

T.C.
VAN YÜZÜNCÜ YIL ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ
İSTATİSTİK ANABİLİM DALI

**MAKİNE ÖĞRENMESİ ALGORİTMALARI KULLANILARAK TWITTER
MOBİL OYUN VERİLERİNDE DUYGU ANALİZİ**

DOKTORA TEZİ

HAZIRLAYAN: Erol KINA
DANIŞMAN: Doç. Dr. Recep ÖZDAĞ

VAN-2022

T.C.
VAN YÜZÜNCÜ YIL ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ
İSTATİSTİK ANABİLİM DALI

**MAKİNE ÖĞRENMEİ ALGORİTMALARI KULLANILARAK TWITTER
MOBİL OYUN VERİLERİNDE DUYGU ANALİZİ**

DOKTORA TEZİ

HAZIRLAYAN: Erol KINA

VAN-2022

KABUL VE ONAY SAYFASI

İstatistik Anabilim Dalı'nda Doç. Dr. Recep ÖZDAĞ'ın danışmanlığında, Erol KINA tarafından sunulan “Makine Öğrenmesi Algoritmaları Kullanılarak Twitter Mobil Oyun Verilerinde Duygu Analizi” isimli bu çalışma Lisansüstü Eğitim ve Öğretim Yönetmeliği'nin ilgili hükümleri gereğince 13/01/2022 tarihinde aşağıdaki jüri tarafından oy birliği ile başarılı bulunmuş ve Doktora tezi olarak kabul edilmiştir.

Başkan:.....

İmza:

Üye:.....

İmza:

Üye:.....

İmza:

Üye:.....

İmza:

Üye:.....

İmza:

Fen Bilimleri Enstitüsü Yönetim Kurulu'nun/....../..... tarih ve sayılı kararı ile onaylanmıştır.

İmza

Enstitü Müdürü

TEZ BİLDİRİMİ

Tez içindeki bütün bilgilerin etik davranış ve akademik kurallar çerçevesinde elde edilerek sunulduğunu, ayrıca tez yazım kurallarına uygun olarak hazırlanan bu çalışmada bana ait olmayan her türlü ifade ve bilginin kaynağına eksiksiz atıf yapıldığını bildiririm.

(İmza)

(Erol KINA)



ÖZET

MAKİNE ÖĞRENMESİ ALGORİTMALARI KULLANILARAK TWITTER MOBİL OYUN VERİLERİNDE DUYGU ANALİZİ

KINA, Erol

Doktora Tezi, İstatistik Anabilim Dalı
Tez Danışmanı: Doç. Dr. Recep ÖZDAĞ
Ocak 2022, 157 sayfa

Bu tez çalışmasında, Türkiye’de ki hem de yabancı ülkelerdeki Twitter kullanıcılarının mobil oyunlar ile alakalı yazdıkları Türkçe ve İngilizce tweetlerin olumlu ve olumsuz bakımdan incelenerek duygu analizinin yapılması ve mobil oyuncu kitlesinin mobil oyunlar hakkındaki pozitif veya negatif bakış açıları esas alınarak mobil oyun üreticilerine mobil oyunlar hakkında fikir vermesi ve yol göstermesi amaçlanmıştır. Bu amaçlar doğrultusunda, Twitter üzerinden elde edilen mobil oyunlarla alakalı 376.431 adet Türkçe ve 1.839.274 adet İngilizce olmak üzere toplam 2.215.708 adet tweet kullanılmıştır.

Lojistik Regresyon, Rastgele Orman Sınıflayıcısı, K-Komşu Sınıflayıcısı, Doğrusal Destek Vektör Sınıflayıcısı, Çok terimli Naïve Bayes Sınıflayıcısı, Bernoulli Naïve Bayes Sınıflayıcısı, Sırt Sınıflayıcısı, Pasif-Agresif Sınıflayıcısı, Çok Katmanlı Algılayıcı Sınıflayıcısı, Oylama Sınıflayıcısı, Evrişimsel Sinir Ağı sınıflayıcısı ve Uzun Kısa Süreli Bellek Sınıflayıcısının kullanıldığı bu tez çalışmasında, İngilizce ve Türkçe veri setleri için ayrı ayrı modeller oluşturularak, sınıflayıcıların performansları karşılaştırılmıştır. TEMSAP-CNNLSTM (CNN-LSTM tabanlı Twitter Mobil Oyun Duyarlılık Analizi Yaklaşımı) adlı bir model geliştirilmiştir. Yapılan analizler sonucunda mobil oyun İngilizce test veri seti için TEMSAP-CNNLSTM modelinde 0.93’lük doğruluk değeri ve Oylama sınıflayıcısında 0.91’lik doğruluk değeri elde edilmiştir. Mobil oyun Türkçe test veri setinde TEMSAP-CNNLSTM modelinde 0.92’lik doğruluk değeri ve oylama sınıflayıcısında 0.90’lık doğruluk değeri elde edilmiştir. TEMSAP-CNNLSTM modelinin ve oylama sınıflayıcısının diğer algoritmalarından daha uzun çalışma süresinde daha başarılı sonuçlar verdiği tespit edilmiştir.

Anahtar kelimeler: Duygu analizi, Oylama sınıflayıcısı, Hibrit, Mobil oyun

ABSTRACT

SENTIMENT ANALYSIS IN TWITTER MOBILE GAME DATAS USING MACHINE LEARNING ALGORITHMS

KINA, Erol

Ph.D. Thesis, Department of Statistics

Supervisor: Assoc. Prof. Dr. Recep ÖZDAĞ

January 2022, 157 pages

The aim of this thesis is to analyze the positive and negative aspects of tweets written by Twitter users in Turkey and abroad about mobile games in Turkish and English, and to analyze mobile games based on the positive or negative views of mobile gamers about mobile games in order to give an idea and guide to mobile game manufacturers. For these purposes, a total of 2.215.708 tweets, 376.431 in Turkish and 1.839.274 in English, related to mobile games were obtained via Twitter.

Logistic Regression, Random Forest classifier, K-Neighbour classifier, Linear Support Vector classifier, Multinomial Naïve Bayes classifier, Bernoulli Naïve Bayes classifier, Ridge classifier, Passive-Aggressive classifier, Multilayer Perceptron classifier, Voting classifier, Convolutional Neural Network classifier, Long-Short-Term Memory classifier used for this thesis and the performance of the classifiers is compared by creating separate models for English and Turkish datasets. A model called TEMSAP-CNNLSTM (CNN-LSTM-based Twitter Mobile Game Sentiment Analysis Approach) was developed. As a result of the analysis, for the English mobile game test dataset, an accuracy value of 0.93 was obtained for the TEMSAP-CNNLSTM model and an accuracy value of 0.91 was obtained for the voting classifier. For the Turkish mobile game test dataset, an accuracy value of 0.92 was obtained for the TEMSAP-CNNLSTM model and an accuracy value of 0.90 was obtained for the voting classifier. It was found that the TEMSAP-CNNLSTM hybrid model and the voting classifier, one of the collective learning methods, produced more successful results in a longer working time than other machine learning algorithms.

Keywords: Sentiment analysis, Voting classifier, Hybrid, Mobile game



ÖN SÖZ

Bu tez çalışmasında her türlü ilgi ve yardımlarını esirgemeyen danışmanım Sayın Doç. Dr. Recep ÖZDAĞ'a, manevi olarak desteklerini hep hissettiren babam Uzm. Dr. Celil KINA'ya, annem Gönül KINA'ya, Eşim Bilg. Prog. Zeliha KINA'ya, çocuklarım Gönül Su'ya ve Rüzgar Ege'ye, kardeşlerim Ziraat Mühendisi Arzu KINA'ya, Türkçe Öğretmeni Canan KINA'ya ve Mimar Celil Birol KINA'ya, çalışma arkadaşlarıma teşekkürlerimi sunarım.

2022
Erol KINA



İÇİNDEKİLER

Sayfa

ÖZET	i
ABSTRACT	iii
ÖN SÖZ.....	v
İÇİNDEKİLER.....	vii
ÇİZELGELER LİSTESİ	ix
ŞEKİLLER LİSTESİ.....	xi
SİMGELER VE KISALTMALAR	xvii
1. GİRİŞ.....	1
1.1. Genel Bilgiler	1
1.2. Mobil Oyun Kavramı ve Tarihçesi.....	6
1.3. Dünyada Mobil Oyunlara Bakış	14
1.4. Türkiye’de Mobil Oyunlara Bakış.....	16
1.5. Tezin Amacı, Hedefleri, Pratik Önemi ve Planlaması.....	18
2. KAYNAK BİLDİRİŞLERİ	21
2.1. Sözlük Tabanlı Çalışmalar.....	21
2.2. Basit Sınıflandırma Algoritma Tabanlı Çalışmalar	22
2.3. Gelişmiş Sınıflandırma Algoritma Tabanlı Çalışmalar	25
3. DUYGU ANALİZİ	31
3.1. Duygu Analizinin Gerekliliği	33
3.2. Duygu Analizinin Önündeki Engeller	34
3.3. Duygu Analizinde Sosyal Medya ve Önemi	35
3.4. Duygu Analizinde Twitter Kullanımı.....	39
3.5. Duygu Analizinde Kullanılan Yaklaşımlar	40
3.6. Duygu Analizinde Doğal Dil İşleme	43
3.6.1. Kelimelere ayırma	45
3.6.2. Morfolojik-ek tabanlı köke indirgeme.....	45
3.6.3. Metin parçası etiketleme.....	46
4. MATERYAL VE YÖNTEM.....	47
4.1. Materyal.....	47
4.1.1. Lojistik regresyon (LR)	47
4.1.2. Rastgele orman (RF).....	49
4.1.3. K-En yakın komşu sınıflayıcısı (KNN).....	51

Sayfa	
4.1.4. Destek vektör makineleri (SVM)	52
4.1.5. Naïve bayes sınıflandırıcısı (NBC).....	54
4.1.6. Yapay sinir ağları (ANN)	56
4.1.7. Sırt sınıflayıcısı (RC).....	59
4.1.8. Pasif agresif sınıflayıcı (PAC).....	59
4.1.9. Evrişimsel sinir ağı modeli (CNN).....	60
4.1.10. Uzun kısa süreli bellek (LSTM)	65
4.2. Yöntem	71
4.2.1. Veri setinin oluşturulması.....	72
4.2.2. Veri birleştirme.....	74
4.2.3. Veri temizleme	76
4.2.4. Verilerin polarize edilmesi	79
4.2.5. Model oluşturma.....	81
4.2.6. Oylama sınıflayıcısı (Voting classifier).....	82
4.2.7. TEMSAP-CNNLSTM.....	84
5. BULGULAR VE TARTIŞMA.....	87
5.1. Çalışma Süreleri Analizi.....	90
5.2. Mobil Oyun İngilizce Veri Seti Analizleri	95
5.2.1. Eğitim veri seti analizi.....	96
5.2.2. Test veri seti analizi.....	101
5.2.3. Tüm veri seti analizi	106
5.2.4. TEMSAP-CNNLSTM modeli analizi	111
5.3. Mobil Oyun Türkçe Veri Seti Analizleri	116
5.3.1. Eğitim veri seti analizi.....	117
5.3.2. Test veri seti analizi.....	122
5.3.3. Tüm veri seti analizi	127
5.3.4. TEMSAP-CNNLSTM modelinin analizi	133
5.4. Ülke ve Yıl Tabanlı Analiz.....	138
6. SONUÇ.....	143
KAYNAKLAR.....	147
ÖZ GEÇMİŞ.....	157

ÇİZELGELER LİSTESİ

Çizelge	Sayfa
Çizelge 1.1. Dijital oyun dünyasında eğilimler	9
Çizelge 1.2. Ülkelere göre mobil oyun kullanıcılarının özellikleri	16
Çizelge 4.1. 2015-2021 yılları için ülkelere göre ulaşılan toplam İngilizce tweet sayıları	75
Çizelge 4.2. 2015-2021 yılları için ulaşılan toplam İngilizce tweet sayıları	75
Çizelge 4.3. 2015-2021 yılları için ulaşılan toplam Türkçe tweet sayıları.....	75
Çizelge 4.4. Ön-işlem adımlarının uygulandığı Türkçe tweet örnekleri	79
Çizelge 4.5. Ön-işlem adımlarının uygulandığı İngilizce tweet örnekleri	79
Çizelge 5.1. Kütüphanelerde kullanılan parametreler ve değerleri	88
Çizelge 5.2. Twitter mobil oyun veri seti için vektör tiplerine göre algoritmaların test, eğitim ve tüm veri setindeki çalışma süreleri (Saniye türünden)	92
Çizelge 5.3. Mobil oyun İngilizce eğitim veri seti için hesaplanan metrik ölçütler.....	97
Çizelge 5.4. Mobil oyun İngilizce test veri seti için hesaplanan metrik ölçütler	102
Çizelge 5.5. Mobil oyun İngilizce Tüm veri seti için hesaplanan metrik ölçütler	107
Çizelge 5.6. Mobil oyun İngilizce veri seti için TEMSAP-CNNLSTM modeli tarafından hesaplanan metrik ölçütler.....	112
Çizelge 5.7. Mobil oyun Türkçe eğitim veri seti için hesaplanan metrik ölçütler	118
Çizelge 5.8. Mobil oyun Türkçe test veri seti için hesaplanan metrik ölçütler	123
Çizelge 5.9. Mobil oyun Türkçe tüm veri seti için hesaplanan metrik ölçütler	128
Çizelge 5.10. Mobil oyun İngilizce veri seti için TEMSAP-CNNLSTM modelinde hesaplanan metrik ölçütler.....	134
Çizelge 5.11. Ülkelere göre 2015-2018 yılları arasında ulaşılan tweetlerin dağılımı	139
Çizelge 5.12. Ülkelere göre 2019-2021 yılları arasında ulaşılan tweetlerin dağılımı	139

ŞEKİLLER LİSTESİ

Şekil	Sayfa
Şekil 1.1. Hagenuk MT-2000 ile Tetris Oyunu	1
Şekil 1.2. Iphone11 ile Hunting Clash Oyunu.....	2
Şekil 1.3. Tek Oyunculu Bir Oyun Örneği (Hackers)	2
Şekil 1.4. Çok Oyunculu Bir Oyun Örneği (Çanak Batak)	2
Şekil 1.5. Dünyadaki akıllı telefon kullanıcılarının sayısının yıllara göre artışı	5
Şekil 1.6. Oynanabilirlik ve oyuncu deneyiminin oyun-oyuncu-tasarım üzerindeki etkisi (Nacke ve ark. (2009)'dan uyarlanmıştır).....	7
Şekil 1.7. Mobil oyun üretim süreçleri.....	8
Şekil 1.8. Nokia 6610 yılan oyunu görseli	11
Şekil 1.9. Sonic Oyunu Görseli	11
Şekil 1.10. Angry Birds oyunundan ekran görüntüsü.	12
Şekil 1.11. Pokemon Go oyunundan ekran görüntüsü	13
Şekil 1.12. Mobil cihazların ve mobil oyunların gelişiminde önemli tarihler.....	14
Şekil 1.13. Mobil oyunların 2012-2021 yılları arasında elde ettiği gelir tablosu.....	14
Şekil 1.14. Deloitte global mobil kullanıcı anket sonuçları 2019	17
Şekil 3.1. Duygu analizi süreç akışı	31
Şekil 3.2. Dünya'da sosyal medya kullanımı	38
Şekil 3.3. Basit bir NLP yapısı	45
Şekil 4.1. LR dağılımı için bir örnek grafik	48
Şekil 4.2. RF için örnek uygulama	50
Şekil 4.3. Karar ağacı için örnek uygulama.....	51
Şekil 4.4. KNN ile sınıflandırma örneği.....	52

Şekil	Sayfa
Şekil 4.5. SVM için destek vektörleri ve hiper düzlemi.....	53
Şekil 4.6. Doğrusal olmayan veri setinin doğrusala çevrilmesi	54
Şekil 4.7. Örnek perceptron (algılayıcı) modeli	56
Şekil 4.8. Yapay sinir ağı yapı blokları ve nöron yapısı	57
Şekil 4.9. Çok katmanlı algılayıcıların yapısı	58
Şekil 4.10. Maksimum havuzlama örneği	62
Şekil 4.11. Ortalama Havuzlama örneği.....	62
Şekil 4.12. Sigmoid fonksiyonun grafiği.....	64
Şekil 4.13. RELU fonksiyonun grafiği.....	64
Şekil 4.14. CNN mimarisi ve katmanları	65
Şekil 4.15. LSTM kapılarının görünümü	66
Şekil 4.16. LSTM hücre yapısı.....	67
Şekil 4.17. LSTM ağının yapısı.....	68
Şekil 4.18. Model tahmin süreci akış diyagramı	68
Şekil 4.19. TP'nin FP'ye oranları esas alınarak (a) %100, (b) %50, (c) %90, (d) %65 AUC için ROC eğrisinin değişim grafikleri.	71
Şekil 4.20. Twitter verisetine uygulanan işlemlerin akış diyagramı	72
Şekil 4.21. Tweetlerin elde edilmesi için kullanılan phyton kodu	73
Şekil 4.22. İngilizce tweetler için kelime bulutu	74
Şekil 4.23. Türkçe tweetler için kelime bulutu.....	74
Şekil 4.24. 2015-2021 yılları için ülkelere göre ulaşılan İngilizce tweet sayısının değişim grafiki.....	76
Şekil 4.25. 2015-2021 yılları için ulaşılan İngilizce tweet sayılarının değişim grafiki.....	76

Şekil	Sayfa
Şekil 4.26. Tweet veri setinin temizlenmesindeki ön-işlem sürecinin grafiksel olarak gösterimi	77
Şekil 4.27. İngilizce mobil oyun verileri için polarite hesaplanması	80
Şekil 4.28. İngilizceye çevrilen 50.000 Türkçe tweet için geçen sürenin değişim grafiği (Her bir satır çubuğu 500 tweet'e karşılık gelmektedir).....	81
Şekil 4.29. VC'nin akış diyagramı	82
Şekil 4.30. Mobil oyun veri seti için modellenen VC akış diyagramı.....	84
Şekil 4.31. TEMSAP-CNNLSTM mimarisi.....	85
Şekil 4.32. Tweet veri setinin %80 eğitim ve %20 test setine ayrılması için phyton taslak kodu	85
Şekil 4.33. Geliştirilen TEMSAP-CNNLSTM'nin akış diyagramı.....	86
Şekil 4.34. Geliştirilen TEMSAP-CNNLSTM'nin phyton taslak kodu.....	86
Şekil 5.1. Anaconda Navigator-JupyterLab 2.2.6 versiyonu web arayüzü.....	89
Şekil 5.2. Tez çalışmasında Scikit-Learn kullanımı ve phyton kodu	89
Şekil 5.3. Tez çalışmasında Keras ve Tensorflow kullanımı ve phyton kodu	89
Şekil 5.4. Tez çalışmasında kelime bulutunun oluşturulması ve phyton kodu	89
Şekil 5.5. Tez çalışmasında Plotly kütüphanesinin kullanımı ve phyton kodu	89
Şekil 5.6. İngilizce ve Türkçe tweetlerdeki test, eğitim ve tüm veri seti için CountVectorizer vektör tipine göre algoritmaların çalışma süreleri.....	93
Şekil 5.7. İngilizce ve Türkçe tweetlerdeki test, eğitim ve tüm veri seti için TF-IDF vektör tipine göre algoritmaların çalışma süreleri	94
Şekil 5.8. JupyterLab platformunda karşılaştırma ölçütlerinin tanımlanması ve Phyton kodu.....	95
Şekil 5.9. Mobil oyun ingilizce veri seti – eğitim veri seti ROC grafiği.....	99
Şekil	Sayfa

Şekil 5.10. Mobil oyun İngilizce eğitim veri setinde CountVectorizer vektör tipine göre TP, TN, FN, FP dağılımı	100
Şekil 5.11. Mobil oyun İngilizce eğitim veri setinde TF-IDF vektör tipine göre TP, TN, FN, FP dağılımı.....	100
Şekil 5.12. Mobil oyun İngilizce veri seti – test veri seti ROC grafiği	104
Şekil 5.13. Mobil oyun İngilizce test veri setinde CountVectorizer için TP, TN, FN ve FP dağılımı.....	105
Şekil 5.14. Mobil oyun İngilizce test veri setinde TF-IDF için TP, TN, FN ve FP dağılımı	105
Şekil 5.15. Mobil oyun İngilizce tüm veri setinde CountVectorizer Vektör Tipine Göre TP, TN, FN, FP Dağılımı	108
Şekil 5.16. Mobil oyun İngilizce tüm veri setinde TF-IDF Vektör Tipine Göre TP, TN, FN, FP Dağılımı	108
Şekil 5.17. Mobil oyun İngilizce veri seti – tüm veri seti ROC grafiği	110
Şekil 5.18. Mobil oyun İngilizce test- eğitim veri seti için tüm algoritmaların P, R ve F değerlerinin karşılaştırılması	111
Şekil 5.19. Validasyon veri setinin oluşturulmasında ve iterasyonunda kullanılan python kodu.....	113
Şekil 5.20. Mobil oyun İngilizce veri setinde TEMSAP-CNNLSTM modeli için hesaplanan TP, TN, FN, FP dağılımı.....	115
Şekil 5.21. Mobil oyun İngilizce test-eğitim-tüm veri seti – TEMSAP-CNNLSTM ROC grafiği.....	116
Şekil 5.22. Mobil oyun Türkçe veri seti için polarite belirlemede kullanılan Python kodu	117
Şekil 5.23. Mobil oyun Türkçe veri seti – eğitim veri seti ROC grafiği	120

Şekil 5.24. Mobil oyun Türkçe eğitim veri seti için CountVectorizer vektör tipine göre TP, TN, FN, FP dağılımı	121
Şekil 5.25. Mobil oyun Türkçe eğitim veri seti için TF-IDF vektör tipine göre TP, TN, FN, FP dağılımı	121
Şekil 5.26. Mobil oyun Türkçe veri seti – test veri seti ROC grafiği.....	125
Şekil 5.27. Mobil oyun Türkçe test veri setinde CountVectorizer vektör tipine göre TP, TN, FN, FP dağılımı	126
Şekil 5.28. Mobil oyun Türkçe test veri setinde TF-IDF vektör tipine göre TP, TN, FN, FP dağılımı	126
Şekil 5.29. Mobil oyun Türkçe veri seti –tüm veri seti ROC grafiği	130
Şekil 5.30. Mobil oyun Türkçe tüm veri setinde CountVectorizer vektör tipine göre TP, TN, FN, FP dağılımı	131
Şekil 5.31. Mobil oyun Türkçe tüm veri setinde TF-IDF vektör tipine göre TP, TN, FN, FP dağılımı	131
Şekil 5.32. Mobil oyun Türkçe test-eğitim veri seti – P, R ve F metrik ölçütlerini karşılaştırılması.....	132
Şekil 5.33. Mobil oyun Türkçe veri seti – TEMSAP-CNNLSTM test-eğitim-tüm veri seti ROC grafiği	136
Şekil 5.34. TEMSAP-CNNLSTM modeli için mobil oyun Türkçe veri setinde hesaplanan TP, TN, FN ve FP dağılımı.....	137
Şekil 5.35. Negatif tweetlerin ülkelere ve yıllara göre değişimi	140
Şekil 5.36. Nötr tweetlerin ülkelere ve yıllara göre değişimi.....	140
Şekil 5.37. Pozitif tweetlerin ülkelere ve yıllara göre değişimi.....	141



SİMGELER VE KISALTMALAR

Bu çalışmada kullanılmış bazı simgeler ve kısaltmalar, açıklamaları ile birlikte aşağıda sunulmuştur.

Simgeler

Açıklama

Acc

Doğruluk (Accuracy)

AUC

Eğri Altındaki Alan (Area Under Curve)

F-Score

F skoru

P

Kesinlik (Precision)

R

Duyarlılık (Recall)

O

Hadamard Ürünü

Kısaltmalar

Açıklama

ANOVA

Varyans Analizi (Analysis Of Variance)

API

Uygulama Programlama Arayüzü (Application Programming Interface)

Bernoulli NBC

Bernoulli Naïve Bayes Sınıflayıcısı

BTK

Bilgi Teknolojileri Kurumu

CNN

Evrişimsel Sinir Ağı

CSV

Virgülle Ayrılmış Değerler (Comma Separated Values)

ÇKA

Çok Katmanlı Algılayıcı

FN

Yanlış Negatif (False Negative)

FP

Yanlış Pozitif (False Positive)

HDFS

Hadoop Dağıtık Dosya Sistemi (Hadoop Distributed File System)

KNN

K-Komşu Sınıflayıcısı

Kısaltmalar**Açıklama****LDA**

Lineer Diskriminant Analizi

Linear SVC

Doğrusal Destek Vektör Sınıflayıcısı

LR

Lojistik Regresyon

LSTM

Uzun Kısa Süreli Bellek (Long Short Term Memory)

MLPC

Çok Katmanlı Algılayıcı Sınıflayıcısı

Multinomial NBC

Çok Terimli Naïve Bayes Sınıflayıcısı

PAC

Pasif-Agresif Sınıflayıcısı

RC

Sırt Sınıflayıcısı (Ridge Classifier)

ReLU

Doğrultulmuş Lineer Birim (Rectified Linear Unit)

RF

Rastgele Orman

ROC

Alıcı İşletim Karakteristiği (Receiver Operating Characteristic)

SCRM

Sosyal Müşteri İlişkileri Yönetimi

TN

Doğru Negatif (True Negative)

TP

Doğru Pozitif (True Positive)

VC

Oylama Sınıflayıcısı (Voting Classifier)

1. GİRİŞ

1.1. Genel Bilgiler

Oyunlar eskiden beri insanların eğlenmek amaçlı seçtikleri bir aktivitedir. Yalnız çocuklar değil, gençler ve yetişkinler de bu aktiviteden faydalanmaktadır. Teknolojinin ilerlemesi ve gelişimiyle beraber oyunlar yerini günden güne teknolojik unsurlara bırakmışlardır. Oyunlar günümüzde sadece bilgisayarlarda, tabletlerde ya da oyun konsollarında değil akıllı telefonlara kadar indirgenmiştir. Teknolojinin her geçen gün hızla ilerlemesi ve oyun alanlarının yetersizliği gibi nedenler, çocukların oyun oynama ve sosyalleşme alışkanlıklarını değiştirmektedir (Aytekin, 2017). Eskiden sokaklarda oynanan oyunların önce evlerimize, daha sonra konsollara ve masaüstü bilgisayarlara ve günümüzde ise tablet ve mobil cihazlara taşındığını görmekteyiz. Önceleri yavaş, sonrasında ise hızlı bir şekilde gelişen mobil oyun platformu tek renkli ve tek sesli teknolojiye başlayarak bugünlerde 3D akıllı telefon teknolojisiyle desteklenmeye başlamıştır. Bu değişim her geçen gün oyun platformuna yeni vizyonlar katarak ilerlemekte ve oyunculara farklı tecrübeler sunmaktadır. 1994 yılında Hagenuk MT-2000 adlı telefonda oynanabilen ve “ilk mobil oyun” olma özelliği taşıyan tetris oyunundan (Şekil 1.1), günümüzdeki akıllı telefonlara (Şekil 1.2) kadar geçen süreçte sayısız oyun geliştirilmiştir.



Şekil 1.1. Hagenuk MT-2000 ile tetris oyunu



Şekil 1.2. Iphone11 ile hunting clash oyunu.

Tek oyuncu (Şekil 1.3) ile bireysel olarak oynanabilen oyunların yanı sıra internet üzerinden oynanabilen oyunlarda birden çok kullanıcı (Şekil 1.4) aynı anda, aynı platformda buluşabilmektedir.



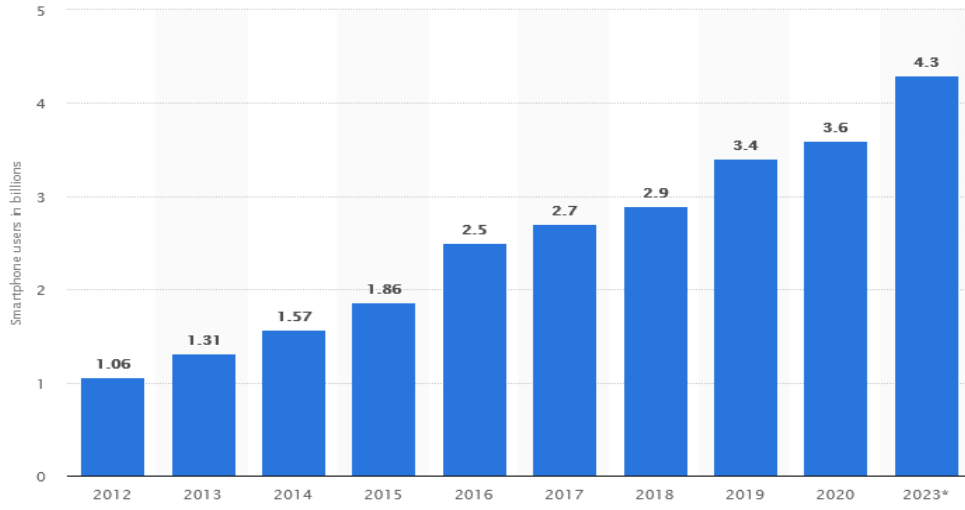
Şekil 1.3. Tek oyunculu bir oyun örneği (Hackers).



Şekil 1.4. Çok oyunculu bir oyun örneği (Çanak Batak).

Mobil oyunlar, tıpkı geçmişte olduğu gibi kullanıcılara eğlence ve eğitim amaçlı sunulan en önemli uygulamalar arasında yer almaktadır. Bireylerin vakit geçirme ve oyunda seviye atlama, sonraki turu görme, oyundaki başarının verdiği sevincin yanı sıra; oyunlarda ödül kazanmak için gösterdiği çaba, rekabet ve bunu sosyal medyaya taşıyıp meydan okuma davranışları kişilerin mobil oynama sürelerine ve oyun oynama isteklerine etki etmektedir. Mobil oyunların sosyal ağ platformları ile bütünleşmesi sonucunda kullanıcı etkileşimi daha da kolaylaşmış ve kullanıcıların daha fazla oyun oynamasına etki etmiştir. Sosyal medya platformlarının da hayatımızda önemli bir yer tutması sonucunda kişiler bu alanda daha kolay iletişime geçebilmekte, gruplar halinde oyunlar oynayabilmekte ve vakitlerinin önemli bir kısmını burada kullanabilmektedir. Sosyal medya üzerinden yapılan yorumlar, paylaşılan içerikler ve beğeniler kitleleri yönlendirebilir hale getirmiştir. Sosyal medya; internetin son yıllarda hızla gelişmesi sonucunda kullanıcıların bu ortamlara çok daha kolay erişebildiği, mobil cihazların taşınabilirliğinin kolaylaşması ve hızlanması sonucunda ise internetin her ortamda kullanabildiği ve böylece YouTube, Facebook, Twitter vb. siteler üzerinden kişilerin kolayca iletişim kurabildiği bir platformdur (Aziz, 2010). Sosyal medya kullanımının artmasıyla birlikte mobil oyunlarla alakalı üretilen içeriklerin sayısı da her saniye artmaktadır. Mobil oyun kavramı hayatımızda her geçen gün daha da etkili olmaktadır. Oyunların özellikle ergenlerde kalıcı etkiler bırakma ihtimali yüksek olması sebebiyle bu alanda gerçekleştirilecek çalışmaların önemli olduğu düşünülmektedir (Kuss ve Griffiths, 2012). Çocuklardan yaşlılara kadar tüm bireylerin ilgisini çeken mobil oyunlar doğal olarak ciddi bir veri trafiğinin oluşmasına da sebep olmaktadır. Mobil oyun sektöründe seri üretim oldukça önemlidir. Kodlanan mobil oyunların oldukça kısa bir sürede piyasaya sürülmesi ve kullanıcılar tarafından tecrübe edilmesi tüm bireyler tarafından istenen bir durumdur. Oyunların yüklenme sayısı vb. istatistiki veriler ile kullanıcıların oyunlara yaptığı yorumlar ve değerlendirmeler oyun üreticileri tarafından izlenmektedir. Büyük veri kavramının önem kazandığı günümüzde bu verilerin analiz edilmesi kaçınılmazdır. Büyük veri bilgisayarların işleyebileceğinden daha da büyük olan verilerdir. Seker (2015) tarafından yapılan çalışmada devamlı artma eğiliminde olan büyük verinin hacim (volume), hız (velocity), çeşitlilik (variety), güvenilirlik (veracity) ve değer (value) olmak üzere 5 özelliğine vurgu yapıldığı bildirilmektedir.

Günümüz teknoloji kullanımını genel olarak ele aldığımızda mobil cihazların kullanımı ve gelişimi öne çıkmaktadır. Bu tez çalışmasının temelini oluşturan mobil cihazlardaki gelişmeler ve yenilikler, mobil oyun dünyasında da dikkat çekici değişimleri ve ilerlemeleri beraberinde getirmektedir. Mobil cihaz sektörü gelişmesini tüm hızıyla sürdürürken, bu sektörde en önemli pay sahibi olan uygulama pazarında rekabet üst düzeydedir. Mobil kullanıcıların bu gelişmeye ve büyümeye kayıtsız kalmayarak mobil cihazlarıyla geçirilen vakitlerin gittikçe arttığı ve yaşam tarzlarına etki ettiği görülmektedir. Kurulan arkadaşlıkların, eğlence anlayışının, alışveriş alışkanlıklarının mobil cihazlara taşındığı görülmektedir. Bu açıdan bakıldığında mobil oyun sektörü, üreticilere sürekli yeni satış fırsatları sunmaktadır. Üreticiler açısından bu satış fırsatlarının değerlendirilmesinde kullanıcı hareketlerinin incelenmesi ve verilerinin takibi ciddi avantajlar sağlamaktadır (Çakır ve ark., 2010). Mobil uygulamaların tamamı esas alındığında her yaş grubuna direk ulaşabilen kategori “oyun ve eğlence” uygulamalarıdır. Farklı kültürel gruplarda insanlar olsalar bile mobil oyunlar ve bunların tartışıldığı sosyal medya platformları kişilerin düşünce vb. paylaşımında bulunduğu yerler olmuşlardır. Ekonomik değeri sürekli artan global oyun dünyası içerisinde Newzoo (2020) tarafından bildirildiğine göre 2020 yılı itibariyle mobil oyunlar %48’lik bir paya sahiptir ve 2022 yılında ise tüm oyun sektörünün en az yarısına sahip olacağı tahmin edilmektedir. Yine Newzoo (2020) tarafından yayınlanan verilerin sonuçlarına göre tüm dünyada yaklaşık 2.5 milyara yakın mobil oyun oynayan kullanıcı olduğuna, ayrıca akıllı telefon ve tabletlerine uygulama yükleyerek kullanan tüm kullanıcıların yaklaşık yarısının da mobil oyun yüklediğine ve oynadığına dikkat çekilmektedir. Online çok oyunculu oyunlar oynayan kişilerin maddi anlamda önemli tutarlarda yatırım yapması kişilerin oyun seçimlerinin farklı olmasına sebep olduğu düşünülmekte ve bu kişilerin 2020 yılı içerisinde mobil oyunlara 77.2 milyar dolar para harcadığı belirtilmektedir. Bu rakamlara göre mobil oyunların tüm oyun sektörü içinde en büyük paya sahip olduğu ve artık oyun sektörünün lideri olduğu görülmektedir. Oyuncular açısından bakıldığında mobil cihazların en çok tercih edilen oyun oynama ortamları olduğu söylenebilir. Statista (2021) tarafından bildirildiğine göre tüm dünyada üç milyarı aşan akıllı telefon kullanıcısı sayısının 2021 itibariyle 3.8 milyara ulaşacağı öngörülmektedir (Şekil1.5).



Şekil 1.5. Dünyadaki akıllı telefon kullanıcılarının sayısının yıllara göre artışı.

Ülkemizdeki Bilgi Teknolojileri Kurumu (BTK)'nın yayınladığı rapora göre 2018 yılı itibariyle 20 milyon mobil oyuncu bulunmaktadır. Bu sayının giderek arttığı vurgulanmış ve yabancı ülkelerdeki firmaların Türkiye’de bulunan bu oyuncu kitlesinin dikkatini çekebilmek için çeşitli yatırımlar yaptığı belirtilmiştir (BTK, 2018). Hem mobil oyun hem de akıllı telefon ile ilgili bu verilere bakıldığında sektördeki pazarın büyüklüğü anlaşılabilmektedir. Aynı zamanda paylaşılan veriler akıllı telefon kullanıcılarının birer oyuncu olarak dikkate alındığını ortaya koymaktadır. Bu durum mobil platformlar için ürün veya hizmet geliştiren firmaları yakından ilgilendirmektedir. Mobil oyun sektörü ile mobil oyuncu kitlesi arasındaki iletişime bağlı olarak daha iyi satış değerleri, kullanıcıların ilgi alanları gibi çeşitli araştırma alanları ortaya çıkmaktadır.

Mobil oyunların kalitesi, oyun deneyimi, kolay kurulması ve boyutu, oyunun sosyal yönden etkisi kullanıcılar açısından önem arz etmektedir. Penttinen ve ark. (2010) tarafından bildirildiğine göre mobil oyun üreticileri hedef kitlenin beklentilerini karşılamak amacıyla ara yüz tasarımı, oyun içi alım servisleri, canlı yardım, ürün tanıtımı, deneme versiyonu, güvenli alışveriş gibi konularda dikkatli davranmaktadır.

Duygu analizi (Görüş madenciliği), doğal dil işleme, istatistik, bilgisayar bilimleri gibi çalışma sahalarından belirli yöntemlerin ve tekniklerin kullanılması ile görüşün sahibi kişinin metnin içerisinde belirttiği duygusu, görüşü, tutumu gibi öznel bilgilerinin belirlenmesini hedefleyen günümüzde popüler bir araştırma alanıdır (Onan, 2017). Duygu

analizi çalışmalarında metinlerin olumlu, olumsuz ya da tarafsız içeriğe sahip olup olmadığı sorgulanır ve analiz edilir. Bu analizin ortaya koyduğu sonuçlara göre bireylerin ya da belirli bir grubun çalışmayla ilgili konu hakkındaki görüşü belirlenmiş olur. Bu çalışmaya indirildiğinde duygu analizi işletmeler açısından piyasaya yeni sürülecek bir ürün için ön pazar araştırması, bir topluluk için alınacak bir kararın olumlu veya olumsuz nasıl tepki alacağı, mobil oyun oynayacak kişilerin önceki yapılan yorumlara göre mobil oyun seçmeye karar vermesi gibi konularda yönlendirme yapabilir.

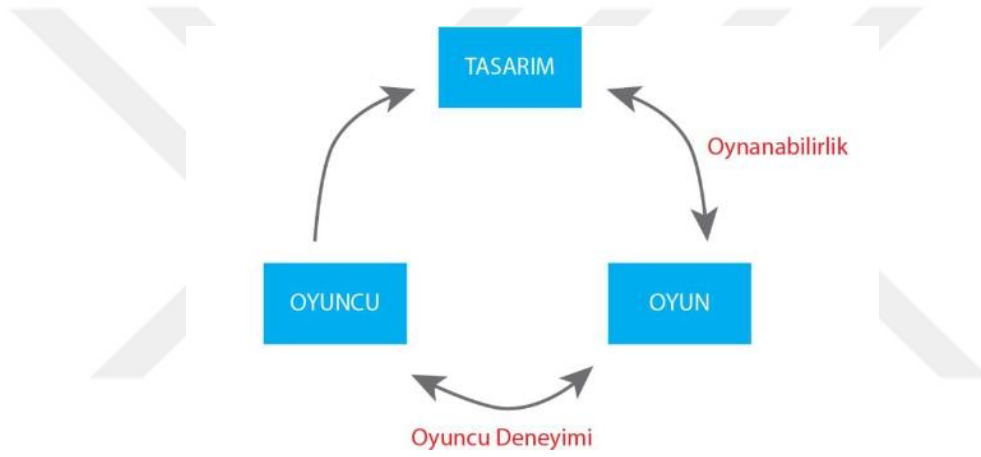
Makine öğrenmesi ile alakalı çalışmalar internet hızının artış göstermesi, mobil cihazların yaygınlaşması ve internetin daha uzun süreler kullanılmasıyla artış göstermiştir. Günümüzde makine öğrenmesi yöntemleriyle duygu analizinden, müşteri odaklı işletme analizlerine kadar çok geniş bir alanda çeşitli akademik çalışmalar yapıldığı görülmektedir. Twitter’da bulunan verinin yüksek hacimde olmasından dolayı bu çalışmaların önemli bir kısmı bu platform üzerinde yoğunlaşmıştır (Gök, 2017). Twitter platformunda birçok alanda yazılmış mevcut yorumların belli bir görüş veya eleştiri taşımasından dolayı duygu analizi çalışmalarında sıklıkla kullanıldığı görülmektedir.

1.2. Mobil Oyun Kavramı ve Tarihçesi

Mobil oyun sektörü teknolojinin gelişmesi ile paralel olarak son yıllarda etkileyici bir biçimde büyümüş ve önemli gelişmeler göstermiştir. Bu sektör özellikle akıllı telefonların artışından dolayı mobil oyun geliştirici üreticiler tarafından önemli bir pazar olarak görülmüştür. Mobil teknolojilerin sosyal hayatımızda oldukça etkili olmaya başlamasıyla insanların yaşam biçimleri ve alışkanlıkları da değişmektedir. Bu gelişmenin diğer alanlara da pozitif etkisi olmuştur. Örneğin tıp eğitiminde, sağlık eğitiminde, kültürel varlıkların sunumunda ya da örgün eğitimde animasyonların kullanımı, oyunların özellikle de oyunlaştırma kavramının altında kullanılır olmuştur. Son yıllarda da artırılmış gerçeklik ve sanal gerçeklik teknolojileri ve uygulamaları, ayrıca etkileşimli sistemler de bu sektörlere paralel olarak katılmış, ivme hızla artmıştır. Bu değişimin etkisi kendisini oyun pazarında güçlü bir şekilde göstermektedir. Oyun oynama şekillerinin ve alışkanlıklarının daha önceki tecrübelerden ayrılarak mobil teknolojilere uyum sağladığı gözle görülen bir geçektir. Bunun etkisiyle mobil oyun oynama tecrübelerinin tekrardan şekillendiği farklı bir oyun

pazarı ve bu pazarın misafirleri olarak da farklı bir oyuncu kitlesi şekillenmektedir. Bu oyuncu kitlesinin istek ve beklentilerine göre de mobil oyunların yapısı şekillenmektedir. Mobil oyun kavramını açıklarken oyun, oyuncu ve tasarım arasındaki ilişkiyi incelemekte fayda vardır.

Mobil oyunların her gün sayıları artan bir oyuncu kitlesine sahip olduğu açıktır. Oyun üreticileri tasarım ve gelişim sürecinde bu kitlenin tercihlerinin önemini dikkate almaktadır. Genel tabiriyle mobil oyun, Tasarımın ve oyuncunun dâhil olduğu bir döngünün içinde birbirine bağlı olarak bulunan ve birbirleriyle olan etkileşimlerinden, “oynanabilirlik” ve “oyuncu deneyimi” kavramlarıyla güçlendirilen bir yapıdadır.



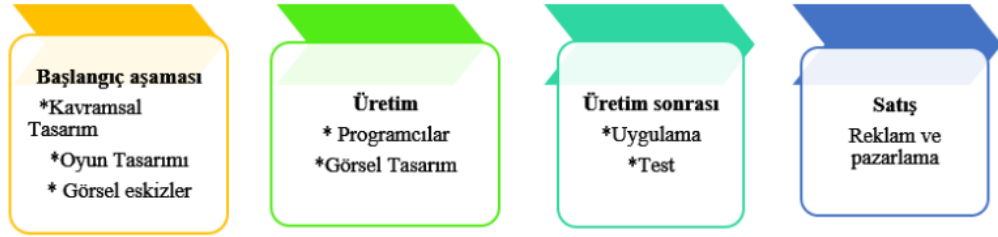
Şekil 1.6. Oynanabilirlik ve oyuncu deneyiminin oyun-oyuncu-tasarım üzerindeki etkisi.

Rollings ve Adams (2003), oynanabilirlik kavramının genel olarak oyun ve tasarımı ile direkt olarak alakalı olduğunu, oyuncunun oynama aşamasındaki deneyimlerini kapsayarak oyunun oynanış (gameplay) kalitesindeki aşamasını ortaya koyduğunu bildirmişlerdir. Oyuncu deneyimi ise oyuncunun kendisiyle ilgilenmekte ve oynama (gaming) seviyesinin daha iyi bir hale getirilmesi ile değerlendirilmektedir. Sánchez ve ark. (2009) oyuncu deneyiminin oynanabilirlik üzerindeki etkisini aşağıda yazılanlarla ilişkili olduğunu belirtmişlerdir:

- Memnuniyet (Satisfaction)
- Öğrenilebilirlik (Learnability)
- Etkililik (Effectiveness)
- Tutulma (Immersion)

- Motivasyon (Motivation)
- Duygu (Emotion)
- Sosyalleşme (Socialization)

Mobil oyun pazarı içerisinde çeşitli iş kollarından birçok uzman kişi bir araya gelmektedir: geliştiriciler, yayıncılar, dağıtıcılar, perakendeciler, yatırımcılar, fikri mülkiyet sahipleri, hukukçular, platform ya da donanım üreticileri ve pazarlama birimi bunlardan bazılarıdır. Oyun geliştirme değişik fikirler, yaratıcılık ve ortak çalışma gerektiren bir süreçtir (Şekil 1.7). İçerisinde kültürel ve ekonomik öğeler bulundurulur (Zackariasson ve Wilson, 2012). Bu öğelerin koordineli olarak işlemesi ortaya çıkacak ürünün kalitesi ve elde edeceği kazanç açısından oldukça önemlidir.



Şekil 1.7. Mobil oyun üretim süreçleri

Castronova (2002) 'nın bildirdiğine göre bizler 1996 yılından itibaren çevrimiçi oyunlara ve sanal dünyaya yapılan büyük bir göçe ev sahipliği yapmaktayız. Bu olağanüstü göçün nedeni olarak gerçek hayatlarından daha fazlasını bu sanal dünyada başarabilmesidir. Ayrıca bu sektörün daha iyi bir yaşam imkânı sunduğunu belirtmiştir. Yakın zamandan şimdiki zamana kadar olan süreçte bile, birçok alanda yenilenerek, oyun geliştiricilerine önemli fırsatlar sunmuştur. Kullanıcı kitlesi genişlemiş, iş modelleri çeşitlenmiş, ödeme biçimindeki kolaylıklardan dolayı anında ve zahmetsiz ticaret yöntemleri geliştirilmiş, oyun platformları genişlemiş ve geline nokta tüm dünya ile etkileşimli bir oyuncu ağı ortaya çıkmıştır. Çizelge 1.1'de dijital oyun dünyasındaki eğilimlerin yakın dönemdeki ve şu zaman dilimindeki karşılaştırmaları verilmiştir.

Çizelge 1.1. Dijital oyun dünyasında eğilimler

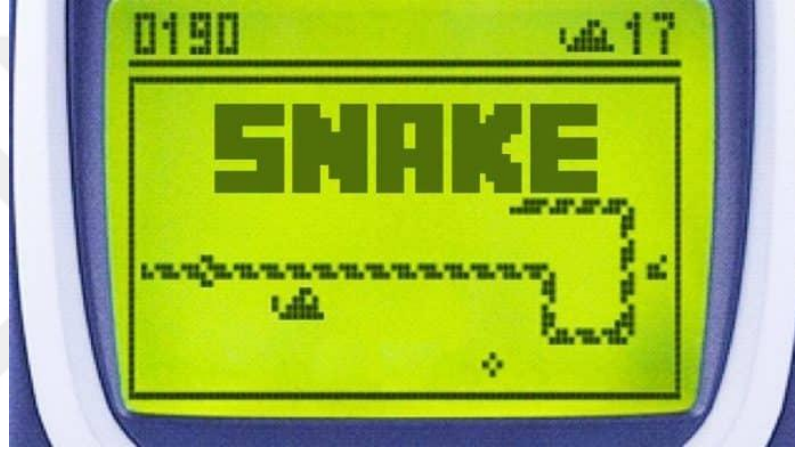
	Yakın Dönem	Şimdi
Kullanıcı Kitlesi	Genellikle sürekli oyuncular	Çocuklar, Yaşlılar, Kadınlar ve Sürekli Oyuncular
İş Modelleri	Kurulu Satış	Kutulu satış dijital satışı üyelik tabanlı oyunlar, oyun içi reklamlar, e-ticaret
Ödeme Biçimi	Nakit, Kredi Kartı	Nakit, kredi kartı, mobil ödeme, kredi kartı entegre mobil cihazlar, e-cüzdanlar (Paypal vs.)
Platform	PC, Konsol	PC, Konsol, internet tarayıcısı, Tablet, Telefon
Etkileşim	Tek Oyunculu, Coğrafi Olarak Sınırlı çoklu-oyuncu	Tüm dünya ile etkileşim

Dijital oyunların en büyük özelliği, bulunduğu zaman dilimindeki tüm teknolojilere dahil olabilesidir. 1980 sonrası kişisel bilgisayarlar, 1990 sonrasında internet, 2000 sonrasında akıllı telefonlar insanlar üzerinde etkiler bırakmışlardır ve dijital oyunlar bu adımların hepsinde kendisine yer bulmayı başarmıştır. Özellikle, işlemcilerin ve ekran kartlarının hızla gelişmesi insanların bilgisayarlar dışında farklı platformlarda deneyim yaşama talepleri, devamlı gelişen oyun içerikleri ve hikâyeleri, her yaşa uygun değişik kitlelere ulaşabilecek oyunların cebimize kadar girmesini sağlamıştır. (Zackariasson ve Wilson, 2012). Bu pazarda çalışan firmalarda yenilikçi yönlerini bu rekabet ortamında göstermişlerdir. Çünkü pazardaki rekabete bağlı olarak, ürünlerinde sürekli olarak güncelleme yapmak durumundadırlar. Daha iyi ve daha kaliteli bir ürün için başlangıç aşamasından, pazarlama yani satış aşamasına kadar titiz ve iş birliği içerisinde bir üretim süreci izlenmelidir.

Dijital oyun, bilgisayarda, oyun konsollarında, cep telefonlarında yer alan, çeşitli teknolojilerle programlandıktan sonra kullanıcılara görsel zemin hazırlanıp, kullanıcı girişi yapılarak oynanan oyunlardır (Çetin, 2013).

Mobil oyun sektörünün öne çıkmasında ve daha başarılı bir görüntü vermesinde, mobil oyun oynayanlar bakımından erişilebilirliğinin daha kolay, telefonların taşınabilirliğinin daha rahat ve oyun üretiminde harcanan paranın daha düşük bütçelerle karşılanabilmesi yatmaktadır (Bose ve Yang, 2011). Mobil oyunu bünyesinde bulunduran ilk cep telefonun ismi 1994 yılında piyasaya çıkmış olan “Hagenuk MT-2000” olarak bilinmektedir. Zamanının en gelişmiş modeli olan bu telefon tarihin ilk dâhili mobil antenini

içeriyordu. Mobil oyunların tarihçesinden söz etmek gerekirse; Mobil oyun anlayışının ortaya çıkmasında ve sonraki süreçte katkı sunan ayrıca mobil oyunlara liderlik eden “Snake” oyunu “Nokia 6610” cep telefonu ile 1997 yılında bizlerle buluşmuştur. Oyun bir pikselin yeşilimsi ekran üzerinde yemine doğru hareket eden bir yılan benzeri çubuğu göstermektedir. 90’ların sonunda telefonlar için kullanılan internetin bedelinin telefonun kendisi kadar yüksek olması Nokia 6610’da “Snake” oyununu oynamanın pahalı bir eğlence olduğunu gösterirdi. Aşağıda “Snake” oyununun ekran görüntüsü mevcuttur (Şekil 1.8).



Şekil 1.8. Nokia 6610 yılan oyunu görseli.

Sürekli ve hızla gelişen teknolojiyle mobil kullanıcıların eğilimlerinin oluşup alışkanlıklarının değişmesi üreticileri yeni oyunlar geliştirmeleri noktasında teşvik etmiştir. 2001 yılına gelene kadar üreticilerden bazıları yatırımlar yapıp bazı gelişmeler sağlamış olsa bile mobil oyun endüstrisinin çok yeni bir sektör olması nedeniyle çok önemli sayılabilecek gelişmeler olmamıştır. 2001 yılında gerçekleştirilen JavaOne konferansında Sega firması mobil cihazlar için Java tabanlı oyunlar hakkında bilgi vermiş ve sonuç olarak “Sonic” adlı verilen oyun Motorola kullanıcılarıyla buluşmuştur. Oyunun ismini de alan bu karakter daha sonradan Sega firmasının maskotu olmuştur (Şekil 1.9).



Şekil 1.9. Sonic oyunu görseli.

2002’de kullanıcıların karşısına çıkan Nokia 3410 ve Siemens M50 Java yazılımını kullanan ilk telefonlardır. Bu dev adımıyla birlikte mobil oyunlarda yeni bir çağ başlamıştır. Artık cep telefonlarına internet sayesinde oyunlar indirilebilir ve silinebilir hale getirmiştir. Oyunlar renkli değildi ve çözünürlükleri oldukça kötü sayılırdı. Örneğin Nokia 3410 sadece 94x64 piksel ölçülerine sahipti. Konsollarda oynanan Tetris, Pacman gibi oyunlar bu dönemde mobil platformlarda yerini almıştı (Brief ve Weiss, 2002).

2003 yılında yaşanan önemli gelişmelerden biri ise artık oyunların renkli hale gelmesidir. Renkli ekranların piyasaya girmesiyle cep telefonları hem mobil oyun sektörüne hem de cep telefonu sektörüne yön vermeye başlamıştır.

Mobil oyun dünyasında ciddi bir değişim ve dönüşüm yaşanmasına neden olacak belki de en önemli adım iphone firmasının piyasaya sürdüğü akıllı telefondur. Yüksek çözünürlüklü ve dokunmatik yüzeyli ekranlar ortaya çıkmıştır. İşlemciler daha hızlı ve güçlü hale gelmiştir. 2008 senesinde Apple firması mobil oyuncular için iTunes programıyla oyunları indirebilme imkânı vermiştir. Bu durum üreticiler için oyunu herhangi bir aracı kullanmadan direk olarak kullanıcıya satabildiği bir platform olması bakımından oldukça önemliydi. Teknoloji, bu dönemde artık mobil oyunlar için yüksek çözünürlük ve kaliteli ses efektleri imkânı vermiştir. Bu büyük değişimler mobil cihazlara bakış açısını değiştirmiş ve mobil oyunların hatta mobil uygulamaların türlerini arttırmıştır (Grønli ve ark., 2014). Mobil oyun pazarının önemli şirketlerinden biri olan Finlandiya menşei Rovio stüdyosu Apple firmasının dikkat çeken üreticilerinden biri olmuştur. Online olarak oynanabilen oyunları mobil platformlar için geliştirmekteydi. Şekil 1.10’da ekran görüntüsü paylaşılan

“Angry Birds” geliřtirdiđi oyunlar ierisinde en ok ilgi gren oyun olmuřtur. (Bhmer, vd., 2011).



řekil 1.10. Angry Birds oyunundan ekran grnts.

Dokunmatik ekranla bađdařık řekilde oynanan ilk oyun olması nemlidir. Oyuncular ilk defa “dokunarak” bir mobil oyunu oynayabilmiřtir.

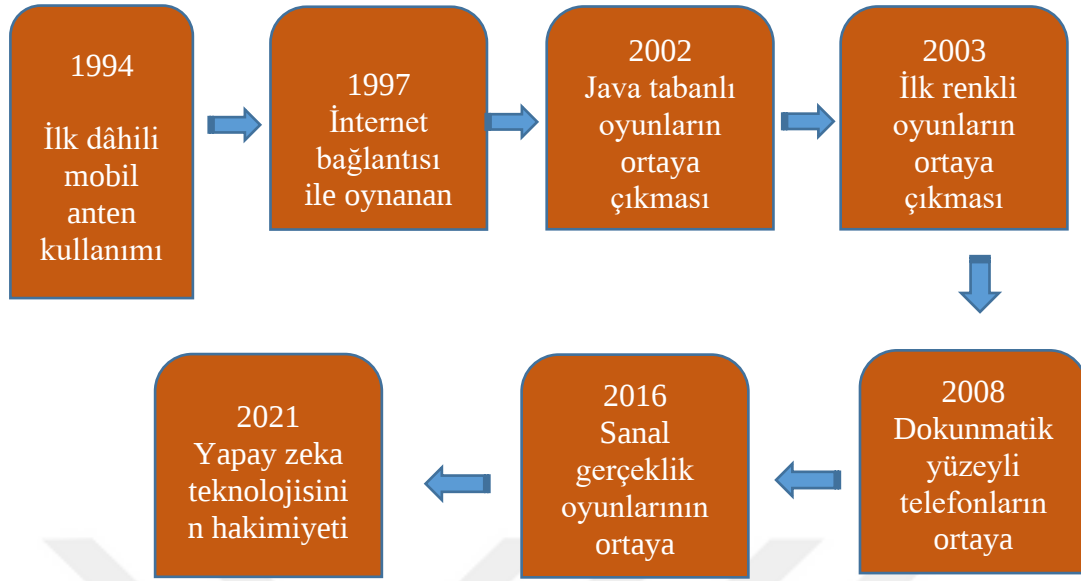
Mobil teknolojilerin byk geliřmeler gsterdiđi bu zaman diliminde mobil oyunlar sanal gereklik oyunlarıyla yeni bir dneme girmiřtir. Sanal gereklik oyunu denilince mobil oyun dnyasında ilk akla gelen oyun “Pokemon GO” adlı Niantic řirketinin geliřtirdiđi oyundur. Arttırılmıř gereklik teknolojisinin geliřmesiyle Pokemon GO 2016 yılı ierisinde iOS ve Android tabanlı sanal gereklik oyunu olarak piyasaya ıkmıřtır. Mobil oyuncular ellerindeki telefonlar ile gerek zamanlı olarak akıllı telefon kameraları aracılıđıyla dıřarda sanal olarak gezen pokemonları yakalamaya alıřıyorlardı. Hibrit-gereklik olarak adlandırılarak, sanal gereklik deneyiminin ilk kez bu oyunla beraber hayatımıza girdiđi sylenebilir. Sokak sokak gezilerek yakında olan pokemonların ayak izlerinin takip edilip poketop yardımı ile yakalanması zerine kurulmuřtur. Kısa zamanda milyonlarca oyuncu bu oyunu deneyimlemiřtir (LeBlanc ve Chaput, 2017). řekil 1.11’de ekran grnts paylařılan bu oyunla grld ki, bir mobil oyunun milyonlara ulařabilmesi artık mmkn hale gelmiřtir.



Şekil 1.11. Pokemon Go oyunundan ekran görüntüsü.

Günümüzde sektörün geldiği son nokta yapay zekânın oynanabilirlik üzerinde etkileriyle birlikte oldukça gerçekçi grafikler, sesler ve oyun deneyimleri sunduğu söylenebilir. Artık milyonlarca insan tarafından indirilen oyunlar, bu oyunlarla ilgili çekilen videoların milyonlarca izleyici tarafından izlenmesi, sosyal medya üzerinden değerlendirilmesi ve yorumlanması, milyarlarca doların bu sektörde pay edilmesi mobil oyunlarla alakalı birçok çalışmanın ve araştırmanın yapılması gerekliliğini ortaya koymaktadır.

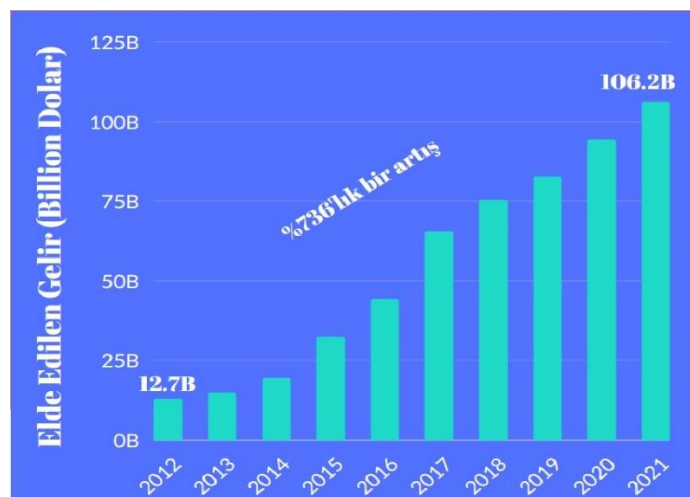
Sonuç olarak teknolojideki bu baş döndürücü değişim ve gelişim mobil oyunlara olan ilgiyi arttırmış ve buna bağlı olarak mobil oyun çeşitliliğinin artmasına sebep olmuştur. Oyuncular oyunlara daha fazla bağlanmış ve oyun oynama süreleri artmıştır. Her geçen gün pazardaki paydan almak isteyen mobil oyun geliştirici firmaların sayısı hızla artmıştır. Bilgisayarların teknolojik gelişmelerle beraber yenilenip akıllı telefonlara indirilmesi kitlelerin eğlenme alışkanlıklarını da kökünden değiştirmiştir. Firmalar oyuncu kitlelerini ellerinde tutabilmek ve oyunların oynanabilirliğini arttırmak için takımlı ya da tekli turnuvalar düzenleyerek ilgi çekmek için ödül verme yoluna gitmişlerdir. Mobil cihazların ve mobil oyunların gelişiminde önemli tarihler Şekil 1.12’de sunulmuştur.



Şekil 1.12. Mobil cihazların ve mobil oyunların gelişiminde önemli tarihler.

1.3. Dünyada Mobil Oyunlara Bakış

Mobil oyun sektörü tüm dünyada en hızlı ve etkili büyüyen sektörlerden bir tanesidir. 2012 yılında 12.7 milyar dolarlık bir pazara sahip olan mobil oyun sektörü, 2021 yılına kadarki 9 yıllık süreçte %736 gibi inanılmaz bir büyüme göstererek 106 milyar doların üzerinde bir pazar payına sahip olmuştur (Şekil 1.13).



Şekil 1.13. Mobil oyunların 2012-2021 yılları arasında elde ettiği gelir tablosu.

Büyüme hızına bakıldığında otomotiv, enerji, sinema gibi sektörleri geride bırakarak müthiş bir atılım göstermiştir. Akıllı telefonların çeşitlenmesi, tabletlerin yaygınlaşması, donanım maliyetlerinin ucuzlaması, internet erişiminin ve yayıncılığının kolaylaşması, artık sadece çocukların değil her yaştan insanın akıllı telefonlar sayesinde oyunlar oynaması, E-spor 'un bir spor dalı olarak kabul edilmesi gibi önemli faktörlerin etkisiyle bu atılımı kısa sayılabilecek bir zaman diliminde yapabilmektedir. Günümüzde dijital oyunlar, kişisel bilgisayarlarda, oyun konsollarında ve mobil platformlarda oynanabilmektedir. Son zamanlarda akıllı telefonların ve tabletlerin daha ince, daha küçük yapılarda üretilmesi ve hızlı olmaları sayesinde pazar payında mobil oyunların ön plana çıktığı net bir şekilde söylenebilmektedir. ABD'li yetişkinlerin 2018 yılında her gün 3 saat 35 dakikalarını mobil cihaz kullanarak geçirdikleri, bu sürenin 23 dakikalık bölümünü yani %11'ini mobil oyun oynayarak değerlendirdikleri yapılan araştırmalardan anlaşılmaktadır. Bu süre haftada yaklaşık 2.7 saat, yılda 5.8 gün ve ortalama bir insanın yaşamının 1.3 yılı anlamına gelmektedir. Hatta araştırmaya göre ağır mobil oyuncu statüsünde değerlendirilen kişilerin haftada 5 saatten daha fazla oyun oynadıkları varsayılarak, ortalama yaşın 38.6 olarak hesaplandığı toplam rakamın %52 sini oluşturan kadınların kullanımda önce oldukları belirtilmektedir. Yılda 12-16 tam gününü mobil oyun oynamaya ayıran bir kitleden bahsedilmektedir.

Dünya genelinde gençlerin mobil oyuncuların en büyük alt kümesini oluşturması beklense de mobil oyun kullanımına orta yaşlı kadınların hâkim olduğu ortaya çıkmıştır. Bugün mobil oyuncuların %63'ü kadınlardan %37'si erkeklerden oluşmaktadır. Ek olarak, mobil kadın oyuncuların %60'ı her gün bir mobil oyun oynarken, mobil erkek oyuncuların %47'si her gün bir mobil oynamaktadır. Ayrıca mobil oyunlar için ödeme yapmaya kadınların erkeklerden daha istekli olduğu ve üreticilerin gözünde erkek oyuncuların kadın oyunculardan %31 daha az değerli olduğu sonucuna varılmıştır. Çizelge 1.2'ye göre dünya ülkelerinin mobil oyun istatistiklerine bakıldığı zaman en fazla mobil oyuna zaman ayıran kitlenin 44 yaş üstü bireylerden oluştuğu görülmektedir. Ayrıca kadınların erkeklerden daha fazla oyun oynadığı ve mobil oyun oynama süresinin günlük 1 ile 3 saat aralığında yoğunlaştığı söylenebilir. Ülkelerde mobil oyun oynayan kitlenin nüfusun %50'sinden fazlası olduğu ortaya çıkmaktadır. İngiltere'de insanların %52'si, Almanya'da insanların %48'i, İtalya'da insanların %59'u, Hollanda'da insanların %49'u, Türkiye'de insanların

%79'u, Rusya'da insanların %53'ü, İsveç'te insanların %53'ü ve Fransa'da insanların %52'si araştırma sonuçlarına göre her gün düzenli olarak mobil oyun oynamaktadırlar (Deloitte, 2019).

Günümüzde akıllı telefon taşıyan kitle aynı zamanda potansiyel mobil oyuncu olarak da değerlendirilmektedir. Günümüz teknolojisinde internetin olduğu her ortamda akıllı telefonlara mobil oyunların indirilmesi çok kolaylaştırılmış ve hızlandırılmıştır. Pandemi sürecinin de etkisiyle bu alanda her kesimden insanın bunu değerlendirildiği gözlemlenmiştir. Yapılan araştırmaya göre Türkiye'de yetişkinlerin %79'u mobil oyun oynamaktadır. Kadınların %81.7'si, Erkeklerin ise %76.5'i bu kitleyi oluşturmaktadır. Macera, bulmaca, yarış, spor ve strateji oyunları oynanan oyun türleri içerisinde başlarda gelmektedirler. Mobil oyun oynayan kitlenin ortalama oyun oynama süresi 4 saatten fazla olarak belirtilmiştir. Deloitte (2019)'a göre Türkiye'nin bu alanda dünyada ilk sırada yer aldığını bildirmiştir (Çizelge 1.2).

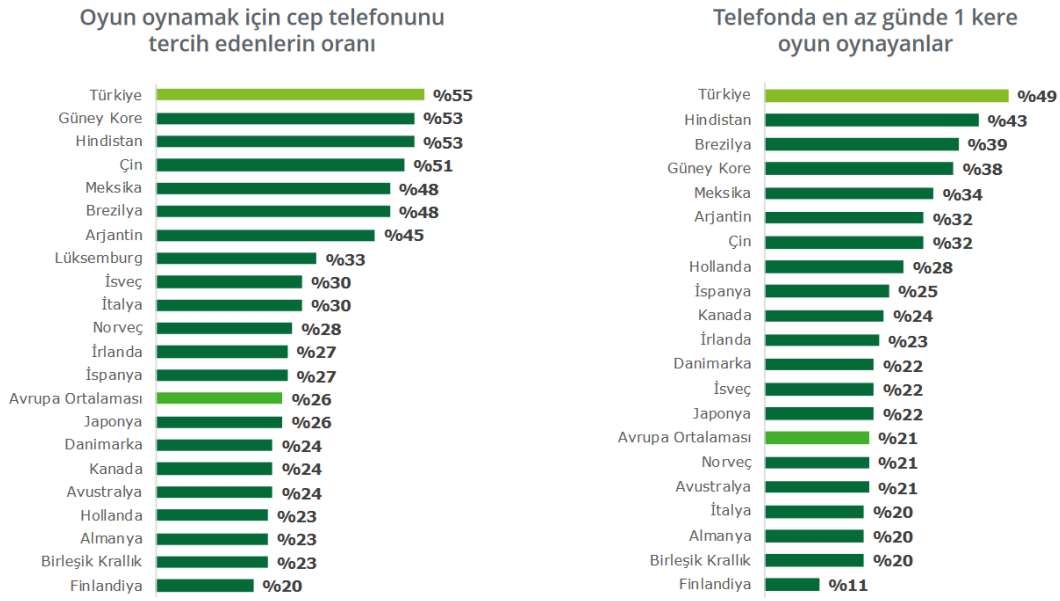
Çizelge 1.2. Ülkelere göre mobil oyun kullanıcılarının özellikleri

Ülke	Yaş Aralığı				Kadın	Erkek	Harcanan Süre			
	16-24	25-34	35-44	44+			0-1 s	1-3 s	3-5 s	5+
İngiltere	%18	%25	%25	%32	%54	%46	%38	%34	%14	%14
Almanya	%18	%25	%21	%36	%52	%48	%44	%31	%13	%12
İtalya	%18	%20	%26	%36	%51	%49	%29	%37	%17	%17
Hollanda	%20	%20	%22	%38	%55	%45	%40	%36	%14	%10
Türkiye	%24	%28	%25	%23	%50	%50	%14	%34	%22	%30
Rusya	%21	%31	%23	%25	%51	%49	%22	%32	%19	%27
İsveç	%18	%23	%22	%37	%55	%45	%29	%38	%17	%16
Fransa	%20	%23	%25	%32	%54	%46	%44	%33	%13	%10

1.4. Türkiye'de Mobil Oyunlara Bakış

Türkiye, dünyada cep telefonu bağımlılığının en yüksek olduğu ülkelerden birisi olarak öne çıkmaktadır (Şekil 1.14). Akıllı telefon kullanan kişi sayısının artış göstermesi ile bağımlılığın gelecekte daha da artacağı öngörülmektedir. Türkiye'de kullanıcılar, cep telefonlarını günde ortalama 70 kere kontrol etmektedirler. Bu, uyanık kalan zaman zarfında yaklaşık her 15 dakikada 1'e denk gelmektedir. Diğer ülkelere kıyasla Türk kullanıcıların video, fotoğraf, oyun gibi uygulamaları daha fazla kullandıkları

görülmektedir. Bunların yanında alışveriş sırasında bilgi almak için telefonlarını sıklıkla kullanmaktadır. Bu eğilim, ilerleyen yıllarda mobil ticaretin ve mobil ödemenin daha da yaygınlaşacağına işaret etmektedir. Bunların ışığında söylenebilir ki; Türkiye pazarı, mobil operatörler için fırsatlara açık bir pazar olarak öne çıkmaktadır (Deloitte, 2019).



Şekil 1.14. Deloitte global mobil kullanıcı anket sonuçları 2019.

Pandemi döneminde Türkiye’de haftada 5 saatten fazla oyun oynayan kişilerin oranı da yüzde 75’ten yüzde 93’e yükselmiştir. Benzer şekilde artık insanlar oyun oynayan diğer oyuncuları izleme konusunda da daha meraklı olmaya başlamış durumda. Oyun meraklılarının yüzde 80’inden fazlası, pandemi süreci boyunca “streaming” videolarını izlediğini ifade etmektedir. Türkiye’de ise, pandemi öncesi dönemdeki yüzde 65’e kıyasla, video oyunu meraklılarının yüzde 94’ü “streaming” videolarını izlemekten hoşlandığı belirtilmiştir.

Türkiye’de firmaların, kitlelerin, üniversitelerin ve derneklerin mobil oyun dünyasına kayıtsız kalmadığı, ilgilerini ve bütçelerini dijital oyun dünyasına kaydıracağı, bu alanda yatırımlar yaptığı bir dönemde elbette sosyal medyada bunun yansımalarını görmemiz çok doğal kabul edilebilir. Deloitte Global Mobil Kullanıcı araştırması (2019), Türkiye’nin Twitter ve Instagram kullanımında birinci sırada olduğunu açıkladı. Ve aynı firmanın verilerine göre her gün cep telefonundan oyun oynayan bir ülke halini aldığımız

söylenmektedir. Türkiye mobil oyuncuları her sene bir önceki seneye göre daha fazla harcama yaparak bu sektörün önünün açık olduğu mesajını herkese vermiştir. Gerek firmalar gerekse sektör paydaşları bu mesajı alarak twitter ve diğer sosyal medya platformlarında değerlendirmeler, yorumlar, analizler, eleştiriler yapmakta ve burada işlenmesi gereken bir veri yığını bırakmaktadırlar (Deloitte Global Mobil Oyun Araştırması, 2019). Yine mobil oyuncularda bu duruma katkı vermekte ve onlarda tartışarak, önererek, eleştirerek twitter ve diğer sosyal medya platformlarında bu veri yığına eklemeler yapmaktadırlar. Ortaya büyük hacimde bir veri yığını çıkmakta ve veriler ışığında araştırmalar, çalışmalar yapılmaktadır.

1.5. Tezin Amacı, Hedefleri, Pratik Önemi ve Planlaması

Bu tez çalışmasının amacı aşağıda sunulmuştur;

- Hem Türkiye’de ki hem de yabancı ülkelerdeki Twitter kullanıcılarının mobil oyunlar ile alakalı attıkları Türkçe ve İngilizce tweetlerin olumlu ve olumsuz bakımdan incelenerek duygu analizinin yapılması,
- Oyuncu kitlesinin mobil oyunlar hakkındaki pozitif veya negatif bakış açıları esas alınarak mobil oyun üreticilerine oyunlar hakkında fikir vermesi ve yol göstermesi amaçlanmıştır.

Bu tez çalışmasında ulaşılabilecek hedefler aşağıda sunulmuştur;

- Twitter verilerinin yıllar bazında indirilmesi, ülkelere göre lojistik olarak verilerin çekilmesi ve ilgili verilerin birleştirilmesi,
- İnsanlar olumsuz duygularını ifade etmek için olumlu metinler kullandıklarından, konuşmalarda ironi ve alaycılığı doğru bir şekilde tespit etmenin zorluğunun giderilmesi,
- Verilerin ön işleme aşamasında temizlenerek, pozitif ve negatif olarak polarize edilmesi,
- Geliştirilen yöntemle makine öğrenmesi algoritmalarından elde edilen sonuç değerlerinin doğruluğunun artırılması,
- Sadece runtime (çalışma süresi) esas alınarak kullanılan hafıza miktarı gibi kriterlere göre en iyi sonucu elde ettiğimiz makine öğrenmesi algoritmasını seçmek yerine,

kolektif öğrenme (ensemble learning) metotlarından, Voting Classifier yöntemini kullanarak elde edilen sonuçların geliştirilebilmesi,

- Makine öğrenmesi algoritmaları kullanılarak verilerin karmaşıklık matrislerine ulaşılması aşamasında eğitim ve test verilerinden optimum sonuçlar elde edilmesi,
- Geliştirilen yöntemle algoritmanın çalışma hızının arttırılması,
- Duygu analizinde başarımların doğruluğun arttırılması hedeflenmiştir.

Bu tez çalışmasının literatürdeki pratik önemi aşağıda sunulmuştur;

- Literatürde Twitter üzerinden toplanan verilerle birçok alanda duygu analizi çalışmaları yapılmış olsa da mobil oyunlar konusunda yapılmış bir duygu analizi çalışmasının olmaması,
- Tez çalışmasında, Voting Classifier yöntemiyle birlikte diğer makine öğrenmesi algoritmalarının karşılaştırılması,
- Standart duygu sınıflandırıcıları ele alındığında, tweetlerin yapılandırılmamış olması, tweetlerden farklı anlamlar çıkarılabileceği, tweetlerin dil kurallarına uygun olmayan bir dilde kullanılması gibi nedenlerle görüş ve tepkileri belirlemede yetersiz kalabilmesinin önüne geçilebilmesi,
- Literatürdeki duygu analizi çalışmalarında TF-IDF vektör tipi kullanılırken, bu tez çalışmasında ise hem TF-IDF hem de CountVectorizer olan iki ayrı vektör tipi kullanılarak sonuçların karşılaştırılması,
- Bu tez çalışmasında, Türkiye’den toplanan mobil oyunlarla ilgili Twitter verileri İngiltere, Güney Kore, Fransa, İtalya, Amerika Birleşik Devletleri, Güney Afrika ve Kanada’dan toplanan verilerle karşılaştırılarak yapılan diğer duygu analizi çalışmalarında olmayan bir çeşitliliğin yakalanması,
- Literatürde mobil oyun Twitter verileri kullanılarak yapılan duygu analizi çalışmalarının çok az sayıda olması bu tezin önemini göstermektedir.

Bu tez çalışması aşağıda sunulduğu üzere planlaması yapılmıştır:

1. Bölümde, mobil oyunlarla alakalı Dünya’dan ve Türkiye’den istatistiksel rakamlar paylaşılmış, ekonomideki önemi, daha ilgili ülkeler, yaş grupları, cinsiyet ve zaman kavramları üzerinden yaklaşımlar üzerine bilgiler verilmiştir.

2. Bölümde, duygu analizi ve kullanılan makine öğrenme algoritmaları ile ilgili literatürdeki çalışmalara vurgu yapılmıştır.

3. Bölümde, duygu analizinin gerekliliği, analizin uygulanmasında karşılaşılan zorluklar, duygu analizinde kullanılan yaklaşımlara değinilmiştir.

4. Bölümde, materyal ve yöntemler açıklanmıştır. Bu tez çalışmasında kullanılan makine öğrenmesi algoritmaları materyal olarak sunulmuş olup, mobil oyun verilerindeki duygu analizinin başarımlarını doğruluğunu artırmak amacıyla geliştirilen makine öğrenmesi algoritmaları ile yeni bir yöntem önerilmiştir.

5. Bölümde makine öğrenmesi algoritmaları kullanılarak ülkeler bazında ve yıllar bazında ayrıştırılan tweetlerin analizleri CountVectorizer ve TF-IDF vektör tipleri esas alınarak yapılmış ve hesaplanan karmaşıklık matrisleri ile ulaşılan değerlendirme ölçütlerinin ROC (Receiver Operating Characteristic) eğrileri çizilmiştir. Ayrıca ülkeler bazında ve yıllar bazında (2015-2021) hesaplanan sonuçlar üzerinden veri görselleştirme işlemi uygulanmıştır. Önerilen yöntemin başarımlarını performansını ölçmek için diğer makine öğrenmesi algoritmaları ile karşılaştırılması yapılarak tartışılmıştır.

6. Bölümde ise, bu tez çalışmasından elde edilen analiz sonuçları sunulmuştur.

2. KAYNAK BİLDİRİŞLERİ

Duygu analizi bağlamında literatürdeki çalışmalara bakıldığında, genel olarak sınıflandırma yöntemi ve sözlük tabanlı teknik kullanıldığı görülmektedir. Sınıflandırma algoritmaları (makine öğrenmesi) tekniğinde bazı sınıflandırma yaklaşımları kullanılmakta ve mevcut metinlerin karşılığı olan duygu kutbu üzerinden öğrenerek, yeni bir metin için belirlenen fikir veya duygunun çıkarımı yapılmaktadır. Literatürdeki duygu analizi çalışmalarında en çok kullanılan makine öğrenmesi algoritmaları aşağıda sunulmuştur:

- Bayes algoritması (Naive Bayes)
- Destek Vektör Makinesi (Support Vector Machines)
- K-En Yakın Komşu Algoritması (K- Nearest Neighbor Algorithm)
- Karar Ağaçları (Decision Tree)
- Yapay Ağlar (Neural Networks)

2.1. Sözlük Tabanlı Çalışmalar

Literatürde sözlük tabanlı teknikler kullanılarak yapılan duygu analizi çalışmaları incelendiğinde, Göçenoğlu (2014), yaptığı çalışmada sosyal ağ verisi toplayarak Türkiye'nin duygu analizini yapmıştır. Duygu analizinde makine öğrenmesi tabanlı Sıralı Minimal Optimizasyon algoritmasını ve sözlük tabanlı yaklaşımı uygulayarak karşılaştırmıştır. Büyük veri kaynağı olarak Twitter API verisi kullanarak duygu analizi çalışmıştır ve çıkan sonuçları Google Maps Javascript API kullanarak görselleştirmiştir.

Liu ve ark. (2019), tarafından yapılan çalışmada çevrimiçi sosyal medya platformlarından turizm acenteleri için yapılan yorumlar veri olarak kullanılarak Çinli Turistlerin Avustralya ile ilgili turistik deneyimlere dair duygu analizi yapılmıştır. Duygu analizi sözlük tabanlı, birliktelik kuralları tabanlı ve semantik analiz teknikleriyle yapılmıştır. Sözlük filtreleme, birlikte oluşum analizi ve anlamsal kümeleme teknikleri üzerine model oluşturarak Çinli Turistlerin pazarlamadaki özellik ve tercihlerinin uluslararası turistik tercihler ve özelliklerle çakışmadığına vurgu yapılmıştır.

Mogaji ve Erkan (2019), İngiltere’de yolcuların tren işleten firmalara yönelik tutum ve deneyimlerini araştırmak amacıyla Twitter üzerinden 1.914.494 tweet toplayarak işletme markalarına yönelik duygu analizi çalışması yapmışlardır. Yapılan çalışmada duygu analizi sözlük tabanlı yaklaşım kullanılarak yapılmıştır. Kütüphane olarak TextBlob kütüphanesi kullanılmıştır. Bu kütüphanedeki sözcüklere karşılık gelen puanlara göre olumlu-olumsuz ve olumlu-olumsuz-nötr olarak değerlendirme yapılmıştır.

Toçoğlu (2018), tez çalışmasında Türkçe metinler üzerinde sözlük tabanlı duygu analizi yapmıştır. 4.709 katılımcıdan 27.350 adet doküman toplandığı bir anket oluşturulmuştur. Sözlük zenginleştirmesi için bi-gram ve kavram hiyerarşisi kullanılmıştır. Sınıflandırmada doğruluğu arttırmak için bi-gram ve tri-gram modellerini kullanarak karşılaştırma yapılmıştır. Yapılan çalışmada, kelime kök-gövde etkeninin duygu üzerindeki etkisine değinilerek bir Türkçe duygu sözlüğü önerilmektedir.

Alaei ve ark. (2017), tarafından yapılan çalışmada turizm sektöründe memnuniyet durumunu ölçmek için SVM ve Naive Bayes makine öğrenmesi algoritmaları kullanılarak duygu analizi yapılmıştır. Çalışmalarında büyük veri tabanlı sistemlere dikkat çekilmeye çalışılmıştır. Kullanılan sözlük tabanlı yaklaşımda olumlu cümlelerin sınıflandırmasının, olumsuz ve nötr cümlelere göre daha başarılı sonuçlar verdiği vurgulanmıştır.

Basit sınıflama algoritmaları kullanılarak yapılan duygu analizi çalışmaları aşağıda sunulmuştur.

2.2. Basit Sınıflandırma Algoritma Tabanlı Çalışmalar

Beyhan (2014), GSM şirketlerinin Türkiye’deki durumunu araştırmak üzere duygu analizine başvurmuştur ve kümelemeye dayalı yeni bir model geliştirmiştir. Bu modelde kümeleme analizi ile duygu analizi sentezi uygulanmıştır. Yoğunluk tabanlı küme dışında kalan verilerinde dahil edildiği yaklaşım kullanılmıştır. Yaptıkları çalışmada Botego Firması üzerinden 3 operatör firmasını içeren belli bir tarih aralığındaki tweetler kullanılarak duygu analizi yapılmıştır. Aynı veri setini kümeleme yaklaşımları ile analiz etmiştir ve bu sonuçları duygu analizi sonuçları ile karşılaştırarak çıkarımlarda bulunmuşlardır.

Cheng ve Jin (2019), yaptıkları çalışmada bir otel uygulamasına ait sosyal medya verileri üzerinde konum, ağırlama ve rahatlık özellikleri bakımından müşterilerin memnuniyetleri analiz edilmiştir. İlgili çalışmada müşteri memnuniyeti duygu analizi yöntemi ile araştırılmıştır ve çalışmada turizm ve otelcilik uygulamalarında büyük verilerin ve görselleştirmenin etkin bir şekilde kullanılabileceği vurgulanmıştır. Çalışmalarında denetimsiz öğrenme yöntemini ve denetimli öğrenme yöntemini ayrı ayrı uygulayarak sonuçları karşılaştırmışlardır. Ayrıca, Jimenez-Marquez ve ark. (2019)'da otelcilik uygulamaları üzerinde çalışma yapmışlardır. Farklı sosyal medya ağları üzerinde veriler toplanarak otel işletmeciliğinde misafirperverlik üzerine yapılan araştırmada büyük veriler denetimli öğrenme yöntemi kullanılarak negatif ve pozitif olarak analiz edilmiştir.

Nasser (2018), tarafından yapılan doktora tezi çalışmasında duygu analizi yöntemleri kullanılarak Arapça dili için bir duygu derlemi yapılmıştır. Yapılan duygu derlemi ile Arapça dili için bir duygu analizi sözlüğü oluşturulması hedeflenmiştir. Bu çalışma kapsamında tüm işlemler geleneksel seri yöntemlerle gerçekleştirilmiştir. Arapça kavram temelli bir duygu sözlüğü üretmek için bir yaklaşım önerilmiştir. Arapça cümleden olumlu ya da olumsuz bir anlam çıkarmak için, anlamsal ayrıştırıcı olarak adlandırılan kural tabanlı bir algoritma önerilerek çalışmalarında uygulanmıştır. SVR (Destek Vektör Regresyon) modeli, SVR-LR (Doğrusal Destek Vektör Regresyon) modeli ile karşılaştırılmıştır.

El Fazziki ve ark. (2017) tarafından Hadoop MapReduce kullanılarak çok etmenli framework önerisinde bulunulmuş ve bir duygu analizi çalışması yapılmıştır. Map/Reduce mimarisi kullanılarak veriler Twitter4J API ile verileri elde etmişlerdir. Yapılan çalışma sosyal müşteri ilişkileri yönetimi (SCRM) üzerinde uygulanmıştır. Sosyal ağ verisi üzerinde müşteri memnuniyetini ölçmek için bir model önerilmiştir.

Pratik olarak, verilerin işlenmesi için hazırlanan modeller ne kadar basit olursa, parametreleri tahmin etmede o kadar az sorun yaşanır. Buna bağlı olarak yanılma eğilimi daha az olur ve daha iyi sonuçlar elde edilir. Bu noktadan hareketle Singh ve ark. (2016) Twitter hashtagleriyle yürüttükleri çalışmalarında, Sırt Sınıflandırıcıya dahil edilen Tikhonov düzenlemesini dikkate alarak, akademi dünyasında mevcut olan en önemli iki sınıflandırıcının doğruluk analizini gerçekleştirmişlerdir. Ridge regresyonu ağırlıkları

küçük tutmak için onları düzenli hale getirerek veriler hakkında aşırı derecede yanılmayı önler. Bununla birlikte tek bir regresyon parametresinin değerini seçmemiz gerektiğinden, model seçimi oldukça basittir ve komplike analizlere gerek yoktur. Sonuç olarak çalışmada, ANOVA ve regresyon analizinin özelliklerini kullanan LDA (Lineer Diskriminant Analizi) ile sınıflandırma işlemi yapılmış ve yüksek performans elde edilmiştir.

Sulaiman ve ark. (2020) eğitimsel alanda yaptıkları çalışmalarında sınav sorularını bilişsel seviyelere göre sınıflandırmak için Naive Bayes, KNN ve Linear SVC kullanmışlardır. Anwar veri seti kullanıldığı araştırmada, toplamda 600 soru mevcuttur (Hatırlama, Anlama, Uygulama, Analiz, Değerlendirme ve Yaratma için sırasıyla 100'er soru). Deneyler iki aşamaya ayrılmıştır. İlk aşama, SVM sınıflandırıcısı için en iyi eşlemeyi bulurken, haritalama için en iyi sonucu veren sınıflandırıcı ikinci aşamada kullanılmıştır. Sonuç, Linear SVC'nin f-ölçümü için %74.8 ve öznel çıkarma olarak TF-IDF ile en iyi sınıflandırıcı olduğunu ve sınıflandırıcının performansını artırmada gerçekten fayda sağladığını göstermektedir. Model, soru sınıflandırması için bilinen en iyi sınıflandırıcılarla (Naive Bayes, KNN) karşılaştırıldığında, Linear SVC'nin diğerleri arasında en iyi sonucu verdiği görülmüştür.

K en yakın komşu (K nearest neighbor-KNN) yöntemi, basit uygulaması ve önemli sınıflandırma performansı nedeniyle veri madenciliği ve istatistikte popüler diğer bir sınıflandırma yöntemidir. Zhang ve ark. (2017) KNN sınıflandırmasına bir eğitim aşamasını dahil ederek, farklı testler için farklı optimal k değerlerini öğrenmek adına bir kTree yöntemi önermişlerdir. Sonuç olarak, önerilen model, geleneksel KNN yöntemlerine nazaran benzer bir çalıştırma maliyetine sahipken, daha yüksek sınıflandırma doğruluğu göstermiştir.

Bir başka KNN çalışması olan Chen ve ark. (2019) tarafından yapılan deneysel araştırmada, DPeak'i iyileştirmek ve büyük ölçekli verileri işleyebilmek için, FastDPeak'in DPeak'teki iki ana adımının (ρ ve δ) hesaplama yoğunluğunun karmaşıklığının nasıl azaltılabileceği üzerine yeni bir model geliştirilmiştir. Her nokta için en yakın komşuları almada örtü ağacının kullanıldığı çalışmada, KNN-Density'nin, yoğunluk hesaplamaları için büyük bir gelişme sağlayan ve DPeak'te tanımlanan orijinal

yoğunluğun yerini alan bir model önerilmiştir. Deneysel sonuçlar, FastDPeak'in etkili olduğunu ve diğer DPeak türevlerinden daha iyi performans gösterdiğini açıklamaktadır.

2.3. Gelişmiş Sınıflandırma Algoritma Tabanlı Çalışmalar

Parveen ve Pandey (2016), yaptıkları çalışmada duygu analizi için Naive-Bayes algoritmasını kullanmışlar, tweetler üzerinde URL kaldırma, özel sembolleri kaldırma, emoji dönüştürme gibi işlemleri uygulayarak elde ettikleri ön işlemeden geçmiş Twitter verilerini kullanarak HDFS mimarisi üzerinde çalışmışlardır. Filmler ile ilgili tweet verilerini filtreleyerek, filmler için çıkarımlarda bulunmuşlardır. Emojilerin çevrilmesi yöntemiyle nötr olarak tespit edilen tweetlerde azalma olduğunu ortaya koymuşlardır.

Herhangi bir konuda karar vermek ya da tahmin yürütmek için, makine öğrenmesi ve algoritmalar aracılığıyla, konuyla ilgili veriler etkili bir şekilde işlenebilir. Günümüzde bilgisayar ortamında toplanan kişilere ait ham bilgiler üzerinde gerçekleştirilmiş pek çok deneysel çalışma mevcuttur. Duygu Analizi ve bu kapsamdaki CNN-LSTM, Ridge Classifier, Linear SVC, Random Forest, KNN araştırmalarıyla birlikte oylama sınıflandırıcısı (Voting Classifier-VC)'na yönelik olarak literatürdeki bazı çalışmalar aşağıda sunulmuştur.

Ortigosa ve ark. (2014), yaptıkları bir araştırmada, Twitter verilerinden öğrencilerin duygularının nasıl açığa çıkarılabileceğini ortaya koymuşlardır. Twitter kullanıcılarının kişiliği hakkında tahminde bulunmak için SentBuk adındaki hibrit bir model ve makine öğrenimi sınıflandırıcıları kullanan araştırmacılar, önerdikleri modelin doğruluğunu %83.27 olarak bulmuşlardır.

Terrana ve ark. (2014), yaptıkları çalışmada Facebook kullanıcıları arasındaki sosyal ilişkileri araştırmak için duygu analizi tekniklerini kullanmışlardır. Araştırmacılar, belirli bir kullanıcının ne hakkında konuştuğunu, en çok kiminle tartıştığını, kimlerle aynı fikirde olduğunu ve kiminle çatışma halinde olduğunu, gerçek zamanlı olarak öğrenmeye çalışmışlardır. Kullanıcı ana sayfasındaki duyguyu ve polariteyi otomatik olarak analiz edecek bir sistem önermişlerdir.

Yine Twitter verileriyle gerçekleştirilen Baj-Rogowska'ya (2017) ait bir başka çalışmada kullanıcıların Uber hakkında 2016 Temmuz ile 2017 Temmuz tarihleri

arasındaki dönemde belirttikleri görüşler toplanmıştır. Duygu analizi, ticari bir yazılım olan ve QDA Miner (nitel veri analizi), WordStat (içerik analizi ve metin madenciliği) ile SimStat (istatistiksel analiz) modüllerinden oluşan ProSuite kullanılarak yapılmıştır. Çalışmada kullanılan programın duygu analizinde yüksek performans gösterdiği ve duygu analizinin, müşterilerin duygularını ortaya koymada mükemmel bir araç olduğu sonucuna varılmıştır.

Jianqiang ve ark. (2018) çalışmalarında, Twitter'den topladıkları tweet'leri, duygu sınıflandırması ve tahmini için derin evrişim sinir ağına entegre etmişlerdir. Denetimsiz öğrenme ile elde edilen bir kelime yerleştirme yöntemini önermişlerdir. Bu yöntem kullanılarak deneylerin sonuçları temel modelle karşılaştırıldığında modelin Twitter duygu analizinde doğruluk ve F1 skorunda daha iyi performans sergilediği görülmüştür.

Twitter verilerini kullanarak makine öğrenimi yoluyla duygu analizinin gerçekleştirildiği Ruz ve ark. (2020) çalışmalarında, doğal afetler veya sosyal hareketler gibi zorlu durumlardaki duygu analizi sorunu ele alınmıştır. Bayes faktör yaklaşımının kullanıldığı bu çalışmada 2010 Şili depremi ve 2017 Katalan bağımsızlık referandumu sırasındaki tweetler, beş sınıflandırıcıyla incelenmiş ve performansları değerlendirilmiştir. Destek vektör makineleri ve rassal orman modelleri kullanılmış olup, doğruluğa ve yüksek performansa bakıldığında, sosyal bilimcilerin, makine öğrenimi tekniklerini Twitter'da paylaşılan içerikler üzerinde kullanabilecekleri ve bu sayede duyguları anlayabilecekleri ortaya çıkmıştır.

Zainuddin ve ark. (2018)'de yaptıkları çalışmada Twitter için yeni bir hibrit duygu sınıflandırması önermişlerdir. Önerdikleri modelle yaptıkları deneyler neticesinde ise modellerindeki sınıflandırma yöntemlerinin, daha önceki sınıflandırma yöntemlerinden doğruluk performansını %71.62 ile %76.55 arasında iyileştirebildiğini görmüşlerdir.

CNN-LSTM başta olmak üzere, diğer ilgili modeller kullanılarak yapılan analiz çalışmalarının özellikle son yıllarda hızlı bir şekilde arttığı, gerçekleştirilen literatür taramasından anlaşılmıştır. Bu çalışmaların büyük çoğunluğu İngilizce dilindeki metinler baz alınarak sürdürülse de farklı diller üzerinde de deneyler yapıldığı görülmüştür.

Arapça, farklı lehçelere sahip bir dil olsa da Ombabi ve ark. (2020), yaptıkları çalışmalarında, makine öğrenmesiyle, diğer dillere nazaran zor bir morfolojik yapısı olan Arapça dilindeki paylaşımlar için duygu analizi gerçekleştirilebileceğini ispat etmişlerdir.

Çalışmalarında CNN ve LSTM mimarisiyle duygu analizi gerçekleştirmişler ve önerdikleri modelin sınıflandırma performansındaki doğruluk oranını %90.75 şeklinde bulmuşlardır.

Vietnam'daki ticari web sitelerinden alınmış içeriklerin veri kümesi olarak baz alındığı başka bir duygu analizi çalışması olan Vo ve ark. (2017) yaptıkları çalışmada CNN ve LSTM'nin avantajları birleştirilmiştir. Bununla birlikte çalışmada makine öğrenmesinden elde edilen değerlendirme sonuçları yanı sıra üç bin insanın yorumları dikkate alınmıştır. Çalışmada olumlu-olumsuz duygu durumları yanı sıra belirsiz duyguların da performansının iyileştirilmesi amaçlanmıştır. Sonuç olarak model, ayrı ayrı SVM, LSTM ve CNN'den daha iyi performans göstermiştir.

İngilizce dilinde gerçekleştirilen çalışmalarda genel olarak başarı elde edilirken, araştırmacılar, diğer dillerde yazılan metinlerin, kendi modelleri kapsamında, duygu analizine dâhil edilemeyeceğini ve farklı hibrit çalışmalara ihtiyaç olduğunu savunmuşlardır. Bu düşünce çerçevesinde sadece İngilizce dili için model geliştiren Jain ve ark. (2021), araştırmalarında, duygu analizi için hibrit bir evrimsel sinir ağı (CNN-LSTM) modeli önermişlerdir. Keras kelime gömme yaklaşımıyla Airlinequality ve Twitter platformundan alınan İngilizce veri setleri üzerinde deneysel analizler gerçekleştirilmiş ve metinlerdeki benzer kelimelerin arasındaki mesafeler sayısal değerlere dönüştürülmüştür. Modelin klasik ML modellerinden daha yüksek doğruluk oranına (%91.3) sahip olduğu görülmüştür.

Wang ve ark. (2019) çalışmalarında, metinlerin birleşme değeri – uyarılma (Valence-Arousal) derecelendirmelerini tahmin etmek için iki bölümden oluşan ağaç yapılı bölgesel bir CNN-LSTM modeli önermişlerdir. Birleşme değeri – uyarılma tahmini yapmak için ağaç yapılı bölgesel CNN- LSTM modeli kullanma prosedürü; bölgesel CNN-LSTM modeli ve bölgesel bölme stratejisi olmak üzere iki bölümden oluşur. Eldeki metinlerin tamamı girdi olarak kullanılmaz. Bunun yerine bölgesel CNN- LSTM modeli, bölgeler içindeki yerel n-gram özelliklerini ve bölgeler arasındaki uzun mesafe bağımlılıklarını çıkarmak için girdi metinlerini birkaç bölgeye ayırır. Bölgeleri daha anlamlı bir şekilde bölmek için ise ağaç yapılı bölge bölme stratejisi kullanılır. Önerilen bu modelle gerçekleştirilen deneylerin sonuçlarıyla daha önceki çalışmaların

regresyonuna bakıldığında, modelin, geleneksel NN tabanlı ve yapılandırılmış yöntemlerden daha iyi performans gösterdiği anlaşılmıştır.

Rehman ve ark. (2019), yaptıkları çalışmada sosyal medya aracılığıyla paylaşılan içeriklerin çokluğuna bağlı olarak duygu analizi problemlerinin ortaya çıkmasını göz önünde bulundurmuş ve bu problemlerin üstesinden gelmek için IMDB ve Amazon üzerinde LSTM ve Hibrit CNN-LSTM modellerini uygulayarak, hibrit bir model önerisinde bulunmuşlardır. İlk olarak veri kümesindeki kelimeleri vektör uzayı olarak başlatmak için word2vec tekniği, kelimeleri vektör temsiline dönüştürmek için ise word2vec gram atlama (skip gram) ve kelime torbası (bag of word) tekniğinden faydalanılmıştır. Deney sonuçları, önerilen Hibrit CNN-LSTM modelinin doğruluğu ve performansının, geleneksel makine öğrenimi tekniklerinden daha iyi olduğunu göstermektedir.

Wang ve ark. (2016) kısa metinlerin duygu analizi için CNN ve LSTM katmanlarının sıralı bir şekilde organize edildiği evrişimli (CNN) ve tekrarlayan (RNN) sinir ağlarının derin bir kombinasyonunu ortaya koymuşlardır. Modeller, veri kümesinde, mevcut modellere göre daha iyi performans göstermiş ve onlardan daha yüksek sınıflandırma doğruluğu sağlamıştır. Önerilen eklemli sinir ağı mimarisinin, doğal dil işleme görevlerinin yanı sıra cümle modellemeye de uygulanabileceği anlaşılmıştır.

Li ve ark. (2020), yaptıkları çalışmada CNN ve LSTM/BiLSTM dallarını paralel bir şekilde birleştiren, sözlük entegreli iki kanallı CNN-LSTM aile modelleri olarak adlandırılan, derin öğrenme tabanlı duygu dolgusu (sentiment padding) yöntemini önermişlerdir. Stanford Sentiment Treebank gibi zorlu veri kümeleri üzerinde yapılan deneylerin sonucunda, önerilen yöntem, birçok temel yöntemden daha iyi performans göstermiştir. Bununla birlikte CNN-LSTM modelinin CNN- BiLSTM modeline kıyasla, Çin turizmine ait veri setinde daha iyi performans gösterdiği bulunmuştur. Deneyler ayrıca, duygu sözlüğü bilgisinin (sentiment lexicon information) ve paralel iki kanallı modelinin, duygu analizi doğruluğunu önemli ölçüde iyileştirdiğine işaret etmektedir.

Zeng ve ark. (2019), yukarıdaki çalışmalardan farklı olarak görünüm tabanlı duygu analizi için etkili bir evrişimli sinir ağı modeli önermişler ve en-boy tabanlı duygu analizine (aspect based sentiment analysis-ABSA) odaklanmışlardır. Çalışmada ABSA'daki sorunları çözmek için GLRC adlı bir model önerilmiştir. SemEval 2014

Restoran Veri Kümeleriyle yapılan deneyler, GLRC'nin dilbilimsel bilgiyi CNN ile birleştirebildiği ve yaklaşımın en-boy tabanlı duygu analizinde mükemmel sonuçlar elde ettiği görülmüştür.

Yousaf ve ark. (2020), yaptıkları çalışmada tweetleri oylama sınıflandırıcısı (Voting Classifier) kullanarak sınıflandırmak için derinlemesine bir karşılaştırmalı performans analizi gerçekleştirmişler ve bunun için LR-SGD kombinasyonundan faydalanmışlardır. Yedim makine öğrenimi modeli, tweetleri mutlu ve mutsuz olmak üzere sınıflandırmıştır. Sonuçlar, tüm modellerin tweetler için çok iyi performans gösterdiğini, ancak önerilen oylama sınıflandırıcısı olan VC (LR-SGD)'nin hem TF hem de TF-IDF kullanarak hepsinden daha iyi performansa sahip olduğu anlaşılmıştır. Önerilen model, %79 Doğruluk, %84 Geri Çağırma ve %81 F1 puanı ile TF-IDF kullanılarak en yüksek sonuçları elde etmektedir.

Sosyal medyada paylaşılan içeriklerin çokluğu, zararlı içeriklerin engellenmesini de zorlaştırmaktadır. Bu sebeple, zararlı içerikleri etkili bir şekilde tespit eden yüksek performanslı sınıflandırma araçlarına ihtiyaç vardır. Bu noktadan hareketle MUCIC ekibi, Nefret Söylemi için önerilen bir Oylama Sınıflandırıcısını (VC) ileri sürmüşlerdir (Balouchzahi ve ark., 2021). Araç, Twitter kullanıcılarını, tweetlerine göre 'Nefret söylemi yayıcısı' veya 'Değil' olarak sınıflandırmak için modellenmiştir. Modellerin performansları, görev düzenleyici tarafından doğruluk metriğine dayalı olarak değerlendirilmiştir. Sonuç olarak VC modelinin İspanyolca ve İngilizce metinler için sırasıyla %83 ve %73 doğruluk elde ettiği görülmüştür. Nefret söylemi ile ilgili İngilizce ve İspanyolca tweetler üzerinde gerçekleştirilmiş

Hölliga ve ark. (2021) araştırmalarında, Twitterda nefret söylemi yayıcıların profilini oluşturmak amaçlanmıştır. İngilizce için, isim öbeği düzeyinde TF-IDF özelliklerine sahip bir Doğrusal Destek Vektör Makinesi kullanılırken, İspanyolca için isim öbeği düzeyinde basit sayımlara sahip bir Ridge Sınıflandırıcıdan faydalanılmıştır. Her iki sınıflandırıcı da bir evrimsel sınır açısından elde edilen özelliklerle beslenmiştir. İngilizce test setinde %73, İspanyolca test setinde %79 doğruluk elde edilmiştir. Her iki dil için genel ortalama doğruluk %76 olarak çıkmıştır. Çalışmada ayrıca her kullanıcıya bir tür "nefret ağırlığı" ekleyen "LoveCount" ve "ProbMean" ek özellikleri geliştirilmiştir.

Son yıllarda oylama sınıflandırıcıları, hastalıkların tanı ve teşhisinde daha hızlı hareket etmek adına sıklıkla kullanılan yöntemler olmuşlardır. Kumar ve ark. (2017), Wisconsin Üniversitesi'nden aldıkları verilerle, iyi huylu ve kötü huylu tümörün sınıflandırılması için çeşitli veri madenciliği tekniklerini kullanmışlardır. Bilgisayarlı Göğüs Tomografisi (BT) taramaları da birtakım hastalıkların tanısında oldukça önemli rol oynar. Halihazırda gündemde olan COVID-19 tanı ve teşhisinde BT görüntülerini işlemek ise oldukça pahalı bir iştir. Makine Öğrenimi teknikleri bu zorluğun üstesinden gelme potansiyeline sahiptir.

El-kenawy ve ark. (2020) AlexNet adlı bir Evrimsel Sinir Ağı (CNN) kullanarak BT taramalarını işlemişler ve bir oylama sınıflandırıcısıyla (Voting Classifier), en çok oy alan sınıfı seçmek için farklı sınıflandırıcıların tahminlerini toplamışlardır. Önerilen öznitelik seçim algoritması (SFS-Guided WOA), literatürde yaygın olarak kullanılan diğer optimizasyon algoritmalarıyla karşılaştırılmış ve performansının, diğer oylama sınıflandırıcılarından daha üstün olduğu görülmüştür.

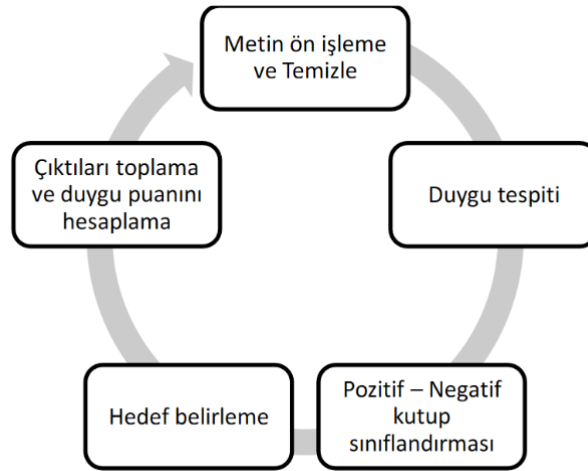
COVID-19 teşhisi için Iwendi ve ark. (2020) yaptıkları çalışmada, AdaBoost algoritmasıyla desteklenen ince ayarlı bir Random Forest modeli önermişlerdir. Modellerinde, vakanın ciddiyetini ve olası sonucunu, hastanın iyileşmesini veya ex olma riskini tahmin etmek için seyahat, sağlık ve demografik verilerini kullanmışlardır. Kullanılan veri seti üzerinde model, %94 doğruluk ve 0.86 F1 puanına sahiptir. Boosted Random Forest algoritmasının, dengesiz veri kümelerinde bile doğru tahminler sağladığı araştırmada keşfedilmiştir. Bununla birlikte yapılan veri analizi sonuçlarına göre, Wuhan yerlileri arasında ölüm oranının yerli olmayanlara ve kadın hastalara göre erkek hastalarda daha yüksek olduğu görülmüştür. Ayrıca hastalıktan etkilenen hastaların büyük bir çoğunluğu 20 ila 70 yaşları arasındadır.

Yine bir başka hastalık türü olan ve çağımızda özellikle yetişkinlerde yaygın olarak görülen diyabet hastalığının tanısı için Kumari ve ark. (2021) en doğru algoritmayı keşfetmeyi amaçlamışlardır. Önerilen oylama sınıflandırıcısı, diyabeti olan ve olmayan hastaların özelliklerini tahmin etmektedir. Pima Indians, diyabet veri setinde %79.08 oranında doğru sonuçlar verirken, aynı model meme kanseri veri setinde %97.02 oranında doğruluğa sahiptir.

3. DUYGU ANALİZİ

Duygu, Türk Dil Kurumu (TDK) sözlüğünde, “belirli bir nesne, olay veya bireylerin insanların iç dünyasında uyandırdığı izlenim” şeklinde ifade edilmiştir. Duygu analizi, Doğal dil işlemenin bir alt disiplini olup içeriklerdeki görüş, duygu, düşünce ve tutumları belirlemek için çalışan bir bilim dalıdır. Geçmişten başlayıp bugüne kadar gelen süreçte insanların duyguları önemli ve üstünde durulan bir yerdedir. Kültürlerin oluşmasında, kültürlerin yok olmasında duygular etkili olmuştur. Bu açıdan bakıldığında duyguların, insanların yaşamında önemli bir yer tuttuğu aşikârdır. Duygu, insanların yaşamları boyunca karşılarına çıkan olayların veya nesnelerin anlam içerip içermediğini gösteren en önemli unsurdur. İnsanların ilgilendikleri konu ya da düşünce üzerindeki duygusal yaklaşımı o konuyu baskın ya da önemsiz kılabilir. Bundan dolayı duygu analizi konusunda çalışmak bir ihtiyaçtır. Hatta duygu analizi günümüz teknolojileri ve gelişmelerine dayanarak üzerinde çalışılması daha fazla ihtiyaç haline gelmiş bir alandır (Eröz, 2011).

Duygu Analizi kişilerin ürünler, konular ve olaylar veya bütün bu başlıkların özellikleri hakkında değerlendirmelerini ve tutumlarını analiz eden bir kavramdır (Satapathy ve ark., 2018). Delen ve ark. (2014) ‘nın duygu analizi süreç akışını Şekil 3.1’de açıklamışlardır.



Şekil 3.1. Duygu analizi süreç akışı.

Duygu Analizi, dijital ortamdaki bütün metin içeriklerinde, makine öğrenmesi tabanlı veya istatistiksel yöntemlerle analizler yapılarak duyguların sınıflandırılması işlemidir (Poria ve ark. 2018). Yani yazılı metinlerin çeşitli bilgisayar yazılımları ile analiz edilmesidir denilebilir. Bu analizler algoritmik veya istatistiksel çözümlemeler ile yapılır. Duygu Analizi üzerine çalışılırken üç teknik seviyeden söz edilebilir.

- 1- Doküman Seviyesi
- 2- Cümle Seviyesi
- 3- Varlık-Görüş Seviyesi

Doküman seviyesinde yapılan analizler için genel olarak tüm metinden pozitif Veya negatif bir fikre ulaşmak hedeflenmektedir. Avantajı uygulamada kolaylık sağlaması dezavantajı ise sadece genel bir duygu polaritesine varılabilesidir (Pang ve Lee, 2008).

Cümle seviyesinde negatif, pozitif ve nötr olarak cümlelerin içerik taşıyıp taşımadıkları tespit edilir. Fikir beyan etmede nötr ifadeler dikkate alınmaz. Bu teknikte ima ya da ironi taşıyan cümlelerin hassasiyetinde sıkıntılarla karşılaşılmaktadır (Satapathy ve ark., 2018).

Varlık – Görüş seviyesinde metin yapıları ile ilgilenilmez, doğrudan pozitif ya da negatif bir duygu taşıyıp taşımadığı ve bir fikir yani hedef gösterip göstermediği üzerinde durulur. Bu analiz seviyesinin amacı metnin yönünü anlamaya çalışmaktır (Satapathy ve ark., 2018).

Duygu analizleri kısaca metinsel veriler üzerinde aşağıda belirtilen amaçları gerçekleştirmek için kullanılmaktadırlar (Kaynar ve ark., 2016).

- Duygu ve görüşlerin sınıflandırılması
- Öznelliğin sınıflandırılması
- Alaycılık ve ironinin tespit edilmesi
- Metin türü ve yazarının tespiti
- SPAM ve sahte yorumların tespiti
- Görüş ve yorumların ortaya çıkarılması
- Görüş ve yorumların özetlenmesi

3.1. Duygu Analizinin Gerekliliği

Firmalar, üreticiler, kamu sektörü günümüzde ürünlerini ve hizmetlerini tanıtmak amacıyla sosyal medyanın gücünden yararlanmaktadırlar. Twitter bu amaçla kullanılan en büyük sosyal medya platformlarından birisidir. Kullanıcılar açısından bakıldığında onlarda ürünler ve hizmetler hakkında bilgi edinmek için bu platformu kullanıyorlar. Sosyal medyanın günümüzde birçok açıdan müşterilerin alışveriş yapma şeklini değiştirdiği söylenebilir. Kullanıcılar çoğu zaman ürün almadan önce sosyal medyadan tüketicilerin yapmış oldukları yorumlarını okuyarak bu yorumları dikkate almaktadırlar. Böylece diğer kullanıcıların ürünler hakkındaki fikirleri insanların o ürünü almadan önce kendileri için en uygun olanı bulmalarına yön verebilmektedir. Üreticilerin kendi ürünleriyle alakalı açıklamaları diğer kullanıcıların aynı ürünle alakalı yaptıkları değerlendirmeler ile karşılaştırıldığında, kullanıcı yorumlarının farklı kişilerce yapılması ve objektif olmasından dolayı daha fazla tüketiciye hitap ettiği söylenebilir. Sonuçta günümüzde insanların tatil için otele gitmeden önce, otel hakkında yorumları okumaları otelin kendi yaptığı tanıtımdan daha etkili olmaktadır. Bilgiyi paylaşma şeklimizin değişmesi aynı zamanda sosyal medya kullanımının sürekli artması, kişilerin yorumlarıyla oluşturduğu verinin de sürekli olarak büyümesini sağlamıştır. Herhangi bir ürün hakkında kullanıcıların yaptıkları çok büyük miktardaki yorumları kişisel imkânlarla okumak ve sonuca varmak mümkün değildir (Vidushi ve Sodhi, 2017). O ürünün olumlu ve olumsuz her şeyini otomatik olarak ortaya koymak ve sonuçlandırma sürecini kısaltmak için otomatik bir sisteme gereksinim vardır.

Bütün bu işleri istekler doğrultusunda yapan ve meslek haline getirmiş şirketler vardır. İşte bu süreçleri yönetmek için, zamandan tasarruf sağlamak için ve gizli kalmış çöp denilen verilerden iyi sonuçlar ve doğru bilgiler elde etmek için duygu analizine ihtiyaç duyulmaktadır. İnsanlar kullandıkları ürünler ve servislerin olumsuz yanlarına sosyal medyada çeşitli tepkiler vermektedirler. Bundan dolayı firmalar da tüketicilerin yorumlarının analizini, kendilerinin yeterlilik, planlama ve karar alma aşamaları için kullanılması kaçınılmazdır (Vidushi ve Sodhi, 2017). Ortalama bir insan okuyucusu ilgili siteleri tanımlamakta ve içindeki fikirleri almakta ve değerlendirmekte zorluk çekecektir. Bundan dolayı otomatik duyarlılık analiz sistemlerine ihtiyaç vardır (Liu, 2012).

3.2. Duygu Analizinin Önündeki Engeller

Duygu Analizi spor, sağlık, eğitim, ekonomi, alışkanlıklar, alışveriş ve sosyal medya 'ya kadar birçok alanda farklı bakış açısı kazandırarak günümüzde popüler bir konumdadır. Fakat aşağıda belirtilen bazı sorunlar Duygu Analizi çalışmaları için sorun teşkil etmektedirler.

İroni Yapılması:

Duygu Analizi çalışmalarında genellikle kelimenin sözlükte tanımlanmış ilk manası kabul edilir. Ancak günlük yaşamımızda kelimelerin asıl anlamı dışında kullanılması, deyimlerin, ironilerin kelimeye başka anlamlar kazandırması yapılan analizleri direkt etkilemektedir. (Pradhan ve ark., 2016; Pozzi ve ark., 2017; Poria ve ark., 2018; Karoui ve ark., 2019).

Anlam farkı:

Duygu barındıran bazı kelimeler farklı alanlarda pozitif, nötr ya da negatif anlam kazanabilir. Örneğin “Bu oyun yeni çıktı, sıcak sıcak!!” ve “Sıcak havalarda oyun oynanmıyor” cümlelerinde geçen sıcak kelimesi ilk cümlede nötr bir polarite, ikinci cümlede negatif bir polarite taşımaktadır. (Pradhan ve ark., 2016; Karouive ark., 2019).

Duygu Sözcüğü Yokluğu:

İçerisinde duygu barındıran bir kelime geçmese dahi cümlelerin olumlu ya da olumsuz bir fikre sahip olduğu durumlar olabilir. “Bu mobil oyun telefonun şarjını çok fazla kullanıyor” cümlesine bakıldığında direkt bir duygu sözcüğü kullanılamaz da kullanıcı mobil oyunun fazlasıyla telefonun şarjını tükettiğinden şikâyet ederek olumsuz bir görüş belirtmektedir. (Pradhan ve ark., 2016; Karouive ark., 2019).

Cümle Yapılarındaki Karışıklık:

Tamamlanmamış, devrik ya da dil kurallarına uymayan birçok yorumun, yazının sosyal medya üzerinden paylaşıldığı görülmektedir. Bu durumun doğal dil işleme ile alakalı araştırmalarda problemlere neden olduğu öngörülebilir. Sözlük temelli yaklaşımlar diğer yaklaşımlara göre bu durum için daha iyi sonuçlar vermektedir. Nedeni ise kelime yapısına bakmasından dolayıdır. (Sharda ve ark., 2014).

Verilerin Gürültülü Olması:

Duygu Analizi sosyal medya ile oldukça iç içedir. Sosyal medya üzerinden paylaşılan neredeyse bütün yorumlar ve yazılar çok farklı karakterler, linkler, dil bilgisine uymayan

ifadeler, kendini tekrar eden harfler ve sayılar gibi problemlerden ötürü gürültülü veriler olarak isimlendirilirler. Bu problemin çözümü için veriler analiz yapılmadan önce ön işleme aşamalarından geçmektedirler. (Pradhan ve ark., 2016; Poria ve ark., 2018).

Olumsuzluk Ekleri:

Olumlu polariteye sahip birçok kelimeyi kendi içerisinde barındıran cümleler, olumsuzluk katan ekleri beraberinde içerdiğinde yanıltıcı olabilmekte ve yanlış sonuçlar döndürebilmektedir. “Birçok süper özelliği ile dünyanın en iyi oyunları samsung modelleri için üretilse de kesinlikle satın almam” cümlesinde olduğu gibi, olumlu yönde polaritesi olan kelimeler ağırlıkta olmasına karşın anlam itibarıyla olumsuzdur. (Pradhan ve ark., 2016).

3.3. Duygu Analizinde Sosyal Medya ve Önemi

Sosyal ağ, sosyal paylaşım sitesi, sosyal iletişim ağları ile sosyal medya kavramları zaman zaman iç içe geçmiştir. Ancak sosyal medya daha kapsamlı bir kavram olarak kabul edilmiş ve sosyal ağ, sosyal medyanın bir aracı ve çeşidi olarak yer almıştır. Bireysel kullanımının hızla yayıldığı internet, 1980’li yıllarda sosyalleşmeye imkân vermeyen tek yönlü ve içe kapanık Web 1.0 teknolojisini kullanmıştır. Önceleri bireyler interneti sadece bilgi edinmek amacıyla kullanılırken, günümüzde Web 2.0 ile yaşanan gelişmeler ve sosyal medya aracılığıyla kullanıcılar çeşitli içerikler oluşturup bu içeriklerle duygu ve düşüncelerini kitlelere ulaştırmaktadırlar. Web 3.0 ve web 4.0 ile birlikte yapay zekânın da bu sistemin bir parçası olarak kullanılması ile birlikte sosyal medya beklentilerin ötesinde bir hızla gelişerek, diğer iletişim araçları ile entegre olabilme becerisi kazanmıştır.

Günümüz teknolojisiyle paralel olarak gelişen ve her geçen gün gelişmeye de devam eden sosyal medya kavramı üzerine sabit bir tanım bulunmamakla birlikte birçok yazar tarafından tanımları yapılmaktadır. Dağıtmaç (2015), sosyal medyayı, bireylerin topluluklar oluşturarak iletişim içine girdikleri ve orada sosyalleştikleri birer sosyal araç olarak tanımlamaktadır. Sosyal medyanın, kullanıcıların günlük hayatlarındaki sosyal ihtiyaçlarına etkileşim biçimiyle karşılık vermesi; kişilerin yaşamlarının ana yerinde konumlanan bir aygıt haline gelmesini sağlamıştır (Hazar, 2011). Kaplan ve Haenlein (2010) bireylerin içerikler oluşturmalarını ve oluşturdukları içerikleri sunmalarına imkân

sağlayan sosyal medyayı teknoloji tabanlı uygulamalar olarak adlandırmaktadır. Barutçu ve Tomaş (2013), yayınlarında bilgi teknolojileriyle birlikte hayatımıza giren ve bununla beraber günlük yaşantımızda geniş bir alana sahip olan sosyal medyayı, her yaş grubundan çok sayıda kullanıcının sosyal çevrelerini genişletme, farklı insanlarla iletişim kurma, bilgi edinme, bilgilerini ve tecrübelerini başkalarıyla paylaşma ve bunların yanında günlük yaşamda boş kalan zamanlarını değerlendirme gibi birçok amaçla kullandığı iletişim platformları olarak tanımlamışlardır. Basit bir ifadeyle sosyal medya internet kullanıcılarının birbirleri ile iletişime geçmeleri, bireysel yorumlarını ve içeriklerini paylaşmaları için sosyal medya platformlarını kullanmasıdır (Kırtış, 2013).

Sosyal medya yaşantımızın birçok alanında sıklıkla kullanılmaktadır. Milyonlarca insanın sosyal medya platformlarını tercih etmesi sosyal medyayı önemli bir iletişim aracı olarak kabul etmemizde en önemli etkidir. Günümüz teknoloji çağında herhangi bir sosyal medya platformunda paylaşılan bir fotoğraf, video ya da yorum çok kısa bir sürede milyonlarca insan tarafından etkileşim alabilirken aynı zamanda kitleler üzerinde büyük etkiler de yaratmaktadır. Sosyal medya platformları, kullanan herkese açık bir profil oluşturabilme, bilgi paylaşabilme, içerik oluşturabilme ve diğer kullanıcılarla bağlantı kurabilme gibi imkanları sunan çevrimiçi ortamlardır (Ranganathan, 2009). Bu ortamlar, sosyalleşme ve yeni bağlantılar kurma, son dakika gelişmelerinden haberdar olma ve daha az çabayla geniş kitlelere ulaşma, zamanı daha etkin kullanma gibi imkanlar sunar (Ghaisani ve ark., 2017). Bu imkanlar sadece kişiler açısından değil, kurumlar ve işletmeler açısından da önemli bir şanstır. Geleneksel medyanın (TV, Gazete, Dergi vb.,) aksine paylaşılan bir içeriğin insanlar tarafından geri dönüşü yeni medyada (internet, sosyal medya, uygulamalar) çok daha hızlı olabilmekte, kamuoyunun reaksiyonu anında gözlemlenebilmektedir (Tuncer, 2014). Sosyal medyanın her an yeni bir paylaşıma elverişli olması, çok az bir maliyetinin olması, ciddi boyutlarda etkileşim sağlaması kurumlara, kişilere önemli avantajlar sağlamaktadır (Bostancı, 2019).

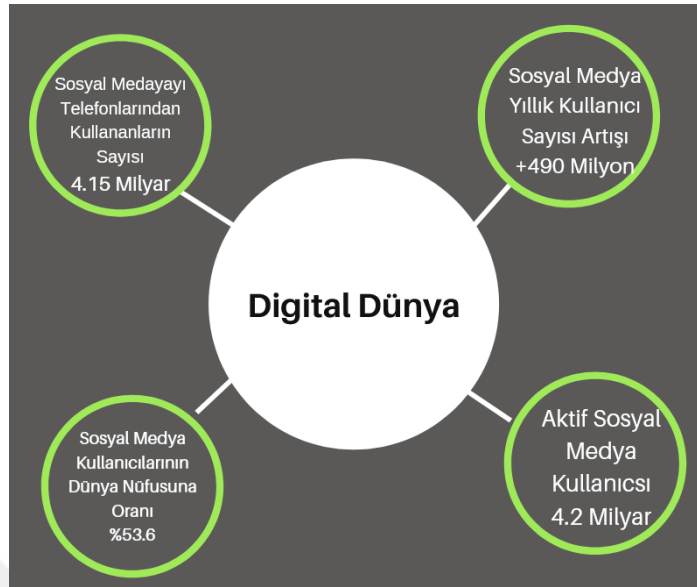
Sosyal medya ile birlikte içinde bulunduğumuz dünya; her geçen gün evrenselleşmektedir. Sonuç olarak zaman, yer ve yakınlık-uzaklık kavramları önemsizleşmektedir. Sosyal medyanın, pratik olması (Hazar, 2011), süre ve birey sınırlaması olmaması (Lynch, 2012), etkileşimli, dinamik ve etkili olması gibi

özellikleri bulunmaktadır. Gürsakal (2009), sosyal medyayı şu kavramlarla açıklamaktadır: İçeriğe katılım yolu ile diğer kullanıcıların ekleme ve yorumlarına imkan vermesi yani katılım kavramı, diğer kullanıcıların içeriklere kolay erişebilmesi ve paylaşabilmeleri yani açıklık, tek yönlü iletişimden çift yönlü iletişime geçilmesi yani karşılıklı konuşma, herhangi bir konu etrafında ya da amaç için birbirini hiç tanımayan kişilerin birlikte hareket etmesi yani topluluk ve siteler arası ya da platformlar arası geçiş olması yani bağlantısallık kavramlarıdır.

Manavcıoğlu (2009), sosyal medyanın özelliklerini aşağıdaki maddelerle özetlemiştir.

- 1- Kullanıcıların zaman ve mekân problemi yaşamadan paylaşım yapabilmesi
- 2- Kullanıcıların üretmiş oldukları içerikleri mobil ya da internet ortamında kolaylıkla paylaşabilmeleri.
- 3- Kullanıcıların diğer kullanıcıların içeriklerini, yorumlarını takip edebilmeleri
- 4- Kullanıcıların hem takip eden hem takip edilen olması
- 5- Samimi bir sohbet mantığına dayanması (Manavcıoğlu, 2009)

Küresel Dijital Rapor (2020) verilerine göre, 7.83 milyar olan Dünya nüfusunun 4.66 milyarı internet kullanmaktadır. Bu nüfusun 4.2 milyarı aktif olarak sosyal medya platformlarını kullanmaktadır. Bu rakamlar Dünya nüfusunun %53.6'sına karşılık gelmektedir. Yine Dünya nüfusunun 4.15 milyarı mobil araçlarından sosyal medyaya erişmektedir. Bir yıllık değişime bakıldığında yıllık bazda 2020 yılı için sosyal medya kullanıcı sayısında 490 milyon kişilik bir artış görülmektedir (Şekil 3.2).



Şekil 3.2. Dünya’da sosyal medya kullanımı.

Küresel Dijital Rapor (2020) verilerine göre, sosyal medya paylaşımı yapan kişilerin toplam nüfusa oranı 2015 yılında %37.7 iken 2020 yılında %52.8 olarak belirtilmiştir. 353 milyon Twitter, 2 milyar 740 milyon Facebook, 2 Milyar 291 milyon Youtube, 1 milyar 221 milyon Instagram kullanıcısının olduğu belirtilen araştırma sonuçları ne kadar büyük bir kitlenin etkileşimde olduğu hakkında bize bilgi vermektedir. Dünya nüfusunun %97’sinin sosyal ağları, mesajlaşma için kullandığı ve %87’sinin sosyal ağlarda etkin olduğu açıklanmıştır. Küresel Dijital Rapor (2020) verilerine göre, Türkiye sosyal medya kullanımında dünya ortalamasının üzerinde olup, 46 ülke arasında 28. sırada yer almaktadır. Dünya’da yaygınlığı ispat edilen sosyal medya kullanımının Türkiye’deki durumu da hemen hemen aynıdır. Hatta rapora göre, sosyal medyada aktif olunan sürenin her geçen yıl arttığı görülmektedir (Küresel Dijital Rapor, 2020).

Ülkemizdeki internet, sosyal medya ve mobil kullanıcı istatistiklerine bakacak olursak; 62 milyon internet kullanıcısı, ülkemiz nüfusunun %74’ünü; 54 milyon sosyal medya kullanıcısı ülkemiz nüfusunun %64’ünü, 7 milyon mobil kullanıcısı, nüfusumuzun %92’sini oluşturmaktadır. Ayrıca ilgili istatistiğe göre; ülkemizdeki internet kullanıcıları günde ortalama olarak 7.5 saat internette vakit geçirmektedir (Dijilopedi, 2020). Bu kadar büyük bir kitlenin kullandığı sosyal medya platformları günümüzde bilgiyi paylaşmanın yanında bilgiyi üreten yerler olmuşlardır. Sosyal medyanın kullanım amacı, yakın ilişkiye sahip bireylerin bir araya gelmesi olmasına rağmen, sosyal medya ile yakın ya da uzak

ilişki demeden herkesle iletişim kurulmaktadır (Uluç ve Yarcı, 2017). Bu açıdan bakıldığında bilgilerin doğrulukları, eksiklikleri, güvenilirlikleri, yönlendirmeleri, amaçları kullanıcılar tarafından sorgulanmaktadır. Yapılan yorumlar, beğeniler, paylaşımlar kullanıcıların içeriğe gösterdikleri tepkiler olarak kabul edilebilir. Bu tepkiler olumlu, olumsuz ya da nötr olarak kabul edilebilir. Sosyal medya platformlarında oluşan bu veri yığınının içerisinde bilgiyi elde etmek, verilerin ayıklanması ve işlenmesi ile olmaktadır.

3.4. Duygu Analizinde Twitter Kullanımı

Twitter, Jack Dorsey, Noah Glass, Biz Stone ve Evan Williams tarafından Mart 2006'da kurulup Temmuz 2006'da kullanıma açılmıştır. Mikroblog sitesi olarak 2006 yılında kurulan Twitter, kullanıcılarına anlık düşüncelerini, güncel konular hakkındaki fikirlerini vb. 140 karakterle sınırlandırarak paylaşımlarını diğer bir deyişle 'tweet' atmalarını sağlamaktadır (Marwick ve Boyd, 2011). Kullanıcı paylaşımlarında 140 kelime ile sınırlı tuttuğu bu sınırlama, 2017 yılı itibarıyla 280 kelimeye çıkartılmıştır. Kullanıcılara blog yazma ve anlık mesaj deneyimi sunmasıyla önemli bir araç haline gelen Twitter, kuruluşunun ardından üç sene içerisinde iş dünyası için de gerekli bir araç olmuştur (Morris, 2009). Ayrıca iş dünyasının yanı sıra hemen hemen tüm devlet kurumları ve bakanların hedef kitleleriyle temas kurmak için bu ortamları tercih ettikleri görülmektedir. Öte yandan mikroblog mantığıyla çalışan Twitter, kullanıcılarına bir nevi yakınlarıyla sürekli iletişim halinde kalmalarına, fotoğraf, video, bağlantı ve mesaj içeren ve "tweet" adı verilen paylaşımlar gerçekleştirmelerine olanak sunmaktadır. Twitterda kullanıcılar ilgi alanlarına göre etiketler üzerinden paylaşım gerçekleştirebilmekte ve bu sayede aynı ya da benzer ilgi alanına sahip topluluk oluşumları meydana gelmektedir (Karabulut ve Küçüksille, 2018). Twitter, aynı konuyu ilgilendiren paylaşımları "#" işaretiyle simgelenen ve "Hashtag" adıyla da bilinen etiketler altında sınıflandırmaktadır. Böylelikle kullanıcılar aynı düşünceye sahip diğer kullanıcılar ile ortak bir paylaşım alanı içerisinde bulunabilmekte, beğeni, yorum ve retweet gibi tepkilerle etkileşim içerisine girebilmektedir (Dağıtmaç, 2015).

Twitterdan toplanan veriler ile doğal dil işleme (NLP) ve veri madenciliği alanında birçok çalışma yapılmıştır. Doğal olarak duygu analizi bu çalışmaların içerisinde yer almaktadır. Twitter üzerinden yapılan duygu analizi konusunda Türkçe olarak az çalışma mevcuttur. (Akgül ve ark. 2016)

3.5. Duygu Analizinde Kullanılan Yaklaşımlar

Duygu Analizi çalışmalarında kullanılan yaklaşımlar makine öğrenmesi yaklaşımı, anlamsal yönelim yaklaşımı ve veri sözlüğü yaklaşımı olmak üzere 3 ana sınıfta incelenebilir.

3.5.1. Duygu Analizinde Makine Öğrenmesi Yaklaşımı

Yapay zekânın gelişmiş bir türü olan makine öğrenmesi, bilgisayar sistemlerinin açık bir biçimde programlanmaksızın istatistiksel teknikler yardımıyla veri işlemesini ve öğrenmesini sağlayan bir yaklaşımdır (Goldberg ve Holland, 1988). Geleneksel yaklaşımlarda bilgisayara veri ve program sağlanmaktadır ve onu işleyip bize çıktı üretilmektedir. Ancak makine öğrenmesi yaklaşımında biz veri ve sonuç sağlamaktayız ve bilgisayarın kendisi girdi-çıkı ilişkisine dayanan bir program çıkarmaktadır. Makine öğrenimi, eğitim veri setinin toplanmasıyla başlamaktadır. Bir sonraki adım da eğitim verilerini kullanarak bir model hazırlamaktadır.

Makine öğrenmesi, denetimsiz (unsupervised), yarı-denetimli (semi-supervised) ve denetimli (supervised) olmak üzere üç başlıkta incelenmektedir.

Denetimsiz Öğrenme:

Denetimsiz öğrenme ile duygu sözlüğü ön bilgi olarak verilmektedir (Onan ve ark., 2016). Başka bir deyişle bir kelime kümesi kullanılmaktadır. Bu kelime kümesi duygu ifadelerine karşılık gelmektedir. Böylelikle metinsel bir ifadenin taşıdığı duygu sistem tarafından otomatik olarak belirlenmektedir. Denetimsiz öğrenmede, etiketlenmemiş veriler kullanıldığından sistemin belirli bir alana bağlı kalması engellenebilmektedir. Denetimsiz öğrenme problemlerine kümeleme problemleri örnek verilebilir (Brownlee, 2016).

Denetimsiz öğrenme modeli hakkında daha fazla bilgi edinmek için hedef veriler altında yatan yapı veya dağılımı modellenmelidir. En önemli denetimsiz öğrenme ise farklı girdi kümeleri oluşturarak uygun kümeye yeni girdi koyabilmektir. Denetimsiz öğrenme problemleri kümeleme ve ilişkilendirme olarak daha fazla gruplandırma yapısına sahiptir.

Denetimsiz öğrenme, bazı basit işletmelerde kullanılan vakalar için aşırı derecede karmaşık olmasına rağmen, insanlar normalde uğraşamayacakları problemleri çözmek için kullanmaktadırlar. Denetimsiz makine öğrenmesi algoritmalarının bazı örnekleri arasında k-aracı kümelemesi, temel ve bağımsız bileşen analizi ve ilişkilendirme kuralları bulunmaktadır.

Yarı Denetimli Öğrenme:

Yarı-denetimli öğrenmede, girdi verilerine karşılık gelen çıktı verilerinden sadece bir kısmı etiketlenmekte ve veriler üzerinde hem denetimsiz hem de denetimli öğrenme algoritmaları problem yapısına göre kullanılabilir. Çok miktarda girdi verisi (X) bulunan ve yalnızca bazı verilerin (Y) etiketli olduğu problemlere ise yarı denetimli öğrenme problemleri denir. Bu problemler hem denetlenen hem de denetlenmeyen öğrenme modeli arasında bulunur. Buna iyi bir örnek olarak yalnızca resimlerin bir kısmının etiketlendiği bir fotoğraf albümü verilebilir. (Örneğin köpek, kedi, kişi) ve çoğunluk etiketsizdir.

Denetimli Öğrenme:

Denetimli öğrenmede ise girdi değerlerine karşılık gelen çıktı değerleri bütünüyle etiketlenmekte, bu sayede girdi ve çıktı arasındaki fonksiyonun eşlenmesi izlenebilmektedir. Bu yöntemlerin düzgün bir şekilde çalışabilmesi için eğitim amacıyla kullanılan veri setinin yeterince büyük olması gerekmektedir (Onan ve ark., 2016). Öğrenme sürecinde etiketlenmiş ve ön tanımlanmış eğitim veri setinin, test veri setini gözetmesinden dolayı bu yöntemler denetimli olarak adlandırılmaktadır. Yani her örnek için hem girdiler hem de o girdiler karşılığında oluşturulması gereken çıktılar sisteme birlikte gönderilmektedirler. Denetlenen bir öğrenme algoritması, eğitim verilerini analiz eder ve yeni örneklerin haritalanması için kullanılabilecek bir çıkarım fonksiyonu üretir. Bu sayede öğrenme algoritmanız, girdilerden ilgili hedeflere bir fonksiyon oluşturur. Eğer hedefler bazı sınıflarda ifade edilirse sınıflandırma problemi olarak adlandırılır. Bu

sınıflandırma problemi doğrusal veya lojistik regresyon ya da çoklu sınıflandırma gibi algoritma modellerini içerir. Buna alternatif olarak hedef uzay sürekli ise regresyon problemi adını alır. Denetimli öğrenme problemleri aslında birer sınıflandırma problemidir. Bu problemler farklı yöntem ve algoritmalar ile çözülebilir.

3.5.2. Duygu analizinde veri sözlüğü yaklaşımı

Duygu analizinde diğer bir yaklaşım veri sözlüğü tabanlı oluşturulan veri sözlükleridir. Bu yaklaşımlar temelde bir duygu sözlüğüne, yani ön derleme yapılmış ve bilinen bir dizi duygu terimi, kavram ve hatta deyimlerden oluşan hazır bir kütüphaneye ihtiyaç duymaktadırlar (Bongirwar, 2015). Bu yaklaşım genelde metinlerin içindeki kelimelerin ontolojik olarak daha derinlemesine anlaşılması gerektiğinde kullanılmaktadırlar. Duygu veri sözlüğü, insanların öznel olarak hisleri ve fikirlerini çıkarımda bulunmak üzere önceden hazırlanmış kelime ve ifade listesine karşılık gelmesiyle duygu analizi yapmaktadır (Palanisamy ve ark., 2013). Bir yazının pozitif ve negatif çoğunluğu onun duygu kutbunun belirlenmesinde büyük etkindir. Bu teknikte eğitim için kullanılması gereken bir veri setine ihtiyaç yoktur. Mesela Gaddarlık kelimesinin duygu değeri -5 değer alırken, harika sözcüğü +5 duygu değeri almaktadır.

3.5.3. Duygu analizinde anlamsal yönelim yaklaşımı

Anlamsal yönelim yaklaşımı, değerlendirmeye tabii tutulan her bir metin içindeki pozitif ve negatif duygu içeren öznel ifadelerin sınıflandırması esasına dayanmaktadır (Xu ve ark., 2012). Bu yaklaşım, araştırmacıya öznel ifadeleri hem negatif ya da pozitif olma, hem de yumuşak ya da sert olma açısından inceleme fırsatı tanıyabilmektedir (Hatzivassiloglou ve McKeown, 1997). Bu yaklaşımda temel olarak iki farklı teknik kullanılmaktadır:

Sözlük Tabanlı Teknikler:

Bu teknikler duygu bilgisi ön tanımlanmış olan bir sözlük (örneğin ‘WordNet’ gibi) yapısına bağlı kalarak metinde yer alan kelimelerin içindeki eş anlamları, zıt anlamları ve hiyerarşileri kullanmaktadır (Jagtap ve Pawar, 2013). Bu özelliği sayesinde sözlük tabanlı teknikler aynı dildeki kelimeler arasında bir ilişki aramada

kullanılabileceği gibi, farklı dillerdeki kelimelerin eşleştirilmesinde de kullanılabilmektedir. Aynı zamanda Latin alfabesinin dışındaki alfabelerle yazılmış olan metinlerin incelenmesinde de bu tekniklerden yararlanılmaktadır (Fu ve Wang, 2010).

Derlem Tabanlı Teknikler:

Metin madenciliğinde Latince temeline dayanan ‘corpus’ kelimesi, dilbilimsel açıdan ‘metin gövdesinin bütünü’ anlamına gelmekte ve bilişim dilinde doğrudan ‘derlem’ olarak yaygın biçimde kullanılmaktadır. Derlem tabanlı teknikler, kelimelerin duygularını saptayabilmek için eş bulunurluk (metinde yan yana olma) örüntülerini bulmayı amaçlamaktadırlar (Bongirwar, 2015).

3.6. Duygu Analizinde Doğal Dil İşleme

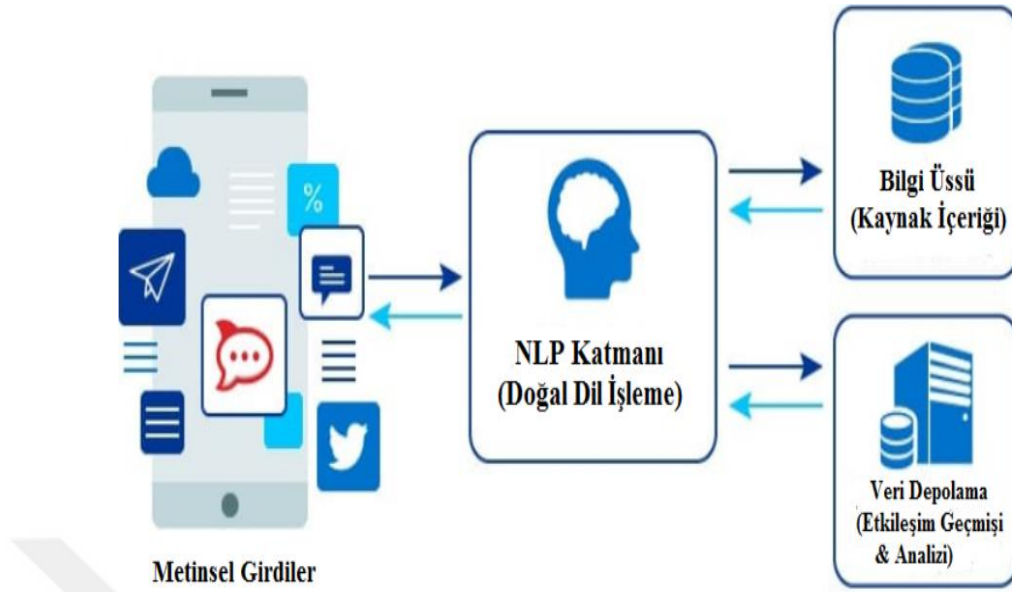
Natural Language Processing (NPL) Türkçe Doğal Dil İşleme (DDİ) olarak adlandırılan yapay zekânın bir dalıdır. Doğal dil işlemede amaçlanan kişilerin yaşamlarındaki günlük döngüde yazarak, konuşarak kullandıkları dili, bilgisayar sistemlerinde birtakım işlemlerden geçirerek bu sistemlerin anlamasını ve anlamlandırabilmesini sağlamaktır. Kısacası insanların kullandıkları dillerle bilgisayar sistemleri arasında anlamsal bir ilişki kurmayı amaçlar.

Doğal dil işleme ile alakalı tarihte ilk çalışma Alan Turing’in yazdığı “Computing Machinery And Intelligence” 1950 yılında yazdığı makale olarak kabul edilir. Sonrasında bu makale içeriği “Turing Test” ismi ile anılmaya başlanmıştır. Doğal dil işleme birçok bilim insanının üzerinde çalışmalarının hala devam ettiği araştırma konusudur.

Doğal diller, insanların fikirleri ifade etmek, aralarında iletişim kurmak için kullandıkları sözcükleri ve yazım kurallarını içermektedir. Kişiler kavramlarla düşünür, bunun üzerinden birbirleriyle haberleşirler. Fakat bilgisayar sistemleri sayısal kurallarla işleyen makinelerdir. Bu makineler sayısal işlemleri çok kısa bir zaman diliminde yapabilmelerine rağmen, insanların basitçe kavrayabildiği karmaşık dilleri yorumlayamazlar. Fakat bilgisayarlar yorumlama konusunda eksik kalsalar bile sınıflandırmada ve çeviride oldukça başarı göstermektedirler. Makine öğrenimi ve bu alandaki gelişmelerle birlikte, bir makineye iletişimi çözme, yazı dilini yorumlama

konularında öğrenme becerisi katmaktır. Makine öğrenmesinin çalıştığı bu alanına “Doğal Dil İşleme” denir. Doğal dil işleme bilgisayar sistemlerinin doğal dil metinlerini veya seslendirmeleri kavrayabilmek için nasıl bir yol izlemesi gerektiğini araştıran bir disiplindir. Bu konuda çalışan araştırmacılar, insanların doğal dili nasıl anladığı ve kullandıklarıyla alakalı elde ettikleri bilgilerle, bilgisayarların doğal dilleri anlaması ve çözümlemesini sağlamak amacıyla değişik araçlar ve teknikler geliştirmektedirler (Chowdhury, 2003). Bilgisayar sistemlerinin doğal dilleri anlamlandırabilmesi ve kavrayabilmesi için, dilin modellenen matematiksel bir çerçeveye oturtulması gerekir.

Doğal Dil işleme (Şekil 3.3), otomatik özetleme, söylev analizi, makine çevirimi, biçimsel bölütleme, soru cevaplama, doğal dil üretimi gibi alanlarda kullanılmaktadır. Otomatik özetlemede bir veya birden fazdan kaynaktan alınan, kısa, önemli kısımların yer aldığı bir yapının oluşmasını sağlayan çalışma alanıdır. Söylev analizi bir söylemin ne istediğini araştırarak yazının yapısını, cümleler arası ilişkileri de gözeterek tanımlamaya çalışır. Makine çevirimi insan yardımı olmadan belirli bir metni hedeflenen dile çevirip yorumlama işlemi yapar. Biçimsel bölütleme dünya üzerindeki çeşitli dillerin yapılarının farklı olması sebebiyle bu dil yapılarını çözümlemeyi hedefler. Soru cevaplama alanında ise kişinin sorduğu bir soruya yapay zekânın yardımıyla veri tabanında tanımlanmış cevaplardan en doğrusuyla cevap verme işlemi yapılmaktadır. Arama motorları bu mantıkla çalışırlar. Makinelerin insan dili ile konuşmasını amaç edinen doğal dil üretimi alanında APPLE-SIRI çalışması örnek gösterilebilir.



Şekil 3.3. Basit bir NLP yapısı.

3.6.1. Kelimelere ayırma

Bir metin içerisindeki yapılandırılmamış verileri, noktalama işaretlerinden ayırarak ‘token’ adı verilen tekli kelime parçalarına ayırma sürecidir. Bu tekli kelime parçalarına ‘fiş’ adı da verilmektedir. Burada belirlenen ‘token’ yapıları, metin içerisindeki bazı ileri düzey söz dizimlerinin analizlerinin girdileri olarak kullanılmaktadır.

3.6.2. Morfolojik-ek tabanlı köke indirgeme

Dokümanlar içindeki türetilmiş ya da eklemeler yapılmış kelimelerin köke indirgenmesi morfolojik indirgeme (lemmatization), ekleri ayrıştırılacak şekilde köke ayrılmasına ek tabanlı köke indirgeme adı verilmektedir. Örneğin ‘dolabımız’ kelimesinin kökü ‘dolap’ iken ek tabanlı kökü ‘dolab’ olacaktır. Google, Yandex gibi arama motorları aranan kelimeleri içerisinde barındıran web siteleri ararken arka planda köklerine indirgenmiş kelime yapılarını kullanmaktadırlar. Bundan dolayıdır ki herhangi bir aramada, türetilmiş bir kelimenin kökü ile eş sesli olan başka bir kelimeyle ilgili aramalarda beraber gösterilmektedir. Buna örnek verecek olursak Google ya da Yandex üzerinde arama kısmına ‘dizin’ yazıldığında, ‘diz’ ile ilgili sonuçların da beraber gelebileceği söylenebilir.

3.6.3. Metin parçası etiketleme

Bu yöntem ile yazılmış olan bir metin, uygun bir yazılım sayesinde okunmakta ve fiil, sıfat, isim gibi metnin içerisinde ayrılan bazı öğelerin anlamca aynı dilde yer alan hangi kelimeye karşılık geldiği bir yazılım vasıtasıyla tespit edilmektedir. Bu yöntem sayesinde fiil, isim ve sıfat gibi kategorilerde ilgili kelimeler toplanabilmekte ve etiketlenebilmektedir.



4. MATERYAL VE YÖNTEM

4.1. Materyal

Tezin bu bölümünde, Twitter'daki mobil oyun verileri ile duygu analizinin gerçekleştirilmesinde kullanılan ve aşağıda isimleri sunulan makine öğrenmesi algoritmaları sunularak her biri ayrıntılı bir şekilde açıklanmıştır.

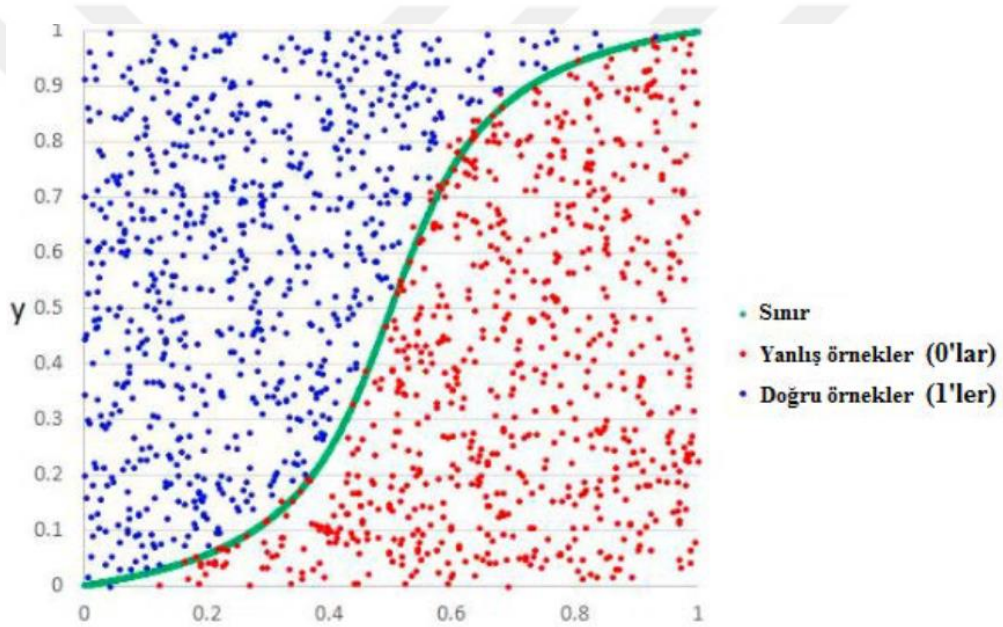
- Lojistik Regresyon (Logistic Regression-LR)
- Rastgele Orman (Random Forest- RF)
- K-Komşu Sınıflayıcısı (K-Neighbors Classifier-KNN)
- Doğrusal Destek Vektör Sınıflayıcısı (Linear Support Vector Classifier-Linear SVC)
- Çok terimli Naïve Bayes Sınıflayıcısı (Multinomial Naïve Bayes Classifier-Multinomial NBC)
- Bernoulli Naïve Bayes Sınıflayıcısı (Bernoulli Naïve Bayes Classifier-Bernoulli NBC)
- Sırt Sınıflayıcısı (Ridge Classifier-RC)
- Pasif-Agresif Sınıflayıcısı (Passive-Aggressive Classifier- PAC)
- Çok Katmanlı Algılayıcı Sınıflayıcısı (Multi-Layer Perceptron Classifier-MLPC)
- Oylama Sınıflayıcısı (Voting Classifier-VC)
- Evrişimsel Sinir Ağı (Convolutional Neural Network-CNN)
- Uzun Kısa Süreli Bellek (Long Short-Term Memory-LSTM)

4.1.1. Lojistik regresyon (LR)

Regresyonun bir formu olan Lojistik regresyon, gözlemleri ayrı bir sınıfa atamak için kullanılan bir sınıflandırma algoritmasıdır ve olasılık hesabıyla çalışır (Norton ve ark. 2019). LR, ismini bağımlı değişkene uygulanan lojistik dönüştürme işleminden almaktadır. Yalnızca iki olası sonuç üretir (Hair ve ark., 2006).

Denetimli bir metot olan lojistik regresyonda amaç, tıpkı doğrusal regresyonda olduğu gibi her bir girdi değişkenini ağırlıklandıran katsayı değerlerini bulmaktır. LR'de farklı olan, çıkış değişkeninin tahmini esnasında fonksiyonun doğrusal olmayan bir

yapıya dönüşmesidir ki bu fonksiyona lojistik fonksiyon adı verilmektedir. Fonksiyon örnekleri sıfır ve bir olmak üzere iki sınıf şeklinde ayrılmakta, bu ayrım lojistik fonksiyonun oluşturduğu sınır yardımıyla yapılmaktadır. Böylelikle fonksiyon yardımıyla ayrılan sınıflar 1 (Doğru, geçti, vb.) veya 0 (Yanlış, kaldı, vb.) olarak kodlanmış verileri içermektedir. Bağımlı değişken “başarılı-başarısız”, “az-orta-çok”, “olumlu-olumsuz” gibi kategorize edilebilecek veri topluluğundan oluşuyorsa bu analiz yöntemi tercih edilebilir. Bağımsız değişkenler için normal dağılıma gereksinim duyulmaması LR’nin uygulanmasını kolaylaştırmaktadır. LR uygulanan bir veri setindeki değişkenlerin dağılımları Şekil 4.1’de gösterilmektedir.



Şekil 4.1. LR dağılımı için bir örnek grafik.

Kleinbaum ve Klein (2002), LR’nin çalışma prensibinde temel alınan noktayı z hangi değeri alırsa alsın lojistik fonksiyon 0 ile 1 arasında bir sonuç üretir olarak açıklamışlardır (Eşitlik 4.1).

$$f(z) = \frac{1}{1+e^{-z}} \quad (4.1)$$

LR’nin amacı, iki yönlü karakteristik (bağımlı değişken = yanıt veya sonuç değişkeni) ile ilgili bir dizi bağımsız (öngörücü veya açıklayıcı) değişken arasındaki ilişkiyi tanımlamak için en uygun (henüz biyolojik olarak makul) modeli bulmaktır. LR,

lojit dönüşümünü tahmin etmek için bir formülün katsayılarını ve standart hatalarını üretir. LR’de, bağımlı değişkenlerin varlığının olasılığını hesaplamak için lojit dönüşümü (Eşitlik 4.2) kullanılır.

$$\text{lojit}(p) = b_0 + b_1x_1 + b_2x_2 + b_3x_3 + \dots \dots b_kx_k \quad (4.2)$$

Eşitlik 4.2’deki x_k bağımsız değişkeninin katsayısı olan b_k değeri, bağımsız değişkenin y bağımlı değişkeni olan etiket sınıfının üzerine yaptığı etkiyi gösterir. p değeri ise, karakteristik ya da bağımlı özelliğinin var olma olasılığıdır. Dolayısıyla, bir bağımlı değişkene ait kategorinin olma olasılığı ile olmama olasılığı arasındaki ilişkiye göre katsayılar belirlenir (Bewick ve ark., 2005).

Lojit dönüşümündeki olasılık hesaplaması için (Eşitlik 4.3) kullanılır. Lojit fonksiyon ile doğrusal bir ilişki kurulmuş olur. Dolayısıyla, tek değişkenli doğrusal bir denklem elde edilmiş olur. Aşağıdaki modelde x bağımsız değişkeni, b_0 olasılık oranını göstermektedir (Taşkın ve Turanlı, 2019).

$$\text{logit } p(x) = \ln \left[\frac{\pi(x)}{1-\pi(x)} \right] = b_0 + b_1x_1 \quad (4.3)$$

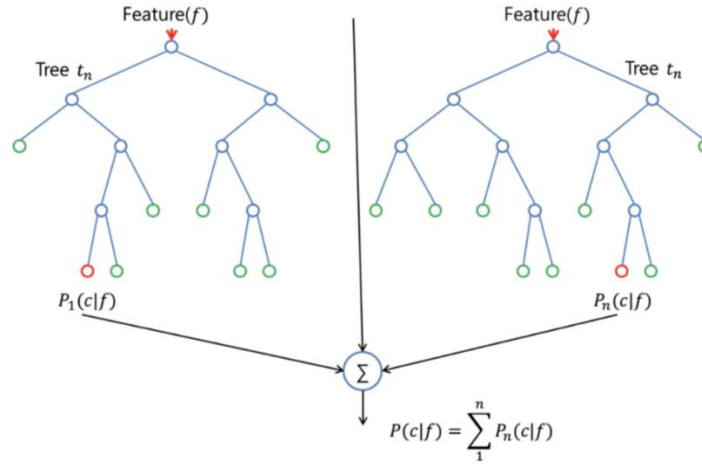
4.1.2. Rastgele orman (RF)

RF, birçok karar ağacının birbirinden bağımsız oluşturulduğu, kullanımı kolay bir makine öğrenmesi algoritmasıdır. Sınıflandırma görevleri için fazlasıyla tercih edilen algoritmalarından biridir. Sınıflandırma problemlerinde uygulandığı gibi regresyon problemlerinde de kullanılabilir olması avantaj sağlamaktadır. Ayrıca karar ağacı algoritmalarıyla karşılaştırıldıklarında daha iyi sonuçlar vermiştir (Akar, 2013).

Karar ağaçları algoritmalarının birçok çeşidi mevcuttur. C4.5, ID3, Hoeffding Tree, Logistic Model Tree ve Random Forest algoritmaları bunlardan bazılarıdır (Aytuğ, 2015). Bu tezde, Karar Ağaçları algoritmalarından Rastgele Orman algoritmasının tercih edilmesinin en temel nedeni aşırı öğrenmeyi daha iyi engelleyebilmesidir. Rastgele Orman ile diğer Karar Ağaçları algoritmaları arasındaki farklar incelendiğinde; bir kullanıcının mobil oyun oynarken ekranda karşısına çıkan reklama dokunup

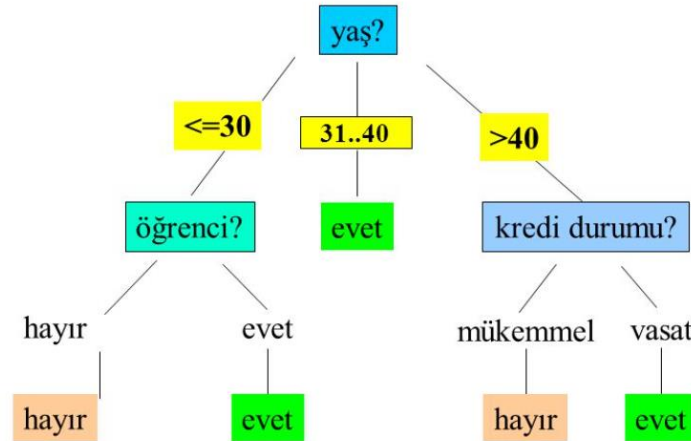
dokunmayacağını tahmin etmek istiyorsak, geçmişte dokunduğu reklamlar ve kararı hakkında bilgi veren bazı verileri toplanmalıdır. Toplanan verinin özellikleri ve sahip olduğu etiketlerin bir karar ağacına koyulması ile kurallar oluşturulur ve böylece bireyin reklama dokunup dokunmayacağı ile ilgili bir tahmin üretilebilir. Rastgele Orman, Karar Ağaçları oluşturabilmek için farklı olarak gözlemleri ve verinin içerdiği özellikleri rastgele olarak seçer ve ortaya çıkan sonuçların ortalamasına dikkate alır. Diğer bir fark ise, Rastgele Orman algoritmasının daha küçük ağaç alt kümeleri oluşturarak çalışması ve böylece aşırı öğrenmeyi daha iyi bir şekilde engellemesidir.

Sadece iki adet karar ağacından oluşan bir Rastgele Orman örneğinde (Şekil 4.2), her bir niteliğe ait karar ağacı yapısı en alt düğümlere kadar indirgenmekte ve yapılardan ortak sistem ağırlıkları (Σ) hesaplanabilmektedir.



Şekil 4.2. RF için örnek uygulama.

Makine öğrenmesinin bir yöntemi olarak önerilen karar ağaçları (Şekil 4.3) (Han ve ark., 2012); her bir yapraksız düğümün (nonleaf node) bir özneliği (attribute) belirttiği, her bir dalın (branch) bir test çıktısını gösterdiği ve her bir yaprak düğümünün (son düğüm) bir sınıf etiketini nitelediği akış şemalarıdır.



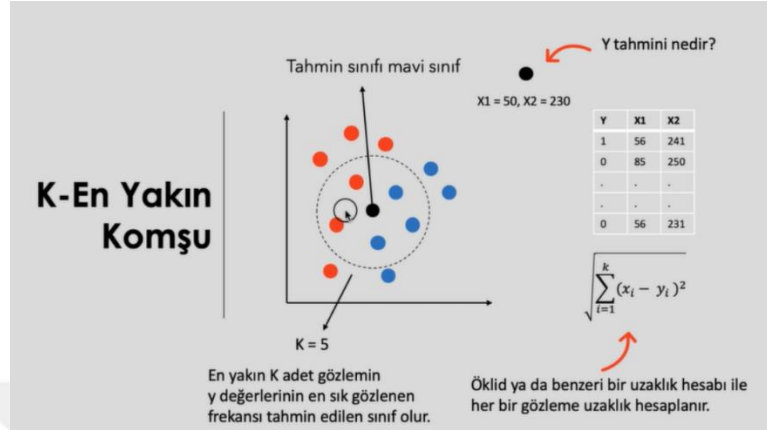
Şekil 4.3. Karar ağacı için örnek uygulama.

4.1.3. K-En yakın komşu sınıflayıcısı (KNN)

İlk kez 1950’li yılların başında geliştirilen KNN, yığın veri setlerinde kullanılan güçlü öğrenme sürecine sahip olan örnek tabanlı bir sınıflandırma algoritmasıdır. Bu algoritmada, ön eğitim işlemine ihtiyaç duyulmazken test verileri her seferinde eğitim kümesi kullanılarak sınıflandırılır (Kaynar ve ark., 2016). KNN sınıflandırıcısı analizlerde kullanılması basit bir yapıdadır. Ayrıca parametrelerin iyi seçilmesi ve uygun sayıda etiketlenmiş eğitim verisini içinde barındırması şartlarını karşılayabiliyorsa diğer makine öğrenmesi algoritmalarına göre daha iyi sonuçlar verebilmektedir. KNN yöntemi, verilerin birbirleriyle karşılaştırılarak öğrenilmesi temeline dayanır (Murphy, 2012). Etiketlenmemiş bir test verisi modele verildiğinde, KNN eğitim veri kümesindeki en yakın k komşuya bakarak çoğunluk durumundaki etikete göre atamasını yapar. (Han ve ark. 2012).

KNN sınıflandırıcısı, Şekil 4.4’te gösterilen etiketli veri seti üzerinden sınıflandırılacak X örneğine en yakın ‘ K ’ tanesini belirleyerek sınıflandırma yapar. Verilen değere en yakın K adet eleman alınarak aralarındaki Öklid uzaklığı (Eşitlik 4.4) hesaplanır. Bu hesaplamada, en yakın değere sahip K tane elamanının sınıfına bakılarak en çok hangi sınıfta eleman mevcutsa ona göre sınıflama yapar (Akyel ve Seçkin, 2011).

$$\text{dist}(X_1, X_2) = \sqrt{\sum_{i=1}^n (x_{1i} - x_{2i})^2} \quad (4.4)$$

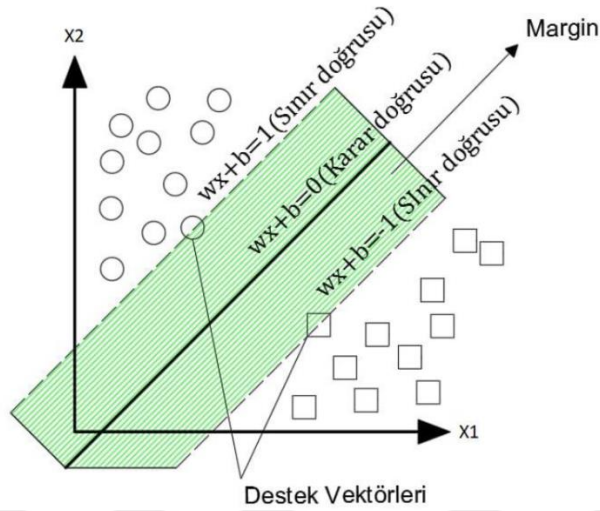


Şekil 4.4. KNN ile sınıflandırma örneği.

4.1.4. Destek vektör makineleri (SVM)

SVM, iki boyutlu uzayda doğrusal, üç boyutlu uzayda düzlemsel ve çok boyutlu uzayda hiperdüzlem şeklindeki ayırma mekanizmaları ile veriyi iki ya da daha çok sınıfa bölme yeteneğine sahiptir (Güran ve ark. 2014). Doğrusal olarak ayrıştırılabilen sınıfların belirlenmesinde yaygın biçimde kullanılan bu yöntem, ‘kernel’ fonksiyonları sayesinde doğrusal olarak ayrıştırılamayan girdi uzayını daha yüksek boyutlu lineer olarak ayrıştırılabilen bu uzaya taşıyarak, doğrusal olmayan verilerin sınıflandırılmasında, yoğunluk tahmininde ve kümeleme amacı ile kullanılmaktadır. SVM birçok uygulama için en iyi sınıflayıcı algoritmalarından biri olarak tanımlanmaktadır (Ballı ve ark. 2018). SVM, denetimli bir öğrenme modeli olup duygu analizi alanında sıkça kullanılmaktadır (Çoban ve Özyer, 2018). SVM, sınıfları birbirinden ayıran doğrusal bir düzlem bulmaya çalışan sınıflandırıcıdır. Eğer iki öznitelige sahip bir veri olursa, sınıfları birbirinden ayıran düzlem değil doğrusal bir çizgi olur. Öznitelik sayısı üç tane olan bir veride ise, sınıfları birbirinden ayırmak için bir düzleme ihtiyaç olacaktır. Eğer öznitelik sayısı dört ve üstü olan bir veride örnekleri birbirinden ayırmak için bir hiper düzleme ihtiyaç olacaktır. SVM, doğrusal düzlemi destek vektörlerine göre tanımlanan maksimum aralık ile bulmaya çalışır (Atasoy ve Tabak, 2018).

SVM'lerin gerçek hedefi, kendilerine en yakın noktalar arasındaki mesafeyi en yüksek seviyeye çıkararak buna uygun hiperdüzlemi bulmaktır. Şekil 4.5'te gösterildiği üzere; mesafenin en yüksek değere ulaşmasıyla optimum düzlem elde edilmektedir ve böylece sınırları gösteren noktalar destek vektörleri olarak isimlendirilmektedir (Karakaynak, 2014).



Şekil 4.5. SVM için destek vektörleri ve hiper düzlemi.

Şekil 4.5'te negatif sınıftaki veriler kare sembollerle, pozitif sınıftaki veriler yuvarlak sembollerle ifade edilmektedir. SVM, pozitif ve negatif sınıfların arasına doğrusal bir düzlem koymaya çalışarak noktalı çizgilerle gösterilmiş kenar boşluğunu en yüksek değerde olacak şekilde ayırmakta ve karar doğrusu bu sınırların arasına yerleştirilmektedir. Yeşil vurgulu marjinin tamamlanması için bu sınırlara en yakın veriler (destek vektörleri) kullanılmaktadır. Buna göre $f(x) = 0$ hiper düzlemini bulmak için gerekli denklem Eşitlik 4.5'de verilmiştir.

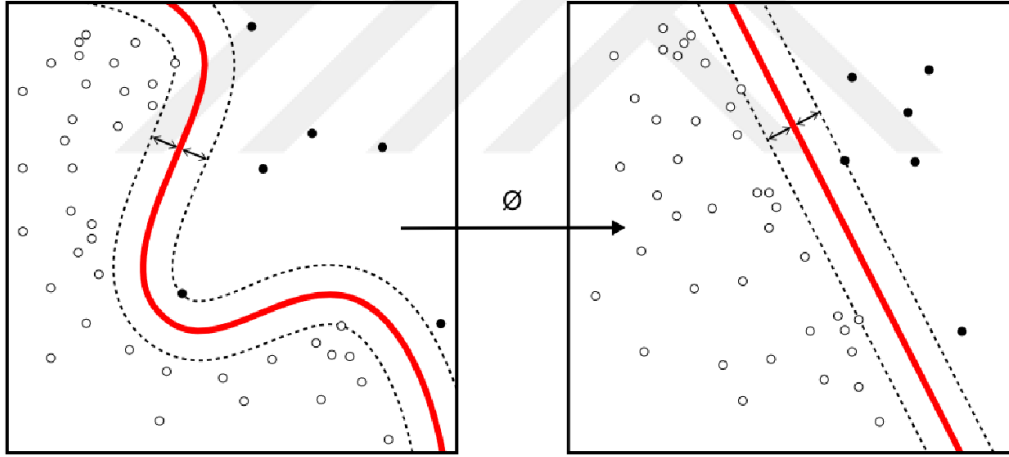
$$f(x) = w^T x + b = \sum_{j=1}^M w_j x_j + b = 0 \quad (4.5)$$

Burada w , M boyutlu bir vektörü, b ise bir sabiti temsil eder.

Doğrusal SVC (Linear SVC): Temel amaç, iki sınıfın verileri arasında bir ayırıcı çizgi veya hiperdüzlem oluşturmaktır. Doğrusal SVC, doğrusal çekirdeklere

dayanmaktadır ve verileri bölerek veya kategorilere ayırarak, onları en uygun hiperdüzleme dönüştürür (Juneja ve Ojha, 2017).

Linear SVC, doğrusal olmayan veri setini doğrusal olarak ayrılabilirliğe temeline dayanır ve verilerin doğrusal olarak ayrıldığı Şekil 4.6’de gösterildiği gibi verileri ayıran bir hiperdüzlem çizer.



Şekil 4.6. Doğrusal olmayan veri setinin doğrusala çevrilmesi.

4.1.5. Naïve bayes sınıflandırıcısı (NBC)

NB sınıflandırıcısı, adını Matematikçi Thomas Bayes’den alan bir sınıflandırma algoritmasıdır. NB sınıflandırıcısı, olasılık ilkelerine göre tanımlanmış bir dizi hesaplama ile sisteme sunulan verilerin sınıfını yani kategorisini tespit etmeyi amaçlar. Naïve Bayes algoritmaları, Bayes teoremine dayanan istatistiksel sınıflandırma algoritmalarıdır ve her bir olayın gerçekleşme olasılıklarına bağlı olarak iki olayın koşullu olasılığını bulmamıza yardımcı olur. Bu algoritma her özneliğin diğerinden bağımsız olduğu ilkesine göre çalışmaktadır.

NB sınıflandırıcısı, aşağıda sunulan iki ihtimalden oluşur:

- Her bir sınıfın oluşturulabilme ihtimali,
- Belirli bir x değeri için her bir sınıftaki durumsal olasılık ihtimali.

Oluşturulan olasılıksal model tahmin işlevi için Naïve Bayes teoremini kullanmaktadır. Özellikle normal dağılım özelliği gösteren veriler için sınıf olasılıkları

oldukça kolay biçimde tahmin edilebilmektedir. Bu algoritmayı aşağıdaki matematiksel modelde (Eşitlik 4.6) ifade etmek mümkündür (Kaynar ve ark., 2016).

$$P\left(\frac{S_i}{x}\right) \times p(x) = p\left(\frac{x}{S_i}\right) \times P(S_i) \quad (4.6)$$

$P(S_i)$: Sınıflandırma yapılacak i sınıfının öncel olasılığı

$P\left(\frac{S_i}{x}\right)$: i sınıfının ardıl olasılığı

$P(x)$: x 'in olasılık yoğunluk fonksiyonu

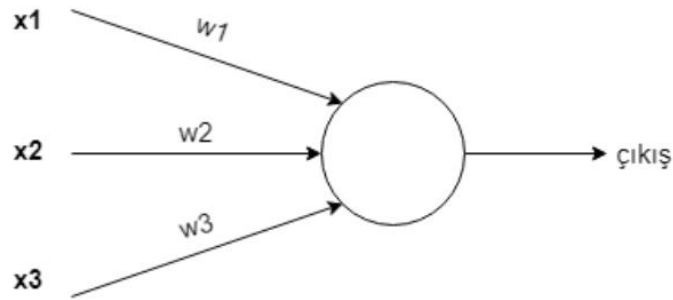
$p\left(\frac{x}{S_i}\right)$: x 'in i sınıfına bağlı koşullu olasılık yoğunluk fonksiyonu

Çok terimli naïve bayes sınıflandırıcısı (Multinomial naïve bayes): Çok terimli Naïve Bayes (NB) sınıflandırıcı, bir NB sınıflandırıcı çeşididir ve metin belgelerinin sınıflandırılması için geliştirilen bir algoritmadır. (Xu ve ark., 2017). Çok terimli Naïve Bayes sınıflandırıcısı, kelimenin bir belgede kaç kez geçtiği anlamına gelen terim sıklığı kavramı üzerinde çalışır. Bu model, kelimenin bir belgede geçip geçmediğini ve geçiyorsa o belgedeki geçme sıklığını belirtmektedir. Bazen, bir belgede bir terim birçok kez kullanılabilir. Terim sıklığını artıran ancak aynı zamanda belgeye potansiyel olarak hiçbir anlam katmayan bu terimler durak kelimeler (stopwords) olabilmektedir. Bu nedenle bu algoritmadan daha iyi doğruluk elde etmek için önceden bu tür kelimelerin çıkarılması gerekmektedir. Naïve Bayes'den farklı olarak sözcükler için çok terimli dağılım kullanan Çok Terimli Naïve Bayes sınıflayıcısı, metin sınıflandırmada Naïve Bayes'e göre daha yüksek başarı oranı sergilemektedir (Tekin ve Tunalı, 2019).

Bernoulli naïve bayes: Bernoulli Naïve Bayes Sınıflandırıcısı, terimin bir belgede bulunup bulunmadığına yönelik, ikili kavram üzerinde çalışmaktadır ancak bu sınıflandırıcı Çok Terimli Naïve Bayes Sınıflandırıcısının aksine, terim sıklığı hakkında bilgi vermemektedir. Ayrıca her bir sözcüğün tekrarlanma sayılarını hesaplamada kullanmaktadır (Singh ve ark., 2019). Bernoulli Naïve Bayes, çok değişkenli Bernoulli dağılımlarına göre dağıtılan veriler için Naïve Bayes eğitim ve sınıflandırma algoritmalarını uygulamaktadır. Örneklerin ikili (1 ve 0) özellik vektörleri olarak elde edilmesi gerekmektedir. Bu sonuca göre, terimin incelenen belgeye ait olması durumunda 1'e eşit, yoksa 0'a eşit bir boolean değer üretmektedir.

4.1.6. Yapay sinir ağıları (ANN)

Bir sinir ağına ana fikri, girdi verilerinin doğrusal kombinasyonlarından özellikler türetmek ve ardından çıktıyı bu özelliklerin doğrusal olmayan bir fonksiyonu olarak modellemektir (Russell ve ark., 2010). Sinir ağıları, tipik olarak, yönlendirilmiş bağlantılarla birbirine bağlanan düğümlerden oluşan bir ağ diyagramı ile temsil edilir. Düğümler katmanlar halinde düzenlenir ve en çok kullanılan sinir ağına yapısı üç katmandan oluşur: Girdi katmanı, gizli katman ve çıktı katmanı (Hastie ve ark., 2001). Düğümler yalnızca bir yönde bağlı olduğundan ileri beslemeli ağ olarak da sınıflandırılır. YSA'da öğrenme nöronlar arasında bulunan bağlantılar ile yapılmaktadır. Bu yapay nörona perceptron (algılayıcı) adı verilmektedir. Nöron, iki adımda bir çıktı değeri üreten basit bir matematiksel modeldir. İlk olarak, nöron girdilerinin ağırlıklı bir toplamını hesaplamakta ve sonrasında çıktısını elde etmek için bu toplama bir aktivasyon fonksiyonu uygulamaktadır (Russell ve ark., 2010). Bir perceptron birçok girdi almakta $x_1, x_2, \dots$ ve tek bir çıktı üretmektedir. Perceptron için girdi ve çıktıları gösteren şekil aşağıda sunulmuştur (Şekil 4.7).



Şekil 4.7. Örnek perceptron (algılayıcı) modeli.

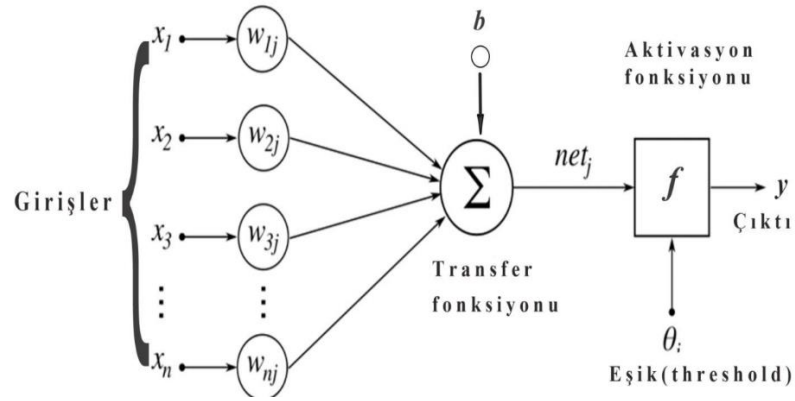
Şekil 4.7’de, perceptron, x_1, x_2, x_3 olmak üzere üç girişe sahiptir. Bu girişler farklı sayılarda olabilmektedir. w_1, w_2, w_3 (ağırlık) değerleri girdilerin çıktı üzerindeki etkisini ayarlayabilmek için kullanılmaktadır. Bu ağırlıklar pozitif, negatif veya nötr olabilmekte ve ağırlığı nötr olan nöronların çıktı üzerinde herhangi bir etkisi bulunmamaktadır.

Nöronun çıktısının 0 veya 1 olması göz önünde bulundurularak Eşitlik 4.7’de belirtilen ağırlıklı toplam değeri, diğer eşik değerlerle karşılaştırılarak belirlenmektedir. Eşik değeri reel bir sayıdır ve Eşitlik 1’deki formül ile yazılabilir:

$$Çıktı = \begin{cases} 0, & \text{eğer } \sum_j w_j x_j \leq \text{eşik(threshold) değeri} \\ 1, & \text{eğer } \sum_j w_j x_j \geq \text{eşik(threshold) değeri} \end{cases} \quad (4.7)$$

Eğitim kümesinin verileri YSA’ya girdi olarak verilerek üretilecek çıktının gerçek etiket değerlerine yaklaşması amaçlanmaktadır. Bu aşamalarda ki geçen süre öğrenme olarak kabul edilmektedir.

YSA’nın yapı blokları ve işlevi olan nöronların yapısı Şekil 4.8’de gösterilmiştir.

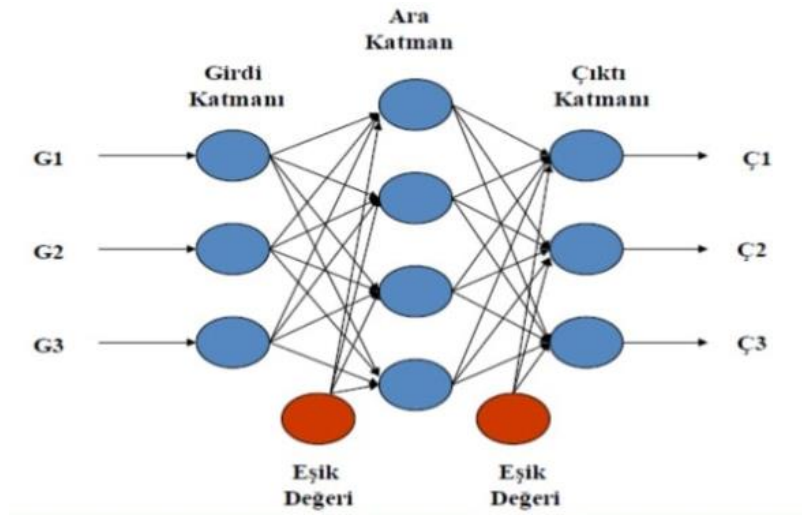


Şekil 4.8. Yapay sinir ağı yapı blokları ve nöron yapısı.

Transfer fonksiyonu: Bu fonksiyon bir nörona gelen net girdiyi Eşitlik 4.8’de bulunan denkleme göre hesaplamaktadır.

$$net = \sum w \times x \quad (4.8)$$

Çok Katmanlı Algılayıcılar Sınıflayıcısı (MLPC): Çok Katmanlı Algılayıcılar (ÇKA) ileri beslemeli YSA’nın bir türüdür (Werbos, 1974). Hatayı ağı yaydığından ‘Hata Yayma Modeli’ ismiyle de bilinmektedir. ÇKA özellikle sınıflandırma ve genelleme yapma durumlarında etkin çalışır. ÇKA’nın yapısı Şekil 4.9’da gösterilmiştir.



Şekil 4.9. Çok katmanlı algılayıcıların yapısı.

Girdi Katmanı: Gelen girdiler alınarak ara katmana gönderilmektedir. Bu katmanda bilgi işlenmemekte ve yapısı değişmeden sonraki katmana gitmektedir.

Ara Katmanı: Girdi katmanından gelen bilgiler burada işlenerek çıkış katmanına gönderilmektedir. Birden fazla ara katman bulunabilmektedir.

Çıkış Katmanı: Gelen bilgileri işleyerek girdilere karşılık ağın ürettiği çıkışları belirlemektedir.

Birçok giriş için bir nöron yeterli gelmediğinde, gizli katmanlar devreye girmektedir. Perception (algılayıcı) doğrusal bir fonksiyonla iki parçaya bölünebilen problemlerde kullanılabilir. Giriş katmanı gelen verileri alarak sayısı en az bir olan ara katmana gönderir. Kullanılan ara katman sayısı probleme göre değişmekte ve ihtiyaca göre ayarlanmaktadır. Her katmanın çıkışı bir sonraki katmanın girişi kabul edilip bu yolla çıkışa ulaşılmaktadır. Her bir nöron bir sonraki katmanda bulunan bütün nöronlara bağlı olup, nöron sayısı ihtiyaca göre belirlenmektedir.

Modelde aktivasyon fonksiyonu olarak herhangi bir matematiksel fonksiyon kullanılabilir. Ancak Sigmoid, tang, lineer, threshold ve hard limiter fonksiyonları en çok kullanılan fonksiyonlardır. ÇKA ile birlikte YSA geniş kullanım alanına ilgi hızlı bir şekilde artmıştır (Almaghrabi ve Chetty, 2020).

4.1.7. Sırt sınıflayıcısı (RC)

Ridge regresyon yöntemi ilk defa 1970 yılında Hoerl ve Kennard tarafından geliştirilmiştir. Ridge sınıflandırıcısı, sıradan doğrusal regresyonu iki açıdan genişleten ve genelleyen doğrusal bir yöntemdir. İlk olarak, eklenen ridge, yöntemin genelleme yeteneklerini geliştirmekte ve ikinci olarak, gerçek değerli etiketlerden farklı olarak ikili etiketlerle ilgilenmektedir (Clausen ve Nickisch, 2018).

En küçük kareler yöntemine benzer bir çözüm yöntemi kullanmaktadır. Bu yöntemde, standart formdaki değişkenlerin oluşturdukları $(X^T X)$ matrisinin köşegen elemanlarına küçük ve pozitif bir sabit eklenmektedir. Pozitif sabitin eklenmesindeki amaç matris sayısını küçültmek içindir. $k=0$ kabul edildiğinde en küçük kareler sonucu ile ridge sonucunun eşdeğer olduğu görülmektedir. Sonrasında regresyon katsayı tahminleri yapılmaktadır. Bu şartlara göre belirlenen ridge regresyon çözümü Eşitlik 4.9'da sunulmuştur (Sakallıoğlu ve Kaçıranlar, 2008).

$$\beta(k) = (X^T X)^{-1} X^T y, \quad k \geq 0 \quad (4.9)$$

Burada $k \geq 0$ yanlılık parametresini temsil etmektedir.

Çoklu regresyon analizinde, bağımlı olmayan değişkenler arasında doğrusal bağlantıların olması sıklıkla oldukça karşılaşılan bir problemdir. Bu problem sonucunda bağımsızlık varsayımının doğruluğu azalmaktadır. Dolayısıyla bağımlı olmayan değişkenlerin, bağımlı değişken üzerindeki etkilerini analiz etmekte zorluklar yaşanmaktadır. Bu problemin çözülebilmesi için önerilen yöntem, modelde bulunan değişkenlerin tümünün regresyon katsayılarını yanı olarak tahmin eden ridge sınıflayıcı yönteminin kullanılmasıdır (Topal ve ark., 2010). Bu yöntemde, veri setinde bulunan metinlerin sınıflandırılması işleminde, biraz yanlılık eklenerek mevcut hata oranının düşürülmesi amaçlanmaktadır.

4.1.8. Pasif agresif sınıflayıcı (PAC)

Çevrimiçi öğrenme algoritmaları arasında en çok kullanılan pasif agresif sınıflayıcı (PAC) tekniği, her yeni örnek için doğrusal bir model sunarak, mevcut ağırlık vektörüne yakın bir ağırlık vektörü bulmayı amaçlamaktadır. (Jorge ve Paredes, 2018).

Pasif-Agresif algoritmalar olarak adlandırılmasının sebebi;

Pasif: Tahmin doğruysa model korunarak ve herhangi bir değişiklik yapılmamaktadır.

Agresif: Tahmin yanlışsa modelde değişiklikler yapılmakta ve yanlışın düzeltebilmesi amaçlanmaktadır.

Pasif-Agresif algoritmalar, bir öğrenme oranı gerektirmesede, bir düzenleme parametresi içermektedirler. Genellikle Twitter gibi sürekli veri akışının olduğu ortamlarda anlık olarak çekilen verilerin yorumlanmasında kullanılabilecek en etkili algoritmalarından birisi olarak kabul edilmektedir.

Çevrimiçi öğrenmede, yeni veriler geldikçe öğrenmeye devam etmesi için arka planda bir makine öğrenimi modeli eğitilmektedir. Verilerin sürekli güncellendiği durumlarda ortaya çıkan problemleri çözmek için kullanılmaktadır.

4.1.9. Evrişimsel sinir ağı modeli (CNN)

Evrişimsel sinir ağı, bir YSA türüdür ve başlangıçta görüntü verilerini işlemek için kullanılmıştır. CNN, birden fazla gizli katmana sahip olup, çok katmanlı algılayıcı yapısında bir derin öğrenme kategorisidir. CNN'nin gizli katmanları oluşturan yapılar aşağıda listelenmiştir.

- Evrişim katmanı (Convolutional layer),
- Havuz katmanı (Pooling layer),
- Aktivasyon fonksiyonu (Activation function)
- Düzleştirme katmanı (Flatten layer)

Evrişim katmanı: Evrişim katmanı, özellik haritaları oluşturmak için evrişimsel filtreler gibi davranan bir dizi nöronudur. Girdi verileri bu katmanda küçük bloklara bölünür ve her blok belirli bir ağırlık kümesiyle birleşir. Aynı ağırlıklara sahip girdi verileri üzerinde evrişimsel filtre kaydırılarak farklı bir öznitelik kümesi elde edilir (Young ve ark., 2018). Evrişim işlemi Eşitlik 4.10'de gösterilmiştir.

$$F_l^k = (A * K_l^k) \quad (4.10)$$

A'nın girdi matrisi olduğu durumlarda K_l^k , k'nıncı katmanın l'ninci evrişim filtresini temsil etmektedir. F_l^k , evrişim katmanının çıkış özellik haritasını ifade etmektedir.

Havuz katmanı: Havuzlama, parametre sayısını azaltmak için, alanın uzaysal boyutunu küçültme işlemidir. Dolayısıyla özellik haritalarının boyutlarını azaltma işlemi de denilmektedir. Bu katman belirlenen alandaki benzer bilgileri tespit edip en baskın sonucu göstermektedir. Bu baskın sonuç, bir evrişim katmanı tarafından oluşturulan özellik haritasının belli bir bölgesinde bulunan özelliklerden belirlenmektedir. Bu nedenle, evrişim katmanı tarafından oluşturulan kesin olarak konumlandırılmış özellikler yerine baskın olarak alınmış özellikler üzerinde daha fazla işlem gerçekleştirilmektedir. Evrişimsel katmanın çıktısı olan özellikler, görüntüde farklı konumlarda ortaya çıkabilmektedir, böylece tam konumları daha az önemli hale gelmektedir. Havuzlama işlemi Eşitlik 4.11'de gösterilmiştir.

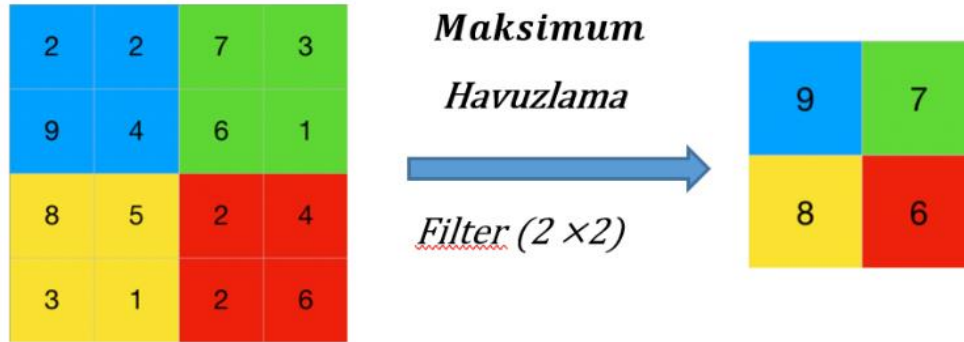
$$P_l = f_p(F^l) \quad (4.11)$$

P_l Ve F^l , l'ninci girdi ve çıktı özellik haritalarını temsil etmektedir. $f_p(.)$ Havuzlama işlemidir.

Havuzlama işlemi maksimum havuzlama (max pooling) ve ortalama havuzlama (average pooling) olarak 2 şekilde incelenebilir.

Maksimum havuzlama (Max pooling)

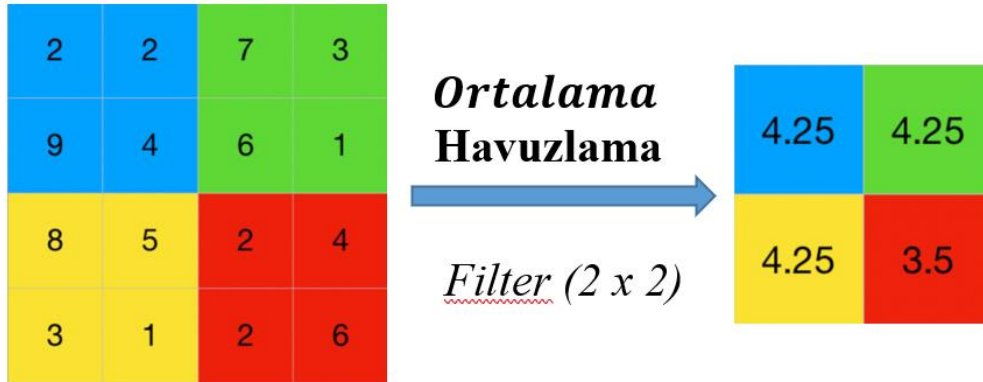
Maksimum havuzlama, filtre tarafından kapsanan özellik haritasının bölgesinden maksimum değeri taşıyan ögeyi seçen bir havuzlama işlemidir. Böylece, maksimum havuzlama katmanından sonraki çıktı, önceki özellik haritasının en belirgin özelliklerini içeren bir özellik haritası olacaktır. Maksimum havuzlama için örnek Şekil 4.10'da gösterilmiştir.



Şekil 4.10. Maksimum havuzlama örneği.

Ortalama havuzlama (Average pooling)

Ortalama havuzlama, filtre tarafından kapsanan özellik haritası bölgesinde bulunan öğelerin ortalamasını hesaplamaktadır. Bu nedenle, maksimum havuzlama, özellik haritasının belirli bir yamasında en belirgin özelliği verirken, ortalama havuzlama, bir yamada bulunan özelliklerin ortalamasını vermektedir. Ortalama havuzlama için örnek Şekil 4.11’de gösterilmiştir.



Şekil 4.11. Ortalama havuzlama örneği.

Aktivasyon fonksiyonu: Aktivasyon fonksiyonu kullanılmayan bir sinir ağı sınırlı öğrenme gücüne sahip bir doğrusal regresyon (linear regression) gibi davranmaktadır. Halbuki sinir ağının doğrusal olmayan durumları da öğrenmesi gerekmektedir. Çünkü sinir ağına öğrenmesi için görüntü, video, yazı ve ses gibi karmaşık gerçek bilgiler de verilmektedir. Çok katmanlı derin sinir ağları bu sayede verilerden anlamlı özellikleri öğrenebilmektedir.

Yapay sinir ağıları, evrensel fonksiyon yakınsayıcıları olarak tasarlanmış ve bu amaçla çalışması istenmektedir. Böylelikle bir fonksiyonu hesaplayarak öğrenebilmektedirler. Doğrusal olmayan aktivasyon fonksiyonlarının sahip olduğu bu özellik sayesinde ağların daha güçlü ve etkili öğrenmesi sağlanabilmektedir.

Aktivasyon fonksiyonu için yapılan işlem Eşitlik 4.12’de gösterilmiştir. $f_A(.)$ Evrişimli katmanın çıkışında kullanılan, F_l^k , doğrusal olmayan bir aktivasyon fonksiyonudur ve T_l^k ise l’ninci girdi için k’nıncı katmanının çıktısıdır.

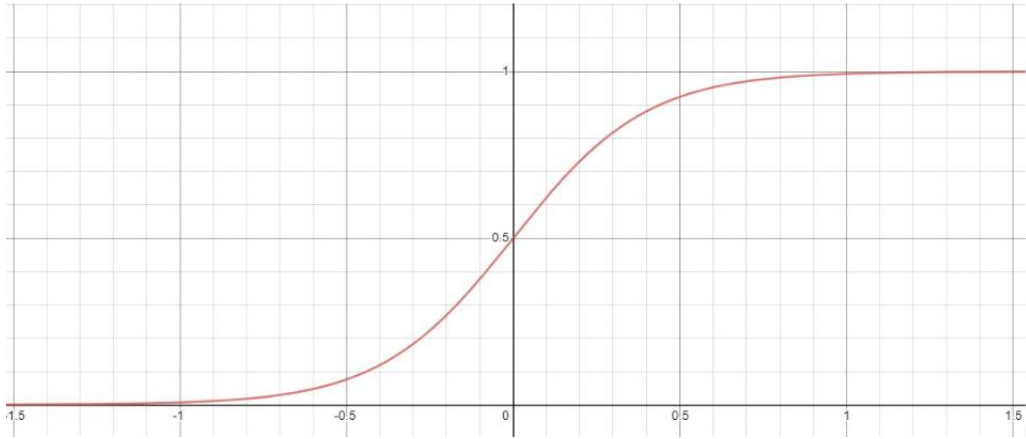
$$T_l^k = f_A(F_l^k) \quad (4.12)$$

Aktivasyon için lineer bir fonksiyon tanımlanmışsa, katmanda kaç nöron olursa olsun lineer bir hesaplama yapılmaktadır. Bu nedenle, doğrusal olmayan bir aktivasyon fonksiyonuna sahip olmak, sinir ağının önemli bir parçasıdır. Sigmoid, çıkışın 0 ile 1 arasında olması istendiğinde kullanılan doğrusal olmayan bir aktivasyon fonksiyonudur ve Eşitlik 4.13’de gösterilmiştir. Sigmoid fonksiyonunun grafiği Şekil 4.12’de verilmiştir.

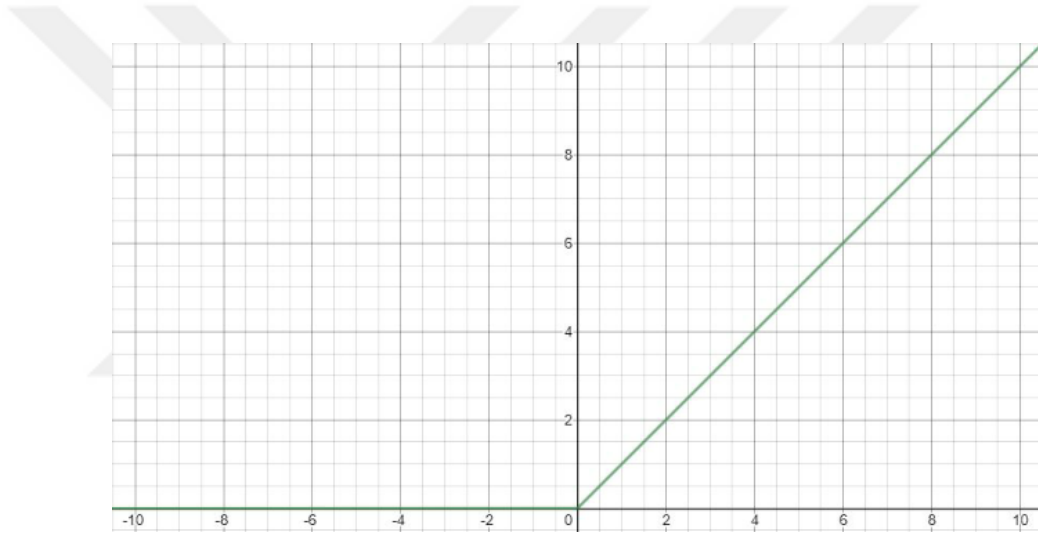
$$\sigma(z) = \frac{1}{1+e^{-z}} \quad (4.13)$$

ReLU, derin sinir ağlarında yaygın olarak kullanılan bir aktivasyon fonksiyonudur ve $[0, \infty)$ aralığını almaktadır. z sıfırdan küçük olduğunda $f(z)$ sıfır olmakta ve z sıfırdan büyük veya sıfıra eşit olduğunda $f(z)$ z ’ye eşit olmaktadır (Eşitlik 4.14). ReLU fonksiyonunun grafiği Şekil 4.13’de verilmiştir.

$$f(z) = \max(0, z) \quad (4.14)$$



Şekil 4.12. Sigmoid fonksiyonun grafiği.



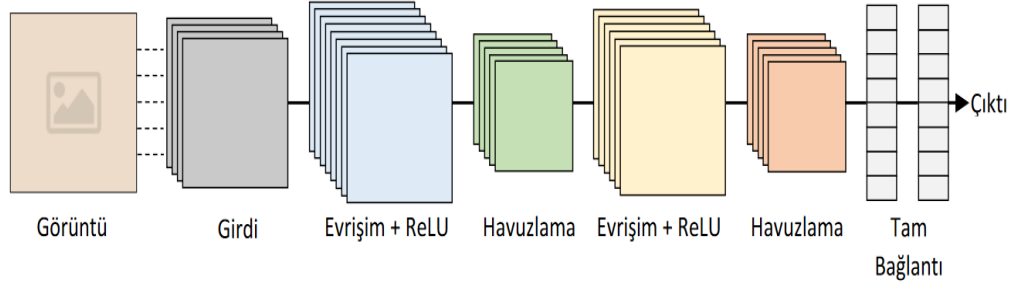
Şekil 4.13. RELU fonksiyonun grafiği.

Sinir ağlarında kullanılan ve doğrusal olmayan iki aktivasyon fonksiyonun şekilleri yukarıda verilmiştir.

Düzleştirme katmanı: Bu katman tam bağlı katmandan önce (Fully Connected Layer) çok boyutlu özellik dizisini bir vektöre dönüştürmektedir.

Tam Bağlı Katman: Tam bağlı katman, özelliklerin doğrusal olmayan bir kombinasyonunu yaparak önceki tüm katmanların çıktısını genel olarak analiz etmektedir. Ağın sonunda uygulanarak sınıflandırma/regresyon görevi için kullanılmaktadır.

Genel olarak bir CNN mimarisi ve katmanları Şekil 4.14'te verilmiştir (Bıçek, 2020).

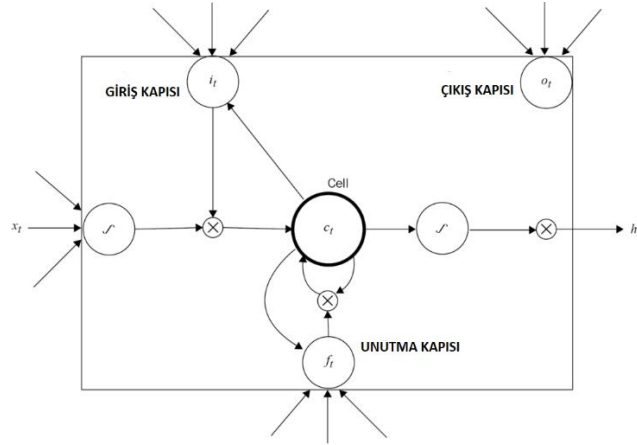


Şekil 4.14. CNN mimarisi ve katmanları.

4.1.10. Uzun kısa süreli bellek (LSTM)

LSTM'yi açıklamadan önce Tekrarlayan Sinir Ağı (RNN) hakkında bilgi verilmesinde fayda görülmektedir. Tekrarlayan Sinir Ağı (RNN), dahili belleğe sahip bir tür sinir ağıdır. Bir önceki adımın çıktısı, giriş olarak mevcut adıma gönderilmektedir. RNN tekrarlayıcıdır çünkü tüm girişlere aynı işlevi görmektedir. Üretilen çıktı, tekrarlayan ağı geri gönderilmektedir. RNN'lerdeki tüm girdiler birbiriyle ilişkilidir ve bazı bilgileri hatırlayan gizli katmanlara sahiptirler. Eğitim aşamasında, tutulan bilgiler ve reddedilen bilgiler tekrarlayan ağırlıklarla öğrenilmelidir. Bu özellik, Uzun kısa süreli bellek (LSTM) adı verilen özel bir RNN türü için önemlidir. RNN'lerin kaybolan gradyanlar problemini çözmek için LSTM kullanılmaktadır. Kaybolan gradyan problemi, derin yapay sinir ağlarının eğitiminin zor, veri ihtiyacının yüksek olmasının ardındaki problemdir. Temel sebebi, aktivasyon fonksiyonlarının ürettiği sonuçların, $[-1,1]$ aralığında oldukları sürece kullanılan algoritmalarda, çarpıla çarpıla sıfıra yakınsaması yani uzun süreli bağımlılıkların yakalanamamasıdır.

Kaybolan gradyan problemini çözmek için RNN'lerin güncellenmiş bir versiyonunu, yani Uzun kısa süreli bellek (LSTM) kullanılmaktadır. LSTM'ler, verilerdeki uzun vadeli bağımlılıkları öğrenme yeteneğine sahiptir. LSTM yapısında 3 çeşit kapı bulunmaktadır. Bu kapılar Şekil 4.15'te gösterilmiştir.



Şekil 4.15. LSTM kapılarının görünümü.

Giriş Kapısı: Giriş değerlerini bellek durumunu güncellemek için ayarlamaktadır (Yazma işlemi gerçekleşir). Sigmoid işlevi, 0,1'den hangi değerlerin geçeceğini seçmektedir (Eşitlik 4.15). Tanh fonksiyonu -1 ile 1 arasında değişen değerlerin ağırlığını belirlemektedir (Eşitlik 4.16).

$$i_t = \sigma(W_i \cdot [h_{t-1}, X_t] + b_i) \quad (4.15)$$

$$C_t = \tanh(W_c \cdot [h_{t-1}, X_t] + b_c) \quad (4.16)$$

Unutma Kapısı: Bloktan atılması gereken bilgiyi belirlemektedir (Reset işlemi gerçekleşir). Eşitlik 4.17'te Sigmoid'ın işlevi, bloktan atılacak ayrıntılara karar vermektir. Önceki durumu (h_{t-1}) ve içerik girdisini (X_t) dikkate almaktadır. C_{t-1} Hücre durumundaki sayılar için 0 ile 1 arasında bir sayı (0 ise atla, 1 ise tut anlamına gelir) üretmektedir.

$$f_t = \sigma(W_f \cdot [h_{t-1}, X_t] + b_f) \quad (4.17)$$

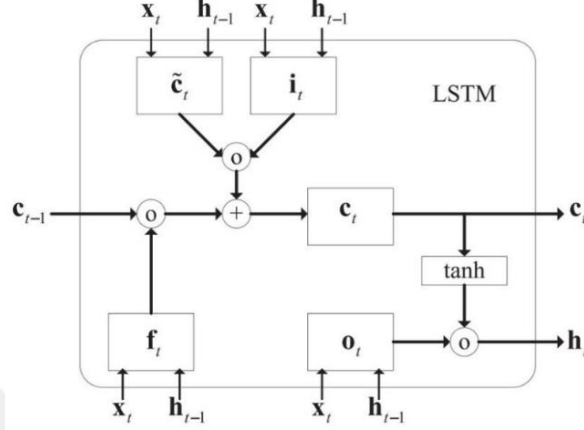
Çıkış Kapısı: Bir sonraki hücrenin girişini (h_{t+1}) belirler. Ayrıca tahmin yapmak için de kullanılır (Okuma işlemi gerçekleşir). Sigmoid'ın işlevi, 0,1'e izin verecek değerleri belirlemektir (Eşitlik 4.18). Tanh'nın işlevi, -1 ile 1 arasında değişen önem derecelerine karar vererek değerlerin ağırlığını belirlemektir (Eşitlik 4.19).

$$o_t = \sigma(W_o \cdot [h_{t-1}, X_t] + b_o) \quad (4.18)$$

$$h_t = o_t * \tanh(C_t) \quad (4.19)$$

Her kapı, okuma/yazmayı kontrol eden ve böylece uzun süreli hafıza fonksiyonunu modele dahil eden bir anahtar gibi davranmaktadır (Messina ve Louradour, 2015).

Örnek bir LSTM hücre yapısı Şekil 4.16’da verilmiştir (Qing ve Niu, 2018).



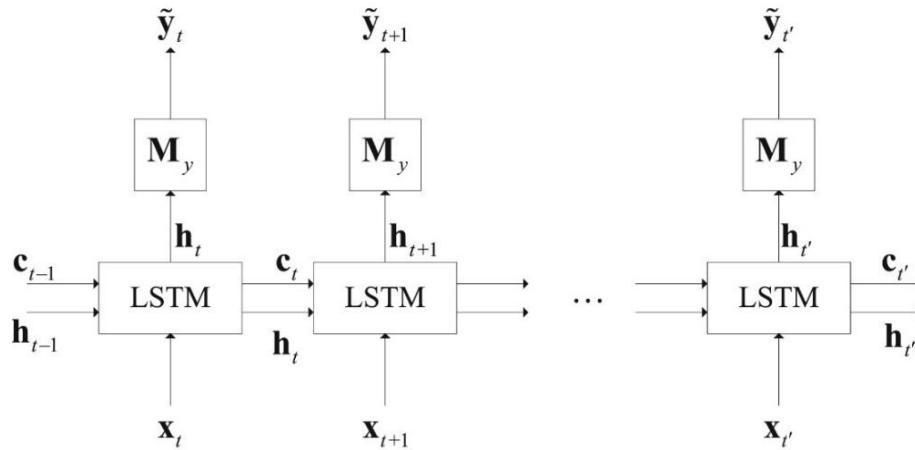
Şekil 4.16. LSTM hücre yapısı.

Şekil 3.15’te t değeri zamanı, i_t giriş kapısını, f_t unutma kapısını, o_t çıkış kapısını, C_t hücre durumunu ifade etmektedir, bilginin kesintisiz olarak bir hücreden diğer hücreye aktığı kanaldır (Messina ve Louradour, 2015). C_t Değeri Eşitlik 4.20’deki denklem ile hesaplanmaktadır.

$$C_t = f_t \circ C_{t-1} + i_t \circ \hat{C}_t \quad (4.20)$$

‘o’ Simgesi Hadamard ürününü ifade etmektedir. Hadamard Ürünü, aynı şekilli matrisler üzerinde çalışmakta ve aynı boyutlarda üçüncü bir matris üretmektedir.

LSTM ağının yapısı Şekil 4.17’de verilmiştir (Qing ve Niu, 2018).

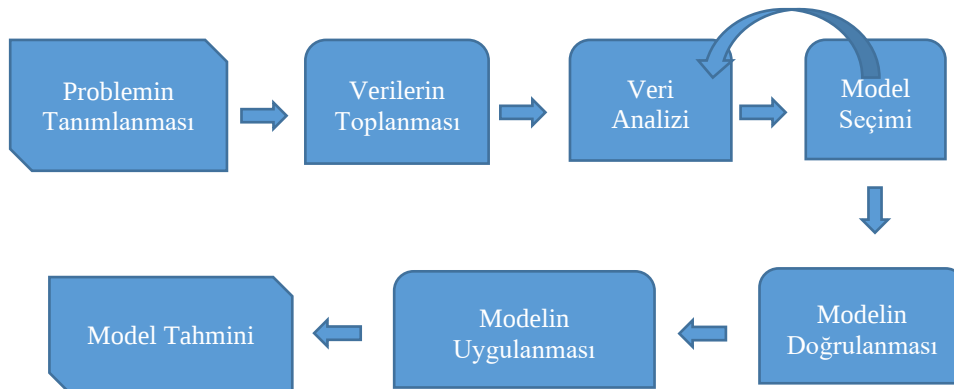


Şekil 4.17. LSTM ağının yapısı.

Modeller için tahmin süreci

Modellerin tahmin süreci kavramını; girdileri olan, bunlara değer kazandıran ve sonucunda çıktı elde edilen faaliyetler serisi veya aktiviteler olarak tanımlamak mümkündür. Montgomery ve ark. (2015) tarafından tahmin sürecindeki faaliyetler aşağıdaki sırasıyla sunulmuş ve akış diyagramı Şekil 4.18’de gösterilmiştir.

- Problemin tanımlanması aşaması,
- Verilerin toplanması aşaması,
- Verilerin analizi aşaması,
- Uygun model seçimi aşaması,
- Modelin doğrulanması aşaması,
- Modelin uygulanması aşaması,
- Modelin tahmini aşaması.



Şekil 4.18. Model tahmin süreci akış diyagramı.

Problemin Tanımlanması:

Amaca doğru bir şekilde ulaşabilmek için en başında problemin tanımlanması gerekir.

Verilerin Toplanması:

Tanımlanan problemin çözümüne yönelik olarak amaca ulaşabilmek için uygun bir veri setinin oluşturulması gerekir.

Verilerin Analizi:

Toplanan verilerin temizlenmesi, ön-işlem aşamalarından geçirilerek uygulanacak model için hazır hale gelmesi aşamasıdır.

Model Seçimi:

Oluşturulan veri setine uygun olarak polaritesi pozitif ve negatif olan veriler ile uyumlu sonuçlar verebilecek modelin seçilmesi aşamasıdır.

Modelin Doğrulanması:

Veri setinden seçilen belli sayıdaki veri için eğitim ve test verileri olarak belirlenmesi, seçilen modelin doğrulanmasında önemli bir aşamadır.

Modelin Uygulanması:

Türkçe ve İngilizce tweetlerin seçilen modellere uygulanmasıdır. Test, eğitim ve tüm veri seti olarak değerlendirilmek üzere modelin başarımlı performansının ölçüldüğü aşamadır.

Modeli Tahmini:

Ulaşılan model başarımlı performansı esas alınarak, seçilen makine öğrenmesi algoritmaları tarafından ulaşılan doğruluk ölçütüdür.

Modeller için değerlendirme ölçütleri

Bu tez çalışmasında, Makine Öğrenmesi algoritmalarını kullanarak başarımlı performanslarının değerlendirilmesinde aşağıda sunulan değerlendirme ölçütleri esas alınmıştır:

Doğruluk (Accuracy): “Sınıflandırıcı bütün örnekleri sınıflandırmada ne kadar başarılıdır” sorusunun cevabıdır. Eşitlik 4.21 tarafından hesaplanır:

$$\text{Doğruluk} = \frac{TP+TN}{TP+TN+FP+FN} \quad (4.21)$$

Burada, TP (True Positive), TN (True Negative), FP (False Positive) ve FN (False Negative) sırasıyla, doğru sınıflandırılmış pozitif tweetleri, yanlış sınıflandırılmış pozitif tweetleri, negatif mobil oyun tweetleri olarak sınıflandırılmış pozitif mobil oyun tweetleri, pozitif mobil oyun tweetleri olarak sınıflandırılmış negatif mobil oyun tweetleri ifade etmektedir.

Kesinlik (Precision): “Ulaşılan bilginin ne kadarı istenilen bilgiyle ilgilidir” sorusunun cevap değeridir. Diğer bir deyişle yanlış pozitifleri eleme kabiliyetidir. Eşitlik 4.22 tarafından hesaplanır:

$$\text{Kesinlik} = \frac{TP}{TP+FP} \quad (4.22)$$

Duyarlılık (Recall): “Ulaşılması gereken bilginin ne kadarına ulaşılmıştır” sorusunun cevap değeridir. Diğer bir deyişle doğru pozitifleri tespit etme kabiliyetidir. Eşitlik 4.23 tarafından hesaplanır:

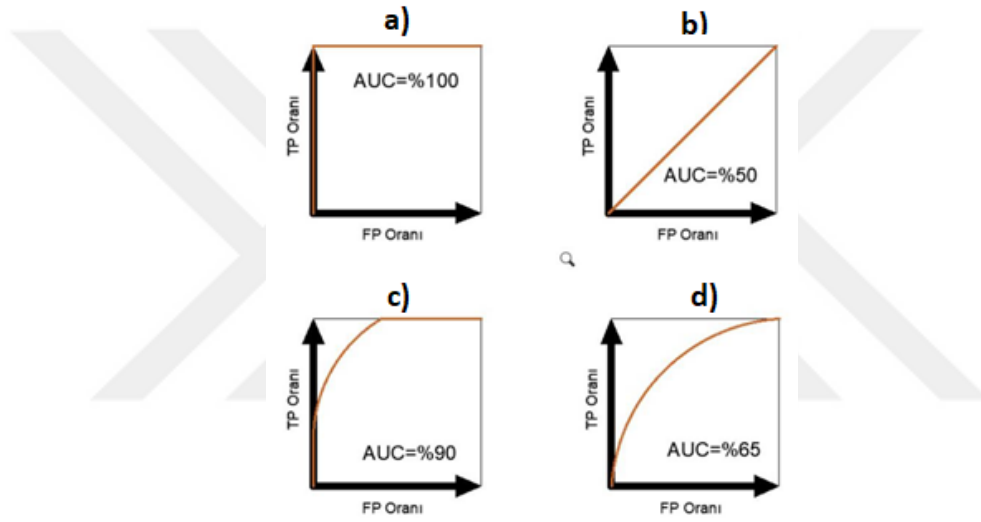
$$\text{Duyarlılık} = \frac{TP}{TP+FN} \quad (4.23)$$

F Skor (F-Measure): Sınıflandırıcının hem kesinliğini hem de duyarlılığını göz önünde bulunduran değerlendirme metriğidir. Kesinlik ve hassasiyet değerlerinin harmonik ortalamasıdır. Eşitlik 4.24 tarafından hesaplanır:

$$F = \frac{\text{Kesinlik} \times \text{Duyarlılık}}{\text{Kesinlik} + \text{Duyarlılık}} \quad (4.24)$$

ROC Eğrisi: Bütün ayırma (pozitif- negatif, doğru-yanlış) işlemlerinde olduğu gibi, kullanılan yöntemler duyarlılık ve kesinlik arasındaki dengeyi inşa etmekle ilgilenmektedir. Veri kümelerindeki olumlu veya olumsuz modeller, denk olacak şekilde bir dağılım göstermediği için, direkt doğruluk ve hassasiyet kriterlerinden önce, ROC eğrisi üzerinden doğruluk ve hassasiyet arasındaki dengeyi yorumlamak gereklidir. Şekil 4.19’da gösterildiği üzere, ROC eğrisi altındaki alan AUC (Area Under Curve) olarak

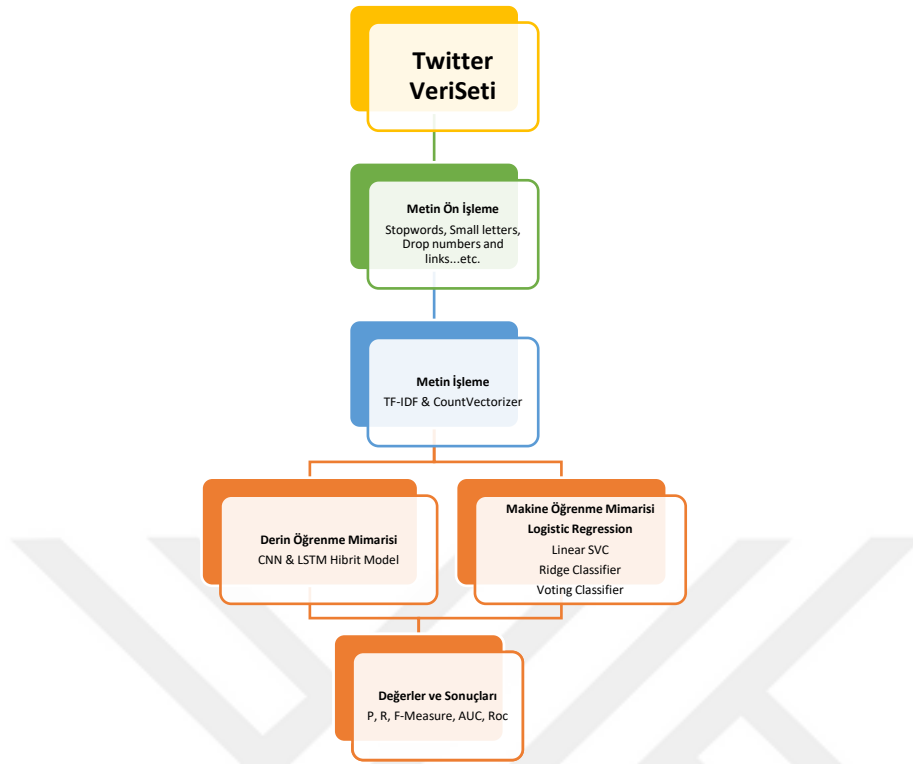
tanımlanır. AUC değeri belirlenirken, *Kesinlik* ve *Duyarlılık* değerleri hesaplanıp *F-Skor* değerini maksimum yapan sınır (threshold) noktası seçilir. ROC eğrisi değişen gruplandırma taban verilerine göre, doğru pozitiflerin sayısının, yanlış pozitiflerin bir işlevi olarak çizilmesiyle oluşur. ROC sayısı ‘1’ olduğunda, pozitifler çok iyi bir şekilde negatiflerden ayrılmış demektir. Şayet ROC sayısı ‘0’ olduğunda ise rastgele bir pozitif bulunamadı manasına gelir. ROC basitçe doğru pozitiflerin yanlış pozitiflere olan kesri olarak söylenebilir. İçeriğindeki alan büyüklüğü arttıkça, makine öğrenmesi modellerinin sınıfları ayırt etme yeteneği artar (Kılıç, 2013).



Şekil 4.19. TP'nin FP'ye oranları esas alınarak (a) %100, (b) %50, (c) %90, (d) %65 AUC için ROC eğrisinin değişim grafikleri.

4.2. Yöntem

Bu tez çalışmasında, mobil oyun ile ilgili Twitter veri setine uygulanan tüm işlem adımlarının akış diyagramı Şekil 4.20'de sunulmuştur.



Şekil 4.20. Twitter veri setine uygulanan işlemlerin akış diyagramı.

Şekil 4.20’ye göre, bu tez çalışmasında, aşağıdaki adımlar esas alınarak mobil oyun veri analizleri gerçekleştirilmiştir.

Adım 1: Twitter üzerinden ulaşılan mobil oyun veri seti çeşitli metin ön işleme aşamalarına tabii tutularak veri birleştirme ve veri temizleme işlemleri gerçekleştirilmiştir.

Adım 2: Metin işleme aşamasında, TF-IDF ve CountVectorizer işlemleri esas alınarak tweet verileri vektörel olarak ifade edilmiştir.

Adım 3: Makine öğrenmesi algoritmaları kullanılarak oluşturulan ve geliştirilen modeller üzerinden veri analizleri gerçekleştirilerek ulaşılan sonuçlar ile başarımlar performansları karşılaştırılmış ve ROC eğrileri çizdirilmiştir.

4.2.1. Veri setinin oluşturulması

Bu tezde, sosyal medya platformu olan Twitter üzerinden çekilen İngilizce ve Türkçe tweetler kullanılmıştır. Toplamda 2.217.708 tweete ulaşılmış olup, bunların 376.431 tanesini Türkçe, 1.839.274 tanesini ise İngilizce tweetler oluşturmuştur. Güney

Kore, İngiltere, Amerika Birleşik Devletleri, Kanada, Fransa, Güney Afrika ve İtalya olmak üzere 7 farklı ülkeden İngilizce tweetlere ulaşılmıştır. Türkçe tweetlere ise, sadece Türkiye’den ulaşılmıştır. Her bir ülkeden tweet verilerine ulaşırken, en kalabalık şehirleri lokasyon olarak işaretlenmiş ve 150 km’lik bir çap tanımlanmıştır. Her ülke için en az 5 şehirden, en fazla ise 15 şehirden mobil oyun tweet verileri alınmıştır. Türkiye’den 2015-2021 yılları arasında her bir yıl için 53.776 tweet elde edilmiş fakat diğer ülkeler için bu durum değişiklikler göstermiştir. Yabancı ülkelere her yıl için sabit bir rakamda tweet elde edilememesinin sebebi, elde edilen tweetlerin kendi yazdığım phyton kodu aracılığıyla temin edilmesinden kaynaklanmaktadır. Tweetlerin ülkelere ve yıllara göre nasıl çekildiği Şekil 4.21’te sunulmuştur.

```
import pandas as pd
import snsrape.modules.twitter as sntwitter
import itertools
import numpy as np

keyword_list = ["mobil oyun", "mobile game"]
tweets_per_city = 20000
km = 150
since_list = ["2015-01-01", "2016-01-01", "2017-01-01", "2018-01-01", "2019-01-01", "2020-01-01", "2021-01-01"]
until_list = ["2015-12-30", "2016-12-30", "2017-12-30", "2018-12-30", "2019-12-30", "2020-12-30", "2021-12-30"]
year_list = [2015, 2016, 2017, 2018, 2019, 2020, 2021]

# Nüfusu yüksek 5 şehir seçtim
city = {"TR": ["İstanbul", "Ankara", "İzmir", "Antalya", "Eskişehir"],
        "ENG": ["London", "Manchester", "Liverpool", "Oxford", "Cambridge"],
        "SKO": ["Seul", "Busan", "Suwon", "Daegu", "Gwangju"],
        "FRA": ["Paris", "Lyon", "Marseille", "Bordeaux", "Nice"],
        "ITA": ["Rome", "Turin", "Milan", "Palermo", "Naples"],
        "USA": ["New York", "Los Angeles", "Chicago", "Houston", "Phoenix"],
        "AFR": ["Cape Town", "Johannesburg", "Durban", "Pretoria", "Rustenburg"],
        "CAN": ["Ottawa", "Montreal", "Toronto", "Vancouver", "Quebec"]}

# verilerin çekilmesini sağlayan yer burası. itertools ile sırayla hepsini gezicek. Tarih ve içerik aldım.
df = pd.DataFrame(itertools.islice(sntwitter.TwitterSearchScraper(
    f'{keyword_list[0]} -filter:links -filter:replies').get_items(), 1))['date', 'content']
```

Şekil 4.21. Tweetlerin elde edilmesi için kullanılan phyton kodu.

Verilerin özelleştirilmesi bakımından Türkçe çekilen tweetler için “mobil oyun”, İngilizce çekilen tweetler için “mobile game” kelimelerini içerisinde barındıran tweetler oluşturulan veri setine dahil edilmiştir. Hem İngilizce hem de Türkçe olarak ulaşılan tweetler için kelime bulutu (Şekil 4.22 ve Şekil 4.23) çıkarılmıştır. Tweet veri seti oluşturulurken, ulaşılan tweetlerin içeriği, yılı ve ülkesi esas alınmıştır ve tüm tweetler CSV formatında kaydedilmiştir.

okunması bazılarının ise istenilen yıl formatında olmaması (2016.0 gibi) nedeniyle bu verilerde düzenleme yapılarak istenilen yıl formatına (2016 gibi) dönüştürülmüştür. Veri birleştirme sonrasında, ulaşılan İngilizce tweetlerin ülkeler bazında dağılımı Çizelge 4.1’de, 2015-2021 yılları arasında ulaşılan tweet sayıları ise Çizelge 4.2’de gösterilmiştir. Çizelge 4.1’in grafiksel gösterimi Şekil 4.24’te, Çizelge 4.2’nin grafiksel gösterimi ise Şekil 4.25’de sunulmuştur. 2015-2021 yılları için ulaşılan toplam Türkçe tweet sayıları Çizelge 4.3’te sunulmuştur.

Çizelge 4.1. 2015-2021 yılları için ülkelere göre ulaşılan toplam İngilizce tweet sayıları

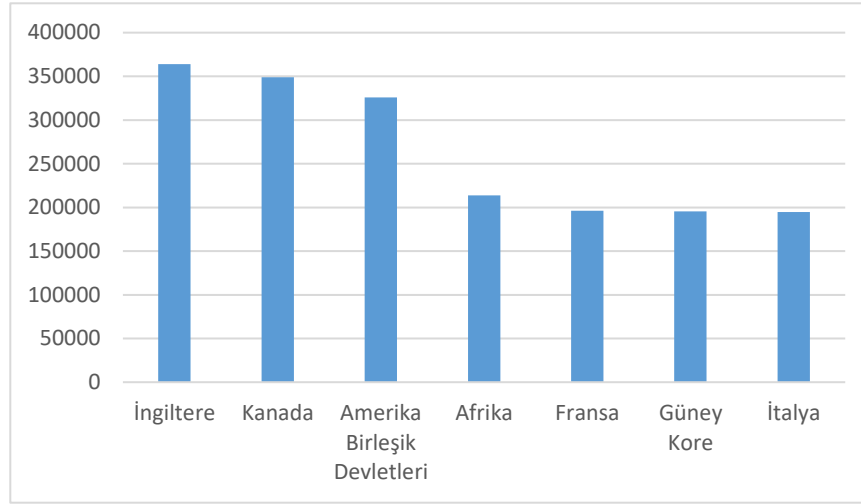
Ülke Adı	Toplam İngilizce Tweet Sayısı
İngiltere	363982
Kanada	349003
Amerika Birleşik Devletleri	325718
Güney Afrika	213941
Fransa	196310
Güney Kore	195562
İtalya	194758

Çizelge 4.2. 2015-2021 yılları için ulaşılan toplam İngilizce tweet sayıları

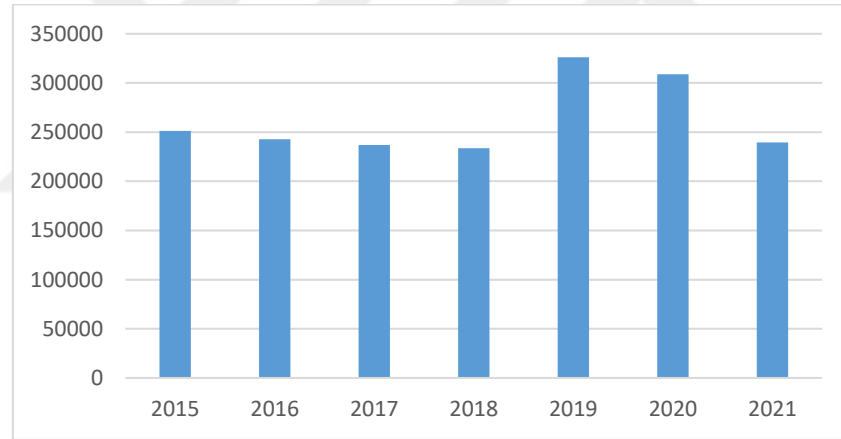
Yıl	Toplam İngilizce Tweet Sayısı
2015	251108
2016	242758
2017	237009
2018	233754
2019	326274
2020	308754
2021	239617

Çizelge 4.3. 2015-2021 yılları için ulaşılan toplam Türkçe tweet sayıları

Yıl	Toplam Türkçe Tweet Sayısı
2015	53776
2016	53776
2017	53776
2018	53776
2019	53776
2020	53776
2021	53775



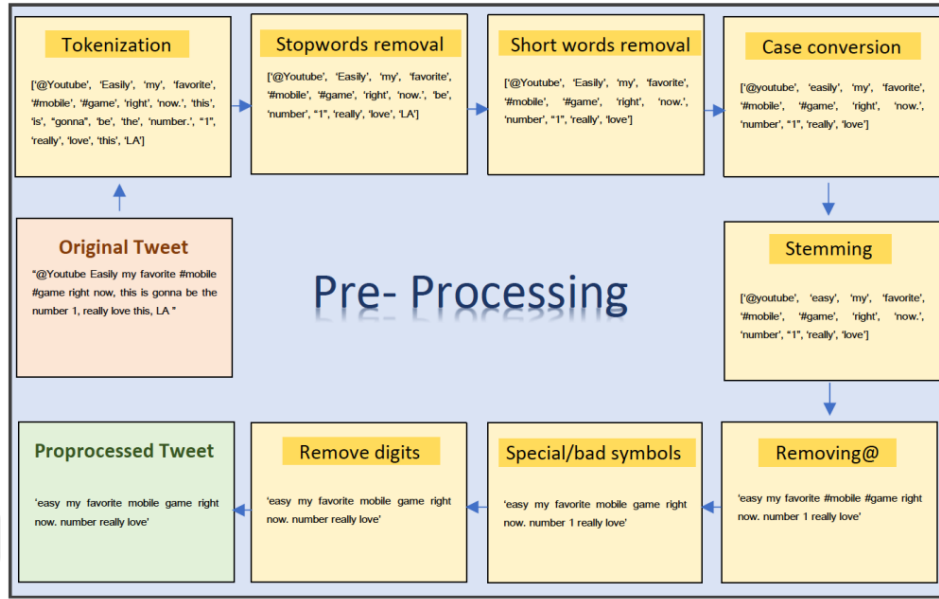
Şekil 4.24. 2015-2021 yılları için ülkelere göre ulaşılan İngilizce tweet sayısının değişim grafiği.



Şekil 4.25. 2015-2021 yılları için ulaşılan İngilizce tweet sayılarının değişim grafiği.

4.2.3. Veri temizleme

Mobil oyun tweet verilerinin birleştirilmesi ile birlikte İngilizce ve Türkçe olan tweetler ayrı CSV formatında kaydedilerek ham tweet verilerinin metin ön-işlem aşamasında işlenmesi sağlanmıştır. Böylelikle polarize edilmeye hazır, kök haline gelmiş, gürültüden temizlenmiş bir veri seti oluşturulabilir. Tweet veri setinin temizlenmesi ve böylece tweetteki kelimelerin kök haline getirilmesinde uygulanan ön-işlem aşamaları Şekil 4.26'da gösterilmiştir.



Şekil 4.216. Tweet veri setinin temizlenmesindeki ön-işleme sürecinin grafiksel olarak gösterimi.

Şekil 4.26'daki mobil oyun tweet verilerinin temizlenmesi sürecinde uygulanan her bir ön-işlem adımı ve açıklaması aşağıda sunulmuştur:

Tokenization:

Bu adımda, tweet metninin sözcüklere, simgelere ve öğelere bölünmesi gerçekleştirilir. Doğru yapılan tokenization işlemi yapılacak analizin verimliliği açısından oldukça önemlidir.

Stop word removal:

Bazı sözcüklerin cümleleri daha okunaklı ve anlamlı hale getirmesine karşın, metnin analizi için bir değer taşımaması nedeniyle kaldırılması daha uygundur.

Short word removal:

Tweet veri setindeki herhangi bir kelime 3 harften az ise metinden çıkarılmalıdır. Korkusuz ve Carus (2020) tarafından yapılan çalışmada Support Vector Machine algoritması kullanılarak tweet'lerdeki kısa kelimelerin doğruluğu etkilendiğini gösteren sonuçlar, bu adımdaki işlemin önemini göstermektedir.

Case conversion:

Tweet verisindeki metnin tamamı küçük ya da büyük harfe dönüştürülmesi adıdır. Büyük harflerin bazı dillerdeki karşılığı farklı olabildiğinden dolayı genellikle bu işlem küçük harflere çevrilerek yapılmaktadır. Özellikle olasılık modelleri, "Game" ve "game" gibi büyük ve küçük olarak yazılmış kelimeleri farklı iki kelime olarak kabul eder ve dolayısıyla kelimeler küçük harfe dönüştürülmezse, sınıflandırıcının verimliliğini bozabilmektedir.

Stemming:

Tweet kelimelerin kök biçimlerine döndürülmesi işlemidir. Örneğin, played, playing ve plays, aynı anlama gelen "play" kelimesinin çeşitleridir. Kök haline getirilme işlemi amaçlanan, karmaşıklığın azaltılması ve sınıflandırıcıların öğrenme kapasitesini olumlu yönde etkilemesidir.

Removing @:

Twitter uygulamasında tweet gönderen her kullanıcı için @ atamaktadır. Tweet kelimesi kök halini aldıktan sonra, ulaşılan tweetteki @ ile başlayan kelimeler kaldırılmalıdır. Metin madenciliğinde herhangi bir katkısı bulunmadığı için bu işarete ihtiyaç duyulmamaktadır.

Removing special / bad symbols:

Tweet verisinden #, {, }, % gibi özel sembollerin kaldırılma işlemi bu adımda gerçekleştirilmiştir. Fakat “❤” gibi bazı emojielerin kaldırılması için ayrıca bir işlem uygulanmıştır.

Remove digits:

Metin analizinde sayıların pozitif ya da negatif yönde herhangi bir değere sahip olmaması dolayısıyla sayısal ifadelerin tweet veri setinden kaldırılması işlemidir. Modellerin eğitiminde karmaşıklığı azaltması nedeniyle önemli bir adımdır.

Tweet verilerinin temizlenmesi aşamasında uygulanan ön-işlem adımları, tweet verilerinin kök haline getirilerek metin-işlem aşamasına hazırlanmasını sağlamıştır. Türkçe ve İngilizce tweetlerin temizlenmesi sonrasında kök halini alan tweet verileri sırasıyla Çizelge 4.4 ve Çizelge 4.5’de gösterilmiştir.

Çizelge 4.4. Ön-işlem adımlarının uygulandığı Türkçe tweet örnekleri

Ham Tweet	Yıl	Ülke	Kök Tweet
mükemmel bir mobil oyun buldum onu oynuyorum 😊	2016	Türkiye	mükemmel mobil oyun bul onu oyna
Hoş bir mobil oyun bulmanın zamanı gelmiş!!	2018	Türkiye	hoş bir mobil oyun bul zaman gel
mobil veri oyun içindeki şeyi indirirsem ücret pahalımı?	2019	Türkiye	mobil veri oyun içi indir ücret pahalı
Oyun dergisi ağustos Sayısından güzel mobil oyun önerisi? @mobiloyun	2017	Türkiye	oyun dergi ağustos sayı güzel mobil oyun öneri
Benim mobil oyun vakit kaybı #iğrenç	2019	Türkiye	ben mobil oyun vakit kaybı iğrenç

Çizelge 4.5. Ön-işlem adımlarının uygulandığı İngilizce tweet örnekleri

Ham Tweet	Yıl	Ülke	Kök Tweet
never understood candy crush hate incredibly!!	2015	İngiltere	never understood candy crush hate incredible
Amazing mobile game and really funny thx	2015	İngiltere	amazing mobile game and really fun thanks
friends we can play mobile game	2018	Güney Kore	friend can play mobile game
getting involved this mobile game good luck	2019	Fransa	getting involved mobile game good luck
@wow hard and bored mobile game #pubg	2019	İtalya	wow hard boring mobile game

4.2.4. Verilerin polarize edilmesi

Tweet verilerin pozitif ya da negatif anlam içerip içermediğini tespiti amacıyla uygulanan işlemdir. Bu tez çalışmasında, nötr tweetler tespit edilerek veri setinden çıkarılmıştır. Tweet verilerinin polarize edilmesinde phyton “TextBlob” kütüphanesi kullanılmıştır. İngilizce tweetlerde herhangi bir problem ile karşılaşılmasına rağmen, Türkçe tweetlerin polaritesinin hesaplanmasında, öncelikle tweetlerin İngilizceye çevrilmesi gerekmektedir. Bunun nedeni “TextBlob” kütüphanesinin sadece İngilizce kelimeler için

polarite hesabı yapılabilmesidir. İngilizce için hazırlanmış çeşitli duygu sözlükleri Türkçe tercüme yöntemleriyle çevrilerek kullanılmış olsa da hem tercüme edilen kelimelerin eksikliği hem de başarılı sonuçlar alınmamış olmasından dolayı bu tez çalışmasında tercih edilmemiştir. Dolayısıyla bahsedilen problemin çözümünde, sözcüklerin anlamlarının birden fazla duygu sözlüğüne bakılarak aranması ve kontrollü çevirme yapılması sağlanarak hatalar daha aza indirgenebilir.

Model oluşturulurken polarize edilmiş 37.042 pozitif, 37.042 negatif tweet olmak üzere toplamda 74.084 adet tweet kullanılmıştır. İngilizce veriler için uygulanan polarite belirleme yöntemi için yazılan phyton kodu Şekil 4.27’de sunulmuştur.

```
en_df.loc[en_df.Polarity >= 0.35, 'Sentiment_Type'] = 1 # Positive
en_df.loc[en_df.Polarity <= -0.35, 'Sentiment_Type'] = -1 # Negative
en_df.loc[(en_df.Polarity > -0.35) & (en_df.Polarity < 0.35), 'Sentiment_Type'] = np.NaN # Neutral
current_df = en_df.loc[en_df['Sentiment_Type'].isnull()==0]

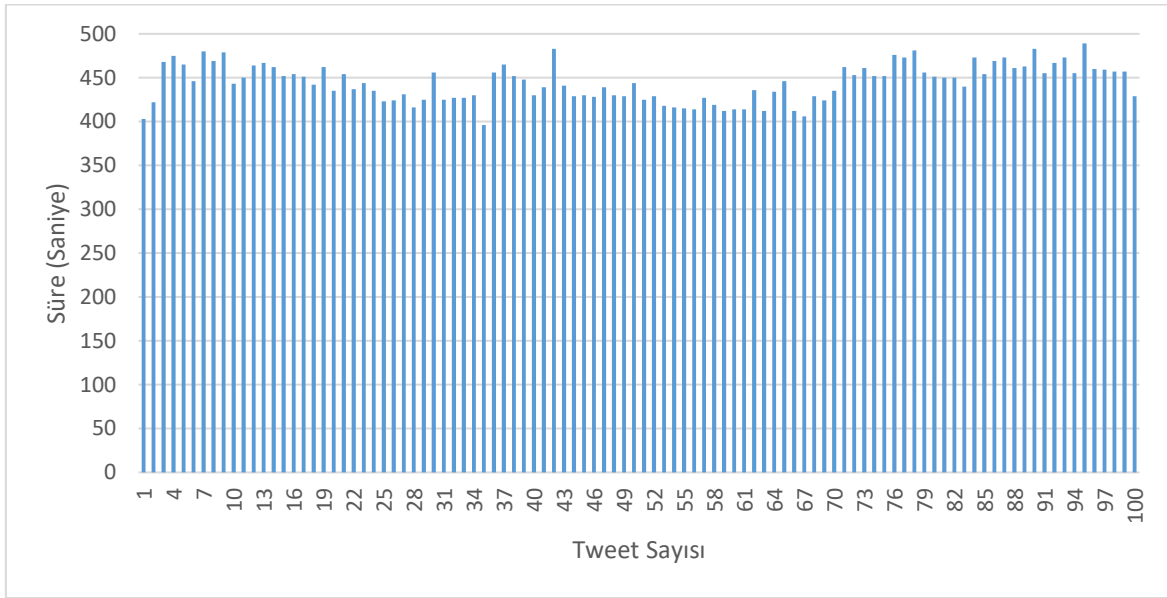
ilk11 = current_df[current_df["Sentiment_Type"]==1][:37042]
ikinci11 = current_df[current_df["Sentiment_Type"]==1]

son_df = pd.concat([ilk11, ikinci11])
son_df.shape

(74084, 6)
```

Şekil 4.27. İngilizce mobil oyun verileri için polarite hesaplanması.

Türkçe tweetlerin polaritesinin hesaplanmasında “TextBlob” kütüphanesini kullanarak İngilizceye çevirme işlemi oldukça vakit almaktadır. Bu tez çalışmasında, 500 tweetin çevirisi ortalama 445 saniyede tamamlanmıştır. 50.000 tweet için ise, ortalama 12.3 saat sürmesi bu işin zorluğunu ortaya koymaktadır. Toplamda 376.431 Türkçe tweetin İngilizceye çevrilmesi durumunda, sistemde kitlenme ve donma gibi problemler ile karşılaşmış ve dolayısıyla sonuç alınamamıştır. Çözüm olarak 7 ayrı dataframe oluşturulmuştur ve 6 tanesine 50.000 tweet ile birlikte son dataframe içerisine ise 76.433 tweet yerleştirilerek bu problemin çözümü sağlanmıştır. 50.000 Türkçe tweetin polarize edilebilmesi sürecinde, İngilizceye çevrilecek Türkçe tweetler için gerekli süre Şekil 4.28’de gösterilmiştir.



Şekil 4.28. İngilizceye çevrilen 50.000 Türkçe tweet için geçen sürenin değişim grafiği
(Her bir satır çubuğu 500 tweete karşılık gelmektedir).

İngilizce tweetlerde polarite hesabı için bir problem oluşmadığından dolayı “TextBlob” kütüphanesinde “sentiment.polarity” fonksiyonu kullanılarak hesaplama yapılmıştır. Polarite işleminin tamamlanması sonucunda, tweetlerin pozitif ve negatif olarak ayrımının yapılmış olması ve nötr tweetlerin veri setine dahil edilmemiş olması gerekmektedir.

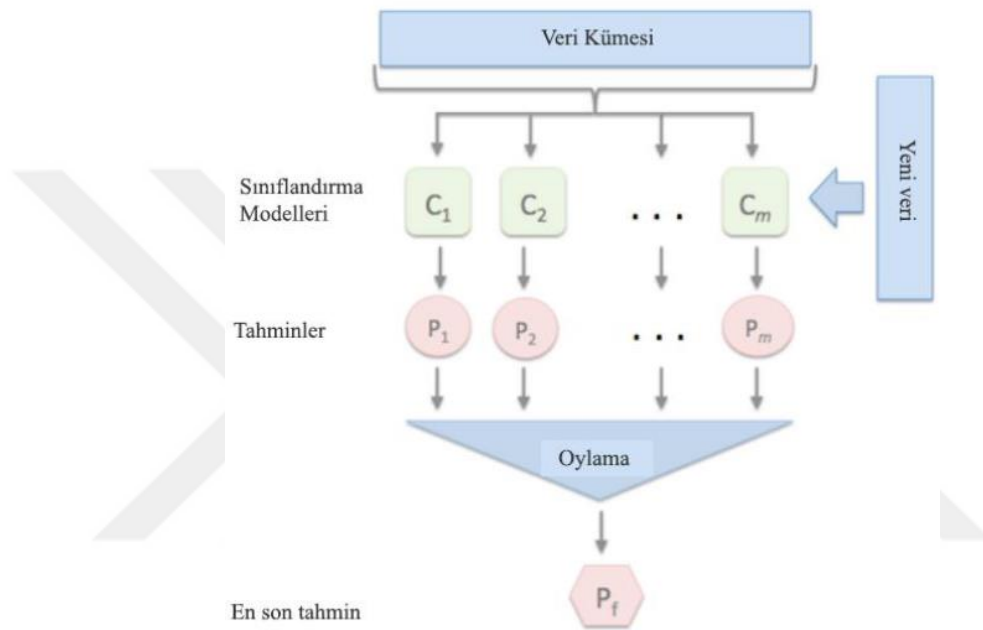
4.2.5. Model oluşturma

Tweetlerin polarize edilmesi sonrasında tweetler CSV formatında kaydedilerek model oluşturulma aşamasına geçilmiştir. Model oluştururken Türkçe tweet veri setinde 11271 pozitif ve 11271 negatif tweet kullanılmıştır. İngilizce tweet veri setinde ise, 37042 pozitif ve 37042 negatif tweet kullanılmıştır. Dolayısıyla, bütün modellerin doğru oluşturulmasında pozitif ve negatif tweet sayılarına dikkat edilmiştir.

Aşağıda sırasıyla, bu tez çalışmasında kullanılan VC’yi oluşturan makine algoritmaları sunularak VC’nin algoritma yapısı ile CNNLSTM tabanlı olarak geliştirilen TEMSAP (Twitter Mobil Game Sentiment Analysis Approach)’ın algoritma yapısı sunulmuştur.

4.2.6. Oylama sınıflayıcısı (Voting classifier)

VC, oylama sınıflayıcısı olarak, oylamaya katılan ve makine öğrenmesi algoritmaları olan kolektif sınıflayıcıların her biri tarafından hesaplanan tahminlerin ortalaması alınarak daha iyi bir tahmine ulaşılmasını sağlayan bir makine öğrenmesi tekniğidir (Şekil 4.29). Literatürdeki birçok çalışma, VC'nin kullanılmasının doğruluk oranını artıracaklarını göstermektedir (Da Silva ve ark., 2014).



Şekil 4.29. VC'nin akış diyagramı.

Çeşitli sınıflandırıcı modeller kendi içerisinde bir tahminde bulunarak (0 ya da 1) oyunu kullanır ve ortalamanın üstünde olan tahmin geçerli sayılarak nihai tahmin açıklanır. Bireysel çalışan makine öğrenmesi algoritmalarının yanlış karar verebilme ihtimalinin düşmesi neticesinde daha doğru sonuçlara ulaşılmaktadır.

Bir T dizisinin içerisinde bireysel sınıflandırıcıların olduğu ve $h_1, h_2, \dots, h_T$ şeklinde ifade edildiğini varsayalım. K sınıflarının olduğu bir dizide her biri c_k olarak etiketlenen k elemanlarının dizideki gösterimi $k \in [1, 2, 3, \dots, K]$ şeklinde olsun. Bir sınıflandırıcı için, $h_t \in [1, 2, 3, \dots, T]$ ise x örneğiyle ilişkili çıktılar $(h_t^1, h_t^2, h_t^3, \dots, h_t^K)$ şeklinde ifade edilerek K boyutlu bir vektör oluşturulur.

h_t^k için sınıf etiketi; $h_t^k \in \{0, 1\}$, sınıf k ise 1 değilse 0 değerini alır. Sınıf olasılıkları; $h_t^k \in \{0, 1\}$, h_t sınıflandırıcısı için $P\left(\frac{c_k}{x}\right)$ olasılığın bir tahminidir.

VC’de tahminlerin birleştirilerek bir sonuca varılması hedeflenmektedir. Her bir sınıf kendi içerisinde en sık karşılaştığı değere göre bir tahminde bulunur. Buna çoğul oylama adı verilir. VC için çoğul oylama ise Eşitlik 4.25’te hesaplanmıştır.

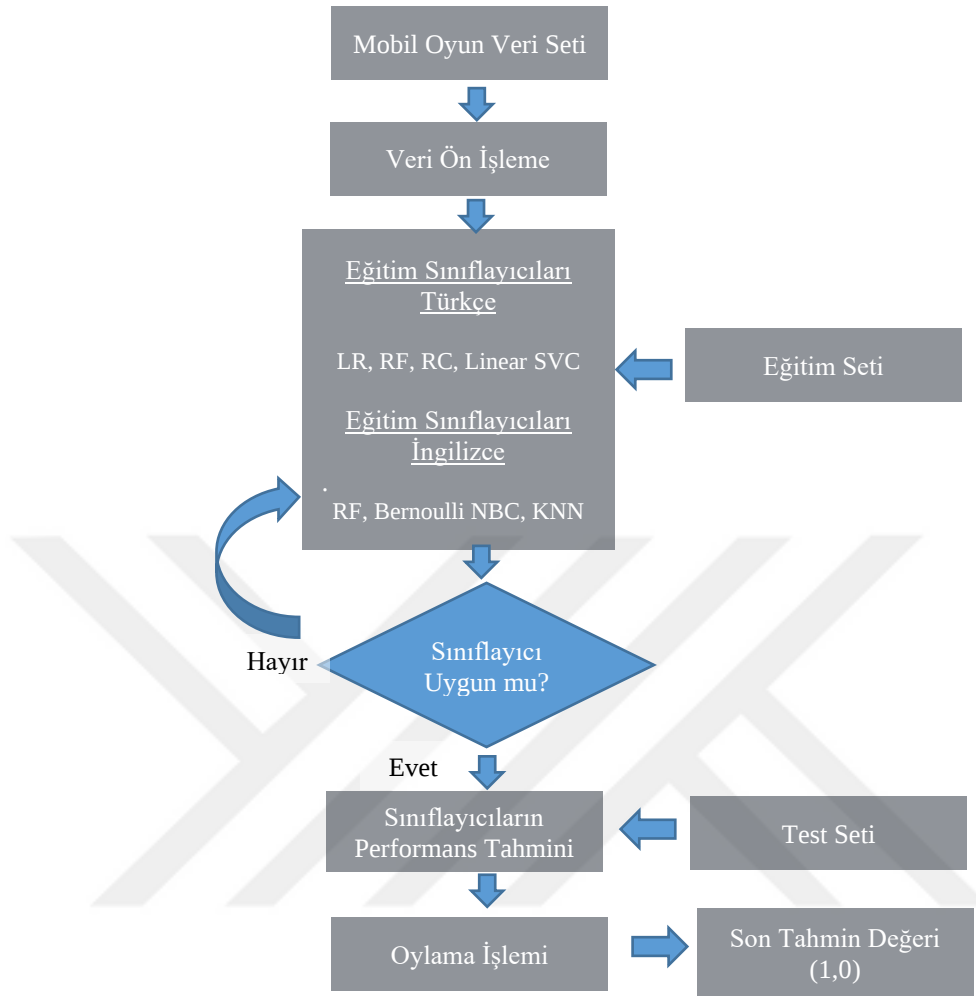
$$H(x) = c_{\text{argmax } j \sum_{i=1}^T h_i^j(x)} \quad (4.25)$$

Sınıflandırıcıların en az %50'sinin aynı sınıfı seçtiği durum ise, yine bir çoğul oylama olarak ifade edilir. Bu durum, Eşitlik 4.26’da hesaplanmıştır.

$$H(x) = c_j = \sum_{i=1}^T h_i^j(x) > \frac{1}{2} \#T \quad (4.26)$$

Burada, #T kullanılan model sayısını ifade etmektedir. c_j Eşitliği sağlanmıyorsa oylama işlemi başarısız olur.

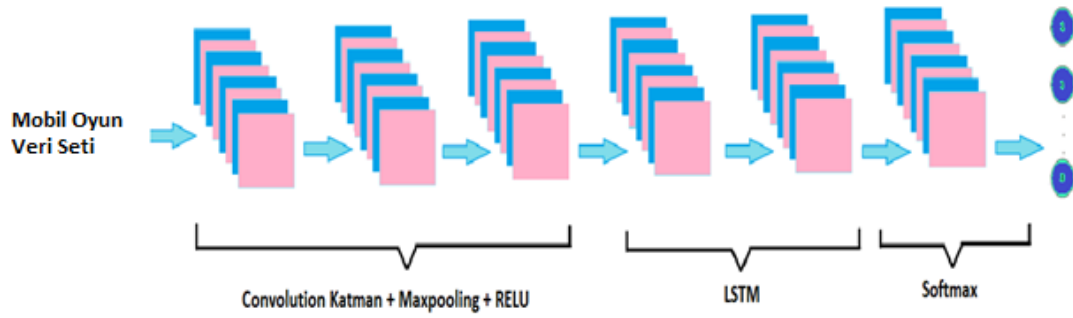
Bu tez çalışmasında, mobil oyun veri setinin VC’ye uygulandığı akış diyagramı Şekil 4.30’da gösterilmiştir. Akış Diyagramında belirtildiği üzere, Türkçe tweetler için kullanılan VC sınıflayıcısında LR, RF, RC ve Linear SVC algoritmaları kullanılmıştır. İngilizce tweetler için kullanılan VC sınıflayıcısında RF, Bernoulli NB ve KNN algoritmaları kullanılmıştır. İngilizce ve Türkçe tweetler için kullanılan VC sınıflayıcısı içerisinde kullanılan algoritmaların farklı olmasının sebebi dil farkından kaynaklı olarak en iyi sonuçları elde edebilmek için uygun algoritmaların tespitinden kaynaklanmaktadır.



Şekil 4.30. Mobil oyun Twitter veri seti için modellenen VC akış diyagramı.

4.2.7. TEMSAP-CNNLSTM

CNN ile LSTM'nin hibrit bir şekilde çalışarak daha iyi sonuçlar alınacağı varsayılarak mobil oyun Twitter verilerinde duygu analizi için **TEMSAP-CNNLSTM** (Twitter Mobile Game Sentiment Analysis Approach based on CNN-LSTM) modeli geliştirilmiştir. CNN-LSTM mimarisi, dizi tahminini desteklemek için LSTM'ler ile birleştirilmiş giriş verilerinde, özellik çıkarımı için CNN katmanlarını kullanır. Bu tez çalışmasında kullanılan **TEMSAP-CNNLSTM** hibrit modelinin mimarisi Şekil 4.31'de gösterilmiştir.

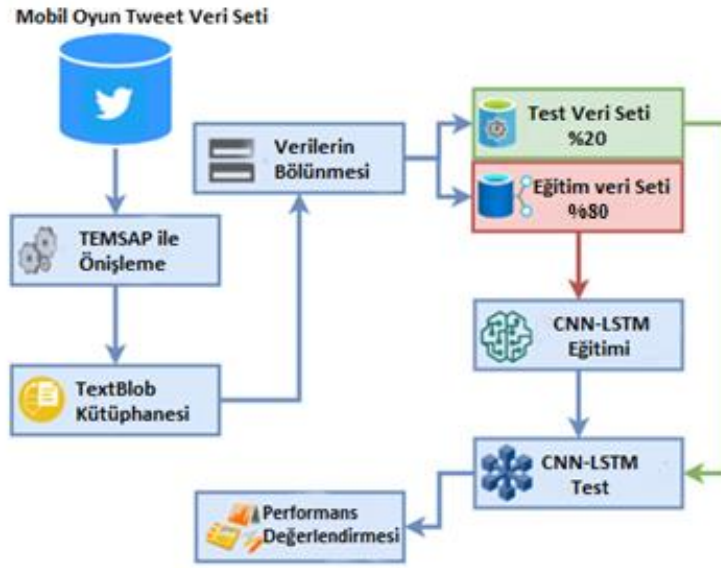


Şekil 4.31. TEMSAP-CNNLSTM mimarisi.

Modellenen TEMSAP-CNNLSTM’de, kısa tweet metin dizilerinin yanında uzun tweet metin dizilerini işleme yeteneğine sahip olarak mobil oyun kullanıcılarının oyunlar ile ilgili yazdıkları tweetlerde gizlenmiş duyguları olumlu ya da olumsuz şekilde otomatik olarak analiz edilmiştir. %50’si olumlu, %50’si olumsuz olmak üzere toplamda 2.217.708 tweet elde edilmiş olup, bunların 376.431 tanesini Türkçe, 1.839.274 tanesini ise İngilizce tweetler oluşturmaktadır. Model oluşturulurken Türkçe tweet veri setinde 11271 pozitif ve 11271 negatif tweet kullanılmıştır. İngilizce tweet veri setinde ise, 37042 pozitif ve 37042 negatif tweet kullanılmıştır. TEMSAP-CNNLSTM’nin geliştirilmesinde “TextBlob” kütüphanesi kullanılmıştır ve böylece, duygu analizi için tweet verileri ön-işleme tabii tutularak alakasız mobil oyun verilerinin temizlenmesi sağlanmıştır. Mobil oyun Tweet veri setinin %80’i eğitim ve %20’i ise test setine ayrılmıştır (Şekil 4.32). Geliştirilen TEMSAP-CNNLSTM tarafından tweet verileri eğitilerek, test veri setinde sırasıyla doğruluk (Accuracy), kesinlik (Precision), duyarlılık (Recall) ve F1-Skor (F1-Measure) metrikleri esas alınarak değerlendirmeler yapılmıştır. Mobil oyun tweet veri setleri kullanılarak duygu analizinin gerçekleştirildiği TEMSAP-CNNLSTM’nin akış diyagramı Şekil 4.33’te, taslak kodu ise Şekil 4.34’de sunulmuştur.

```
x_train, x_val, y_train, y_val = train_test_split(x_train, y_train,
                                                test_size = 0.2,
                                                random_state = 18)
```

Şekil 4.32. Tweet veri setinin %80 eğitim ve %20 test setine ayrılması için python taslak kodu.



Şekil 4.33. Geliştirilen TEMSAP-CNNLSTM'nin akış diyagramı.

```
def create_conv_model():
    model_conv = Sequential()
    model_conv.add(Embedding(vocabulary_size, 50, input_length=50))
    model_conv.add(Dropout(0.8))
    model_conv.add(Conv1D(32, 5, activation='relu'))
    model_conv.add(MaxPooling1D(pool_size=4))
    model_conv.add(LSTM(50))
    model_conv.add(Dense(1, activation='sigmoid'))
    model_conv.compile(loss='binary_crossentropy', optimizer='adam', metrics=['accuracy'])
    return model_conv

model_conv = create_conv_model()
model_conv.fit(x_train, np.array(y_train.astype("float32")), validation_data=(x_val, np.array(y_val.astype("float32"))), epochs = 3)
```

Şekil 4.34. Geliştirilen TEMSAP-CNNLSTM'nin python taslak kodu.

5. BULGULAR VE TARTIŞMA

Tezin bu bölümünde, sırasıyla;

- Uygulama aşamasında esas alınan varsayımlar,
- Kullanılan modellerin analiz süreleri,
- Modellerin Eğitim, Test ve Tüm veri olarak ayrı ayrı TN, TP, FN, FP, P, R, AUC, Acc metriklerine göre karşılaştırmaları,
- Modellerin ROC grafikleri sunulmuştur ve ulaşılan sonuçlar tartışılmıştır.

Varsayımlar

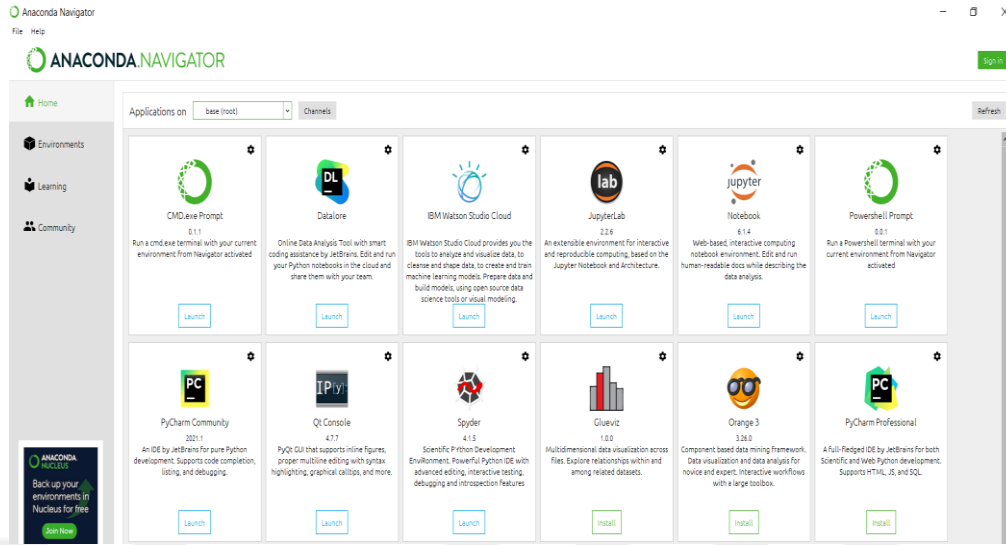
Tezin uygulaması sürecinde kullanılan programlama dilleri, kütüphaneler ve araçlar belirlenmiştir.

- Komut satırı kullanmadan, uygulamayı başlatmayı, yönetmeyi sağlayan ve bir grafik kullanıcı arabirimi olarak Anaconda Navigator-JupyterLab 2.2.6 versiyonu (Şekil 5.1) kullanılmıştır.
- Kodlama platformu olarak python 3.8.5 programlama dili kullanılarak verilerin analizleri gerçekleştirilmiştir.
- Python programlama dilinde bazı denetimli ve denetimsiz öğrenme algoritmalarını sunan ve bir makine öğrenmesi kitaplığı olan Scikit-Learn (Şekil 5.2) kullanılmıştır.
- Derin öğrenme kütüphanesi olarak Keras, açık kaynaklı bir matematik kütüphanesi olarak ta TensorFlow kullanılmıştır (Şekil 5.3).
- Veri setinin ön-işlem aşamasında, makine öğrenmesi ve doğal dil işleme kütüphanesi olarak NLTK (The Natural Language Toolkit) kullanılmıştır.
- Python'da verileri saklamak ve analiz etmek için Pandas kütüphanesi kullanılmıştır.
- Python'da dizi ve matris işlemleri için NumPy kitaplığı kullanılmıştır.
- Veri setinde geçen bir kelimenin geçme sıklığını temsil etmek için veri görselleştirme tekniği olarak Kelime Bulutu (WordCloud) kullanılmıştır (Şekil 5.4).
- Tweet metinlerinin olumlu ya da olumsuz bir içeriğe sahip olduğunu tespit etmek amacıyla TextBlob kütüphanesi kullanılmıştır.
- Çizgi, dağılım, alan, çubuk veya ROC grafikleri oluşturmak amacıyla Plotly kütüphanesi kullanılmıştır (Şekil 5.5).

- Twitterda mobil oyunlarla alakalı 376.431 adet Türkçe ve 1.839.274 adet İngilizce olmak üzere toplam 2.215.708 adet tweet kullanılmıştır.
- Probleme uygulanan modeller TDF-IDF ve CountVectorizer olmak üzere iki farklı vektör tipine göre incelenmiştir.
- ROC eğrisi altındaki alanların (AUC) yorumlanmasında aşağıdaki derecelendirmeler (Dirican, 2001; Kanık ve Erden, 2003) esas alınmıştır:
 - %90-%100, mükemmel
 - %80-%90, iyi
 - %70-%80, orta
 - %60-%70, zayıf
 - %50-%60, başarısız
- Phyton; 1.80Ghz CPU, 8GB RAM ve 256GB HDD donanımına sahip bir bilgisayar kullanılarak çalıştırılmıştır.
- JupyterLab platformu için Çizelge 5.1'deki parametreler ve değerleri esas alınmıştır.

Çizelge 5.1. Kütüphanelerde kullanılan parametreler ve değerleri

Parametre Adı	Değeri
Convolution Layer Aktivasyon Fonksiyonu	Activation.RELU
Havuz tipi	Maximum
Çıkış Katmanı Aktivasyon Fonksiyonu	Activation.SOFTMAX
LSTM Katmanı Aktivasyon Fonksiyonu	Activation.RELU
Çıkış Katmanı Kayıp Fonksiyonu	Activation.SIGMOID
Çıkış Katmanı Aktivasyon Fonksiyonu	Activation.SOFTMAX



Şekil 5.1. Anaconda Navigator-JupyterLab 2.2.6 versiyonu web arayüzü.

```
from sklearn.feature_extraction.text import TfidfVectorizer, CountVectorizer

from sklearn.metrics import accuracy_score, roc_auc_score, roc_curve
from sklearn.model_selection import train_test_split
from sklearn.metrics import confusion_matrix
from sklearn.metrics import precision_recall_fscore_support
from sklearn import metrics

from sklearn.linear_model import LogisticRegression
from sklearn.svm import LinearSVC
from sklearn.naive_bayes import MultinomialNB, BernoulliNB
from sklearn.linear_model import RidgeClassifier
from sklearn.linear_model import PassiveAggressiveClassifier
from sklearn.neural_network import MLPClassifier
from sklearn.neighbors import KNeighborsClassifier
from sklearn.ensemble import VotingClassifier
from sklearn.ensemble import RandomForestClassifier

from sklearn.feature_selection import SelectFromModel
```

Şekil 5.2. Tez çalışmasında Scikit-Learn kullanımı ve python kodu.

```
# Keras
import tensorflow
from tensorflow.keras.preprocessing.text import Tokenizer
from tensorflow.keras.preprocessing.sequence import pad_sequences
from tensorflow.keras.models import Sequential
from tensorflow.keras.layers import Dense, Flatten, LSTM, Conv1D, MaxPooling1D, Dropout, Activation
from tensorflow.keras import keras
from tensorflow.keras.layers import Embedding
from tensorflow.keras import Sequential
from tensorflow.keras.layers import Conv2D, Flatten, Dense
```

Şekil 5.3. Tez çalışmasında Keras ve Tensorflow kullanımı ve python kodu.

```

#Creating the text variable
text = " ".join(title for title in en_df.kok_tweet)

# Creating word_cloud with text as argument in .generate() method
word_cloud = WordCloud(collocations = False, background_color = 'white').generate(text)

# Display the generated Word Cloud
plt.imshow(word_cloud, interpolation='bilinear')

plt.axis("off")

plt.show()

```

Şekil 5.4. Tez çalışmasında kelime bulutunun oluşturulması ve python kodu.

```

## Plotly
import plotly.offline as py
import plotly.graph_objs as go
py.init_notebook_mode(connected=True)

```

Şekil 5.5. Tez çalışmasında Plotly kütüphanesinin kullanımı ve python kodu.

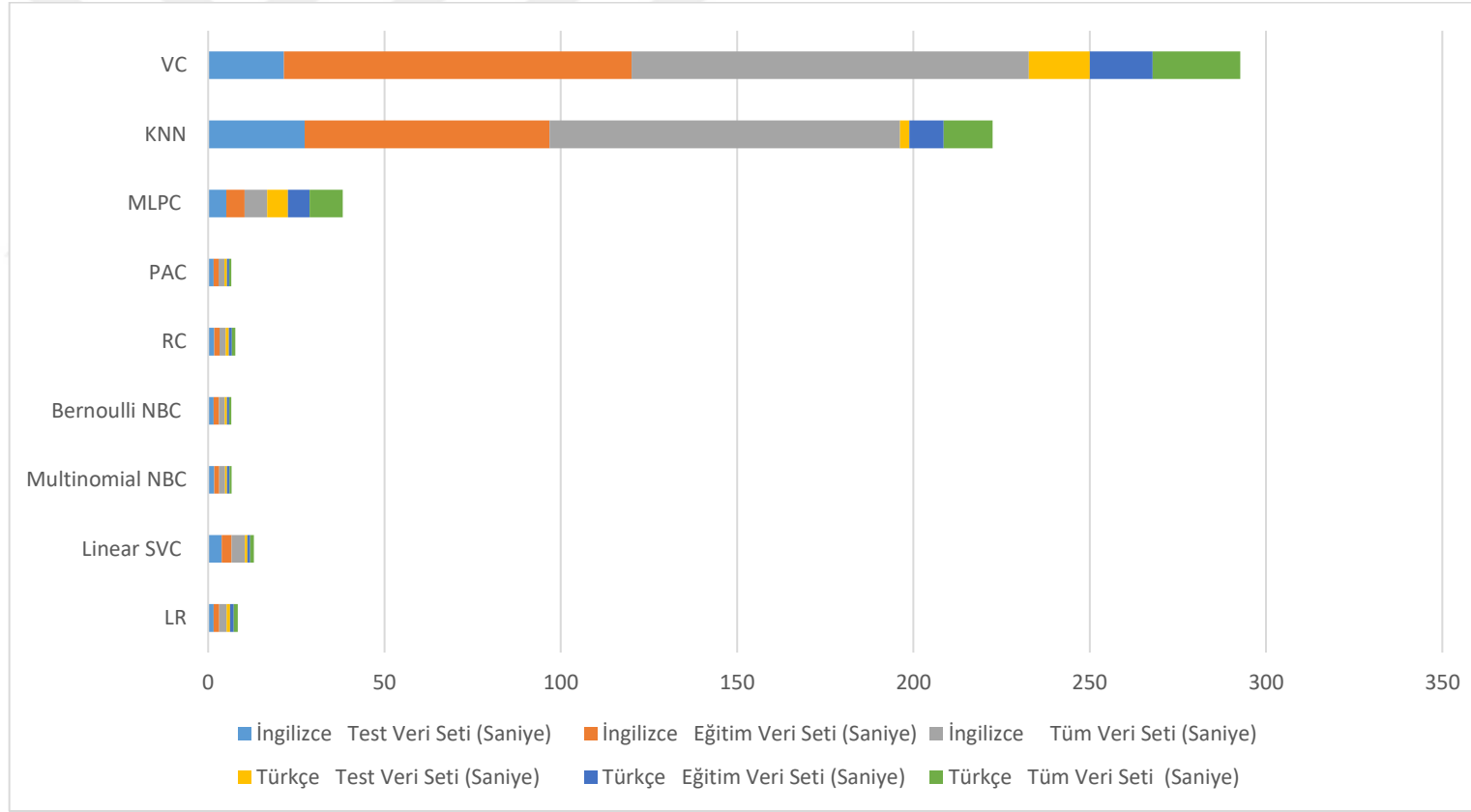
5.1. Çalışma Süreleri Analizi

Tez çalışmasında, modellerin çalışma süreleri tespit edilerek, daha hızlı sonuç veren modellerin belirlenmesi sağlanmıştır. Türkçe ve İngilizce tweetler için vektör tiplerine göre test verilerinin makine öğrenmesi algoritmalarında çalışma süreleri hesaplanarak Çizelge 5.2’de karşılaştırılması da grafik olarak Şekil 5.6 ve 5.7’de gösterilmiştir. CountVectorizer vektör tipine esas alındığında; VC modeli toplamda 292.69 saniye ile analizini diğer basit sınıflandırma algoritmalarına göre en geç tamamlayan algoritma olmuştur. KNN modeli toplamda 222.47 saniye ile ikinci sırada, MLPC modeli toplamda 38.18 saniye ile üçüncü sırada tamamlamıştır. TF-IDF vektör tipi esas alındığında ise; VC modeli toplamda 297.86 saniye ile yine analizini en geç tamamlayan algoritma olmuştur. KNN modeli toplamda 277.67 saniye ile ikinci sırada, MLPC modeli toplamda 38.74 saniye ile üçüncü sırada tamamlamıştır. Geliştirilen TEMSAP-CNNLSTM modeli ise, tüm modeller içerisinde 285.5 saniye ile VC modelinden daha hızlı, diğer basit makine öğrenme algoritmalarından ise daha yavaş bir

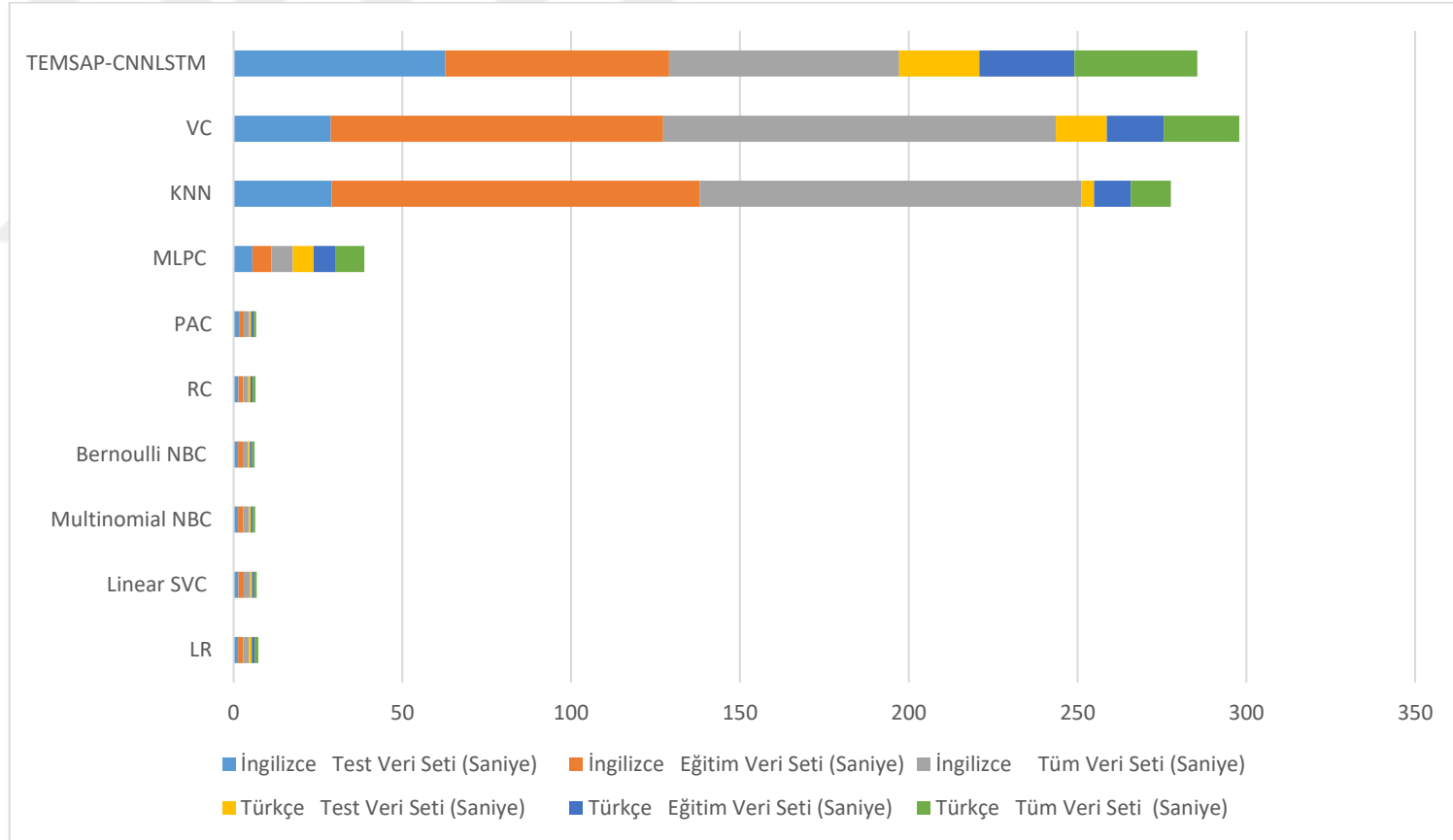
sürede analizini tamamlamıştır. VC modelini oluşturan diğer sınıflama algoritmaları; verinin eğitilmesi, test verilerinin sonuçlarına göre oylama işlemini gerçekleştirmeleri ve tüm bu sonuçlara göre bir tahmin üretmelerinden dolayı analizini geç tamamlamaktadır. TEMSAP-CNNLSTM modeli, evrişim katmanı, havuz katmanı ve kullanmış olduğu ReLU aktivasyon fonksiyonuyla düzleştirme katmanında yaptığı işlemlerin yanında LSTM katmanında uzun metinlerin analizi için gereken sürenin fazla olmasından kaynaklı olarak VC modelinden sonra en yavaş çalışan model olmuştur. Diğer kullanılan basit makine öğrenmesi algoritmaları matematiksel işlemler ve öğrenme için fazla süreye ihtiyaç duymamaktadır. Dolayısıyla, VC modeli ile TEMSAP-CNNLSTM modellerinin diğer basit makine öğrenmesi algoritmalarına göre daha yavaş çalışması normal kabul edilir. Modellerin tümünde eğitim veri setlerinin çalışma sürelerinin daha uzun olduğu tespit edilmiştir. Nedenlerinden ilki, model oluşturmak için kullanılan verinin %80’ni eğitim veri setini oluşturmaktadır ve bu sebeple verinin fazla olması süreyi artırmaktadır. İkincisi ise, eğitim veri seti kullanılarak öğrenme işleminin gerçekleştirilmesidir ve bu sebeple analiz sürecinin uzamasına neden olmuştur.

Çizelge 5.2. Twitter mobil oyun veri seti için vektör tiplerine göre algoritmaların test, eğitim ve tüm veri setindeki çalışma süreleri (Saniye türünden)

Vektör tipi	Algoritma	İngilizce			Türkçe		
		Test Veri Seti	Eğitim Veri Seti	Tüm Veri Seti	Test Veri Seti	Eğitim Veri Seti	Tüm Veri Seti
CountVectorizer	LR	1.58	1.54	2.08	1.02	1.02	1.14
	Linear SVC	3.81	2.86	3.76	0.72	0.72	1.07
	Multinomial NBC	1.68	1.41	1.71	0.55	0.65	0.63
	Bernoulli NBC	1.57	1.43	1.73	0.59	0.63	0.63
	RC	1.69	1.52	1.72	0.92	0.9	0.93
	PAC	1.54	1.39	1.7	0.68	0.59	0.67
	MLPC	5.11	5.19	6.42	5.93	6.11	9.42
	KNN	27.34	69.51	99.33	2.7	9.67	13.92
	VC	21.52	98.53	112.66	17.34	17.76	24.88
Tfidf Vectorizer	LR	1.36	1.45	1.79	0.75	0.83	1.15
	Linear SVC	1.46	1.51	1.92	0.59	0.62	0.75
	Multinomial NBC	1.37	1.43	1.73	0.55	0.6	0.7
	Bernoulli NBC	1.33	1.42	1.56	0.56	0.6	0.75
	RC	1.4	1.38	1.63	0.61	0.67	0.79
	PAC	1.57	1.41	1.66	0.58	0.76	0.68
	MLPC	5.58	5.6	6.33	6.25	6.31	8.67
	KNN	29.08	108.93	113.12	3.77	10.83	11.94
	VC	28.78	98.44	116.36	15.07	16.82	22.39
TEMSAP-CNNLSTM		62.7	66.3	68.1	23.8	28.2	36.4



Şekil 5.6. İngilizce ve Türkçe tweetlerdeki test, eğitim ve tüm veri seti için CountVectorizer vektör tipine göre algoritmaların çalışma süreleri.



Şekil 5.7. İngilizce ve Türkçe tweetlerdeki test, eğitim ve tüm veri seti için TF-IDF vektör tipine göre algoritmaların çalışma süreleri.

5.2. Mobil Oyun İngilizce Veri Seti Analizleri

Güney Afrika, Kanada, İngiltere, Fransa, İtalya, Güney Kore ve Amerika Birleşik Devletleri olmak üzere toplamda 7 ülkeden 2015-2021 yılları arasında toplam 1.839.274 İngilizce tweete ulaşarak Mobil oyun İngilizce veri seti (Çizelge 4.1 ve Çizelge 4.2) oluşturulmuştur. Yapılan analizler için LR, Linear SVC, Multinomial NBC, Bernoulli NBC, RC, PAC, MLPC, KNN, VC ve geliştirilen TEMSAP-CNNLSTM modelleri kullanılmıştır. Ayrıca RF, Bernoulli NBC ve KNN sınıflama algoritmaları birlikte kullanılarak önerilen VC modeli oluşturulmuştur. Kullanılan tüm modeller TF-IDF ve CountVectorizer vektör tipleri esas alınarak karşılaştırılmıştır. Karşılaştırma ölçütleri olan TN, TP, FN, FP, P, R, F, AUC ve Accuracy (Doğruluk) değerleri tablo halinde sunulmuş ve ROC eğrileri grafiksel olarak ayrıca gösterilmiştir. Hibrit bir model olarak geliştirilen TEMSAP-CNNLSTM modelinde derin öğrenme yöntemleri kullandığı için diğer tüm modellerden ayrı olarak veri setlerine uygulanmıştır. Mobil oyun İngilizce tweet verisinin %80'i eğitim (59.267 tweet) ve %20'si test (14.817 tweet) verisi olarak ayrılmıştır. Test, eğitim ve tüm veri olmak üzere üç farklı veri seti kullanılarak analiz yapılmıştır.

Mobil oyun İngilizce veri seti için doğruluk ölçütlerinin tanımlanmasında kullanılan python kodu Şekil 5.8'de sunulmuştur.

```
en_df_dict = {"Model Adı": [], "Vektör Tipi": [], "Sonuç": [], "Süre": [], "True-Negative": [],
              "True-Positive": [], "False-Negative": [], "False-Positive": [], "P": [], "R": [], "F": [], "auc": []}

for i, vector_name in enumerate(en_df_acc):
    for j, model_name in enumerate(en_df_acc[vector_name]):
        for k in range(len(en_df_acc[vector_name][model_name])):

            en_df_dict["Vektör Tipi"].append(vector_name)
            en_df_dict["Model Adı"].append(model_name)
            en_df_dict["Sonuç"].append(en_df_acc[vector_name][model_name][k])
            en_df_dict["Süre"].append(en_df_time[vector_name][model_name][k])
            en_df_dict["True-Negative"].append(en_df_tn[vector_name][model_name][k])
            en_df_dict["True-Positive"].append(en_df_tp[vector_name][model_name][k])
            en_df_dict["False-Negative"].append(en_df_fn[vector_name][model_name][k])
            en_df_dict["False-Positive"].append(en_df_fp[vector_name][model_name][k])
            en_df_dict["P"].append(pre[vector_name][model_name][k])
            en_df_dict["R"].append(rec[vector_name][model_name][k])
            en_df_dict["F"].append(f1[vector_name][model_name][k])
            en_df_dict["auc"].append(auc_score[vector_name][model_name][k])

en_df_sonucclar = pd.DataFrame(en_df_dict)
en_df_sonucclar
```

Şekil 5.8. JupyterLab platformunda karşılaştırma ölçütlerinin tanımlanması ve python kodu.

5.2.1. Eğitim veri seti analizi

Twitterdan ulařılan mobil oyun İngilizce eğitim veri seti için her iki vektör tipi esas alınarak kullanılan modeller tarafından hesaplanan doğruluk ölçütleri Çizelge 5.3’de gösterilmiştir.



Çizelge 5.3. Mobil oyun İngilizce eğitim veri seti için hesaplanan metrik ölçütler

Vektör Tipi	Model Adı	Acc	TN	TP	FN	FP	P	R	F	AUC %
CountVectorizer	LR	0.89	26799	26026	3598	2844	0.90149	0.878544	0.889869	89.13
	Linear SVC	0.83	24734	24747	4877	4909	0.83447	0.83537	0.834919	83.49
	Multinomial NBC	0.7	21473	20239	9385	8170	0.71242	0.683196	0.6975	70.38
	Bernoulli NBC	0.81	23117	25187	4437	6526	0.79422	0.850223	0.821266	81.5
	RC	0.86	23967	27219	2405	5676	0.82745	0.918816	0.870743	86.37
	PAC	0.68	20411	19994	9630	9232	0.68412	0.674926	0.67949	68.17
	MLPC	0.84	23335	27003	2621	6308	0.81063	0.911524	0.858123	84.94
	KNN	0.88	25307	27012	2612	4336	0.86168	0.911828	0.886046	88.28
	VC	0.92	25576	28955	669	4067	0.87684	0.977417	0.924401	92.01
Tfidf Vectorizer	LR	0.89	24764	28201	1423	4879	0.85251	0.951965	0.899496	89.37
	Linear SVC	0.83	22906	26718	2906	6737	0.79863	0.901904	0.847128	83.73
	Multinomial NBC	0.7	20462	21574	8050	9181	0.70148	0.728261	0.714619	70.93
	Bernoulli NBC	0.81	23523	24833	4791	6120	0.80228	0.838273	0.819882	81.59
	RC	0.86	24177	26966	2658	5466	0.83146	0.910275	0.869086	86.29
	PAC	0.68	21581	19023	10601	8062	0.70234	0.642148	0.670899	68.51
	MLPC	0.84	23899	26348	3276	5744	0.82102	0.889414	0.853847	84.78
	KNN	0.88	24697	27777	1847	4946	0.84885	0.937652	0.891045	88.54
	VC	0.92	26902	28007	1617	2741	0.91086	0.945416	0.927814	92.65

Çizelge 5.3'deki metrik ölçütlerine göre;

R; pozitif tweetlerin doğru tespit edilme oranını vermektedir. İki vektör tipi içinde Multinomial NBC ve PAC modelleri hariç bütün modellerde %83'ün üzerinde bir başarı oranı elde edilmiştir. Voting Classifier modeli *R* değeri için CountVectorizer vektör tipinde 0.97'lik, TF-IDF vektör tipinde ise 0.94'lük bir orana ulaşarak en başarılı model olmuştur. Dolayısıyla VC modeli pozitif tweetleri en doğru belirleyen model olarak tespit edilmiştir.

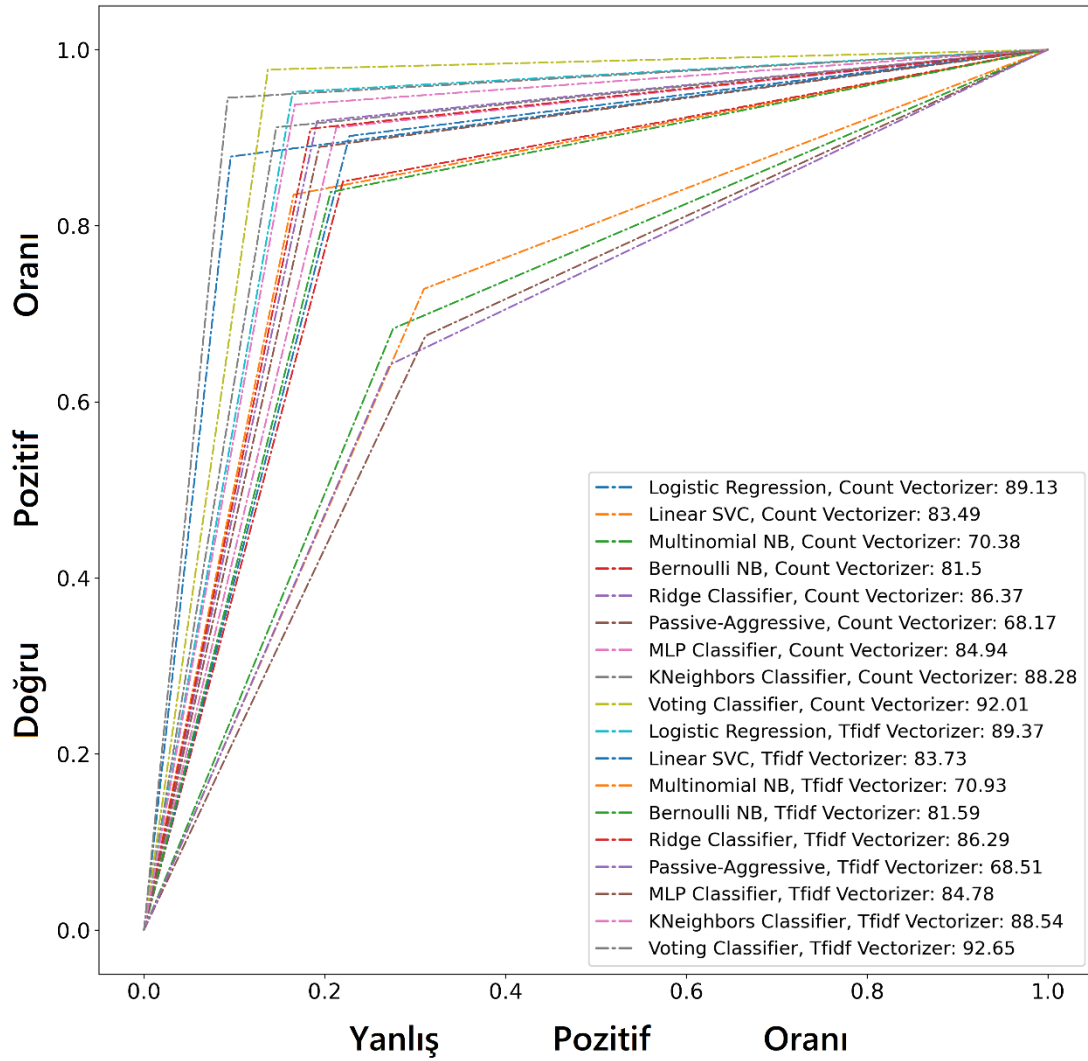
P; pozitif olarak belirlenen tweetlerin gerçekten kaçının pozitif olduğunu göstermektedir. CountVectorizer vektör tipine göre; PAC modelinin 0.68'lik ve Multinomial NBC modelinin 0.71'lik oranlarla hesaplanan en düşük değerli iki model olduğu tespit edilmiştir. LR modelinin 0.90'lık ve VC modelinin 0.87'lik oranlarla hesaplanan en iyi değerleri iki model olduğu tespit edilmiştir. TF-IDF'de VC 0.91 ile diğer tüm modellerden daha iyi bir sonuç elde edilmiştir.

F; *R* ve *P* değerleri arasında doğruluğun tespit edilme oranı ve gerçekten kaç doğru tahmin yapıldığı konusunda tutarsız durumların görülebilme ihtimaline binaen hesaplanan bir ölçüttür. *F* incelemesinde, iki vektör tipi içinde VC modelinin 0.92'lik oranla en iyi değere ulaştığı, PAC ve Multinomial NBC ise en düşük değere ulaştığı tespit edilmiştir.

ROC; Sınıflandırma modelinin başarısını belirlemede faydalanan bir ölçüttür. AUC değeri 1'e yakınsadığı ölçüde modelin başarısı artmaktadır. Hesaplamada, *R* ve FPR (False Positive Rate) değerine göre esas alınmaktadır. FPR değeri pozitif olmayan tweetlere hangi oranda hatalı yaklaşıldığını göstermektedir. Hesaplanan AUC'ye göre, iki vektör tipi içinde VC modeli %92'lik bir mükemmel orana ulaşmıştır. *R*, *P* ve *F* değerleri dikkate alındığında, PAC ve Multinomial NBC modellerinin en düşük AUC değerine ulaştığı tespit edilmiştir. Hiçbir modelde AUC başarısız bir orana ulaşmamıştır. Karşılaştırılan modellerin ROC grafiği şekil 5.9'da gösterilmiştir.

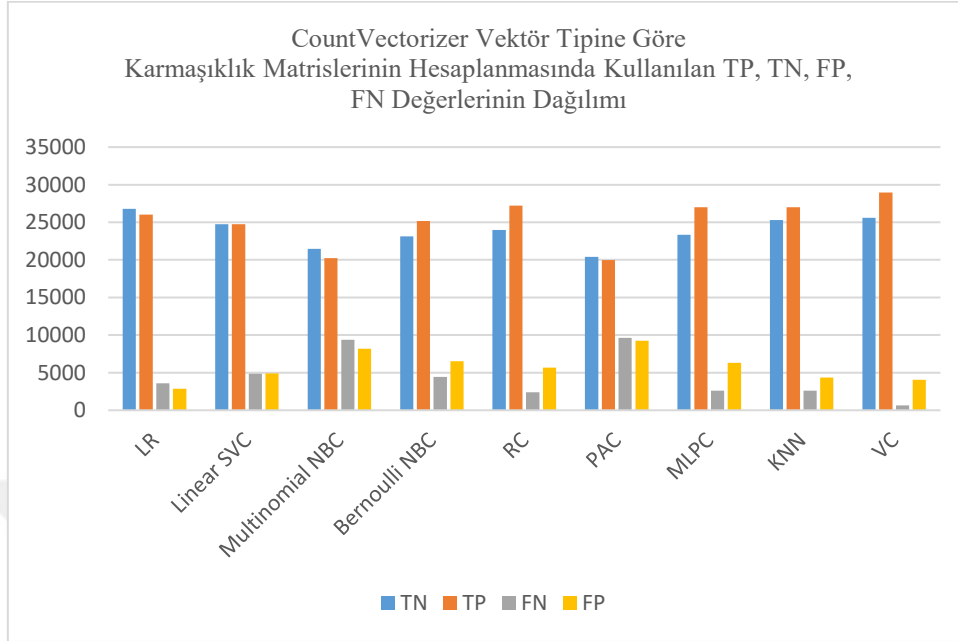
Acc; doğru tespit edilmiş pozitif tweetlerin ve doğru tespit edilmiş negatif tweetlerin toplamının tüm tweet sayısına bölünmesi ile hesaplanmakta ve modelin doğruluk oranını temsil etmektedir. Hesaplanan Acc'ye göre, her iki vektör tipinde de VC modeli 0.92'lik oran ile en iyi model olarak tespit edilmiştir.

Bu sonuçlara göre, mobil oyun İngilizce eğitim veri seti için VC modeli tarafından yapılan doğruluk ölçütü hesaplamalarında optimum sonuçları ulaşıldığı tespit edilmiş ve veri setinin %80'nini oluşturan eğitim veri seti ile karşılaştırılan modellerin öğrenme sürecinde genel olarak başarılı olduğu belirlenmiştir.

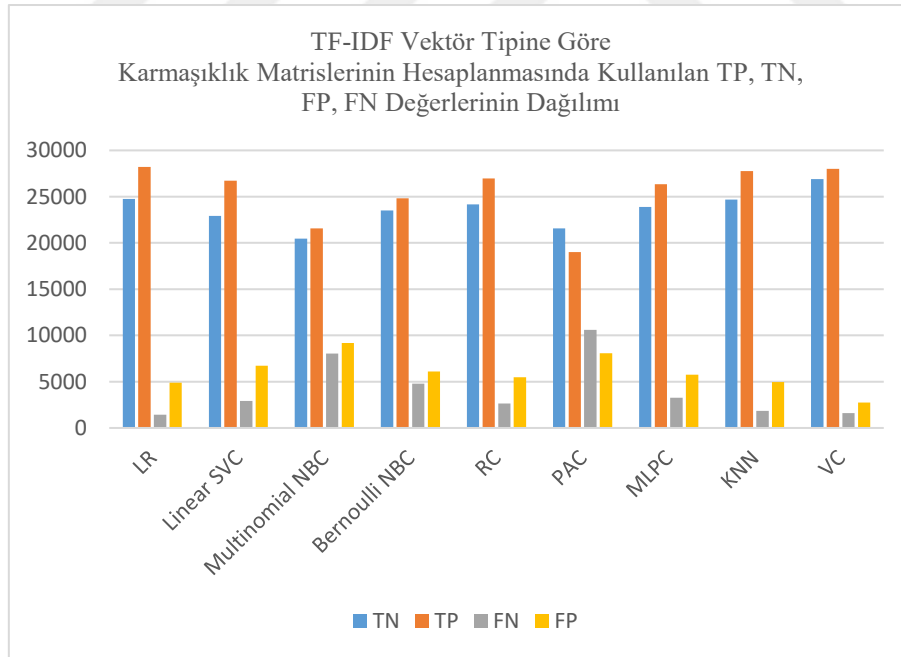


Şekil 5.9. Mobil oyun İngilizce veri seti – eğitim veri seti ROC grafiği.

Mobil oyun İngilizce eğitim veri seti kullanılarak karmaşıklık matrislerinin hesaplanmasında kullanılan TN, TP, FN, FP değerlerinin CountVectorizer’da dağılımı şekil 5.10’da TF-IDF’de dağılımı ise Şekil 5.11’de sunulmuştur. FN (Yanlış hesaplanan olumsuz tweetler) ve FP (Yanlış hesaplanan olumlu tweetler) değerlerinin, TP (Doğru olumlu tweetler) ve TN (Doğru olumsuz tweetler) değerlerine göre oldukça düşük olması; modelin iyi çalıştığını göstermekle birlikte polaritelerinin iyi düzenlendiğini de ispatlamaktadır. TP ve TN değerlerinin optimum olduğu model VC olarak tespit edilmiştir. Dolayısıyla, karmaşıklık matrislerinin hesaplanmasında en doğru sonuçlar VC modeli tarafından ulaşılmıştır.



Şekil 5.10. Mobil oyun İngilizce eğitim veri setinde CountVectorizer vektör tipine göre TP, TN, FN, FP dağılımı.



Şekil 5.11. Mobil oyun İngilizce eğitim veri setinde TF-IDF vektör tipine göre TP, TN, FN, FP dağılımı.

5.2.2. Test veri seti analizi

Twitter'daki mobil oyun İngilizce test veri seti için her iki vektör tipi esas alınarak kullanılan modeller tarafından hesaplanan doğruluk ölçütleri Çizelge 5.4'te gösterilmiştir.



Çizelge 5.4. Mobil oyun İngilizce test veri seti için hesaplanan metrik ölçütler

Vektör Tipi	Model Adı	Acc	TN	TP	FN	FP	P	R	F	AUC %
CountVectorizer	LR	0.88	5837	7211	207	1562	0.82195	0.972094	0.890741	88.05
	Linear SVC	0.82	5119	7134	284	2280	0.75781	0.961714	0.847671	82.68
	Multinomial NBC	0.7	3847	6548	870	3552	0.64832	0.882717	0.747573	70.13
	Bernoulli NBC	0.79	4861	6894	524	2538	0.73092	0.929361	0.818278	79.32
	RC	0.83	5119	7234	184	2280	0.76035	0.97519	0.854476	83.35
	PAC	0.67	3581	6412	1006	3818	0.62678	0.864383	0.726654	67.42
	MLPC	0.82	4829	7363	55	2570	0.74127	0.992585	0.848711	82.26
	KNN	0.86	5402	7398	51	1966	0.79049	0.99296	0.88299	86.71
	VC	0.91	6133	7397	21	1266	0.85386	0.997169	0.919967	91.3
Tfidf Vectorizer	LR	0.88	5801	7313	105	1598	0.82067	0.985845	0.895707	88.49
	Linear SVC	0.82	5083	7105	313	2316	0.75417	0.957805	0.843874	82.24
	Multinomial NBC	70.19	3682	6718	700	3717	0.64379	0.905634	0.75259	70.16
	Bernoulli NBC	0.79	4514	7297	121	2885	0.71666	0.983688	0.829204	79.69
	RC	0.83	5070	7366	52	2329	0.75977	0.99299	0.860866	83.91
	PAC	0.67	3599	6438	980	3800	0.62883	0.867888	0.72927	67.72
	MLPC	0.82	5002	7281	137	2397	0.75232	0.981531	0.851778	82.88
	KNN	0.86	5617	7133	285	1782	0.80011	0.961579	0.873446	86.04
	VC	0.91	6175	7396	22	1224	0.858	0.997034	0.922309	91.58

Çizelge 5.4'teki metrik ölçütlerine göre;

R için VC her iki vektör tipinde de 0.99'luk bir orana ulaşarak en başarılı model olmuştur. MLPC, LR ve KNN modelleri her iki vektör tipi içinde 0.96'nın üzerine ulaşarak başarılı değerler elde etmişlerdir. Genel olarak R değeri için karşılaştırılan tüm modeller her iki vektör tipi içinde 0.86'nın üzerinde oranlara ulaşmışlardır. En düşük değerler ilk olarak 0.86'lık oranla PAC, ikinci olarakta 0.88'lik oranla Multinomial NBC modellerinde tespit edilmiştir.

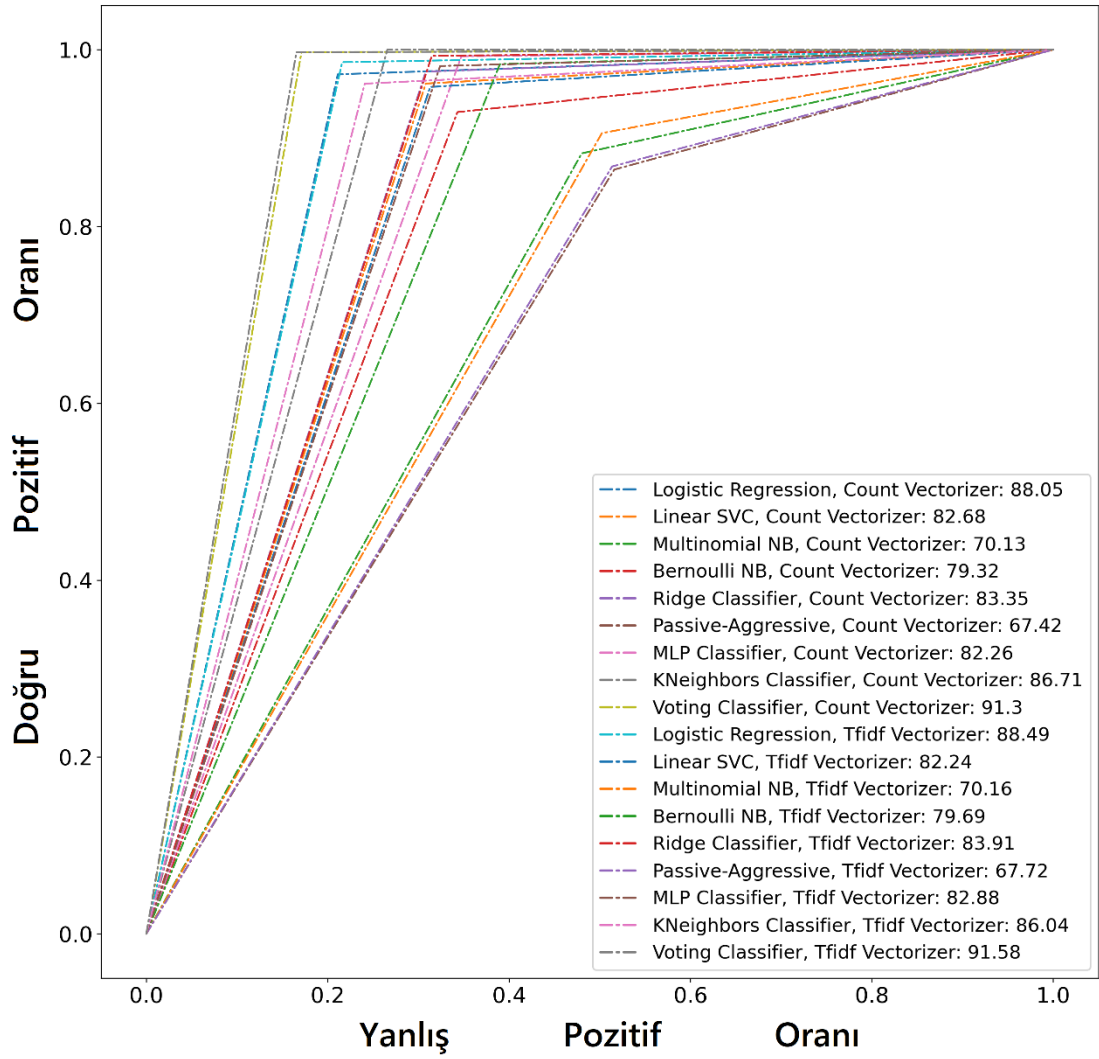
P için VC modeli her iki vektör tipinde 0.85'lik oranlarla en başarılı model olduğu tespit edilmiştir. PAC 0.62'lik ve Multinomial NBC 0.64'lük oranlarla diğer tüm modellere kıyasla en başarısız modeller olarak tespit edilmiştir.

F için VC modelinin CountVectorizer'da 0.91'lik, TF-IDF'de ise 0.92'lik oranlarla en başarılı olduğu tespit edilmiştir. PAC ve Multinomial NBC, en başarısız modeller olmuşlardır. Diğer karşılaştırılan modellerde ise, her iki vektör tipi için 0.81 ve üzerinde oranlarla başarılı sonuçlara ulaşıldığı tespit edilmiştir.

AUC için VC modelinin CountVectorizer'da %91.3'lük, TF-IDF'de ise %91.58'lik oranlarla mükemmel oranlara ulaşıldığı tespit edilmiştir. R, P, F metrik değerler dikkate alındığında, PAC ve Multinomial NBC modellerinde en düşük AUC değerlerine ulaşılmıştır. PAC modeli için zayıf (%67), Multinomial NBC için orta (%70) oranlara ulaşılmıştır. Karşılaştırılan modellerin ROC grafiği Şekil 5.12'de gösterilmiştir.

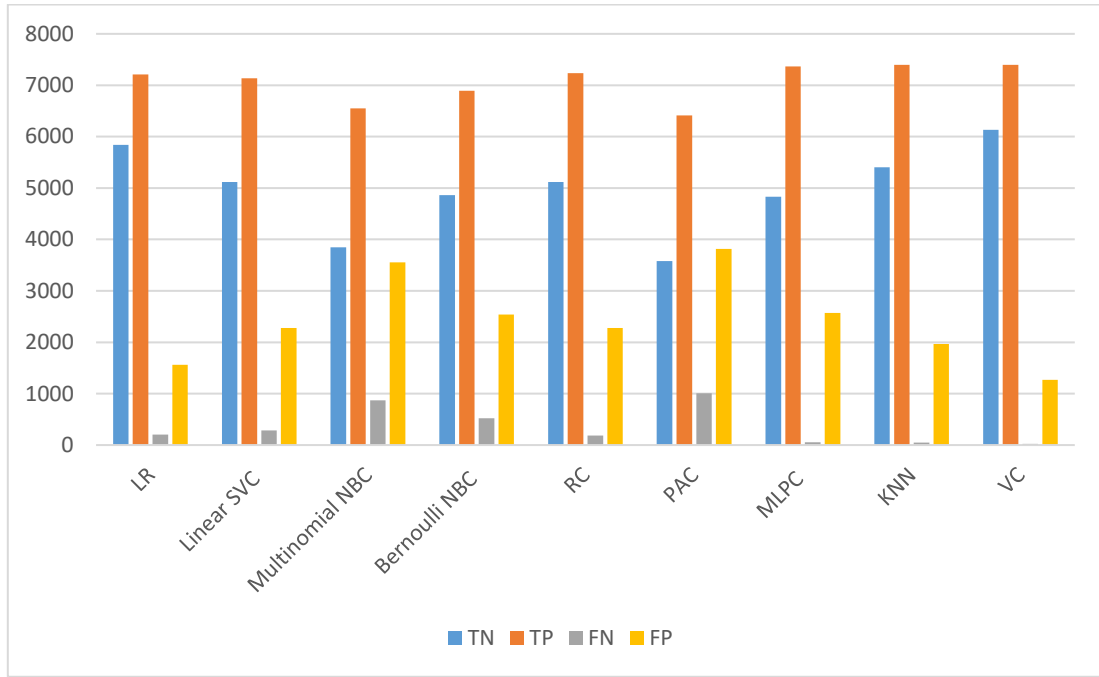
Acc için VC 0.91 ile en başarılı model olarak tespit edilmiştir.

Twitterdan oluşturulan mobil oyun İngilizce veri setindeki test veri setinin analizi sonucunda, LR, RF, RC ve Linear SVC modellerinin işbirlikçi olarak çalışmasıyla oluşturulan VC modeli tarafından en iyi sonuçlara ulaşıldığı tespit edilmiştir. Dolayısıyla Twitter mobil oyun test verilerinin duygu analizinde VC'nin bir model olarak önerilebileceği sonucuna varılmıştır. Mobil oyun veri setinin %20'sini oluşturan test veri seti üzerinden ulaşılan bu sonuçlara göre, mobil oyun eğitim veri setinin öğrenme sürecinde başarılı olduğu tespit edilmiştir.

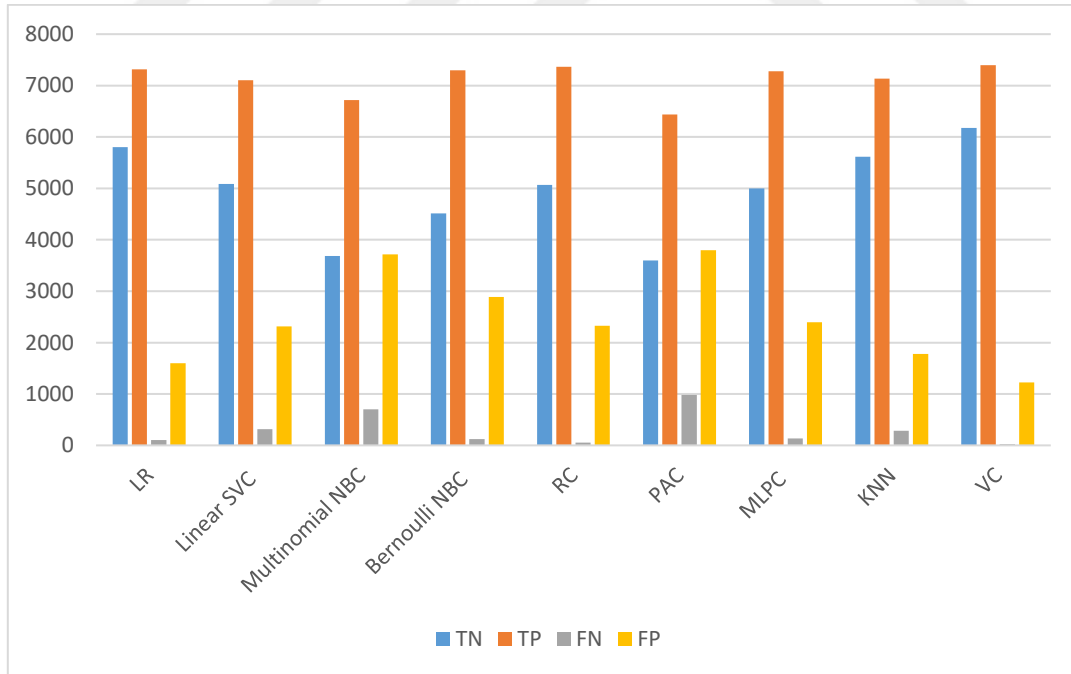


Şekil 5.12. Mobil oyun İngilizce veri seti – test veri seti ROC grafiği.

Mobil oyun İngilizce test veri setinde karmaşıklık matrislerinin hesaplanmasında kullanılan TN, TP, FN, FP değerlerinin CountVectorizer için dağılımı Şekil 5.13’de, TF-IDF için dağılımı ise Şekil 5.14’te gösterilmiştir.



Şekil 5.13. Mobil oyun İngilizce test veri setinde CountVectorizer için TP, TN, FN ve FP dağılımı.



Şekil 5.14. Mobil oyun İngilizce test veri setinde TF-IDF için TP, TN, FN ve FP dağılımı.

FN ve FP deęerlerinin TP ve TN deęerlerine gre olduka dřk olması oluřturulan modelin iyi alıřtıęını gstermekle birlikte, polaritelerinin iyi dzenlendięini de ispatlamaktadır. TF-IDF iin hesaplanan P, R, F, AUC ve Acc deęerlerinde daha iyi sonulara ulařıldıęı, fakat anlamlı bir farkın oluřmadıęı tespit edilmiřtir. VC modeli ile TP ve TN zerinden optimum deęerlere ulařıldıęı ve bylece karmařıklık matrislerinin hesaplanmasında bařarılı sonulara ulařıldıęı tespit edilmiřtir.

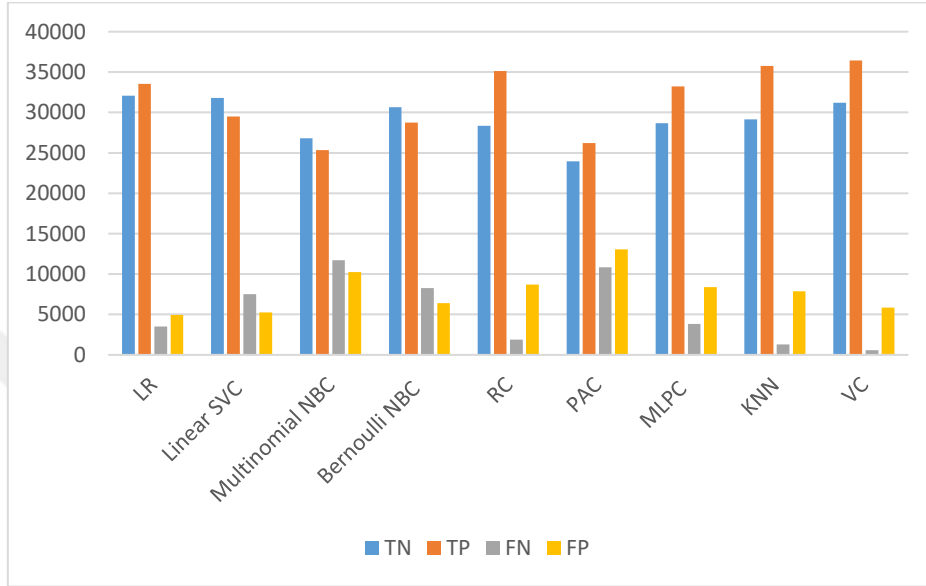
5.2.3. Tm veri seti analizi

Mobil oyun İngilizce tm veri seti iin her iki vektr tipi esas alınarak kullanılan modeller tarafından hesaplanan doęruluk ltleri izelge 5.5'te gsterilmiřtir.

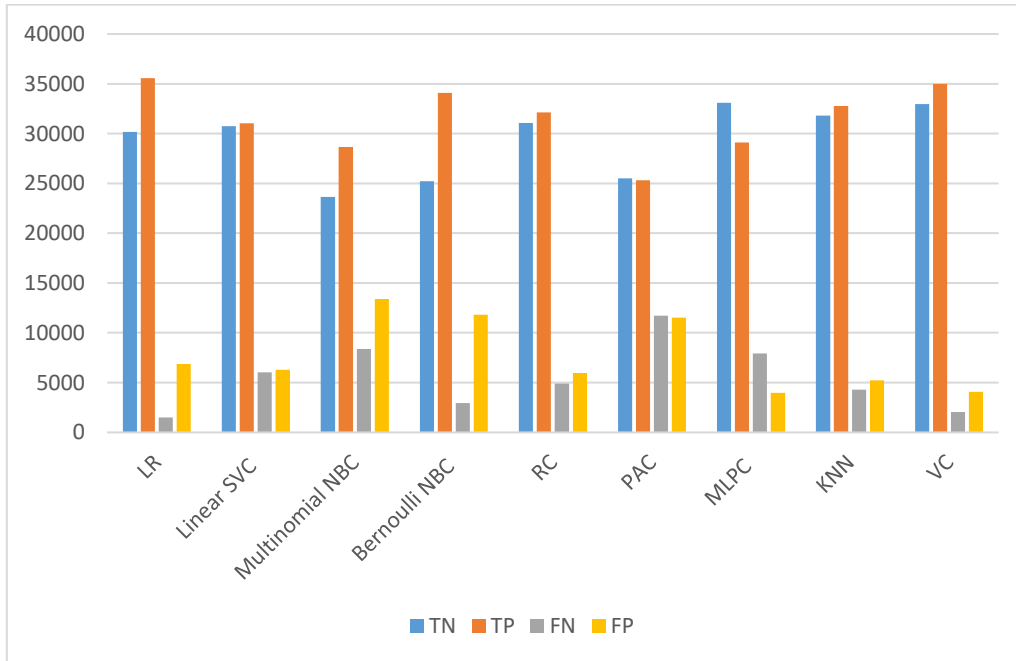
Çizelge 5.5. Mobil oyun İngilizce Tüm veri seti için hesaplanan metrik ölçütler

Vektör Tipi	Model Adı	Acc	TN	TP	FN	FP	P	R	F	AUC %
Count Vectorizer	LR	0.88	32093	33526	3516	4949	0.87137	0.90508	0.887906	88.57
	Linear SVC	0.82	31785	29521	7521	5257	0.84884	0.79696	0.822082	82.75
	Multinomial NBC	0.7	26810	25327	11715	10232	0.71225	0.683737	0.697703	70.38
	Bernoulli NBC	0.8	30651	28766	8276	6391	0.81822	0.776577	0.796853	80.2
	RC	0.85	28337	35146	1896	8705	0.80149	0.948814	0.86895	85.69
	PAC	0.67	23977	26210	10832	13065	0.66735	0.707575	0.686871	67.74
	MLPC	0.83	28658	33226	3816	8384	0.79851	0.896981	0.844886	83.53
	KNN	0.87	29163	35752	1290	7879	0.81942	0.965174	0.886343	87.62
	VC	0.91	31210	36447	595	5832	0.86206	0.983937	0.918974	91.32
Tfidf Vectorizer	LR	0.88	30180	35555	1487	6862	0.83823	0.959856	0.894926	88.73
	Linear SVC	0.83	30746	31025	6017	6296	0.8313	0.837562	0.83442	83.38
	Multinomial NBC	0.7	23658	28661	8381	13384	0.68167	0.773743	0.724796	70.62
	Bernoulli NBC	0.8	25228	34101	2941	11814	0.7427	0.920603	0.822136	80.08
	RC	0.85	31082	32131	4911	5960	0.84353	0.86742	0.855309	85.33
	PAC	0.68	25502	25312	11730	11540	0.68686	0.683332	0.685089	68.59
	MLPC	0.83	33080	29111	7931	3962	0.8802	0.785891	0.830378	83.95
	KNN	0.87	31815	32759	4283	5227	0.8624	0.884374	0.873247	87.16
	VC	0.91	32972	35001	2041	4070	0.89583	0.9449	0.919711	91.75

Mobil oyun İngilizce tüm veri setinde karmaşıklık matrislerinin hesaplanmasında kullanılan TN, TP, FN, FP değerlerinin CountVectorizer vektör tipine göre dağılımı Şekil 5.15'te, TF-IDF vektör tipine göre dağılımı ise Şekil 5.16'da sunulmuştur.



Şekil 5.15. Mobil oyun İngilizce tüm veri setinde CountVectorizer Vektör Tipine Göre TP, TN, FN, FP Dağılımı.



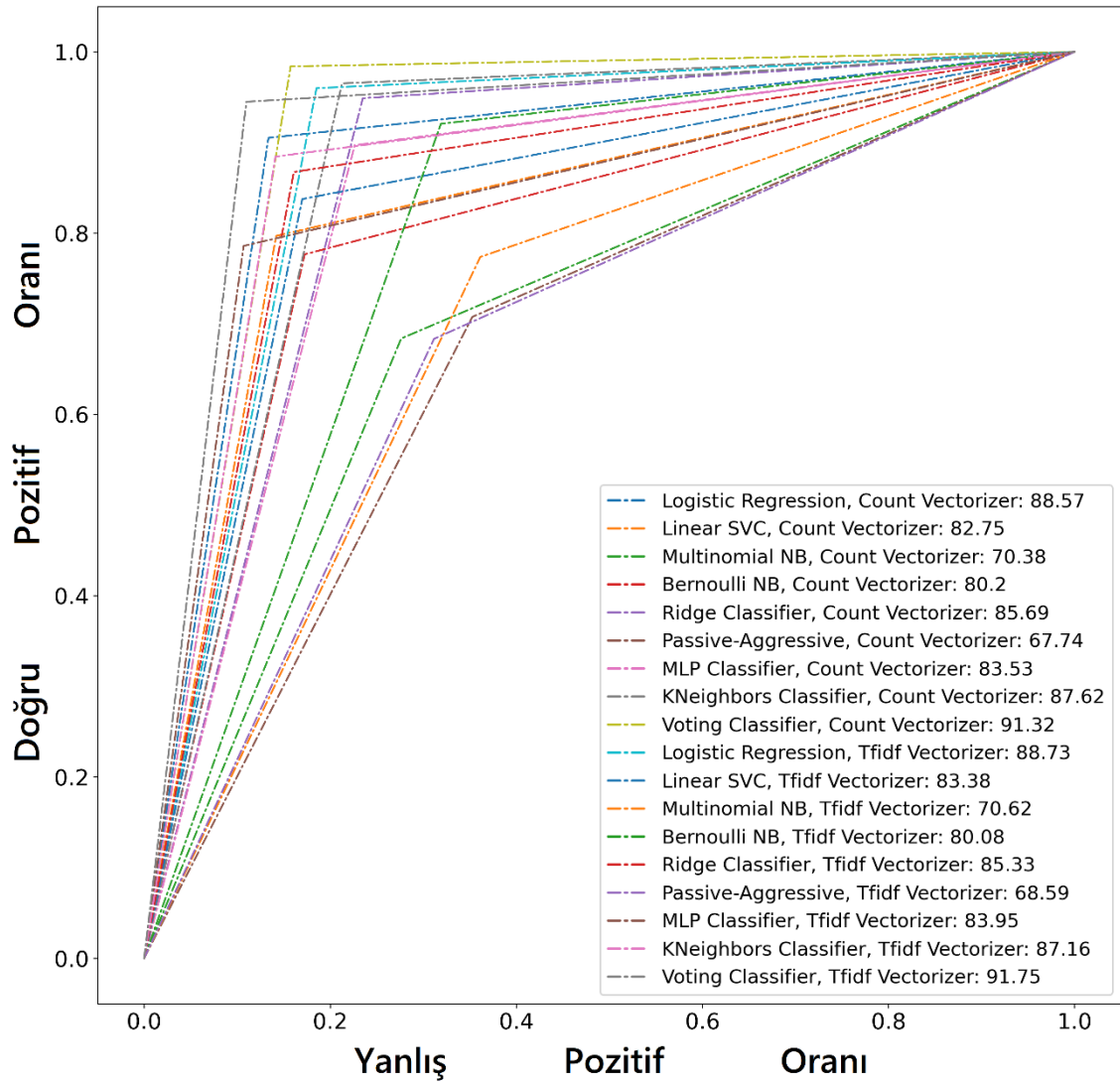
Şekil 5.16. Mobil oyun İngilizce tüm veri setinde TF-IDF Vektör Tipine Göre TP, TN, FN, FP Dağılımı.

Çizelge 5.5'deki metrik ölçütlerine göre;

Twitterdaki mobil oyun İngilizce tüm veri seti için hesaplanan sonuçların eğitim ve test veri setlerinin sonuçları ile tutarlılık gösterdiği tespit edilmiştir. AUC için VC modelinin CountVectorizer'da %91.32'lik, TF-IDF'de ise %91.7'lik oranlarla mükemmel değerlere ulaştığı tespit edilmiştir. R, P, F metrik değerler dikkate alındığında, PAC modeli için zayıf (%68), Multinomial NBC için orta (%70) değerlere ulaşılmıştır. R, P, F için VC modeli ile en iyi değerlere ulaşıldığı tespit edilmiştir. Acc için VC modelinde 0.91'e ulaşılarak en başarılı model olarak belirlenmiştir. Karşılaştırılan modellerin ROC grafiği Şekil 5.17'de sunulmuştur.

Şekil 5.15- 5.16'daki TP, TN, FN ve FP dağılımları dikkate alındığında, FN ve FP değerlerinin TP ve TN değerlerine göre oldukça düşük düzeyde olması oluşturulan modelin iyi çalıştığını göstermekle birlikte, polaritelerinin iyi düzenlendiğini de ispatlamaktadır. P, R, F, AUC ve Acc metrik ölçümler için TF-IDF'de daha iyi değerlere ulaşıldığı, fakat anlamlı bir farkın oluşmadığı tespit edilmiştir.

Sonuç olarak; eğitim veri setiyle eğitilen modeller tarafından test veri setinde başarılı sonuçlara ulaşıldığı ve tüm veri setinde ulaşılan oranla ile tutarlı olduğu tespit edilmiştir. Veri setinin pozitif ve negatif olarak eşit sayıda tweet içermesi ve polaritelerinin mümkün olduğunca doğruya yakın belirlenmiş olması başarılı sonuçlara ulaşmada önemli etkenlerdir. Hesaplanan metrik ölçümlerine göre önerilen VC ile mobil oyun İngilizce tüm veri setinde en iyi sonuçlara ulaştığı tespit edilmiştir.



Şekil 5.17. Mobil oyun İngilizce veri seti – tüm veri seti ROC grafiği.

Tüm modeller için Çizelge 5.3 ve Çizelge 5.4’de hesaplanan P, R ve F metrik ölçümlerinin test ve eğitim veri setine göre karşılaştırılması Şekil 5.18’de gösterilmiştir.



Şekil 5.18. Mobil oyun İngilizce test- eğitim veri seti için tüm algoritmaların P, R ve F değerlerinin karşılaştırılması.

5.2.4. TEMSAP-CNNLSTM modeli analizi

TEMSAP-CNNLSTM modeli için ROC grafiklerinde tek bir sonuç alınacağından, mobil oyun İngilizce test-eğitim-tüm veri seti ROC eğrileri aynı grafikte (Şekil 5.20) gösterilmiştir. Ayrıca, test-eğitim-tüm veri seti için hesaplanan TN, TP, FN, FP, P, R, F, AUC ve Acc metrik ölçütler de aynı çizelgede (Çizelge 5.6) sunulmuştur. Geliştirilen TEMSAP-CNNLSTM modelinde, VC ve diğer basit makine öğrenme algoritmalarına uygulanan veri seti esas alınarak pozitif ve negatif tweet sayıları eşit tutulmuştur.

Çizelge 5.6. Mobil oyun İngilizce veri seti için TEMSAP-CNNLSTM modeli tarafından hesaplanan metrik ölçütler

Veri Seti	Acc	TN	TP	FN	FP	P	R	F	AUC %
Test	0.93	6665	7212	247	693	0.984307	0.966885	0.93882	93.64
Eğitim	0.97	7041	23334	324	372	0.986304	0.987626	0.94715	96.81
Test + Eğitim (Tüm)	0.94	36507	33356	3686	535	0.984214	0.900491	0.94049	94.3

Çizelge 5.6’da hesaplanan metrik ölçütlere göre;

Test veri setindeki TP, TN, FN, FP değerlerinin ve eğitim veri setindeki TP, TN, FN, FP değerlerinin toplamının, tüm veri setindeki TP, TN, FN, FP değerlerinin toplamına eşit olmadığı tespit edilmiştir. TEMSAP-CNNLSTM’de eğitim veri seti için 74.084 tweetin %80’ni (59.267 Tweet), test veri seti içinse %20’si (14.817 Tweet) ayrılmıştır. Çizelge 5.6’da test veri setindeki TP, TN, FN, FP değerlerinin toplamı 14.817 iken, eğitim veri setindeki TP, TN, FN, FP değerlerinin toplamı 31.071 olarak tespit edilmiştir. Eğitim veri setinde TP, TN, FN, FP toplamı 59.267 adettir. Eksik olan toplam 28.196 TP, TN, FN, FP ile Validasyon veri seti (Validation Data Set) oluşturulmuştur. Validasyon veri seti, Eğitim veri seti içerisinde seçilmiş ve optimizasyon sürecinde kullanılmamıştır. Değerlerde mevcut olan eksik tweet sayıları bu durumdan kaynaklanmıştır.

Her iterasyon sonunda (epoch) Validasyon veri seti ve test veri seti ile modelin başarımı ayrı ayrı ölçülmüştür. Test veri seti öğrenme sürecinde adım adım modelin başarımını ölçen tamamen bağımsız bir test kümesi olduğu için öğrenme işleminin hiçbir basamağında kullanılmamıştır. Buna karşı Validasyon veri seti her bir epoch işleminde eğitim veri setinin farklı bir parçası olacak şekilde yenilenmekte ve modelin öğrenme sürecine doğrudan katkı sağlamaktadır. Bu işlemde amaç; TEMSAP-CNNLSTM modeli eğilirken test verisinden bağımsız olarak aşırı öğrenmenin önüne geçilebilmesidir. Validasyon veri setinin oluşturulmasında ve epoch işleminde kullanılan python kodu Şekil 5.19’da gösterilmiştir.

```
### Create sequence
vocabulary_size = 1000
tokenizer = Tokenizer(num_words=vocabulary_size)
tokenizer.fit_on_texts(son_df["kok_tweet"])
sequences = tokenizer.texts_to_sequences(son_df["kok_tweet"])
data = pad_sequences(sequences, maxlen=50)

x_train, x_test, y_train, y_test = train_test_split(data, son_df["Sentiment_Type"],
                                                    test_size = 0.2,
                                                    random_state = 18)

x_train, x_val, y_train, y_val = train_test_split(x_train, y_train,
                                                  test_size = 0.2,
                                                  random_state = 18)

def create_conv_model():
    model_conv = Sequential()
    model_conv.add(Embedding(vocabulary_size, 50, input_length=50))
    model_conv.add(Dropout(0.8))
    model_conv.add(Conv1D(32, 5, activation='relu'))
    model_conv.add(MaxPooling1D(pool_size=4))
    model_conv.add(LSTM(50))
    model_conv.add(Dense(1, activation='sigmoid'))
    model_conv.compile(loss='binary_crossentropy', optimizer='adam', metrics=['accuracy'])
    return model_conv

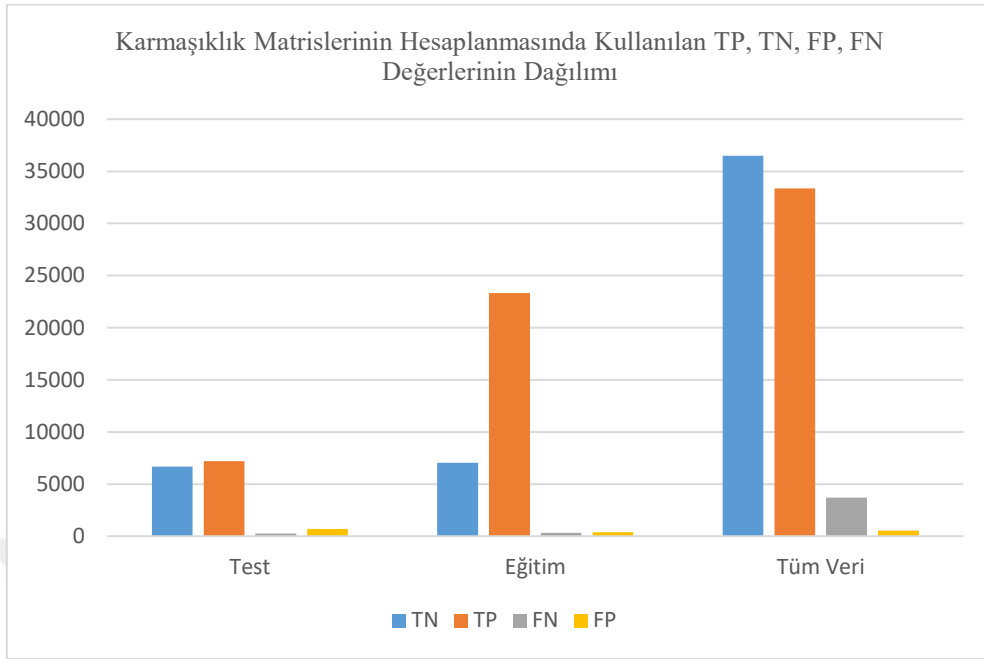
model_conv = create_conv_model()
model_conv.fit(x_train, np.array(y_train.astype("float32")), validation_data=(x_val, np.array(y_val.astype("float32"))), epochs = 3)
```

Şekil 5.19. Validasyon veri setinin oluşturulmasında ve iterasyonunda kullanılan python kodu

TEMSAP-CNNLSTM modeli için hesaplanan metrik ölçütleri esas alındığında;

- 1- Mobil oyun İngilizce test veri setinde Acc için hesaplanan 0.93 ile diğer tüm modellerden daha başarılı olmuştur. Diğer modellerde mobil oyun İngilizce veri seti için en yakın değerin 0.91 ile VC'ye ait olduğu belirlenmişti (Çizelge 5.4). P, F ve AUC metrik ölçütleri için TEMSAP-CNNLSTM modelinin diğer tüm modellerden daha başarılı olduğu, R için ise TEMSAP-CNNLSTM'nin VC'ye kıyasla birbirine yakın değerler aldığı tespit edilmiştir.
- 2- Mobil oyun İngilizce eğitim veri setinde Acc için 0.97 ile diğer tüm modellerden daha başarılı olmuştur. Diğer modellerde mobil İngilizce oyun veri seti için en yakın değerin 0.92 ile VC modeline ait olduğu belirlenmişti (Çizelge 5.3). P, R, F ve AUC metrik ölçütleri için kıyaslanan diğer tüm modellere göre TEMSAP-CNNLSTM modeli daha başarılı sonuçlara ulaşıldığı tespit edilmiştir.
- 3- Mobil oyun İngilizce tüm veri setinde Acc için 0.94 ile diğer tüm modellerden daha başarılı olmuştur. Diğer modellerde mobil oyun İngilizce veri setinde en yakın değerin 0.91 ile VC'a ait olduğu belirlenmişti (Çizelge 5.5). P, F ve AUC metrik ölçütleri için TEMSAP-CNNLSTM modelinin diğer tüm modellere kıyasla daha başarılı sonuçlara ulaştığı tespit edilmiştir. Fakat, R metriği için 0.98 ile VC modeli optimuma ulaşırken, TEMSAP-CNNLSTM modeli ise 0.90 değerine ulaşmıştır.

Mobil oyun İngilizce veri setinde TEMSAP-CNNLSTM modeli için, karmaşıklık matrislerinin hesaplanmasında kullanılan TN, TP, FN, FP metrik ölçütlerinin dağılımı Şekil 5.20'de gösterilmiştir.



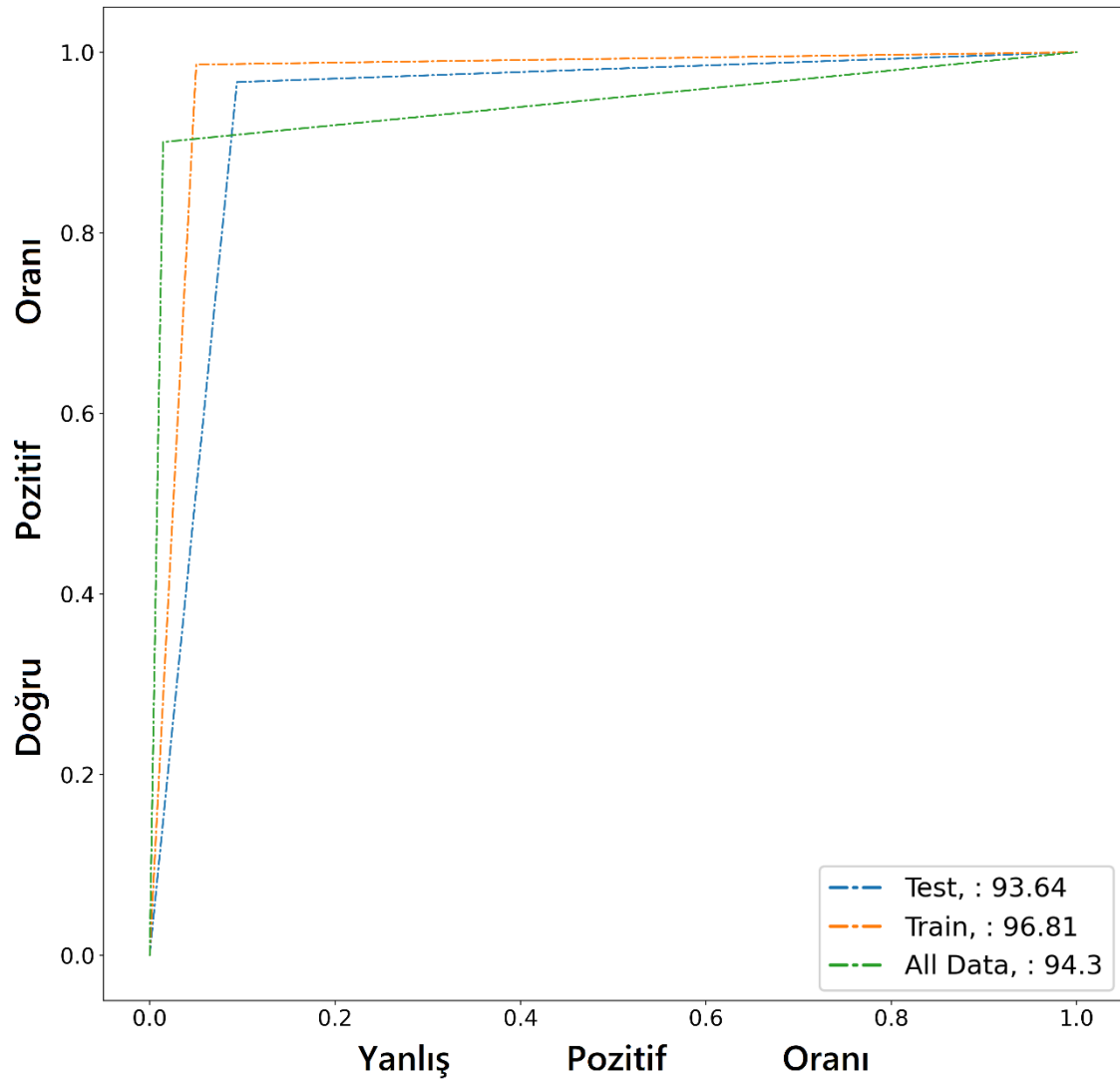
Şekil 5.20. Mobil oyun İngilizce veri setinde TEMSAP-CNNLSTM modeli için hesaplanan TP, TN, FN, FP dağılımı.

FN ve FP metriklerinin TP ve TN'ye göre oldukça düşük değerde olması TEMSAP-CNNLSTM modelinin başarısını kanıtlamakla birlikte, polaritelerinin iyi düzenlendiğini de ispatlamaktadır. Diğer tüm modeller ile kıyaslandığında TEMSAP-CNNLSTM modelinde;

- Mobil oyun İngilizce test veri seti için FP (%4.91) ve FN (%1.66),
- Mobil oyun İngilizce eğitim veri seti için FP (%1.19) ve FN (%1.04),
- Mobil oyun İngilizce tüm veri seti için FP (%0.72) ve FN (%4.97) ölçülen en düşük metrikler olmuştur.

TEMSAP-CNNLSTM modelinde, FN ve FP metrik ölçütlerine en yakın VC modeli ile ulaşılmıştır. TF-IDF ve CountVectorizer için diğer modellerin FN ve FP metrikleri dikkate alındığında, TEMSAP-CNNLSTM modelinin çok altında kaldığı tespit edilmiştir.

Mobil oyun İngilizce veri setinde test-eğitim-tüm veri setleri için, TEMSAP-CNNLSTM modeli ile ulaşılan ROC grafiği Şekil 5.21'de gösterilmiştir.



Şekil 5.21. Mobil oyun İngilizce test-eğitim-tüm veri seti – TEMSAP-CNNLSTM ROC grafiği.

5.3. Mobil Oyun Türkçe Veri Seti Analizleri

Türkiye’den 2015-2021 yılları arasında elde edilen 376.431 Tweet baz alınarak analizler yapılmıştır (Çizelge 3.3). Yapılan analizler için “LR”, “Linear SVC”, “Multinomial NBC”, “Bernoulli NBC”, “RC”, “PAC”, “MLPC”, “KNN”, “VC” ve geliştirilen “TEMSAP-CNNLSTM” modelleri kullanılmıştır. RF, LR, RF, RC algoritmaları birlikte kullanılarak VC algoritması oluşturulmuştur. Tüm modellerin karşılaştırılmasında, TF-IDF ve CountVectorizer vektör tipleri esas alınarak sırasıyla

Acc, TN, TP, FN, FP, P, R, F ve AUC metrik ölçütler hesaplanmış ve ROC eğrileri grafiksel olarak çizdirilmiştir. Model oluşturulmasında, polarize edilen 11.271 pozitif ve 11.271 negatif tweet olmak üzere toplam 22.542 tweet kullanılmıştır. Mobil oyun Türkçe veri seti için uygulanan polarite belirleme yönteminde kullanılan python kodu Şekil 5.22’de gösterilmiştir.

```
tr_df.loc[tr_df.Polarity >= 0.35, 'Sentiment_Type'] = 1 # Positive
tr_df.loc[tr_df.Polarity <= -0.35, 'Sentiment_Type'] = -1 # Negative
tr_df.loc[(tr_df.Polarity > -0.35) & (tr_df.Polarity < 0.35), 'Sentiment_Type'] = np.NaN # Neutral
current_df = tr_df.loc[tr_df['Sentiment_Type'].isnull()==0]

ilk11 = current_df[current_df["Sentiment_Type"]==1][:11271]
ikinci11 = current_df[current_df["Sentiment_Type"]==1]

son_df = pd.concat([ilk11, ikinci11])
son_df.shape

(22542, 3)

son_df.head()
```

Şekil 5.22. Mobil oyun Türkçe veri seti için polarite belirlemede kullanılan python kodu.

TEMSAP-CNNLSTM modeli hibrit bir model olarak derin öğrenme yöntemleri kullandığı için diğer tüm modeller ile aynı veri seti kullanılarak karşılaştırılmıştır. Dolayısıyla hesaplanan metrik ölçütlerin diğer modeller ile karşılaştırılmasında farklı ROC grafikleri kullanılmıştır.

Mobil oyun Türkçe veri setinin %80’i eğitim verisi (18.033 tweet) ve %20’si test verisi (4.509 tweet) olarak ayrılmıştır. Yapılan analizler mobil oyun Türkçe test, mobil oyun Türkçe eğitim ve mobil oyun Türkçe tüm veri olmak üzere üç farklı şekilde analiz edilmiştir.

5.3.1. Eğitim veri seti analizi

Mobil oyun Türkçe eğitim veri seti için her iki vektör tipi esas alınarak kullanılan modeller tarafından hesaplanan metrik ölçütler Çizelge 5.7’de gösterilmiştir.

Çizelge 5.7. Mobil oyun Türkçe eğitim veri seti için hesaplanan metrik ölçütler

Vektör Tipi	Model Adı	Acc	TN	TP	FN	FP	P	R	F	AUC %
CountVectorizer	LR	0.86	7122	8552	549	1810	0.82532	0.939676	0.878795	86.85
	Linear SVC	0.81	6473	8234	867	2459	0.77004	0.904735	0.831969	81.47
	Multinomial NBC	0.69	5374	7166	1935	3558	0.66822	0.787386	0.722925	69.45
	Bernoulli NBC	0.77	6327	7711	1390	2605	0.74748	0.847269	0.794252	77.78
	RC	0.81	6835	7869	1232	2097	0.78958	0.86463	0.825405	81.49
	PAC	0.65	5243	6655	2446	3689	0.64337	0.731238	0.684494	65.91
	MLPC	0.75	5898	7777	1324	3034	0.71936	0.854521	0.781137	75.74
	KNN	0.86	6921	8602	499	2011	0.81052	0.94517	0.872679	86
	VC	0.92	7694	9009	92	1238	0.87918	0.989891	0.931259	92.56
Tfidf Vectorizer	LR	0.87	6787	8985	116	2145	0.80728	0.987254	0.88824	87.36
	Linear SVC	0.82	6063	8834	267	2869	0.75485	0.970662	0.849259	82.47
	Multinomial NBC	0.7	5256	7522	1579	3676	0.67173	0.826502	0.74112	70.75
	Bernoulli NBC	0.78	6238	7897	1204	2694	0.74563	0.867706	0.802051	78.3
	RC	0.83	7028	8097	1004	1904	0.80962	0.889682	0.847764	83.83
	PAC	0.64	4676	7011	2090	4256	0.62226	0.770354	0.688432	64.69
	MLPC	0.76	5746	8000	1101	3186	0.71518	0.879024	0.788682	76.12
	KNN	0.88	7141	8811	290	1791	0.83107	0.968135	0.894381	88.38
	VC	0.93	7943	8999	102	989	0.90098	0.988792	0.942846	93.9

Çizelge 5.7'deki metrik ölçütlere göre;

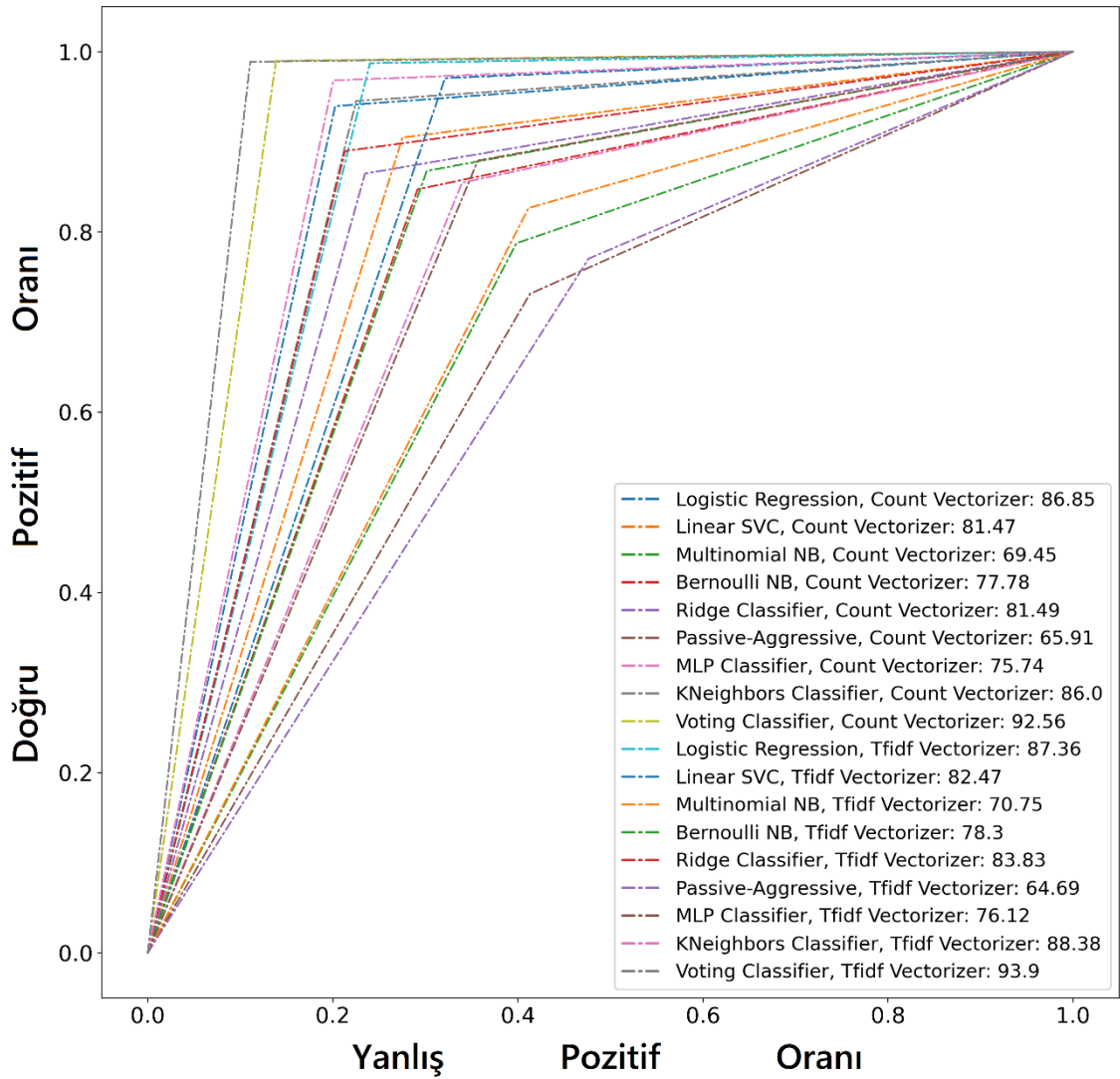
R için; CountVectorizer esas alınarak Multinomial NBC ile 0.78'lik, PAC ile 0.73'lük, Bernoulli NBC ile 0.84'lük, RC ile 0.86'lık, MLPC ile 0.85'lik, LR ile 0.93'lük, Linear SVC ile 0.90'lık, KNN ile 0.94'lük ve VC ile 0.98'lik metrik değerlere ulaşılmıştır. TF-IDF esas alındığında ise, sadece Multinomial NBC modeli 0.82'lik bir orana ulaşarak artış göstermiş olup, diğer modellerin oranlarında ise anlamlı bir değişim gözlenmemiştir. VC, her iki vektör tipine göre hesaplanan 0.98 ile en iyi metrik ölçüt değerine ulaşan model olmuştur. Dolayısıyla VC, pozitif tweetleri en doğru belirleyen model olarak belirlenmiştir.

P için; her iki vektör tipi esas alındığında, PAC ve Multinomial NBC modellerinin diğer modellere kıyasla metrik ölçütünde düşük değerlerin hesaplandığı tespit edilmiştir. VC, LR ve KNN modelleri başarılı olarak değerlendirilirken, VC modeli CountVectorizer için 0.87'lik, TF-IDF için ise %90'lık en yüksek metrik değerlere ulaşmıştır. Ayrıca her iki vektör tipi için Linear SVC, Bernoulli NBC, RC ve MLPC modellerinin LR ve KNN modellerine kıyasla daha düşük metrik değerlere ulaştığı tespit edilmiştir.

F için; her iki vektör tipi esas alındığında, VC ile CountVectorizer için 0.93'lük, TF-IDF için ise 0.94'lük metrik ölçüm değerlerine ulaşarak en iyi model olduğu tespit edilmiştir. PAC modeli ise, her iki vektör tipi için 0.68'lik düşük ölçüm metrik değerine ulaşmıştır. Diğer modellerde ise, bu metrik ölçütü açısından birbirine yakın değerler hesaplanmıştır.

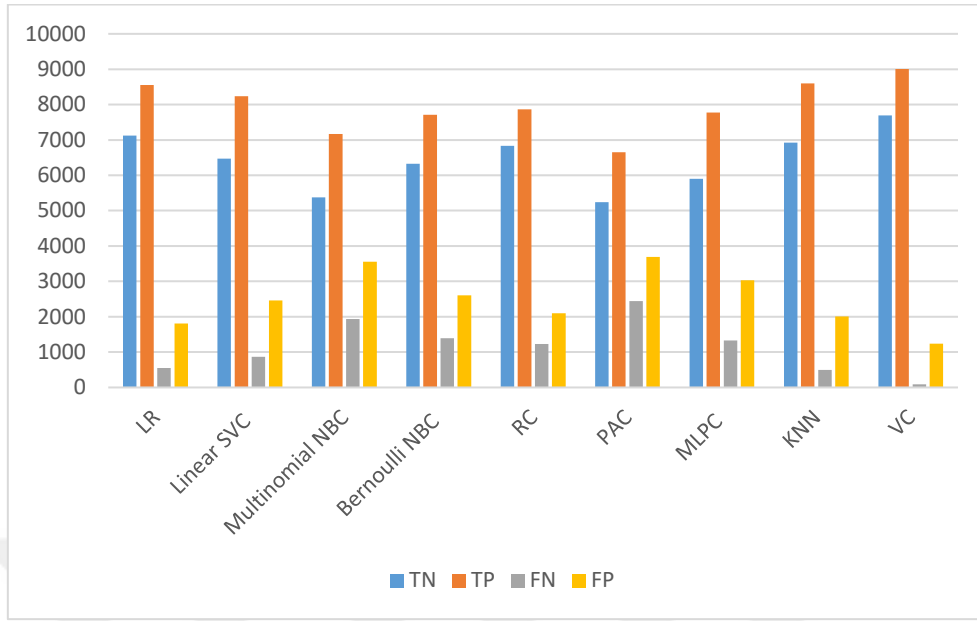
AUC için; VC modeli CountVectorizer için %92.56'lık, TF-IDF için ise %93.9'luk metrik ölçüt değerlerine ulaşarak mükemmel olarak derecelendirilmiştir. *R*, *P* ve *F* metrik ölçütleri dikkate alındığında, PAC ve Multinomial NBC modellerinin en düşük ölçüm metrik değerlerine sahip modeller oldukları tespit edilmiştir. Karşılaştırılan modellerin hiç birisinde düşük ölçüm metriği hesaplanmamıştır. CountVectorizer için 0.93'lük ve TF-IDF için ise 0.94'lük bir metrik ölçüt değerlerine ulaşan VC'nin diğer modellere kıyasla en başarılı model olduğu belirlenmiştir.

Karşılaştırılan modellerin ROC grafiği Şekil 5.23'de gösterilmiştir.

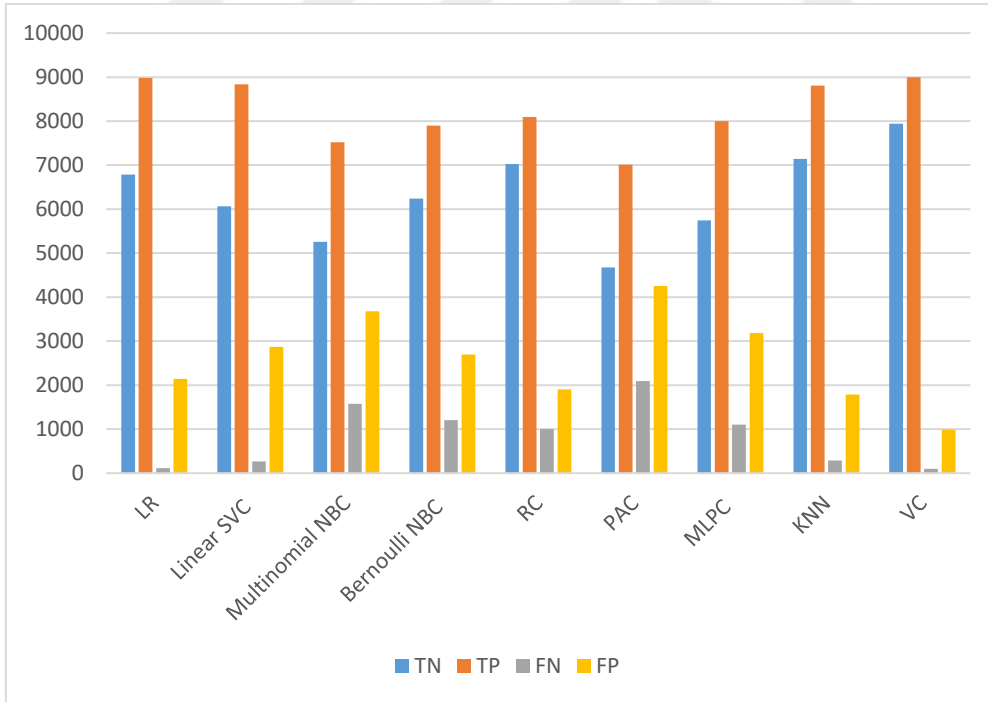


Şekil 5.23. Mobil oyun Türkçe veri seti – eğitim veri seti ROC grafiği.

Mobil oyun Türkçe eğitim veri seti için karmaşıklık matrislerinin hesaplanmasında kullanılan TN, TP, FN, FP değerlerinin CountVectorizer vektör tipine göre dağılımı Şekil 5.24'te, TF-IDF vektör tipine göre dağılımı ise Şekil 5.25'te gösterilmiştir.



Şekil 5.24. Mobil oyun Türkçe eğitim veri seti için CountVectorizer vektör tipine göre TP, TN, FN, FP dağılımı.



Şekil 5.25. Mobil oyun Türkçe eğitim veri seti için TF-IDF vektör tipine göre TP, TN, FN, FP dağılımı.

Şekil 5.24'te ve 5.25'te Mobil oyun Türkçe eğitim veri seti için hesaplanan sonuçlar incelendiğinde hem CountVectorizer hem de TF-IDF vektör tiplerine göre FN ve FP metrik ölçüt değerlerinin TP ve TN değerlerine kıyasla oldukça düşük olması oluşturulan modelin başarılı çalıştığını göstermekle birlikte, polaritelerinin iyi düzenlendiğini de ispatlamaktadır. TF-IDF için P, R, F, AUC ve Acc ölçüm metrik değerlerinde daha başarılı sonuçlara ulaşıldığı, fakat anlamlı bir farkın oluşmadığı tespit edilmiştir.

5.3.2. Test veri seti analizi

Mobil oyun Türkçe test veri seti için her iki vektör tipi esas alındığında, karşılaştırılan modeller tarafından hesaplanan metrik ölçütler Çizelge 5.8'de sunulmuştur.

Çizelge 5.8. Mobil oyun Türkçe test veri seti için hesaplanan metrik ölçütler

Vektör Tipi	Model Adı	Acc	TN	TP	FN	FP	P	R	F	AUC %
Count Vectorizer	LR	0.86	2132	1750	420	207	0.89423	0.806451	0.848073	85.9
	Linear SVC	0.79	1982	1606	564	357	0.81814	0.740092	0.777159	79.37
	Multinomial NBC	0.66	1588	1424	746	751	0.65471	0.656221	0.655466	66.76
	Bernoulli NBC	0.77	1732	1764	406	607	0.74399	0.812903	0.776921	77.67
	RC	0.8	1734	1907	263	605	0.75916	0.878801	0.814609	81.01
	PAC	0.63	1329	1537	633	1010	0.60346	0.708294	0.651685	63.82
	MLPC	0.75	1608	1777	393	731	0.70853	0.818894	0.759726	75.32
	KNN	0.85	1868	1974	196	471	0.80736	0.909677	0.855471	85.42
	VC	0.9	2038	2032	138	301	0.87098	0.936405	0.902509	90.39
Tfidf Vectorizer	LR	0.85	1783	2077	93	556	0.78883	0.957142	0.864876	85.97
	Linear SVC	0.78	1671	1888	282	668	0.73865	0.870046	0.798984	79.22
	Multinomial NBC	0.67	1524	1536	634	815	0.65334	0.707834	0.679495	67.97
	Bernoulli NBC	0.77	1558	1936	234	781	0.71255	0.892165	0.792306	77.91
	RC	0.81	1653	2001	169	686	0.74470	0.922119	0.823965	81.44
	PAC	0.63	1277	1598	572	1062	0.60075	0.736405	0.661697	64.12
	MLPC	0.73	1572	1741	429	767	0.69418	0.802304	0.744335	73.72
	KNN	0.84	1804	1993	177	535	0.78837	0.918433	0.848446	84.49
	VC	0.9	2087	2014	156	252	0.88879	0.92811	0.908025	91.02

Çizelge 5.8’de hesaplanan metrik ölçütlerine göre;

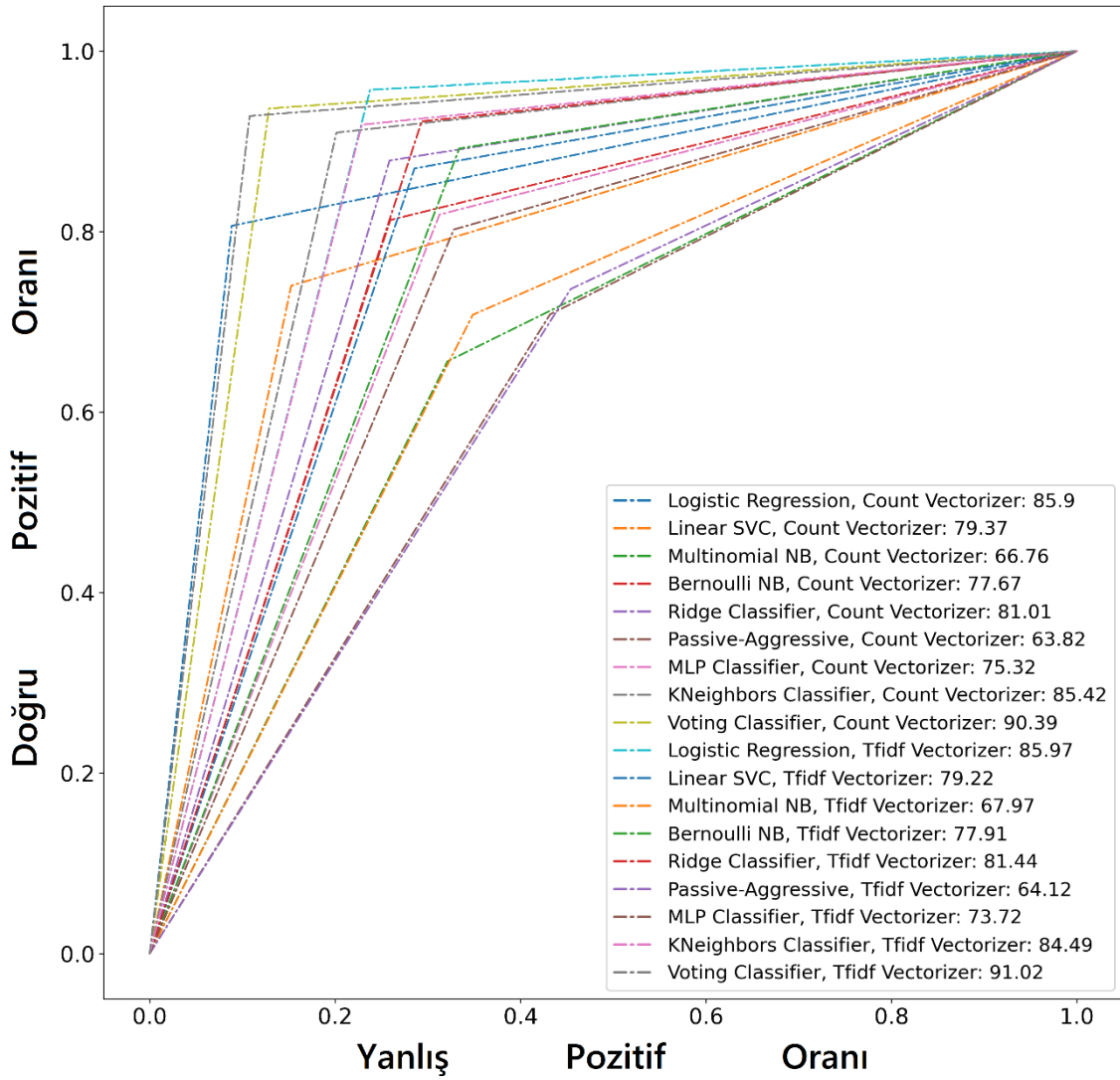
R için; CountVectorizer esas alındığında, Multinomial NBC için 0.65’lik ve PAC için 0.70’lik ölçüm metrik değerlerine ulaşıldığı tespit edilmiştir. Bernoulli NBC ile 0.81’lik, LR ile 0.80’lik, MLPC ile 0.81’lik, RC ile 0.87’lik, Linear SVC ile 0.74’lük, KNN ile 0.90’lık ve VC ile 0.93’lük ölçüm metrik değerlerine ulaşıldığı tespit edilmiştir. TF-IDF esas alındığında ise, LR ile 0.80’den 0.95’e, RC ile 0.87’ten 0.92’ye, Multinomial NBC ile 0.65’ten 0.70’e, Bernoulli NBC ile 0.81’den 0.89’a, Linear SVC ile 0.74’ten 0.87’ye artarak daha başarılı ölçüm metrik değerlerine ulaşılmıştır. Her iki vektör tipi esas alındığında, VC’nin en başarılı model olduğu tespit edilmiştir. Dolayısıyla VC modeli, pozitif tweetleri en başarılı belirleyen model olarak tespit edilmiştir.

P için; her iki vektör tipine göre PAC ve Multinomial NBC modellerinin diğer modellere kıyasla düşük metrik ölçütünün hesaplandığı tespit edilmiştir. VC, LR ve Linear SVC ile başarılı bir metrik ölçütüne ulaşılırken, VC ile CountVectorizer için 0.87’lik, TF-IDF için ise 0.88’lik en iyi ölçüm metriğine ulaşılmıştır. Her iki vektör tipi esas alındığında KNN, Bernoulli NBC, RC ve MLPC’nin LR ve Linear SVC’ye kıyasla daha düşük ölçüm metriğine ulaştığı tespit edilmiştir.

F için; her iki vektör tipi esas alındığında, VC ile 0.90’lık en iyi ölçüm metriğine ulaşılmıştır. PAC ve Multinomial NBC’de ise en düşük ölçüm metriği tespit edilmiştir. LR ve KNN ise 0.84 – 0.86 aralığında ölçüm metriğine ulaşarak VC’ye en yakın başarı ölçütü tespit edilmiştir.

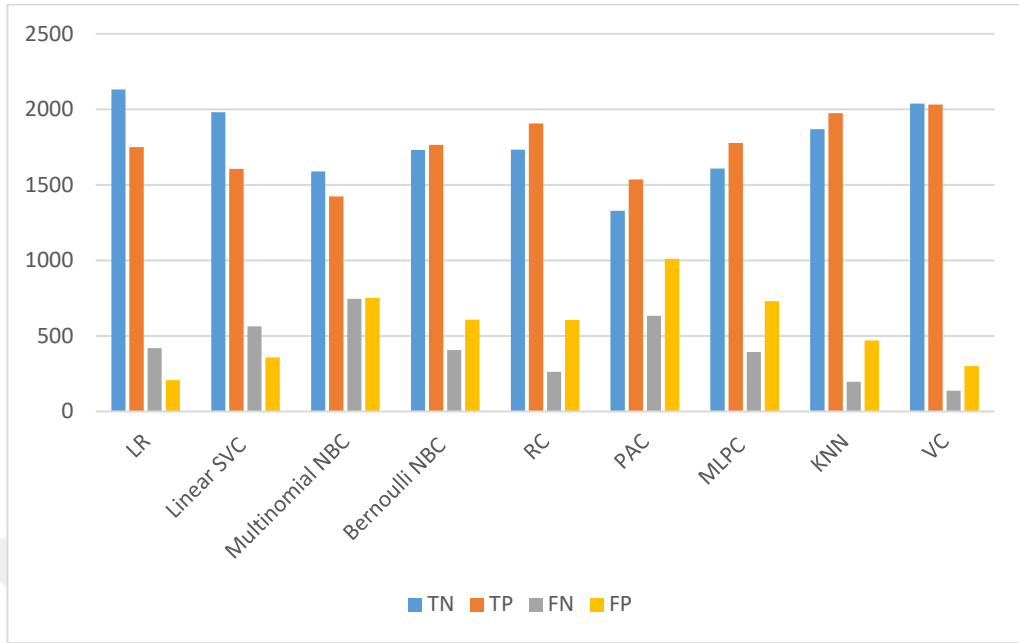
AUC için; VC ile CountVectorizer için %90.39’luk, TF-IDF için ise %91.02’lik metrik ölçümleri hesaplanmıştır ve mükemmel olarak derecelendirilmiştir. LR, RC ve KNN ise %80’nin üzerinde ölçüm metriğine ulaşarak başarı olarak derecelendirilmişlerdir. R, P, F metrik ölçütleri dikkate alındığında, PAC ve Multinomial NBC ile en düşük ölçüm metriği hesaplanmıştır ve böylece zayıf olarak derecelendirilmişlerdir. Karşılaştırılan modellerin hiçbirisi bu metrik ölçütü için başarısız olarak derecelendirilmemiştir. Dolayısıyla her iki vektör tipi esas alındığında 0.90’lık ölçüm metriği ile VC’nin en başarılı model olduğu tespit edilmiştir.

Mobil oyun Türkçe test veri seti kullanılarak karşılaştırılan tüm modeller için ROC grafiği Şekil 5.26’da gösterilmiştir.

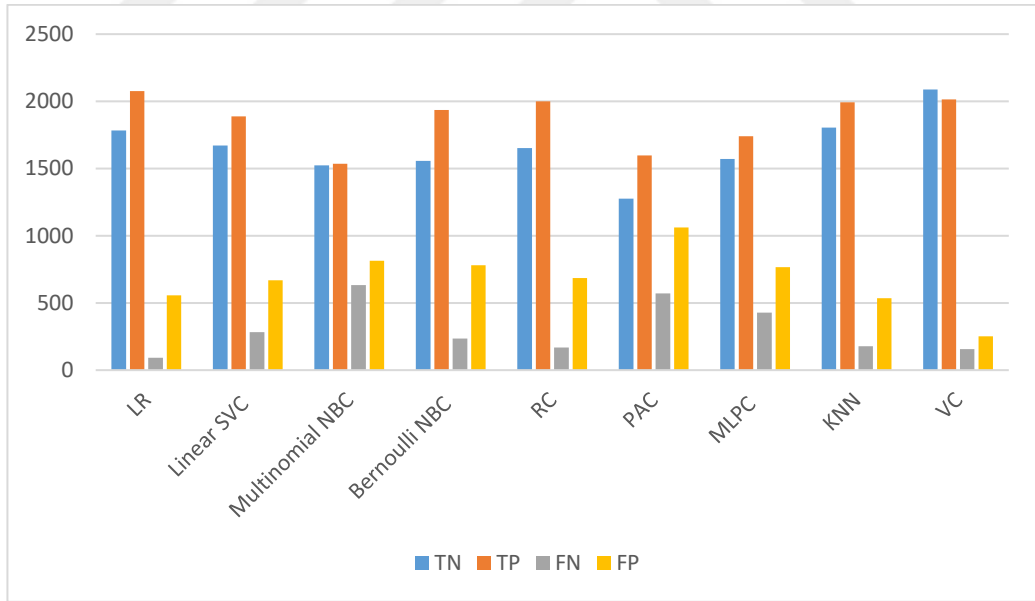


Şekil 5.26. Mobil oyun Türkçe veri seti – test veri seti ROC grafiği

Mobil oyun Türkçe test veri setinde karmaşıklık matrislerinin hesaplanması için kullanılan TN, TP, FN, FP değerlerinin CountVectorizer'a göre dağılımı Şekil 5.27'da TF-IDF'ye göre dağılımı ise Şekil 5.28'de sunulmuştur.



Şekil 5.27. Mobil oyun Türkçe test veri setinde CountVectorizer vektör tipine göre TP, TN, FN, FP dağılımı.



Şekil 5.28. Mobil oyun Türkçe test veri setinde TF-IDF vektör tipine göre TP, TN, FN, FP Dağılımı.

Şekil 5.27 ve 5.28'e göre hem CountVectorizer hem de TF-IDF esas alındığında FN ve FP metrik ölçütlerinin TP ve TN'ye kıyasla oldukça düşük olması oluşturulan

modelin başarılı olduğunu göstermekle birlikte, polaritelerinin de doğru bir şekilde düzenlendiğini ispatlamaktadır. TF-IDF’de hesaplanan P, R, F, AUC ve Acc için daha iyi sonuçlara ulaşıldığı, fakat anlamlı bir farkın oluşmadığı tespit edilmiştir. R için anlamlı bir farkın oluştuğu tespit edilmiş olup tek başına bu değer ile TF-IDF’nin CountVectorizer’den daha başarılı olduğu sonucuna ulaşamayacağı düşünülmektedir.

5.3.3. Tüm veri seti analizi

Mobil oyun Türkçe test veri seti için her iki vektör tipi esas alınarak kullanılan modeller tarafından hesaplanan metrik ölçütler Çizelge 5.9’da gösterilmiştir.



Çizelge 5.9. Mobil oyun Türkçe tüm veri seti için hesaplanan metrik ölçütler

Vektör Tipi	Model Adı	Acc	TN	TP	FN	FP	P	R	F	AUC %
CountVectorizer	LR	0.86	9265	10247	1024	2006	0.83628	0.909147	0.871195	86.56
	Linear SVC	0.8	8481	9768	1503	2790	0.77783	0.866648	0.819841	80.96
	Multinomial NBC	0.67	6588	8585	2686	4683	0.64705	0.761689	0.699702	67.31
	Bernoulli NBC	0.77	8276	9123	2148	2995	0.75285	0.809422	0.78011	77.18
	RC	0.82	8276	10253	1018	2995	0.77393	0.909679	0.836331	82.2
	PAC	0.64	6256	8288	2983	5015	0.62302	0.735338	0.674534	64.52
	MLPC	0.75	7603	9472	1799	3668	0.72085	0.840386	0.776043	75.75
	KNN	0.84	8883	10100	1171	2388	0.80878	0.896105	0.850204	84.21
	VC	0.91	9962	10747	524	1309	0.89142	0.953509	0.921421	91.87
Tfidf Vectorizer	LR	0.86	9506	10001	1270	1765	0.84999	0.887321	0.868255	86.54
	Linear SVC	0.81	8158	10121	1150	3113	0.76477	0.897968	0.826035	81.09
	Multinomial NBC	0.67	6508	8724	2547	4763	0.64685	0.774021	0.704741	67.57
	Bernoulli NBC	0.77	7625	9783	1488	3646	0.7285	0.867979	0.792145	77.22
	RC	0.82	7992	10565	706	3279	0.76315	0.937361	0.841329	82.32
	PAC	0.64	6359	8110	3161	4912	0.62279	0.719545	0.667682	64.19
	MLPC	0.76	7615	9578	1693	3656	0.72374	0.849791	0.781718	76.27
	KNN	0.86	8470	10944	327	2801	0.79622	0.97098	0.87496	86.12
	VC	0.93	9915	11124	147	1356	0.89135	0.986957	0.936718	93.33

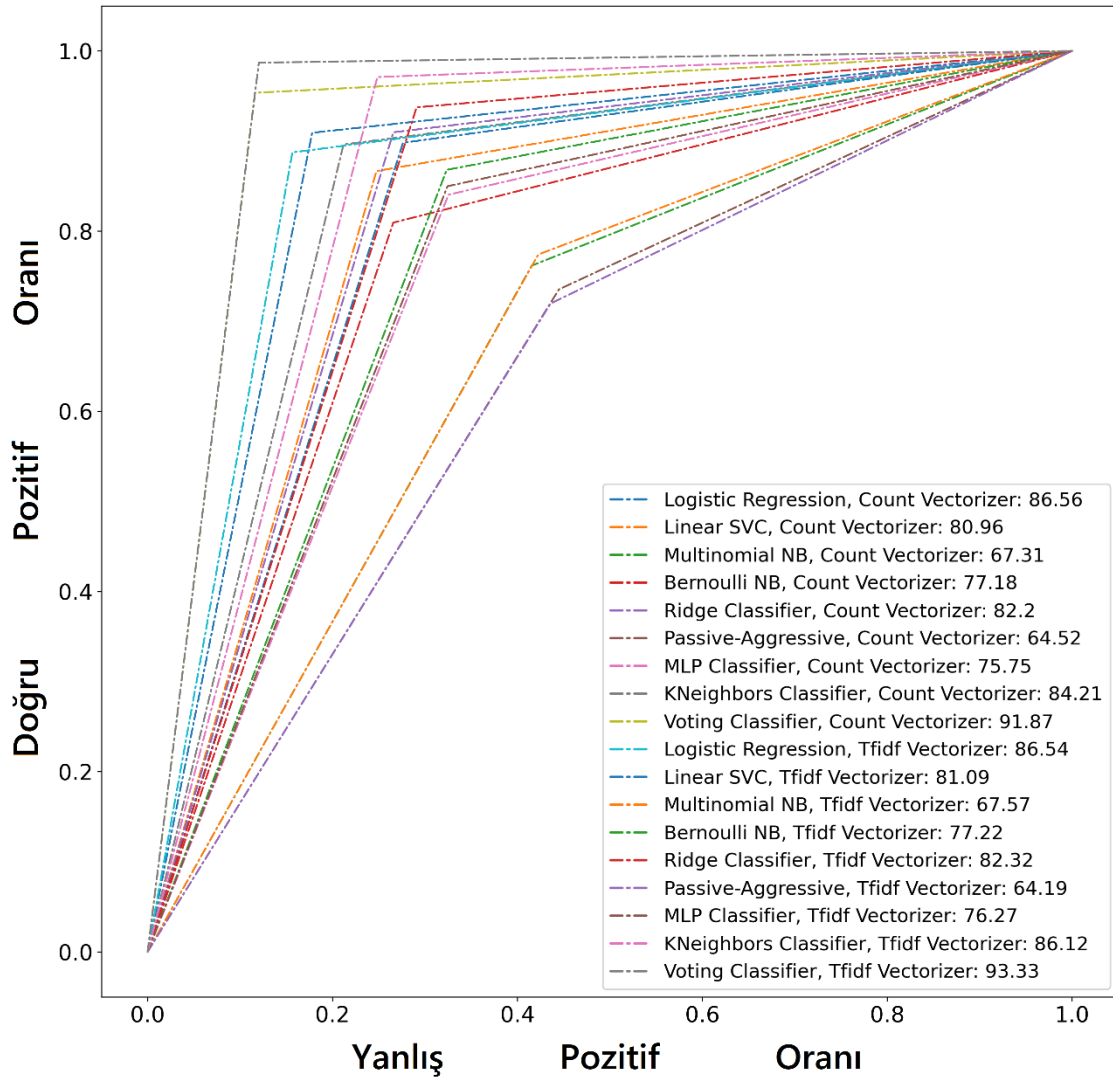
Çizelge 5.9’da hesaplanan metriklere göre;

R için; CountVectorizer’da Multinomial NBC ile 0.76’lık, PAC ile 0.73’lük, Bernoulli NBC ile 0.80’lik, LR ile 0.90’lık, MLPC ile 0.84’lük, RC ile 0.90’lık, Linear SVC ile 0.86’lık, KNN ile 0.89’luk ve VC ile 0.95’lik ölçüm metriğine ulaşıldığı tespit edilmiştir. TF-IDF’nin CountVectorizer’a kıyasla daha başarı değerlere ulaşılmıştır. KNN ve VC sırasıyla 0.95 ve 0.98 ölçüm metrik değerleri ile en yüksek başarı oranlarının hesaplandığı modeller oldukları tespit edilmiştir. Ayrıca VC’nin her iki vektör tipine göre en başarılı model olduğu tespit edilmiştir. Dolayısıyla VC modeli için pozitif tweetleri en doğru belirleyen model olduğu belirlenmiştir.

P için; her iki vektör tipi esas alındığında PAC ile 0.62 ve Multinomial NBC ile 0.64 ölçüm metriği hesaplanmıştır ve böylece diğer modellere kıyasla daha düşük metrik ölçütü tespit edilmiştir. VC, LR ve Linear SVC modelleri başarılı olarak değerlendirilirken, VC modeli her iki vektör tipine göre 0.89’luk ölçüm metriği ile en iyi değerlere ulaşmıştır. VC’ye en yakın ölçüm metriği LR ile ulaşarak en iyi ikinci model olduğu tespit edilmiştir. KNN, MLPC, RC ve Linear SVC’nin 0,72 ile 0.80 aralığındaki yakın oranlarda ölçüm metriğine ulaşılmıştır.

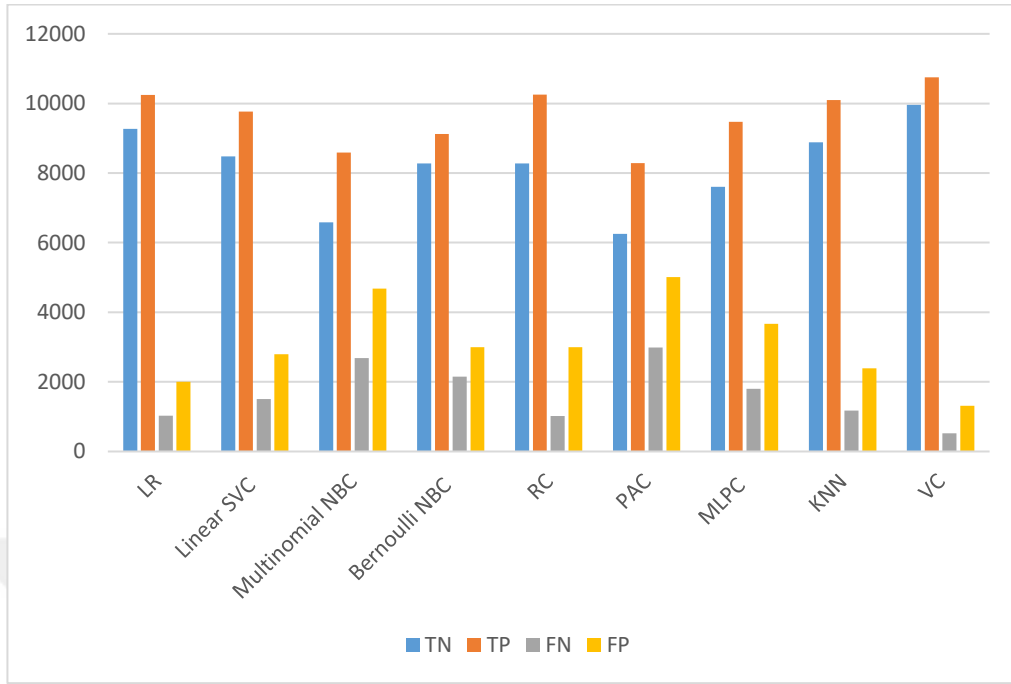
F için; VC’de CountVectorizer için 0.92’lik, TF-IDF için 0.93’lük ölçüm metriğine ulaşarak en başarılı model olduğu tespit edilmiştir. PAC ile Multinomial NBC’de her iki vektör tipine göre en düşük ölçüm metriğine ulaşılmıştır. KNN’de CountVectorizer için 0.85’lik, TF-IDF için 0.87’lik ölçüm metriğine ulaşarak VC’dan sonraki en yüksek başarı tespit edilmiştir.

AUC için; VC’de CountVectorizer için %91.87’lik, TF-IDF için %93.33’lük ölçüm metriğine ulaşarak mükemmel olarak derecelendirilmiştir. LR, Linear SVC ve KNN modellerinde %80’nin üzerinde ölçüm metriğine ulaşarak başarılı olarak derecelendirilmişlerdir. *R*, *P*, *F* metrik ölçümleri dikkate alındığında PAC ve Multinomial NBC modelleri her iki vektör tipinde sırasıyla %64’lük ve %67’lik ölçüm metrikleri hesaplanmıştır ve böylece en zayıf derecelendirmeye sahip oldukları tespit edilmiştir. Karşılaştırılan modellerin hiçbirisi *AUC* ölçüm metriğinde başarısız olarak kabul edilen %60’ın altında bir metrik ölçümünün hesaplanmadığı tespit edilmiştir. Doğruluk (*Acc*) metrik ölçütü CountVectorizer için 0.91 ve TF-IDF için 0.93 olarak hesaplanan VC’nin en başarılı olan model olduğu belirlenmiştir. Karşılaştırılan modellerin ROC grafiği Şekil 5.29’da gösterilmiştir.

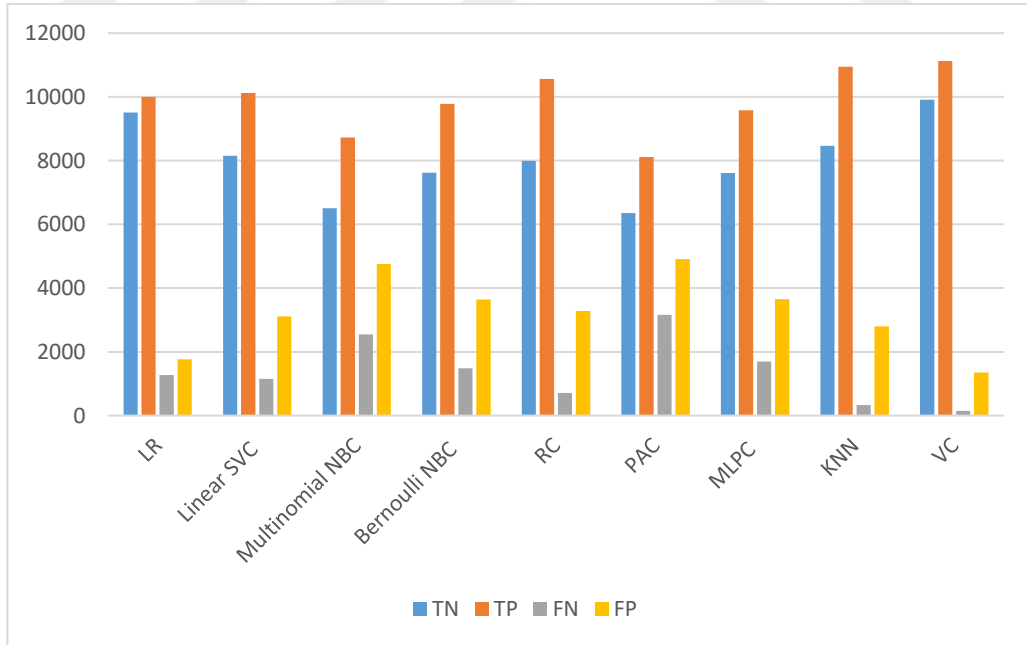


Şekil 5.29. Mobil oyun Türkçe veri seti –tüm veri seti ROC grafiği.

Mobil oyun Türkçe tüm veri setinde karmaşıklık matrislerinin hesaplanmasında kullanılan TN, TP, FN, FP metrik ölçütlerinin CountVectorizer için dağılımı Şekil 5.30’da TF-IDF için dağılımı ise Şekil 5.31’de gösterilmiştir.



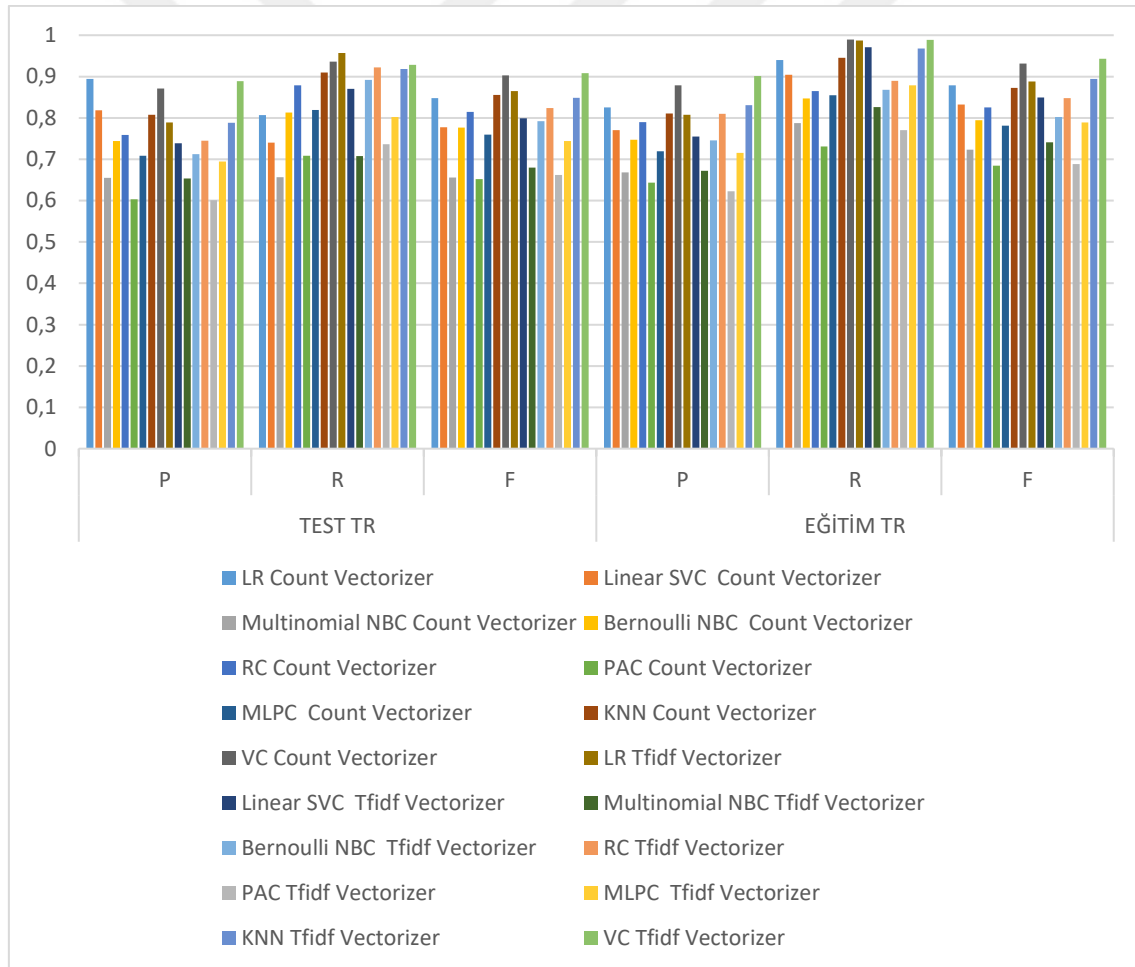
Şekil 5.30. Mobil oyun Türkçe tüm veri setinde CountVectorizer vektör tipine göre TP, TN, FN, FP dağılımı.



Şekil 5.31. Mobil oyun Türkçe tüm veri setinde TF-IDF vektör tipine göre TP, TN, FN, FP dağılımı.

Şekil 5.30 ve 5.31'deki sonuçlar incelendiğinde hem CountVectorizer hem de TF-IDF için FN ve FP metrik ölçütlerinin TP ve TN'ye göre oldukça düşük olması oluşturulan modelin başarılı olduğunu göstermekle birlikte, polaritelerinde iyi düzenlendiğini ispatlamaktadır. TF-IDF için hesaplanan P, R, F, AUC ve Acc metrik ölçütleri dikkate alındığında daha iyi sonuçlar elde edildiği, fakat anlamlı bir farkın oluşmadığı tespit edilmiştir. VC'de en düşük ölçüt metriği olarak FN, PAC'da en yüksek ölçüt metriği FN olarak tespit edilmiştir.

Çizelge 5.7 ve Çizelge 5.8'de sunulan Acc ölçütleri esas alındığında tüm modeller için hesaplanan P, R ve F metrik ölçütlerinin test-eğitim veri seti için karşılaştırılması Şekil 5.32'de gösterilmiştir.



Şekil 5.32. Mobil oyun Türkçe test-eğitim veri seti – P, R ve F metrik ölçütlerini karşılaştırılması.

5.3.4. TEMSAP-CNNLSTM modelinin analizi

Geliştirilen “TEMSAP-CNNLSTM” modelinin ROC grafiklerinde tek bir sonuç alınacağından, mobil oyun Türkçe test-eğitim-tüm veri seti ROC eğrileri tek bir grafik altında sunulmuştur. Ayrıca TN, TP, FN, FP, P, R, F, AUC ve Acc ölçüm metrikleri için hesaplanan sonuçlar eğitim, test ve tüm veri setleri esas alınarak Çizelge 5.10’da sunulmuştur. VC ve diğer makine öğrenme algoritmalarına uygulanan veri seti TEMSAP-CNNLSTM için de uygulanmış olup pozitif ve negatif tweet sayıları eşit sayıda tanımlanmıştır.



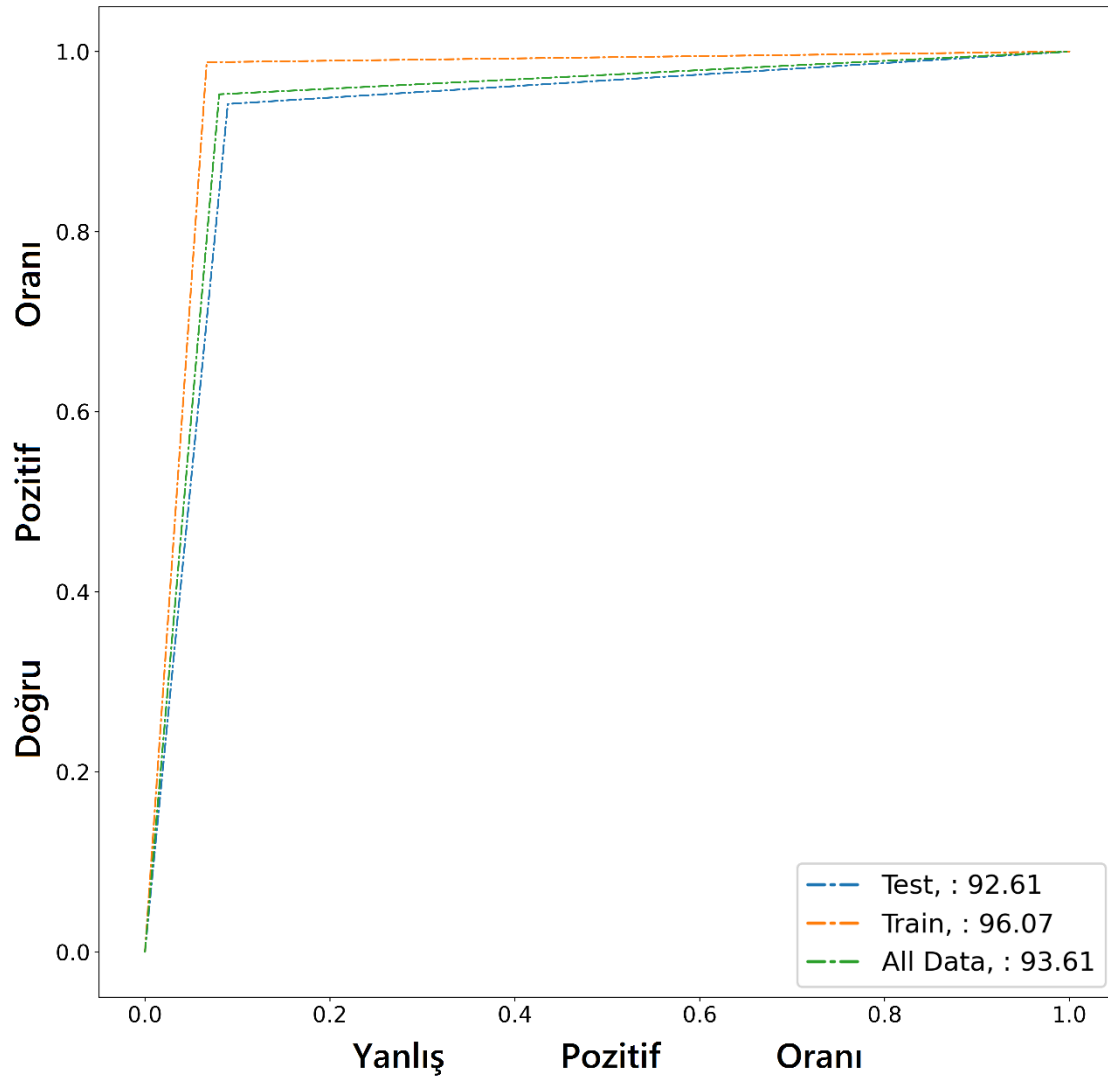
Çizelge 5.10. Mobil oyun İngilizce veri seti için TEMSAP-CNNLSTM modelinde hesaplanan metrik ölçütler

Veri Seti	Acc	TN	TP	FN	FP	P	R	F	AUC %
Test	0.92	2072	2103	130	204	0.911573	0.941782	0.92643	92.61
Eğitim	0.96	6712	7148	86	480	0.937073	0.988111	0.96192	96.07
Test + Eğitim (Tüm)	0.93	10366	10736	535	905	0.922257	0.952533	0.93715	93.61

Çizelge 5.10’da hesaplanan metrik ölçütlerine göre;

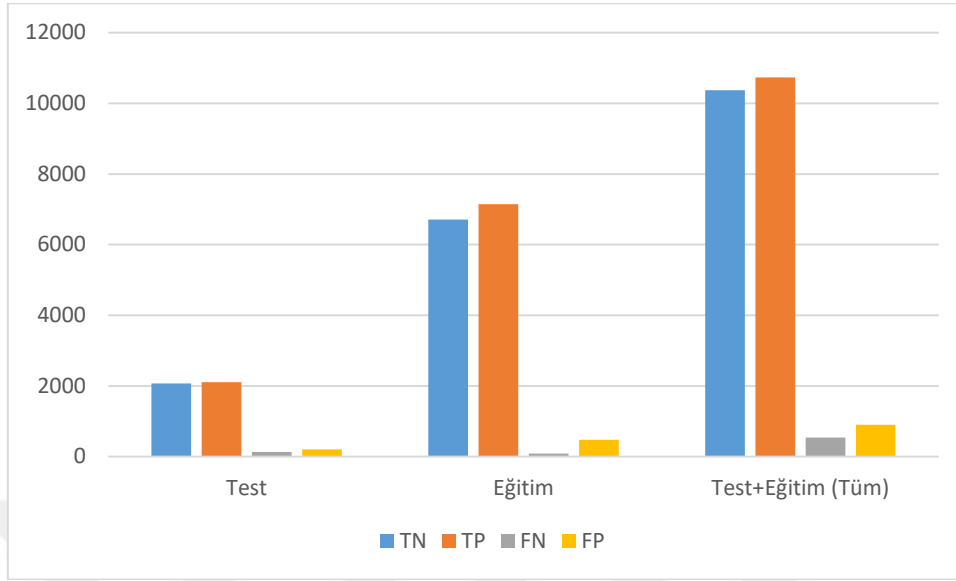
- TEMSAP-CNNLSTM tarafından mobil oyun Türkçe test veri setinde Acc için hesaplanan 0.92’lik ölçüm metriği ile diğer tüm modellerden daha iyi bir başarıma ulaşılmıştır. Diğer modellere kıyasla mobil oyun Türkçe veri setinde TEMSAP-CNNLSTM ile ulaşılan Acc metriğine en yakın ölçümün 0.90 ile VC tarafından ulaşıldığı tespit edilmiştir (Çizelge 5.8). Ayrıca P, R, F ve AUC metrikleri için TEMSAP-CNNLSTM’nin diğer tüm modellere kıyasla en iyi başarımlar oranlarına ulaştığı tespit edilmiştir.
- TEMSAP-CNNLSTM tarafından mobil oyun Türkçe eğitim veri setinde Acc için hesaplanan 0.96’lik ölçüm metriği ile diğer tüm modellerden daha iyi bir başarıma ulaşılmıştır. Diğer modellere kıyasla mobil oyun Türkçe veri seti için TEMSAP-CNNLSTM ile ulaşılan Acc metriğine en yakın ölçümün 0.93 ile VC tarafından ulaşıldığı tespit edilmiştir (Çizelge 5.7). Ayrıca P, R, F ve AUC ölçüm metrikleri için TEMSAP-CNNLSTM’nin diğer tüm modellere kıyasla en iyi başarımlar oranlarına ulaştığı tespit edilmiştir.
- TEMSAP-CNNLSTM tarafından mobil oyun Türkçe tüm veri setinde Acc için hesaplanan 0.93’lik ölçüm metriği ile diğer tüm modellerden daha iyi bir başarıma ulaşılmıştır. Diğer modellere kıyasla mobil oyun Türkçe veri setinde TEMSAP-CNNLSTM ile ulaşılan Acc metriğine en yakın ölçümün TF-IDF’de 0.91 ve Countvectorizer’da 0.93 ile VC tarafından ulaşıldığı tespit edilmiştir (Çizelge 5.5). Ayrıca P, F ve AUC ölçüm metrikleri için TEMSAP-CNNLSTM’nin diğer tüm modellere kıyasla en iyi başarımlar oranlarına ulaştığı tespit edilmiştir. Fakat R için 0.98 ile VC tarafından en iyi ölçüm metriğine ulaşılırken, TEMSAP-CNNLSTM tarafından 0.95’lik bir ölçüm metriğine ulaşılmıştır.

Mobil oyun Türkçe veri setindeki test-eğitim-tüm veri seti için TEMSAP-CNNLSTM için çizdirilen ROC grafiği Şekil 5.33’de gösterilmiştir.



Şekil 5.33. Mobil oyun Türkçe veri seti – TEMSAP-CNNLSTM test-eğitim-tüm veri seti ROC grafiği.

Mobil oyun Türkçe veri setinde TEMSAP-CNNLSTM için karmaşıklık matrislerinin hesaplanmasında kullanılan TN, TP, FN ve FP metrik ölçütlerinin dağılımı Şekil 5.34'te gösterilmiştir.



Şekil 5.34. TEMSAP-CNNLSTM modeli için mobil oyun Türkçe veri setinde hesaplanan TP, TN, FN ve FP dağılımı.

FN ve FP ölçüm metriklerinin TP ve TN'ye kıyasla oldukça düşük olması geliştirilen TEMSAP-CNNLSTM'nin başarılı olduğunu göstermekle birlikte, polaritelerinin iyi düzenlendiğini de ispatlamaktadır. Diğer tüm modeller ile kıyaslandığında;

- TEMSAP-CNNLSTM'nin mobil oyun Türkçe test veri setinde en düşük ölçüm metriği olarak FP (%4.52) ve FN (%2.88) için hesaplandığı tespit edilmiştir.
- TEMSAP-CNNLSTM'nin mobil oyun Türkçe eğitim veri setinde en düşük ölçüm metriği olarak FP (%3.32) ve FN (%0.59) için hesaplandığı tespit edilmiştir.
- TEMSAP-CNNLSTM'nin mobil oyun Türkçe tüm veri setinde en düşük ölçüm metriği olarak FP (%4.01) ve FN (%2.37) için hesaplandığı tespit edilmiştir.

Dolayısıyla TEMSAP-CNNLSTM için hesaplanan FN ve FP metrik ölçütlerine en yakın VC ile ulaşıldığı, TF-IDF ve CountVectorizer vektör tipleri esas alındığında VC ile hesaplanan FN ve FP ölçüm metriklerinin TEMSAP-CNNLSTM'den düşük olduğu belirlenmiştir.

5.4. Ülke ve Yıl Tabanlı Analiz

Bu tez çalışmasında ulaşılan veriler esas alınarak ülkeler, yıllara göre (2015-2016-2017-2018-2019-2020-2021) pozitif ve negatif tweet sayılarıyla kıyaslanmışlardır. 2021 yılına ait veriler mart ayına kadar alındığından dolayı bazı ülkelere ait veriler diğer yıllara göre daha az sayıda kalmıştır. Bu analiz, hangi yılda, hangi ülkelerin mobil oyunlar hakkında olumlu ya da olumsuz bir fikre sahip olduğu konusunda fikir verebilir. Bu analiz raporuna nötr tweetlerde dahil edilmiş olup daha iyi yorum yapabilme imkânı sunulması amaçlanmıştır.

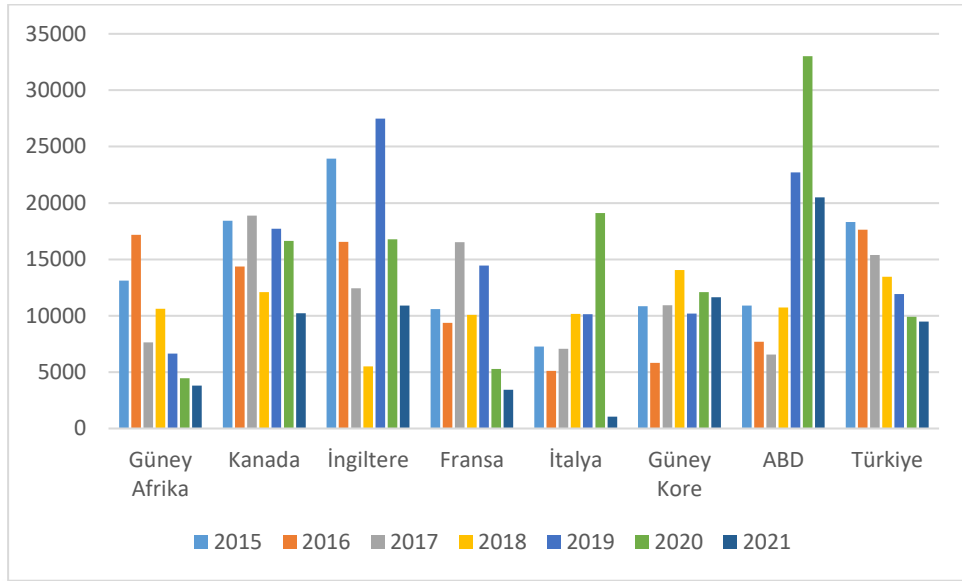
Güney Afrika, Kanada, İngiltere, Fransa, İtalya, Güney Kore ve Amerika Birleşik Devletlerinden mobil oyun ile alakalı İngilizce tweetler çekilmiş olup, Türkiye’den sadece Türkçe tweetler çekilmiştir. Çizelge 5.11 ve çizelge 5.12’de 2015 ve 2021 yılları arasında ülkelere göre tweetlerin olumlu-nötr-olumsuz olma durumları göz önünde bulundurularak tweetler sunulmuştur. İngilizce olarak 7 yabancı ülkelere çekilen tweetler göz önünde bulundurulduğunda 2015-2021 yılları arasında 602.626 olumsuz, 532.171 Nötr ve 704.477 adet olumlu olmak üzere toplamda 1.839.274 adet tweet elde edilmiştir. Türkçe olarak Türkiye’den çekilen tweetler göz önünde bulundurulduğunda 2015-2021 yılları arasında 96.107 olumsuz, 116.053 nötr ve 164.271 olumlu olmak üzere toplamda 376.431 adet tweet elde edilmiştir. İngilizce olarak elde edilen tweetlerin %32.76’sı olumsuz, %28.93’ü nötr ve %39.30’u olumlu olarak ayrılmıştır. Türkçe olarak elde edilen tweetlerin %25.53’ü olumsuz, %30.82’si nötr ve %43.63’ü olumlu olarak ayrılmıştır. 2015-2021 yılları arasında elde edilen tweetler göz önünde bulundurulduğunda Türkçe tweetlerin, İngilizce tweetlere göre daha fazla olumlu tweet ve daha az olumsuz tweet barındırdığı görülmüştür. İngilizce ve Türkçe tweetler elde edilen bütün ülkelerin yıllar bazında attıkları negatif tweetlerin değişimi Şekil 5.35’te, nötr tweetlerin değişimi Şekil 5.36’da ve pozitif tweetlerin değişimi Şekil 5.37’de sunulmuştur.

Çizelge 5.11. Ülkelere göre 2015-2018 yılları arasında ulaşılan tweetlerin dağılımı

Ülke	2015			2016			2017			2018		
	Negatif	Nötr	Pozitif	Negatif	Nötr	Pozitif	Negatif	Nötr	Pozitif	Negatif	Nötr	Pozitif
Güney Afrika	13121	7644	12316	17168	10575	21976	7629	6119	16262	10618	12109	12977
Kanada	18438	10254	11125	14358	15563	11584	18886	15681	29593	12093	10551	18612
İngiltere	23922	20494	21541	16539	19368	26900	12441	7805	15337	5518	6352	8389
Fransa	10603	4371	10009	9355	9246	9009	16530	10331	11680	10071	8287	15778
İtalya	7258	7341	7013	5108	5163	10005	7083	7152	13002	10161	10383	10042
Güney Kore	10850	10775	15371	5814	4656	7494	10925	2264	10139	14049	9882	12509
ABD	10915	8965	8782	7684	9945	5248	6572	2380	9198	10742	9451	15180
Türkiye	18321	13744	21711	17631	14530	21615	15375	16933	21468	13463	18124	22189

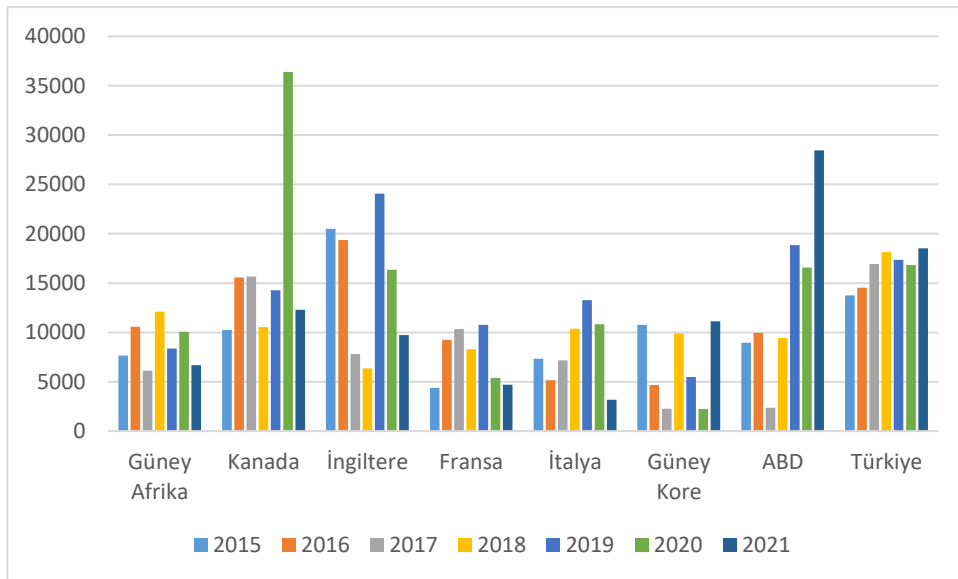
Çizelge 5.12. Ülkelere göre 2019-2021 yılları arasında ulaşılan tweetlerin dağılımı

Ülke	2019			2020			2021		
	Negatif	Nötr	Pozitif	Negatif	Nötr	Pozitif	Negatif	Nötr	Pozitif
Güney Afrika	6634	8360	10260	4457	10057	5593	3797	6684	9585
Kanada	17728	14274	20814	16627	36410	21680	10210	12287	12235
İngiltere	27488	24082	32049	16768	16350	21022	10891	9714	21012
Fransa	14449	10763	14022	5268	5386	7010	3433	4690	6019
İtalya	10134	13262	13340	19115	10834	17121	1038	3190	7013
Güney Kore	10179	5472	10060	12099	2240	5024	11638	11121	13001
ABD	22714	18861	21329	33021	16570	26102	20487	28457	33115
Türkiye	11928	17363	24485	9900	16827	27049	9489	18532	25754



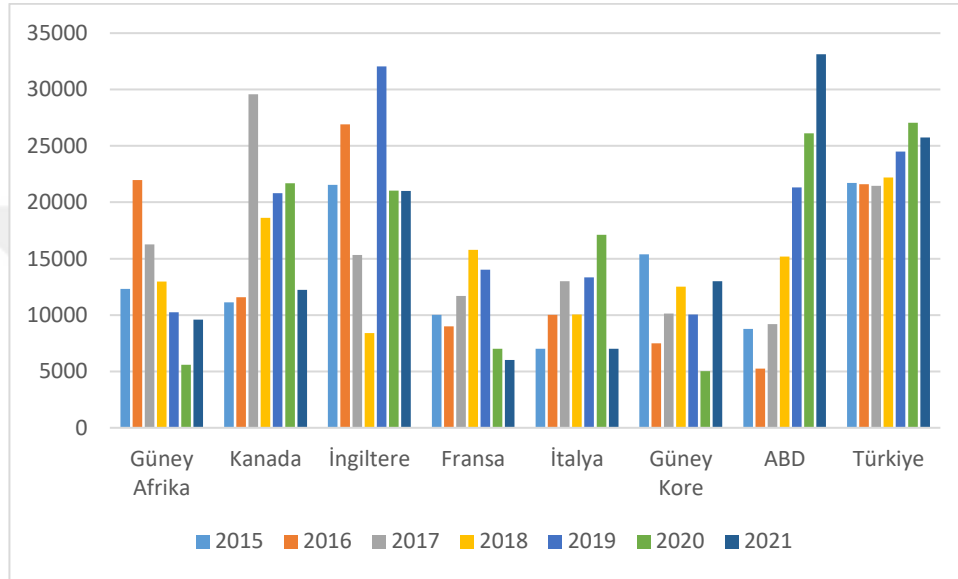
Şekil 5.35. Negatif tweetlerin ülkelere ve yıllara göre değişimi.

Şekil 5.35'deki veriler dikkate alındığında, Güney Afrika'da her geçen yıl mobil oyunlarla alakalı paylaşılan negatif tweetlerde azalma olduğu tespit edilmiştir. Amerika Birleşik Devletleri'nde son yıllarda mobil oyunlarla alakalı paylaşılan negatif tweetlerde artış olduğu tespit edilmiştir. İtalya'da benzer şekilde 2021 yılına kadar artış olduğu dikkat çekmektedir. Kanada ve İngiltere üzerinden değerlendirme yapılırsa son 3 senede bir azalma olduğu tespit edilmiştir. Türkiye'de ise 2015 yılından günümüze mobil oyunlarla alakalı paylaşılan negatif tweetlerin düzenli olarak azaldığı tespit edilmiştir.



Şekil 5.36. Nötr tweetlerin ülkelere ve yıllara göre değişimi.

Şekil 5.36'daki veriler dikkate alındığında, genel olarak bütün ülkelerde mobil oyunlarla alakalı paylaşılan nötr tweetlerde önemli bir artış ya da düşüş tespit edilmemiştir. Amerika Birleşik Devletleri'nde son 3 yıldaki değerlerin önceki yıllara göre daha fazla olduğu, Türkiye'de ise 2015 yılından günümüze mobil oyunlarla alakalı paylaşılan nötr tweetlerin azda olsa arttığı tespit edilmiştir.



Şekil 5.37. Pozitif tweetlerin ülkelere ve yıllara göre değişimi.

Şekil 5.37'deki veriler dikkate alındığında, Güney Afrika'da 2021 yılına kadar her geçen yıl mobil oyunlarla alakalı paylaşılan pozitif tweetlerde azalma olduğu tespit edilmiştir. Amerika Birleşik Devletleri'nde 2015 yılında günümüze mobil oyunlarla alakalı paylaşılan pozitif tweetlerde düzenli olarak artış olduğu tespit edilmiştir. Fransa'da son 3 senede paylaşılan pozitif tweetlerin oranında azalma olduğu tespit edilmiştir. Türkiye'de ise 2015 yılından günümüze mobil oyunlarla alakalı paylaşılan pozitif tweetlerin düzenli olarak arttığı tespit edilmiştir.



6. SONUÇ

Tezin bu bölümünde, elde edilen bulgulara dayalı olarak sonuçlara yer verilmiştir.

Bu tez çalışmasında Twitter üzerinden elde edilen mobil oyunlarla ilgili İngilizce ve Türkçe tweetler kullanılarak bir duygu analizi çalışması yapılmıştır. Mobil oyun veri setine ait değişkenlerinin kullanıldığı bu duygu analizi çalışmasının, oyuncu kitlesinin mobil oyunlar hakkındaki pozitif veya negatif bakış açıları esas alınarak mobil oyun üreticilerine oyunlar hakkında fikir vermesi ve yol göstermesi açısından katkı sağlayacağı düşünülmektedir. Literatürde Twitter üzerinden toplanan verilerle birçok alanda duygu analizi çalışmaları yapılmış olsa da mobil oyunlar konusunda yapılmış bir tez çalışmasının olmaması, TF-IDF ve CountVectorizer vektör tiplerinin kullanılarak sonuçların karşılaştırılması, Türkiye’den toplanan mobil oyunlarla ilgili Twitter verileri İngiltere, Güney Kore, Fransa, İtalya, Amerika Birleşik Devletleri, Güney Afrika ve Kanada’dan toplanan verilerle karşılaştırılarak yapılan diğer duygu analizi çalışmalarında olmayan bir çeşitliliğin yakalanması tezin önemini göstermektedir.

Bu çalışmada Anaconda Navigator-JupyterLab 2.2.6 versiyonu ile phyton kodu yazılarak elde edilen mobil oyunlarla alakalı İngilizce ve Türkçe tweetler basit makine öğrenme algoritmaları (LR, Linear SVC, Multinomial NBC, Bernoulli NBC, RC, PAC, MLPC, KNN), önerilen kolektif öğrenme algoritması (VC) ve geliştirilen hibrit model (TEMSAP-CNNLSTM) ile analiz edilerek duygu analizi çalışması yapılmıştır.

Güney Afrika, Kanada, İngiltere, Fransa, İtalya, Güney Kore ve Amerika Birleşik Devletlerinden mobil oyun ile alakalı İngilizce tweetler elde edilerek, Türkiye’den sadece Türkçe tweetler veri setine eklenmiştir. 2015 ve 2021 yılları arasında İngilizce olarak toplamda 1.839.274 adet tweet, Türkçe olarak 376.431 adet tweet elde edilmiştir. Elde edilen tüm tweetler veri ön işleme adımlarından geçirilerek kök tweetler elde edilmiştir. Polariteleri hesaplanan tweetler pozitif – negatif ve nötr olarak ayrılmış daha sonrasında nötr tweetler dışarıda bırakılarak eşit sayıda negatif ve pozitif tweetler kullanılarak modeller oluşturulmuş, bu modeller kullanılarak makine öğrenme algoritmalarıyla analizler yapılmıştır. Türkçe tweetler önce İngilizce’ye çevrilmiş, polariteleri hesaplanmıştır. Bunun için TextBlob kütüphanesinden yararlanılmıştır. Türkçe tweetlerin polaritesinin hesaplanmasında TextBlob kütüphanesini kullanarak İngilizceye çevirme

işlemi oldukça vakit almıştır. Bu tez çalışmasında, 500 tweetin çevirisi ortalama 445 saniyede tamamlanmıştır. 50.000 tweet için ise, ortalama 12.3 saat sürmesi bu işin zorluğunu ortaya koymaktadır. Toplamda 376.431 Türkçe tweetin İngilizceye çevrilmesi durumunda, sistemde kitlenme ve donma gibi problemler ile karşılaşmış ve dolayısıyla sonuç alınamamıştır. Çözüm olarak 7 ayrı dataframe oluşturulmuş ve 6 tanesine 50.000 tweet ile birlikte son dataframe içerisine ise 76.431 tweet yerleştirilerek bu problemin çözümü sağlanmıştır.

Diğer çalışmalardan farklı olarak kolektif öğrenme algoritması için İngilizce için ve Türkçe için farklı sınıflandırma algoritmaları kullanılmıştır. İngilizce Tweetler için kullanılan kolektif öğrenme algoritması için RF, Bernoulli NBC, KNN sınıflama algoritmaları beraber kullanılmıştır. Türkçe Tweetler için kullanılan kolektif öğrenme algoritması için RF, LR, Linear SVC, RC sınıflama algoritmaları beraber kullanılmıştır. Hem İngilizce hem de Türkçe Tweetler için önerilen bu kolektif öğrenme algoritması diğer basit sınıflama algoritmalarına göre en iyi sonuçları vermiş ve başarılı olmuştur. Mobil oyun İngilizce test veri setinde Acc değeri incelendiğinde 0.91'lik değerle VC modeli hem TF-IDF hem de CountVectorizer vektör tipinde en başarılı ikinci model olmuştur. Metrik değerler incelendiğinde, VC modeli, 0.858'lik P değeri, 0.997'lik R değeri, 0.922'lik F değeri ve %91.58'lik AUC değeri ile diğer basit makine öğrenme algoritmalarına göre daha başarılı olmuştur. En başarılı model ise 0.93'lük Acc değeri ile TEMSAP-CNNLSTM modeli olmuştur. Metrik değerler incelendiğinde, TEMSAP-CNNLSTM modeli, 0.984'lük P değeri, 0.966'lık R değeri, 0.938'lik F değeri ve %93.64'lük AUC değeri ile en başarılı sonuçları elde etmiştir. Mobil oyun Türkçe test veri setinde Acc değeri incelendiğinde 0.90'lık değerle VC modeli hem TF-IDF hem de CountVectorizer vektör tipinde en başarılı ikinci model olmuştur. Metrik değerler incelendiğinde, VC modeli, 0.888'lik P değeri, 0.928'lik R değeri, 0.908'lik F değeri ve %91.02'lik AUC değeri ile diğer basit makine öğrenme algoritmalarına göre daha başarılı olmuştur. En başarılı model ise 0.92'lik Acc değeri ile TEMSAP-CNNLSTM modeli olmuştur. Metrik değerler incelendiğinde, TEMSAP-CNNLSTM modeli, 0.911'lik P değeri, 0.941'lik R değeri, 0.926'lık F değeri ve %92.61'lik AUC değeri ile en başarılı sonuçları elde etmiştir. Geliştirilen TEMSAP-CNNLSTM modeli kolektif öğrenme algoritmasından (VC) ve diğer basit öğrenme algoritmalarından çok daha iyi sonuçlar

elde edilmiş olup P, R, F, auc ve doğruluk metriklerinde elde edilen başarılı değerlerle desteklenmiştir. Veriler CountVectorizer ve TF-IDF vektör tiplerine göre değerlendirilmiş ve sonuçları karşılaştırılmıştır. Mobil oyun İngilizce test veri seti için ve mobil oyun Türkçe test veri seti için en başarılı sonuçlar TF-IDF vektör tipinde elde edilmiştir.

Bu tez çalışmasında karmaşıklık matrisleri elde edilerek ROC eğrileri çizdirilmiştir. Birçok çalışmada sadece test verileri üzerinde uygulanan karmaşıklık matrisleri bu çalışmada eğitim ve tüm veri setinde de uygulanarak aralarındaki farklar incelenmiştir. Mobil oyun İngilizce test veri seti incelendiğinde, LR %88.49, Linear SVC %82.24, Multinomial NBC %70.16, Bernoulli NBC %79.69, RC %83.91, PAC %67.72, MLPC %82.88, KNN %86.04, VC %91.58, TEMSAP-CNNLSTM %93.64'lük AUC değerleri elde etmişlerdir. Sonuçlara bakıldığında TEMSAP-CNNLSTM modelinin VC modeli ile birlikte mobil oyun İngilizce test veri seti için pozitif ve negatif tweetleri ayırt etme gücünde en başarılı modeller oldukları tespit edilmiştir. Mobil oyun Türkçe test veri seti incelendiğinde, LR %85.97, Linear SVC %79.22, Multinomial NBC %67.97, Bernoulli NBC %77.91, RC %81.44, PAC %64.12, MLPC %73.72, KNN %84.49, VC %91.02, TEMSAP-CNNLSTM %92.61'lik AUC değerleri elde etmişlerdir. Sonuçlara bakıldığında TEMSAP-CNNLSTM modelinin VC modeli ile birlikte mobil oyun Türkçe test veri seti için pozitif ve negatif tweetleri ayırt etme gücünde en başarılı modeller oldukları tespit edilmiştir.

Bütün makine öğrenme modellerinin çalışma süreleri karşılaştırılmış ve en doğru sonuçların elde edildiği TEMSAP-CNNLSTM ve VC modellerinin diğer modellere göre daha yavaş çalıştığı görülmüştür. VC modelini oluşturan diğer sınıflama algoritmaları; verinin eğitilmesi, test verilerinin sonuçlarına göre oylama işlemini gerçekleştirmeleri ve tüm bu sonuçlara göre bir tahmin üretmelerinden dolayı analizini geç tamamlamışlardır. TEMSAP-CNNLSTM modeli, evrişim katmanı, havuz katmanı ve kullanmış olduğu ReLU aktivasyon fonksiyonuyla düzleştirme katmanında yaptığı işlemlerin yanında LSTM katmanında uzun metinlerin analizi için gereken sürenin fazla olmasından kaynaklı olarak VC modelinden sonra en yavaş çalışan model olmuştur. İngilizce mobil oyun test veri seti incelendiğinde, VC modeli CountVectorizer vektör tipinde 21.52 saniyede, TF-IDF vektör tipinde ise 28.78 saniyede işlemi tamamlamıştır. TEMSAP-

CNNLSTM modeli ise 62.7 saniyede işlemi tamamlamıştır. Türkçe mobil oyun test veri seti incelendiğinde, VC modeli CountVectorizer vektör tipinde 17.34 saniyede, TF-IDF vektör tipinde ise 15.04 saniyede işlemi tamamlamıştır. TEMSAP-CNNLSTM modeli ise 23.8 saniyede işlemi tamamlamıştır. Bu çalışmada Türkçe tweetler üzerinde uygulanan modellerin çalışma süreleri daha kısa olarak hesaplanmıştır. Bunun temel nedeni verinin büyüklüğü ile alakalıdır.

Çalışmada karşılaşılan zorluklar göz önünde bulundurulduğunda, çalışmanın yapılacağı dil seçilirken mevcut kütüphanelerin iyi araştırılması gerekmekte ve çalışmaya yararlı olacak kütüphanelerin kullanılabildiği diller ile alakalı geliştirme yapılması önerilmektedir. Zıt duyguları bir arada bulunduran ve değerlendirme dışı bırakılan tweetlere yönelik de duygu tespitinde zorluklarla karşılaşmış ve ilerde geliştirilmesi gerektiği düşünülmektedir. Hibrit ve kolektif öğrenme modellerinin çalışma sürelerinin kısaltılması ayrıca çok büyük verilerle çalışılırken kullanılan sistemlerin donanımının yetersiz kalmasından dolayı bulut sistem üzerinden çalışmaların yapılması önerilmektedir.

KAYNAKLAR

- Akar, Ö., 2013. *Rastgele Orman Sınıflandırıcısına Doku Özellikleri Entegre Edilerek Benzer Spektral Özellikteki Tarımsal Ürünlerin Sınıflandırılması* (doktora tezi). KTÜ, Fen Bilimleri Enstitüsü, Trabzon.
- Akgül, E. S., Ertano, C., Diri, B., 2016. Twitter verileri ile duygu analizi. *Pamukkale University Journal of Engineering Sciences*, 22(2): 106-111.
- Akyel, N., ve Seçkin, K., 2011. K-en yakın komşuluk algoritmasının hile denetiminde kullanımı. *Muhasebe ve Vergi Uygulamaları Dergisi*, 51: 21-39.
- Alaei, A. R., Becken, S., Stantic, B., 2017. Sentiment analysis in tourism: Capitalizing on big data. *Journal of Travel Research*, 58(2): 175-191.
- Almaghrabi, M., & Chetty, G., 2020. Improving Sentiment Analysis in Arabic and English Languages by Using Multi-Layer Perceptron Model (MLP). *IEEE 7th International Conference on Data Science and Advanced Analytics (DSAA)*, 6-9 October 2020, Sydney. 745-746.
- Al-Mekhlal, M., Khwaja, A. A., 2019. A synthesis of big data definition and characteristics. *International Conference on Computational Science and Engineering (CSE) and IEEE International Conference on Embedded and Ubiquitous Computing (EUC)*, 1-3 August 2019, New York. 314-322.
- Atasoy, N., Tabak D., 2018. Face Recognition Application Development Using Ssupport Vector Machines. *Engineering Sciences (NWSAENS)*, 13(2): 119-127.
- Aytekin, G., 2017. Dijital Oyunlar ve Bireyler Üzerindeki Etkileri. <https://www.guvenliweb.org.tr/blog-detay/dijital-oyunlar-ve-bireyler-uzerindeki-etkileri>. Erişim Tarihi: 17.01.2022.
- Aytuğ, O., 2015. Şirket İflaslarının Tahminlenmesinde Karar Ağacı Algoritmalarının Karşılaştırmalı Başarım Analizi. *Bilişim Teknolojileri Dergisi*, 8(1): 9-19.
- Aziz A., 2010. *İletişime Giriş*, 3.baskı. Hiperlink Yayınları. İstanbul, Türkiye. 134
- Baj-Rogowska, A., 2017. *Sentiment analysis of Facebook posts: The Uber case. Eighth International Conference on Intelligent Computing and Information Systems (ICICIS)*, 5-7 December 2017, Cairo. 391-395.
- Ballı, S., Sağbaş, E. A., Korukoğlu, S., 2018. Akıllı telefon hareket algılayıcıları ile ulaşım türlerinin gerçek zamanlı tespit edilmesi. *International Conference on Intelligent Transportation Systems (BANU-ITSC'18)*, 19-21 Nisan 2018, Balıkesir. 21.
- Balouchzahi, F., Shashirekha, H. L., Sidorov, G., 2021. HSSD: Hate Speech Spreader Detection using N-grams and Voting Classifier. *Conference and Labs of the Evaluation Forum*, September 21-24 (CLEF), Bucharest. 1-8.
- Barutçu, S., ve Tomaş, M., 2013. Sürdürülebilir Sosyal Medya Pazarlaması ve Sosyal Medya Pazarlaması Etkinliğinin Ölçümü. *Journal of Internet Applications & Management/İnternet Uygulamaları ve Yönetimi Dergisi*, 4(1): 5-23.
- Beyhan, H. D., 2014. *Sosyal Medya Üzerinden Metin Madenciliği ve Duygu Analizi ile Pazar Değerlendirme* (yüksek lisans tezi). İstanbul Teknik Üniversitesi, Fen Bilimleri Enstitüsü, İstanbul.
- Bewick, V., Cheek, L., Ball, J., 2005. Statistics review 14: Logistic regression. *Critical care*, 9(1): 1-7.

- Bıçek, E., & Çelik, E. H., 2020. Fotovoltaik Sistemin Güç Üretiminin Meteorolojik Değişkenler ile Modellenmesi: Van Yüzüncü Yıl Üniversitesi Örneği. *Yüzüncü Yıl Üniversitesi Fen Bilimleri Enstitüsü Dergisi*, 25(3): 135-146.
- Bongirwar, V. K., 2015. A Survey on Sentence Level Sentiment Analysis. *International Journal of Computer Science Trends and Technology (IJCT)*, 3(3): 110–113.
- Bose, I., and Yang, X., 2011. Enter The Dragon: Khillwar's Foray Into The Mobile Gaming Market of China. *Communications of the Association for Information Systems*, 29: 551–564.
- Bostancı, M. (2019). *Sosyal medya: Dün Bugün Yarın*. 1. Baskı. Palet Yayınları. Konya, Türkiye.
- Böhmer, M., Hecht, B., Schöning, J., Krüger, A., & Bauer, G., 2011. Falling asleep with Angry Birds, Facebook and Kindle: a large scale study on mobile application usage. *In Proceedings of the 13th international conference on Human computer interaction with mobile devices and services*, 30August 2011 – 02 September 2011, New York. 47-56.
- Brief, A. P., Weiss, H. M., 2002. Organizational behavior: *Affect in the workplace*. *Annual review of psychology*, 53(1): 279-307.
- Brownlee, J., 2016. Supervised and Unsupervised Machine Learning Algorithms., <https://machinelearningmastery.com/supervised-and-unsupervised-machine-learning-algorithms/> Erişim Tarihi: 04.01.2022.
- BTK , 2018. Bilgi Teknolojileri Kurumu. <https://www.btk.gov.tr/haberler/1-milyarin-uzerinde-insan-dijital-oyun-oyunuyor> , Ankara. Erişim Tarihi: 17.01.2022
- Buyya, R., Calheiros, R. N., Dastjerdi, A. V., 2016. *Big data: Principles and Paradigms*. 1st Edition, Morgan Kaufmann, USA. 3-223.
- Castronova, E., 2002. On Virtual Economies. *The International Journal of Computer Game Research*, 3(2): 752.
- Chen, Y., Hu, X., Fan, W., Shen, L., Zhang, Z., Liu, X., Li, H., 2019. Fast density peak clustering for large scale data based on kNN. *Knowledge-Based Systems*, 187: 1-7.
- Cheng, M.,J and Jin, X., 2019. What do Airbnb users care about? An analysis of online review comments. *International Journal of Hospitality Management*, 76: 58-70.
- Clausen, L. B. N., and Nickisch, H., 2018. Automatic Classification of Auroral Images From the Oslo Auroral THEMIS (OATH) Data Set Using Machine Learning. *Journal of Geophysical Research: Space Physics*, 7: 1-23.
- Chowdhury, G. G., 2003. Natural language processing. *Annual review of information science and technology*, 37(1): 51–89.
- Cox, M., and Ellsworth, D., 1997. Application-controlled demand paging for out-of-core visualization. *In: Proceedings of the IEEE Visualization Conference*. 21 January 1997, Phoenix. 235-244.
- Çakır, M., Çakır, F., ve Usta, G., 2010. Üniversite öğrencilerinin tüketim tercihlerini etkileyen faktörlerin belirlenmesi. *Organizasyon ve Yönetim Bilimleri Dergisi*, 2(2): 87-94.
- Çetin, E., 2013. *Tanımlar ve temel kavramlar, Eğitsel dijital oyunlar*. 1. Baskı, Pagem Akademi, Ankara. 2-18.

- Çoban, Ö., Özyer, G. T., 2018. Word2vec and clustering based twitter sentiment analysis. *International Conference on Artificial Intelligence and Data Processing (IDAP)*. 28-30 September 2018, Malatya. 1-5.
- Dağıtmaç, M., 2015. *Sosyal medya bizi neden kullanır*. 1. Baskı, Metamorfoz Yayıncılık, İstanbul. 216.
- Da Silva, N. F., Hruschka, E. R., Hruschka Jr, E. R., 2014. Tweet sentiment analysis with classifier ensembles. *Decision Support Systems*, **66**: 170-179.
- Delen D., Sharda R., Turban E., 2014. *Business Intelligence and Analytics : System for Decision Support*, 10th Edition, Pearson Education Limited, UK. 322-326.
- Deloitte, 2019. Global Mobil Oyun Araştırması. Ankara. <https://www2.deloitte.com/tr/tr/pages/about-deloitte/press-releases/deloitte-global-mobil-kullanici-arastirmasi-2019-sonuclari-aciklandi.html> Erişim Tarihi: 03.01.2022.
- Dijilopedi, 2020. Dünya İnternet, Sosyal Medya ve Mobil Kullanım İstatistikleri, İstanbul. <https://dijilopedi.com/2020-dunya-internet-sosyal-medya-ve-mobil-kullanim-istatistikleri/> Erişim tarihi: 15.10.2021.
- Dirican, A., 2001. Evaluation of the diagnostic test's performance and their comparisons. *Cerrahpaşa J Med*, **32**(1): 25-30.
- El Fazziki, A., Ennaji, F. Z., Sadiq, A., Benslimane, D., Sadgal, M., 2017. A multi-agent based social crm framework for extracting and analysing opinions. *Journal of Engineering Science and Technology*, **12**(8): 2154-2174.
- El-kenawy, E.-S. M., Ibrahim, A., Mirjalili, S., Eid, M. M., Hussein, S. E., 2020. Novel Feature Selection and Voting Classifier Algorithms for COVID-19 Classification in CT Images. *IEEE Access*, **8**: 1–19.
- Eröz, S. S., 2011. *Duygusal Zeka ve İletişim Arasındaki İlişki: Bir Uygulama* (Doktora Tezi). Uludağ Üniversitesi, Sosyal Bilimler Enstitüsü, Bursa.
- Fu, G. ve X. Wang., 2010. Chinese Sentence-Level Sentiment Classification Based on Fuzzy Sets. *23rd International Conference on Computational Linguistics: Posters*, 23-27 August 2010, Pekin. 312–319.
- Gantz, J., Reinsel, D., 2012. The digital universe in 2020: Big data, bigger digital shadows, and biggest growth in the far east. *IDC iView: IDC Analyze the Future*, **2007(2012)**: 1-16.
- Ghaisani, A. P., Handayani, P. W., and Munajat, Q., 2017. Users' Motivation in Sharing Information on Social Media. *Procedia Computer Science*, **124**: 530-535. <https://doi.org/10.1016/j.procs.2017.12.186>
- Goldberg, D. E. ve J. H. Holland., 1988. Genetic Algorithms and Machine Learning. *Machine Learning*, **3**(2): 95–99.
- Göçenoğlu, M., 2014. *Sosyal Ağ Verileri Kullanılarak Türkiye'nin Duygu Analizinin Görselleştirilmesi* (Yüksek Lisans Tezi). Karabük Üniversitesi, Fen Bilimleri Enstitüsü, Karabük.
- Gök, M., 2017. Makine Öğrenmesi Yöntemleri İle Akademik Başarının Tahmin Edilmesi. *Gazi Üniversitesi Fen Bilimleri Dergisi Part C: Tasarım ve Teknoloji*, **5**: 139-148.
- Grønli, T. M., Hansen, J., Ghinea, G., Younas, M., 2014. Mobile application platform heterogeneity: Android vs Windows Phone vs iOS vs Firefox OS. *28th*

- International Conference on Advanced Information Networking and Applications**. 13-16 May 2014, Victoria. 635-641.
- Güran, A., Uysal, M., Doğruöz, Ö., 2014. Destek Vektör Makineleri Parametre Optimizasyonunun Duygu Analizi Üzerindeki Etkisi. *DEÜ Mühendislik Fakültesi Mühendislik Bilimleri Dergisi*, **48**: 87–88.
- Gürsakan, N., 2009. *Sosyal Ağ Analizi*. 1. Baskı, Dora Yayıncılık, Bursa. 528.
- Hair, J. F., Black, W. C., Babin, B., Anderson, R. E., Tatham, R. L., 2006. *Multivariate data analysis*. 6th edition, Upper Saddle River, New Jersey. 924.
- Han, J., Kamber, M. and Pei, J., 2012. *Data Mining Concepts and Techniques*. 3rd Edition, Morgan Kaufmann Publishers, United States. 703.
- Hastie, T., Friedman, J., Tibshirani, R., 2001. Prototype Methods and Nearest-Neighbors. *The Elements of Statistical Learning*, **3**(1): 411-435.
- Hatzivassiloglou, V., K. R. McKeown., 1997. Predicting the Semantic Orientation of Adjectives. *35th Annual Meeting of the Association for Computational Linguistics and Eighth Conference of the European Chapter of the Association for Computational Linguistics*, July 1997, Madrid. 174–181.
- Hazar, M., 2011. Sosyal medya bağımlılığı-bir alan çalışması. *İletişim, Kuram ve Araştırma Dergisi*, **(32)**: 151-176.
- Hölliga, J., Leea, Y. S., Seemanna, N., Geierhosa, M., 2021. Effective Detection of Hate Speech Spreaders on Twitter. *Conference and Labs of the Evaluation Forum (CLEF)*, 21-24 September 2021, Bucharest. 1-11.
- Hurwitz, J. S., Nugent, A., Halper, F., Kaufman, M., 2013. *Big Data for Dummies*. 1st Edition, John Wiley & Sons, USA. 25-110.
- Iwendi, C., Bashir, A. K., Peshkar, A., Sujatha, R., Chatterjee, J. M., Pasupuleti, S., Jo, O., 2020. COVID-19 Patient Health Prediction Using Boosted Random Forest Algorithm. *Frontiers in Public Health*, **8**: 1-9.
- Jagtap, V., K. Pawar, K., 2013. Analysis of Different Approaches to Sentence-Level Sentiment Classification. *International Journal of Scientific Engineering and Technology*, **2**(3): 164–170.
- Jain, P. K., Saravanan, V., Pamula, R., 2021. A Hybrid CNN-LSTM: A Deep Learning Approach for Consumer Sentiment Analysis Using Qualitative User-Generated Contents. *ACM Transactions on Asian and Low-Resource Language Information Processing*, **20**(5): 1–15.
- Jianqiang, Z., Xiaolin, G., Xuejun, Z., 2018. Deep Convolution Neural Networks for Twitter Sentiment Analysis. *IEEE Access*, **6**: 23253–23260.
- Jimenez-Marquez, J. L., Gonzalez-Carrasco, I., Lopez-Cuadrado, J. L., Ruiz-Mezcua, B., 2019. Towards a big data framework for analyzing social media content. *International Journal of Information Management*, **44**: 1-12.
- Jorge, J., and Paredes, R., 2018. Passive-Aggressive online learning with nonlinear embeddings. *Pattern Recognition*, **79**: 162–171.
- Juneja, P., Ojha, U., 2017. Casting online votes: To predict offline results using sentiment analysis by machine learning classifiers. *8th International Conference on Computing, Communication and Networking Technologies (ICCCNT)*, 3-5 July 2017, Delhi. 1-6.

- Kanık, E. A., Erden, S., 2003. Tanı Testlerinin değerlendirilmesinde ROC (Receive Operating Characteristics) Eğrisinin Kullanımı. *Mersin Üniversitesi Tıp Fakültesi Dergisi*, 3: 260-264.
- Kaplan, A. M., Haenlein, M., 2010. Users of the world, unite! The challenges and opportunities of social media. *Business Horizons*, 53(1): 59-68.
- Karabulut, Y. E., ve Küçüksille, E. U., 2018. Twitter profesyonel izleme ve analiz aracı. *Teknik Bilimleri Dergisi*, 8(2): 17-24.
- Karakaynak, Z., 2014. *Destek vektör makineleri ile sınıflandırma ve görüntü tanıma üzerine bir uygulama* (yüksek lisans tezi). Mimar Sinan Güzel Sanatlar Üniversitesi, Fen Bilimleri Enstitüsü, İstanbul.
- Karoui, J., Zitoun, F. B., Moriceau, V., 2019. *Automatic Detection of Irony*. 1st Edition, John Wiley & Sons, USA. 189.
- Kaynar, O., Y. Görmez, M. Yıldız ve A. Albayrak., 2016. Makine Öğrenmesi Yöntemleri ile Duygu Analizi. *International Artificial Intelligence and Data Processing Symposium*, 17-18 September 2021, Malatya. 234-241.
- Kılıç S., 2013. Klinik Karar Vermede ROC Analizi. *Journal of Mood Disorders*, 3: 135-140.
- Kirtiç, K., 2013. *Pazarlama İlkeleri Global ve Yönetimsel Yaklaşım*. 1. Baskı, İstanbul Gelişim Üniversitesi Yayınları, İstanbul. 139.
- Kleinbaum, D. G., Klein, M., 2002. *Logistic Regression A Self-learning Text*. 1st Edition, New York: Springer, USA. 701.
- Korkusuz, R., Carus, A., 2020. Futbol Müsabakaları ile İlgili Tweetlerin Anlık Duygu Analizi. *Avrupa Bilim ve Teknoloji Dergisi, (Özel Sayı)*: 386-396.
- Kumar, U. K., Nikhil, M. B. S., Sumangali, K., 2017. Prediction of breast cancer using voting classifier technique. *International Conference on Smart Technologies and Management for Computing, Communication, Controls, Energy and Materials (ICSTM)*, 2-4 August 2017, Chennai. 108-114.
- Kumari, S., Kumar, D., Mittal, M., 2021. Yumuşak oylama sınıflandırıcısı kullanılarak diabetes mellitusun sınıflandırılması ve tahmini için bir topluluk yaklaşımı. *International Journal of Cognitive Computing in Engineering*, 2: 40-46.
- Kuss, D. J., Griffiths, M. D., 2012. Internet gaming addiction: A systematic review of empirical research. *International Journal of Mental Health and Addiction*, 10(2): 278-296.
- Küresel Dijital rapor, 2020. Digital 2020: Global Digital Overview <https://datareportal.com/reports/digital-2020-global-digital-overview> Erişim tarihi: 11.01.2022
- Laney, D., 2001. 3D data management: Controlling data volume, velocity and variety. *META Group Research Note*, 6(70): 1.
- LeBlanc, A. G., Chaput, J. P., 2017. Pokémon Go: a game changer for the physical inactivity crisis?. *Preventive medicine*, 101: 235-237.
- Li, W., Zhu, L., Shi, Y., Guo, K., Zheng, Y., 2020. User reviews: Sentiment analysis using lexicon integrated two-channel CNN-LSTM family models. *Applied Soft Computing*, 94: 1-11.
- Liu, B., 2012, *Sentiment Analysis and Opinion Mining*, 1st Edition, Morgan & Claypool Publishers, Chicago. 168.

- Liu, Y., Huang, K., Bao, J., Chen, K., 2019. Listen to the voices from home: An analysis of Chinese tourists' sentiments regarding Australian destinations. *Tourism Management*, 71: 337-347.
- Lynch, J. G., 2012. *Perceived stress and the buffering effect of perceived social support on facebook* (doctoral dissertation), Antioch University, New England..
- Manavcıoğlu, K., 2009. İnternette Kullanıcıların Oluşturduğu ve Dağıttığı İçeriklerin Etik Açısından İncelenmesi: Sosyal Medya Örnekleri. *Medya ve Etik Sempozyumu*, 7-9 Ekim 2009, Elazığ. 7-9.
- Marwick, A., Boyd, D., 2011. To see and be seen: celebrity practice on twitter. *The International Journal of Research into New Media Technologies*, 17(2): 139-158.
- Messina R., Louradour J., 2015. Segmentation-free handwritten Chinese text recognition with LSTM-RNN. *13th International Conference on Document Analysis and Recognition (ICDAR)*, 23-26 August 2015, Tunis. 171-175.
- Mogaji, E., Erkan, I., 2019. Insight into consumer experience on UK train transportation services. *Travel Behaviour and Society*, 14: 21-33.
- Morris, T., 2009. *Sams teach yourself twitter in 10 minutes*. 1st Edition, Sams Publishing, USA. 173.
- Murphy, K. P., 2012. *Machine learning: a probabilistic perspective*. 1st Edition, MIT press, Cambridge. 1104.
- Nasser, A., 2018. *Kavram-Tabanlı Duygu Analizi Sistemleri İçin Büyük Ölçekli Arapça Duygu Derlemi ve Sözlüğü Oluşturulması* (doktora tezi). Hacettepe Üniversitesi, Fen Bilimleri Enstitüsü, Ankara.
- Newzoo, 2020. Global Games Market Report. Amsterdam. <https://newzoo.com/insights/articles/newzoo-games-market-numbers-revenues-and-audience-2020-2023/> Erişim Tarihi: 17.01.2022
- Norton, E. C., Dowd, B. E., Maciejewski, M. L., 2019. Marginal Effects—Quantifying the Effect of Changes in Risk Factors in Logistic Regression Models. *JAMA*, 321(13): 1304.
- Ombabi, A. H., Ouarda, W., Alimi, A. M., 2020. Deep learning CNN–LSTM framework for Arabic sentiment analysis using textual information shared in social networks. *Social Network Analysis and Mining*, 10(53): 1-13.
- Onan A., 2017. Sentiment Analysis On Twitter Messages Based On Machine Learning Methods. *Yönetim Bilişim Sistemleri Dergisi*, 3(2): 1-14.
- Onan, A., Korukoğlu, S., Bulut, H., 2016. Ensemble of keyword extraction methods and classifiers in text classification. *Expert Systems with Applications*, 57: 232-247.
- Ortigosa, A., Martín, J. M., Carro, R. M., 2014. Sentiment analysis in Facebook and its application to e-learning. *Computers in Human Behavior*, 31: 527–541.
- Ousterhout, J. K., 1998. Scripting: Higher level programming for the 21st century. *Computer*, 31(3): 23-30.
- Palanisamy, P., Yadav, V., Elchuri, H., 2013. Serendio: Simple and Practical lexicon based approach to Sentiment Analysis. *Second Joint Conference on Lexical and Computational Semantics*. June 2013, Atlanta. 543-548.
- Pang, B. and Lee, L., 2008. Opinion mining and sentiment analysis. *Foundations and Trends® in Information Retrieval*, 2(1–2): 1-135.

- Parveen, H., Pandey, S., 2016. Sentiment analysis on Twitter Data-set using Naive Bayes algorithm. **2nd international conference on applied and theoretical computing and communication technology (iCATccT)**. 21-23 July 2016, Karnataka. 416-419.
- Penttinen, E., Rossi, M., Tuunainen, V. K., 2010. Mobile games: Analyzing the needs and values of the consumers. **Journal of Information Technology Theory and Application (JITTA)**, 11(1): 5-22.
- Poria, S., Hussain, A., Cambria, E., 2018. **Multimodal sentiment analysis**. 1st Edition, Springer International Publishing, UK. 214. <https://doi.org/10.1007/978-3-319-95020-4>
- Pozzi, F. A., Fersini, E., Messina, E., Liu, B., 2017. Challenges of sentiment analysis in social networks: an overview. **Sentiment analysis in social networks**, 1-11. <https://doi.org/10.1016/B978-0-12-804412-4.00001-2>
- Pradhan, V. M., Vala, J., Balani, P., 2016. A survey on sentiment analysis algorithms for opinion mining. **International Journal of Computer Applications**, 133(9): 7-11.
- PYPL, 2021. Programlama dilleri popülerlik indeksi, <https://pypl.github.io/PYPL.html> Erişim Tarihi: 17.01.2022
- Qing, X., Niu, Y., 2018. Hourly day-ahead solar irradiance prediction using weather forecasts by LSTM. **Energy**, 148: 461-468.
- Ranganathan, C., Dhaliwal, J. S., & Teo, T. S. H., 2009. Estudio sobre la privacidad de los datos personales y la seguridad de la información en las redes sociales online. **International Journal of Electronic Commerce**, 9: 127-162.
- Rehman, A. U., Malik, A. K., Raza, B., Ali, W., 2019. A hybrid CNN-LSTM model for improving accuracy of movie reviews sentiment analysis. **Multimedia Tools and Applications**, 78(18): 26597-26613.
- Rollings, A., Adams, E., 2003. **Andrew Rollings and Ernest Adams on game design**. 1st Edition, New Riders Publishers, UK. 400.
- Russell, S. J., Norvig, P., Davis, E., 2010. **Artificial intelligence: A modern approach**. Global 4th Edition, Prentice Hall Publishers, USA. 1069
- Ruz, G. A., Henríquez, P. A., Mascareño, A., 2020. Sentiment analysis of Twitter data during critical events through Bayesian networks classifiers. **Future Generation Computer Systems**, 106: 92-104.
- Sakallıoğlu, S., Kaçıranlar, S., 2008. A new biased estimator based on ridge estimation. **Statistical Papers**, 49(4): 669-689.
- Sánchez, J. L. G., Zea, N. P., Gutiérrez, F. L., 2009. From usability to playability: Introduction to player-centred video game development process. **International Conference on Human Centered Design** . 19-24 July 2009, San Diego. 65-74.
- Satapathy, R., Cambria, E., Hussain, A. 2018. Sentiment Analysis in the Bio-Medical Domain. **Springer International Publishing AG, Gwerbestrasse**, 11: 6630.
- Seker, S. E., 2015. Büyük Veri ve Büyük Veri Yaşam Döngüleri. **YBS Ansiklopedi**, 2(3): 10-17.
- Sharda, R., Delen, D., Turban, E., Aronson, J., Liang, T., 2014. **Business intelligence and analytics, System for Decision Support**. 10th Edition, Pearson Education Limited, UK. 688.
- Singh, G., Kumar, B., Gaur, L., Tyagi, A., 2019. Comparison between multinomial and Bernoulli naïve Bayes for text classification. **International Conference on**

- Automation, Computational and Technology Management (ICACTM)**. 24-26 April 2019, Uttar Pradesh. 593-596.
- Singh, A., Prakash, B. S., Chandrasekaran, K., 2016. A comparison of linear discriminant analysis and ridge classifier on Twitter data. **International Conference on Computing, Communication and Automation (ICCCA)**. 29-30 April 2016, Noida. 133-138.
- Statista, 2021. Business Data Platform. New York. <https://www.statista.com/statistics/330695/number-of-smartphone-users-worldwide/> Erişim Tarihi: 05.01.2022
- Steele, C., 2005. No easy Rider? The scholar and the future of the research library by Fremont Rider: A review article. **Journal of Librarianship and Information Science**, 37(1): 45-51.
- Sulaiman, S., Wahid, R. A., Ariffin, A. H., Zulkifli, C. Z., 2020. Question classification based on cognitive levels using linear svc. **Test Eng Manag**, 83: 6463-6470.
- Tekin, M. C., Tunalı, V. 2019. Prioritization of software development demands with text mining techniques. **Pamukkale University Journal of Engineering Sciences**, 25(5): 615-620.
- Terrana, D., Augello, A., Pilato, G., 2014. Facebook users relationships analysis based on sentiment classification. **IEEE International Conference on Semantic Computing**. 16-18 June 2014, California. 290-296.
- Toçoğlu, M. A., 2018. **Türkçe Metinlerde Sözlük Tabanlı Duygu Analizi** (doktora tezi). Dokuz Eylül Üniversitesi, Fen Bilimleri Enstitüsü, İzmir.
- Topal, M., Eydurhan, E., Yağanoğlu, A. M., SÖNMEZ, A., & Keskin, S. (2010). Çoklu doğrusal bağlantı durumunda ridge ve temel bileşenler regresyon analiz yöntemlerinin kullanımı. **Atatürk Üniversitesi Ziraat Fakültesi Dergisi**, 41(1): 53-57.
- Tuncer, E. (2014). **Sosyal medya imparatorluğu, Patron**. 1. Baskı, Akis Yayınları, İstanbul. 304.
- Taşkın, C., Turanlı, M., 2019. TFRS 9 ve Temerrüt Olasılığı Modellemesi. **Sosyal Bilimler Araştırma Dergisi**, 8(1): 273-284.
- Uluç, G. ve Yarcı, A., 2017. Sosyal medya kültürü. **Dumlupınar Üniversitesi Sosyal Bilimler Dergisi**, (52): 88-102.
- Vidushi, S. G., 2017. Sentiment mining of online reviews using machine learning algorithms. **Int. J. Eng. Devel. Res. IJEDR**, 5(2): 1321-1334.
- Vo, Q. H., Nguyen, H. T., Le, B., Nguyen, M. L., 2017. Multi-channel LSTM-CNN model for Vietnamese sentiment analysis. **9th international conference on knowledge and systems engineering (KSE)**. 19-21 October 2017, Vietnam. 24-29.
- Wang, X., Jiang, W., Luo, Z., 2016. Combination of convolutional and recurrent neural network for sentiment analysis of short texts. **Proceedings of COLING 2016, the 26th international conference on computational linguistics: Technical papers**. December 2016, Osaka. 2428-2437.
- Wang, J., Yu, L. C., Lai, K. R., Zhang, X., 2019. Tree-structured regional CNN-LSTM model for dimensional sentiment analysis. **IEEE/ACM Transactions on Audio, Speech, and Language Processing**, 28: 581-591.

- Werbos P., 1974. Beyond regression: New tools for prediction and analysis in the behavioral sciences (doctoral dissertation). Harvard University, MA institute, USA.
- Xu, S., Li, Y., Wang, Z., 2017. Bayesian multinomial Naïve Bayes classifier to text classification. *Advanced multimedia and ubiquitous engineering*. 1st Edition, Springer, Singapore. 347-352.
- Xu, T., Peng, Q., Cheng, Y., 2012. Identifying the semantic orientation of terms using S-HAL for sentiment analysis. *Knowledge-Based Systems*, **35**: 279-289.
- Young, T., Hazarika, D., Poria, S., Cambria, E., 2018. Recent trends in deep learning based natural language processing. *IEEE Computational intelligence magazine*, **13**(3): 55-75.
- Yousaf, A., Umer, M., Sadiq, S., Ullah, S., Mirjalili, S., Rupapara, V., Nappi, M., 2020. Emotion recognition by textual tweets classification using voting classifier (LR-SGD). *IEEE Access*, **9**: 6286-6295.
- Zackariasson, P., Wilson, T. L., 2012. *The video game industry: Formation, present state, and future*. 1st Edition, Routledge Publishers, London. 282.
- Zainuddin, N., Selamat, A., Ibrahim, R., 2018. Hybrid sentiment classification on twitter aspect-based sentiment analysis. *Applied Intelligence*, **48**(5): 1218-1232.
- Zeng, D., Dai, Y., Li, F., Wang, J., Sangaiah, A. K., 2019. Aspect based sentiment analysis by a linguistically regularized CNN with gated mechanism. *Journal of Intelligent & Fuzzy Systems*, **36**(5): 3971-3980.
- Zhang, S., Li, X., Zong, M., Zhu, X., Wang, R., 2017. Efficient kNN classification with different numbers of nearest neighbors. *IEEE transactions on neural networks and learning systems*, **29**(5): 1774-1785.

ÖZ GEÇMİŞ

Lise öğrenimini 1999 yılında Antakya’da tamamladı. Lisans eğitimini 2005 yılında Çankaya Üniversitesi Bilgisayar Mühendisliği Bölümü’nde (İngilizce) tamamladı.

2008 yılında, Van Yüzüncü Yıl Üniversitesi Özalp Meslek Yüksekokulu’nda Öğretim Görevlisi olarak akademik hayatına başladı. 2015 yılında, Van Yüzüncü Yıl Üniversitesi Sağlık Bilimleri Enstitüsü Fizyoloji Anabilim Dalı’nda Yüksek Lisans eğitimini tamamladı. 2015 yılında, Van Yüzüncü Yıl Üniversitesi Fen Bilimleri Enstitüsü İstatistik Anabilim Dalı’nda Doktora eğitimine başladı.

Evli ve iki çocuk babasıdır, iyi derecede İngilizce bilmektedir.

T.C
VAN YÜZÜNCÜ YIL ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ
DOKTORA TEZİ ORJİNALLİK RAPORU

Tarih: 18/01/2022

Tez Başlığı / Konusu: Makine Öğrenmesi Algoritmaları Kullanılarak Twitter Mobil Oyun Verilerinde Duygu Analizi

Yukarıda başlığı/konusu belirlenen tez çalışmamın Kapak sayfası, Giriş, Ana bölümler ve Sonuç bölümlerinden oluşan toplam 106 sayfalık kısmına ilişkin, 18/01/2022 tarihinde şahsım/tez danışmanım tarafından ithenticate intihal tespit programından aşağıda belirtilen filtreleme uygulanarak alınmış olan orijinallik raporuna göre, tezimin benzerlik oranı %9 dur.

Uygulanan filtreler aşağıda verilmiştir:

- Kabul ve onay sayfası hariç,
- Teşekkür hariç,
- İçindekiler hariç,
- Simge ve kısaltmalar hariç,
- Gereç ve yöntemler hariç,
- Kaynakça hariç,
- Alıntılar hariç,
- Tezden çıkan yayınlar hariç,
- 7 kelimedenden daha az örtüşme içeren metin kısımları hariç

Van Yüzüncü Yıl Üniversitesi Lisansüstü Tez Orijinallik Raporu Alınması ve Kullanılmasına İlişkin Yönergeyi inceledim ve bu yönergede belirtilen azami benzerlik oranlarına göre tez çalışmamın herhangi bir intihal içermediğini; aksinin tespit edileceği muhtemel durumda doğabilecek her türlü hukuki sorumluluğu kabul ettiğimi ve yukarıda vermiş olduğum bilgilerin doğru olduğunu beyan ederim.

Gereğini bilgilerinize arz ederim.

18/01/2022

Adı Soyadı: EROL KINA

Öğrenci No: 159102128

Anabilim Dalı: İstatistik Anabilim Dalı

Programı: İstatistik

Statüsü: Y. Lisans ☐

Doktora ☒

DANIŞMAN ONAYI
UYGUNDUR

Doç. Dr. Recep ÖZDAĞ

ENSTİTÜ ONAYI
UYGUNDUR

(Ad,Soyad,Unvan)