D CNN and RNN approach. *Bahçeşehir Üniversitesi, Lisansüstü Eğitim Enstitüsü, Yapay Zeka Ana Bilim Dalı*.
- Christian, H., Agus, M. P., & Suhartono, D. (2016). Single Document Automatic Text Summarization Using Term Frequency-Inverse Document Frequency (TF-IDF). *ComTech: Computer, Mathematics and Engineering Applications*, 7(4), 285-294.

- Custer, C. (2020). 15 Python libraries for data science you should know. *Dataquest*. Erişim adresi: <https://www.dataquest.io/blog/15-python-libraries-for-data-science/> [Ziyaret Tarihi: 10.12.2024].
- Çabuk, M., (2023), Makine Öğrenmesi ile E-Ticaret Ürün Yorumlarının Analizi, Yüksek Lisans, *Celal Bayar Üniversitesi, Fen Bilimleri Enstitüsü, Yazılım Mühendisliği Anabilim Dalı* Manisa, 12-75.
- Çataltaş, M., Doğramacı, S., Yumuşak, S., and Öztoprak, K., (2020), Extraction of Product Defects and Opinions from Customer Reviews by Using Text Clustering and Sentiment Analysis, *IEEE International Conference on Big Data (Big Data)*, Atlanta, ABD, 4529-4534.
- Çelik, M. Y., (2011), Biyoistatistik Bilimsel Araştırma SPSS Nasıl?, Kişisel Yayınlar, 561, Diyarbakır.
- Çınar, A., (2020), Sınıflandırma algoritmaları ile bir metin madenciliği uygulaması: Veri madenciliği ve makine öğrenmesi, temel kavramlar, algoritmalar, uygulamalar, Çağlayan Kitapevi ve Eğitim Çözümleri Tic. A.Ş., İstanbul, 105-140.
- Çoban, Ö., (2016), Metin sınıflandırma teknikleri ile Türkçe Twitter duygu analizi, Yüksek Lisans Tezi, *Atatürk Üniversitesi Fen Bilimleri Enstitüsü*, Erzurum.
- Demircan, M., Seller, A., Abut, F., & Akay, M. F. (2021). Developing Turkish Sentiment Analysis Models Using Machine Learning And E-Commerce Data. *International Journal of Cognitive Computing in Engineering*, 2, 202-207.
- Deniz, E. (2023). Türkçe E-Ticaret Müşteri Yorumlarının Derin Öğrenme İle Çok Etiketli Analizi. *Kırıkkale Üniversitesi, Fen Bilimleri Enstitüsü, Bilgisayar Bilimleri Ana Bilim Dalı*.
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 1, 4171-4186. <https://doi.org/10.18653/v1/N19-1423>
- Dey, R., & Salem, F. M. (2017). Gate-variants of gated recurrent unit (GRU) neural networks. *2017 IEEE 60th International Midwest Symposium on Circuits and Systems (MWSCAS)*, 1597-1600.
- Emirali, E. (2022). bi-TEZAT: BiLSTM Yöntemiyle Türkçe Şikayet Metinlerinde Zaman İfadelerinin Tespit Edilmesi. *Yıldız Teknik Üniversitesi, Fen Bilimleri Enstitüsü, Bilgisayar Mühendisliği Ana Bilim Dalı*.
- Erdaş, A., & Güzel, M. S. (2022). Medikal veri üzerinde makine öğrenmesi yöntemlerinin karşılaştırılması. *Journal of Engineering Sciences and Design*, 10(1), 1-10.

- Erkan, S. (2023). E-ticaret sitelerinde yer alan yorumların BERT ve ALBERT dil modelleri ile analizi. *İstanbul Beykent Üniversitesi, Lisansüstü Eğitim Enstitüsü, Bilgisayar Mühendisliği Ana Bilim Dalı*.
- Faraggi, D., & Reiser, B. (2002). Estimation of the area under the ROC curve. *Statistics in Medicine*, 21(20), 3093-3106.
- Gers, F. A., Schmidhuber, J., & Cummins, F. (2000). Learning to forget: Continual prediction with LSTM. *Neural Computation*, 12(10), 2451-2471.
- Gezici, G., & Yanıkoğlu, B. (2018). Sentiment analysis in Turkish. In *Theory and Applications of Natural Language Processing* (pp. 255–271).
- Gordin, D. M. (2015). *Scientific Babel: How science was done before and after global English*. Chicago, IL: University of Chicago Press.
- Graves, A. (2012). *Supervised sequence labelling with recurrent neural networks*. Springer Science & Business Media.
- Greff, K., Srivastava, R. K., Koutník, J., Steunebrink, B. R., & Schmidhuber, J. (2017). LSTM: A search space odyssey. *IEEE Transactions on Neural Networks and Learning Systems*, 28(10), 2222-2232.
- Han, J., & Moraga, C. (1995). The influence of the sigmoid function parameters on the speed of backpropagation learning. *International Workshop on Artificial Neural Networks*, 195–201.
- Kahraman, S. (2023). Web kazıma ve makine öğrenmesi yöntemleri kullanılarak fiyat tahminleme: İkinci el araç piyasasında bir örnek. Yüksek Lisans, *Sakarya Üniversitesi, Fen Bilimleri Enstitüsü, Bilişim Sistemleri Mühendisliği Ana Bilim Dalı*.
- Karakuş, B. A., Talo, M., Hallaç, I., & Aydın, G. (2018). Evaluating deep learning models for sentiment classification. *Concurrency and Computation: Practice and Experience*, 30(21), e4783. <https://doi.org/10.1002/cpe.4783>
- Korkusuz, R. (2018). Futbola ilişkin Twitter paylaşımlarının duygu analizi. Yüksek Lisans, *Trakya Üniversitesi, Fen Bilimleri Enstitüsü, Edirne*, 13-21.
- Kuzucu, K. (2015). Müşteri memnuniyeti belirlemek için metin madenciliği tabanlı bir yazılım aracı. *Maltepe Üniversitesi, Fen Bilimleri Enstitüsü, İstanbul*.
- Küçük, Z. (2023). Matris hatası nedir? - Confusion Matrix. *Womaneng*. Erişim adresi: <https://womaneng.com/matris-hatasi-nedir-confusion-matrix/> [Ziyaret Tarihi: 05.11.2024].
- Kızıllırmak, E. (2020). İngilizce-Türkçe çeviri metinlerde Levenshtein uzaklığı ile desteklenmiş çapa tabanlı cümle eşleme. Yüksek Lisans Tezi, *Maltepe Üniversitesi, Lisansüstü Eğitim Enstitüsü, İstanbul*.

- Malkoç, B. (2012). Temel bilimler ve mühendislik eğitiminde programlama dili olarak Python. *XIV. Akademik Bilişim Konferansı Bildirileri*, 201.
- Mayda, I., & Korkmaz, M. (2018). Sentiment analysis with Turkish adjective dictionary. *Innovations in Intelligent Systems and Applications Conference (ASYU)*, Adana, 4-6 Ekim, 1-6. IEEE.
- Mengutayci, Ü., & Temurtas, H. (2021). Yapay Sinir Ağları ile Türkçe Otel Yorumlarının Sınıflandırılması. *International Black Sea Coastline Countries Scientific Research Symposium - VI*, Nisan, 683-687.
- Özkan, Y., (2008), Veri Madenciliği Yöntemleri, 1. Baskı, Papatya Yayıncılık Eğitim, İstanbul.
- Öztürk, Y., Kılınç, H. Ç., & Polat, A., (2022), Hibrit Parçacık Sürüsü Optimizasyonu ile Geçitli Tekrarlayan Birim Modeli Kullanılarak Nehir Akım Tahmini: Ceyhan Havzası Örneği, *Avrupa Bilim ve Teknoloji Dergisi*, 41, 202-210.
- Patro, V. M. ve Patra, M. R., (2014), Augmenting Weighted Average with Confusion Matrix to Enhance Classification Accuracy, *Transactions on Machine Learning and Artificial Intelligence*, 2 (4), 77-91.
- Pervan, N. and Keleş, H. Y., (2017), Sentiment Analysis Using a Random Forest Classifier on Turkish Web Comments, *Communications Faculty of Sciences University of Ankara Series A2-A3 Physical Sciences and Engineering*, 59 (2), 69-79.
- Polat, B., (2021), Türkçe Ürün Yorumları Verisi ile Duygu Analizi, Yüksek Lisans, Bilgisayar Mühendisliği, Ankara, 11-55.
- Rastogi, N., Yadav, K., Divya, K., and Perwej, Y., (2022), Sentimental Analysis on Web Scraping Using Machine Learning Method, *Journal of Information and Computational Science*, 12 (8), 24-29.
- Rumelli, M., Akkuş, D., Kart, Ö., and Isik, Z., (2019), Türkçe Metinlerde Makine Öğrenmesi Algoritmaları ile Duygu Analizi, *2019 Innovations in Intelligent Systems and Applications Conference (ASYU)*, İzmir, 31 Ekim – 02 Kasım, IEEE, 15.
- Sar, K. T., (2021), “Yapay sinir ağları ve BERT dil modeli kullanılarak zaman bazlı duygu analizi: WhatsApp yeni gizlilik sözleşmesine yönelik yorumların araştırılması”, Yüksek Lisans Tezi, *Dokuz Eylül Üniversitesi Sosyal Bilimler Enstitüsü*, İzmir.
- Taşkıran, S. F., (2021), Doğal Dil İşleme ile Akademik Metinlerin Kümelenmesi, Yüksek Lisans, *Konya Teknik Üniversitesi Bilgisayar Mühendisliği Anabilim Dalı*, Konya.
- Tomak, L. ve Yüksel, B., (2009), İşlem Karakteristik Eğrisi Analizi ve Eğri Altında Kalan Alanların Karşılaştırılması, *Journal of Experimental and Clinical Medicine*, 27 (2).

- Toçoğlu, M. A., Öztürkmenoğlu, O., and Alpkoçak, A., (2019), Emotion Analysis from Turkish Tweets Using Deep Neural Networks, *IEEE Access*, 7, 183061-183069.
- Turan, F. (2018). ASP. NET MVC 4 teknolojilerini kullanarak bir e-ticaret sitesi uygulamasının geliştirilmesi. Yüksek Lisans Tezi, *Beykent Üniversitesi Sosyal Bilimler Enstitüsü, İşletme Yönetimi Ana Bilim Dalı*.
- Uçar, K. T., (2020), BERT Modeli ile Türkçe Metinlerde Sınıflandırma Yapmak, Web adresi: <https://medium.com/@ktoprakucar/bert-modeli-ile-t%C3%BCrk%C3%A7e-metinlerde-s%C4%B1n%C4%B1fland%C4%B1rma-yapmak-260f15a65611> [Ziyaret Tarihi: 08.10.2024].
- Yelmen, İ., (2016), “Doğal dil işleme yöntemleriyle Türkçe sosyal medya verileri üzerinde duygu analizi”, Yüksek Lisans Tezi, *İstanbul Aydın Üniversitesi Fen Bilimleri Enstitüsü*, İstanbul, 21-33.
- Yeşilyurt, M. (2024). SQL Server DML işlemleri üzerinde performans ve kod yazımı açısından C# ve Python programlama dillerinin karşılaştırılması. Yüksek Lisans Tezi, *Atatürk Üniversitesi Sosyal Bilimler Enstitüsü, Yönetim Bilişim Sistemleri Ana Bilim Dalı*.
- Yönet, E. (2023). Şiir kategorisinin doğal dil işleme yöntemleri kullanılarak tahmin edilmesi. Yüksek Lisans Tezi, *Konya Teknik Üniversitesi Lisansüstü Eğitim Enstitüsü, Bilgisayar Mühendisliği Ana Bilim Dalı*.
- Yüceoğlu, B., (2017), Scikit-Learn ile Veri Analitiğine Giriş [online], Web adresi: <http://www.veridefteri.com/2017/11/23/scikit-learn-ile-veri-analitigine-giris/> [Ziyaret Tarihi: 25.08.2024].
- Yıldırım, S., (2018), Twitter verileriyle duygu analizi ve Türkçe duygu kütüphanesi, Yüksek Lisans Tezi, *Bahçeşehir Üniversitesi Fen Bilimleri Enstitüsü*, İstanbul, 517.