



SAMSUN ÜNİVERSİTESİ
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ

**SOĞAN MİMARİSİNDE METİN İÇERİKLERİNİN YAPAY ZEKÂ DESTEKLİ
MODELLER İLE DEĞERLENDİRİLMESİ VE DAĞITIMI**

Semih Osman SAKA

YÜKSEK LİSANS TEZİ

Yazılım Mühendisliği Anabilim Dalı

Danışman: Doç. Dr. Zafer CÖMERT

Ağustos 2024

01/08/2024

ETİK İLKE VE KURALLARA UYGUNLUK BEYANNAMESİ

Bu tezin bana ait, özgün bir çalışma olduğunu; çalışmamın hazırlık, veri toplama, analiz ve bilgilerin sunumu olmak üzere tüm aşamalarında bilimsel etik ilke ve kurallara uygun davrandığımı; bu çalışma kapsamında elde edilen tüm veri ve bilgiler için kaynak gösterdiğimi ve bu kaynaklara kaynakçada yer verdiğimi; bu çalışmanın Samsun Üniversitesi tarafından kullanılan “bilimsel intihal tespit programı”yla tarandığını ve hiçbir şekilde “intihal içermediğini” beyan ederim. Herhangi bir zamanda, çalışmamla ilgili yaptığım bu beyana aykırı bir durumun saptanması durumunda, ortaya çıkacak tüm ahlaki ve hukuki sonuçları kabul ettiğimi bildiririm.

Semih Osman SAKA

ÖNSÖZ

Bu tez çalışması, Yüksek Lisans derecemi almak için hazırladığım bir araştırmanın sonucudur. Tez çalışmamın konusunun belirlenmesinde ve hazırlanma sürecinin her aşamasında bilgilerini, tecrübelerini ve değerli zamanlarını esirgemeyerek bana her fırsatta rehberlik eden ve yardımcı olan değerli danışman hocam Doç. Dr. Zafer CÖMERT'e sabırları ve özverileri için sonsuz teşekkürlerimi sunarım.

Ağustos 2024

Semih Osman SAKA



İÇİNDEKİLER

ÖNSÖZ.....	III
ÖZET.....	VII
ABSTRACT	VIII
Tablolar Listesi.....	IX
Şekiller Listesi.....	X
Kısaltmalar Dizini	XI

BİRİNCİ BÖLÜM

1. Giriş.....	1
1.1. Problem	1
1.2. Araştırmanın Amacı	2
1.3. Araştırmanın Önemi	3
1.4. Sayıtlar ve Araştırma Soruları	3
1.5. Sınırlılıklar	4

İKİNCİ BÖLÜM

2. Literatür Taraması.....	6
----------------------------	---

ÜÇÜNCÜ BÖLÜM

3. Materyal ve Yöntemler.....	14
3.1. Veri Seti.....	14
3.2. Doğal Dil İşleme.....	15
3.3. Veri Önleme	17
3.4. Veri setinin eğitim ve test seti olarak ayrılması.....	17
3.5. Derin Öğrenme	18
3.6. Derin Öğrenme Modellerinde Sık Kullanılan Katmanlar.....	18
3.6.1. Giriş Katmanı	18
3.6.2. Embedding Katmanı	19

3.6.3.	Dropout Katmanı	19
3.6.4.	LSTM Katmanı.....	19
3.6.5.	Tam Bağlı Katman	21
3.6.6.	Batch Normalizasyon Katmanı	21
3.7.	Optimizasyon	22
3.7.1.	Stochastic Gradient Descent (SGD).....	22
3.7.2.	Momentum	22
3.7.3.	AdaGrad.....	23
3.7.4.	RMSProp.....	23
3.7.5.	Adaptive Moment Estimation (Adam)	23
3.8.	Kayıp Fonksiyonu.....	24
3.8.1.	Kategorik Çapraz Entropi.....	24
3.8.2.	İkili Çapraz Entropi.....	24
3.8.3.	Ortalama Kare Hata	24
3.9.	Model Doğrulama ve Değerlendirme	25
3.9.1.	Hata matrisi	25
3.9.2.	ROC Eğrileri.....	26

DÖRDÜNCÜ BÖLÜM

4.	DeneySEL Sonuçlar ve Bulgular	27
4.1.	Veri Seti Analizi.....	27
4.2.	Yorum Uzunluğu Analizi	29
4.3.	N-gram Analizi	30
4.4.	Bileşene Yönelik Duygu Analizi	33
4.5.	Konu Analizi	35
4.6.	Duygu Analizi için CNN tabanlı Model Önerisi	36
4.6.1.	Model yapısı ve hiperparametreler.....	36

4.6.2.	Model Eğitimi	39
4.6.3.	Erken Durdurma.....	40
4.6.4.	İkili Sınıflandırma	42
4.6.5.	Karşılaştırmalı Sonuçlar.....	44
4.7.	Transfer Öğrenme ile Duygu Analizi	46
4.7.1.	nnlm-en-dim50.....	46
4.7.2.	nnlm-en-dim128.....	48
4.7.3.	USE Modeli	49
4.7.4.	Karşılaştırmalı Sonuçlar.....	51
BEŞİNCİ BÖLÜM		
5.	Sonuçlar, Tartışma ve Öneriler	52
5.1.	Literatürdeki ilgili Çalışmaların Karşılaştırılması	52
5.2	Sonuçlar ve Öneriler	57
KAYNAKÇA		59
ÖZGEÇMİŞ.....		62

ÖZET

SOĞAN MİMARİSİNDE METİN İÇERİKLERİNİN YAPAY ZEKÂ DESTEKLİ MODELLER İLE DEĞERLENDİRİLMESİ VE DAĞITIMI

Semih Osman SAKA

Yazılım Mühendisliği Anabilim Dalı

Samsun Üniversitesi, Lisansüstü Eğitim Enstitüsü, Ağustos 2024

Danışman: Doç. Dr. Zafer CÖMERT

Bu tez çalışması, müşteri geri bildirimlerinin yapay zekâ destekli modellerle analiz edilmesi ve elde edilen içgörülerin soğan mimarisi kullanılarak dağıtılmasını amaçlamaktadır. Araştırma, özellikle havayolu taşımacılığı sektöründe müşteri memnuniyetini artırmak ve hizmet kalitesini iyileştirmek için yenilikçi yöntemler sunmayı hedeflemektedir. Tez kapsamında, doğal dil işleme (NLP) teknikleri ve derin öğrenme modelleri kullanılarak metin tabanlı müşteri yorumları analiz edilmiştir. Çalışmanın birinci amacı, müşteri yorumlarından anlamlı içgörüler elde etmek olup, bigram, trigram analizleri ve duygu analizi gibi yöntemlerle müşterilerin yoğunlaştığı konular ve duygusal tepkileri belirlemektir. İkinci amaç, duygu analizi ile müşteri memnuniyetini değerlendirerek pozitif, negatif veya nötr duyguların sınıflandırılmasını sağlamaktır. Üçüncü amaç, elde edilen analiz sonuçlarının soğan mimarisi ile etkin bir şekilde dağıtılmasıdır. Bu mimari, yazılım geliştirme süreçlerinde katmanlı ve modüler bir yaklaşım sunarak uygulamaların sürdürülebilirliğini ve esnekliğini artırır. Son olarak, dördüncü amaç, müşteri geri bildirimlerinden elde edilen verilerle hizmet kalitesini iyileştirmek için karar destek sistemleri geliştirmektir. Araştırma kapsamında kullanılan veri seti, havayolu sektöründe müşteri deneyimlerine ait metin tabanlı yorumları içermektedir. Bu yorumlar, doğrulama durumu, içerik, değerlendirme puanı, tavsiye durumu ve duygu analizi gibi bilgiler içermektedir. Veri seti, duygu analizi için ön işleme tabi tutularak eğitim ve test seti olarak ayrılmıştır. Derin öğrenme modelleri ve optimizasyon algoritmaları kullanılarak elde edilen sonuçlar, hizmet kalitesinin artırılması ve müşteri memnuniyetinin iyileştirilmesi için stratejik öneriler sunmaktadır.

Anahtar Sözcükler: Doğal dil işleme, Duygu analizi, Derin öğrenme, Sınıflandırma

ABSTRACT

EVALUATION AND DISTRIBUTION OF TEXT CONTENTS WITH MODELS SUPPORTED BY ARTIFICIAL INTELLIGENCE TECHNIQUES IN ONION ARCHITECTURE

Semih Osman SAKA

Software Engineering Graduate Program

Samsun University, Graduate School, August 2024

Supervisor: Assoc. Prof. Dr. Zafer COMERT

This thesis aims to analyze customer feedback using AI-supported models and distribute the obtained insights through onion architecture. The research specifically targets the airline transportation sector to enhance customer satisfaction and improve service quality by providing innovative methods. The thesis employs natural language processing (NLP) techniques and deep learning models to analyze text-based customer reviews. The first objective of the study is to extract meaningful insights from customer reviews by identifying topics and emotional responses through methods such as bigram, trigram analyses, and sentiment analysis. The second objective is to evaluate customer satisfaction through sentiment analysis, categorizing the reviews into positive, negative, or neutral sentiments. The third objective is to effectively distribute the analysis results using onion architecture. This architecture offers a layered and modular approach in software development, enhancing the sustainability and flexibility of applications. Lastly, the fourth objective is to develop decision support systems to improve service quality based on the data obtained from customer feedback. The dataset used in the research comprises text-based comments on airline experiences, including information such as verification status, review content, rating, recommendation status, and sentiment analysis. The dataset is pre-processed for sentiment analysis and split into training and testing sets. Using deep learning models and optimization algorithms, the results provide strategic recommendations for enhancing service quality and improving customer satisfaction.

Keywords: Natural language processing, Sentiment analysis, Deep learning, Classification

Tablolar Listesi

Tablo 3.1 Verilerin duygu durumuna göre dağılımı.....	14
Tablo 3.2 Performans metrikleri ve kısa açıklamaları	25
Tablo 4.1 Bigrams frekans tablosu.....	31
Tablo 4.2 Trigram frekans tablosu	32
Tablo 4.3 Model eğitiminde kullanılan hiperparametreler ve değerleri.....	38
Tablo 4.4 Karşılaştırmalı performans sonuçları.....	45
Tablo 4.5 Transfer öğrenmeye ilişkin karşılaştırmalı sonuçlar.....	51
Tablo 5.1 İlgili çalışmaların karşılaştırılması.....	52



Şekiller Listesi

Şekil 4.1 Verilerin duygu durumuna göre dağılımı.....	27
Şekil 4.2 Veri setindeki yorumlar ve puanlara ilişkin dağılım.....	28
Şekil 4.3. Puanlar ve müşteri duygu durumları.....	28
Şekil 4.4 Veri setine ait kelime bulutu	29
Şekil 4.5 Yorum uzunluğu dağılımı	30
Şekil 4.6 Bigrams grafiği.....	31
Şekil 4.7 Trigrams grafiği	32
Şekil 4.8 Uçuş deneyimi için duygu dağılımı	33
Şekil 4.9 Yemek servisi için duygu dağılımı	34
Şekil 4.10 Müşteri hizmetleri için duygu dağılımı.....	35
Şekil 4.11 Konu başlıklarının keşfedilmesi.....	36
Şekil 4.12 Önerilen duygu analizi modeli	37
Şekil 4.13 Modele ait doğruluk ve kayıp grafikleri.....	39
Şekil 4.14 Hata matrisi	40
Şekil 4.15 Erken durdurma yaklaşımı ile eğitim.....	40
Şekil 4.16 Erken durdurma yaklaşımı ile elde edilen hata matrisi	41
Şekil 4.17 İkili sınıflandırma için doğruluk ve kayıp eğrileri	42
Şekil 4.18 İkili sınıflandırma için hata matrisi	43
Şekil 4.19 ROC eğrisi.....	44
Şekil 4.20 Deneysel çalışmalara ait karşılaştırmalı sonuçlar	46
Şekil 4.21 nnlm-en-dim50 modeline ilişkin doğruluk ve kayıp eğrileri	47
Şekil 4.22 nnlm-en-dim50 hata matrisi	47
Şekil 4.23 nnlm-en-dim50 ROC eğrisi.....	48
Şekil 4.24 nnlm-en-dim128 modeline ait doğruluk ve kayıp eğrileri	48
Şekil 4.25 nnlm-en-dim128 modeline ait hata matrisi	49
Şekil 4.26 nnlm-en-dim128 modeline ait ROC eğrisi.....	49
Şekil 4.27 USE modele ait doğruluk ve kayıp eğrileri.....	50
Şekil 4.28 USE modele ait hata matrisi.....	50
Şekil 4.29 USE modele ait ROC eğrisi	51

Kısaltmalar Dizini

ANN	Artificial Neural Network (Yapay Sinir Ağları)
API	Application Programming Interface (Uygulama Programlama Arayüzü)
ASP	Active Server Pages (Etkin Sunucu Sayfaları)
BERT	Bidirectional Encoder Representations from Transformers (İki Yönlü Kodlayıcı Temsillerden Transformatör)
Bi-LSTM	Bidirectional Long Short-Term Memory Networks (Çift Yönlü Uzun-Kısa Vadeli Bellek)
CNN	Convolutional Neural Networks (Evrişimli Sinir Ağları)
CRFA	Context-Relevant Features
CRNN	Convolutional Recurrent Neural Network (Evrişimli Tekrarlayan Sinir Ağları)
CSV	comma separated values (virgülle ayrılmış değerler)
DK	Durdurma Kelimeleri
DNN	Deep Neural Network (Derin Sinir Ağları)
DT	Decision Forest (Karar ağacı)
EDI	Electronic Data Interchange (Elektronik Veri Değişimi)
EFT	Elektronik Fon Transferi
ETKK	Elektronik Ticaret Kordinasyon Kurulu
GNN	Graph neural network (Grafik Sinir Ağları)
HTML	HyperText Markup Language (Hiper Metin İşaretleme Dili)
IMDB	Internet Movie Database (İnternet Film Veritabanı)
ISO	International Organization for Standardization (Uluslararası Standartlar Teşkilatı)
Knn	k-Nearest Neighbors (K En Yakın Komşu)
LCM	Label Confusion Model (Etiket Karmaşası Modeli)
LR	Logistic regression (Lojistik Regresyon)
LSTM	Long Short-Term Memory (Uzun-Kısa Vadeli Bellek)
LSTM	Long-Short-Term Memory (Uzun Kısa Süreli Bellek)
MLP	Multilayer Perceptron (Çok Katmanlı Algılayıcı)
MVC	Model-View-Controller (Model - Görünüm - Denetleyici)
MY-15130	Türkçe Müşteri Yorumları Veri Seti
NB	Naive Bayes
NER	Named Entity Recognition (Adlandırılmış Varlık Tanımı)
NET	Network Enabled Technologies (Ağa Bağlı Teknolojiler)
NLP	Natural Language Processing (Doğal Dil İşleme)
PCA	Principal Component Analysis (Temel Bileşen Analizi)
R8	Reuters- 21578 Dataset
RCNN	Recurrent Convolutional Neural Networks (Tekrarlayan Evrişimli Sinir Ağları)
RF	Random Forest (Rastgele Ağaç)
RNN	Recurrent Neural Network (Tekrarlayan Sinir Ağları)
SGD	Stochastic Gradient Descent (Stokastik Gradyan İnişi)
SMO	Sequential Minimal Optimization (Ardışık Minimal Optimizasyon)
SST	Stanford Sentiment Treebank Dataset

SVM	Support Vector Machine (Destek Vektör Makinesi)
TF-IDF	Terim Frekansı-Ters Doküman Frekansı
TTC-3600	Türkçe Metin Sınıflandırma Veri Seti
TTC-4900	Türkçe Haber Metinleri Veri Seti
Web	World Wide Web (Dünya Çapında Ağ)



BİRİNCİ BÖLÜM

1. Giriş

Yazının icadı dünya tarihinde bilimin sanatın ve insanlığın gelişiminde önemli katkısı olmuştur. Önceleri duvarlara, kil tabletlere yazılan yazılar kâğıdın icadıyla kâğıda yazılarak saklanmıştır. Bilgisayar çağına geldiğimizde yazılar artık dijitalleşerek depolanmaya başlamıştır. Teknolojik gelişmelerle birlikte veri saklama kapasitesi veriye ulaşma hızı da artmaktadır. Veri tabanında saklanan bu veri bir anlam ifade etmez. Ancak bu veri dağı, belirli bir amaç doğrultusunda sistematik olarak işlenir ve analiz edilirse, değersiz görülen veri yığnında, amaca yönelik sorulara cevap verebilecek çok değerli bilgilere ulaşılabilir (Özekes 2003).

Metin sınıflandırma, doğal dil işleme (NLP) alanının önemli bir alt dalı olup, çeşitli metinlerin belirli kategorilere ayrılması işlemidir. Günümüzde dijital verinin hızla artmasıyla birlikte, metin sınıflandırma teknikleri büyük bir önem kazanmıştır. Sosyal medya platformlarından haber sitelerine, akademik makalelerden müşteri geri bildirimlerine kadar geniş bir yelpazede kullanılan metin verileri, doğru ve etkili bir şekilde sınıflandırıldığında değerli içgörüler sunmaktadır. Bu bağlamda, metin sınıflandırma yöntemlerinin geliştirilmesi ve iyileştirilmesi, bilgiye erişimi kolaylaştırmak ve karar destek sistemlerini güçlendirmek adına kritik bir rol oynamaktadır.

Günlük hayatta karşılaşılan; metin dosyaları, web sayfalarında yer alan forum verileri, e-posta içerikleri, makaleler, bloglar ve açık uçlu anket cevapları gibi verilere bakıldığında çoğunlukla verilerin yapısal olmayan metinsel veriler olduğu görülmektedir (Göker ve Tekedere 2017). Bu veriler işlenerek anlamlı hale getirilebilir. Metin madenciliği metin verilerinden yapısallaştırılmış veri elde etmeyi amaçlar (Seker 2015). Metin madenciliği çalışmaları metini veri olarak kabul eder. Metin madenciliği; önceden bilinmeyen ve önemli olan bilgilerin keşfedilmesi amacıyla çok sayıda dokümanı analiz eden bir teknolojidir (Karadağ ve Takçı 2010).

1.1. Problem

Günümüzde işletmeler, müşteri geri bildirimlerini anlamak ve hizmet kalitesini artırmak için büyük miktarda metin verisi toplamakta ve analiz etmektedir. Bu verilerin etkili bir şekilde işlenmesi ve anlamlı içgörüler elde edilmesi, işletmelerin rekabet avantajı kazanmasına yardımcı olmaktadır. Soğan mimarisi, yazılım geliştirme süreçlerinde katmanlı ve modüler bir

yaklaşım sunarak, uygulamaların sürdürülebilirliğini ve esnekliğini artırmaktadır. Bu tez çalışması, müşteri geri bildirimlerinin yapay zekâ destekli modellerle analiz edilmesi ve elde edilen içgörülerin soğan mimarisi kullanılarak dağıtılmasını amaçlamaktadır.

Bu çalışma, küçük/orta ölçekli metin verilerinin yapay zekâ ve NLP teknikleri kullanılarak analiz edilmesi ve bu analiz sonuçlarının soğan mimarisi kullanılarak etkin bir şekilde dağıtılmasını hedeflemektedir. Müşteri geri bildirimlerinden anlamlı içgörüler elde etmek için çeşitli NLP teknikleri ve makine öğrenimi algoritmaları bu çerçevede kullanılır. Elde edilen içgörüler, soğan mimarisi çerçevesinde geliştirilen bir yazılım mimarisi aracılığıyla ilgili paydaşlara sunulur.

1.2. Araştırmanın Amacı

Müşteri Geri Bildirimlerinden Anlamlı İçgörüler Elde Etmek

Bu araştırmanın ilk amacı, müşteri geri bildirimlerinden anlamlı içgörüler elde etmektir. NLP teknikleri kullanılarak, müşteri yorumlarındaki sıkça karşılaşılan problemler ve memnuniyet alanları belirlenir. Bu süreçte bigram ve trigram analizleri, birlikte görünme (co-occurrence) analizi ve duygu analizi gibi yöntemler kullanılarak, müşterilerin hangi konularda yoğunlaştığı ve duygusal tepkileri ortaya konulur. Bu içgörüler, müşteri deneyimlerini anlamak ve hizmet kalitesini iyileştirmek için kritik bilgiler sağlar.

Duygu Analizi ile Müşteri Memnuniyetini Değerlendirmek

İkinci amaç, müşteri yorumlarının duygusal analizini gerçekleştirerek müşteri memnuniyetini değerlendirmektir. Yorumların pozitif, negatif veya nötr duygularla ilişkilendirilmesi, genel müşteri memnuniyetinin ölçülmesine olanak tanır. Bu analiz, müşteri geri bildirimlerinin yalnızca sayısal verilerle değil, duygusal tepkilerle de değerlendirilmesini sağlar. Sonuçlar, hizmet kalitesindeki güçlü ve zayıf yönlerin belirlenmesine yardımcı olarak, müşteri memnuniyetini artırmak için stratejik iyileştirmeler yapılmasına katkıda bulunur.

Soğan Mimarisi ile Etkili Veri Dağıtımını Sağlamak

Üçüncü amaç, elde edilen analiz sonuçlarının soğan mimarisi kullanılarak etkin bir şekilde dağıtılmasını sağlamaktır. Soğan mimarisi, yazılım geliştirme süreçlerinde katmanlı ve modüler bir yaklaşım sunarak, uygulamaların sürdürülebilirliğini ve esnekliğini artırır. Bu araştırmada, müşteri geri bildirimlerinden elde edilen içgörülerin ilgili paydaşlara hızlı ve etkili bir şekilde ulaştırılması hedeflenmektedir. Bu amaç doğrultusunda, analiz sonuçlarının farklı

kullanıcı gruplarına uygun formatlarda sunulması ve veri dağıtımının optimize edilmesi planlanmaktadır.

Hizmet Kalitesini İyileştirmek için Karar Destek Sistemleri Geliştirmek

Dördüncü amaç, müşteri geri bildirimlerinden elde edilen verilerle hizmet kalitesini iyileştirmek için karar destek sistemleri geliştirmektir. Yapay zekâ ve makine öğrenimi algoritmaları kullanılarak, müşteri memnuniyetini artırmak için stratejik önerilerde bulunacak sistemler tasarlanacaktır. Bu sistemler, müşteri geri bildirimlerini analiz ederek, hizmet süreçlerinde iyileştirilmesi gereken alanları belirler ve yöneticilere bu konularda bilgi sağlar. Bu sayede, müşteri deneyimini sürekli olarak izlemek ve iyileştirmek mümkün olacaktır.

1.3. Araştırmanın Önemi

Bu araştırma, müşteri geri bildirimlerinin yapay zekâ destekli modellerle analiz edilmesi ve elde edilen içgörülerin soğan mimarisi kullanılarak etkin bir şekilde dağıtılmasını sağlayarak, hizmet kalitesini artırmak için önemli katkılar sunar. Orta/küçük ölçekli metin verilerinin doğru analiz edilmesi, işletmelerin müşteri memnuniyetini artırmak için stratejik kararlar almasını kolaylaştırır. Soğan mimarisi ile geliştirilen yazılım çözümleri, esnek ve sürdürülebilir yapılarıyla veri dağıtımını optimize eder. Bu araştırma, havayolu taşımacılığı sektöründe müşteri memnuniyetini ve hizmet kalitesini artırmak için yenilikçi yaklaşımlar sunarak, literatüre değerli bilgiler kazandırır.

1.4. Sayıtlar ve Araştırma Soruları

Tez çalışması kapsamında ortaya koyulan araştırma konuları aşağıda sıralanmıştır:

- **Müşteri yorumlarının NLP teknikleri ile analizi doğru ve güvenilir sonuçlar verir:** Bu sayıtlı, doğal dil işleme tekniklerinin müşteri geri bildirimlerinde sıkça bahsedilen konuları ve duyguları doğru bir şekilde tespit edebileceğini varsayar.
- **Duygu analizi, müşteri memnuniyetini ölçmek için etkili bir yöntemdir:** Bu sayıtlı, müşteri yorumlarının pozitif, negatif ve nötr olarak sınıflandırılmasının, genel hizmet kalitesini değerlendirmede etkili bir yöntem olduğunu varsayar.
- **Soğan mimarisi, analiz sonuçlarının hızlı ve etkin bir şekilde dağıtılmasını sağlar:** Bu sayıtlı, soğan mimarisinin katmanlı yapısının, müşteri geri bildirimlerinden elde edilen içgörülerin ilgili paydaşlara hızlı ve etkili bir şekilde ulaştırılmasını kolaylaştıracağını varsayar.
- **Karar destek sistemleri, hizmet kalitesini iyileştirmek için stratejik öneriler sunabilir:** Bu sayıtlı, yapay zekâ ve makine öğrenimi algoritmalarının kullanıldığını

karar destek sistemlerinin, müşteri memnuniyetini artırmak ve hizmet kalitesini iyileştirmek için değerli stratejik önerilerde bulunabileceğini varsayar.

1.5. Sınırlılıklar

Bu tez çalışmasını kısıtlayan unsurlar aşağıda kısaca açıklanmıştır:

1. Veri Kalitesi ve Erişimi:

Bu tez çalışmasında kullanılan müşteri yorumları veri seti, sadece belirli bir havayolu şirketine ait olup, veri kalitesi ve kapsamı sınırlı olabilir. Yorumların doğruluğu, detay seviyesi ve temsiliyeti, analiz sonuçlarını etkileyebilir. Ayrıca, veri setinin sadece belirli bir zaman dilimini kapsaması, uzun vadeli trendleri ve değişimleri incelemeyi zorlaştırabilir.

2. NLP ve Makine Öğrenimi Modellerinin Sınırlılıkları:

NLP ve makine öğrenimi teknikleri, metin analizinde güçlü araçlar olsa da, dilin karmaşıklığı ve bağlamsal anlamları tam olarak yakalayamayabilirler. Özellikle, ironik, alaycı veya çok anlamlı ifadelerin doğru bir şekilde analiz edilmesi zor olabilir. Kullanılan modellerin doğruluğu ve genel performansı, eğitim verilerinin kalitesi ve miktarına bağlıdır.

3. Genelleştirme Problemleri:

Bu çalışma, Air Canada müşterilerinin geri bildirimlerini temel alarak yapılmış olup, elde edilen sonuçlar diğer havayolu şirketleri veya farklı sektörler için doğrudan genelleştirilemez. Farklı şirketlerin ve sektörlerin müşteri profilleri, beklentileri ve geri bildirimleri farklılık gösterebilir.

4. Yazılım Mimarisi ve Uygulama Sınırlılıkları:

Soğan mimarisi, esnek ve modüler yapısıyla birçok avantaj sunsa da, uygulama sürecinde teknik sınırlamalar ve zorluklarla karşılaşılabilir. Geliştirilen yazılım mimarisinin performansı, ölçeklenebilirliği ve entegrasyon yetenekleri, sistemin genel başarısını etkileyebilir. Ayrıca, uygulamanın gerçek dünya kullanımında karşılaşılabilecek operasyonel ve teknik problemler, teorik çerçevenin ötesinde sınırlılıklar oluşturabilir.

5. Etik ve Gizlilik Konuları:

- Müşteri yorumlarının analizi sırasında, kişisel verilerin gizliliği ve etik kullanım ilkeleri göz önünde bulundurulmalıdır. Veri anonimleştirme ve güvenlik önlemleri alınsa bile, müşteri bilgilerinin korunması ve etik kurallara uyulması gereklidir. Bu sınırlılıklar, verinin toplanması, işlenmesi ve sonuçların paylaşılması aşamalarında dikkatle yönetilmelidir.

Bu sınırlılıklar, tez çalışmasının kapsamını ve sonuçlarının geçerliliğini etkileyebilir. Bu sınırlılıkları göz önünde bulundurarak, elde edilen bulguların dikkatli bir şekilde değerlendirilmesi ve gelecekteki çalışmalar için öneriler geliştirilmesi önemlidir.



İKİNCİ BÖLÜM

2. Literatür Taraması

İnternet kullanımının artmasıyla birlikte, çeşitli kaynaklardan büyük miktarda metin verisi üretilmektedir. Bu metin verileri yapılandırılmamış olup, doğal dil yapıları içerir, bu da verilerden anlamlı bilgiler çıkarmayı zorlaştırır. Bu durum, NLP alanında duygu sınıflandırması ve doğal dil çıkarımı için derin öğrenme kullanımına yönelik araştırmaları artırmıştır (Jang vd. 2020).

Günümüz bilgi çağında, internetin en büyük bilgi kaynağı olarak kabul edilmesiyle, elektronik ortamdaki metinlerin artması metin madenciliği ve makine öğrenimi konularını önemli hale getirmiştir. Bu çalışma, farklı makine öğrenmesi yöntemleri kullanarak Türkçe haber metinlerini sınıflandırmaktadır. Destek Vektör Sınıflandırıcısı (SVM), Rastgele Orman (RF) ve Naive Bayes (NB) Sınıflandırıcısı gibi yöntemler kullanılarak haber metinleri analiz edilmiş ve en başarılı performansın %91 doğruluk oranı ile NB sınıflandırıcısı tarafından elde edildiği görülmüştür (Uslu ve Akyol 2021).

Tuzcu, çevrimiçi kitap satış sitesinin kullanıcı yorumları üzerinde duygu analizi yapmıştır. Çalışma kapsamında, Python programlama dili kullanılarak Çok Katmanlı Algılayıcı (MLP) algoritması uygulamıştır. Ayrıca, RapidMiner yazılımı kullanılarak NB, SVM ve Lojistik Regresyon (LR) algoritmaları uygulamıştır. Kitap satış web sitesinden HTML çözümleme yöntemi ile alınan 91309 okuyucu yorumu bu veri setinden alınarak 1400 adet 700 pozitif 700 negatif yorum ile bir veri kümesi oluşturmuştur. Bu çalışma sonucunda MLP %89, NB %77,57, SVM %80,93 ve LR 84,07 doğruluk oranı elde edilmiştir. Kullanıcı yorumlarının daha doğru analiz edilmesi ve duygu analizi çalışmalarının etkinliğinin artırılması için demografik bilgiler ve daha büyük veri setleri kullanılması gerektiği önerilmiştir (Tuzcu 2020).

Acı ve Çırak, konvolüsyonel sinir ağları (CNN) ve Word2Vec metodu kullanılarak Turkish Text Classification 3600 (TTC-3600) veri kümesi üzerinde metin sınıflandırma çalışması yapmış ve aynı veri kümesini kullanarak daha önce yapılan çalışmalar ile kıyaslamışlardır. Her iki CNN modeli de Python dili ve Tensorflow arayüzünün Keras Kütüphanesi kullanılarak kodlanmıştır. Deney sonuçlarına göre, birinci CNN modeli gövdelenmiş veride %93,3, ham veride ise %90,1 doğruluk değeri elde etmiştir. İkinci CNN modeli ise gövdelenmiş veri üzerinde %90,8, ham veri üzerinde %90,2 doğruluk performansı

göstermiştir. Bu çalışma ile CNN ve Word2Vec yöntemlerinin Türkçe haber metinlerini sınıflandırmada klasik makine öğrenmesi yöntemlerine göre daha yüksek performans sağladığını tespit etmişlerdir (Acı ve Çırak 2019).

Demircan ve arkadaşları sosyal medya ve e-ticaret sitelerinde yayınlanan metinler üzerinden duygu analizini makine öğrenme yöntemleri ile gerçekleştirmek amacıyla hepsiburada.com e-ticaret sitesinden alınan ürün incelemeleri ve puanları kullanılarak Türkçe duygu analizi modelleri geliştirmişlerdir. İncelemeler, pozitif, negatif ve nötr olarak sınıflandırılmış ve SVM, RF, karar ağaçları (DT), LR ve K en yakın komşu (Knn) algoritmaları kullanılmıştır. Yapılan çalışmalar sonucunda SVM ve RF tabanlı modellerin pozitif, negatif ve nötr yorum sınıflamada en iyi performansı gösterdiği tespit edilmiştir (Demircan vd. 2021).

Chen ve arkadaşları yaptıkları çalışmada ABD yüksek mahkeme belgelerinin otomatik olarak sınıflandırılması için kullanılan makine öğrenmesi ve derin öğrenme yöntemlerini karşılaştırmışlardır. 30.000 adet ABD dava belgesi üzerinde yapılan deneyler, anahtar kelimeler ve temel bileşen analizi (PCA) kullanarak çıkarılan alan kavramlarına RF sınıflandırıcısının, derin öğrenme modelleri (Word2Vec, GloVe, BERT ile TextCNN, BiLSTM ve BiLSTM-Attention) karşısında üstün performans gösterdiğini ortaya koymuştur. RF sınıflandırıcısı, doğruluk, geri çağırma, kesinlik ve F1 skoru açısından daha iyi sonuçlar elde etmiştir. Bu çalışmada yasal metin sınıflandırma sistemleri için en uygun makine öğrenmesi stratejilerini seçmede rehberlik etmeyi amaçlamışlardır (Chen vd. 2022).

Barfar yaptığı çalışmada çevrimiçi propagandanın tespiti ve açıklaması için dilbilimsel ve oyun teorisine dayalı bir yaklaşım geliştirmiştir. ABD'deki 39 propagandacı ve 30 güvenilir haber kaynağından toplanan yaklaşık 205.000 makale ile oluşturulan veri seti, propaganda içeriğini dilsel özellikler kullanarak puanlayan modellerin geliştirilmesinde kullanılmıştır. Çalışmada LightGBM algoritması kullanılarak oluşturulan model, C-indeksi 0.9 ve F1 skoru 0.84 ile üstün performans göstermiştir. Bu model ile makalelerin propagandist içerik puanlarını hesaplamak için dilsel özelliklerin katkısını açıklayarak kullanıcı güvenini artırmayı amaçlamıştır. Çalışma, gelecekte anti-aşı propagandası ve diğer bilgi kirliliği türlerinin tespiti ve açıklanmasında kullanılabilecek bir yaklaşım sunmaktadır (Barfar 2022).

Kısa metinler yeterli bağlamsal bilgi içermedikleri ve çok anlamlı kelimeler barındırdıkları için sınıflandırmada zorluk yaratmaktadır. Liu ve arkadaşları bağlama dayalı özellikleri çok aşamalı dikkat ağı ile birleştiren yöntem Context-Relevant Features (CRFA)

adını verdikleri yöntemini önermişlerdir. Yaptıkları çalışmalar sonucunda bu yöntemin CNN ve tekrarlayan sinir ağıları (RNN) tabanlı kısa metin sınıflandırma yaklaşımlarına kıyasla daha yüksek doğruluk ve verimlik sağladığını tespit etmişlerdir (Liu, Li, ve Hu 2022).

Ruiz ve Bedmar ilaç incelemelerinde duygu analizi için derin öğrenme mimarilerini karşılaştırmışlardır. Yaptıkları çalışmada drugs.com üzerinden aldıkları 215,063 ilaç incelemesini kullanmışlardır. Derin öğrenme modelleri Basit CNN, Bi-LSTM, CNN ve LSTM kombinasyonları ve BERT ile Bi-LSTM kullanmışlardır. Yaptıkları çalışmalar sonucunda duygu analizi için bu modeller karşılaştırılmış BERT tabanlı modellerin en iyi performans sağladığı tespit edilmiştir. Fakat bu modellerin eğitim süresi diğerlerine göre daha uzun olduğunu gözlemlemişlerdir (Colón-Ruiz ve Segura-Bedmar 2020).

Metin sınıflandırmada derin öğrenme yöntemleri kullanılarak hibrit modeller geliştirilmektedir. Wang ve arkadaşları derin öğrenme yöntemlerinden CNN veya RNN kullanarak hibrit bir model olan Konvolüsyonel Tekrarlayan Sinir Ağı (CRNN) modelini önermişlerdir. Yaptıkları çalışmalarda iki adet Çince veri seti ve beş adet İngilizce veri seti üzerinde çalışmalar yapmışlardır. Önerdikleri bu yöntemi fastText ve HAN gibi diğer sınıflandırma yöntemleri ile karşılaştırmışlardır. CRNN modeli ile hem İngilizce hem de Çince metin verileri üzerinde daha yüksek sınıflandırma doğruluğu elde etmişlerdir (Wang vd. 2019).

Metin madenciliği ve doğal dil işleme alanlarında anahtar kelime çıkarma önemli bir araştırma alanıdır. Onan ve arkadaşları bilimsel metin belge sınıflandırmasında en sık kullanılan anahtar kelime çıkarma yöntemlerinin (ölçüye dayalı anahtar kelime çıkarımı, terim frekansı-ters cümle frekansı tabanlı anahtar kelime çıkarımı, eş-oluşum istatistiksel bilgi tabanlı anahtar kelime çıkarımı, eksantriklik tabanlı anahtar kelime çıkarımı ve TextRank algoritması) metin sınıflandırma ve topluluk öğrenmesi yöntemleri üzerindeki tahmin performansı incelemişlerdir. Çalışmada, Boosting, Bagging, Dagging, Random Subspace ve Voting gibi beş popüler topluluk yöntemi ve NB, SVM, LR ve RF temel öğrenme algoritmaları olarak kullanılmışlardır. Bu çalışmada veri seti olarak Reuters-21578 veri seti kullanılmıştır. Araştırma sonucunda en sık kullanılan anahtar kelime çıkarma yöntemleri ile birlikte Bagging topluluk öğrenimi yöntemi ve RF algoritması birlikte kullanıldığında en yüksek ortalama tahmin performansı (%93,80) sağlanmıştır (Onan, Korukoğlu, ve Bulut 2016).

Göker ve Tekdere FATİH projesi ile ilgili yapılan yorumların otomatik değerlendirisi ile ilgili bir çalışma yapmışlardır. İnternet ortamında bu proje ilgili yapılan 444 yorum toplanmış ve bu yorumlar pozitif ve negatif olarak gruplandırmışlardır. NB, Knn, Ardışıl Minimal Optimizasyon (SMO) ve RBF Network algoritmalarını uygulamışlardır. Çalışma sonucunda doğruluk oranını en yüksek SMO algoritması olduğunu gözlemlemişlerdir (Göker ve Tekdere 2017).

Kılınç ve arkadaşları metin madenciliği tekniklerini kullanarak akademik makalelerin sınıflandırılmasını incelemişlerdir. Yaptıkları çalışmada <https://www.researchgate.net> web sitesi üzerinden, geliştirilen yazılım aracı kullanılarak belirlenen 50 farklı dergiden 2000 adet makaleye ait özetleri ile oluşturmuşlardır. Oluşturdukları veri setinde “Materials Science & Engineering” ve “Social Sciences & Humanities” olarak iki adet sınıf bulunmaktadır. Bu veri setinin %70’ini eğitim %30’unu test için kullanmışlardır. R studio kullanarak yapılan deneysel araştırmalar sonucunda en iyi doğruluk oranı %96.67 ile Knn algoritması olduğunu gözlemlemişlerdir (Kılınç vd. 2016).

Bilimsel araştırmacılar her yıl büyük miktarda belge üretmektedir. Belge sayısının sürekli artmasıyla geleneksel gözetimli öğrenme teknikleri performans kaybı yaşamaktadır. Kowsari ve arkadaşları belgeleri hiyerarşik bir yapıda sınıflandıran HDLTex yöntemini geliştirmişlerdir. Bu yöntemde derin sinir ağları (DNN), RNN ve CNN kullanılmışlardır. Web of Science’den 7 ana kategoriden elde edilen 46985 makale veri seti olarak kullanılmıştır. Derin öğrenme yöntemlerinin, belge sınıflandırmada geleneksel yöntemlere kıyasla önemli iyileştirmeler sağlayabileceğini gözlemlemişlerdir. Özellikle hiyerarşik sınıflandırma için derin öğrenme mimarilerinin kullanımı, genel performansı artırmakta ve daha esnek sınıflandırma olanakları sunduğunu tespit etmişlerdir (Kowsari vd. 2017).

Alparslan ve Dursun makine öğrenme teknikleri ve CNN kullanarak Türkçe metinler üzerine sınıflandırma çalışması yapmıştır. Yaptıkları bu çalışmada TTC-4900 (Türkçe haber metinleri) ve MY-15130 (Türkçe müşteri yorumları) veri setlerini kullanmışlardır. Türkçe metin sınıflandırma üzerine yapılan bu çalışmada CNN tabanlı derin öğrenme modelinin NB, Knn, SVM ve RF algoritmaları ile yapılan sınıflandırmaya göre daha başarılı olduğunu gözlemlemişlerdir (Alparslan ve Dursun 2023).

Guo ve arkadaşları metin sınıflandırmada etiketleme ile ilgili yeni bir model önermişlerdir. Etiket Karmaşası Modeli (LCM) adını verdikleri bu modelde daha iyi bir etiket

dağılımı oluşturarak sınıflandırma performansını artırmayı hedeflemişlerdir. Yöntemi değerlendirmek için üç adet İngilizce ve iki adet Çince olmak üzere beş farklı veri seti kullanmışlardır. Yaptıkları çalışma sonucunda LCM'nin mevcut sınıflandırma modellerinin performansını artırdığını gözlemlemişlerdir. Özellikle karışık veri setlerinde önerdikleri bu modelin daha etkili olduğunu tespit etmişlerdir (Guo vd. 2021)

Makine öğrenimi tabanlı metin sınıflandırma algoritmalarını Türkçe dil bağlamında incelemek üzere Alzoubi ve arkadaşları bir çalışma yapmıştır. Bir şirketten aldıkları 225,239 müşteri taleplerini veri seti olarak kullanmışlardır. Veriyi hazırlık SVM, NB, Uzun-Kısa Süreli Bellek (LTSM), RF ve LR algoritmaları kullanılarak sınıflandırmışlardır. Yaptıkları çalışma sonucunda en yüksek doğruluk oranını %84 ile LTSM algoritmasında elde etmişlerdir. Ayrıca verileri hazırlıklı ve hazırlıksız olarak ayrı ayrı sınıflandırmış hazırlıklı verilerle daha iyi sonuçlar elde etmişlerdir (Alzoubi, Topcu, ve Erkaya 2023).

Kaşıkcı ve Gökçen kullanıcıların e-ticaret sitelerini bulmalarını kolaylaştırmak amacıyla metin madenciliği yöntemlerini kullanarak bir uygulama geliştirmişlerdir. İnternette yer alan dizin sitelerinden alışveriş kategorisinden topladıkları siteleri inceleyerek “e-ticaret sitesi” ve “e-ticaret sitesi değil” şeklinde iki sınıf oluşturarak gruplandırmışlardır. Bu veriler üzerinde Knn ve NB sınıflandırma algoritmaları kullanılmışlardır. NB algoritmasında %85,30 Knn algoritmasında ise %83,87 doğruluk oranı elde etmişlerdir (Kaşıkcı ve Gökçen 2013).

Sosyal medyanın kullanımının artmasıyla birlikte kullanıcıların oluşturdukları içeriklerde rahatsız edici, saldırgan içeriklerde de artış olmaktadır. Sosyal medyanın saldırgan dilini tespit etmek amacıyla Hajibabae ve arkadaşları bir çalışma yapmıştır. Twitter API üzerinden elde ettikleri 24783 tweeti nefret söylemi, saldırgan dil ve hiçbirisi olarak sınıflandırmışlardır. En çok kullanılan kelime gömme yöntemleri olan TF-IDF, Word2Vec ve FastText kullanmışlardır. Sınıflandırıcı olarak Gaussian Naive Bayes, DT, LR, RF, AdaBoost, SVM, Gradient Boosting ve MLP algoritmalarını kullanmışlardır. TF-IDF gömme yöntemi ile AdaBoost, SVM ve MLP sınıflandırıcılarında %95 ile en yüksek F1 skoru elde etmişlerdir. En düşük performansı NB sınıflandırıcısında elde etmişlerdir (Hajibabae vd. 2022).

Gürbüz ve arkadaşları oteller için paylaşılan çevre ile alakalı yorumları analiz etmek amacıyla metin madenciliğini kullanmışlardır. Antalya bölgesi otelleri ile ilgili internet üzerinden veri kazıma yöntemi 211.790 adet yorumu toplamışlardır. Topladıkları bu verileri ön işledikten sonra çevreci ve çevreci olmayan yorumlar olarak sınıflandırmışlardır. Yaptıkları

alışma sonucunda %94,59 ile en yksek bařarı oranına DT sınıflandırma algoritmasında ulaşmışlardır. (Grbz vd. 2024)

Neogi ve arkadaşları sosyal medya zerinden paylaşılan kamuoyu duygularını analiz etmek amacıyla bir alışma yapmışlardır. Hindistan'da yapılan ifti protestolarını ve yansımalarını incelemek amacıyla twitter zerinden 9061 benzersiz kullanıcıdan 18bin tweet toplamışlardır. TextBlob ktphanesi kullanarak tweetlerin pozitif, negatif ve ntr duygu durumlarını belirlemişlerdir. Bag of Words ve TF-IDF teknikleri kullanılarak tweet'ler sayısal verilere dnřtrlmř ve NB, DT, RF, SVM ile sınıflandırma yapılmıştır. Yaptıkları alışma sonucunda Bag of Words tekniğinin TF-IDF'ye gre daha bařarılı olduėunu gzlemlemişlerdir. RF algoritması %96,6 doėruluk oranı ile diėer sınıflandırma yntemlerine gre daha bařarılı olduėunu tespit etmişlerdir (Neogi vd. 2021).

elik ve Ko Trke haber sitelerinden elde edilen verileri sınıflandırma amacıyla bir alışma yapmıştır. 6 farklı haber sitesinden bilim ve teknoloji, eėitim, kltr ve sanat, saėlık, siyaset ve spor kategorilerinden her bir kategoride 2000 adet haber metni olmak zere 12000 adet haber metnini alarak bir veri seti oluřturmuşlardır. n iřlemeden geirip vektr modeli ıkarılan veri seti %70 eėitim %30 test dokmanı olarak ayırmışlardır. TF-IDF, Word2Vec ve FastText yntemleri ile elde ettikleri vektr modellerini SVM, NB, LR, RF ve Yapay Sinir Aėları (ANN) ile sınıflandırmışlardır. alışma sonucunda FastText vektr modeli ve SVM sınıflandırma yntemiyle %95.75 en yksek bařarı oranı izlendiėi grlmřtr (elik ve Ko 2021).

Soni ve arkadaşları ikili ve ok sınıflı metin sınıflandırma problemleri iin TextConvoNet adını verdikleri yeni bir CNN tabanlı model nermişlerdir. nerdikleri model iki boyutlu konvlsyonel filtre kullanarak cmle ii ve cmleler arası n-gram zellikleri ıkarmaktadır. İkili ve ok sınıflı veriler bulunan Binary Stanford Sentiment Treenbank (SST), Amazon Review, Reuters- 21578 (R8), Twitter User Airline Sentiment ve Covid Tweets veri setlerini kullanmışlardır. Makine ėrenimi yntemleri olarak NB, DT, RF, SVM ve ayrıca derin ėrenme tabanlı LSTM, VDCNN, BiLSTM+Attention, BERT, HAN, BerConvoNet ve CNN-BiLSTM modelleriyle karřılařtırmışlardır. Yaptıkları deneyler sonucunda nerdikleri modelin diėerlere kıyasla daha stn bir model gsterdiėini gzlemlemişlerdir (Soni, Chouhan, ve Rathore 2023).

Jang ve arkadaşları LSTM ve CNN'nin avantajlarını birleřtiren ve ek olarak dikkat mekanizması kullanan bir Bi-LSTM+CNN hibrit modeli nermişlerdir. Yaptıkları bu

çalışmada IMDB film yorumları veri setini kullanmışlardır. Bu veri setinde filmler hakkında pozitif ve negatif etiketli toplam 50000 veri bulunmaktadır. Bu çalışmada yaptıkları deneyler sonucunda önerilen modelin mevcut modellerden daha yüksek doğruluk ve F1 skoruna sahip olduğunu gözlemlemişlerdir (Jang vd. 2020).

Lai ve arkadaşları geleneksel makine öğrenme yöntemleri ile metin sınıflandırma yerine RCNN ile bir model önermişlerdir. Bu modelde kelime temsillerini öğrenirken bağlamsal yapıyı yakalamak için tekrarlayan bir yapı uygulamış ve ayrıca metin sınıflandırmadaki önemli kelimeleri belirlemek için maksimum havuzlama katmanını kullanmışlardır. Yaptıkları bu deneyde 20Newsgroups (İngilizce) Fudan Set (Çince) , Anthology Network (ACL) (İngilizce) , SST (İngilizce) olmak üzere 4 ayrı veri seti üzerinde çalışmışlardır. Yaptıkları deneylerde geleneksel yöntemler LR ve SVM ile kıyaslamışlardır. Yapılan deneyler sonucunda 20Newsgroups veri setinde %33 ve Fudan set veri setinde %19 hata oranının düştüğünü gözlemlemişlerdir (Lai vd. 2015).

Dijital belgelerin büyümesiyle birlikte metin sınıflandırma tekniklerinde de hızlı bir ilerleme olmaktadır. Gasparetto ve arkadaşları geleneksel makine öğrenme yöntemlerini ve derin öğrenme yöntemlerini incelemek amacıyla bir çalışma yapmıştır. EnWiki-100 ve RCV1-57 veri setlerini kullanarak NB, SVM, FstText, BiLSTM, XML-CNN, BERT ve XML-R modellerinde deney yapmışlardır. Transformer tabanlı dil modelleri, tüm metin sınıflandırması görevlerinde üstün performans sergilemiştir. BERT ve XLNet'in çoğu veri setinde en iyi performansları elde ettiği tespit etmişlerdir. (Gasparetto vd. 2022).

Farklı dillerde metin sınıflandırma ile ilgili çeşitli çalışmalar yapılmaktadır. Chen ve arkadaşları Lao dilindeki metinleri Knn ile sınıflandırmak amacıyla bir çalışma yapmıştır. Lao dilinde yazılmış 2411 haber makalesini toplamışlardır. Toplam 7 kategoride bulunan bu haber metinlerini 1490 eğitim 921 test verisi olarak ayırmışlardır. K 4 olarak seçilmiş ve en iyi sonuçları bu değerde elde etmişlerdir. Deney sonucunda %71.4 doğruluk oranı elde etmişlerdir. Deney sonucunda Knn tabanlı Lao haber metin sınıflandırma yönteminin veri normalizasyonu ve veri boyut azaltma işlemleri ile iyi performans gösterdiğini tespit etmişlerdir (Chen vd. 2020).

Metin madenciliğinde, özellik seçimi, büyük miktardaki özellik uzayını azaltmak ve sınıflandırma doğruluğunu artırmak için yaygın bir yöntemdir. Bahassine ve arkadaşları Arapça metin sınıflandırması için ki-Kare yöntemini geliştirerek ImpCHI adını verdikleri yeni bir model önermişlerdir. Open source Arabic corpora veri setinden 5070 farklı metin belgesi

kullanmışlardır. Toplam 6 sınıfa ait belgelerin olduğu veri setini 4057 eğitim 1013 test olarak ikiye ayırmışlardır Bilgi kazancı, karşılıklı bilgi, ki-Kare ve önermiş oldukları ImpCHI yöntemlerini DT ve SVM sınıflandırma kullanarak bir deney yapmışlardır. Yaptıkları deney sonucunda ImpCHI yöntemi ve SVM sınıflandırıcısının birlikte kullanılmasıyla doğruluk, hatırlama ve f-ölçütlerinde diğerlerine göre daha iyi sonuçlar elde etmişlerdir. Bu model için elde edilen en iyi f-ölçütü değeri, 900 özellik sayısında %90.50 olarak gözlemlemişlerdir (Bahassine vd. 2020).

Reusens ve arkadaşları metin sınıflandırması için beş farklı görevi (sahte haber sınıflandırması, konu tespiti, duygu tespiti, kutupluluk tespiti ve alay tespiti) kapsayan kapsamlı bir kıyaslama çalışması yapmışlardır. Bu çalışmada 20 veri seti ve 11 model mimarisini içeren bu çalışma, 42.800 algoritma çalıştırmasını değerlendirerek, basit yöntemlerin belirli durumlarda karmaşık modellerle rekabet edebileceğini gözlemlemişlerdir. Sonuçlar, özellikle sahte haber ve konu tespitinde, basit modellerin karmaşık sinir ağlarına göre daha iyi performans gösterebildiğini ortaya koymaktadır. Ayrıca, küçük veri setlerinde model karmaşıklığı ile performans arasında negatif bir ilişki olduğu bulunludur. Bu çalışmada, metin sınıflandırma alanında daha verimli ve şeffaf yöntemlerin önemini vurgulamışlardır (Reusens vd. 2024)

Mikro blog sitelerinde, yorumlarda ve soru cevap platformlarında kısa metinler sıkça karşılaşılan ve hızla büyüyen bir veri türüdür. Wang ve arkadaşları CNN'nin kısa metin sınıflandırmadaki eksiklerini gidermek için N-gram ve CNN teknolojisini birleştiren yenilikçi bir model önermektedirler. Yaptıkları deneylerde İngilizce ve Çince veri setlerinde, önerilen yöntemin geleneksel makine öğrenme algoritmaları ve diğer CNN modellerine göre daha yüksek doğruluk oranlarına sahip olduğunu göstermektedir (Wang vd. 2020).

Sosyal medyadaki artan zararlı yorumları sınıflandırmak amacıyla Bhattacharya ve arkadaşları LSTM ve CNN mimarilerini birleştiren hibrit bir yöntem önermişlerdir. Önerilen bu yöntemde CNN bileşeni, yorumlardaki yerel kalıpları ve özellikleri tanımlarken, LSTM bileşeni, dilin ardışık akışını anlamaya yardımcı olmaktadır. Veri setinde 2087 zararlı 28404 adet zararsız toplam 30491 yorumdan oluşmaktadır. Veri seti %70 eğitim %20 test ve %10 doğrulama olarak ayrılmıştır. Çalışma sonucunda CNN ve LSTM hibrit modelinde %98,91 doğruluk oranı elde edilmiştir (Bhattacharya vd. 2024).

ÜÇÜNCÜ BÖLÜM

3. Materyal ve Yöntemler

3.1. Veri Seti

Veri seti, hava yolu deneyimlerine ait müşteri yorumlarını içermektedir. Her bir satırda, bir müşterinin doğrulama durumu, yorum içeriği, değerlendirme puanı, önerip önermediği ve duygu durumu gibi bilgiler yer almaktadır. Özellikle "verify" sütunu müşteri yorumunun doğruluğunu, "review" sütunu müşteri yorum metnini, "rating" sütunu müşterinin verdiği puanı, "recommended" sütunu müşterinin havayolunu tavsiye edip etmediğini ve "sentiment" sütunu yorumun duygu analizini göstermektedir. Bu veri seti, müşteri memnuniyeti ve havayolu hizmet kalitesini değerlendirmek için kullanılabilecek zengin bir bilgi kaynağı sunmaktadır.

Bu veri seti, metin tabanlı müşteri yorumlarının duygu analizi için kullanılmaktadır. Duygu analizi, müşterilerin geri bildirimlerinde ifade edilen duygusal tonları belirlemek için metin madenciliği ve doğal dil işleme tekniklerinin kullanılmasıdır. Bu tür analizler, havayolu şirketlerinin müşteri memnuniyetini anlamalarına, hizmet kalitesini değerlendirmelerine ve müşteri deneyimlerini iyileştirmelerine yardımcı olmaktadır.

Veri seti, havayolu sektöründe müşteri deneyimlerinin analizi için önemli bir kaynak sağlamaktadır. Özellikle, metin verilerinin sınıflandırılması ve analiz edilmesi, işletmelerin müşteri memnuniyeti üzerine stratejik kararlar almasına olanak tanır. Duygu analizi, müşteri geri bildirimlerinden anlamlı bilgiler çıkarmayı ve bu bilgileri operasyonel ve stratejik iyileştirmeler için kullanmayı hedefler.

Tablo 3.1 Verilerin duygu durumuna göre dağılımı

Duygu Durumu	Kayıt Sayısı
Olumsuz	705
Olumlu	303
Nötr	92

Tablo 3.1 Verilerin duygu durumuna göre dağılımı incelendiğinde, toplamda 1100 kaydın bulunduğu gözlemlenmiştir. Kayıtların 705'i olumsuz duygu durumunu (veri setindeki sayısal temsil, 0), 303'ü pozitif duygu durumunu (veri setindeki sayısal temsil, 2) ve 92'si nötr duygu durumunu (veri setindeki sayısal temsil, 1) yansıtmaktadır. Bu dağılım, müşteri yorumlarının büyük çoğunluğunun olumsuz veya pozitif olduğunu göstermektedir. Olumsuz duygu durumu en yüksek orana sahipken, pozitif ve nötr yorumlar daha az sıklıkla görülmektedir. Bu bulgu,

havayolu hizmet kalitesinin genel olarak müşteri beklentilerini karşıladığı, ancak iyileştirme alanlarının da mevcut olduğunu ortaya koymaktadır.

İncelenen veri seti, müşteri yorumlarına dayalı duygu durumu analizi içermektedir ve akademik ve teknik araştırmalar için açık erişimlidir. Veri seti, İngilizce dilinde müşteri yorumlarını içermektedir. Bu veri setine, araştırmacıların ve ilgili kullanıcıların erişimine sunulmuş olup, ilgili internet¹ adresinden indirilebilir. Bu açık erişimli kaynak, havayolu hizmet kalitesi ve müşteri memnuniyeti üzerine çalışmalar yapmak isteyenler için değerli bir bilgi kaynağı sunmaktadır.

3.2. Doğal Dil İşleme

NLP, insan dilinin bilgisayarlar tarafından anlaşılması, işlenmesi ve üretilmesini amaçlayan bir yapay zekâ alt dalıdır. NLP, dilbilim, bilgisayar bilimi ve yapay zekâ tekniklerinin birleşimini kullanarak metin ve konuşma verilerini analiz eder. Bu alan, makine çevirisi, duygu analizi, metin sınıflandırma, adlandırılmış varlık tanıma ve dil modelleme gibi çeşitli uygulamaları içerir. NLP sistemleri, büyük dil verilerini işleyerek dil kalıplarını öğrenir ve anlam çıkarır. İstatistiksel ve derin öğrenme yöntemleri, özellikle sinir ağları, NLP'nin gelişiminde kritik rol oynamaktadır. Bu sayede, NLP insan-bilgisayar etkileşiminde devrim yaratmaktadır.

NLP, metinle ilgili birçok görevi kapsayan genel bir terimdir. Metnin büyük/küçük harfini değiştirme (örneğin, büyük harfleri küçük harfe dönüştürme) gibi basit görevlerden, diller arası çeviri ve kelime anlamı ayrımı (aynı yazılışa sahip bir kelimenin bağlama göre anlamını çıkarma) gibi karmaşık görevlere kadar her şey NLP kapsamına girer. NLP alanında karşılaşılan bazı önemli görevler şunlardır:

- **Durdurma Kelimesi Çıkarma**—Durdurma kelimeleri, metin korpuslarında sıkça rastlanan ve bilgi içeriği düşük olan kelimelerdir (örneğin "ve", "bu", "bir", "ben" vb.). Genellikle bu kelimeler, metnin anlamsal içeriğine çok az katkıda bulunur veya hiç bulunmaz. Özellik uzayını azaltmak için, birçok işlemde metin modeliyle çalışmadan önce durdurma kelimeleri erken aşamalarda çıkarılır. Bu, modelin daha az ve daha anlamlı verilerle çalışmasını sağlayarak, performansını ve işlem verimliliğini artırır.

¹ <https://data.mendeley.com/datasets/pc6fxc95h5/1> (Erişim Tarihi: 29.07.2024)

- **Lemmatizasyon**, modelin işlemekte olduğu özellik uzayını azaltmak için kullanılan bir tekniktir. Bu süreç, belirli bir kelimeyi kök formuna dönüştürür (örneğin, "otobüsler" kelimesini "otobüs", "yürüdü" kelimesini "yürü", "gitti" kelimesini "git" olarak dönüştürmek gibi). Lemmatizasyon, dilbilimsel kurallara ve kelimenin bağlamına dayanarak kök formunu belirler. Bu, kelime haznesinin boyutunu küçülterek, modelin öğrenmesi gereken veri boyutunu azaltır ve dolayısıyla hesaplama yükünü hafifletir. Lemmatizasyon, doğal dil işleme uygulamalarında yaygın olarak kullanılan etkili bir ön işleme yöntemidir ve veri setlerinin daha yönetilebilir ve analiz edilebilir hale getirilmesini sağlar.
- **Sözcük Türü (PoS) Etiketleme—PoS etiketleme**, belirli bir metindeki her sözcüğü bir sözcük türü ile etiketleme işlemidir (örneğin, isim, fiil, sıfat vb.). Bu işlem, metindeki kelimelerin dilbilgisel işlevlerini belirlemek amacıyla gerçekleştirilir. Penn Treebank projesi, en popüler ve kapsamlı PoS etiket listelerinden birini sağlar. Bu etiketler, metin analizi ve doğal dil işleme uygulamalarında yaygın olarak kullanılır. Penn Treebank PoS etiket listesinin tam bir versiyonuna erişmek için internet adresini² ziyaret edebilirsiniz. Sözcük türü etiketleme, dilin yapısını anlamak ve doğal dil işleme modellerinin doğruluğunu artırmak için kritik bir adımdır.
- **Adlandırılmış Varlık Tanıma (NER)—NER**, metinden çeşitli varlıkları (örneğin, kişi/şirket isimleri, coğrafi konumlar vb.) çıkarmakla sorumludur. Bu teknik, metin içindeki özel isimleri, yer adlarını ve diğer belirgin varlıkları tanımlayarak yapılandırılmamış metni daha anlamlı ve analiz edilebilir hale getirir. Adlandırılmış varlık tanıma, doğal dil işleme alanında önemli bir yer tutar ve özellikle bilgi çıkarımı, metin sınıflandırma ve arama motoru optimizasyonu gibi uygulamalarda geniş bir kullanım alanı bulur.
- **Dil Modellemesi—Dil modellemesi**, 1'den $w-1$ 'e kadar olan kelimeler verildiğinde n 'inci kelimeyi tahmin etme görevidir. Dil modellemesi, ilgili veriler üzerinde eğitilerek şarkılar, film senaryoları ve hikâyeler üretmek için kullanılabilir. Dil modellemesi için gereken verilerin oldukça erişilebilir olması (yani, etiketlenmiş veriye ihtiyaç duymaması) nedeniyle, genellikle karar destek NLP modellerine dil anlayışı kazandırmak için ön eğitim yöntemi olarak

² https://www.ling.upenn.edu/courses/Fall_2003/ling001/penn_treebank_pos.html (Erişim Tarihi:29.07.2024)

kullanılır. Bu teknik, geniş veri kümeleri üzerinde çalışarak dilin yapısını ve kurallarını öğrenir, böylece çeşitli metin üretim ve analiz uygulamalarında yüksek performans sağlar.

- **Duygu Analizi**—Duygu analizi, verilen bir metnin duygusal içeriğini belirleme görevidir. Örneğin, bir duygu analiz aracı, ürün veya film yorumlarını analiz ederek ürünün ne kadar iyi olduğunu belirten bir puan üretebilir. Bu teknik, metinlerin pozitif, negatif veya nötr duygusal tonlarını tespit ederek, müşteri geri bildirimleri, sosyal medya gönderileri ve diğer metin tabanlı veri kaynakları üzerinde değerli içgörüler sağlar. Duygu analizi, pazarlama stratejileri, müşteri hizmetleri ve halkla ilişkiler gibi alanlarda yaygın olarak kullanılır.

3.3. Veri Ön İşleme

NPL’de, metin ön işleme adımları önem arz eder. İlk olarak, tüm karakterler küçük harfe dönüştürülür (örneğin, "Ahmet" → "ahmet"). Bu işlem, büyük-küçük harf duyarlılığını ortadan kaldırır. Noktalama işaretleri ve rakamlar metinden çıkarılarak, anlamsal olmayan unsurlar temizlenir. Durdurma kelimeleri ("ve", "bu", "bir" gibi) metinden kaldırılarak, veri seti daha bilgi dolu hale getirilir. Son olarak, lemmatizasyon uygulanır; bu işlem, kelimeleri kök formuna dönüştürerek ("yürüyordu" → "yürü"), kelime çeşitliliğini azaltır ve modelin öğrenme sürecini iyileştirir.

3.4. Veri setinin eğitim ve test seti olarak ayrılması

Veri setinin eğitim ve test seti olarak ayrılması, makine öğrenimi modellerinin performansını değerlendirmek için kritik bir adımdır. Eğitim seti, modelin öğrenme süreci boyunca kullanılarak model parametrelerinin ayarlanmasını ve örüntülerin tanınmasını sağlar. Bu süreçte model, eğitim setindeki veriler üzerinden belirli bir görevi gerçekleştirmeyi öğrenir ve bu verilere dayalı olarak kararlar alır. Test seti ise, modelin eğitim süreci tamamlandıktan sonra, genel performansını ve genelleme yeteneğini değerlendirmek amacıyla kullanılır. Eğitim setine hiç maruz kalmayan verilerden oluşan test seti, modelin daha önce görmediği verilerle ne kadar iyi sonuçlar üretebildiğini belirler. Bu çalışmada veri setinin %75’i eğitim, %25’i ise test seti olarak ayrılmıştır. Bu oran, modelin yeterli miktarda veriyle eğitilmesini sağlarken, test seti için de yeterli sayıda örnek bırakarak modelin genelleme yeteneğinin objektif bir şekilde değerlendirilmesine olanak tanır. Eğitim ve test setlerinin bu şekilde ayrılması, farklı modellerin aynı veriler üzerinde performanslarını karşılaştırmayı kolaylaştırır ve modellerin gerçek dünya uygulamalarındaki etkinliğini daha doğru bir şekilde ölçmemizi sağlar.

3.5. Derin Öğrenme

Derin öğrenme, yapay sinir ağlarının birçok katmandan oluşan karmaşık yapılar kullanarak veriden otomatik olarak özellikler çıkarıp öğrenme sürecini optimize eden bir makine öğrenimi dalıdır. Büyük veri kümeleri üzerinde etkili performans gösterir ve özellikle görüntü tanıma, doğal dil işleme, ses tanıma ve otonom sürüş gibi alanlarda yaygın olarak kullanılır. Derin öğrenme, tıbbi teşhis, finansal tahmin, oyun geliştirme ve öneri sistemleri gibi çeşitli uygulamalarda da kullanılarak, insan benzeri öğrenme ve karar verme yetenekleriyle teknoloji ve bilimde devrim yaratmaktadır.

Derin öğrenme, metin sınıflandırma alanında önemli ilerlemeler kaydeden bir teknolojidir. Metin sınıflandırma, metin verilerini belirli kategorilere ayırma işlemidir ve derin öğrenme modelleri bu süreçte önemli bir rol oynar. Derin öğrenme yöntemleri, büyük veri kümelerinden karmaşık dil örüntülerini öğrenme kapasitesine sahiptir. Örneğin, CNN ve RNN gibi modeller, metin verilerinin anlamını ve bağlamını yakalayıp sınıflandırma görevlerinde yüksek doğruluk sağlar. Ayrıca, Transformer tabanlı modeller, özellikle BERT ve GPT, bağlamsal dil anlayışında üstün performans gösterir. Bu yöntemler, duygu analizi, konu tespiti ve spam filtresi gibi çeşitli uygulamalarda kullanılarak, metin sınıflandırmanın etkinliğini ve doğruluğunu artırır.

3.6. Derin Öğrenme Modellerinde Sık Kullanılan Katmanlar

Derin öğrenme modellerinde metin sınıflandırma amacıyla sıkça kullanılan katmanlar arasında Embedding, Convolutional ve Attention katmanları yer alır. Embedding katmanı, kelimeleri sayısal vektörlere dönüştürerek modelin dildeki anlam ilişkilerini öğrenmesini sağlar. Convolutional katmanlar, metindeki n-gram özelliklerini yakalayıp yerel bilgi çıkarımı yapar. Attention katmanları ise, önemli bilgi bölgelerine odaklanarak modelin performansını artırır ve bağlamsal anlamı daha iyi yakalar. Bu katmanlar, derin öğrenme modellerinin metin verilerini daha etkin ve doğru bir şekilde sınıflandırmasını sağlayan temel yapı taşlarıdır.

3.6.1. Giriş Katmanı

Bir derin öğrenme modelinde giriş katmanı, modelin aldığı ham verinin ilk işlendiği katmandır ve modelin geri kalanına veri sağlamak için gerekli ön işlemleri yapar. Giriş katmanı, veriyi modelin anlayabileceği bir formata dönüştürür. Bu katman, metin verileri için kelime dizilerini sayısal vektörlere, görüntüler için piksel değerlerini, ses verileri için ise dalga formu özelliklerini temsil edebilir. Giriş katmanı, veri türüne bağlı olarak Embedding katmanı,

piksel tabanlı tensörler veya spektral özellikler gibi çeşitli formatlardan oluşabilir. Bu, modelin farklı veri türleriyle çalışmasını ve öğrenme sürecine başlamasını sağlar.

3.6.2. Embedding Katmanı

Embedding katmanı, metin verilerini sayısal vektörlere dönüştürerek derin öğrenme modellerinin dildeki anlam ilişkilerini öğrenmesini sağlar. Her kelimeyi düşük boyutlu bir vektörle temsil eden bu katman, kelimeler arasındaki benzerlikleri ve bağlamı yakalamada etkilidir. Embedding katmanı, yüksek boyutlu ve seyrek kelime temsillerini yoğun ve anlamlı vektörlerle değiştirerek, modelin daha hızlı ve verimli öğrenmesini sağlar. Metin işlemede, dil modellerinin performansını artırır ve anlam ilişkilerini daha iyi kavrayarak sınıflandırma, duygu analizi ve makine çevirisi gibi görevlerde önemli avantajlar sunar.

3.6.3. Dropout Katmanı

Dropout katmanı, derin öğrenme modellerinde ezberlemeyi önlemek için kullanılan bir düzenleme tekniğidir. Bu katman, her eğitim adımında belirli bir oranda nöronları rastgele olarak devre dışı bırakır, yani sıfırlanır. Bu süreç, modelin belirli nöronlara veya özelliklere aşırı bağımlı olmasını engeller ve genel performansını artırır. Eğitim sırasında her ileri geçişte farklı bir alt ağı eğitilmesini sağlayarak, modelin daha dayanıklı ve genelleştirilebilir olmasına katkıda bulunur. Test aşamasında ise tüm nöronlar aktif hale getirilir ve dropout uygulanmaz. Dropout, özellikle derin sinir ağlarında yaygın olarak kullanılır.

SpatialDropout1D katmanı, özellikle metin verisi gibi tek boyutlu veri dizileri üzerinde çalışan derin öğrenme modellerinde kullanılan bir düzenleme tekniğidir. Bu katman, belirli oranlarda rastgele olarak belirli özellik haritalarını sıfırlayarak modelin aşırı öğrenmesini engellemeye yardımcı olur. Ancak klasik Dropout katmanının aksine, SpatialDropout1D, tüm özellik haritalarını (örneğin, belirli bir kelime vektöründeki tüm değerleri) aynı anda düşürür. Bu, uzaysal (spatial) yapıyı korur ve modelin önemli örüntüleri daha iyi öğrenmesini sağlar. Bu katman, CNN ve RNN gibi modellerde yaygın olarak kullanılır.

3.6.4. LSTM Katmanı

LSTM katmanı, özellikle sıralı veri ile çalışan derin öğrenme modellerinde yaygın olarak kullanılan bir tür RNN katmanıdır. LSTM'ler, uzun vadeli bağımlılıkları öğrenme ve koruma yetenekleri sayesinde zaman serileri, metin verileri ve diğer ardışık veriler üzerinde başarılı performans gösterirler. LSTM katmanı, unutma kapısı, giriş kapısı ve çıkış kapısı olmak üzere üç temel bileşen içerir. Bu kapılar, hücre durumu ve gizli durum arasındaki bilgi akışını kontrol

ederek önemli bilgilerin korunmasını ve gereksiz bilgilerin unutulmasını sağlar. LSTM'ler, dil modelleme, metin tahmini ve ses tanıma gibi uygulamalarda yaygın olarak kullanılır.

LSTM'nin temel birimi olan hafıza hücresi, bilgiyi zaman içinde saklar ve manipüle eder. Hücre durumu C_t olarak adlandırılır. LSTM, bilgiyi hücre durumu boyunca kontrol etmek için üç farklı kapı kullanır:

- Unutma Kapısı (Forget Gate): Hücre durumundan hangi bilgilerin atılacağını belirler.
- Giriş Kapısı (Input Gate): Yeni bilgi eklemek için hücre durumunu günceller.
- Çıkış Kapısı (Output Gate): Hücre durumunun hangi kısmının çıktı olarak kullanılacağını belirler.

Her kapı, sigmoid aktivasyon fonksiyonu (σ) ve bir ağırlık matrisi kullanarak hesaplanır. Unutma kapısı f_t şu şekilde hesaplanır:

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f) \quad (3.1)$$

Burada W_f unutma kapsının ağırlık matrisi, b_f unutma kapısının bias vektörü, h_{t-1} önceki gizli durum, x_t şu an ki girdi ve σ sigmoid aktivasyon fonksiyonudur. Giriş kapısı $\tilde{C}_t$ ise aşağıdaki gibi hesaplanır:

$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i) \quad (3.2)$$

$$\tilde{C}_t = \tanh(W_C \cdot [h_{t-1}, x_t] + b_C)$$

Burada W_i giriş kapısının ağırlık matrisi, b_i giriş kapısının bias vektörü, W_C hücre durumu adaylarının ağırlık matrisi, b_C hücre durumu adaylarının bias vektörü, $\tanh$ hiperbolik tanjant fonksiyonudur.

Hücre durumu C_t ise aşağıdaki gibi güncellenir:

$$C_t = f_t \cdot C_{t-1} + i_t \cdot \tilde{C}_t \quad (3.3)$$

Çıkış kapısının çıktısı o_t ve güncellenmiş gizli durum h_t şu şekilde hesaplanır:

$$o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o) \quad (3.4)$$

$$h_t = o_t \cdot \tanh(C_t)$$

Burada W_0 çıkış kapısının ağırlık matrisi, b_o çıkış kapısının bias vektörüdür.

3.6.5. Tam Bağlı Katman

Tam bağlı (dense) katmanı, yapay sinir ağlarında en temel yapı taşlarından biridir. Her nöronun bir önceki katmandaki tüm nöronlara bağlı olduğu bir katmandır. Dense katmanı, girdileri alır, ağırlıklarla çarpar, ardından bias ekler ve aktivasyon fonksiyonu ile işlenmiş çıktıları üretir. Dense katmanının matematiksel ifadesi aşağıdaki gibidir:

$$y = f(W \cdot x + b) \quad (3.5)$$

Burada x giriş vektörü, W ağırlık matrisi olup girişlerin her bir nöronla olan bağlantı ağırlıklarını içerir. b bias vektörü, f genellikle doğrusal olmayan bir aktivasyon fonksiyonu (örneğin ReLU, sigmoid, tanh), y çıkış vektörüdür.

Tama bağlı katman, her nöronun bir önceki katmandaki tüm nöronlardan bilgi alması nedeniyle, yüksek esneklik ve öğrenme kapasitesi sağlar. Bu katman, sınıflandırma, regresyon ve diğer birçok makine öğrenimi görevinde yaygın olarak kullanılır ve modelin doğrusal olmayan ilişkileri öğrenmesini mümkün kılar.

3.6.6. Batch Normalizasyon Katmanı

BatchNormalization katmanı, derin öğrenme modellerinin eğitim sürecini hızlandırmak ve stabil hale getirmek için kullanılan bir normalizasyon tekniğidir. Bu katman, her minibatch'teki her katmanın aktivasyonlarını yeniden ölçeklendirir ve merkezler, böylece modelin öğrenme sürecinde daha tutarlı ve hızlı bir ilerleme sağlanır.

BatchNormalization katmanında öncelikle ortalama ve varyans hesaplanır:

$$\begin{aligned} \mu_B &= \frac{1}{m} \sum_{i=1}^m x_i \\ \sigma_B^2 &= \frac{1}{m} \sum_{i=1}^m (x_i - \mu_B)^2 \end{aligned} \quad (3.6)$$

Burada x_i minibatch'teki bir örneğin girişi, μ_B minibatch'in ortalaması, σ_B^2 minibatch'in varyansı, m minibatch'teki örnek sayısıdır.

Bir sonraki adımda normalizasyon işlevi uygulanır:

$$\hat{x} = \frac{x_i - \mu_B}{\sqrt{\sigma_B^2 + \epsilon}} \quad (3.7)$$

Burada $\hat{x}$ normalize edilmiş değer, ϵ küçük bir sabit olup, sıfıra bölmeyi önler.

Son adımda ölçeklendirme ve kaydırma işlevi uygulanır:

$$y_i = \hat{\gamma}x + \beta \quad (3.8)$$

Burada γ öğrenilebilir ölçeklendirme parametresi, β öğrenilebilir kaydırma parametresidir.

BatchNormalization katmanı, her minibatch'teki aktivasyonları normalize ederek, öğrenme hızını artırır ve modelin daha derin katmanlar öğrenmesine olanak tanır. Ayrıca, bu katman, ağırlıkların aşırı büyük veya küçük olmasından kaynaklanan sorunları azaltarak modelin genel performansını iyileştirir. Bu nedenle, BatchNormalization, özellikle derin sinir ağları için yaygın olarak kullanılır.

3.7. Optimizasyon

Derin öğrenme, özellikle CNN gibi ağlarda çeşitli optimizasyon algoritmaları kullanılarak modellerin eğitim sürecini iyileştirir ve hızlandırır. Bu optimizasyon algoritmaları, modelin ağırlıklarını güncelleyerek kayıp fonksiyonunu minimize etmeyi amaçlar. İşte yaygın olarak kullanılan bazı optimizasyon algoritmaları:

3.7.1. Stochastic Gradient Descent (SGD)

SGD, her eğitim örneği veya küçük bir minibatch kullanarak ağırlıkları iteratif olarak günceller. Bu yöntem, büyük veri setlerinde hesaplama maliyetini düşürür ve daha hızlı bir öğrenme süreci sağlar. Ancak, SGD'nin doğruluk ve kararlılık sorunları olabilir.

$$\theta = \theta - \eta \cdot \nabla J(\theta; x_i, y_i) \quad (3.9)$$

Burada θ modelin ağırlıkları, η öğrenme oranı, ∇J kayıp fonksiyonunun gradyanıdır.

3.7.2. Momentum

Momentum, SGD'nin iyileştirilmiş bir versiyonudur. Önceki gradyanların hızlanma etkisini kullanarak güncellemeleri hızlandırır ve sarsıntıları azaltır.

$$\begin{aligned} v_t &= \gamma v_{t-1} + \eta \nabla J(\theta) \\ \theta &= \theta - v_t \end{aligned} \quad (3.10)$$

Burada v_t momentum terimi, γ momentum katsayısıdır.

3.7.3. AdaGrad

AdaGrad, her parametre için ayrı bir öğrenme oranı kullanarak, nadiren güncellenen parametrelerin daha büyük adımlarla güncellenmesini sağlar.

$$\theta = \theta - \frac{\eta}{\sqrt{G_{t,ii} + \epsilon}} \cdot \nabla J(\theta) \quad (3.11)$$

Burada G_t geçmiş gradyanların karelerinin toplamı, (numerik istikrar için) küçük bir sabittir.

3.7.4. RMSProp

RMSProp, AdaGrad'ın bir varyasyonudur. Gradyanların karelerinin hareketli ortalamasını kullanarak, öğrenme oranlarının adaptif olarak ayarlanmasını sağlar.

$$\begin{aligned} E[g^2]_t &= \gamma E[g^2]_{t-1} + (1 - \gamma) g_t^2 \\ \theta &= \theta - \frac{\eta}{\sqrt{E[g^2]_t + \epsilon}} \cdot \nabla J(\theta) \end{aligned} \quad (3.12)$$

Burada $E[g^2]_t$ gradyanların karelerinin beklenen değeridir.

3.7.5. Adaptive Moment Estimation (Adam)

Adam, Momentum ve RMSProp'un birleşimidir. Hem gradyanların ortalamasını hem de ikinci momentini kullanarak adaptif öğrenme oranları sağlar.

$$\begin{aligned} m_t &= \beta_1 m_{t-1} + (1 - \beta_1) g_t \\ v_t &= \beta_2 v_{t-1} + (1 - \beta_2) g_t^2 \\ \hat{m} &= \frac{m_1}{1 - \beta_1^t} \\ \hat{v} &= \frac{v_1}{1 - \beta_2^t} \\ \theta &= \theta - \frac{\eta \hat{m}}{\sqrt{\hat{v}} + \epsilon} \end{aligned} \quad (3.13)$$

Burada m_t moment, v_t ikinci moment, β_1 ve β_2 moment katsayılarıdır.

Bu optimizasyon algoritmaları, CNN'ler gibi derin öğrenme ağlarının daha hızlı, kararlı ve etkili bir şekilde öğrenmesini sağlayarak model performansını önemli ölçüde artırır.

3.8. Kayıp Fonksiyonu

Derin öğrenme modellerinde, özellikle CNN gibi ağlarda kullanılan kayıp fonksiyonu (loss function), modelin tahminlerinin doğruluğunu değerlendirmek ve ağırlıklarını optimize etmek için temel bir ölçüttür. Kayıp fonksiyonu, modelin ürettiği çıktı ile gerçek değer arasındaki farkı hesaplar ve bu farkı minimize etmek için kullanılır. CNN'ler, görüntü sınıflandırma, nesne algılama gibi görevlerde yaygın olarak kullanıldığından, uygun kayıp fonksiyonunun seçilmesi model performansı için kritik öneme sahiptir.

3.8.1. Kategorik Çapraz Entropi

Kategorik çapraz entropi, sınıflandırma problemlerinde yaygın olarak kullanılır. Modelin çıktıları ile gerçek etiketler arasındaki farkı ölçer. Bu fonksiyon, modelin yanlış tahminlerde cezalandırılmasını sağlar.

$$L = - \sum_{i=1}^N y_i \log(\hat{y}_i) \quad (3.14)$$

Burada y_i gerçek etiket, $\hat{y}_i$ modelin tahmin ettiği olasılıktır.

3.8.2. İkili Çapraz Entropi

İkili sınıflandırma problemlerinde kullanılır. Her sınıf için ayrı ayrı hesaplanır ve toplam kayıp değeri olarak kullanılır.

$$L = - \frac{1}{N} [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)] \quad (3.15)$$

Burada y_i gerçek etiket (0 veya 1), $\hat{y}_i$ modelin tahmin ettiği olasılıktır.

3.8.3. Ortalama Kare Hata

Regresyon problemlerinde yaygın olarak kullanılır. Tahmin edilen değer ile gerçek değer arasındaki kare farkların ortalamasını alır.

$$L = \frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2 \quad (3.16)$$

Burada y_i gerçek değer, $\hat{y}_i$ tahmin edilen değerdir.

3.9. Model Doğrulama ve Değerlendirme

3.9.1. Hata matrisi

Hata matrisi, sınıflandırma modellerinin performansını değerlendirmek için kullanılan temel bir araçtır. Matris, modelin gerçek sınıflarla tahmin edilen sınıflar arasındaki ilişkileri tablo şeklinde gösterir. Hata matrisi dört ana bileşenden oluşur: Doğru Pozitifler (TP), Doğru Negatifler (TN), Yanlış Pozitifler (FP), ve Yanlış Negatifler (FN). TP ve TN, modelin doğru sınıflandırdığı örnekleri temsil ederken, FP ve FN, modelin yanlış sınıflandırdığı örnekleri temsil eder. Çalışmamızda, hata matrisinde olumsuz yorumlar Doğru Pozitifler (TP) olarak, olumlu yorumlar ise Doğru Negatifler (TN) olarak ifade edilmiştir. TP, modelin olumsuz yorumları doğru bir şekilde tespit ettiği durumu gösterirken, TN, modelin olumlu yorumları doğru bir şekilde sınıflandırdığı durumu yansıtır. Bu yapı, modelin olumsuz yorumlardaki başarı oranını (duyarlılık) ve olumlu yorumlardaki başarı oranını (özgüllük) ayrı ayrı değerlendirmeyi mümkün kılar. Yanlış Pozitifler (FP), olumlu olarak sınıflandırılan ancak gerçekte olumsuz olan yorumları, Yanlış Negatifler (FN) ise olumsuz olarak sınıflandırılan ancak gerçekte olumlu olan yorumları ifade eder. Bu durum, modelin her iki sınıf için performansını detaylı bir şekilde analiz etmeyi sağlar.

Bu matris, modelin duyarlılık, özgüllük, doğruluk ve F1 skoru gibi performans metriklerinin hesaplanmasına olanak tanır ve modelin hangi sınıflarda daha fazla hata yaptığını analiz etmeyi sağlar. Tablo 3.2’de performans metrikleri açıklanmıştır.

Tablo 3.2 Performans metrikleri ve kısa açıklamaları

Metrik	Formül	Kısa açıklama
Doğruluk (Accuracy)	$\frac{TP + TN}{TP + TN + FP + FN}$	Modelin tüm doğru tahminlerinin toplam tahminlere oranı.
Hata (Error)	$\frac{FP + FN}{TP + TN + FP + FN}$	Modelin tüm yanlış tahminlerinin toplam tahminlere oranı.
Duyarlılık (Sensitivity)	$\frac{TP}{TP + FN}$	Modelin doğru bir şekilde olumsuz yorumları tespit etme oranı.
Özgüllük (Specificity)	$\frac{TN}{TN + FP}$	Modelin doğru bir şekilde olumlu yorumları tespit etme oranı.

Keskinlik (Precision)	$\frac{TP}{TP + FP}$	Olumsuz olarak tahmin edilen yorumların gerçekten olumsuz olma oranı.
Yanlış Pozitif Oranı	$\frac{FP}{FP + TN}$	Olumlu olarak tahmin edilen, ancak gerçekte olumsuz olan yorumların oranı.
F1 Skoru	$2 \frac{Duyarlılık * Keskinlik}{Duyarlılık + Keskinlik}$	Modelin duyarlılık ve keskinlik oranlarının harmonik ortalaması; dengeyi ifade eder.

3.9.2. ROC Eğrileri

ROC (Receiver Operating Characteristic) eğrisi, bir sınıflandırma modelinin performansını değerlendirmek için kullanılan grafiksel bir araçtır. Yatay ekseninde Yanlış Pozitif Oranı (FPR), dikey ekseninde Doğru Pozitif Oranı (TPR) yer alır. Modelin farklı eşik değerlerinde TPR ve FPR'yi gösterir. Eğrinin altındaki alan (AUC), modelin ayırıştırma yeteneğini ifade eder. ROC eğrisi, özellikle dengesiz veri setlerinde ve ikili sınıflandırma problemlerinde faydalıdır, çünkü modelin farklı eşik değerlerinde nasıl performans gösterdiğini görselleştirir ve optimum eşik değerinin belirlenmesine yardımcı olur.

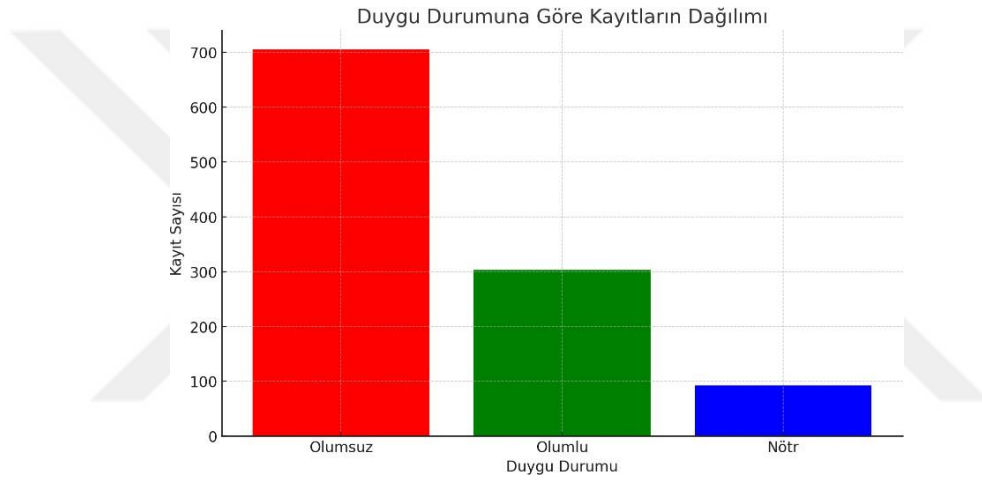
DÖRDÜNCÜ BÖLÜM

4. Deneysel Sonuçlar ve Bulgular

Tez çalışmasındaki tüm deneysel çalışmalar HP İş istasyonu üzerinde gerçekleştirilmiştir. Bu iş istasyonu Intel(R) Xeon(R) Gold 6132 CPU @ 2.60 GHz işlemci ve NVIDIA® Quadro® P6000 24GB ekran kartına sahiptir.

4.1. Veri Seti Analizi

Deneysel çalışmaların daha anlamlı bir zemin üzerinde oturtulmasını sağlamak amacıyla öncelikle veri setine ait bulgular paylaşılmıştır. Şekil 4.1’de veri setindeki örneklerin dağılımına yer verilmiştir.

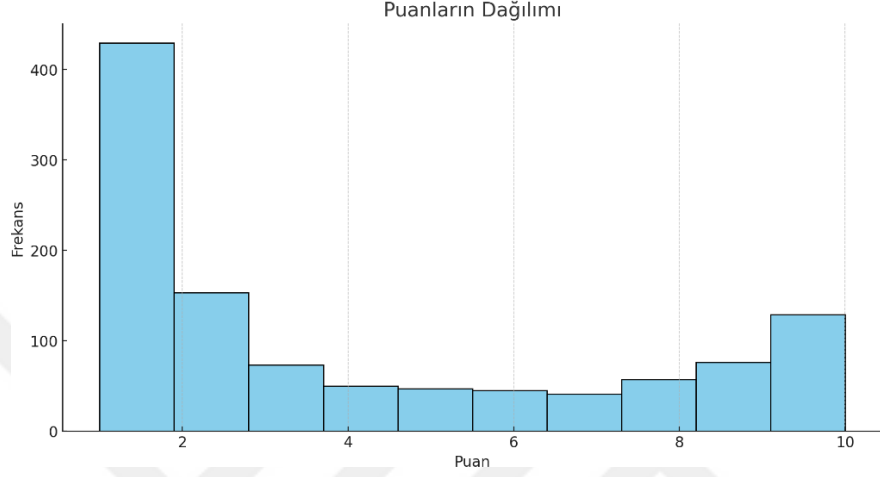


Şekil 4.1 Verilerin duygu durumuna göre dağılımı

Veri setinin duygu durumu dağılımını gösteren grafik incelendiğinde, 705 kayıt olumsuz, 303 kayıt olumlu ve 92 kayıt nötr olarak sınıflandırılmıştır. Olumsuz duygu durumu, toplam kayıtların en büyük kısmını oluştururken, olumlu ve nötr duygu durumları daha düşük oranlarda gözlemlenmiştir. Bu dağılım, müşteri memnuniyetinin iyileştirilmesi gereken alanların olduğunu göstermektedir. Özellikle, olumsuz geri bildirimlerin yüksek olması, havayolu hizmet kalitesinin belirli konularda müşteri beklentilerini karşılamadığını işaret etmektedir. Bulgular, müşteri hizmetleri ve operasyonel süreçlerin gözden geçirilmesi gerektiğini ortaya koymaktadır.

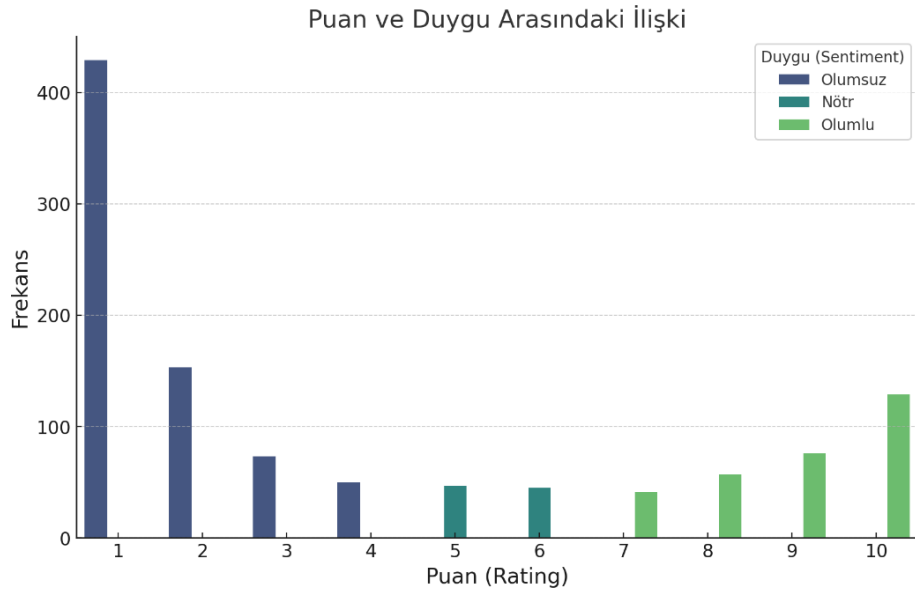
Şekil 4.2’de veri setindeki yorumlar ve puanlara ilişkin dağılım sunulmuştur. Grafikte, müşteri geri bildirimlerinin puan dağılımı gösterilmektedir. En yüksek yoğunluk 1 puanda, bu da kullanıcıların çoğunlukla çok olumsuz deneyimler yaşadığını belirtmektedir. 2 ve 10 puanları, diğer belirgin yoğunluklardır, ancak 10 puan pozitif deneyimleri, 2 puan ise

genellikle olumsuz deneyimleri temsil etmektedir. Orta düzeyde puanlar (4, 5, 6) ise düşük yoğunluktadır, bu da kullanıcıların genellikle aşırı uçlarda puanlama eğiliminde olduğunu göstermektedir. Bu dağılım, müşteri memnuniyeti konusunda önemli uçurumlar olduğunu ve hizmet kalitesinin tutarlılığında ciddi farklılıklar bulunduğunu göstermektedir. Bu bulgular, hizmet süreçlerinde iyileştirmeler yapılması gerektiğine işaret etmektedir.

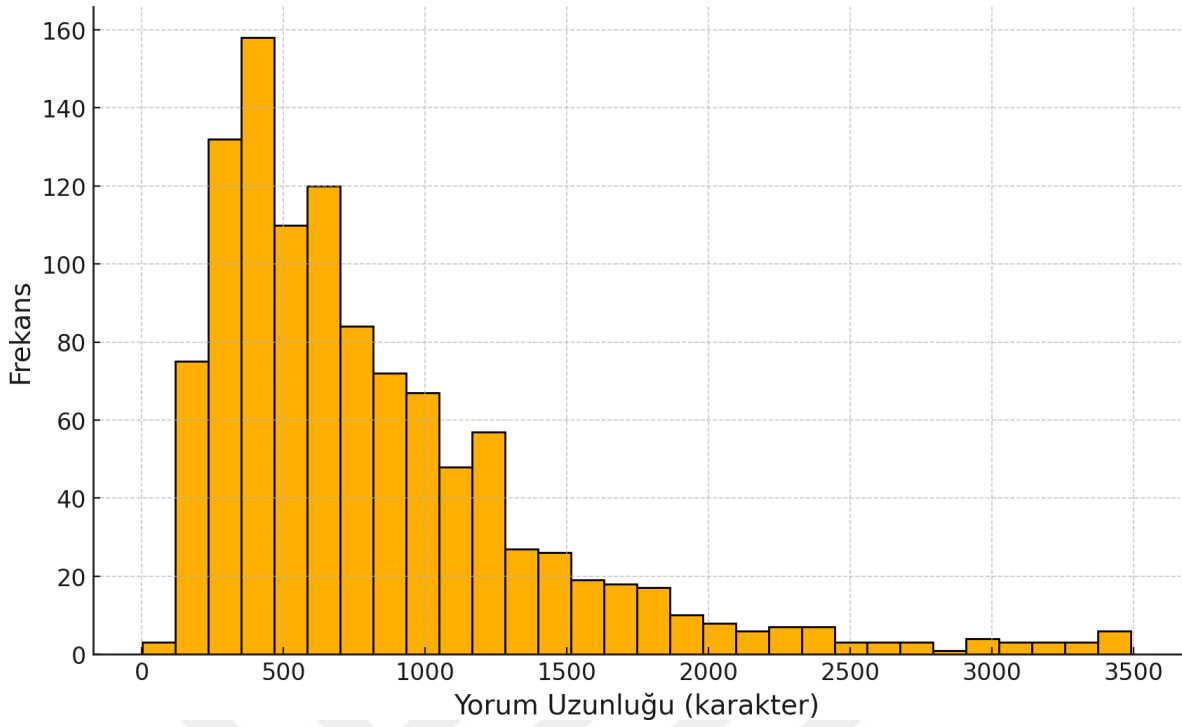


Şekil 4.2 Veri setindeki yorumlar ve puanlara ilişkin dağılım

Bu noktada puan durumu ve müşteri duygu durumu arasındaki ilişkinin de incelenmesi gerekmektedir. Şekil 4.3’de puanlar ile birlikte müşteri duygu durumları arasındaki ilişkiye yer verilmiştir.



Şekil 4.3. Puanlar ve müşteri duygu durumları



Şekil 4.5 Yorum uzunluğu dağılımı

Şekil 4.5’de yorum uzunluğu dağılımına yer verilmiştir. Yorum uzunluğu analizine göre, müşteri yorumları geniş bir karakter yelpazesine sahiptir, ortalama yorum uzunluğu 820 karakterdir. Yorum uzunlukları 3 ile 3491 karakter arasında değişmektedir. Medyan değer 638 karakter olup, yorumların çoğunluğunun 406 ile 1060 karakter arasında yoğunlaştığı görülmektedir. Bu dağılım, müşterilerin genellikle ayrıntılı geri bildirim verdiğini göstermektedir. Yorum uzunluğu, müşteri memnuniyeti hakkında önemli bilgiler sunar; uzun yorumlar genellikle daha güçlü duygular ve ayrıntılı deneyimler içerir. Bu bulgular, hizmet kalitesini iyileştirmek için derinlemesine analiz yapılması gereken alanları belirlemede faydalıdır.

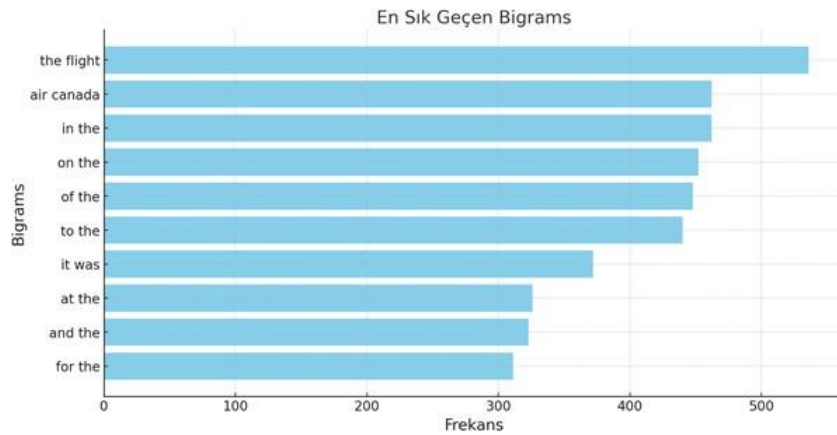
4.3. N-gram Analizi

N-gram analizi, bir metin içinde yan yana bulunan n sayıda kelime grubunu (bigrams: iki kelime, trigrams: üç kelime) inceleyen bir tekniktir. Bu analiz, metindeki yaygın kelime öbeklerini ve bunların frekanslarını belirlemeyi sağlar. N-gram analizi, müşteri yorumları gibi büyük metin veri setlerinde belirli şikâyetleri, memnuniyet alanlarını veya dikkat çeken konuları tespit etmek için kullanılır. Bu sayede, hizmet kalitesini artırmak için önemli ipuçları elde edilir ve müşteri geri bildirimleri daha etkili bir şekilde değerlendirilir. Ayrıca, doğal dil işleme ve makine öğrenimi uygulamalarında da yaygın olarak kullanılır.

Tablo 4.1 Bigrams frekans tablosu

Bigram	Frekans
the flight	536
air canada	462
in the	462
on the	452
of the	448
to the	440
it was	372
at the	326
and the	323
for the	311

Tablo 4.1’de müşteri yorumlarında en sık geçen on bigramı ve bunların frekanslarını göstermektedir. "The flight" bigramı 536 kez geçerek en yaygın kullanılan öbek olmuştur, bu da uçuş deneyimlerinin sıkça tartışıldığını gösterir. "Air Canada" ve "in the" bigramları 462 kez geçmiştir, bu da havayolu şirketi hakkında yoğun geri bildirim olduğunu belirtir. Diğer yaygın bigramlar arasında "on the", "of the", ve "to the" yer almakta, her biri 440 kezden fazla geçmektedir. Bu bigramlar, müşterilerin hizmetle ilgili genel deneyimlerini ve sıkça karşılaşılan konuları anlamamıza yardımcı olmaktadır. Şekil 4.6’da bigrams grafiğine yer verilmiştir.

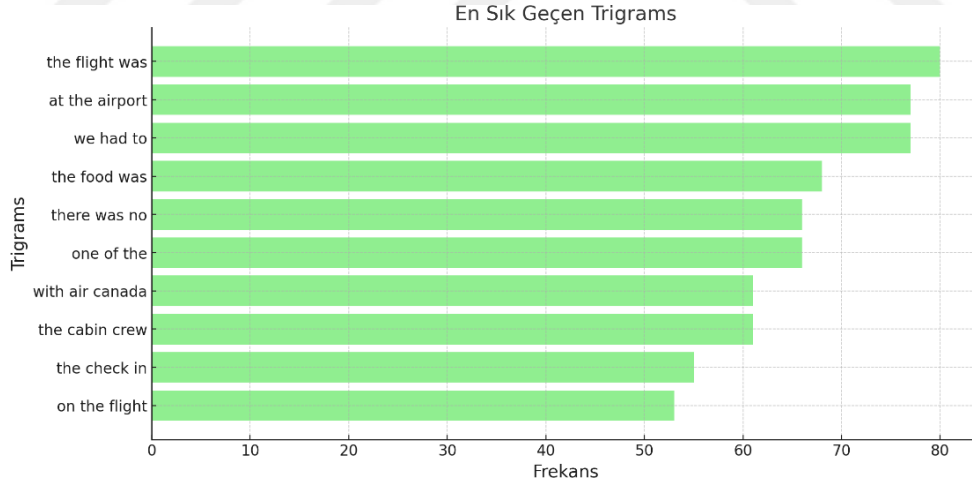


Şekil 4.6 Bigrams grafiği

Tablo 4.2 Trigram frekans tablosu

Trigram	Frekans
the flight was	80
at the airport	77
we had to	77
the food was	68
there was no	66
one of the	66
with air canada	61
the cabin crew	61
the check in	55
on the flight	53

Tablo 4.2’de müşteri yorumlarında en sık geçen on trigramı ve bunların frekanslarını göstermektedir.



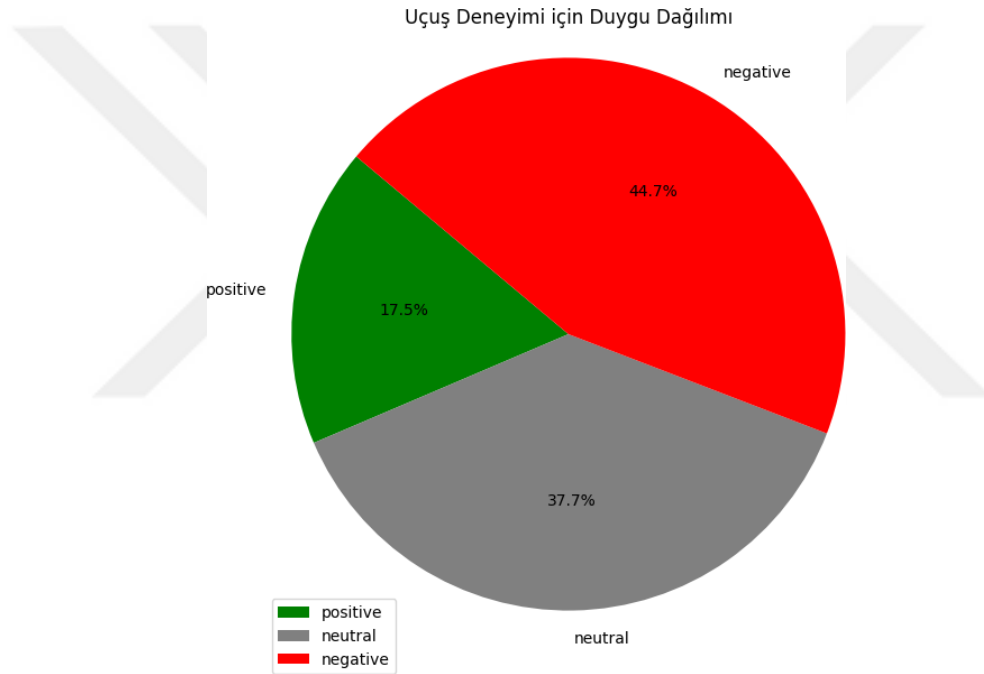
Şekil 4.7 Trigrams grafiği

"The flight was" trigramı 80 kez ile en yaygın kullanılan öbek olup, uçuş deneyimlerine dair detaylı geri bildirimlerin olduğunu belirtir. "At the airport" ve "we had to" trigramları 77 kez geçmiştir, bu da havaalanı deneyimleri ve zorunlu eylemlerin sıkça tartışıldığını gösterir. "The food was" ve "there was no" trigramları, hizmet kalitesi ve eksikliklere yönelik yorumları yansıtmaktadır. Diğer önemli trigramlar arasında "with air canada", "the cabin crew", ve "the

check in" yer almakta, bu da müşteri deneyimlerinin çeşitli yönlerine dair geri bildirimlerin olduğunu göstermektedir.

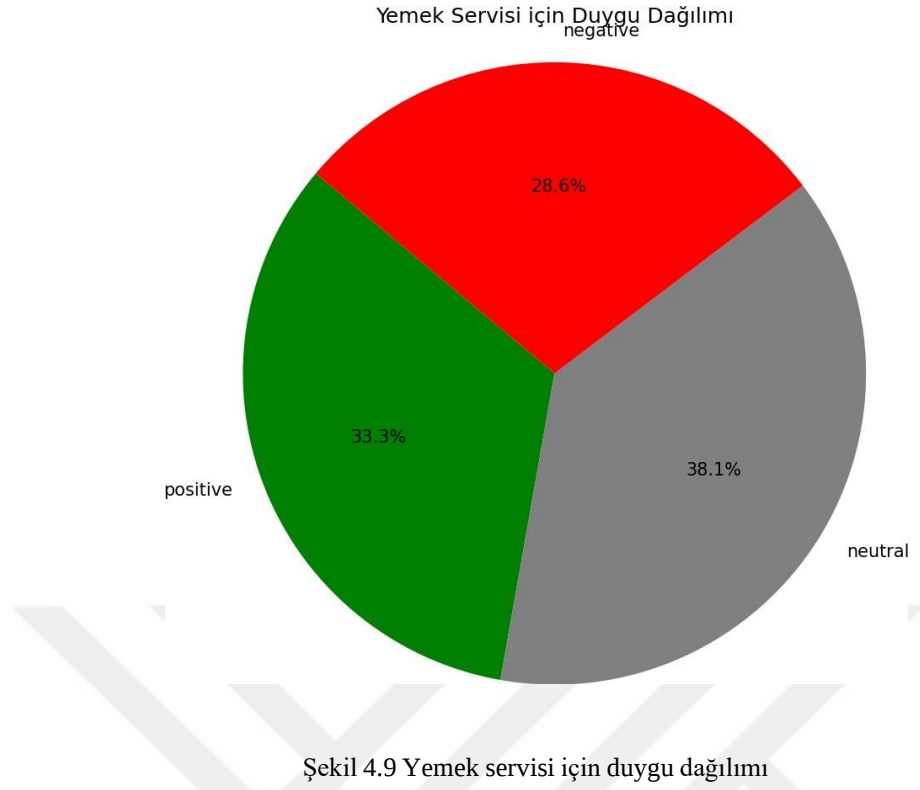
4.4. Bileşene Yönelik Duygu Analizi

Bu analizde, müşteri geri bildirimlerindeki belirli bileşenlere yönelik duygu analizi gerçekleştirilmiştir. İlk olarak, yorumlar "uçuş deneyimi", "yemek servisi" ve "müşteri hizmetleri" olarak üç ana bileşene ayrılmıştır. Her bileşen için ilgili anahtar kelimeler belirlenmiş ve bu anahtar kelimeleri içeren cümleler tespit edilmiştir. Ardından, VADER duygu analiz aracı kullanılarak her cümlenin duygu skoru hesaplanmış ve pozitif, negatif veya nötr olarak sınıflandırılmıştır. Sonuçlar, her bileşen için pozitif, negatif ve nötr duyguların yüzdesel dağılımını gösteren pasta grafiklerle görselleştirilmiştir.



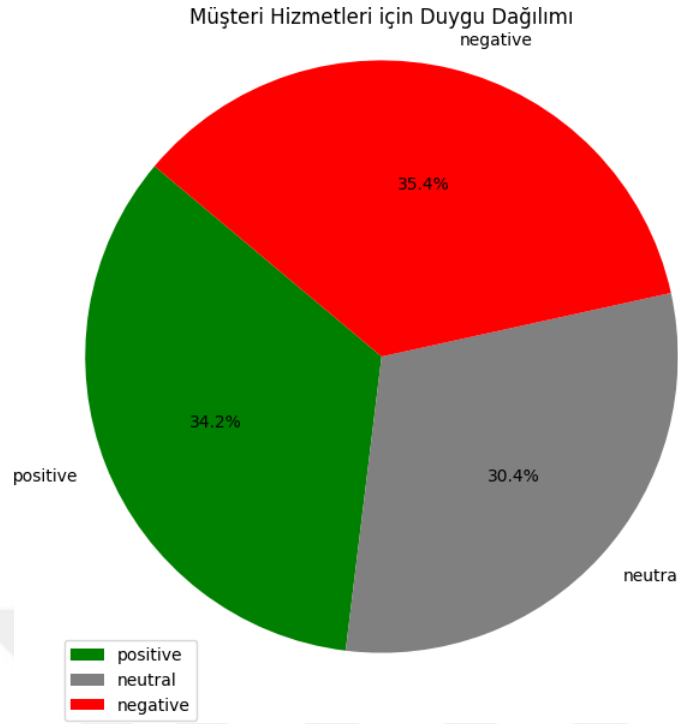
Şekil 4.8 Uçuş deneyimi için duygu dağılımı

Şekil 4.8’de uçuş deneyimi için duygu dağılımı grafiğine yer verilmiştir, uçuş deneyimi bileşeni için müşteri yorumlarındaki duygu dağılımı gösterilmektedir. Yüzde 44.7 ile negatif duygu en yüksek oranı temsil ederken, yüzde 37.7 ile nötr ve yüzde 17.5 ile pozitif duygu oranları takip etmektedir. Bu dağılım, müşterilerin uçuş deneyimleri hakkında genellikle olumsuz geri bildirimlerde bulunduğunu göstermektedir. Nötr yorumlar da önemli bir paya sahiptir, bu da bazı müşterilerin uçuş deneyimleri hakkında kararsız veya belirsiz görüşlere sahip olduğunu gösterir. Pozitif yorumlar ise en düşük oranda olup, uçuş deneyiminden memnun kalan müşteri sayısının daha az olduğunu ortaya koymaktadır. Bu bulgular, uçuş deneyimi alanında iyileştirmeler yapılması gerektiğini işaret etmektedir.



Şekil 4.9’da yemek servisi için duygu dağılımı grafiğine yer verilmiştir, yemek servisi bileşeni için müşteri yorumlarındaki duygu dağılımı gösterilmektedir. Yüzde 38.1 ile nötr duygu en yüksek oranı temsil ederken, yüzde 33.3 ile pozitif ve yüzde 28.6 ile negatif duygu oranları takip etmektedir. Bu dağılım, müşterilerin yemek servisi hakkında genellikle kararsız veya nötr görüşlere sahip olduğunu göstermektedir. Pozitif yorumlar, müşterilerin yemek servisinden memnun olduklarını belirtirken, negatif yorumlar ise memnuniyetsizliklerin de önemli bir paya sahip olduğunu ortaya koymaktadır. Bu bulgular, yemek servisi alanında hem olumlu hem de olumsuz geri bildirimlerin dikkate alınarak iyileştirmeler yapılması gerektiğini işaret etmektedir.

Şekil 4.10’da müşteri hizmetleri için duygu dağılımı grafiğine yer verilmiştir, müşteri hizmetleri bileşeni için müşteri yorumlarındaki duygu dağılımı gösterilmektedir. Yüzde 35.4 ile negatif duygu en yüksek oranı temsil ederken, yüzde 34.2 ile pozitif ve yüzde 30.4 ile nötr duygu oranları takip etmektedir. Bu dağılım, müşteri hizmetleri hakkında hem olumlu hem de olumsuz geri bildirimlerin neredeyse eşit derecede önemli olduğunu göstermektedir. Nötr yorumlar da dikkate değer bir paya sahiptir, bu da bazı müşterilerin müşteri hizmetleri hakkında kararsız veya belirsiz görüşlere sahip olduğunu göstermektedir. Bu bulgular, müşteri hizmetleri alanında iyileştirmeler yapılması gerektiğini ve bu konudaki memnuniyetin artırılması için stratejiler geliştirilmesi gerektiğini işaret etmektedir.



Şekil 4.10 Müşteri hizmetleri için duygu dağılımı

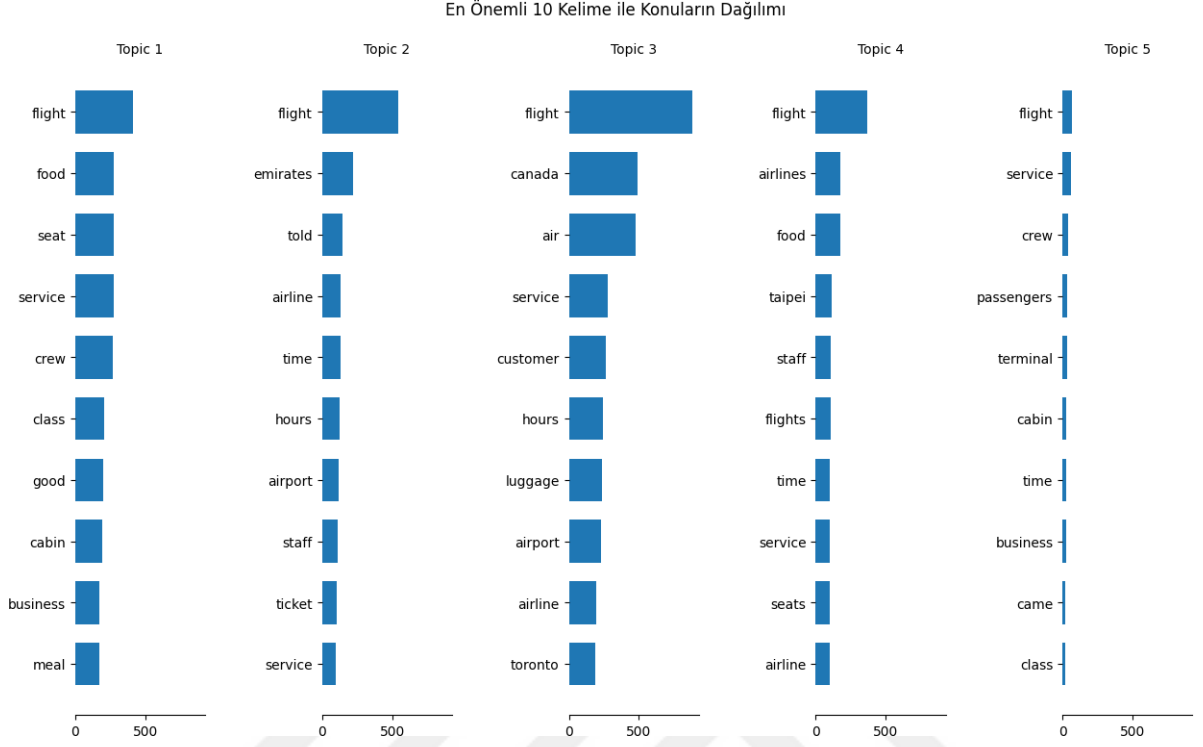
4.5. Konu Analizi

Latent Dirichlet Allocation (LDA) Yöntemi, metin verilerinde gizli konuları belirlemek için kullanılan bir makine öğrenimi tekniğidir. LDA, her belgenin birden fazla konuyu içerdiğini ve her konunun belirli kelimelerle temsil edildiğini varsayarak, metinleri konu dağılımlarına göre sınıflandırır.

Yapılan analiz, müşteri yorumlarında geçen kelimeleri kullanarak yorumların gizli temalarını veya konularını belirlemektir. LDA, her yorumun birden fazla konuya ait olabileceği varsayımıyla çalışır ve her konunun belirli kelimelerle temsil edildiği bir dağılım çıkarır. Bu yöntem, büyük metin verilerinde yaygın temaları keşfetmek ve metinlerin hangi konular etrafında yoğunlaştığını anlamak için kullanılır. Sonuç olarak, müşteri yorumlarının hangi ana temalar etrafında toplandığı ve bu temaların öne çıkan kelimelerle nasıl ifade edildiği ortaya konur.

Şekil 4.11’de müşteri yorumları üzerinden LDA yöntemi kullanılarak belirlenen beş ana konu ve bu konuların öne çıkan kelimeleri gösterilmektedir. Her bir konu, ilgili kelimelerle temsil edilmektedir. Örneğin, "Topic 1" için "flight", "food", "seat", "service" ve "crew" gibi kelimeler öne çıkmaktadır. Diğer konular da benzer şekilde, "airline", "time", "hours", "staff", "terminal" gibi kelimelerle temsil edilmektedir. Bu bulgular, müşteri yorumlarının belirli ana temalar etrafında toplandığını ve bu temaların hizmet kalitesi, uçuş deneyimi, yiyecek ve

iecek servisi gibi unsurları kapsadığını gstermektedir. Bu analiz, mřteri memnuniyetinin anahtar bileřenlerini anlamak iin nemli igrler saėlamaktadır.



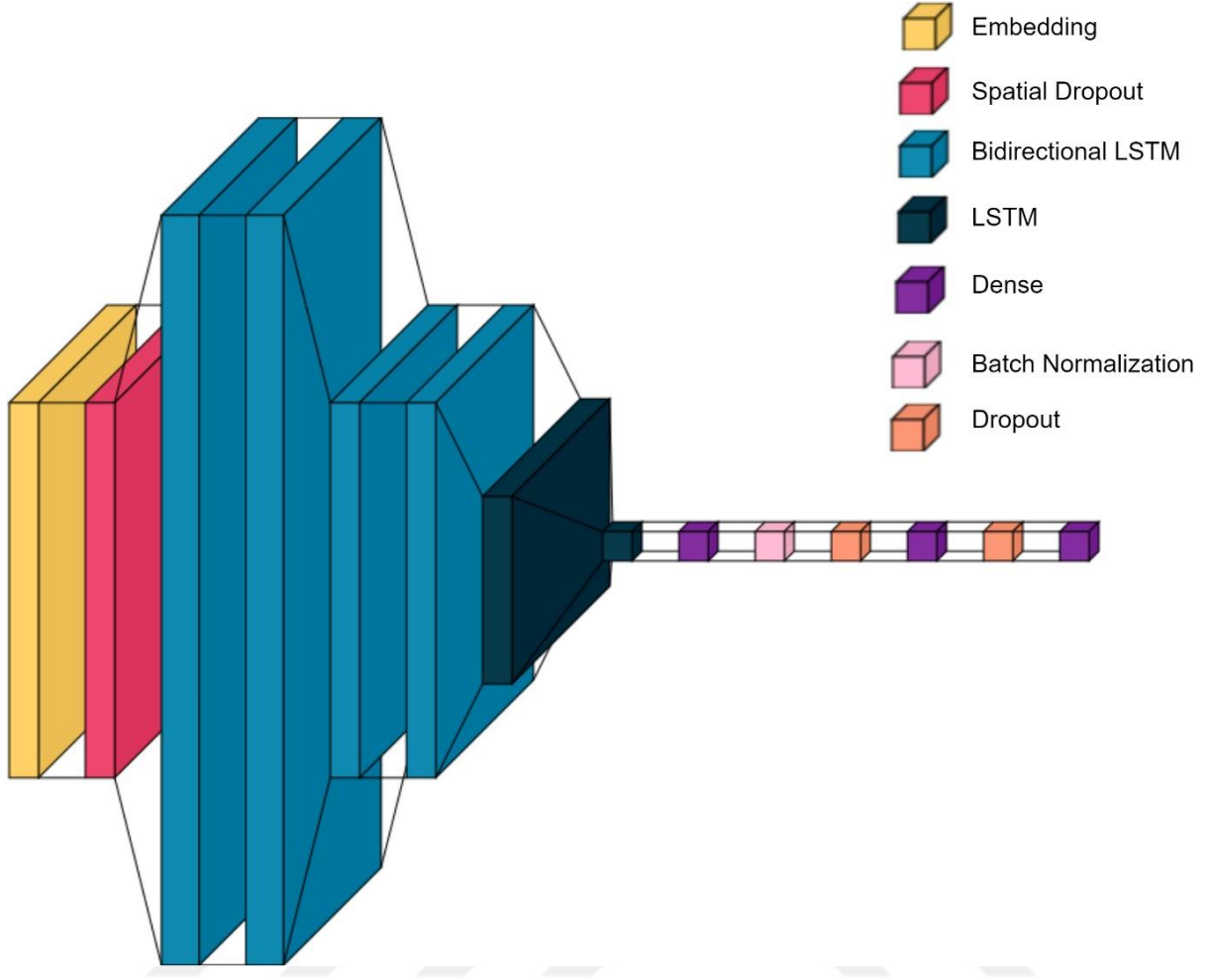
řekil 4.11 Konu bařlıklarının keřfedilmesi

4.6. Duygu Analizi iin CNN tabanlı Model nerisi

4.6.1. Model yapısı ve hiperparametreler

řekil 4.12’de nerilen modele ait blok diyagrama yer verilmiřtir. nerilen model, sıralı (Sequential) bir yapıda olup metin sınıflandırma iin optimize edilmiřtir. nerilen modeldeki blokların kısaca iřlevlerini incelendiėinde:

- **Embedding Katmanı (Sarı):** Bu katman, kelimeleri 128 boyutlu vektrlerle temsil eder. Metnin her kelimesi, belirli bir vektre dnřtrlr ve bu vektrler modelin giriř verisi olarak kullanılır.
- **Spatial Dropout Katmanı (Kırmızı):** Overfitting’i nlemek iin kullanılır. Bu katman, belirli oranlarda girdileri rastgele sıfırlayarak modelin genel performansını artırır.
- **Bidirectional LSTM Katmanları (Aık Mavi):** Bu katmanlar, metnin baėlamsal bilgilerini her iki ynde de iřler. İleri ve geri LSTM’lerin ıktıları birleřtirilerek daha zengin bir temsil elde edilir. Modelde drt adet Bidirectional LSTM katmanı bulunmaktadır.



Şekil 4.12 Önerilen duygu analizi modeli

- **Dense Katmanları (Mor):** Bu katmanlar, modelin öğrendiği özellikleri daha yüksek seviyede temsil eder. İlk dense katman 64 nöron içerir ve ReLU aktivasyon fonksiyonu kullanır. İkinci dense katman ise 32 nöron içerir ve yine ReLU aktivasyon fonksiyonu kullanır.
- **Batch Normalization Katmanı (Pembe):** Bu katman, eğitim sürecini stabilize eder ve eğitim süresini hızlandırır. Dense katmanından sonra gelir ve modelin öğrenme performansını artırır.
- **Dropout Katmanları (Turuncu):** Overfitting'i önlemek için kullanılır. Dense katmanlarından sonra gelir ve belirli oranlarda nöronları rastgele sıfırlar.
- **Çıkış Katmanı (Softmax):** Bu katman, modelin tahminlerini üç sınıfa (pozitif, negatif, nötr) olasılık dağılımı şeklinde verir. Softmax aktivasyon fonksiyonu kullanılarak her sınıf için olasılık hesaplanır.

Model eğitiminde oldukça önemli bir yere sahip olan hiperparametrelere ise Şekil 4.12’de verilmiştir.

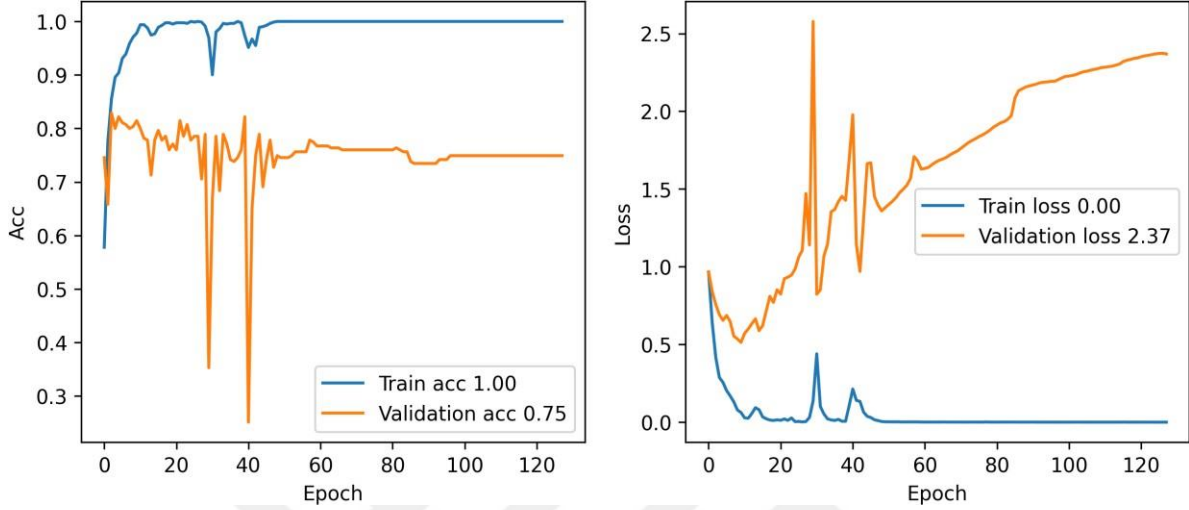
Tablo 4.3 Model eğitiminde kullanılan hiperparametreler ve değerleri

Hiperparametre Adı	Hiperparametrenin Kısa Tanımı	Değeri
Embedding giriş boyutu	Embedding katmanının girdi boyutu	$\text{len}(\text{word_index}) + 1$
Embedding çıkış boyutu	Embedding katmanının çıktı boyutu	128
Embedding giriş uzunluğu	Giriş sekanslarının uzunluğu	100
Spatial Dropout Rate	Spatial Dropout katmanında dropout oranı	0.2
LSTM Units (1)	Birinci Bidirectional LSTM katmanındaki LSTM hücre sayısı	128
LSTM Units (2)	İkinci Bidirectional LSTM katmanındaki LSTM hücre sayısı	128
LSTM Units (3)	Üçüncü Bidirectional LSTM katmanındaki LSTM hücre sayısı	64
LSTM Units (4)	Dördüncü Bidirectional LSTM katmanındaki LSTM hücre sayısı	64
LSTM Units (5)	Beşinci LSTM katmanındaki LSTM hücre sayısı	64
LSTM Units (6)	Altıncı LSTM katmanındaki LSTM hücre sayısı	32
LSTM Dropout Rate	Tüm LSTM katmanlarındaki dropout oranı	0.2
Dense Units (1)	Birinci dense katmanındaki nöron sayısı	64
Dense Units (2)	İkinci dense katmanındaki nöron sayısı	32
Dense Activation	Dense katmanlarındaki aktivasyon fonksiyonu	ReLU
Batch Normalization	Batch Normalization katmanı varlığı	Evet
Dropout Rate (Dense 1)	Birinci dense katmanından sonra uygulanan dropout oranı	0.2
Dropout Rate (Dense 2)	İkinci dense katmanından sonra uygulanan dropout oranı	0.2
Output Activation	Çıkış katmanındaki aktivasyon fonksiyonu	Softmax
Loss Function	Modelin derlenmesinde kullanılan kayıp fonksiyonu	Sparse Categorical Crossentropy
Optimizer	Modelin derlenmesinde kullanılan optimizasyon algoritması	Adam
Metrics	Modelin değerlendirilmesinde kullanılan metrik	Accuracy

Bu çerçevede önerilen model, metin sınıflandırma görevini optimize etmek için dikkatlice seçilmiş hiperparametreler doğrultusunda eğitilmiştir. Tüm veri setinin %75'i eğitim seti olarak kullanılırken, verilerin %25'i ise test kümesi olarak kullanılmıştır.

4.6.2. Model Eğitimi

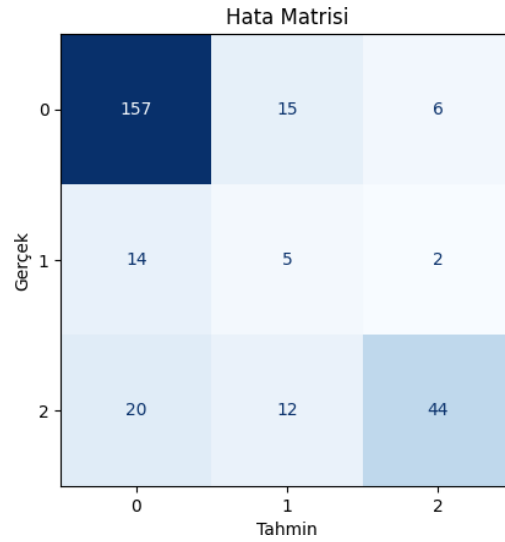
Şekil 4.13’de model eğitime ilişkin olarak doğruluk ve kayıp grafiklerine yer verilmiştir. Bu grafikler, modelin eğitim verisi üzerinde aşırı uyum gerçekleştirdiğini ve doğrulama verisi üzerinde yeterli seviyede genelleme yapamadığını ortaya koymaktadır.



Şekil 4.13 Modele ait doğruluk ve kayıp grafikleri

Grafikte kayıp değerinin test veri seti için epoch ilerledikçe artması, modelin doğrulama verisi üzerinde aşırı uyum (overfitting) gerçekleştirdiğini göstermektedir. Bu durum, modelin eğitim verisindeki örüntüleri aşırı derecede öğrenmesi ve doğrulama verisindeki genel örüntüleri yakalayamaması sonucunda meydana gelir. Aşırı uyum, modelin genelleme yeteneğini ciddi şekilde azaltır ve doğrulama setinde düşük performansla sonuçlanır. Epoch sayısı arttıkça doğrulama kaybının artması, modelin doğrulama verisi üzerindeki hatalarının arttığını ve modelin doğrulama verisi için uygun hale gelmediğini gösterir. Bu durum, modelin genelleme kabiliyetini iyileştirmek için düzenleme teknikleri ve erken durdurma stratejilerinin uygulanması gerektiğine işaret eder.

Şekil 4.14’de eğitim sonucunda elde edilen hata matrisine yer verilmiştir. Hata matrisine ilişkin bulgulara göre, modelin performansı sınıflar arasındaki karışıklıkları göstermektedir. Sıfırıncı sınıfta (Olumsuz) 157 doğru tahmin yapılmış, ancak 15 ve 6 örnek sırasıyla birinci (Nötr) ve ikinci (Olumlu) sınıflara yanlış tahmin edilmiştir. Birinci sınıfta (Nötr sınıf) sadece 5 doğru tahmin yapılırken, 14 örnek sıfırıncı (Olumsuz) ve 2 örnek ikinci sınıfa (Olumlu) yanlış sınıflandırılmıştır. İkinci sınıfta (Olumlu sınıf) ise 44 doğru tahmin yapılmış, ancak 20 ve 12 örnek sırasıyla sıfırıncı (Olumsuz) ve birinci (Nötr) sınıflara yanlış tahmin edilmiştir.

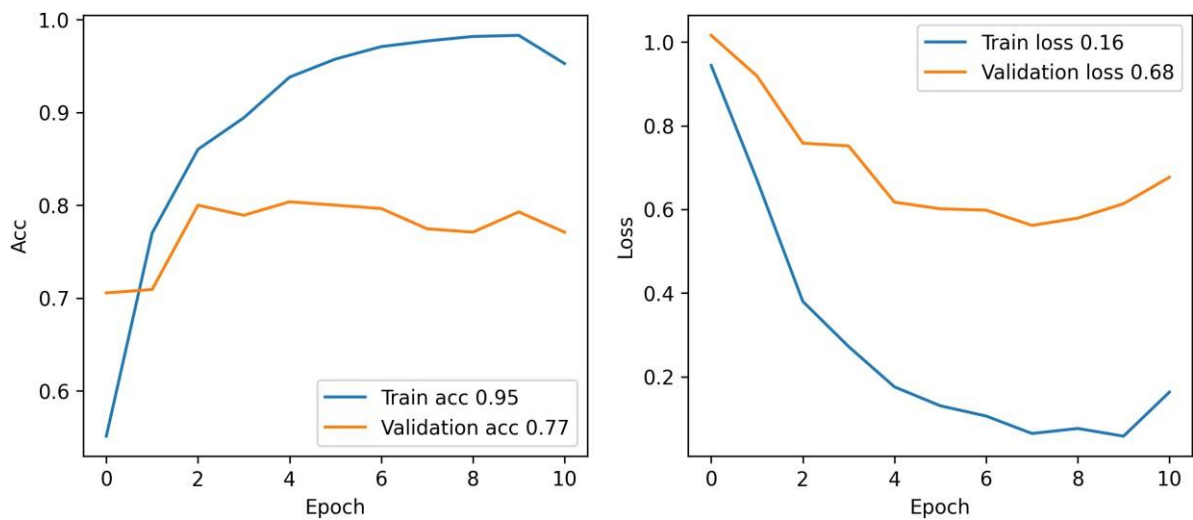


Şekil 4.14 Hata matrisi

Bu sonuçlar, modelin özellikle birinci sınıfı (Nötr) ayırt etmede zorlandığını ve bu sınıf için daha fazla iyileştirme gerektiğini göstermektedir.

4.6.3. Erken Durdurma

Erken durdurma (Early stopping), model eğitimi sırasında doğrulama kaybı veya doğrulama doğruluğu gibi bir performans metriğinde iyileşme durduğunda eğitimi erken durduran bir tekniktir. Bu yöntem, aşırı uyum (overfitting) riskini azaltarak modelin genelleme yeteneğini artırmayı hedefler.

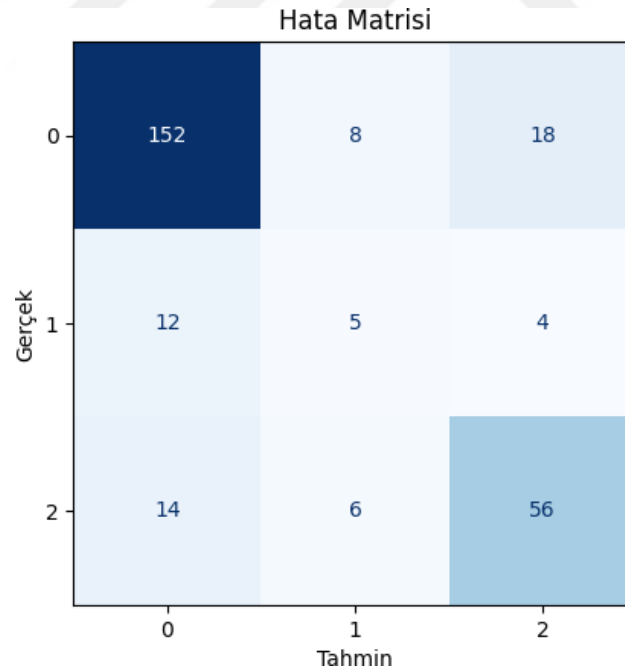


Şekil 4.15 Erken durdurma yaklaşımı ile eğitim

Deneysel çalışmada uygulanan erken durdurma yaklaşımı, modelin doğrulama kaybı (val_loss) metriğini izlemektedir ve bu metrik belirli bir süre (patience=3) boyunca iyileşmezse eğitimi durdurur. Bu yaklaşım, modelin aşırı uyum yapmasını engeller ve

genelleme yeteneğini artırır. Eğitimin durdurulması durumunda, modelin ağırlıkları en iyi doğrulama performansını gösterdiği ana geri yüklenir (restore_best_weights=True). Bu, modelin doğrulama setindeki en iyi performansını korumasını sağlar ve gereksiz yere fazla eğitimin önüne geçer. Bu yaklaşım, eğitim sürecini daha verimli hale getirir ve aşırı uyum riskini minimize eder.

Şekil 4.15’de erken durdurma yaklaşımı ile elde edilen doğruluk ve kayıp eğrilerine yer verilmiştir. Grafiklere göre, eğitim doğruluğu (train accuracy) hızla artarak 0.95 seviyesine ulaşmış ve bu seviyede sabitlenmiştir. Ancak, doğrulama doğruluğu (validation accuracy) yaklaşık 0.77 seviyesinde kalmış ve dalgalanmalar göstermektedir. Bu durum, modelin eğitim verisi üzerinde iyi performans gösterdiğini, ancak doğrulama verisinde genelleme yapmada zorluk çektiğini işaret eder. Eğitim kaybı (train loss) sürekli olarak azalmış ve 0.16 seviyesine düşmüştür, bu da modelin eğitim verisi üzerinde hatasız çalıştığını gösterir. Buna karşılık, doğrulama kaybı (validation loss) 0.68 seviyesinde kalmış ve belirli bir noktadan sonra artmaya başlamıştır. Bu bulgular, modelin aşırı uyum (overfitting) yaptığını ve doğrulama verisi üzerinde genelleme yeteneğinin sınırlı olduğunu göstermektedir. Erken durdurma yaklaşımı, bu durumu önlemek için etkili olmuştur.



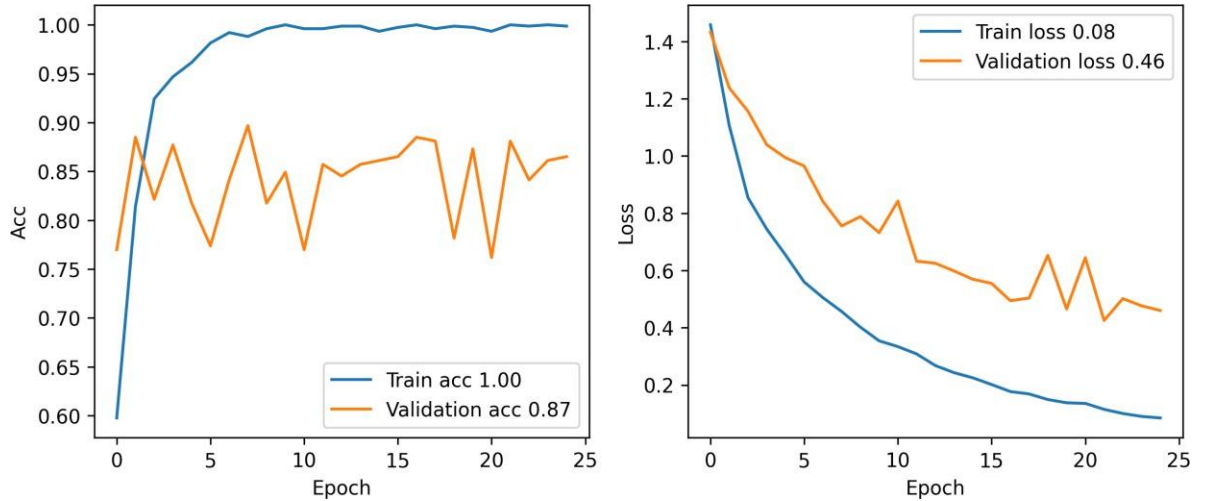
Şekil 4.16 Erken durdurma yaklaşımı ile elde edilen hata matrisi

Birinci deneysel çalışmada 0.75 olarak elde edilen doğruluk değeri, ikinci deneysel çalışma sonucunda 0.77’ye ulaşmıştır.

4.6.4. İkili Sınıflandırma

Bu deneysel çalışmada, veri setinde nötr örneklerin öğrenme için yeterli sayıda olmadığı dikkate alınmıştır. Nötr örneklerin sayısal yetersizliği, modelin bu sınıfları doğru şekilde öğrenmesini ve genelleme yapmasını zorlaştırmaktadır. Bu nedenle, bu çalışma kapsamında ikili sınıflandırma hedeflenmiş ve nötr örnekler analiz dışı bırakılmıştır. Bu yaklaşım, modelin pozitif ve negatif sınıflar üzerinde daha doğru ve etkili sonuçlar üretmesini sağlamak amacıyla tercih edilmiştir. Böylelikle, modelin eğitim ve değerlendirme süreçlerinde daha dengeli ve güvenilir performans göstermesi beklenmektedir.

Ayrıca yeni modelde, aşırı uyum sorununu önlemek ve modelin genelleme yeteneğini artırmak için bir dizi optimizasyon yaklaşımı uygulanmıştır. İlk olarak, dropout oranları artırılarak (0.2'den 0.3'e yükseltilmiştir) nöronların rastgele devre dışı bırakılması sağlanmış, böylece modelin aşırı uyum yapması engellenmiştir. İkinci olarak, L2 regularizasyonu eklenmiş ve bu sayede modelin ağırlıklarına ceza uygulanarak aşırı uyum riskinin azaltılması hedeflenmiştir. Ayrıca, ikinci deneysel çalışmada olduğu gibi erken durdurma yaklaşımı kullanılarak modelin doğrulama kaybı belirli bir süre iyileşmediğinde eğitim süreci durdurulmuş ve en iyi ağırlıkların geri yüklenmesi sağlanmıştır. Son olarak, ReduceLROnPlateau kullanılarak doğrulama kaybı iyileşmediğinde öğrenme oranı dinamik olarak azaltılmış, böylece modelin daha iyi genelleme yapması desteklenmiştir. Bu optimizasyon teknikleri, modelin performansını ve genelleme yeteneğini artırmayı amaçlamaktadır.

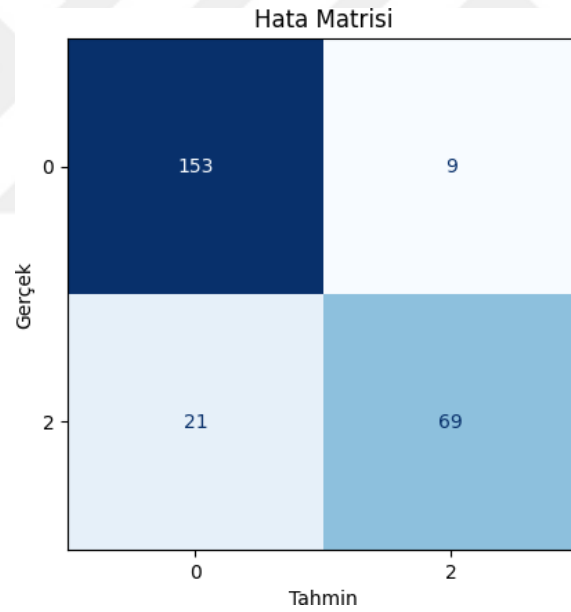


Şekil 4.17 İkili sınıflandırma için doğruluk ve kayıp eğrileri

Şekil 4.17’de yeni modele ilişkin elde edilen doğruluk ve kayıp eğrilerine yer verilmiştir. Grafiklere göre, eğitim doğruluğu hızla 1.00 seviyesine ulaşmış ve bu seviyede sabit kalmıştır, bu da modelin eğitim verisi üzerinde mükemmel bir performans gösterdiğini göstermektedir. Ancak, test doğruluğu yaklaşık 0.87 seviyesinde dalgalanmalar göstermektedir, bu durum modelin doğrulama verisi üzerinde genelleme yapmada hala zorluk çektiğini işaret eder.

Eğitim kaybı sürekli olarak azalmış ve 0.08 seviyesine düşmüştür, bu da modelin eğitim verisi üzerinde neredeyse hatasız çalıştığını gösterir. Buna karşılık, doğrulama kaybı belirgin bir düşüş göstermiş ve 0.46 seviyesine inmiştir, ancak doğrulama kaybı eğitimin ilerleyen aşamalarında dalgalanma eğilimi göstermektedir. Bu durum, modelin doğrulama verisi üzerinde daha iyi performans gösterdiğini, ancak hala bazı dalgalanmalar ve aşırı uyum (overfitting) belirtileri olabileceğini göstermektedir.

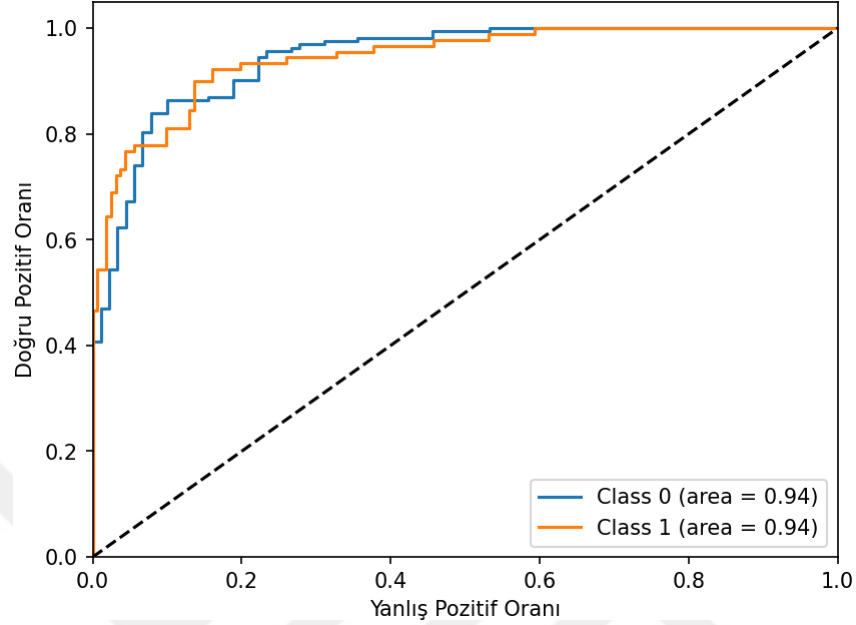
Genel olarak, modelin eğitim verisi üzerinde mükemmel performans sergilediği, ancak doğrulama verisinde genelleme yeteneğinin hala iyileştirilmesi gerektiği görülmektedir.



Şekil 4.18 İkili sınıflandırma için hata matrisi

Şekil 4.18’de ise ikili sınıflandırma için elde edilen hata matrisine yer verilmiştir. Bu hata matrisine göre, modelin olumsuz (0) ve olumlu (2) sınıfları ayırt etme performansı analiz edilmiştir. Olumsuz sınıfta (0), 153 örnek doğru olarak sınıflandırılmış, ancak 9 örnek yanlış bir şekilde olumlu (2) sınıfa atanmıştır. Bu, modelin olumsuz sınıfı oldukça iyi ayırt ettiğini göstermektedir. Olumlu sınıfta (2) ise 69 örnek doğru olarak tanımlanmış, ancak 21 örnek yanlış bir şekilde olumsuz (0) olarak tahmin edilmiştir. Bu, modelin olumlu sınıfı ayırt etmede daha fazla zorlandığını ve bu sınıfta daha fazla hata yaptığını ortaya koymaktadır. Genel

olarak, model olumsuz sınıfı daha yüksek doğrulukla tahmin ederken, olumlu sınıf için genelleme yeteneğinin iyileştirilmesi gerekmektedir. Bu durumun temel nedeni veri setindeki sınıflar arası dengesiz dağılımdır.



Şekil 4.19 ROC eğrisi

Şekil 4.19’da ROC eğrisi sunulmuştur. Modelin her iki sınıf için de AUC değerleri 0.94 olarak hesaplanmıştır, bu da modelin hem Olumsuz hem de Olumlu sınıfları ayırt etmede yüksek performans gösterdiğini ifade eder. ROC eğrisinin 45 derece çizgisine olan uzaklığı, modelin rastgele tahmin yapmaktan çok daha iyi performans sergilediğini ortaya koyar. Bu sonuçlar, modelin genel olarak sınıflandırma görevinde oldukça etkili olduğunu göstermektedir.

4.6.5. Karşılaştırmalı Sonuçlar

Bu bölümde üç farklı deneysel çalışma gerçekleştirilmiş ve deneysel çalışmaların sonuçları çeşitli performans metrikleri dikkate alınarak değerlendirilmiştir. Deneyler boyunca doğruluk, hata oranı, duyarlılık, özgüllük, keskinlik, yanlış pozitif oranı ve F1 skoru gibi önemli metrikler kullanılmıştır. İlk deneyde elde edilen sonuçlar, ikinci ve üçüncü deneylerle karşılaştırılarak modelin performansındaki iyileşmeler gözlemlenmiştir. Üçüncü deneyde, modelin doğruluk ve duyarlılık gibi metriklerde en yüksek değerlere ulaşması, yapılan optimizasyon ve iyileştirme çalışmalarının etkisini açıkça göstermektedir. Bu bölümde sunulan karşılaştırmalı sonuçlar, modelin genelleme yeteneği ve doğrulama verisi üzerindeki performansını daha iyi anlamamıza olanak tanımaktadır. Sonuçlar Tablo 4.4’de sunulmuştur.

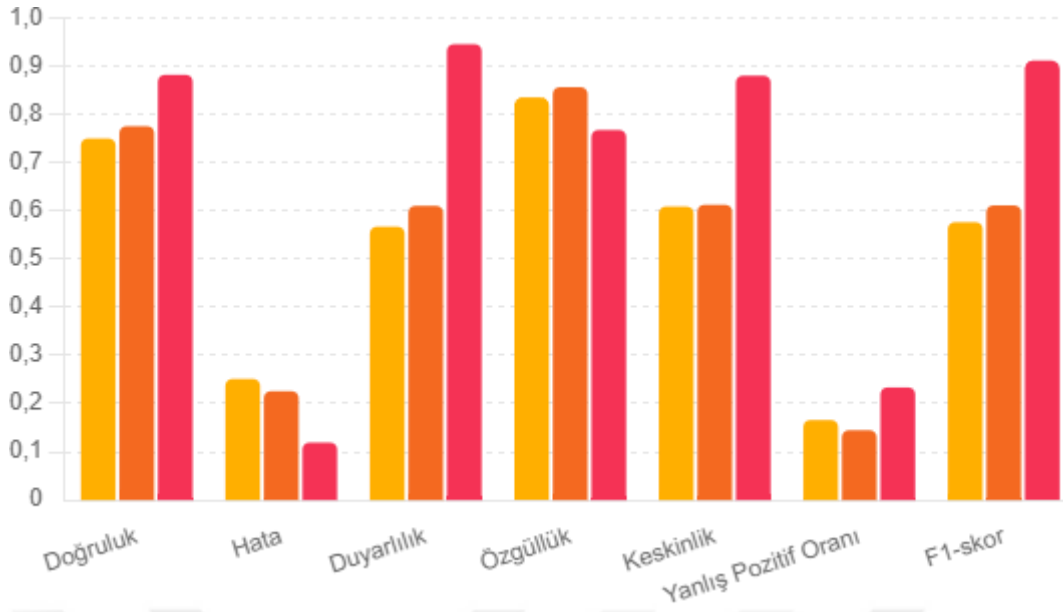
Tablo 4.4 Karşılaştırmalı performans sonuçları

Metrikler	#1	#2	#3
Doğruluk	0.7491	0.7745	0.8810
Hata	0.2509	0.2255	0.1190
Duyarlılık	0.5664	0.6096	0.9444
Özgüllük	0.8343	0.8554	0.7667
Keskinlik	0.6081	0.6117	0.8793
Yanlış Pozitif Oranı	0.1657	0.1446	0.2333
F1-skor	0.5757	0.6104	0.9107

Her bir deney, modelin performansını değerlendiren doğruluk, hata oranı, duyarlılık, özgüllük, keskinlik, yanlış pozitif oranı ve F1 skoru gibi çeşitli metrikler üzerinden analiz edilmiştir. İlk deneyde modelin doğruluk oranı %74.91 iken, ikinci deneyde bu oran %77.45'e ve üçüncü deneyde %88.10'a yükselmiştir. Benzer şekilde, duyarlılık metrikleri de sırasıyla %56.64, %60.96 ve %94.44 olarak artış göstermiştir, bu da modelin pozitif sınıfları ayırt etme yeteneğinin önemli ölçüde iyileştiğini göstermektedir.

Özgüllük metrikleri incelendiğinde, ikinci deneyde %85.54 ile en yüksek değere ulaşılmış, ancak üçüncü deneyde %76.67'ye düşmüştür. Bu durum, modelin negatif sınıfları ayırt etme yeteneğinin ikinci deneyde daha iyi olduğunu, ancak genel doğruluk ve duyarlılık açısından üçüncü deneyin daha üstün olduğunu göstermektedir. Keskinlik ve F1 skoru gibi diğer metrikler de üçüncü deneyde belirgin bir iyileşme sergilemiştir; özellikle F1 skoru %91.07'ye yükselmiştir.

DeneySEL çalışmalarındaki bu farklılıklar, uygulanan optimizasyon teknikleri ve hiperparametre ayarlarının model performansı üzerindeki etkisini ortaya koymaktadır. Üçüncü deneyde kullanılan optimizasyon stratejileri, modelin genelleme yeteneğini artırarak daha dengeli ve yüksek performanslı bir sonuç elde edilmesini sağlamıştır. Bu bulgular, modelin hem akademik hem de pratik uygulamalarda etkinliğini artırmak için gerekli olan iyileştirme ve optimizasyon süreçlerinin önemini vurgulamaktadır. Şekil 4.20'de model performansındaki iyileşme açıkça görülmektedir.



Şekil 4.20 Deneyisel çalışmalara ait karşılaştırmalı sonuçlar

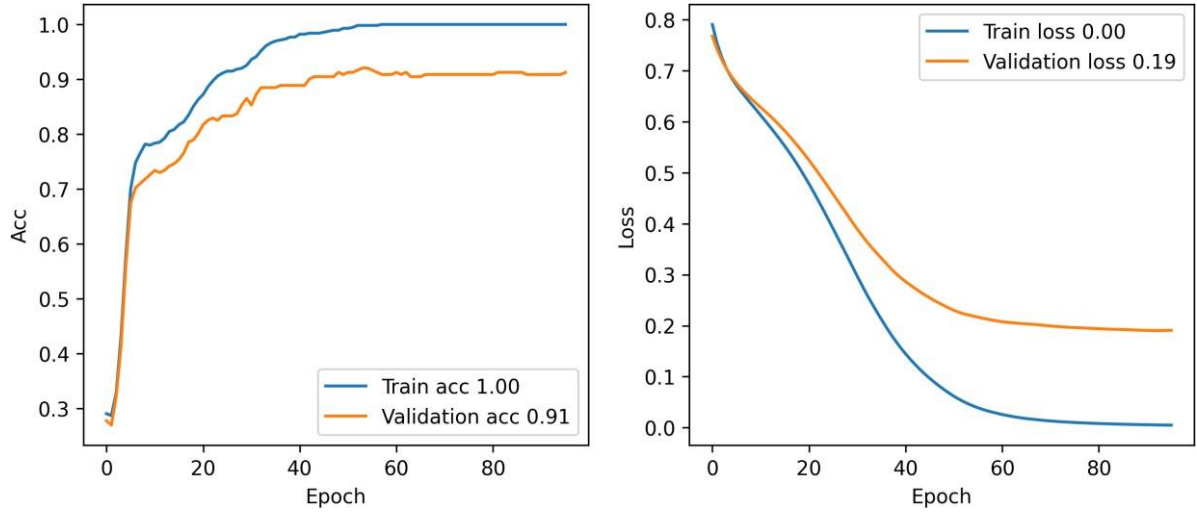
4.7. Transfer Öğrenme ile Duygu Analizi

Bu bölümde, transfer öğrenme yaklaşımı kullanılarak geniş veri setleri üzerinde önceden eğitilmiş modellerin yeniden eğitimi gerçekleştirilmiştir. Bu sayede, geniş veri setleri ile eğitilmiş olan modellerin, mevcut problem alanına uyarlanması sağlanmıştır. Transfer öğrenme, büyük ve çeşitli veri setlerinde eğitilmiş modellerin, belirli bir görev veya veri kümesi üzerinde yeniden kullanılarak, performansın artırılmasına olanak tanır. Bu yöntem, özellikle sınırlı veri setleri için etkili olup, mevcut modelin genelleme yeteneğini ve doğruluğunu önemli ölçüde iyileştirmektedir.

4.7.1. nnlm-en-dim50

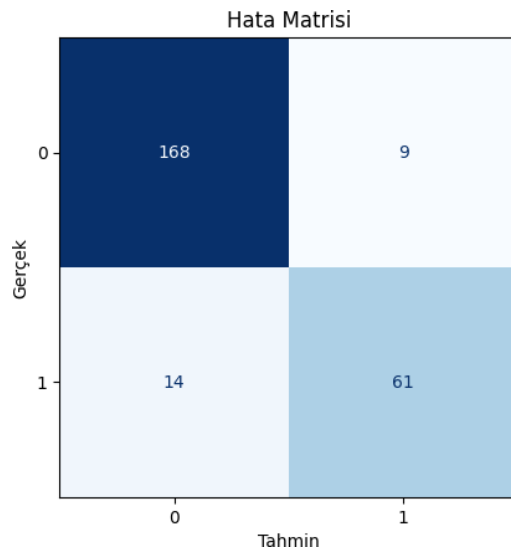
Google'ın TensorFlow Hub'da sağladığı "**nnlm-en-dim50**" modeli, İngilizce dilinde 50 boyutlu bir metin gömme (embedding) modelidir. Bu model, metin verilerini 50 boyutlu yoğun (dense) vektörlere dönüştürerek NLP görevlerinde kullanılır. Model, geniş bir metin veri kümesi üzerinde eğitilmiştir ve kelime anlamlarını temsil eden vektörler oluşturur. Bu vektörler, metin sınıflandırma, duygu analizi ve benzeri NLP görevlerinde kullanılabilir. Model, yüksek doğruluk ve genelleme yeteneği sunar, bu nedenle çeşitli NLP uygulamalarında tercih edilmektedir. Modelin URL adresi ³ paylaşılmıştır.

³ <https://tfhub.dev/google/nnlm-en-dim50/2>. (Erişim Tarihi:29.07.2024)



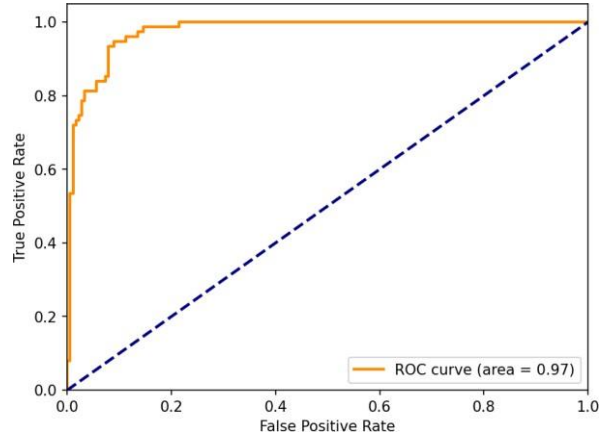
Şekil 4.21 nnlm-en-dim50 modeline ilişkin doğruluk ve kayıp eğrileri

Şekil 4.21’de, modelin eğitim ve doğrulama performansını göstermektedir. Sol grafikte, eğitim doğruluğu hızla artarak 1.00 seviyesine ulaşmış ve sabitlenmiştir, bu modelin eğitim verisi üzerinde mükemmel performans sergilediğini gösterir. Doğrulama doğruluğu ise 0.91 seviyesinde kalmış ve küçük dalgalanmalar göstermiştir, bu modelin genelleme yeteneğinin iyi olduğunu fakat eğitim verisi kadar yüksek olmadığını gösterir. Sağ grafikte, eğitim kaybı hızla azalarak 0.00 seviyesine düşmüştür, bu modelin eğitim verisinde hatasız çalıştığını gösterir. Doğrulama kaybı ise 0.19 seviyesinde kalmış ve sabitlenmiştir, bu modelin doğrulama verisinde makul bir performans sergilediğini gösterir.



Şekil 4.22 nnlm-en-dim50 hata matrisi

Modele ait hata matrisine Şekil 4.22’de yer verilmiştir. Görüldüğü üzere sıfırdan eğitilmiş modellere kıyasla önceden eğitilmiş modelin genelleştirme performansı oldukça yüksektir.

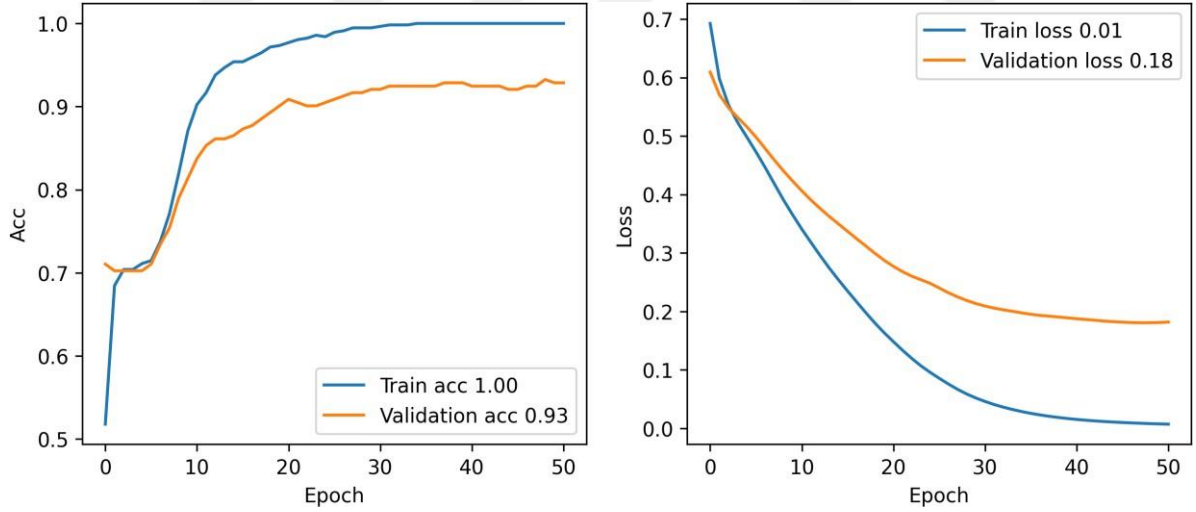


Şekil 4.23 nnlm-en-dim50 ROC eğrisi

Şekil 4.23’de modele ait ROC eğrisine yer verilmiştir. Eğrinin altındaki alan AUC değeri 0.97 olarak elde edilmiştir.

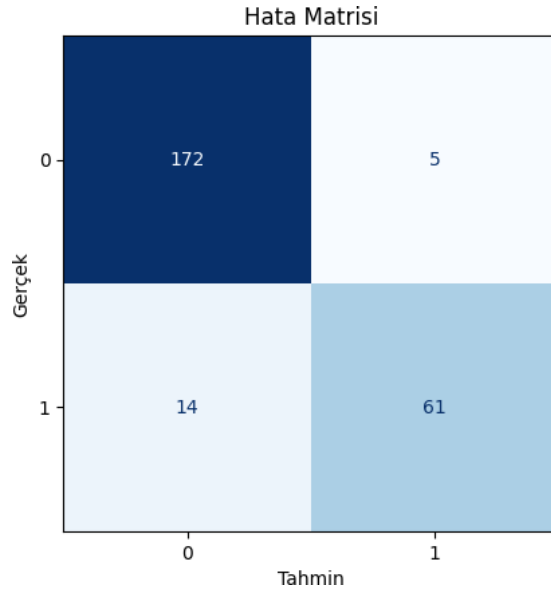
4.7.2. nnlm-en-dim128

Bu model, metin verilerini 128 boyutlu yoğun (dense) vektörlere dönüştürerek NLP görevlerinde kullanılır. nnlm-en-dim128 modeli de yine transfer öğrenme yaklaşımı ile eğitilmiştir. Modele ait doğruluk ve kayıp eğrilerine Şekil 4.24’de yer verilmiştir.

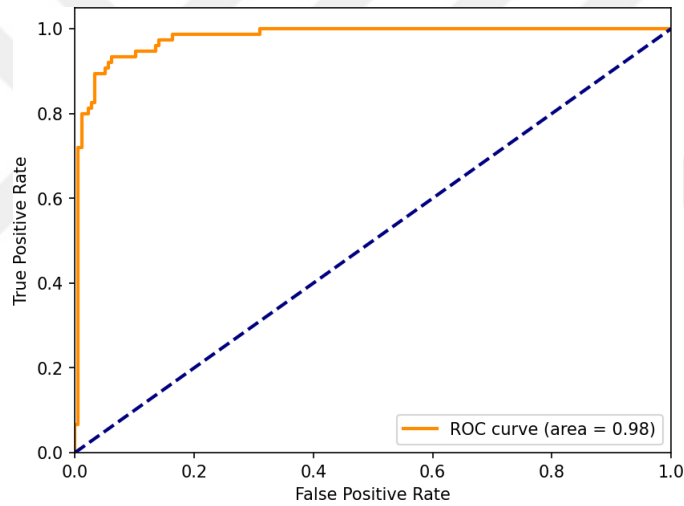


Şekil 4.24 nnlm-en-dim128 modeline ait doğruluk ve kayıp eğrileri

Eğitim ve doğrulama eğrileri incelendiğinde modelin yeterli seviyede bir genelleştirme performansı sağladığı görülmektedir. Modele ait hata matrisine Şekil 4.25’de yer verilmiştir. Benzer şekilde model ait ROC eğrisi Şekil 4.26’da sunulmuştur. Modelin %93 doğruluğa 0.98 AUC değerine ulaştığı görülmüştür.



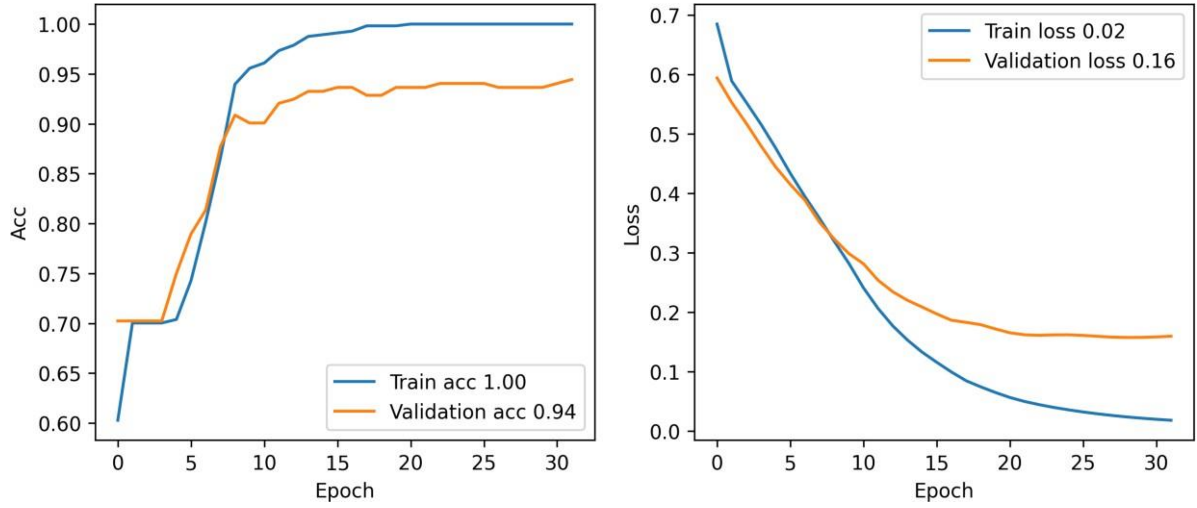
Şekil 4.25 nnlm-en-dim128 modeline ait hata matrisi



Şekil 4.26 nnlm-en-dim128 modeline ait ROC eğrisi

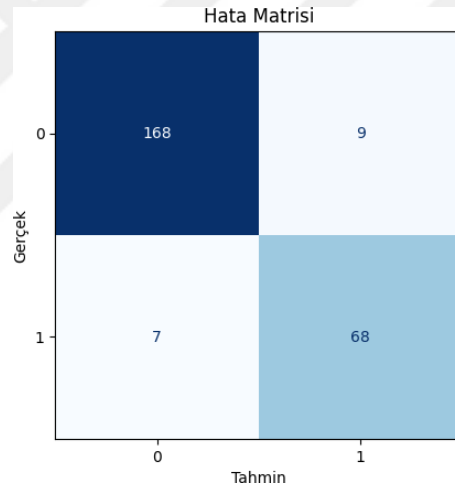
4.7.3. USE Modeli

Universal Sentence Encoder (USE), Google tarafından geliştirilen ve TensorFlow Hub'da sunulan bir cümle gömme modelidir. Model, metin verilerini 512 boyutlu vektörlere dönüştürerek NLP görevlerinde kullanılmasını sağlar. USE, cümleler ve paragraflar arasındaki anlamsal benzerlikleri yakalamak için eğitilmiştir ve geniş bir metin veri kümesi üzerinde eğitilmiştir. Bu model, çeşitli NLP görevlerinde, örneğin metin sınıflandırma, duygu analizi ve bilgi getirme işlemlerinde yüksek performans sunar. Anahtar özellikleri arasında yüksek doğruluk, genelleme yeteneği ve kullanım kolaylığı bulunur. USE, transfer öğrenme yaklaşımıyla, farklı dil uygulamalarında etkili sonuçlar elde etmeyi mümkün kılar.



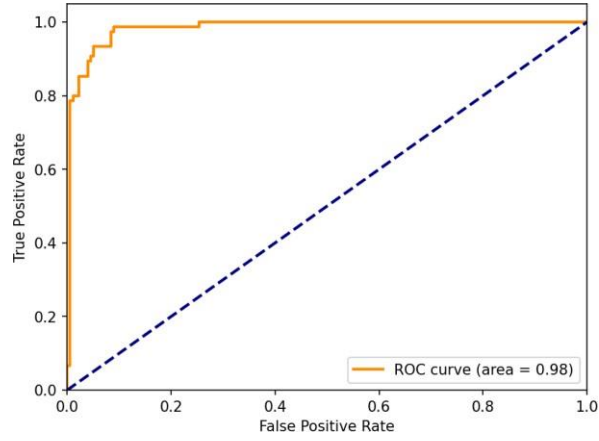
Şekil 4.27 USE modele ait doğruluk ve kayıp eğrileri

Şekil 4.27’de bu modelin transfer öğrenme yaklaşımı ile yeniden eğitilmesine ilişkin doğruluk ve kayıp eğrilerine yer verilmiştir.



Şekil 4.28 USE modele ait hata matrisi

Şekil 4.28’de ise elde edilen hata matrisi sunulmuştur. Modelin 0.94 sınıflandırma doğruluğu sağlandığı görülmüştür. ROC eğrisi ise Şekil 4.29’da sunulmuştur ve AUC değeri 0.98 olarak elde edilmiştir.



Şekil 4.29 USE modele ait ROC eğrisi

4.7.4. Karşılaştırmalı Sonuçlar

Tablo 4.5 Transfer öğrenmeye ilişkin karşılaştırmalı sonuçlar

Metrikler	nnlm-en-dim50	nnlm-en-dim128	USE
Doğruluk	0.9087	0.9246	0.9365
Hata	0.0913	0.0754	0.0635
Duyarlılık	0.9492	0.9718	0.9492
Özgüllük	0.8133	0.8133	0.9067
Keskinlik	0.9231	0.9247	0.9600
Yanlış Pozitif Oranı	0.1867	0.1867	0.0933
F1-skor	0.9359	0.9477	0.9545

Tablo 4.5’de transfer öğrenme yaklaşımı ile eğitilen modellerin performans sonuçları raporlanmıştır. Transfer öğrenme yaklaşımı ile eğitilen modellerinin karşılaştırmalı performansları gösterilmiştir. USE modeli, %93.65 doğruluk ve %6.35 hata oranıyla en yüksek performansı sergilemiştir. Duyarlılık ve F1-skor değerleri sırasıyla %94.92 ve %95.45 olup, bu modelin pozitif sınıfları doğru bir şekilde tanımlamada üstün olduğunu göstermektedir. Özgüllük değeri %90.67 ile en yüksek olup, yanlış pozitif oranı %9.33 ile en düşüktür. nnlm-en-dim128 modeli de yüksek performans sergilemiş, %92.46 doğruluk, %97.18 duyarlılık ve %94.77 F1-skor değerleriyle dikkat çekmiştir. nnlm-en-dim50 modeli ise bu iki modelin gerisinde kalmıştır.

BEŞİNCİ BÖLÜM

5. Sonuçlar, Tartışma ve Öneriler

5.1. Literatürdeki ilgili Çalışmaların Karşılaştırılması

Bu bölümde literatürde yer alan ilgili çalışmaların kullanılan yöntemler, veri setleri ve deneysel farklılıkları dikkate alınarak bir karşılaştırma sunulmuştur.

Tablo 5.1 İlgili çalışmaların karşılaştırılması

Metotlar	Veri setinin kısa tanımı	Örnek sayısı	Sınıf sayısı	Erişilebilirlik	Doğruluk
CNN ve Word2Vec Metodu (Acı ve Çırak 2019)	TTC-3600	3600	6	Erişime Açık	%93,3
CNN + TF-IDF + DK Filtre + Kök Bulma (Alparslan ve Dursun 2023)	TTC-4900	4900	7	Erişime Açık ⁴	%91,7
	MY-15130	15170	2	Erişime Açık ⁵	%94
LSTM (Alzoubi vd. 2023)	Özel veri seti	225239	14	Erişime Açık Değil	%84
FastText + SVM (Çelik ve Koç 2021)	Özel veri seti	12000	6	Erişime Açık Değil	%95,75
Knn (Chen vd. 2020)	Özel Veri Seti	2411	7	Erişime Açık değil	%71,4

⁴ <https://www.kaggle.com/datasets/savasy/ttc4900> (Erişim Tarihi: 29.07.2024)

⁵ <https://www.kaggle.com/datasets/mujdatcabuk/eticaret-urun-yorumlari/> (Erişim Tarihi: 29.07.2024)

BERT (Gasparetto vd. 2022)	EnWiki-100	329600	100	Erişime açık ⁶	%75,28
	RCV1-57	758100	57	Erişime açık	%73,48
SMO (Göker ve Tekedere 2017)	Özel Veri Seti	444	2	Erişime açık değil	%88,73
Adaboost + TF IDF (Hajibabaei vd. 2022)	Özel Veri Seti	24783	3	Erişime Açık değil	%94
Bi LSTM (Jang vd. 2020)	IMDB	50000	2	Erişime açık⁷	%91,41
Knn (Kaşıkçı ve Gökçen 2013)	Özel veri seti	279	2	Erişime açık değil	%83,87
RCNN (Lai vd. 2015)	20News	13919	4	Erişime açık	%96,49
	Fudan	19636	20	Erişime açık	%95,20
	ACL	202979	5	Erişime açık	%49,19
	SST	11855	5	Erişime açık ⁸	%47,21
RT + Bag of Words (Neogi vd. 2021)	Özel Veri Seti	18000	3	Erişime açık değil	%96,62

⁶ <https://gitlab.com/distratn/dsi-nlp-publib> (Erişim Tarihi: 29.07.2024)

⁷ <https://www.kaggle.com/datasets/lakshmi25npathi/imdb-dataset-of-50k-movie-reviews> (Erişim Tarihi: 29.07.2024)

⁸ <https://nlp.stanford.edu/sentiment/> (Erişim Tarihi: 29.07.2024)

TextConvoNet (Soni vd. 2023)	Binary SST-2 Dataset	8732	2	Erişime açık ⁹	%82,20
	Amazon Review for Sentiment Analysis Dataset	20000	2	Erişime açık ¹⁰	%90,90
	Reuters-21578	7112	8	Erişime açık ¹¹	%99,20
	Twitter User Airline Sentiment Dataset	13000	3	Erişime açık ¹²	%88,40
	Covid Tweets Dataset	7500	5	Erişime açık ¹³	%83,90
MLP (Tuzcu 2020)	Özel Veri Seti	1400	2	Erişime açık değil	%86
CRNN (Wang vd. 2019)	Paper Data Set 1	200000	5	Erişime açık	%98,45

⁹ <https://www.kaggle.com/datasets/jgggikmf/binary-sst2-dataset> (Erişim Tarihi: 29.07.2024)

¹⁰ <https://www.kaggle.com/datasets/bittlingmayer/amazonreviews> (Erişim Tarihi: 29.07.2024)

¹¹ <https://www.kaggle.com/datasets/weipengfei/ohr8r52> (Erişim Tarihi: 29.07.2024)

¹² <https://www.kaggle.com/datasets/crowdfower/twitter-airline-sentiment> (Erişim Tarihi: 29.07.2024)

<https://www.kaggle.com/datasets/datatattle/covid-19-nlp-text-classification> (Erişim Tarihi: 29.07.2024)

¹³ <https://www.kaggle.com/datasets/datatattle/covid-19-nlp-text-classification> (Erişim Tarihi: 29.07.2024)

	Paper Data Set 2	400000	10	Erişime açık	%95,55
	AG's news corpus	127600	4	Erişime açık	%93,38
	Sogou news corpus	510000	5	Erişime açık	%96,57
	DBPedia ontology dataset	630000	14	Erişime açık	%98,89
	Yelp Review Full	700000	5	Erişime açık	%92,66
	Yahoo! Answers dataset	1460000	10	Erişime açık	%72,53

Tablo 5.1’de verilen çalışmalarda kullanılan farklı veri setleri, sınıflandırma modellerinin performansını önemli ölçüde etkileyebilmektedir. Örneğin, erişime açık TTC-3600 ve EnWiki-100 veri setleri, büyük ve dengeli veri yapıları sayesinde yüksek doğruluk oranları sunarken, özel veri setleri genellikle daha sınırlı ve dengesiz yapılar içerdiğinden model performansını olumsuz etkileyebilir. Alzoubi vd. (2023) tarafından kullanılan 225239 örnekli özel veri seti, sınıf sayısının fazlalığı ve verinin çeşitliliği nedeniyle %84 doğruluk oranı sağlarken, Gasparetto vd. (2022) tarafından kullanılan büyük ölçekli EnWiki-100 veri seti %75,28 doğruluk oranına ulaşmıştır. Benzer şekilde, IMDB ve Reuters-21578 gibi erişime açık veri setleri, geniş örneklem ve kapsamlı sınıflar içermeleri sayesinde daha güvenilir performans ölçütleri sunar. Bu bağlamda, veri setlerinin erişilebilirliği, büyüklüğü ve heterojenliği,

sınıflandırma modellerinin genelleme yeteneğini ve doğruluk oranlarını doğrudan etkilemektedir.

Tablo 5.1'deki çalışmalarda kullanılan makine öğrenmesi, derin öğrenme ve hibrit yöntemlerin her birinin kendine özgü avantajları ve dezavantajları bulunmaktadır. Makine öğrenmesi yöntemleri, özellikle SVM ve Naive Bayes gibi klasik algoritmalar, daha küçük veri setlerinde ve düşük hesaplama maliyetlerinde yüksek performans sunar; ancak, büyük ve karmaşık veri setlerinde sınırlı kalabilirler. Derin öğrenme yöntemleri, CNN ve LSTM gibi algoritmalar, büyük veri setlerinde üstün performans ve genelleme yeteneği sağlar, ancak yüksek hesaplama gücü ve uzun eğitim süreleri gerektirir. Hibrit yöntemler, örneğin CRNN ve Attention mekanizmalarıyla zenginleştirilmiş modeller, her iki yaklaşımın da avantajlarını birleştirerek daha yüksek doğruluk oranları sunar, ancak bu modellerin tasarımı ve optimizasyonu karmaşık olabilir. Bu bağlamda, Alparslan ve Dursun'un (2023) CNN+TF-IDF hibriti %94 doğruluk sağlarken, Jang vd. (2020) BiLSTM modeli %91,41 doğruluk oranına ulaşmıştır. Bu yöntemlerin seçiminde veri setinin boyutu, çeşitliliği ve problemin doğası belirleyici rol oynamaktadır.

Tablo 5.1'deki performans metrikleri, kullanılan veri setlerinin boyutu, sınıf sayısı ve erişilebilirliği ile yakından ilişkilidir. Örneğin, CRNN (Wang vd. 2019) modelinin geniş ölçekli veri setleri üzerinde elde ettiği yüksek doğruluk oranları (%98,45 ve %95,55), büyük veri kümelerinin derin öğrenme modellerinin genelleme yeteneğini artırdığını göstermektedir. Buna karşılık, daha küçük ve özel veri setleriyle çalışan yöntemler, veri kısıtlamaları nedeniyle genellikle daha düşük doğruluk oranlarına ulaşmaktadır. Örneğin, Chen vd. (2020) tarafından kullanılan KNN modeli %71,4 doğruluk oranı ile daha düşük performans sergilemiştir. Ayrıca, erişime açık veri setleri, genellikle daha geniş kullanıcı toplulukları tarafından optimize edildikleri için daha yüksek doğruluk oranları sağlar; bu durum, BERT (Gaspardo vd. 2022) modelinin EnWiki-100 veri seti ile elde ettiği %75,28 doğruluk oranında gözlemlenebilir. Sonuç olarak, veri setlerinin büyüklüğü, çeşitliliği ve erişilebilirliği, kullanılan modellerin performansını doğrudan etkilemektedir ve bu faktörler göz önünde bulundurularak metrikler değerlendirilmelidir.

Çalışmalardaki metodlar, veri setleri ve parametrelerdeki farklılıklar nedeniyle bire bir karşılaştırma yapmak zordur. Her çalışma, belirli veri setleri ve özel parametre ayarları kullanılarak gerçekleştirilmiştir, bu da sonuçların doğrudan karşılaştırılabilirliğini kısıtlamaktadır. Örneğin, Acı ve Çırak (2019) TTC-3600 veri seti ile %93,3 doğruluk oranı elde ederken, Alzoubi vd. (2023) özel veri seti kullanarak %84 doğruluk oranına ulaşmıştır. Bu

farklılıklar, veri setlerinin boyutu, sınıf sayısı ve veri kalitesi gibi faktörlerden kaynaklanabilir. Ayrıca, kullanılan modellerin türü (makine öğrenmesi, derin öğrenme veya hibrit) ve bu modellerin eğitiminde kullanılan hiperparametreler de sonuçları etkileyen önemli unsurlardır. Dolayısıyla, çalışmalardaki metriklerin karşılaştırılmasında, her bir çalışmanın özgün koşullarını ve kullanılan veri setlerinin özelliklerini dikkate almak gereklidir. Bu bağlamda, her çalışmanın kendi bağlamında değerlendirilmesi ve genelleme yapılırken dikkatli olunması önemlidir.

5.2 Sonuçlar ve Öneriler

Bu çalışmanın ana hedefleri, havayolu müşteri yorumlarının duygu analizi için farklı yapay zekâ ve derin öğrenme modellerinin performansını değerlendirmek ve bu değerlendirmelerden elde edilen bulgularla müşteri memnuniyetini artıracak stratejik öneriler sunmaktır. Çalışmada CNN-LSTM, nnlm-en-dim50, nnlm-en-dim128 ve USE modelleri kullanılmış ve bu modellerin metin sınıflandırma görevlerindeki başarıları detaylı bir şekilde incelenmiştir.

CNN-LSTM modeli %92.06 doğruluk oranı ile oldukça başarılı bir performans göstermiştir. Ancak, bu modelin duyarlılığı %97.18 iken özgüllüğü %80.00 olarak hesaplanmıştır. Yanlış pozitif oranı %20.00 olan modelin F1 skoru ise %94.51'dir. Bu model, yüksek doğruluk oranına sahip olmasına rağmen bazı durumlarda aşırı uyum problemi göstermiştir. nnlm-en-dim50 modeli %90.87 doğruluk oranı ile biraz daha düşük performans sergilemiştir. Bu modelin duyarlılığı %94.92 ve özgüllüğü %81.33'tür. Yanlış pozitif oranı %18.67 olan modelin F1 skoru %93.59 olarak hesaplanmıştır. nnlm-en-dim50 modeli, test doğruluğu ve kaybı açısından daha istikrarlı bir performans göstermiştir. nnlm-en-dim128 modeli %92.46 doğruluk oranı ile en yüksek doğruluğa sahip modellerden biri olmuştur. Bu modelin duyarlılığı %97.18, özgüllüğü %81.33 ve yanlış pozitif oranı %18.67'dir. F1 skoru %94.77 olan nnlm-en-dim128 modeli, büyük embedding boyutu sayesinde daha zengin özellikler yakalayarak performansını artırmıştır. USE modeli %95.63 doğruluk oranı ile en yüksek performansı sergilemiştir. Bu modelin duyarlılığı %97.74, özgüllüğü %90.67 ve yanlış pozitif oranı %9.33'tür. F1 skoru %96.92 olan USE modeli, diğer modellere göre daha düşük hata oranı ve yüksek doğruluk ile öne çıkmıştır.

Sonuç olarak, CNN-LSTM, nnlm-en-dim50, nnlm-en-dim128 ve USE modelleri duygu analizi için etkili araçlar olarak değerlendirilmiştir. USE modeli, en yüksek doğruluk oranı ve en düşük hata oranı ile en iyi performansı göstermiştir. nnlm-en-dim128 modeli de yüksek doğruluk oranı ile dikkat çekmektedir. Bu sonuçlar, işletmelerin müşteri memnuniyetini

artırmak ve hizmet kalitesini iyileştirmek için duygu analizi teknolojilerini etkin bir şekilde kullanmalarına yardımcı olacak önemli içgörüler sunmaktadır.

Gelecek çalışmalar için öneriler şunlardır: Farklı veri setleri ve daha geniş örneklem büyüklükleri ile modellerin performanslarının yeniden değerlendirilmesi, modellerin hiperparametre optimizasyonlarının daha detaylı incelenmesi ve farklı yapay zekâ tekniklerinin entegrasyonu ile duygu analizinde daha yüksek doğruluk oranlarının elde edilmesi mümkün olabilir. Ayrıca, müşteri yorumlarının dil ve kültürel farklılıklarını göz önünde bulunduran modellerin geliştirilmesi de müşteri memnuniyeti analizinde önemli bir adım olacaktır. Bu öneriler, duygu analizi çalışmalarının daha ileriye taşınması ve işletmelerin müşteri geri bildirimlerinden daha derinlemesine içgörüler elde etmesi için yol gösterici nitelikte olacaktır.



KAYNAKÇA

- Acı, Çiğdem, ve Adem Çırak. 2019. “Türkçe Haber Metinlerinin Konvolüsyonel Sinir Ağları ve Word2Vec Kullanılarak Sınıflandırılması”. *Bilişim Teknolojileri Dergisi* 12(3):219–28. doi: 10.17671/gazibtd.457917.
- Alparslan, Güler, ve Mahir Dursun. 2023. “Konvolüsyonel Sinir Ağları Tabanlı Türkçe Metin Sınıflandırma”. *Bilişim Teknolojileri Dergisi* 16(1):21–31. doi: 10.17671/gazibtd.1165291.
- Alzoubi, Yehia Ibrahim, Ahmet E. Topcu, ve Ahmed Enis Erkaya. 2023. “Machine Learning-Based Text Classification Comparison: Turkish Language Context”. *Applied Sciences (Switzerland)* 13(16). doi: 10.3390/app13169428.
- Bahassine, Said, Abdellah Madani, Mohammed Al-Sarem, ve Mohamed Kissi. 2020. “Feature selection using an improved Chi-square for Arabic text classification”. *Journal of King Saud University - Computer and Information Sciences* 32(2):225–31. doi: 10.1016/j.jksuci.2018.05.010.
- Barfar, Arash. 2022. “A linguistic/game-theoretic approach to detection/explanation of propaganda”. *Expert Systems with Applications* 189(October 2021):116069. doi: 10.1016/j.eswa.2021.116069.
- Bhattacharya, Sampurna, Bhavani Gowri Shankar, B. Pitchaimanickam, ve A. Alice Nithya. 2024. “Toxic Comments Classification using LSTM and CNN”. Ss. 214–21 içinde *2024 3rd International Conference on Applied Artificial Intelligence and Computing (ICAAIC)*. IEEE.
- Çelik, Özer, ve Burak Can Koç. 2021. “TF-IDF, Word2vec ve Fasttext Vektör Model Yöntemleri ile Türkçe Haber Metinlerinin Sınıflandırılması”. *Deu Muhendislik Fakultesi Fen ve Muhendislik* 23(67):121–27. doi: 10.21205/deufmd.2021236710.
- Chen, Haihua, Lei Wu, Jiangping Chen, Wei Lu, ve Junhua Ding. 2022. “A comparative study of automated legal text classification using random forests and deep learning”. *Information Processing and Management* 59(2):102798. doi: 10.1016/j.ipm.2021.102798.
- Chen, Zhuo, Lan Jiang Zhou, Xuan Da Li, Jia Nan Zhang, ve Wen Jie Huo. 2020. “The Lao text classification method based on KNN”. *Procedia Computer Science* 166:523–28. doi: 10.1016/j.procs.2020.02.053.
- Colón-Ruiz, Cristóbal, ve Isabel Segura-Bedmar. 2020. “Comparing deep learning architectures for sentiment analysis on drug reviews”. *Journal of Biomedical Informatics* 110(February):103539. doi: 10.1016/j.jbi.2020.103539.
- Demircan, Murat, Adem Seller, Fatih Abut, ve Mehmet Fatih Akay. 2021. “Developing Turkish sentiment analysis models using machine learning and e-commerce data”. *International Journal of Cognitive Computing in Engineering* 2(October):202–7. doi: 10.1016/j.ijcce.2021.11.003.
- Gasparetto, Andrea, Matteo Marcuzzo, Alessandro Zangari, ve Andrea Albarelli. 2022. “Survey on Text Classification Algorithms: From Text to Predictions”. *Information (Switzerland)* 13(2):1–39. doi: 10.3390/info13020083.
- Göker, Hanife, ve Hakan Tekedere. 2017. “FATİH Projesine Yönelik Görüşlerin Metin Madenciliği Yöntemleri İle Otomatik Değerlendirilmesi”. *Bilişim Teknolojileri Dergisi* 291–99. doi: 10.17671/gazibtd.331041.
- Guo, Biyang, Songqiao Han, Xiao Han, Hailiang Huang, ve Ting Lu. 2021. “Label Confusion

- Learning to Enhance Text Classification Models”. Ss. 12929–36 içinde *35th AAAI Conference on Artificial Intelligence, AAAI 2021*. C. 14B.
- Gürbüz, Meryem, Dilek Sürmeli, Kamil Taşkın, ve Halil İbrahim Cebeci. 2024. “Otellere için paylaşılan çevre ile alakalı yorumların metin madenciliği ile analizi: Antalya otelleri üzerine bir araştırma”. *Business & Management Studies: An International Journal* 12(1):218–39. doi: 10.15295/bmij.v12i1.2369.
- Hajibabae, Parisa, Masoud Malekzadeh, Mohsen Ahmadi, Maryam Heidari, Armin Esmailzadeh, Reyhaneh Abdolazimi, ve James H. J. Jones. 2022. “Offensive Language Detection on Social Media Based on Text Classification”. *2022 IEEE 12th Annual Computing and Communication Workshop and Conference, CCWC 2022* 92–98. doi: 10.1109/CCWC54503.2022.9720804.
- Jang, Beakcheol, Myeonghwi Kim, Gaspard Harerimana, Sang Ug Kang, ve Jong Wook Kim. 2020. “Bi-LSTM model to increase accuracy in text classification: Combining word2vec CNN and attention mechanism”. *Applied Sciences (Switzerland)* 10(17). doi: 10.3390/app10175841.
- Karadağ, Anıl, ve Hidayet Takçı. 2010. “Metin Madenciliği ile Benzer Haber Tespiti”. *Akademik Bilişim* (July).
- Kaşıkçı, Tuğba, ve Hadi Gökçen. 2013. “Metin Madenciliği İle E-Ticaret Sitelerinin Belirlenmesi”. *Bilişim Teknolojileri Dergisi* 7(1):0. doi: 10.12973/bid.2014.
- Kılınç, Deniz, Emin Borandağ, Fatih Yücalar, Volkan Tunalı, Macit Şimşek, ve Akın Özçift. 2016. “KNN Algoritması ve R Dili ile Metin Madenciliği Kullanılarak Bilimsel Makale Tasnifi”. *Marmara Fen Bilimleri Dergisi* 28(3):89–94. doi: 10.7240/mufbed.69674.
- Kowsari, Kamran, Donald E. Brown, Mojtaba Heidarysafa, Kiana Jafari Meimandi, Matthew S. Gerber, ve Laura E. Barnes. 2017. “HDLTex: Hierarchical Deep Learning for Text Classification”. *Proceedings - 16th IEEE International Conference on Machine Learning and Applications, ICMLA 2017* 2017-Decem:364–71. doi: 10.1109/ICMLA.2017.0-134.
- Lai, Siwei, Liheng Xu, Kang Liu, ve Jun Zhao. 2015. “Recurrent convolutional neural networks for text classification”. Ss. 65–82 içinde *Proceedings of the National Conference on Artificial Intelligence*. C. 3.
- Liu, Yingying, Peipei Li, ve Xuegang Hu. 2022. “Combining context-relevant features with multi-stage attention network for short text classification”. *Computer Speech and Language* 71(June 2021):101268. doi: 10.1016/j.csl.2021.101268.
- Neogi, Ashwin Sanjay, Kirti Anilkumar Garg, Ram Krishn Mishra, ve Yogesh K. Dwivedi. 2021. “Sentiment analysis and classification of Indian farmers’ protest using twitter data”. *International Journal of Information Management Data Insights* 1(2):100019. doi: 10.1016/j.jjime.2021.100019.
- Onan, Aytuğ, Serdar Korukoğlu, ve Hasan Bulut. 2016. “Ensemble of keyword extraction methods and classifiers in text classification”. *Expert Systems with Applications* 57:232–47. doi: 10.1016/j.eswa.2016.03.045.
- Özekes, Serhat. 2003. “Veri Madenciliği Modelleri ve Uygulama Alanları”. *İstanbul Ticaret Üniversitesi Fen Bilimleri Dergisi* 2(3):65–82.
- Reusens, Manon, Alexander Stevens, Jonathan Tonglet, Johannes De Smedt, Wouter Verbeke, Seppe vanden Broucke, ve Bart Baesens. 2024. “Evaluating text classification: A benchmark study”. *Expert Systems with Applications* 254(March):124302. doi:

10.1016/j.eswa.2024.124302.

- Seker, Sadi Evren. 2015. “Metin Madenciliği (Text Mining)”. *YBS Ansiklopedisi* 2(3):30–32.
- Soni, Sanskar, Satyendra Singh Chouhan, ve Santosh Singh Rathore. 2023. “TextConvoNet: a convolutional neural network based architecture for text classification”. *Applied Intelligence* 53(11):14249–68. doi: 10.1007/s10489-022-04221-9.
- Tuzcu, Seda. 2020. “Çevrimiçi Kullanıcı Yorumlarının Duygu Analizi ile Sınıflandırılması Classification of Online User Reviews with Sentiment Analysis”. *Estudam Bilişim Dergisi* 1(2):1–5.
- Uslu, Osman, ve Serel Akyol. 2021. “Türkçe Haber Metinlerinin Makine Öğrenmesi Yöntemleri Kullanılarak Sınıflandırılması”. 2(1):15–20.
- Wang, Haitao, Jie He, Xiaohong Zhang, ve Shufen Liu. 2020. “A Short Text Classification Method Based on N -Gram and CNN”. *Chinese Journal of Electronics* 29(2):248–54. doi: 10.1049/cje.2020.01.001.
- Wang, Ruishuang, Zhao Li, Jian Cao, Tong Chen, ve Lei Wang. 2019. “Convolutional Recurrent Neural Networks for Text Classification”. Ss. 1–6 içinde *2019 International Joint Conference on Neural Networks (IJCNN)*. C. 32. IEEE.

ÖZGEÇMİŞ

Adı Soyadı: Semih Osman SAKA

Eğitim ve Mesleki Geçmişi:

- Samsun Üniversitesi, Lisansüstü Eğitim Enstitüsü, Yazılım Mühendisliği (YL) (11.08.2021-), Samsun
- İstanbul Üniversitesi, Mühendislik Fakültesi, Bilgisayar Mühendisliği Bölümü Mezunu (30.05.2017), İstanbul
- Hoca Ahmet Yesevi Uluslararası Türk-Kazak Üniversitesi, Beşeri Bilimler Fakültesi, Eğitimde Ölçme ve Değerlendirme Tezsiz Yüksek Lisans (09.02.2015)
- Süleyman Demirel Üniversitesi, Teknik Eğitim Fakültesi, Bilgisayar Sistemleri Öğretmenliği (26.08.2009)
- Samsun Üniversitesi, Uzaktan Eğitim Uygulama ve Araştırma Merkezi, Tekniker (Ağustos 2020), Samsun
- Boğaziçi Üniversitesi, Bilgi İşlem Merkezi, Tekniker (Mart 2011-Ağustos 2020), İstanbul

Yayınlar:

Classification of Apple Diseases and Pests using The Google.com Powered Teachable Machine Journal of Agricultural Faculty of Gaziosmanpaşa University (JAFAG), 41(2), 66-71. <https://doi.org/10.55507/gopzfd.1287389>

Sinir Ağı Dil Modelleri ve Evrensel Cümle Kodlayıcı Kullanarak Havayolu Müşteri Yorumlarının Metin Tabanlı Duygu Analizinin Gerçekleştirilmesi (Editörde)

Sentiment Analysis based on Text with Universal Sentence Encoder and CNN-LSTM Models (IDAP 2024)