

T.C.
BİLECİK ŞEYH EDEBALI ÜNİVERSİTESİ
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ
İŞLETME ANABİLİM DALI

**METİN MADENCİLİĞİ VE DUYGU ANALİZİ ÜZERİNE BİR ÇALIŞMA:
DEPREM UYGULAMALARI KULLANICI YORUMLARI**

YÜKSEK LİSANS TEZİ

EDA EMİR ÖZ

TEZ DANIŞMANI
DOÇ. DR. GÖZDE KOCA

BİLECİK, 2025

10758724

T.C.
BİLECİK ŞEYH EDEBALI ÜNİVERSİTESİ
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ
İŞLETME ANABİLİM DALI

**METİN MADENCİLİĞİ VE DUYGU ANALİZİ ÜZERİNE BİR ÇALIŞMA:
DEPREM UYGULAMALARI KULLANICI YORUMLARI**

YÜKSEK LİSANS TEZİ

EDA EMİR ÖZ

TEZ DANIŞMANI
DOÇ. DR. GÖZDE KOCA

BİLECİK, 2025

10758724

BEYAN

“Metin Madenciliği ve Duygu Analizi Üzerine Bir Çalışma: Deprem Uygulamaları Kullanıcı Yorumları” adlı yüksek lisans tezinin hazırlık ve yazımı sırasında bilimsel araştırma ve etik kurallarına uyduğumu, başkalarının eserlerinden yararlandığım bölümlerde bilimsel kurallara uygun olarak atıfta bulunduğumu, kullandığım verilerde herhangi bir tahrifat yapmadığımı, tezin herhangi bir kısmının Bilecik Şeyh Edebali Üniversitesi veya başka bir üniversitede başka bir tez çalışması olarak sunulmadığını, aksinin tespit edileceği muhtemel durumlarda doğabilecek her türlü hukuki sorumluluğu kabul ettiğimi ve vermiş olduğum bilgilerin doğru olduğunu beyan ederim.

Bu çalışmanın, Bilimsel Araştırma Projeleri (BAP), TÜBİTAK veya benzeri kuruluşlarca desteklenmesi durumunda; projenin ve destekleyen kurumun adı proje numarası ile birlikte, ETİK KURUL onayı alınması durumunda ise ETİK KURUL tarih karar ve sayı bilgilerinin beyan edilmesi gerekmektedir.			
DESTEK ALINMIŞTIR	☐	DESTEK ALINMAMIŞTIR	X
Destek alındı ise;			
Destekleyen kurum;			
Desteğin Türü		Proje Numarası	
1- BAP (Bilimsel Araştırma Projesi)			
2- TÜBİTAK			
Diğer;			
ETİK KURUL onayı var ise;			
ETİK KURUL karar tarih/sayı:			

Eda EMİR ÖZ

.../.../2025

İmza

ÖN SÖZ

Bu tez çalışmasının yazılmasında, çalışmamı sahiplenerek takip eden danışmanım Sayın Doç. Dr. Gözde KOCA'ya ve Doç. Dr. ÖZÜM EĞİLMEZ'e değerli katkı ve emekleri için teşekkürlerimi ve saygılarımı sunarım.

Tez savunma sürecinde çalışmamın son haline gelmesindeki değerli katkıları adına Sayın jüri üyelerine teşekkürlerimi ve saygılarımı sunarım.

Son olarak bugünlere ulaşmamdaki emekleri adına değerli eşime ve aileme teşekkür ederim.

Eda EMİR ÖZ

2025



ÖZET

METİN MADENCİLİĞİ VE DUYGU ANALİZİ ÜZERİNE BİR ÇALIŞMA:

DEPREM UYGULAMALARI KULLANICI YORUMLARI

Türkiye, jeolojik yapısı itibarıyla deprem riski yüksek bir ülkedir. Özellikle 2023 Kahramanmaraş Depremleri gibi büyük felaketler, toplumda deprem bilincinin artırılması gerektiğini açıkça ortaya koymuştur. Depremler hem can hem de mal kaybına yol açabilecek potansiyele sahip doğal afetlerdir ve insanlar hayat kurtarma, maddi zararları azaltma, psikolojik olarak hazırlıklı olmak gibi nedenlerle depremleri önceden bilmek isterler. Telefonlardaki deprem uygulamaları da hem erken uyarı sağlamak hem de depremle ilgili bilgileri hızlı ve güvenilir bir şekilde kullanıcılara iletmek için önemli araçlar haline gelmiştir. Kullanıcıların, kullandıkları deprem uygulamaları hakkındaki geri bildirimleri, geri bildirimin taşıdığı duygunun analiz edilerek yorumlanabilmesi için önemli bir veridir. Bu veriler analiz edilerek elde edilen bilgiler ise uygulamanın geliştirilmesi ve kullanıcı deneyiminin iyileştirilmesi açısından uygulama geliştiricileri için oldukça önemlidir. Bu çalışmada, denetimli makine öğrenmesi yaklaşımı kullanılarak en çok yorum alan deprem uygulamalarına yapılan kullanıcı yorumlarının duygu analizi yapılmıştır. Duygu analizi için denetimli makine öğrenmesi sınıflandırma algoritmalarından Naive Bayes, Sıralı Minimal Optimizasyon ve k-en yakın komşu algoritmaları kullanılmıştır. Google Play ve App Store üzerinden Instant Data Scraper veri çıkarma aracı kullanılarak toplanan veriler olumlu, olumsuz ve nötr olarak etiketlenerek sınıflandırılmıştır. Bu çalışmada “veri dağılımının” ve “öznitelik seçiminin” sınıflandırma üzerindeki etkileri WEKA 3.8.6 yazılımı kullanılarak incelenmiştir. Çalışma sonucunda veri kümeleri arasında dengesiz veri kümesinin, dengeli veri kümesinden daha iyi performans sağladığı gözlemlenmiştir. Dengeli veri kümesinde analiz yapılırken öznitelik seçiminin yapılmasının, öznitelik seçiminin yapılmadığı modellere göre kayda değer bir performans değişikliğine neden olmadığı ve dengesiz veri kümesinde analiz yapılırken öznitelik seçiminin yapılmasının, öznitelik seçiminin yapılmadığı modellere göre daha iyi performans gösterdiği gözlemlenmiştir. Dengesiz veri kümesinde öznitelik seçimi yapılarak elde edilen %94,9 sınıflandırma doğruluğu oranı ile en iyi performans gösteren algoritma Sıralı Minimal Optimizasyon olmuştur. Ayrıca elde edilen bulgular, kullanıcıların büyük çoğunluğunun deprem uygulamalarına yönelik memnuniyet düzeyinin görece yüksek olduğunu ve genel kullanıcı deneyiminin pozitif bir eğilim sergilediğini ortaya koymaktadır.

Anahtar Kelimeler: Veri madenciliği, Duygu analizi, Naive bayes, Sıralı minimal optimizasyon, Deprem uygulamaları.

ABSTRACT

A STUDY ON TEXT MINING AND SENTIMENT ANALYSIS: USER COMMENTS OF EARTHQUAKE APPS

Turkey, due to its geological structure, is a country at high risk of earthquakes. Major disasters such as the 2023 Kahramanmaraş Earthquake have clearly demonstrated the need to increase public awareness of earthquakes. Earthquakes are natural disasters with the potential to cause both loss of life and property, and people want to know about earthquakes in advance for reasons such as saving lives, reducing material damage, and being psychologically prepared. Earthquake apps on phones have become important tools for both providing early warnings and delivering earthquake-related information to users quickly and reliably. User feedback on earthquake apps they use provides crucial data for analyzing and interpreting the sentiment conveyed in the feedback. The information obtained by analyzing this data is crucial for app developers in developing apps and improving the user experience. In this study, sentiment analysis of user comments on the most commented earthquake apps was conducted using a supervised machine learning approach. Supervised machine learning classification algorithms such as Naive Bayes, Sequential Minimal Optimization, and k-nearest neighbor algorithms were used for sentiment analysis. Data collected using the Instant Data Scraper data extraction tool available on Google Play and the App Store were classified as positive, negative, and neutral. In this study, the effects of "data distribution" and "feature selection" on classification were investigated using WEKA 3.8.6 software. The study observed that an imbalanced dataset performed better than a balanced dataset. It was observed that performing feature selection while analyzing a balanced dataset did not cause a significant performance change compared to models without feature selection, and performing feature selection while analyzing an imbalanced dataset performed better than models without feature selection. The best-performing algorithm, with a classification accuracy rate of 94.9% achieved by performing feature selection on an imbalanced dataset, was Sequential Minimal Optimization. Furthermore, the findings reveal that the majority of users' satisfaction with earthquake applications is relatively high and the overall user experience exhibits a positive trend.

Keywords: Data mining, Sentiment analysis, Naive bayes, Sequential minimal optimization, Earthquake applications.

İÇİNDEKİLER

	Sayfa
ÖN SÖZ.....	i
ÖZET.....	ii
ABSTRACT.....	iii
İÇİNDEKİLER.....	iv
TABLolar LİSTESİ.....	vii
ŞEKİLLER LİSTESİ.....	ix
KISALTMALAR VE SİMGELER LİSTESİ.....	x
1. GİRİŞ.....	1
1.1. Araştırmanın Amacı.....	2
1.2. Araştırmanın Önemi.....	2
1.3. Araştırmanın Soruları.....	3
1.4. Araştırmanın Sınırlılıkları.....	3
1.5. Araştırmanın Yöntemi.....	4
1.6. Araştırmanın Organizasyonu.....	5
2. ARAŞTIRMANIN KAVRAMSAL ÇERÇEVESİ.....	6
3. METİN MADENCİLİĞİ.....	7
3.1. Metin Madenciliği Süreci.....	8
4. DUYGU ANALİZİ.....	10
4.1. Duygu Analizi Süreci.....	10
4.2. Duygu Sınıflandırma Yöntemleri.....	12
4.2.1. Sözlük Tabanlı Yöntemler.....	12
4.2.2. Makine Öğrenmesi Tabanlı Yöntemler.....	12
4.2.2.1. Naive Bayes.....	13

4.2.2.2. Sıralı Minimal Optimizasyon.....	13
4.2.2.3. En Yakın Komşu (kNN).....	14
4.2.2. Derin Öğrenme Tabanlı Yöntemler.....	16
5. LİTERATÜR.....	17
5.1. Araştırmanın Literatüre Katkısı.....	20
6. METODOLOJİ.....	22
6.1. Araştırmanın Amacı.....	22
6.2. Destek Araçlar.....	22
6.2.1. Instant Data Scraper.....	22
6.2.2. ITU Turkish NLP Web Servis API.....	22
6.2.3. WEKA 3.8.6.....	22
6.3. Araştırmanın Örneklemi ve Verilerin Analiz için Hazırlanması.....	23
6.4. Model Başarısını Değerlendirme.....	30
6.4.1. Karşıtlık Matrisi.....	30
6.4.2. Sınıflandırma Doğruluğu.....	32
6.4.3. Duyarlılık (Recall)	32
6.4.4. Kesinlik (Precision)	32
6.4.5. Doğru Pozitif Oranı (True Positive Rate)	32
6.4.6. Yanlış Pozitif Oranı (False Positive Rate)	33
6.4.7. F-Ölçütü (F1 Score)	32
6.4.8. Kappa İstatistiği.....	33
6.5. Araştırmanın Bulguları.....	34
6.5.1. Dengeli Veri Kümesi ve Öznitelik Seçiminin Yapılmaması.....	34
6.5.1.1. Naive Bayes Sınıflandırma Yöntemi ile Analiz.....	34

6.5.1.2. SMO Sınıflandırma Yöntemi ile Analiz.....	34
6.5.1.3. kNN ($k=1$) Sınıflandırma Yöntemi ile Analiz.....	35
6.5.2. Dengeli Veri Kümesi ve Öznitelik Seçiminin Yapılması.....	36
6.5.2.1. Naive Bayes Sınıflandırma Yöntemi ile Analiz.....	36
6.5.2.2. SMO Sınıflandırma Yöntemi ile Analiz.....	37
6.5.2.3. kNN ($k=1$) Sınıflandırma Yöntemi ile Analiz.....	37
6.5.3. Dengesiz Veri Kümesi ve Öznitelik Seçiminin Yapılmaması.....	39
6.5.3.1. Naive Bayes Sınıflandırma Yöntemi ile Analiz.....	39
6.5.3.2. SMO Sınıflandırma Yöntemi ile Analiz.....	39
6.5.3.3. kNN ($k=1$) Sınıflandırma Yöntemi ile Analiz.....	40
6.5.4. Dengesiz Veri Kümesi ve Öznitelik Seçiminin Yapılması.....	41
6.5.4.1. Naive Bayes Sınıflandırma Yöntemi ile Analiz.....	41
6.5.4.2. SMO Sınıflandırma Yöntemi ile Analiz.....	42
6.5.4.3. kNN ($k=1$) Sınıflandırma Yöntemi ile Analiz.....	42
6.6. Araştırmanın Çıktıları.....	44
6.7. Tartışma.....	47
6.7.1. Bulguların Literatür ile Karşılaştırılması.....	49
7. SONUÇ.....	51
8. ÖNERİLER.....	53
KAYNAKÇA.....	55
EKLER.....	62

TABLÖLAR LİSTESİ

	Sayfa
Tablo 6.1. Veri Çekilen Uygulamalar.....	23
Tablo 6.2. Türkçe Durak Kelimeler.....	27
Tablo 6.3. Veri Etiketleme ve Veri Önişleme Örnekleri.....	28
Tablo 6.4. Karşıtlık (Hata) Matrisi.....	31
Tablo 6.5. Naive Bayes Sınıflandırma Yöntemine ait Hata Matrisi.....	34
Tablo 6.6. Naive Bayes ile Yapılan Modelin Değerlendirme Ölçütleri.....	34
Tablo 6.7. SMO Sınıflandırma Yöntemine ait Hata Matrisi.....	35
Tablo 6.8. SMO ile Yapılan Modelin Değerlendirme Ölçütleri.....	35
Tablo 6.9. kNN ($k=1$) (Chebyshev) Sınıflandırma Yöntemine ait Hata Matrisi.....	35
Tablo 6.10. kNN ($k=1$) (Chebyshev) ile Yapılan Modelin Değerlendirme Ölçütleri...	35
Tablo 6.11. kNN ($k=1$) (Öklid) ile Sınıflandırma Yöntemine ait Hata Matrisi.....	36
Tablo 6.12. kNN ($k=1$) (Öklid) ile Yapılan Modelin Değerlendirme Ölçütleri.....	36
Tablo 6.13. Naive Bayes Sınıflandırma Yöntemine ait Hata Matrisi.....	36
Tablo 6.14. Naive Bayes ile Yapılan Modelin Değerlendirme Ölçütleri	37
Tablo 6.15. SMO Sınıflandırma Yöntemine ait Hata Matrisi.....	37
Tablo 6.16. SMO ile Yapılan Modelin Değerlendirme Ölçütleri.....	37
Tablo 6.17. kNN ($k=1$) (Chebyshev) Sınıflandırma Yöntemine ait Hata Matrisi.....	38
Tablo 6.18. kNN ($k=1$) (Chebyshev) ile Yapılan Modelin Değerlendirme Ölçütleri...	38
Tablo 6.19. kNN ($k=1$) (Öklid) ile Sınıflandırma Yöntemine ait Hata Matrisi.....	38
Tablo 6.20. kNN ($k=1$) (Öklid) ile Yapılan Modelin Değerlendirme Ölçütleri.....	38
Tablo 6.21. Naive Bayes Sınıflandırma Yöntemine ait Hata Matrisi.....	39
Tablo 6.22. Naive Bayes ile Yapılan Modelin Değerlendirme Ölçütleri.....	39
Tablo 6.23. SMO Sınıflandırma Yöntemine ait Hata Matrisi.....	39
Tablo 6.24. SMO ile Yapılan Modelin Değerlendirme Ölçütleri.....	40
Tablo 6.25. kNN ($k=1$) (Chebyshev) Sınıflandırma Yöntemine ait Hata Matrisi.....	40

Tablo 6.26. kNN ($k=1$) (Chebyshev) ile Yapılan Modelin Değerlendirme Ölçütleri...	40
Tablo 6.27. kNN ($k=1$) (Öklid) ile Sınıflandırma Yöntemine ait Hata Matrisi.....	40
Tablo 6.28. kNN ($k=1$) (Öklid) ile Yapılan Modelin Değerlendirme Ölçütleri.....	41
Tablo 6.29. Naive Bayes Sınıflandırma Yöntemine ait Hata Matrisi.....	41
Tablo 6.30. Naive Bayes ile Yapılan Modelin Değerlendirme Ölçütleri	41
Tablo 6.31. SMO Sınıflandırma Yöntemine ait Hata Matrisi.....	42
Tablo 6.32. SMO ile Yapılan Modelin Değerlendirme Ölçütleri.....	42
Tablo 6.33. kNN ($k=1$) (Chebyshev) Sınıflandırma Yöntemine ait Hata Matrisi.....	43
Tablo 6.34. kNN ($k=1$) (Chebyshev) ile Yapılan Modelin Değerlendirme Ölçütleri...	43
Tablo 6.35. kNN ($k=1$) (Öklid) ile Sınıflandırma Yöntemine ait Hata Matrisi.....	43
Tablo 6.36. kNN ($k=1$) (Öklid) ile Yapılan Modelin Değerlendirme Ölçütleri.....	43
Tablo 6.37. Analiz Sonuçlarının Karşılaştırılması.....	44

ŞEKİLLER LİSTESİ

	Sayfa
Şekil 2.1. Çalışmanın Modeli.....	6
Şekil 3.2. Metin Madenciliğinin İlişkili Olduğu Alanlar.....	8
Şekil 4.3: CRISP-DM Referans Modelinin Evreleri.....	11
Şekil 6.4: Olumlu Yorumlarda En Çok Kullanılan 10 Kelime	25
Şekil 6.5: Olumsuz Yorumlarda En Çok Kullanılan 10 Kelime	25
Şekil 6.6: İTÜ Türkçe Doğal Dil İşleme Aracı ile Metin Normalizasyonu.....	26
Şekil 6.7. Dengeli ve Dengesiz Veri Kümeleri.....	29

KISALTMALAR VE SİMGELER LİSTESİ

ARFF: Attribute Relation File Format (Öznitelik İlişkisi Dosya Biçimi)

CRISP-DM: Cross Industry Standard Process for Data Mining (Veri Madenciliği için Endüstriler Arası Standart Süreç)

KNN: K-Nearest Neighbors (K-En Yakın Komşu)

LP: Logistic Regression (Lojistik Regresyon)

MLP: Multi-Layer Perceptron (Çok Katmanlı Algılayıcı)

NB: Naives Bayes

NLP: Natural Language Processing (Doğal Dil İşleme)

QP: Quadratic Programming (İkinci Dereceden-Karmaşık Programlama)

SMO: Sequential Minimal Optimization (Sıralı Minimal Optimizasyon)

SVM: Support Vector Machine (Destek Vektör Makinesi)

WEKA: Waikato Environment for Knowledge Analysis (Waikato Bilgi Analizi Ortamı)

1.GİRİŞ

Depremeler hem can hem de mal kaybına yol açabilecek potansiyele sahip ciddi doğal afetlerdir ve ülkemizin jeolojik yapısı itibarıyla deprem riski oldukça yüksektir (Türkoğlu, 2001; USGS, 2021). İnsanlar için bir depremi önceden bilmek, en önemli ve en basit haliyle güvenli bir şekilde tahliye olabilmeleri için bir fırsat sağlar (Allen vd., 2009). İkincil olarak da insanlar deprem gibi doğal afetler konusunda bilinçli olmak ve hazırlıklı olmak isterler. Bu, paniği azaltabilir ve insanlara bir miktar kontrol hissi verebilir, evlerini ve iş yerlerini güçlendirebilir, değerli eşyalarını güvenli bir yere taşıyabilir ya da altyapıyı güçlendirebilirler. Teknolojinin gelişmesiyle birlikte geliştirilen deprem uygulamaları insanların depremlerle ilgili bilgi sağladıkları ve erken uyarılar almak için kullandıkları etkili araçlar haline gelmiştir (Rahadian vd., 2018). Deprem uygulamaları, özellikle bazı ülkelerde, depremler meydana geldiğinde hızlı bir şekilde uyarı yapabilir. Bu tür uygulamalar, deprem dalgalarının daha hızlı ulaşan P-waves (ön dalgalar) ile yıkıcı S-waves (ana dalgalar) arasında geçen süreyi hesaplayarak kullanıcılara birkaç saniye önceden uyarı sağlayabilir (Allen ve Melgar, 2019). Eğer bir deprem algılanırsa, kullanıcılar birkaç saniye önceden uyarı alabilir, bu da onları güvenli bir yere geçme ya da güvenlik önlemleri alma konusunda bilgilendirir. Deprem meydana geldikten sonra, bu tür uygulamalar, depremin büyüklüğünü, merkez üssünü, derinliğini ve zamanını bildirir. Ayrıca, depremle ilgili resmî açıklamalar, afet raporları ve kurtarma bilgileri de sunulabilir. Deprem uygulamaları, kullanıcıların yakınlarıyla iletişim kurmalarına yardımcı olabilir. Bazı uygulamalar, "güvendedim" butonu gibi özelliklerle, kullanıcıların ailelerine ya da arkadaşlarına, depremden sonra güvende olduklarını bildirmelerini sağlar. Ayrıca, acil durum kitleri ve güvenli alanlar hakkında rehberlik de verebilirler.

Rogers'ın Korunma Motivasyonu Teorisi bireylerin risk algısı ve koruyucu davranış motivasyonlarını açıklamak için kullanılan önemli bir psikolojik çerçevedir (Rogers, 1975). Teoriye göre, bireyler bir tehdidin ciddiyetini ve kendi etkileniş olasılıklarını değerlendirdikten sonra, bu tehditle başa çıkmak için koruyucu davranışlarda bulunma motivasyonuna sahiptir. Deprem uygulamaları, kullanıcıların deprem riskini somut olarak algılamalarını sağlamakta, erken uyarı ve güvenlik rehberliği gibi özellikler aracılığıyla bireylerin kapasite algısını güçlendirmektedir. Böylece, Korunma Motivasyonu Teorisi çerçevesinde, bu uygulamalar hem tehdit algısını hem de koruyucu davranış motivasyonunu artırarak, bireylerin deprem öncesi ve sırasındaki hazırlık düzeylerini yükseltmektedir.

Deprem uygulamaları, kullanıcılara deprem güvenliği hakkında deprem öncesinde, sırasında ve sonrasında nasıl hareket etmeleri gerektiği konusunda bilgiler sunar. Ayrıca bu uygulamalar deprem hazırlığı yapmaya yönelik hatırlatmalar da yapabilir. Uygulamalar, kullanıcılara bölgesel deprem aktivitesini gösteren haritalar ve grafikler sunarak, mevcut sismik hareketler hakkında bilgi verir. Bu, insanların çevrelerindeki deprem olaylarıyla ilgili daha fazla bilgi sahibi olmalarını sağlar. Sonuç olarak, telefonlardaki deprem uygulamaları hem erken uyarı sağlamak hem de depremle ilgili bilgileri hızlı ve güvenilir bir şekilde kullanıcılara iletmek için önemli ve etkili araçlar haline gelmişlerdir.

Kullanıcıların, kullandıkları deprem uygulamaları hakkındaki geri bildirimleri, uygulama ile ilgili yaşadıkları deneyimleri anlamamızı sağlar. Geri bildirimler, taşıdıkları duygunun metin madenciliği yöntemleriyle analiz edilebilmesi ve yorumlanabilmesi açısından önemli bir veri kaynağı oluşturmaktadırlar. Metin madenciliği çalışmaları, metni veri kaynağı olarak kabul eden veri madenciliği (data mining) çalışmasıdır (Şeker, 2015). Metin madenciliğinin önemli bir alanı olan duygu analizi, elde edilen verilerin yorumlanması ve uygulamaların kullanıcı odaklı geliştirilmesi için kritik öneme sahiptir (Pang & Lee, 2008).

1.1. Araştırmanın Amacı

Bu çalışmanın temel amacı, deprem uygulamalarına ilişkin kullanıcı yorumlarını metin madenciliği ve duygu analizi yöntemleriyle inceleyerek, kullanıcıların algılarını, memnuniyet düzeylerini ve eleştirel geri bildirimlerini sistematik olarak ortaya koymaktır. Böylelikle, uygulamaların güçlü ve zayıf yönleri belirlenerek, kullanıcı deneyiminin iyileştirilmesine ve afet yönetiminde dijital çözümlerin daha etkin hale getirilmesine katkı sağlanması hedeflenmektedir.

1.2. Araştırmanın Önemi

Doğal afetler, özellikle depremler, toplumun hızlı bilgiye erişimini ve güvenilir yönlendirmelere ulaşmasını zorunlu kılmaktadır. Bu bağlamda geliştirilen mobil uygulamalar, afet anında ve sonrasında kritik bir rol üstlenmektedir. Ancak bu uygulamaların başarısı, yalnızca teknik yeterlilikleriyle değil, aynı zamanda kullanıcıların deneyim ve geri bildirimleriyle de doğrudan ilişkilidir. Kullanıcı yorumlarının sistematik olarak analiz edilmesi, uygulamaların etkililiğini artıracak geliştirme alanlarını ortaya koymak açısından büyük önem taşımaktadır. Bu araştırmanın önemini, iki boyutta değerlendirebiliriz:

Akademik Katkı: Metin madenciliği ve duygu analizi yöntemlerinin afet yönetimi bağlamında uygulanması, literatürde sınırlı sayıda ele alınmış bir konu olup, çalışmanın özgün bilimsel katkısını oluşturmaktadır.

Toplumsal Katkı: Deprem uygulamalarına yönelik kullanıcı deneyimlerinin analiz edilmesi, afet hazırlık, müdahale ve iyileştirme süreçlerinde daha etkin dijital çözümlerin geliştirilmesine zemin hazırlamakta ve toplumun afetlere karşı direncini artırmaktadır.

1.3. Araştırmanın Soruları

Bu çalışmada temel olarak veri dağılımının ve öznitelik seçmenin sınıflandırma üzerindeki etkisi araştırılmaktadır. “Veri kümesinin dengeli veya dengesiz tercih edilmesi sınıflandırma algoritmalarının performanslarını etkiler mi? ve “Öznitelik seçiminin yapıp yapılmaması sınıflandırma algoritmalarının performanslarını etkiler mi?” temel sorularına ek olarak “Naive Bayes, SMO ve kNN algoritmaları, deprem uygulamalarına ilişkin kullanıcı yorumlarını sınıflandırmada hangi performans düzeylerine sahiptir?”, “Sınıflandırma doğruluğu, kappa istatistiği, duyarlılık, kesinlik ve F-ölçütü gibi metrikler açısından algoritmalar arasında anlamlı farklar var mıdır?”, “Hangi algoritma, Türkçe kullanıcı yorumlarının dilsel özellikleri karşısında daha yüksek genellenebilirlik göstermektedir?” ve “Kullanıcıların deprem uygulamalarına yönelik memnuniyet düzeyleri nedir?” sorularına da yanıt aranmıştır.

1.4. Araştırmanın Sınırlılıkları

Araştırmanın sınırlılıkları çeşitli boyutlarda ele alınabilir. İlk olarak, *veri kaynağı sınırlılığı* söz konusudur. Çalışmada yalnızca deprem uygulamalarına yapılan kullanıcı yorumları incelenmiş olup, sosyal medya, forumlar veya haber sitelerindeki paylaşımlar değerlendirmeye alınmamıştır. Bu durum, elde edilen bulguların yalnızca belirli bir platform türüne özgü olmasına neden olmaktadır.

İkinci olarak, *dil sınırlılığı* bulunmaktadır. Araştırmaya yalnızca Türkçe dilinde yazılmış yorumlar dahil edilmiş, diğer dillerdeki kullanıcı geri bildirimleri kapsam dışında bırakılmıştır. Bu nedenle, farklı kültür ve dil gruplarına ait kullanıcı deneyimleri ve tutumları bu çalışmada temsil edilmemektedir.

Üçüncü olarak, *zaman sınırlılığı* dikkate alınmalıdır. Yorumlar belirli bir tarih aralığında (09.03.2025 – 19.03.2025) toplanmıştır. Farklı dönemlerde meydana gelen depremler veya güncellemeler sonrasında yapılan kullanıcı yorumlarının analize dahil edilmemesi, bulguların zamana bağlı değişkenlik göstermesi ihtimalini artırmaktadır.

Dördüncü sınırlılık, *algoritma sınırlılığıdır*. Çalışmada yalnızca üç denetimli öğrenme algoritması (Naive Bayes, SMO ve kNN) değerlendirilmiştir. Literatürde yaygın olarak kullanılan diğer algoritmalar (örneğin Random Forest, Decision Tree, Support Vector Machine'in farklı varyasyonları veya Derin Öğrenme yöntemleri) kapsam dışı bırakılmıştır. Bu durum, daha geniş algoritma karşılaştırmalarının yapılabilmesi açısından çalışmanın kapsamını daraltmaktadır.

Beşinci olarak, *veri dengesi sınırlılığı* söz konusudur. Veri kümesinde olumlu, olumsuz ve nötr yorumların dağılımı tam anlamıyla dengeli değildir. Bu durum, özellikle dengesiz sınıflarda sınıflandırma performansını olumsuz yönde etkileyebilmekte ve elde edilen doğruluk oranlarının sınıflar arası farklılık göstermesine neden olabilmektedir.

Altıncı sınırlılık, *öznel yorum sınırlılığıdır*. Kullanıcı yorumları bireylerin öznel deneyimlerine dayalı olduğundan, yorumların duygusal tonu, kullanılan dil ve bağlamdan bağımsız olarak değişkenlik gösterebilir. Bu durum, duygu analizi sürecinde belirli bir düzeyde belirsizlik ve yorum farklılığına yol açabilmektedir.

Son olarak, *araç sınırlılığı* da araştırmanın kapsamını etkilemektedir. Tüm analizler yalnızca WEKA 3.8.6 yazılımı kullanılarak gerçekleştirilmiştir. Farklı yazılımlar (Python, R, RapidMiner vb.) ya da bu yazılımlar üzerinde çalışan farklı kütüphaneler (Scikit-learn, TensorFlow, PyTorch vb.) kullanılsaydı, algoritmaların farklı implementasyonları ve optimizasyon teknikleri sayesinde farklı sonuçlar elde edilebilirdi.

Bu sınırlılıklar, araştırma bulgularının belirli koşullar altında değerlendirilebileceğini ve genellenebilirlik açısından dikkatli olunması gerektiğini ortaya koymaktadır. Ancak aynı zamanda, gelecekte yapılacak çalışmalara farklı veri kaynakları, daha geniş algoritma yelpazesi ve farklı yazılım ortamlarıyla yeni araştırma imkânları sunduğu da söylenebilir.

1.5. Araştırmanın Yöntemi

Bu çalışmada, denetimli makine öğrenmesi yaklaşımı kullanılarak en çok yorum alan deprem uygulamalarına yapılan kullanıcı yorumlarının duygu analizi yapılmıştır. Bu çalışma için veriler Google Play ve App Store'da yer alan en çok indirilen ve en çok değerlendirme yapılan 8 deprem uygulamasına (Deprem Ağı, Deprem Bilgi Sistemi, Deprem Uyarılarım, Last Quake, Deprem Ağı Pro, Deprem Türkiye, Deprem Uyarılarım Pro ve Deprem – Earthquake) ait kullanıcı yorumları Instant Data Scraper veri çıkarma aracıyla toplanmıştır. İlk oluşturulan veri kümesi olumlu, olumsuz ve nötr olarak etiketlenmiştir. Veriler analiz sürecine hazırlanmış, olumlu ve olumsuz etiketli verilerden dengeli ve dengesiz veri kümeleri oluşturulmuş ve veri

kümeleri WEKA 8.6.3 programında Navie Bayes, SMO ve kNN sınıflandırma algoritmaları kullanılarak analiz edilmiştir.

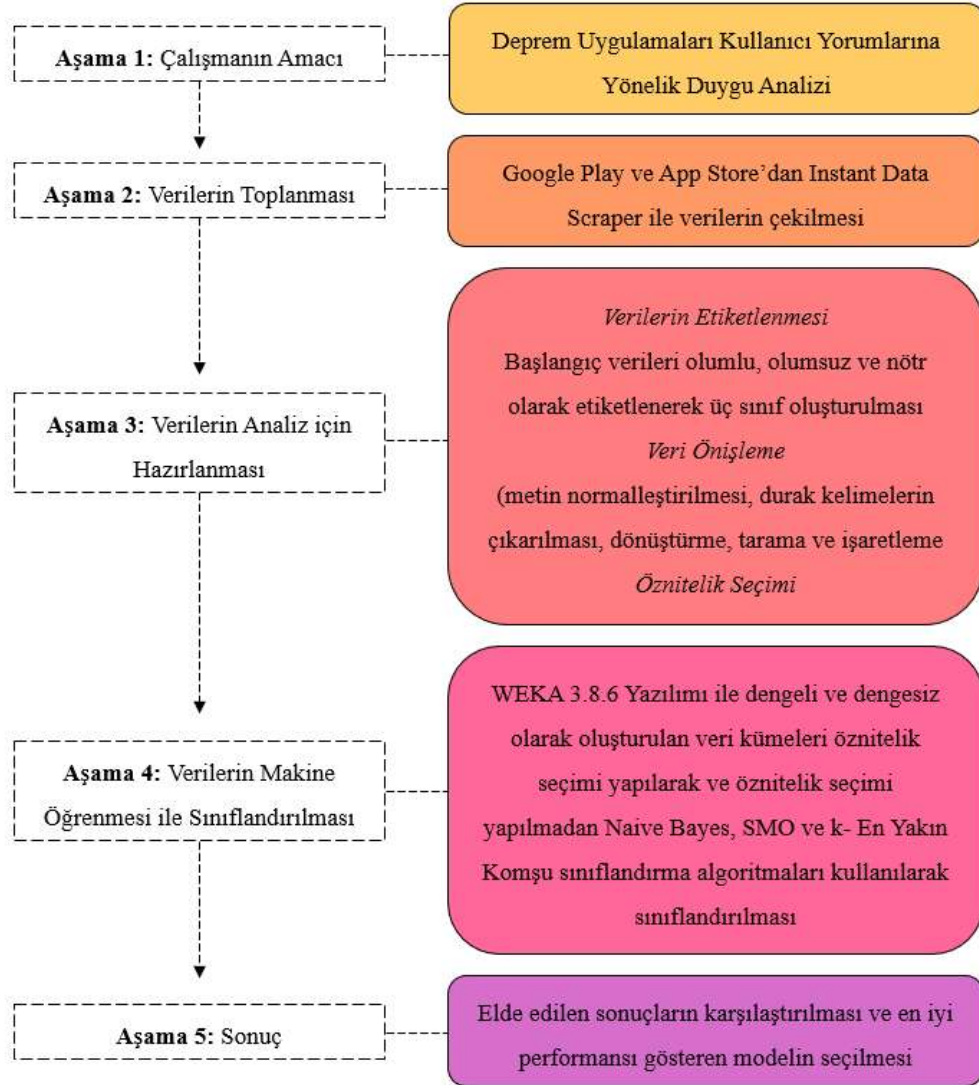
1.6. Araştırmanın Organizasyonu

Çalışmamız 8 bölümden oluşmaktadır. Araştırmanın Kavramsal Çerçevesi 2. bölümde çizilmiştir. Metin Madenciliği (3. bölüm) ve Duygu Analizi (4. bölüm) hakkında bilgi verilerek 5. bölümde Literatürde yer alan diğer çalışmalar incelenmiştir. Araştırmanın amacı, destek araçlar, araştırmanın örnekleme ve verilerin analiz için hazırlanması, araştırmanın bulguları, çıktıları ve tartışma kısımlarının yer aldığı Metodoloji (6. bölüm), Sonuç (7. bölüm) ve Öneriler (8. bölüm) bölümleriyle çalışma tamamlanmıştır.



2. ARAŞTIRMANIN KAVRAMSAL ÇERÇEVESİ

Çalışmanın modeli Şekil 2.1'de gösterildiği gibi 5 aşamadan oluşmaktadır. Birinci aşama, çalışmanın genel amacıdır. 2. Aşama, Google Play ve App Store'da yer alan kullanıcı yorumlarının Instant Data Scraper veri çıkarma aracı ile toplanmasıdır. Toplanan bu verilerin etiketlenmesi (olumlu, olumsuz ve nötr olarak etiketlenerek üç sınıf oluşturulması), veri önışlemenin yapılması (metin normalleştirilmesi, durak kelimelerin çıkarılması, dönüştürme, tarama ve işaretleme) ve öznitelik seçimini süreçleri, verilerin analiz için hazırlanmasının yer aldığı 3. Aşamadır. Verilerin makine öğrenmesi ile sınıflandırılması (4. Aşama) ise WEKA 3.8.6 Yazılımı ile dengeli ve dengesiz olarak oluşturulan veri kümeleri öznitelik seçimi yapılarak ve öznitelik seçimi yapılmadan Naive Bayes, SMO ve k- En Yakın Komşu sınıflandırma algoritmaları kullanılarak sınıflandırılmasını sağlayarak, elde edilen sonuçlar karşılaştırılacak ve en iyi performansı gösteren model seçilecektir (Aşama 5).



Şekil 2.1. Çalışmanın Modeli

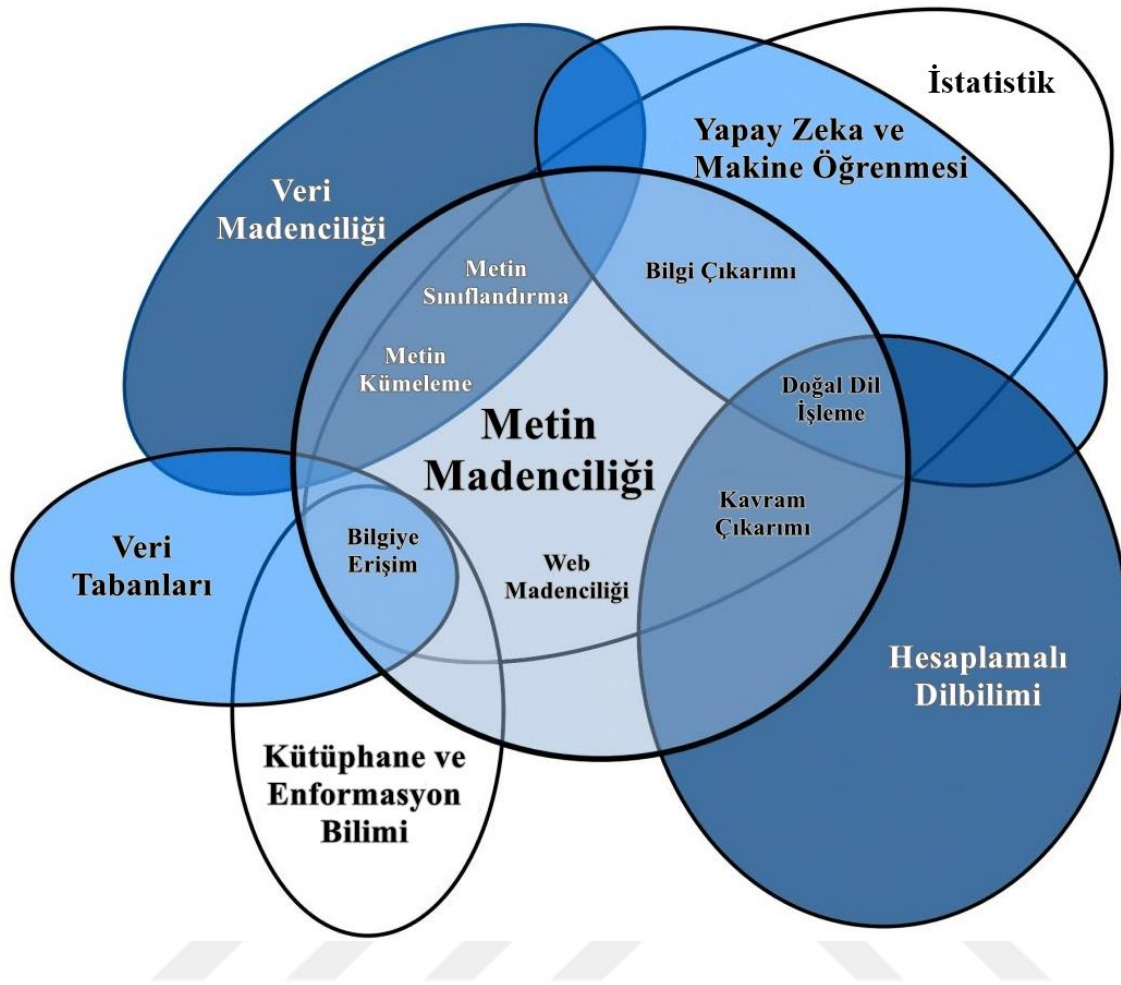
3. METİN MADENCİLİĞİ

Günümüzde bilgi üretimi ve paylaşımı büyük ölçüde dijital ortamlarda gerçekleşmektedir. İnternet, sosyal medya platformları, çevrim içi haber siteleri, e-posta arşivleri ve kurumsal veri tabanları her gün milyarlarca metin tabanlı veri üretmektedir. Bu büyük hacimli veriler, yapılandırılmamış formda olduklarından klasik veri işleme yöntemleri ile doğrudan analiz edilemez. Bu noktada devreye giren metin madenciliği (text mining), metinlerden anlamlı bilgilerin çıkarılmasını sağlayan disiplinler arası bir araştırma alanı olarak önem kazanmıştır.

Metin madenciliği, yapılandırılmamış veya yarı-yapılandırılmış metin verilerinden istatistiksel, dilbilimsel ve yapay zekâ temelli yöntemlerle bilgi çıkarımı yapmayı amaçlayan bir süreçtir (Feldman ve Sanger, 2006). Başka bir ifadeyle metin madenciliği, büyük ölçekli metin verilerinin işlenebilir hale getirilmesi ve bunlardan örüntülerin, ilişkilerin, eğilimlerin keşfedilmesi sürecidir.

Bilgiye erişim (IR:Information Retrieval) ve bilgi çıkarımı (IE:Information Extraction) metin madenciliğinin iki köküdür (Miner vd., 2012). Metin madenciliğinin ilk aşaması olan bilgiye erişim, bireyin belirli bir bilgi ihtiyacını karşılamak amacıyla, yapısal açıdan düzensiz nitelik taşıyan metin veya belgelerin geniş bir kaynak havuzu içerisinde belirlenmesi sürecidir. Bilgi çıkarımı ise büyük ölçekli veri yığınları içerisinde anlamlı ve özet bilgilerin elde edilmesi sürecidir. Akademik makalelerden araştırmacının belirlediği anahtar kavramlarla ilişkili çalışmaların otomatik olarak çıkarılması veya sosyal medya paylaşımlarından, “deprem” anahtar kelimesi girildiğinde ilgili içeriklerin ve öne çıkan duygu eğilimlerinin belirlenmesi bilgi çıkarımı örnekleridir.

Miner vd. (2012) metin madenciliğinin; veri madenciliği, yapay zekâ ve makine öğrenmesi, veri tabanları, kütüphane ve enformasyon bilimi ve hesaplamalı dilbilimi alanlarıyla yakından ilişki içerisinde olduklarını ve bu ilişkilerden metin sınıflandırma, metin kümeleme, bilgiye erişim, bilgi çıkarımı, kavram çıkarımı, doğal dil işleme ve web madenciliği gibi uygulama alanları meydana gelmiştir. Aşağıdaki Venn şemasında (Şekil 3.2) metin madenciliğinin diğer alanlarla kesişimi ve bu kesişme noktalarında bulunan metin madenciliği uygulama alanları gösterilmektedir.



Şekil 3.2. Metin Madenciliğinin İlişkili Olduğu Alanlar

Kaynak: Miner vd. (2012: 31)

3.1. Metin Madenciliği Süreci

Metin madenciliği süreci, yapılandırılmamış metin verilerinden anlamlı ve işlevsel bilgilerin elde edilmesini sağlayan sistematik ve çok aşamalı bir yöntemler bütünüdür. Bu süreç genel olarak dört temel aşamadan oluşmaktadır:

Metin Ön İşleme (Text Preprocessing): Analize uygun veri yapısı oluşturmak amacıyla gerçekleştirilen bu aşamada, durak kelimelerin (stopwords) çıkarılması, noktalama işaretlerinin temizlenmesi, kök veya gövde bulma (stemming, lemmatization) gibi dilsel normalizasyon işlemleri uygulanmaktadır (Hotho vd., 2005).

Özellik Çıkarımı ve Temsili (Feature Extraction and Representation): Ön işlenmiş metinlerin sayısal forma dönüştürülmesiyle analiz edilebilir bir yapı elde edilmektedir. Bu bağlamda en sık kullanılan yöntemler arasında Bag-of-Words ve TF-IDF bulunmakta, ayrıca

Word2Vec, GloVe ve BERT gibi ileri düzey temsil teknikleri de çağdaş yaklaşımlar arasında yer almaktadır (Aggarwal ve Zhai, 2012).

Modelleme ve Analiz (Modeling and Analysis): Sayısal olarak temsil edilen metin verileri üzerinde denetimli veya denetimsiz makine öğrenmesi yöntemleri uygulanmaktadır. Bu aşamada sınıflandırma, kümeleme, duygu analizi ve konu modelleme gibi tekniklerden yararlanılmaktadır (Feldman ve Sanger, 2006).

Sonuçların Yorumlanması (Interpretation of Results): Analizler sonucunda elde edilen bulgular, karar verme süreçlerine, bilimsel araştırmalara veya uygulama alanlarına entegre edilerek anlamlandırılmakta ve değerlendirilmektedir.

Sonuç olarak, metin madenciliği süreci bu aşamaların bütüncül ve etkileşimli bir biçimde ilerlemesiyle, büyük ölçekli metin verilerinden stratejik, bilimsel ve pratik açıdan değerli bilgilerin ortaya çıkarılmasını mümkün kılmaktadır.

4. DUYGU ANALİZİ

Duygu analizi, metin madenciliğinin en önemli alt alanlarından biri olarak, bireylerin belirli bir konu, olay veya nesneye ilişkin duygu, tutum, kanaat ve deneyimlerinin sistematik olarak incelenmesini amaçlamaktadır (Farhadloo ve Rolland, 2016). Literatürde bu kavram; *duygu durumu analizi*, *duygu sınıflandırması*, *fikir madenciliği* ve *kanaat çıkarımı* gibi farklı terimlerle de anılmaktadır. Bu bağlamda duygu analizi, yalnızca metinlerin içerdiği sözcüklerin yüzeysel değerlendirilmesiyle sınırlı kalmayıp, aynı zamanda bağlam, sözdizimsel yapı ve anlamsal ilişkiler üzerinden bireylerin ifade ettikleri duyguların daha doğru şekilde anlaşılmasını hedeflemektedir.

Agarwal vd. (2011) göre duygu analizi, metinlerde yer alan duyguların otomatik biçimde çıkarılması ve sınıflandırılmasına yönelik olarak doğal dil işleme (NLP), metin analizi ve hesaplamalı yöntemlerin bütünleşik biçimde kullanılmasını içermektedir. Bu yöntemler sayesinde duygu analizi; sosyal medya içeriklerinin incelenmesi, müşteri geri bildirimlerinin değerlendirilmesi, politik eğilimlerin belirlenmesi, sağlık alanında psikolojik duygu durumlarının tespiti ve pazarlama stratejilerinin geliştirilmesi gibi çok geniş bir uygulama alanına sahiptir (Liu, 2012).

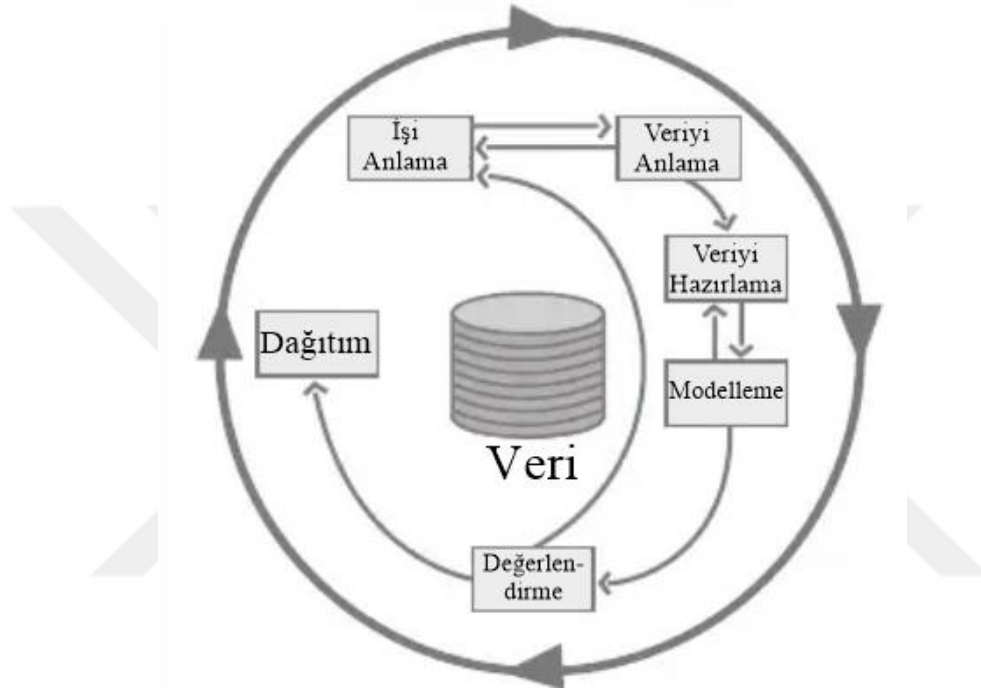
Sonuç olarak duygu analizi, günümüzde hem akademik çalışmalarda hem de endüstriyel uygulamalarda karar verme süreçlerini destekleyen güçlü bir araç olarak kabul edilmektedir.

4.1. Duygu Analizi Süreci

Duygu analizi, kullanıcıların metin tabanlı ifadelerinden (örneğin sosyal medya paylaşımları, yorumlar veya geri bildirimler) olumlu, olumsuz ya da nötr tutumlarının ortaya çıkarılmasını amaçlayan doğal dil işleme (NLP) tabanlı bir yöntemdir. Günümüzde özellikle afet yönetimi, e-ticaret ve kamu hizmetleri gibi kritik alanlarda kullanıcıların duygu ve düşüncelerinin sistematik olarak analiz edilmesi, karar vericilere önemli katkılar sağlamaktadır. Bu bağlamda, duygu analizi sürecinin sistematik ve tekrarlanabilir bir şekilde yürütülmesi için CRISP-DM (Cross Industry Standard Process for Data Mining) süreç modeli uygun bir çerçeve sunmaktadır. Bu çalışmada, duygu analizi CRISP-DM süreç modelinin aşamaları temel alınarak gerçekleştirilecek ve böylece hem metodolojik tutarlılık hem de akademik geçerlilik sağlanacaktır.

CRISP-DM (Cross Industry Standard Process for Data Mining), veri madenciliği projelerinin planlanması, yürütülmesi ve yönetilmesi için geliştirilmiş, teknolojiden bağımsız bir süreç modelidir. Bu model, herhangi bir sektörde uygulanabilir olup, veri madenciliği

projelerini daha kolay, hızlı, tekrarlanabilir ve yönetilebilir hale getirmeyi amaçlamaktadır (Mariscal vd., 2010; Wirth ve Hipp, 2000). CRISP-DM, bir veri madenciliği projesinde gerçekleştirilmesi gereken faaliyetleri tanımlamakta, neyin yapılması gerektiğini ve hangi adımlardan kaçınılması gerektiğini ortaya koymaktadır. Model, veri madenciliği projeleri için altı aşamadan oluşan döngüsel bir yaşam döngüsü tanımlar (Şekil 4.3). Bu aşamalar ardışık olarak ilerleyebileceği gibi, proje gereksinimlerine bağlı olarak tekrarlı biçimde de uygulanabilir (Wirth ve Hipp, 2000).



Şekil 4.3. CRISP-DM Referans Modelinin Evreleri

Kaynak: Chapman vd., (2000: 9)

Sürecin temel aşamaları Chapman vd. (2002), Azevedo ve Santos (2008) ve Wirth ve Hipp (2000) tarafından şu şekilde özetlenmiştir:

İş Anlama: Projenin hedeflerinin ve gereksinimlerinin işletme perspektifinden anlaşılmasına odaklanır. Bu aşamada veri bilimci, işletme hedeflerini veri madenciliği problemine dönüştürerek bir proje planı geliştirir.

Veri Anlama: İlk veri toplama faaliyetleri ile veriye aşinalık kazanılması sağlanır. Veri bilimci, veriler üzerinde ön analizler gerçekleştirerek ilk hipotezleri oluşturur.

Veri Hazırlama: Modelleme aşamasında kullanılacak uygun bir veri kümesi oluşturulur. Bu aşamada veri seçimi, temizleme ve dönüştürme işlemleri gerçekleştirilir. Süreç, önceden belirlenmiş bir sıraya bağlı olmaksızın birden çok kez yinelenabilir.

Modelleme: Uygun modelleme teknikleri belirlenir, farklı yöntemler karşılaştırılır ve parametreler optimize edilir. Bu aşamada veri kümesinde eksiklikler veya sorunlar fark edildiğinde önceki aşamalara dönülmesi gerekebilir.

Değerlendirme: Oluşturulan modellerin doğruluğu ve iş hedefleriyle uyumluluğu test edilir. Veri bilimci, modelin önceden tanımlanan ölçütleri karşılayıp karşılamadığını değerlendirir.

Dağıtım: Geliştirilen model, kullanıcıların faydalanabileceği şekilde uygulanır. Bu aşama, yalnızca bir rapor hazırlanmasından, tüm kuruluş genelinde kullanılabilecek tekrarlanabilir bir sistemin (ör. e-ticaret öneri sistemi) entegrasyonuna kadar geniş bir yelpazeyi kapsayabilir.

CRISP-DM aşamaları genel olarak ardışık ilerlemekle birlikte, model döngüsel bir yapıya sahiptir ve geriye dönüşlere izin verir. Ancak yinelemelerin nasıl ve hangi koşullarda gerçekleştirileceğine ilişkin kesin bir süreç tanımı bulunmamaktadır (Wirth ve Hipp, 2000).

4.2. Duygu Sınıflandırma Yöntemleri

Duygu sınıflandırma, metin madenciliği ve doğal dil işleme (NLP) alanlarında önemli bir araştırma konusudur. Amaç, bireylerin yazılı ifadelerinden duygu durumlarını (ör. olumlu, olumsuz, nötr) otomatik olarak belirlemektir. Literatürde bu amaçla kullanılan yöntemler genel olarak sözlük tabanlı yöntemler, makine öğrenmesi tabanlı yöntemler ve derin öğrenme tabanlı yöntemler olmak üzere üç grupta incelenmektedir (Pang ve Lee, 2008; Liu, 2012).

4.2.1. Sözlük Tabanlı Yöntemler

Bu yaklaşımda, duygu yönelimi belirli bir duygu sözlüğü aracılığıyla hesaplanır. Metinde geçen kelimeler, pozitif veya negatif anlam içeren kelime listeleriyle eşleştirilir (Taboada vd., 2011). Bu yöntem hızlı ve basit olmakla birlikte, bağlamı dikkate almaması nedeniyle ironi, sarkazm veya çok anlamlılık durumlarında sınırlı kalmaktadır.

4.2.2. Makine Öğrenmesi Tabanlı Yöntemler

Makine öğrenmesi tabanlı yöntemler, duygu sınıflandırmada en yaygın kullanılan yaklaşımlardan biridir. Bu yöntemlerde etiketlenmiş veri kümeleri üzerinde modeller eğitilir ve yeni metinlerin duygu sınıfları tahmin edilir.

Makine öğrenmesi tabanlı yöntemler, öznitelik seçimi ve uygun parametre optimizasyonu ile daha yüksek doğruluk, duyarlılık ve kesinlik değerleri elde edebilir. Naive Bayes (NB), Destek Vektör Makineleri (SVM/SMO), k-En Yakın Komşu (kNN) ve Karar Ağaçları / Rastgele Ormanlar en yaygın kullanılan algoritmalar arasındadır. Literatürde özellikle SMO'nun, duygu sınıflandırma görevlerinde genellikle diğer geleneksel algoritmalarından daha iyi performans gösterdiği belirtilmektedir (Pang ve Lee, 2008).

4.2.2.1. Naive Bayes (NB)

Duygu analizi sürecinde yararlandığımız denetimli makine öğrenme yaklaşımlarından bir olan Naive Bayes, olasılıksal bir sınıflandırma algoritmasıdır. Naive Bayes sınıflandırıcılarının olasılıksal modeli Bayes teoremine dayanır ve “Naive” (naif/saf) sıfatı, bir veri kümesindeki özelliklerin karşılıklı olarak bağımsız olduğu varsayımından gelir (Raschka, 2014). Yani, bir veri noktasının belirli bir sınıfa ait olma olasılığını hesaplarken her özneliğin diğerlerinden bağımsız olduğunu kabul ederek en yüksek olasılığı olan sınıfa atar. Büyük boyutlu (çok öznelikli) veri kümelerinde (metin madenciliği, spam filtreleme, hastalık tahmini gibi) etkilidir. Bayes modeli basit bir modeldir. Bu basitliği ile birlikte, karışık diğer yöntemlere oranla gayet iyi sonuçlar vermektedir. Bu sebeple de yaygın bir kullanım alanı bulmuştur (Arslan, 2018). Naive Bayes algoritmasının basitliği, uygulanabilirliği ve büyük veri setlerinde sağladığı doğruluk, Aghila (2010) tarafından da vurgulanan temel avantajları arasında yer almaktadır.

4.2.2.2. Sıralı Minimal Optimizasyon (SMO)

Sıralı Minimal Optimizasyon (SMO: Sequential Minimal Optimization), Support Vector Machines (SVM) eğitiminde ortaya çıkan büyük boyutlu sayısal ikinci dereceden programlama (QP: quadratic programming) problemlerini, iki Lagrange çarpanı (α) içeren en küçük alt problemler hâline indirerek hızlı ve verimli bir şekilde çözen optimizasyon algoritmasıdır. Bu yöntem, John Platt tarafından 1998 yılında Microsoft Research'te geliştirilmiştir.

SMO, yakınsamayı sağlamak için Osuna teoremini kullanarak genel QP problemini QP alt problemlerine ayırır. Önceki yöntemlerin aksine, SMO her adımda mümkün olan en küçük optimizasyon problemini çözmeyi tercih eder (Platt, 1998). Bu alt problemler analitik olarak çözülerek, ileri düzey QP çözümleri veya matris işlemleri gerektirmeden işlem tamamlanır. Bu özelliği sayesinde, ekstra bellek kullanımı olmaz; işlem zamanı ve bellek maliyeti, tablo

boyutuna göre bünyede lineer ile kuadratik ölçekli olur. Özellikle lineer SVM'ler ve seyrek veri (sparse) yapılar için son derece hızlıdır; bazı durumlarda önceki "chunking" yöntemlerine göre 1000 kat daha hızlı sonuç verebilir

4.2.2.3. En Yakın Komşu (kNN)

Denetimli öğrenmede, bir veri setindeki sınıflar arasındaki dengesiz örnek sayısı, algoritmaların azınlık sınıfından bir örneği çoğunluk sınıfından biri olarak sınıflandırmasına neden olabilir. Bu problemi çözmek amacıyla, kNN algoritması diğer dengeleme yöntemlerine bir temel sağlar (Beckmann vd, 2015). Yeni bir örneği sınıflandırırken, eğitim verisindeki k , en yakın komşusuna bakar ve komşuların çoğunluğu hangi sınıfta ise yeni örnek o sınıfa atanır.

kNN, Veri dağılımının şekline uyum sağlayan bir karar yüzeyi oluşturarak, eğitim kümesi büyük veya temsili olduğunda iyi doğruluk oranları elde etmelerini sağlar. kNN ilk olarak Fix ve Hodges (1989) tarafından tanıtılmış ve olasılık yoğunluklarının güvenilir parametrik tahminlerinin bilinmediği veya belirlenmesinin zor olduğu durumlarda ayırıcı analiz yapma ihtiyacıyla geliştirilmiştir. kNN, parametrik olmayan bir tembel öğrenme algoritmasıdır. Temel veri dağılımı hakkında herhangi bir varsayımda bulunmadığı için parametrik değildir. Gerçek dünyadaki pratik verilerin çoğu, yapılan tipik teorik varsayımlara (örneğin, Gauss karışımları, doğrusal ayrılabilirlik vb.) uymaz (Beckmann vd, 2015). kNN gibi parametrik olmayan algoritmalar bu durumlar için daha uygundur (Dasarathy, 1991; Duda vd., 2001). Aynı zamanda tembel bir algoritma olarak da kabul edilir. Tembel bir algoritma, var olmayan veya minimum düzeyde bir eğitim aşamasıyla çalışır, ancak maliyetli bir test aşaması vardır. kNN için bu, eğitim aşamasının hızlı olduğu, ancak test aşamasında tüm eğitim verilerine ihtiyaç duyulduğu veya en azından en temsili verilere sahip bir alt kümenin mevcut olması gerektiği anlamına gelir. Bu, tüm destek dışı vektörleri atabileceğiniz Destek Vektör Makineleri (SVM) gibi diğer tekniklerle çelişir.

kNN algoritması aşağıdaki adımlara göre gerçekleştirilir (Beckmann vd, 2015):

1. Bir x_i örneği ile T eğitim kümesinin tüm örnekleri arasındaki mesafeyi (genellikle Öklid) hesaplayın,
2. En yakın k komşuyu seçin,

3. x_i örneği, en yakın k komşu arasında en sık görülen sınıfa göre sınıflandırılır (etiketlenir). Ayrıca, sınıflandırma kararını ağırlıklandırmak için komşuların mesafesini kullanmak da mümkündür.

k değeri eğitim verisine bağlıdır. Küçük bir k değeri, gürültünün sonuç üzerinde daha fazla etkiye sahip olacağı anlamına gelir. Büyük bir değer, hesaplama açısından maliyetli hale getirir ve kNN'nin arkasındaki temel felsefeyi çürütür: Birbirine yakın noktalar benzer yoğunluklara veya sınıflara sahip olabilir. Literatürde genellikle k için tek değerler bulunur, normalde $k = 5$ veya $k = 7$ 'dir. Dasarathy (1991), büyük veri kümelerinde Bayes sınıflandırıcısına çok yakın bir performans elde etmeyi sağlayan $k = 3$ olduğunu bildirir. Bu çalışmada ise $k = 1$ tercih edilmiştir. $k = 1$ durumunda sınıflandırılacak örnek, eğitim kümesindeki yalnızca en yakın tek bir komşuya bakılarak etiketlenir. Yani hangi metrik seçildiyse, mesafesi en küçük olan komşu dikkate alınır. “En yakın komşuların” belirlenmesinde ise uzaklık ölçüleri (distance metrics) kullanılır. kNN algoritması için en yaygın kullanılan uzaklık ölçüleri Öklid uzaklığı (Euclidean Distance) ve Chebyshev uzaklığı (Chessboard / Maximum Distance)’dır.

Öklid Uzaklığı (Euclidean Distance)

Öklid uzaklığı, n -boyutlu uzayda iki nokta arasındaki doğrusal mesafeyi ölçen en temel metrik olup, geometri temelli ilk açıklamaları Öklid’in Elementler adlı eserine dayanmaktadır (Euclid, 300 BCE/1956; O’Connor ve Robertson, 1999). Kavram, adını antik Yunan matematikçisi Öklid’den almaktadır. Öklid Uzaklığı en çok kullanılan uzaklık ölçüsüdür ve 2 boyutlu bir düzlemde klasik “cetvelle ölçülen” düz mesafe gibidir. İki boyutlu düzlemde Pisagor teoremi kullanılarak tanımlanır:

$$D_{Euclidean}(x, y) = \sqrt{\sum_{i=0}^n (x_i - y_i)^2} \quad \text{formülü ile hesaplanır.}$$

Chebyshev Uzaklığı (Chebyshev Distance)

Chebyshev uzaklığı, iki nokta arasındaki mesafeyi koordinatlar arasındaki mutlak farkların en büyüğüyle tanımlayan bir L_∞ normudur (O’Connor ve Robertson, 2002) ve adını, 19. yüzyılda yaşamış Rus matematikçi Pafnuty Lvovich Chebyshev’den almaktadır. Chebyshev

uzaklığı, iki nokta arasındaki farkların en büyüğünü mesafe olarak alır yani eksenlerden hangisinde en büyük fark varsa, o mesafe alınır. Bu sebeple Maksimum uzaklık veya “Chessboard Distance” olarakta adlandırılmaktadır. Chebyshev uzaklığı,

$$D_{Chebyshev}(x, y) = \max_i |x_i - y_i| \quad \text{formülü ile hesaplanır.}$$

F4.2.3. Derin Öğrenme Tabanlı Yöntemler

Derin öğrenme tabanlı yöntemler, duygu sınıflandırmada son yıllarda öne çıkan güçlü yaklaşımlardır. Bu yöntemler, metinlerdeki karmaşık anlam ilişkilerini ve bağlam bilgilerini öğrenmek için yapay sinir ağlarının çok katmanlı yapılarından yararlanır.

Yinelemeli Sinir Ağları (RNN, LSTM, GRU): Metinlerdeki sıralı bilgiyi yakalamada etkilidir. LSTM (Long Short-Term Memory) ve GRU (Gated Recurrent Unit), uzun bağımlılıkları öğrenme konusunda klasik RNN'lere göre daha başarılıdır (Hochreiter ve Schmidhuber, 1997; Chung vd., 2014).

Evrişimli Sinir Ağları (CNN): Başlangıçta görüntü işleme için geliştirilmiş olsa da, kelime öbekleri ve n-gram temelli duygu özelliklerini yakalamada oldukça etkilidir (Kim, 2014).

Dönüştürücü (Transformer) Tabanlı Modeller: Özellikle BERT, RoBERTa ve GPT gibi önceden eğitilmiş dil modelleri, duygu sınıflandırmada en güncel ve başarılı yaklaşımlar arasındadır. Bu modeller, bağlama duyarlı kelime temsilleri üreterek klasik makine öğrenmesi ve RNN/CNN tabanlı yöntemlerden daha yüksek başarı sağlamaktadır (Devlin vd, 2019).

Derin öğrenme yöntemlerinin en önemli avantajı, manuel öznitelik mühendisliğine duyulan ihtiyacı azaltması ve büyük veri kümelerinden otomatik olarak anlamlı temsiller öğrenebilmesidir. Bununla birlikte, yüksek hesaplama gücü ve geniş veri ihtiyacı dezavantajları arasında sayılabilir.

Duygu sınıflandırma yöntemleri üç genel başlık altında ele alsak da bazı çalışmalarda sözlük tabanlı ve makine öğrenmesi tabanlı yöntemler bir araya getirilmekte ya da derin öğrenme modelleri ile makine öğrenmesi sınıflandırıcıları birlikte kullanılmaktadır. Bu hibrit yaklaşımlar, özellikle sınırlı veri bulunan durumlarda daha dengeli sonuçlar verebilmektedir (Young vd., 2018).

5. LİTERATÜR

Sosyal medya platformlarının yaygınlaşması, kullanıcıların görüşlerini, memnuniyet düzeylerini ve tecrübelerini çevrimiçi ortamlarda yoğun bir şekilde paylaşmasına imkân tanımış; bu durum, metin tabanlı büyük veri kaynaklarının ortaya çıkmasına ve bu verilerin sistematik olarak analiz edilmesi gerekliliğine yol açmıştır. Duygu analizi, sosyal medya, e-ticaret ve çevrimiçi platformlarda kullanıcıların görüşlerini anlamak ve sınıflandırmak amacıyla yaygın olarak kullanılan bir metin madenciliği yöntemidir. Literatürde, Twitter, e-ticaret siteleri, çevrimiçi yorum platformları, sağlık verileri ve kamu projelerine ilişkin metinler gibi farklı veri kaynakları kullanılarak gerçekleştirilen çalışmalar, çeşitli makine öğrenmesi algoritmalarının performanslarını karşılaştırmalı olarak incelemiştir. Bu bağlamda, Naive Bayes, Destek Vektör Makineleri (SVM), k-En Yakın Komşu (kNN), Lojistik Regresyon, Çok Katmanlı Algılayıcı (MLP) ve derin öğrenme gibi yöntemler sıklıkla tercih edilmiş; farklı öznitelik çıkarma teknikleri (Bag-of-Words, TF-IDF, sözlük tabanlı yaklaşımlar vb.) ile elde edilen temsil biçimlerinin sınıflandırma başarımı üzerindeki etkileri de kapsamlı biçimde araştırılmıştır. Çalışmalar, yalnızca algoritma seçiminden değil, aynı zamanda veri kümesinin dengesi, boyutu ve niteliğinden de model performansının doğrudan etkilendiğini ortaya koymaktadır. Aşağıda metin madenciliği ve duygu analizi yapılan benzer çalışmalara yer verilmiştir.

Go vd. (2009) Twitter verilerini ilk kez kullanarak yaptıkları duygu analizinde sorunu ikili bir sınıflandırma problemi olarak ele almışlar ve tweetleri olumlu veya olumsuz olarak sınıflandırmıştır. Verilerin son işlenmesinden sonra, olumlu ifadeler içeren ilk 800.000 tweet ve olumsuz ifadeler içeren 800.000 tweet olmak üzere toplam 1.600.000 eğitim tweeti kullanılmıştır. 177 olumsuz tweet ve 182 olumlu tweet'ten oluşan bir set manuel olarak işaretlenmiştir. Makine öğrenimi algoritmaları (Naive Bayes, Maximum Entropy ve SVM) kullanılarak modeller karşılaştırılmış ve modeller ifade verileriyle eğitildiğinde %80'in üzerinde doğruluğa sahip olduğunu raporlanmıştır.

Amir Asiaee vd. (2012) çalışmalarında toplu tweetler yerine tweet başına duygu analizi için genel bir çerçeve sunarak makine öğrenmesi algoritmalarının performansını incelemiştir. Basit kelime çantası özellik vektörüyle tweet başına sınıflandırma görevleri için yüksek doğruluk elde edilebileceğini göstermişler, seyrek özellik vektörünün daha düşük boyutlu bir uzaya yansıtılmasının hesaplama açısından faydalı olduğunu ve sınıflandırma doğruluğunu önemli ölçüde etkilemediğini de göstermişlerdir.

Poyraz (2012) çalışmasında tıp alanında veri madenciliği uygulamaları için 683 hasta verisinden oluşan bir veri kümesi kullanmıştır. WEKA programıyla oluşturulan karar ağacı algoritması olan ve temeli ID3 ve C4.5 algoritmalarına dayanan J48, Bayes sınıflandırma algoritmalarından Naive-Bayes, regresyon tabanlı algoritmalarından lojistik regresyon ve örnek tabanlı sınıflandırma algoritmalarından Kstar algoritmaları kullanılarak oluşturulan modeller karşılaştırılmıştır. Lojistik regresyon %96.92 ile en doğru sonucu vermiş ve %96.33 oranı ile NaiveBayes ikinci en iyi sonucu çıkarmıştır.

Dey vd. (2016), film ve otel yorumlarında Naive Bayes ve k-NN yöntemlerini karşılaştırmıştır. Film yorumlarında Naive Bayes daha iyi performans sağlarken, otel yorumlarında her iki algoritmanın da benzer doğruluk oranlarına ulaştığı görülmüştür.

Göker & Tekedere (2017) çalışmalarında, FATİH projesine yönelik internet ortamında yer alan yorumları makine öğrenmesi algoritmalarını uygulayarak Visual Studio C# programı ile sınıflandırmışlardır. FATİH projesine yönelik 444 görüş içeren metin dosyasındaki metinler tf-idf ağırlıklandırma yöntemi ile vektörel olarak temsil edilerek sınıflandırma algoritmalarının model başarımları karşılaştırılmış ve en yüksek performansı gösteren algoritmanın Sıralı Minimal Optimizasyon (%88,73) olduğu rapor edilmiştir.

Mesri (2017) çalışmasında, Knime programını kullanarak bir bankanın mobil uygulamasına yapılan kullanıcı yorumlarını toplamış ve bu program aracılığıyla veriyi analiz için hazır hale getirmiştir. Hazırlanan veriler WEKA yazılımı kullanılarak duygu analizi gerçekleştirilmiştir. Kullanıcı yorumlarını olumlu, olumsuz ve öneri şeklinde sınıflamış; model sınıflandırma doğruluğu oranının %84,74 olduğunu belirtmiştir.

Karamanlı (2019) tarafından yapılan çalışmada bir e-ticaret sitesinden en çok satılan üç farklı markaya ait cep telefonlarına yapılan yorumlar veri olarak toplamıştır. Makine öğrenmesi sınıflandırma algoritmalarıyla duygu analiz sonucu tahmin edilmiştir. DVM, Naive Bayes ve kNN algoritmaları arasından en iyi sonuç hem pozitif hem de negatif tahminleme dikkate alındığında DVM algoritmasına ait olduğu raporlanmıştır.

Işık (2019) tarafından yapılan çalışmada e-ticaret markalarına yapılan sosyal medya yorumları Twitter API kullanılarak toplanmıştır. Veri dağılımının ve öznelik seçimi yapılmasının Naive Bayes, Sıralı Minimal Optimizasyon ve k-En Yakın Komşu sınıflandırma algoritmaları üzerindeki etkisini incelemek üzere WEKA yazılımını kullanarak, dengesiz veri kümesinin, dengeli veri kümesine göre daha iyi performans sağladığını belirtmiştir. %93,52 sınıflandırma doğruluğu ile kNN en iyi performansı gösteren sınıflandırma yöntemi olmuştur.

Tuzcu (2020) tarafından yapılan çalışmada, çevrimiçi bir kitap satış sitesinin kullanıcı yorumları üzerinde duygu analizi yapılmış ve Python programlama dili kullanılarak Çok Katmanlı Algılayıcı algoritması uygulanmıştır. Daha sonra RapidMiner kullanılarak, aynı veri seti üzerinde Naive Bayes, Destek Vektör Makineleri ve Lojistik Regresyon algoritmaları uygulanarak algoritmaların sınıflandırma başarıları karşılaştırılmıştır. Toplam sınıflandırma başarısı en yüksek algoritmanın Çok Katmanlı Algılayıcı olduğu görülmektedir. Pozitif yorumları en iyi sınıflandıran algoritma Destek Vektör Makineleri ve negatif yorumları en iyi sınıflandıran algoritma ise Çok Katmanlı Algılayıcı olmuştur.

Budak (2021) tarafından yapılan çalışmada dünyanın en büyük hava yolu ağı olan Star Alliance'a üye 26 hava yolu şirketini değerlendirmek için, TripAdvisor'daki müşteri yorumları ve sayısal skorları kriter olarak ele almış ve duygu analizinden çıkan sonuçları DVM, Naive Bayes ve Derin Öğrenme Algoritmalarını kullanarak sınıflandırma tahmini yapmıştır. En yüksek doğruluk oranının Derin Öğrenme yöntemi ile elde edildiği görülmüştür.

Kaur ve Malik (2021), United Airways, Jet Blue, American Airways gibi altı havayolu şirketine ait Twitter verileri üzerinde SVM, Naïve Bayes ve Rastgele Orman gibi makine öğrenimi algoritmalarını kullanarak duygu analizi yapmıştır. En yüksek başarı, %90'nın üzerinde doğruluk oranıyla SVM algoritmasından elde edilmiştir.

George ve Srividhya (2022), dengeli ve dengesiz veri kümelerinde ensemble yöntemlerle duygu analiz performansını karşılaştırmış; sonuçlar, dengesiz veri kümelerinde sınıflandırma performansının belirgin şekilde farklılaştığını göstermiştir.

Azhar vd. (2023), k-NN algoritmasını sözlük tabanlı yöntemle birlikte kullanarak Twitter duygu analizinde doğruluk, kesinlik ve duyarlılık metriklerini karşılaştırmıştır. k-NN + sözlük kombinasyonunun %77, Naive Bayes'in ise %81 doğruluk sağladığı bildirilmiştir.

Dhiyaulhaq ve Gunawan (2023), yarı denetimli öğrenme çerçevesinde Hızlı Tren projesine ilişkin kullanıcılardan gelen yorumlarını analiz etmiştir. SVM'in %86,9 ile daha iyi ve kararlı performans sonuçları elde ettiğini ve TF-IDF çıkarma özelliğini kullanan Destek Vektör Makinesi yönteminin Naïve Bayes yöntemi ve Word2vec çıkarma özelliğiyle karşılaştırıldığında daha iyi olduğu sonucuna varmıştır.

Munawaroh ve Alamsyah (2023), COVID-19 aşısıyla ilgili Twitter duygu analizinde SVM, Naive Bayes ve kNN algoritmalarını karşılaştırmış; SVM'in %96,3, Naive Bayes'in %94 ve kNN'in %91 doğruluk oranına sahip olduğunu raporlamıştır.

Marttin (2025) tarafından yapılan çalışmada KADES uygulamasına yapılan kullanıcı yorumları veri olarak toplanmış ve WEKA programı kullanılarak elde edilen Naive Bayes, SMO VE kNN algoritmalarının sonuçları karşılaştırılmıştır. Dengeli ve dengesiz veri kümelerindeki analiz sonuçlarına göre SMO, diğer sınıflandırma yöntemlerinden daha başarılı olduğu raporlanmıştır.

Yapılan çalışmalar, duygu analizi alanında tek bir algoritmanın her problem senaryosunda mutlak üstünlük sağlamadığını göstermektedir. Bununla birlikte, SVM ve Naive Bayes algoritmaları, farklı veri kümelerinde genellikle yüksek doğruluk oranları ile öne çıkarken; lojistik regresyon, kNN, MLP ve derin öğrenme yaklaşımları belirli bağlamlarda daha üstün performans sergileyebilmektedir. Örneğin, metin boyutunun kısa olduğu ve veri dengesizliğinin düşük olduğu senaryolarda SVM yüksek başarı sağlarken; daha karmaşık, çok boyutlu veya dengesiz veri setlerinde derin öğrenme ve lojistik regresyon yöntemleri avantajlı olabilmektedir. Ayrıca, öznitelik çıkarma tekniklerinin ve veri ön işleme adımlarının optimizasyonu, sınıflandırma doğruluğunda belirleyici bir rol oynamaktadır. Dengeli ve dengesiz veri kümeleri ile gerçekleştirilen deneyler, veri dağılımının performans üzerindeki etkisini açıkça ortaya koymuş; özellikle dengesiz veri kümelerinde uygun yeniden örnekleme (resampling) yöntemlerinin kullanılması gerektiğini göstermiştir. Genel olarak, literatür, duygu analizinde başarı elde etmenin yalnızca algoritma seçiminden ibaret olmadığını; veri kümesi özellikleri, öznitelik mühendisliği ve model optimizasyonunun birlikte ele alınmasının kritik öneme sahip olduğunu ortaya koymaktadır.

5.1. Araştırmanın Literatüre Katkısı

Duygu analizi ve makine öğrenmesi algoritmalarının performanslarının incelendiği çalışmaların büyük çoğunluğu, genellikle Twitter paylaşımları, e-ticaret yorumları, film veya otel incelemeleri gibi günlük hayata yönelik veri kaynakları üzerinde yoğunlaşmıştır. Bu bağlamda, farklı veri kümesi dağılımları ve öznitelik seçim tekniklerinin sınıflandırma algoritmalarına etkisi çeşitli yönleriyle değerlendirilmiştir. Ancak, afet yönetimi ve deprem uygulamaları bağlamında yapılan çalışmalar sınırlı sayıdadır.

Bu araştırmanın özgün katkısı, deprem gibi kritik bir doğal afet senaryosunu temel alması ve kullanıcıların deprem uygulamalarına yönelik yorumlarını veri kaynağı olarak incelemesidir. Böylelikle, sadece algoritmaların performansları değil, aynı zamanda afete hazırlık süreçlerinde kullanılan dijital araçların etkililiği de değerlendirilmiş olmaktadır. Çalışmanın kullandığı veri seti, Google Play ve App Store üzerinden toplanan deprem

uygulaması kullanıcı yorumları olup, bu yönüyle literatürde nadir işlenen bir kaynak özelliği taşımaktadır.

Elde edilen bulgular, hem teknik düzeyde (veri kümesi dağılımı ve öznitelik seçiminin sınıflandırma performansına etkisi) hem de uygulamalı düzeyde (deprem uygulamalarının kullanıcı memnuniyet düzeyleri) katkı sunmaktadır. Ayrıca, bu çalışma yalnızca akademik alana değil, aynı zamanda uygulama geliştiricilerine ve afet yönetimi kurumlarına da somut öneriler sağlayarak toplumsal fayda perspektifiyle de değer taşımaktadır.



6. METODOLOJİ

6.1. Araştırmanın Amacı

Bu çalışmada sınıflardaki veri dağılımının ve öznitelik seçmenin sınıflandırma üzerindeki etkisi araştırılmaktadır. “Veri kümesinin dengeli veya dengesiz tercih edilmesi sınıflandırma algoritmalarının performansları etkiler mi? ve “Öznitelik seçiminin yapıldığı durumlarda veya özniteliğin seçilmediği durumlarda sınıflandırma algoritmalarının performansları etkilenir mi?” sorularına yanıt aranmıştır.

6.2. Destek Araçlar

Bu başlık altında, çalışmamızda kullandığımız veri madenciliği, doğal dil işleme ve makine öğrenmesi araçları hakkında kısa bilgiler paylaşılmıştır.

6.2.1. Instant Data Scraper

Instant Data Scraper, herhangi bir web sitesi için otomatikleştirilmiş bir veri çıkarma aracıdır. Bir HTML sayfasında hangi verilerin en alakalı olduğunu tahmin etmek için yapay zekayı kullanır ve Excel'e veya CSV dosyasına (XLS, XLSX, CSV) kaydedilmesine izin verir.

Yapay zekâ tarafından algılanan veriler, kullanıcı tarafından yapılan seçimlerle özelleştirilerek daha fazla doğruluğa sahip analizler elde edilebilmektedir. Kullanıcıların herhangi bir kodlama, json veya xml becerisine sahip olmalarına gerek olmadan hizmet veren bu veri çıkarma aracı ücretsiz olduğu için kullanıcı dostudur.

6.2.2. ITU Turkish NLP Web Servis API

İstanbul Teknik Üniversitesi Doğal Dil İşleme Grubu tarafından geliştirilen "İTÜ Türkçe NLP Web Servisi" bir doğal dil işleme platformu sunarak SaaS (Software as a service: Servis olarak yazılım) olarak çalışır ve ön işleme, morfoloji, sözdizimi ve varlık tanıma gibi birçok katmanda en son teknoloji NLP araçlarını sunar (Eryiğit, 2014).

6.2.3. WEKA 3.8.6

WEKA (Waikato Environment for Knowledge Analysis: Waikato Bilgi Analizi Ortamı), araştırmacıların makine öğrenimindeki en son tekniklere kolayca erişmesini sağlayacak birleşik bir çalışma ortamına duyulan ihtiyaçtan doğmuş, araştırmacılara ve uygulayıcılara kapsamlı bir makine öğrenimi algoritmaları ve veri ön işleme araçları koleksiyonu sunmayı amaçlayarak kullanıcıların yeni veri kümeleri üzerinde farklı makine öğrenimi yöntemlerini hızla denemelerine ve karşılaştırmalarına olanak tanımaktadır. WEKA, veri madenciliği ve makine öğreniminde çığır açan bir sistem olarak kabul edilmektedir (Hall vd., 2009: 10).

WEKA açık kaynaklı ve ücretsiz olduğu için herkes tarafından indirip kullanabilmektedir. Explorer, experimenter gibi kullanıcı dostu arayüzleri sayesinde kodlama bilmeden makine öğrenmesi yapılabilmektedir. Sınıflandırma, kümeleme, regresyon, öznitelik seçimi ve ilişki kuralları için çok sayıda hazır algoritma barındırmakta ve veri ön işleme sürecinde yer alan eksik verilerin temizlenmesi, özniteliklerin dönüştürülmesi, filtreleme gibi işlemler kolayca yapılmaktadır. Sonuçların tekrarlanabilirliği yüksek olması nedeniyle akademik makalelerde sıkça tercih edilmektedir.

6.3. Araştırmanın Örneklemi ve Verilerin Analiz için Hazırlanması

Bu çalışma için veriler Google Play ve App Store’da yer alan en çok indirilen değerlendirme yapılan 8 deprem uygulamasına (Deprem Ağı, Deprem Bilgi Sistemi, Deprem Uyarılarım, Last Quake, Deprem Ağı Pro, Deprem Türkiye, Deprem Uyarılarım Pro ve Deprem – Earthquake) ait kullanıcı yorumları Instant Data Scraper veri çıkarma aracıyla toplanmıştır. 09.03.2025 - 19.03.2025 tarihleri arasında toplanan 2017 ile 2025 yılları arasında yapılmış 9.271 yorum içerisinde Türkçe dilinde olan 8.405 yorum, çalışmada kullanılan başlangıç veri setini oluşturmaktadır (Tablo 6.1). Veri seti 2020 Malatya, 2020 İzmir (Samos depremi) ve 2023 Kahramanmaraş depremleri gibi Türkiye’de son yıllarda yaşanan büyük depremler sonrasında kullanıcılar tarafından yapılmış yorumları da içermektedir.

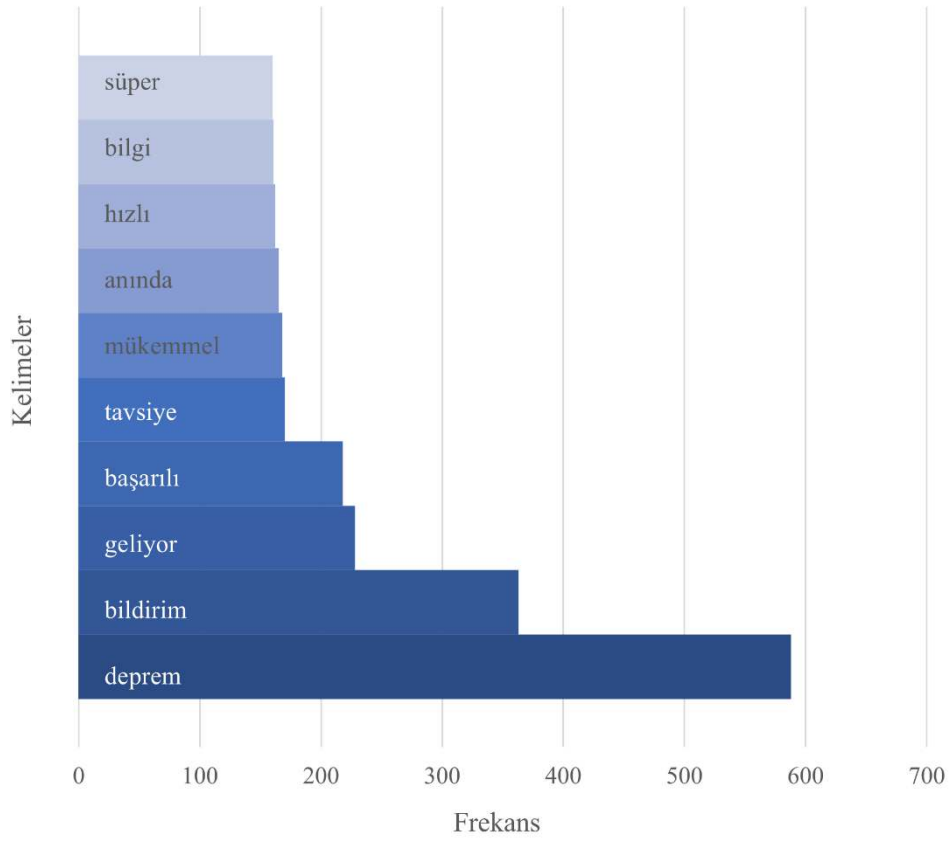
Tablo 6.1. Veri Çekilen Uygulamalar

<u>Uygulama Adı</u>	<u>Platform</u>	<u>Veri Tarih Aralığı</u>	<u>Çekilen Yorum Sayısı</u>
Deprem Ağı	Google Play ve App Store	2019-2025	2519
Deprem Uyarılarım	Google Play ve App Store	2019-2025	1931
Deprem Bilgi Sistemi	Google Play ve App Store	2023-2025	1677
Deprem Ağı Pro	Google Play	2020-2025	857
Deprem – Earthquake	App Store	2020-2025	550
Deprem Türkiye	Google Play ve App Store	2021-2025	407
Last Quake	Google Play	2017-2025	347
Deprem Uyarılarım Pro	Google Play	2020-2025	97
<u>Toplam Yorum Sayısı:</u> 8405			

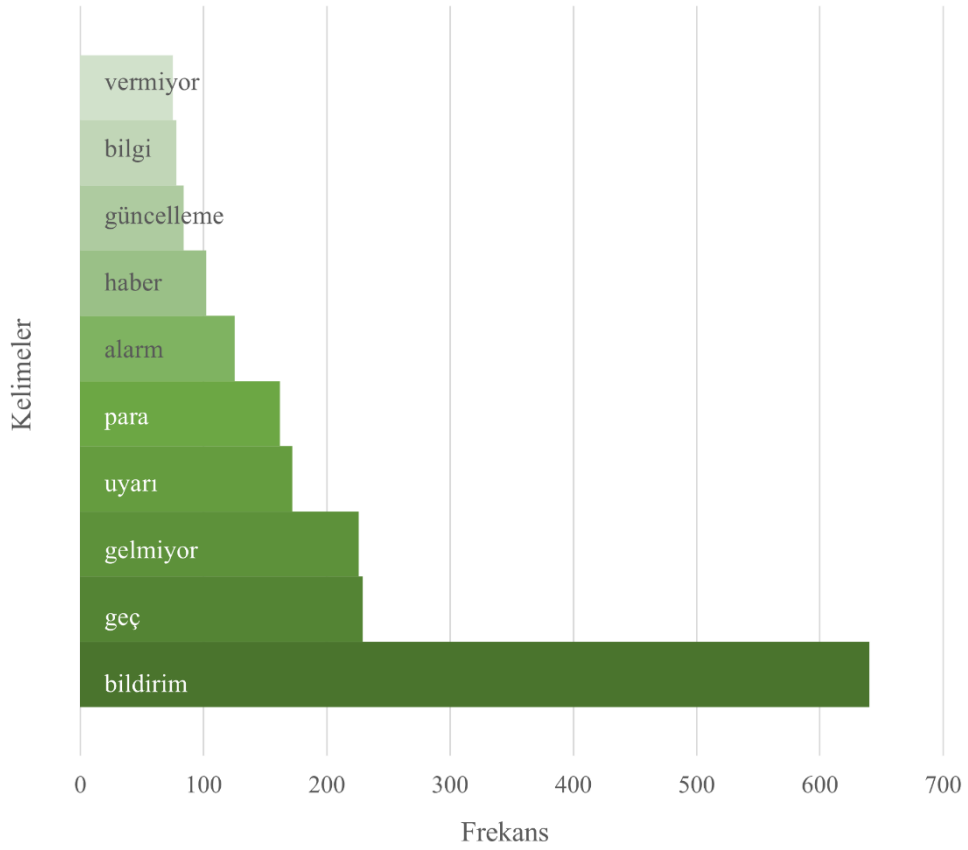
Veri Etiketleme: Veri kümesinde yer alan yorumlar, uygulama hakkında pozitif bir duygu taşıyorsa "olumlu" olarak etiketlenirken, negatif bir duygu taşıyorsa "olumsuz" olarak etiketlenmiştir. Pozitif veya negatif bir duygu taşımayan veya duygunun uygulamaya mı yoksa uygulama geliştiricisine mi yönelik olduğu açık olmayan yorumlar ise nötr olarak etiketlenmiştir. Pozitif etiketli yorumlar, pozitif ifadeler içerir; "başarılı", "harika", "iyi", "mükemmel", "çok beğendim" gibi ve kullanıcı, yorumunda uygulamadan memnun olduğunu ifade eder. Örneğin "Çok hızlı bir şekilde bilgi geliyor ve bence güzel bir uygulama" yorumunda "güzel" kelimesi pozitif bir anlam taşıyor ve genel olarak memnuniyet duygusu ifade ediliyor. Negatif etiketli yorumlar ise negatif ifadeler içerir; "yetersiz", "geç", "kötü", "beğenmedim", "hayal kırıklığı" gibi ve kullanıcı, yorumunda uygulama ile ilgili memnuniyetsizlik duyduğunu dile getirir. Örneğin "Güncellemeler yeteriz, bildirimler geç geliyor" yorumunda "yetersiz" ve "geç" gibi olumsuz kelimeler, kullanıcı şikayetini ve memnuniyetsizliğini göstermektedir. Pozitif veya negatif bir duygu taşımayan "Allah büyüktür veya imana namaza sarılana birsey olmaz" gibi ve yorumun uygulamaya mı uygulama geliştiricisine mi yapıldığı belli olmayan "Millet kendini vatani kurtardı sanıyor AFAD'dan bilgi alıyorsun zaten" gibi yorumlar nötr yorum olarak etiketlenmiştir. "Yaşadığımız il için ayrıcalıklı bildirim özelliği tanımlanmalıdır" şeklindeki öneri içeren yorumlarda yine nötr olarak etiketlenmiştir. Farklı duyguları ifade etmek için kullanılan "kızgın surat", "gülen surat" gibi emoji ve ":", "D" gibi simgelerle ifade edilen yorumlar veri kümesinden çıkarılmıştır. Etiketleme işlemi sonucunda 4.283 yorum olumlu, 1.583 yorum olumsuz ve 2.539 yorum nötr olarak etiketlenmiştir. Olumlu ve olumsuz yorumlarda en çok kullanılan 10 kelime Şekil 6.4 ve Şekil 6.5'deki gibidir.

Veri Önileme; metin normalizasyonunun yapıldığı, durak kelimelerin çıkarılarak gerekli dönüştürme, tarama ve işaretlemelerin yapıldığı bir süreçtir.

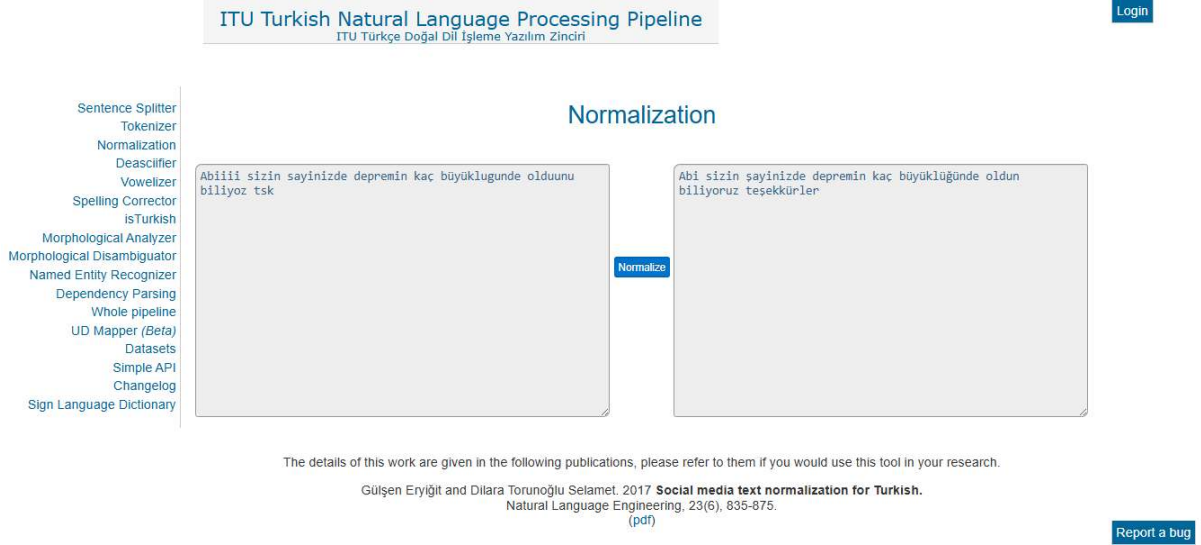
Metin normalizasyonu, doğal dil işleme (NLP) ve metin madenciliği süreçlerinde oldukça önemli bir adımdır. Bu işlem, ham metnin daha temiz ve analiz edilebilir hale gelmesini sağlar. İTÜ Türkçe Doğal Dil İşleme aracında yer alan "Normalization" sekmesi kullanılarak araştırma için topladığımız veriler içindeki metinler normalleştirilmiştir (Şekil 6.6). Metin normalizasyonu, daha iyi sonuçlar elde etmek ve model performansını artırmak için kritik bir adımdır.



Şekil 6.4. Olumlu Yorumlarda En Çok Kullanılan 10 Kelime



Şekil 6.5. Olumsuz Yorumlarda En Çok Kullanılan 10 Kelime



Şekil 6.6. İTÜ Türkçe Doğal Dil İşleme Aracı ile Metin Normalizasyonu

İTÜ Türkçe Doğal Dil İşleme aracının yaptığı düzeltmeler sonrası veriler tekrar gözden geçirilerek İTÜ Türkçe Doğal Dil İşleme aracının düzenlemediği bazı kelimeler el ile düzenlenmiştir. Şekil 6.6’da yer alan “Abiiii sizin sayinizde depremin kaç büyüklugunde olduunu biliyoz tsk” yorumu İTÜ Türkçe Doğal Dil İşleme aracı ile “Abi sizin şayinizde depremin kaç büyüklüğünde oldun biliyoruz teşekkürler” şeklinde normalize edilmiştir. Fakat cümle içerisinde yer alan “sayinizde” ve “olduunu” kelimelerinin doğru şekilde normalleştirilmediği gözükmemektedir. Bu yüzden bu kelimeler “Abi sizin sayenizde depremin kaç büyüklüğünde olduğunu biliyoruz teşekkürler” şeklinde manuel olarak düzenlenmiştir.

Durak kelimelerin çıkarılması, kullanıcıların yorum yaparken sık kullandığı zarflar, zamirler, edatlar, selamlaşmalar ve tek başına bir anlam ifade etmeyen kelimelerin veri kümesinden çıkarılmasını ifade etmektedir. Bu kelimelerin çıkarılması, modelin daha doğru ve verimli çalışmasına olanak tanır. Durak kelimeler çıkarıldıktan sonra metin üzerinde yapılacak olan işlem veya modelin analiz ettiği veriler daha net, anlamlı ve odaklanmış olur.

Veri ön işleme sürecinin bir parçası olan durak kelimelerin çıkarılması için Can vd. (2008: 420) tarafından yapılan çalışmadaki Türkçe durak kelimelerden (Tablo 6.2) yararlanılmıştır. Bu kelimeler dışında ayrıca zarflar (hızlıca, yavaşça, dikkatlice. vb.), zamirler (ben, sen, o, biz, siz, onlar), edatlar (ile, için, gibi, ama, vb.), selamlaşmalar (merhaba, selam, günaydın, vb.) ve diğer gereksiz ifadelerde (kolay gelsin, sağ olun, teşekkür ederim, vb.) veri kümesinden çıkarılmıştır.

Tablo 6.2. Türkçe Durak Kelimeler

ama	böylece	eden	hiç	mi	olsun	tarafından
ancak	bu	ederek	hiçbir	mu	olup	üzere
arada	buna	edilecek	için	mü	olur	var
ayrıca	bundan	ediliyor	ile	nasıl	olursa	vardı
bana	bunlar	edilmesi	ilgili	ne	oluyor	ve
bazı	bunları	ediyor	ise	neden	ona	veya
belki	bunların	eğer	işte	nedenle	onlar	ya
ben	bunu	etmesi	itibaren	o	onları	yani
beni	bunun	etti	itibariyle	olan	onların	yapacak
benim	burada	ettiği	kadar	olarak	onu	yapılan
beri	çok	ettiğini	karşın	oldu	onun	yapılması
bile	çünkü	gibi	kendi	olduğu	öyle	yapıyor
bir	da	göre	kendilerine	olduğunu	oysa	yapmak
birçok	daha	halen	kendini	olduklarını	pek	yaptı
biri	de	hangi	kendisi	olmadı	rağmen	yaptığı
birkaç	değil	hatta	kendisine	olmadığı	sadece	yaptığını
biz	diğer	hem	kendisini	olmak	şey	yaptıkları
bize	diye	henüz	ki	olması	siz	yerine
bizi	dolayı	her	kim	olmayan	şöyle	yine
bizim	dolayısıyla	herhangi	kimse	olmaz	şu	yoksa
böyle	edecek	herkesin	mı	olsa	şunları	zaten

Kaynak: Can vd., (2008: 420)

Dönüştürme aşaması, kullanıcı yorumların tamamının küçük karakterlerle değiştirilerek Türkçe karakterlerin (ç, ğ, i, ö, ş, ü) İngilizce karakterlerle (c, g, ı, o, s, u) değiştirildiği, “Dandikmiş, app drop yedi, vb.” gibi küfürlü, alaycı, kinayeli ve mecazlı ifadeler kaldırılarak herhangi bir anlama gelmeyen “qwerty, asdfghj, asasasaddas” gibi kelimelerin veri setinden çıkarılmasını ifade etmektedir. Ayrıca bu aşamada kullanıcı yorumları HTML, CSV ve XML etiketlerinden de temizlenmiştir.

Tarama ve işaretleme aşamasında kullanıcı yorumlarında yer alan tüm noktalama işaretleri, özel karakterler (emojiler, semboller, simgeler) ve sayılar kaldırılmıştır.

Veri etiketleme ve veri ön işleme süreci sonrasına ait örneklerle Tablo 6.3’de yer verilmiştir.

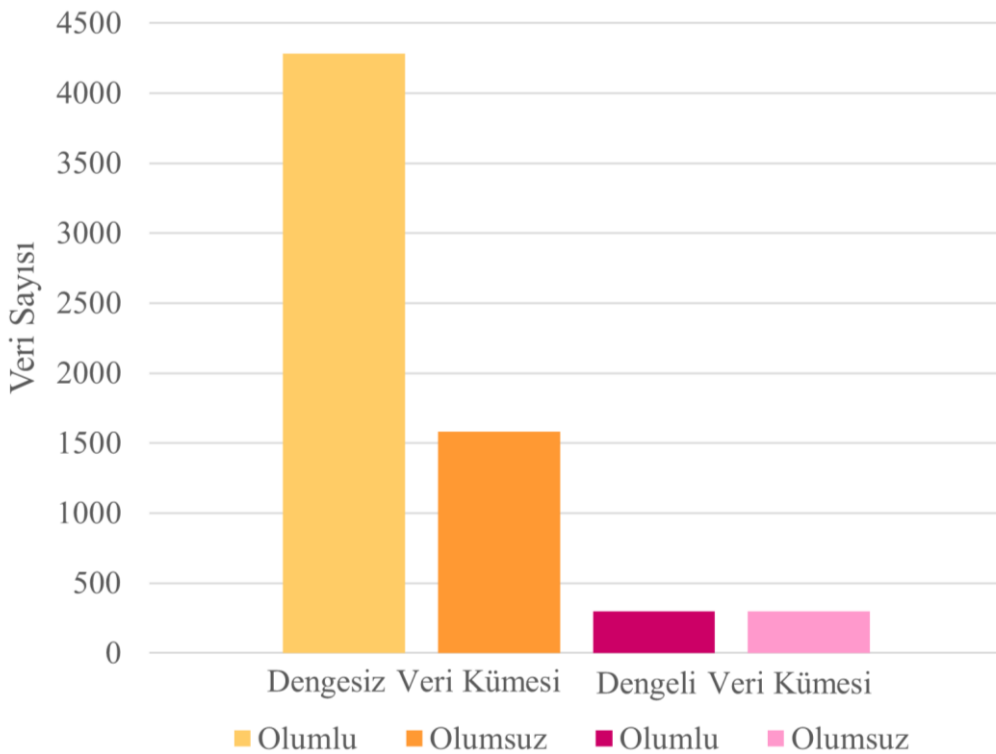
Tablo 6.3. Veri Etiketleme ve Veri Önişleme Örnekleri

<i>Veri (Yorum)</i>	<i>Veri Etiketleme</i>	<i>Veri Önişleme Sonrası</i>
Anlık haberdar oluyorum. Çok kısa sürede uyarı geliyor. Başarılı bir uygulama	Olumlu	anlik haberdar oluyorum kısa surede uyari geliyor basarili uygulama
Çok başarılı elinize sağlık... iyi bir uygulama çok güzel haberlerden ve internetten önce telefonumuza Nerede deprem oldu Hemen düşüyor çok güzel Elinize sağlık. Başarılı. Düdük iyi olmuş :)	Olumlu	basarili elinize saglik iyi uygulama guzel haberlerden internetten once telefonumuza nerede deprem hemen dusuyor guzel elinize saglik basarili duduk iyi olmus
Emeği geçenlere teşekkür ederim. Güzel uygulama yapılmış. Kötü diyenlere aldırmayın.	Olumlu	emegi gecenlere guzel uygulama yapilmis kotu diyenlere aldirmayin
uygulama çok iyi kesinlikle tavsiye ederim en ufak bir deprem de uyarı veriyor	Olumlu	uygulama iyi kesinlikle tavsiye ederim ufak depremde uyari veriyor
Yıllardır kullanıyorum. Çok faydalı.	Olumlu	yillardir kullaniyorum faydali
zamanında uyarı veriyor	Olumlu	zamaninda uyari veriyor
Deprem uyarıları çok geç geliyor	Olumsuz	deprem uyarilari gec geliyor
Farklı uygulamaya kıyasla deprem bildirimleri 15-20 dakika sonra geliyor	Olumsuz	farkli uygulamaya kiyasla deprem bildirimleri dakika sonra geliyor
Gerçekten saçma sapan bi program. Olur olmadık gece vaktinde test alarmı veriyor. 5üstü şiddeti alarm kurmama rağmen 1 de ötüyor. Saçmalığın daniskası!!!	Olumsuz	gercekten sacma sapan program olmadik gece vaktinde test alarmi veriyor üstü siddeti alarm kurmama otuyor sacmalik
Hem uyarı vermiyor hem de ayarlardan ayarladığım sismik ağ bildirimlerindeki deprem büyüklük oranları devamlı sıfırlanıyor.Yani bir işe yaramıyor uygulama.	Olumsuz	uyari vermiyor ayarlardan ayarladigim sismik ag bildirimlerindeki deprem buyukluk oranlari devamli sifirlaniyor yaramiyor uygulama
Son güncellemeden sonra berbat olmuş	Olumsuz	son guncellemeden sonra berbat olmus
yanlış bildirim geliyor ! sürekli panikte başka birsey değil ! ! !	Olumsuz	yanlis bildirim geliyor surekli panikten baska
Millet kendini vatani kurtardı sanıyor AFAD'dan bilgi alıyorsun zaten	Nötr	Veri kümesine dahil edilmemiştir.
Çok fena sallandık hala artcılar devam ediyor	Nötr	Veri kümesine dahil edilmemiştir.
hepimize Allah yardım etsin inşallah Amin	Nötr	Veri kümesine dahil edilmemiştir.

Veri Kümesinin Seçilmesi: Bu çalışmada yukarıda bahsedilen veri etiketleme ve veri ön işleme süreçlerinden geçen 4.283 olumlu etiketli yorum ve 1.583 olumsuz etiketli yorum kullanılarak oluşturulan dengeli ve dengesiz olmak üzere 2 ayrı veri kümesi kullanılmıştır.

Dengesiz veri kümesi oluşturulurken elde edilen 4.283 olumlu ve 1.583 olumsuz olmak üzere toplam 5.866 yorum veri kümesine dahil edilmiştir.

Dengeli veri kümesi oluşturulurken 300 olumlu ve 300 olumsuz olmak üzere toplam 600 yorum veri kümesine dahil edilmiştir. Dengeli veri kümesi oluşturma sürecinde, tekrar eden örneklerin dâhil edilmesi öznitelik uzayındaki çeşitliliği sınırlandırmakta ve bu durum öznitelik sayısının azalmasına yol açmaktadır. Buna karşılık, içerik bakımından farklı verilerin eklenmesi öznitelik çıkarım aşamasında daha yüksek temsil gücü sağlayarak öznitelik sayısının artmasına katkıda bulunmaktadır. Bu sebeple metin ön işleme ve öznitelik çıkarımı aşamalarında veri çeşitliliğinin korunması amacıyla dengeli veri kümesi 300 olumlu yorum ve 300 olumsuz yorum ile sınırlandırılmıştır. Kümeye dâhil edilen yorumlar rastgele seçilmiş olmakla birlikte, bu rastgelelik sürecinde tekrar eden verilerin kullanılmamasına dikkat edilmiştir.



Şekil 6.7. Dengeli ve Dengesiz Veri Kümeleri

Veri Kümesini Özniteliklerine Ayırma her bir veriyi onu tanımlayan niteliklere bölerek analiz edilebilir bir özellik matrisi elde etmeyi ve bu sayede veri madenciliği, makine öğrenmesi ya da istatistiksel modelleme süreçlerinde kullanılabilecek yapısal bir veri formu oluşturmayı ifade eder.

Bu çalışmada WEKA 3.8.6 programı kullanılarak Dengeli ve Dengesiz veri kümelerini özniteliklerine ayrılmıştır. Program içerisinde Experimenter – Analys – File – Dengeli ve Dengesiz veri kümeleri için oluşturulan .arff uzantılı dosya (Dengeli ve Dengesiz veri kümeleri için iki ayrı dosya oluşturulmuştur) – Open Explorer – filters – unsupervised – attribute – StringToWordVector yolu izlenerek özniteliğe ayırma tamamlanmıştır. WEKA 3.8.6 programında wordToKeep değeri default olarak 1000 kabul edilerek dengesiz veri kümesi 1150 özniteliğe ayrılmaktadır. Dengeli veri kümesi, dengesiz veri kümesine göre daha küçük olduğu için wordsToKeep değeri 1000 olduğunda model karmaşılaşmaktadır. Bu yüzden wordToKeep değeri 300 kabul edilerek dengeli veri kümesi 655 özniteliğe ayrılmaktadır.

Öznitelik Seçimi: WEKA 3.8.6’da öznitelik seçimi için yerleşik araçlar bulunmaktadır. Bu işlem, modelin daha az ama daha anlamlı özniteliklerle eğitilmesini sağlar ve çoğu zaman daha iyi doğruluk vermektedir. WEKA 3.8.6 programında Explorer – Select attributes sekmesinden Attribute Evaluator (özniteliklerin önemini nasıl ölçmek istediğin) ve Search Method (en iyi öznitelik alt kümesini nasıl arayacağın) seçilerek öznitelik seçimi tamamlanmaktadır. Sezer (2018: 76) tarafından yapılan bir çalışmada Attribute Evaluator olarak öznitelikleri birlikte değerlendiren ve birbirini tekrar edenleri elemeye çalışan *CfsSubsetEval* ve bilgi kazancını normalize ederek dengesiz sınıflarda tercih edilen *GainRatioAttributeEval* ve Search Method olarak öznitelikleri tek tek sırayan *Ranker* kullanmıştır. Bu çalışmada da en iyi sonuçları verdiği için *CfsSubsetEval* – *GainRatioAttributeEval* ve *Ranker* kullanılarak Dengeli veri kümesi için 40 ve Dengesiz veri kümesi için 256 öznitelik seçimi yapılmıştır.

6.4. Model Başarısını Değerlendirme

Çalışmada elde edilen sonuçların değerlendirilmesi için farklı yöntem ve metrikler kullanılabilir. Bu başlık altında karışıklık matrisi, sınıflandırma doğruluğu, duyarlılık (recall), kesinlik (precision), doğru pozitif oranı (TPR), yanlış pozitif oranı (FPR), F-Ölçütü ve kappa istatistiği incelenmiştir.

6.4.1. Karışıklık (Hata) Matrisi

Karışıklık matrisi (confusion matrix), sınıflandırma algoritmalarının performansını değerlendirmede kullanılan temel araçlardan biridir. Bu matris, modelin gerçek sınıflarla

tahmin edilen sınıfları karşılaştırarak doğru ve hatalı sınıflandırmaları tablo halinde sunar (Tablo 6.4). Böylece doğruluk, duyarlılık, kesinlik ve diğer performans ölçütlerinin hesaplanmasına temel oluşturur (Han, Kamber ve Pei, 2012).

Tablo 6.4. Hata Matrisi

<u>Gerçek</u>	<u>Tahmin</u>		
	<u>Pozitif</u>	<u>Negatif</u>	Toplam
<u>Pozitif</u>	TP	FN	Toplam Pozitif Gerçekler
<u>Negatif</u>	FP	TN	Toplam Negatif Gerçekler
Toplam	Toplam Pozitif Tahminler	Toplam Negatif Tahminler	Tüm Veri Sayısı

Sınıfların pozitif ve negatif olarak önceden belirlenmesi durumunda, karşılık matrisinde yer alan değerler aşağıdaki gibi tanımlanmaktadır:

TP (True Positive: Doğru Pozitif): Doğru olarak tahmin edilen pozitif örnek sayısı

TN (True Negative: Doğru Negatif): Doğru olarak tahmin edilen negatif örnek sayısı

FP (False Positive: Yanlış Pozitif): Yanlış olarak tahmin edilen pozitif örnek sayısı

FN (False Negative: Yanlış Negatif): Yanlış olarak tahmin edilen negatif örnek sayısı

Toplam Pozitif Gerçekler = Doğru Pozitif + Yanlış Negatif = TP + FN

Toplam Negatif Gerçekler = Yanlış Pozitif + Doğru Negatif = FP + TN

Toplam Pozitif Tahminler = Doğru Pozitif + Yanlış Pozitif = TP + FP

Toplam Negatif Tahminler = Yanlış Negatif + Doğru Negatif = FN + TN

Toplam Veri Sayısı (N) = TP + FP + TN + FN

6.4.2. Sınıflandırma Doğruluğu (Doğruluk Oranı)

Sınıflandırma doğruluğu (accuracy), doğru tahmin edilen örneklerin toplam örnek sayısına oranı olarak tanımlanır. En sık kullanılan performans ölçütlerinden biri olmasına karşın, özellikle sınıfların dengesiz olduğu veri setlerinde yanıltıcı sonuçlar verebilir. Bu nedenle tek başına değerlendirme için yeterli görülmemektedir (Witten, Frank ve Hall, 2016).

$$\text{sınıflandırma doğruluğu} = \frac{TN + TP}{N}$$

Hata oranı, modelin yanlış sınıflandırdığı örneklerin oranını ifade eder.

$$\text{hata oranı} = 1 - \text{doğruluk oranı}$$

$$\text{hata oranı} = 1 - \frac{TN + TP}{N}$$

6.4.3. Duyarlılık (Recall)

Duyarlılık (recall), modelin pozitif sınıfa ait örnekleri ne ölçüde doğru yakaladığını gösterir. Gerçek pozitiflerin, tüm gerçek pozitiflere oranı olarak hesaplanır.

$$\text{duyarlılık} = \frac{TP}{TP + FN}$$

6.4.4. Kesinlik (Precision)

Kesinlik (precision), modelin pozitif olarak sınıflandırdığı örneklerin ne kadarının gerçekten pozitif olduğunu ölçer. Gerçek pozitiflerin, tahmin edilen tüm pozitiflere oranı ile hesaplanır.

$$\text{kesinlik} = \frac{TP}{TP + FP}$$

6.4.5. Doğru Pozitif Oranı (True Positive Rate)

Doğru pozitif oranı (TPR), duyarlılıkla eş anlamlıdır ve gerçek pozitiflerin doğru şekilde sınıflandırılma oranını ifade eder. Bu ölçüt, sınıflandırıcının pozitif örnekleri yakalama başarısının doğrudan bir göstergesidir (Han, Kamber ve Pei, 2012).

$$\text{doğru pozitif oranı} = \frac{TP}{TP + FN}$$

6.4.6. Yanlış Pozitif Oranı (False Positive Rate)

Yanlış pozitif oranı (FPR), negatif örneklerin yanlışlıkla pozitif sınıfa atanma oranını ifade eder.

$$\text{yanlış pozitif oranı} = \frac{FP}{FP + TN}$$

6.4.7. F-Ölçütü (F1 Score)

F-ölçütü, kesinlik ve duyarlılığın harmonik ortalamasını alarak model performansını daha dengeli biçimde değerlendirmeye olanak tanır. Özellikle sınıfların dengesiz dağıldığı durumlarda, tek başına doğruluktan daha güvenilir bir ölçüt sunar (Witten, Frank ve Hall, 2016).

$$F - \text{Ölçütü} = 2 \times \frac{\text{keskinlik} \times \text{duyarlılık}}{\text{keskinlik} + \text{duyarlılık}}$$

6.4.8. Kappa İstatistiği

Kappa istatistiği (Cohen's Kappa), sınıflandırma doğruluğunu tesadüfi başarıdan ayırarak değerlendirir. Böylece modelin gerçek performansını daha doğru biçimde ölçer. Literatürde 0.60 üzerindeki Kappa değerleri kabul edilebilir, 0.80 üzeri ise oldukça güçlü sınıflandırma başarısı olarak değerlendirilmektedir (Landis ve Koch, 1977: 165).

Kappa değeri (k) -1 ile +1 arasında bir değer alabilir. Sim ve Wright (2005: 259) bulunan bu değeri şöyle yorumlamışlardır:

$k = +1$: iki gözlemcinin sonuçları birbiri ile tamamen uyumludur.

$k = 0$: iki gözlemci arasındaki uyum şansa bağlıdır.

$k = -1$: iki gözlemcinin sonuçları birbirinin tersini değerlendirmektedir.

Kılıç (2015: 142), kappa katsayısını hesaplarken iki olasılıktan bahsetmiş ve şu şekilde hesaplamıştır:

$P_r(a)$: İki gözlemci için gözlemlenen uyumların toplam orantısı

$P_r(e)$: Bu uyumun şansa bağlı ortaya çıkma olasılığıdır.

$$k = \frac{P_r(a) - P_r(e)}{1 - P_r(e)}$$

6.5. Araştırmanın Bulguları

Bu çalışmada WEKA 3.8.6 programı kullanılarak makine öğrenmesi temelli denetimli öğrenme algoritmaları uygulanmış ve modellerin etkinliği doğruluk, kesinlik ve F-ölçütü gibi metriklerle karşılaştırılmıştır.

6.5.1. Dengeli Veri Kümesi ve Öznitelik Seçiminin Yapılmaması

Dengeli veri kümesi, 655 özniteliğin tamamı dikkate alınarak üç ayrı sınıflandırma yöntemi ile modellenmiştir.

6.5.1.1. Naive Bayes Sınıflandırma Yöntemi ile Analiz

Dengeli veri kümesi, 655 özniteliğin tamamı dikkate alınarak Naive Bayes sınıflandırma yöntemi ile analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisi (Tablo 6.5) ve model değerlendirme ölçütleri (Tablo 6.6) aşağıda yer almaktadır.

Tablo 6.5. Naive Bayes Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	231	69	300
<u>Gerçek: Negatif</u>	65	235	300
<u>Toplam</u>	296	304	600

Tablo 6.6. Naive Bayes ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u>	<u>Hata</u>						<u>Kappa</u>
<u>Doğruluğu(%)</u>	<u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>İstatistiği</u>
77,06	22,94	0,77	0,77	0,77	0,23	0,77	0,66

6.5.1.2. SMO Sınıflandırma Yöntemi ile Analiz

Dengeli veri kümesi, 655 özniteliğin tamamı dikkate alınarak SMO sınıflandırma yöntemi ile analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisi (Tablo 6.7) ve model değerlendirme ölçütleri (Tablo 6.8) aşağıda yer almaktadır.

Tablo 6.7. SMO Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	238	62	300
<u>Gerçek: Negatif</u>	87	213	300
<u>Toplam</u>	325	275	600

Tablo 6.8. SMO ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u>	<u>Hata</u>					<u>F-</u>	<u>Kappa</u>
<u>Doğruluğu(%)</u>	<u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>Ölçütü</u>	<u>İstatistiği</u>
79,84	20,16	0,80	0,80	0,80	0,20	0,80	0,73

6.5.1.3. kNN ($k=1$) Sınıflandırma Yöntemi ile Analiz

Dengeli veri kümesi, 655 öznetiliğin tamamı dikkate alınarak kNN ($k=1$) sınıflandırma yöntemi ile Chebyshev ve Öklid olmak üzere iki farklı uzaklık ölçütü kullanılarak analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisileri (Tablo 6.9 ve Tablo 6.11) ve model değerlendirme ölçütleri (Tablo 6.10 ve Tablo 6.12) aşağıda yer almaktadır.

Tablo 6.9. kNN ($k=1$) (Chebyshev) ile Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	274	26	300
<u>Gerçek: Negatif</u>	273	27	300
<u>Toplam</u>	547	53	600

Tablo 6.10. kNN ($k=1$) (Chebyshev) ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u>	<u>Hata</u>					<u>F-Ölçütü</u>	<u>Kappa</u>
<u>Doğruluğu(%)</u>	<u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>		<u>İstatistiği</u>
62,4	37,6	0,62	0,62	0,62	0,38	0,45	0,65

Tablo 6.11. kNN ($k=1$) (Öklid) ile Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	209	91	300
<u>Gerçek: Negatif</u>	113	187	300
<u>Toplam</u>	322	278	600

Tablo 6.12. kNN ($k=1$) (Öklid) ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u> <u>Doğruluğu(%)</u>	<u>Hata</u> <u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>Kappa</u> <u>İstatistiği</u>
62,9	37,1	0,63	0,63	0,63	0,37	0,63	0,62

6.5.2. Dengeli Veri Kümesinde Öznitelik Seçiminin Yapılması

Dengeli veri kümesi, CfsSubsetEval – GainRatioAttributeEval ve Ranker kullanılarak 655 öznitelik arasından seçilen 40 öznitelik dikkate alınarak üç ayrı sınıflandırma yöntemi ile modellenmiştir.

6.5.2.1. Naive Bayes Sınıflandırma Yöntemi ile Analiz

Dengeli veri kümesi, CfsSubsetEval – GainRatioAttributeEval ve Ranker kullanılarak seçilen 40 öznitelik dikkate alınarak Naive Bayes sınıflandırma yöntemi ile analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisi (Tablo 6.13) ve model değerlendirme ölçütleri (Tablo 6.14) aşağıda yer almaktadır.

Tablo 6.13. Naive Bayes Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	222	78	300
<u>Gerçek: Negatif</u>	102	198	300
<u>Toplam</u>	324	276	600

Tablo 6.14. Naive Bayes ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u>	<u>Hata</u>						<u>Kappa</u>
<u>Doğruluğu(%)</u>	<u>Oran (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>İstatistiği</u>
75,30	24,7	0,75	0,75	0,72	0,28	0,75	0,69

6.5.2.2. SMO Sınıflandırma Yöntemi ile Analiz

Dengeli veri kümesi, CfsSubsetEval – GainRatioAttributeEval ve Ranker kullanılarak seçilen 40 öznelik dikkate alınarak SMO sınıflandırma yöntemi ile analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisi (Tablo 6.15) ve model değerlendirme ölçütleri (Tablo 6.16) aşağıda yer almaktadır.

Tablo 6.15. SMO Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	273	27	300
<u>Gerçek: Negatif</u>	163	137	300
<u>Toplam</u>	436	167	600

Tablo 6.16. SMO ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u>	<u>Hata</u>						<u>Kappa</u>
<u>Doğruluğu(%)</u>	<u>Oran (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>İstatistiği</u>
76,16	23,84	0,76	0,76	0,76	0,24	0,76	0,71

6.5.2.3. kNN ($k=1$) Sınıflandırma Yöntemi ile Analiz

Dengeli veri kümesi, CfsSubsetEval – GainRatioAttributeEval ve Ranker kullanılarak seçilen 40 öznelik dikkate alınarak kNN ($k=1$) sınıflandırma yöntemi ile Chebyshev ve Öklid olmak üzere iki farklı uzaklık ölçütü kullanılarak analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisileri (Tablo 6.17 ve Tablo 6.19) ve model değerlendirme ölçütleri (Tablo 6.18 ve Tablo 6.20) aşağıda yer almaktadır.

Tablo 6.17. kNN ($k=1$) (Chebyshev) ile Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	279	21	300
<u>Gerçek: Negatif</u>	160	145	300
<u>Toplam</u>	436	166	600

Tablo 6.18. kNN ($k=1$) (Chebyshev) ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u> <u>Doğruluğu(%)</u>	<u>Hata</u> <u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>Kappa</u> <u>İstatistiği</u>
68,89	31,11	0,69	0,72	0,69	0,31	0,69	0,68

Tablo 6.19. kNN ($k=1$) (Öklid) ile Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	274	26	300
<u>Gerçek: Negatif</u>	111	189	300
<u>Toplam</u>	385	215	600

Tablo 6.20. kNN ($k=1$) (Öklid) ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u> <u>Doğruluğu(%)</u>	<u>Hata</u> <u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>Kappa</u> <u>İstatistiği</u>
76,12	23,88	0,76	0,76	0,76	0,24	0,75	0,61

6.5.3. Dengesiz Veri Kümesinde Öznitelik Seçiminin Yapılmaması

Dengesiz veri kümesi, 1150 özniteliğin tamamı dikkate alınarak üç ayrı sınıflandırma yöntemi ile modellenmiştir.

6.5.3.1. Naive Bayes Sınıflandırma Yöntemi ile Analiz

Dengesiz veri kümesi, 1150 özniteliğin tamamı dikkate alınarak Naive Bayes sınıflandırma yöntemi ile analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisi (Tablo 6.21) ve model değerlendirme ölçütleri (Tablo 6.22) aşağıda yer almaktadır.

Tablo 6.21. Naive Bayes Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	356	3927	4283
<u>Gerçek: Negatif</u>	609	974	1583
<u>Toplam</u>	365	4901	5866

Tablo 6.22. Naive Bayes ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u>	<u>Hata</u>						<u>Kappa</u>
<u>Doğruluğu(%)</u>	<u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>İstatistiği</u>
89,57	10,43	0,90	0,90	0,90	0,56	0,90	0,62

6.5.3.2. SMO Sınıflandırma Yöntemi ile Analiz

Dengesiz veri kümesi, 1150 özniteliğin tamamı dikkate alınarak SMO sınıflandırma yöntemi ile analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisi (Tablo 6.23) ve model değerlendirme ölçütleri (Tablo 6.24) aşağıda yer almaktadır.

Tablo 6.23. SMO Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	342	3941	4283
<u>Gerçek: Negatif</u>	660	923	1583
<u>Toplam</u>	1002	4864	5866

Tablo 6.24. SMO ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u> <u>Doğruluğu(%)</u>	<u>Hata</u> <u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>Kappa</u> <u>İstatistiği</u>
92,02	7,98	0,92	0,92	0,92	0,80	0,89	0,87

6.5.3.3. kNN ($k=1$) Sınıflandırma Yöntemi ile Analiz

Dengesiz veri kümesi, 1150 özniteliğin tamamı dikkate alınarak kNN ($k=1$) sınıflandırma yöntemi ile Chebyshev ve Öklid olmak üzere iki farklı uzaklık ölçütü kullanılarak analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisileri (Tablo 6.25 ve Tablo 6.27) ve model değerlendirme ölçütleri (Tablo 6.26 ve Tablo 6.28) aşağıda yer almaktadır.

Tablo 6.25. kNN ($k=1$) (Chebyshev) ile Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	38	4245	4283
<u>Gerçek: Negatif</u>	276	1307	1583
<u>Toplam</u>	314	5552	5866

Tablo 6.26. kNN ($k=1$) (Chebyshev) ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u> <u>Doğruluğu(%)</u>	<u>Hata</u> <u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>Kappa</u> <u>İstatistiği</u>
91,19	8,81	0,91	0,91	0,91	0,76	0,87	0,80

Tablo 6.27. kNN ($k=1$) (Öklid) ile Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	104	4179	4283
<u>Gerçek: Negatif</u>	298	1285	1583
<u>Toplam</u>	402	5464	5866

Tablo 6.28. kNN ($k=1$) (Öklid) ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u> <u>Doğruluğu(%)</u>	<u>Hata</u> <u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>Kappa</u> <u>İstatistiği</u>
93,31	6,69	0,93	0,92	0,93	0,55	0,92	0,74

6.5.4. Dengesiz Veri Kümesinde Öznitelik Seçiminin Yapılması

Dengesiz veri kümesi için CfsSubsetEval – GainRatioAttributeEval ve Ranker kullanılarak 1550 öznitelik arasından seçilen 256 öznitelik dikkate alınarak üç ayrı sınıflandırma yöntemi ile modellenmiştir.

6.5.4.1. Naive Bayes Sınıflandırma Yöntemi ile Analiz

Dengesiz veri kümesi için CfsSubsetEval – GainRatioAttributeEval ve Ranker kullanılarak 1150 öznitelik arasından seçilen 256 öznitelik dikkate alınarak Naive Bayes sınıflandırma yöntemi ile analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisi (Tablo 6.29) ve model değerlendirme ölçütleri (Tablo 6.30) aşağıda yer almaktadır.

Tablo 6.29. Naive Bayes Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	1427	2856	4283
<u>Gerçek: Negatif</u>	167	1416	1583
<u>Toplam</u>	1594	4272	5866

Tablo 6.30. Naive Bayes ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u> <u>Doğruluğu(%)</u>	<u>Hata</u> <u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>Kappa</u> <u>İstatistiği</u>
91,75	8,25	0,92	0,91	0,92	0,63	0,90	0,66

6.5.4.2. SMO Sınıflandırma Yöntemi ile Analiz

Dengesiz veri kümesi için CfsSubsetEval – GainRatioAttributeEval ve Ranker kullanılarak 1150 öznitelik arasından seçilen 256 öznitelik dikkate alınarak SMO sınıflandırma yöntemi ile analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisi (Tablo 6.31) ve model değerlendirme ölçütleri (Tablo 6.32) aşağıda yer almaktadır.

Tablo 6.31. SMO Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	1522	2761	4283
<u>Gerçek: Negatif</u>	143	1440	1583
<u>Toplam</u>	1665	4201	5866

Tablo 6.32. SMO ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u>	<u>Hata</u>						<u>Kappa</u>
<u>Doğruluğu(%)</u>	<u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>İstatistiği</u>
94,90	5,1	0,94	0,94	0,94	0,50	0,95	0,82

6.5.4.3. kNN ($k=1$) Sınıflandırma Yöntemi ile Analiz

Dengesiz veri kümesi, CfsSubsetEval – GainRatioAttributeEval ve Ranker kullanılarak 1150 öznitelik arasından seçilen 256 öznitelik dikkate alınarak kNN ($k=1$) sınıflandırma yöntemi ile Chebyshev ve Öklid olmak üzere iki farklı uzaklık ölçütü kullanılarak analiz edilmiştir. WEKA 3.8.6 programı kullanılarak elde edilen hata matrisileri (Tablo 6.33 ve Tablo 6.35) ve model değerlendirme ölçütleri (Tablo 6.34 ve Tablo 6.36) aşağıda yer almaktadır.

Tablo 6.33. kNN ($k=1$) (Chebyshev) ile Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	1498	2785	4283
<u>Gerçek: Negatif</u>	52	1531	1583
<u>Toplam</u>	1550	4316	5866

Tablo 6.34. kNN ($k=1$) (Chebyshev) ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u> <u>Doğruluğu(%)</u>	<u>Hata</u> <u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>Kappa</u> <u>İstatistiği</u>
92,01	7,99	0,92	0,92	0,92	0,48	0,92	0,78

Tablo 6.35. kNN ($k=1$) (Öklid) ile Sınıflandırma Yöntemine ait Hata Matrisi

	<u>Tahmin: Pozitif</u>	<u>Tahmin: Negatif</u>	<u>Toplam</u>
<u>Gerçek: Pozitif</u>	98	4185	4283
<u>Gerçek: Negatif</u>	149	1434	1583
<u>Toplam</u>	247	5619	5866

Tablo 6.36 kNN ($k=1$) (Öklid) ile Yapılan Modelin Değerlendirme Ölçütleri

<u>Sınıflandırma</u> <u>Doğruluğu(%)</u>	<u>Hata</u> <u>Oranı (%)</u>	<u>Duyarlılık</u>	<u>Kesinlik</u>	<u>TPR</u>	<u>FPR</u>	<u>F-Ölçütü</u>	<u>Kappa</u> <u>İstatistiği</u>
93,64	6,36	0,94	0,93	0,94	0,52	0,96	0,72

6.6. Araştırmanın Çıktıları

Bu çalışmada, WEKA 3.8.6 yazılımı kullanılarak makine öğrenmesine dayalı denetimli öğrenme algoritmalarından Naive Bayes, SMO ve k-NN ($k=1$) yöntemleri uygulanmış, elde edilen modellerin performansına ilişkin tüm çıktılar Tablo 6.37’de sunulmuştur.

Tablo 6.37. Analiz Sonuçlarının Karşılaştırılması

Sınıflandırma Yöntemi	Veri Kümesi	Öznitelik Seçimi	Sınıflandırma Doğruluğu(%)	Hata Oranı (%)	Duyarlılık	Kesinlik	TP Oranı	FP Oranı	F-Ölçütü	Kappa İstatistiği
Naives Bayes	Dengeli	Yok	77,06	22,94	0,77	0,77	0,77	0,23	0,77	0,66
		Var	75,30	24,7	0,75	0,75	0,75	0,28	0,75	0,69
	Dengesiz	Yok	89,57	10,43	0,90	0,90	0,90	0,56	0,90	0,62
		Var	91,75	8,25	0,92	0,91	0,92	0,63	0,90	0,66
SMO	Dengeli	Yok	79,84	20,16	0,80	0,80	0,80	0,20	0,80	0,73
		Var	76,16	23,84	0,76	0,76	0,76	0,24	0,76	0,71
	Dengesiz	Yok	92,02	7,98	0,92	0,92	0,92	0,80	0,89	0,87
		Var	94,90	5,1	0,94	0,94	0,94	0,50	0,95	0,82
kNN ($k=1$) Chebysey	Dengeli	Yok	62,4	37,6	0,62	0,62	0,62	0,38	0,62	0,65
		Var	68,89	31,11	0,69	0,72	0,69	0,31	0,69	0,68
	Dengesiz	Yok	91,19	8,81	0,91	0,91	0,91	0,76	0,87	0,80
		Var	92,01	7,99	0,92	0,92	0,92	0,48	0,92	0,78
kNN ($k=1$) Öklid	Dengeli	Yok	62,9	37,1	0,63	0,63	0,63	0,37	0,63	0,62
		Var	76,12	23,88	0,76	0,76	0,76	0,24	0,75	0,61
	Dengesiz	Yok	93,31	6,69	0,93	0,92	0,93	0,55	0,92	0,74
		Var	93,64	6,36	0,94	0,93	0,94	0,52	0,96	0,72

Sınıflandırma algoritmalarının başarısının değerlendirilmesinde en temel ölçütlerden biri sınıflandırma doğruluğu (accuracy) oranıdır. Ancak kabul edilebilir doğruluk oranı, veri kümesinin özelliklerine ve problem alanına bağlı olarak değişiklik göstermektedir. Literatürde genellikle %70'in üzerinde elde edilen doğruluk oranları kabul edilebilir düzeyde görülmekte, %80–90 aralığı başarılı, %90'ın üzeri ise oldukça yüksek başarı olarak değerlendirilmektedir (Han vd.,2012; Witten vd., 2016). Doğal dil işleme, duygu analizi, sosyal medya verisi gibi karmaşık problemlerde %70–90 doğruluk oranının kabul edilebilir olduğunu söyleyebiliriz.

Bu çalışmada dengeli veri kümesine ait sınıflandırma doğruluğu ortalaması % 72,3 iken dengesiz veri kümesine ait sınıflandırma doğruluğu ortalaması %92,3'tür. Dengeli veri kümesinde sınıflandırma doğruluğu oranı en yüksek model %79,84 ile SMO, sınıflandırma doğruluğu oranı en düşük model %62,4 ile kNN ($k=1$) Chebyshev uzaklık ölçütü ile yapılan analize aittir. Dengesiz veri kümesinde sınıflandırma doğruluğu oranı en yüksek model %94,9 ile SMO analizine aittir. Dengesiz veri kümesinde sınıflandırma doğruluğu oranı en düşük model ise %89,57 ile Naive Bayes olmuştur.

Dengeli veri kümesinde öznitelik seçiminin yapılmadığı modellerde sınıflandırma doğruluğu en yüksek %79,87 oranı ile SMO analizine aittir. Sonra sırasıyla %77,06 oranıyla Naive Bayes, %62,9 oranıyla kNN ($k=1$) Öklid ve son sırada %62,4 oranıyla kNN ($k=1$) Chebyshev uzaklık ölçütü ile yapılan analize aittir.

Dengeli veri kümesinde öznitelik seçiminin yapıldığı modellerde sınıflandırma doğruluğu en yüksek %76,16 oranı ile SMO analizine aittir. Sonra sırasıyla %73,12 oranıyla kNN ($k=1$) Öklid, %72,3 oranıyla Naive Bayes ve son sırada %68,89 oranıyla kNN ($k=1$) Chebyshev uzaklık ölçütü ile yapılan analize aittir.

Dengesiz veri kümesinde öznitelik seçiminin yapılmadığı modellerde sınıflandırma doğruluğu en yüksek %93,31 oranı ile kNN ($k=1$) Öklid uzaklık ölçütü ile yapılan analize aittir. Sonra sırasıyla %92,02 oranıyla SMO, %91,19 oranıyla kNN ($k=1$) Chebyshev ve son sırada %89,57 oranıyla Naive Bayes analizine aittir.

Dengesiz veri kümesinde öznitelik seçiminin yapıldığı modellerde sınıflandırma doğruluğu en yüksek %94,9 oranı ile SMO analizine aittir. Sonra sırasıyla %93,64 oranıyla kNN ($k=1$) Öklid, %92,01 oranıyla SMO kNN ($k=1$) Chebyshev ve son sırada %91,75 oranıyla Naive Bayes analizine aittir.

Bu çalışmada elde ettiğimiz yukarıdaki sonuçlara göre dengesiz veri kümesinin, dengeli veri kümesine göre daha iyi performans sağladığı gözlemlenmiştir. Dengeli veri kümesinin performansının düşük olmasının sebebi veri sayısının yetersiz kalmış olması olabilir.

Dengeli veri kümesinde analiz yapılırken öznitelik seçiminin yapılmasının, öznitelik seçiminin yapılmadığı modellere göre kayda değer bir performans değişikliğine neden olmadığı gözlemlenmiştir.

Dengesiz veri kümesinde analiz yapılırken öznitelik seçiminin yapılmasının, öznitelik seçiminin yapılmadığı modellere göre daha iyi performans sağladığını söyleyebiliriz. Dengesiz veri kümesinde öznitelik seçimi yapılarak elde edilen %94,9 sınıflandırma doğruluğu oranı ile SMO en iyi performansı gösteren sınıflandırma algoritması olmuştur.

Sınıflandırma modellerinin performansı yalnızca doğruluk oranı ile değil, duyarlılık (Recall), kesinlik (Precision), yanlış pozitif oranı (FPR) ve F-ölçütü (F1-score) gibi ek metriklerle de değerlendirilmelidir. Literatürde genellikle 0.70 ve üzerindeki Precision, Recall ve F1 değerleri kabul edilebilir; 0.80 üzeri iyi, 0.90 ve üzeri ise yüksek başarı olarak nitelendirilmektedir. Yanlış pozitif oranının ise mümkün olduğunca düşük tutulması, özellikle kritik uygulamalarda %5'in altına indirilmesi önerilmektedir (Powers, 2011; Han vd., 2012). Bu çalışmada duyarlılık, kesinlik ve F-ölçütü değerleri kNN ($k=1$) Chebyshev ve kNN ($k=1$) Öklid için %62 ve %63'ken, diğer modellerde %75 ile %96 değerleri arasındadır. Yanlış pozitif oranlarının da kabul edilebilir düzeyde olduğunu söyleyebiliriz. Analiz sonucunda elde ettiğimiz duyarlılık, kesinlik, yanlış pozitif oranı (FPR) ve F-ölçütü kriterlerinin değerlerine baktığımızda dengesiz veri kümesi, dengeli veri kümesine göre daha iyi bir performans gösterdiğini söylemek mümkündür. Öznitelik seçimi yapılmasının, öznitelik seçimi yapılmamasına göre ise kayda değer bir etkisi olmamıştır.

Kappa istatistiği, sınıflandırma doğruluğunu şansa bağlı başarıdan ayırarak modelin gerçek performansını ortaya koymaktadır. Literatürde 0.60'ın üzerindeki değerler kabul edilebilir düzeyde, 0.80'in üzerindeki değerler ise oldukça güçlü bir sınıflandırma başarısı olarak değerlendirilmektedir (Landis ve Koch, 1977). Bu çalışmada kappa istatistiği sonuçlarına baktığımızda veri kümesinin dağılımına ve öznitelik seçiminin yapıp yapılmamasına bağlı olarak değişiklik gösterdiğini söyleyebiliriz. Veri kümesi dağılımı ve öznitelik seçimine göre SMO algoritmasının sınıflandırma başarısının diğerlerine göre yüksek olduğunu söyleyebiliriz.

6.7. Tartışma

Bu çalışmada deprem uygulamalarına yapılan kullanıcı yorumlarının sınıflandırılması için Naive Bayes, SMO ve kNN algoritmalarının performansları farklı veri kümesi dağılımları (dengeli/dengesiz) ve öznitelik seçimi durumları açısından karşılaştırılmıştır. Elde edilen bulgular, literatürde metin madenciliği ve duygu analizi alanında yapılan çalışmalarla uyumlu sonuçlar ortaya koymuştur. Genel olarak sınıflandırma doğruluğu açısından dengesiz veri kümesinin daha başarılı olduğu, SMO algoritmasının ise her iki veri kümesi türünde de diğer algoritmalarla kıyasla öne çıktığı görülmüştür.

Çalışmanın en dikkat çekici bulgularından biri, dengesiz veri kümesinin %92,3 ortalama doğruluk oranıyla dengeli veri kümesine kıyasla daha yüksek performans göstermesidir. Dengeli veri kümesinde ortalama doğruluğun %72,3'te kalması, literatürde küçük ve dengeli veri kümelerinin sınıflandırma modelleri açısından öğrenme kapasitesini sınırladığına yönelik bulgularla örtüşmektedir (Han vd., 2012; Witten vd., 2016). Özellikle doğal dil işleme ve duygu analizi gibi karmaşık problemlerde veri setinin büyüklüğü ve çeşitliliği, modelin genelleme başarısı üzerinde kritik bir rol oynamaktadır. Bu bağlamda, dengesiz veri kümesinde örnek sayısının fazla olması, modellerin daha yüksek doğruluk oranları elde etmesine katkı sağlamıştır.

Öznitelik seçimi, sınıflandırma algoritmalarının performansını artırmada sıklıkla önerilen bir yöntemdir. Ancak bu çalışmada öznitelik seçiminin etkisi, veri kümesinin dağılımına göre farklılık göstermiştir. Dengeli veri kümesinde öznitelik seçiminin kayda değer bir iyileşme sağlamadığı, hatta bazı durumlarda sınıflandırma doğruluğunu düşürdüğü gözlemlenmiştir. Bu durum, düşük sayıda örneğe sahip veri setlerinde öznitelik seçiminin modelin öğrenme kapasitesini sınırlayabileceğini göstermektedir. Öte yandan, dengesiz veri kümesinde öznitelik seçiminin yapılmasıyla birlikte SMO algoritmasının %94,9 doğruluk oranına ulaşması, daha büyük veri setlerinde öznitelik seçiminin daha etkin bir şekilde katkı sağlayabileceğini ortaya koymaktadır.

Algoritmaların performansları karşılaştırıldığında, SMO'nun hem dengeli hem de dengesiz veri kümelerinde en iyi sonuçları verdiği görülmüştür. Bu durum, destek vektör makinelerinin (SVM) yüksek boyutlu ve karmaşık veri setlerinde güçlü sınıflandırma yetenekleri sunduğunu ortaya koyan literatürle uyumludur (Cortes ve Vapnik, 1995). Naive Bayes algoritması, özellikle dengesiz veri kümesinde orta düzeyde performans göstermiştir. Bu sonuç, Naive Bayes'in basit ve hızlı bir algoritma olmasına rağmen veri setindeki karmaşık ilişkileri yeterince iyi yakalayamadığını düşündürmektedir. kNN algoritması ise, özellikle

Chebyshev uzaklık ölçütü ile en düşük performansı göstermiştir. Bu bulgu, kNN'nin uzaklık ölçütüne duyarlılığını ve veri dağılımına göre performansının değişebileceğini göstermektedir (Han vd., 2012).

Performans değerlendirmesinde yalnızca doğruluk oranına odaklanmak yanıltıcı olabileceğinden, ek ölçütler olan duyarlılık (Recall), kesinlik (Precision), yanlış pozitif oranı (FPR) ve F-ölçütü (F1-score) de incelenmiştir. Bulgulara göre, SMO ve Naive Bayes modellerinde bu değerlerin çoğunlukla %75–96 aralığında olduğu, kNN'de ise bazı durumlarda %62–63 seviyelerinde kaldığı görülmüştür. Literatürde genellikle %70 ve üzeri değerlerin kabul edilebilir, %80 üzerinin başarılı, %90 ve üzerinin ise yüksek başarı olarak değerlendirildiği dikkate alındığında (Powers, 2011), bu çalışmada elde edilen metriklerin kabul edilebilir düzeyde olduğu söylenebilir. Ayrıca yanlış pozitif oranlarının düşük seviyelerde kalması, özellikle afet yönetimi gibi kritik uygulamalarda büyük önem arz etmektedir.

Kappa istatistiği sonuçları da dikkate alındığında, SMO'nun diğer algoritmalarla kıyasla daha güvenilir sınıflandırmalar yaptığı görülmüştür. Literatürde 0.60 üzerindeki değerlerin kabul edilebilir, 0.80 üzerindeki değerlerin ise oldukça güçlü sınıflandırma başarısı olarak değerlendirildiği göz önüne alındığında (Landis ve Koch, 1977), bu çalışmada elde edilen Kappa değerleri SMO'nun üstünlüğünü bir kez daha ortaya koymaktadır. Bu bulgu, yalnızca doğruluk oranına değil, aynı zamanda şansa bağlı sınıflandırma başarısını ayırt etmeye yarayan ek metriklerin de dikkate alınmasının gerekliliğini ortaya koymaktadır.

Sonuç olarak, elde edilen bulgular deprem uygulamalarına yönelik kullanıcı yorumlarının sınıflandırılmasında SMO algoritmasının güçlü bir alternatif olduğunu göstermektedir. Dengesiz veri kümelerinin daha yüksek başarı sunması, afet gibi kritik alanlarda veri çeşitliliği ve miktarının artırılmasının model başarısını doğrudan etkileyebileceğini ortaya koymaktadır. Bununla birlikte, dengeli veri kümelerinde düşük örnek sayısı nedeniyle sınıflandırma doğruluğunun sınırlı kalması, gelecekte yapılacak çalışmalarda daha büyük ve dengeli veri setlerinin oluşturulmasının önemini göstermektedir. Ayrıca, farklı öznitelik seçimi yöntemlerinin denenmesi ve algoritma parametrelerinin optimize edilmesi ile daha yüksek sınıflandırma performansları elde edilebileceği öngörülmektedir.

Deprem uygulamalarına yönelik memnuniyet düzeyi açısından analiz sonuçlarına bakıldığında, kullanıcıların büyük çoğunluğunun deprem uygulamalarına yönelik memnuniyet düzeyinin yüksek olduğunu ortaya koymaktadır. Olumlu etiketlenmiş yorumların oranı %54,94 iken, olumsuz etiketlenmiş yorumların oranı %18,83 olarak belirlenmiştir. Nötr yorumlar ise toplamın %30,2'sini oluşturmaktadır. Olumlu geri bildirimlerin olumsuz geri bildirimlere

oranla yaklaşık 2,7 kat daha fazla olması, genel kullanıcı deneyiminin pozitif bir eğilim sergilediğini göstermektedir. Bu durum, uygulamaların temel işlevlerini büyük ölçüde başarıyla yerine getirdiğini ve kullanıcıların ihtiyaçlarını karşılamada etkili olduğunu düşündürmektedir. Literatürde, mobil afet uygulamalarına yönelik yüksek memnuniyet düzeylerinin, kullanıcıların uygulamaları daha sık kullanma ve bu uygulamalara güven duyma eğilimini artırdığı belirtilmektedir (Palen & Liu, 2007; Tan vd., 2017). Bu çalışmanın bulguları da söz konusu görüşü desteklemektedir. Görece yüksek memnuniyet oranı, kullanıcıların deprem uygulamalarını düzenli olarak kullanma motivasyonunu ve uygulamalara yönelik güven algısını güçlendirebilecek bir unsur olarak değerlendirilebilir.

Bununla birlikte, %18,83'lük olumsuz geri bildirim oranı, uygulamaların geliştirilmesi açısından dikkate alınması gereken önemli bir veri kaynağıdır. Literatürde benzer çalışmalarda da (Paul vd., 2021) olumsuz yorumların genellikle erken uyarı süresinin yetersizliği, yanlış veya gecikmeli bildirimler, kullanıcı arayüzünün karmaşıklığı, çevrimdışı kullanım eksikliği ve afet sonrası bilgilendirme yetersizlikleri gibi alanlarda yoğunlaştığı belirtilmektedir. Bu çalışmanın bulguları da söz konusu sorun alanlarını doğrulamaktadır.

6.7.1. Bulguların Literatür ile Karşılaştırılması

Bu araştırmada elde edilen sonuçlar, literatürde yer alan çeşitli çalışmaların bulguları ile karşılaştırıldığında büyük ölçüde paralellik göstermektedir. Özellikle, hem dengeli hem de dengesiz veri kümelerinde Sıralı Minimal Optimizasyon (SMO) algoritmasının en yüksek doğruluk oranlarını vermesi, Göker & Tekedere (2017), Dhiyaulhaq & Gunawan (2023) ve Martin (2025) tarafından rapor edilen sonuçlarla örtüşmektedir. Ayrıca, kNN algoritmasının özellikle küçük veri setlerinde veya uzaklık ölçütüne duyarlı senaryolarda düşük performans sergilediğine dair bulgular, bu çalışmada dengeli veri kümesinde kNN ($k=1$) Chebyshev uzaklık ölçütü ile elde edilen %62,4 doğruluk oranı ile desteklenmektedir.

Dengeli ve dengesiz veri kümelerinin performans karşılaştırmasında ise bu çalışma, dengesiz veri kümesinin daha yüksek doğruluk sağladığını ortaya koymuştur (%92,3 ortalama doğruluk), ki bu durum Işık (2019) ile George & Srividhya (2022)'nin benzer bulgularıyla uyumludur. Bununla birlikte, bazı literatür çalışmalarında dengeli veri setlerinin daha iyi sonuçlar verdiği görülmekte olup (ör. bazı derin öğrenme uygulamaları), bu farklılık veri kümesinin boyutu, sınıf dağılımı ve kullanılan öznitelik seti gibi faktörlerden kaynaklanabilir.

Öznitelik seçiminin etkisi açısından, dengeli veri kümesinde bu işlemin doğruluk oranlarında belirgin bir değişim yaratmadığı gözlemlenmiştir. Bu bulgu, Amir Asiaee vd.

(2012) tarafından rapor edilen, boyut indirgeme veya öznitelik seçimi tekniklerinin doğruluk oranını önemli ölçüde değiştirmedığı yönündeki sonuçlarla tutarlıdır. Dengesiz veri kümesinde ise öznitelik seçiminin sınırlı da olsa performans artışı sağladığı tespit edilmiştir.

Ek değerlendirme metrikleri (duyarlılık, kesinlik, F-ölçütü) bakımından, çalışmada elde edilen değerlerin kNN Chebyshev ve kNN Öklid dışındaki tüm modellerde %75 ile %96 aralığında olduğu ve literatürde kabul edilen “iyi” veya “yüksek” başarı düzeylerini karşıladığı görülmektedir (Powers, 2011). Yanlış pozitif oranlarının da kritik uygulamalar için önerilen %5 eşiğinin altında olduğu belirlenmiştir (Han vd., 2012).

Genel olarak, bu çalışma literatürdeki baskın eğilimleri doğrulamakta, ancak dengesiz veri kümelerinin üstün performans göstermesi bakımından bazı çalışmalardan ayrılmaktadır. Bu durum, veri miktarının yetersizliği, veri setinin yapısı ve öznitelik özellikleri gibi faktörlerin model başarımı üzerinde belirleyici olabileceğini göstermektedir.

7. SONUÇ

Bu çalışmada, deprem uygulamalarına yönelik kullanıcı yorumları üzerinde metin madenciliği ve duygu analizi yöntemleri kullanılarak Naive Bayes, SMO ve kNN algoritmalarının sınıflandırma performansları karşılaştırılmıştır. Çalışmada dengeli ve dengesiz veri kümeleri üzerinde öznitelik seçimi yapılmış ve yapılmamış durumlar incelenmiştir. Bulgular, sınıflandırma doğruluğu açısından dengesiz veri kümesinin daha başarılı olduğunu, SMO algoritmasının ise her iki veri kümesi türünde de diğer algoritmalara kıyasla en yüksek performansı sergilediğini göstermektedir.

Dengeli veri kümesinde ortalama doğruluk oranı %72,3'te kalırken, dengesiz veri kümesinde bu oran %92,3'e ulaşmıştır. Bu durum, örnek sayısının artmasıyla birlikte modelin öğrenme kapasitesinin ve genelleme başarısının yükseldiğini göstermektedir. Ayrıca, dengesiz veri kümesinde öznitelik seçiminin yapılması performansı artırmış ve SMO algoritması %94,9 ile en yüksek başarıyı sağlamıştır. Bu bulgu, daha büyük veri setlerinde öznitelik seçiminin sınıflandırma performansına olumlu katkı sağlayabileceğini göstermektedir.

Algoritmaların karşılaştırılmasında SMO'nun üstünlüğü açıkça görülmüş, Naive Bayes'in orta düzeyde, kNN'nin ise özellikle Chebyshev uzaklık ölçütü ile düşük performans gösterdiği ortaya çıkmıştır. Ek performans ölçütleri olan Precision, Recall, F1 ve FPR değerleri incelendiğinde, SMO ve Naive Bayes'in genellikle kabul edilebilir ve yüksek düzeylerde performans sergilediği, kNN'nin ise bazı durumlarda düşük seviyelerde kaldığı belirlenmiştir. Kappa istatistiği sonuçları da SMO'nun daha güvenilir ve güçlü bir sınıflandırma başarısına sahip olduğunu desteklemiştir.

Bu bulgular ışığında, çalışmanın temel katkıları şu şekilde özetlenebilir:

1. Afet yönetimi ve deprem uygulamaları bağlamında kullanıcı yorumlarının analizi yapılmış ve sınıflandırma algoritmalarının etkinliği ortaya konmuştur.
2. Veri kümesi dağılımının (dengeli/dengesiz) sınıflandırma performansına etkisi incelenmiş ve dengesiz veri kümesinin daha başarılı sonuçlar verdiği gözlemlenmiştir.
3. Öznitelik seçiminin etkisi analiz edilmiş, dengesiz veri kümesinde performansı artırdığı, ancak dengeli veri kümesinde belirgin bir iyileşme sağlamadığı görülmüştür.
4. SMO algoritmasının üstünlüğü özellikle vurgulanmış, afet yönetimi bağlamında metin sınıflandırma görevlerinde güçlü bir alternatif olduğu sonucuna varılmıştır.

Yapılan analizler sonucunda, olumlu etiketlenmiş yorum sayısı 4.283, olumsuz etiketlenmiş yorum sayısı ise 1.583 olarak belirlenmiştir. Bu durum, olumlu geri bildirimlerin olumsuz geri bildirimlere oranla yaklaşık 2,7 kat daha fazla olduğunu göstermektedir. Elde edilen bulgular, kullanıcıların büyük çoğunluğunun deprem uygulamalarına yönelik memnuniyet düzeyinin görece yüksek olduğunu (tüm yorumların %54,94'ü) ve genel kullanıcı deneyiminin pozitif bir eğilim sergilediğini ortaya koymaktadır.

Her ne kadar genel memnuniyet düzeyi yüksek olsa da 1.583 olumsuz yorum (tüm yorumların % 18,83'ü) uygulamaların geliştirilmesi açısından dikkate alınması gereken önemli bir geri bildirim kaynağını oluşturmaktadır. Literatür bulguları ve benzer çalışmalardan elde edilen çıkarımlar ışığında, bu olumsuz geri bildirimlerin şu başlıklar altında yoğunlaştığı değerlendirilmektedir:

Erken uyarı süresinin yetersizliği: Deprem bildirimlerinin daha hızlı ve güvenilir bir şekilde iletilmesine yönelik iyileştirmelerin yapılmaması.

Yanlış veya yanıltıcı bildirimler: Algoritmaların doğruluk oranlarının düşük olması.

Kullanıcı arayüzü sorunları: Daha sade, anlaşılır ve hızlı kullanılabilir bir arayüz tasarımının benimsenmemiş olması.

Çevrimdışı kullanım imkânı: İnternet bağlantısının bulunmadığı durumlarda temel bilgilere erişim sağlanamaması.

Kapsamlı bilgilendirme: Deprem sonrası güvenli alanların, acil durum iletişim bilgilerinin ve ilk yardım yönergelerinin kullanıcılara sunulmaması.

Olumlu geri bildirim oranının yüksek olması, kullanıcıların uygulamaları kullanma sıklığını artıracak ve uygulamalara yönelik güven algısını güçlendirebilecek bir faktör olarak değerlendirilebilir. Bu durum, uygulamaların temel işlevlerini büyük ölçüde başarıyla yerine getirdiğini göstermektedir. Bununla birlikte, 1.583 olumsuz yorum (%18,83) ise kullanıcı memnuniyetini sınırlayan faktörleri ortaya koymak açısından dikkate alınmalıdır. Olumsuz geri bildirimlerin analizi, uygulamaların teknik performansının, kullanıcı deneyiminin ve acil durum bilgilendirme yeteneklerinin geliştirilmesi açısından önemli fırsatlar barındırmakta ve erken uyarı süresinin yetersizliği, yanlış veya gecikmeli bildirimler, kullanıcı arayüzü karmaşıklığı ve deprem sonrası bilgilendirme eksiklikleri gibi alanlarda geliştirme gereksinimlerine işaret etmektedir.

8. ÖNERİLER

Bu araştırmada elde edilen bulgular, belirli sınırlılıklar çerçevesinde değerlendirilmiş olmakla birlikte, gelecekte yapılacak çalışmalara yönelik çeşitli öneriler sunulabilir.

Öncelikle, veri çeşitliliğinin artırılması önerilmektedir. Bu çalışmada yalnızca deprem uygulamalarına yapılan kullanıcı yorumları incelenmiştir. İlerleyen çalışmalarda sosyal medya platformları, haber siteleri, çevrimiçi forumlar ve diğer dijital mecralardaki kullanıcı yorumlarının da analize dahil edilmesi, daha geniş ve temsil gücü yüksek veri kümelerinin oluşturulmasına katkı sağlayacaktır.

İkinci olarak, dil çeşitliliği dikkate alınabilir. Araştırmada yalnızca Türkçe yorumlar kullanılmıştır. Farklı dillerdeki yorumların analize dahil edilmesi, kültürler arası karşılaştırmalar yapılmasına imkân tanıyacak ve kullanıcı deneyimlerinin evrensel yönleri hakkında daha kapsamlı bulgular elde edilmesini sağlayacaktır.

Üçüncü olarak, algoritma çeşitliliğinin genişletilmesi önem arz etmektedir. Bu çalışmada Naive Bayes, SMO ve kNN algoritmalarının performansı karşılaştırılmıştır. Gelecekte yapılacak araştırmalarda Karar Ağacı, Rastgele Orman, Gradyan Artırma, XGBoost, Destek Vektör Makinelerinin farklı varyasyonları veya derin öğrenme temelli yöntemler (ör. LSTM, CNN, Transformer tabanlı modeller) kullanılarak karşılaştırmalı analizlerin yapılması, sınıflandırma başarımı açısından daha güçlü sonuçlar sağlayabilir.

Dördüncü olarak, dengesiz veri kümeleri için yeni yaklaşımlar araştırılabilir. Bu çalışmada kullanılan veri kümesinde yorumların dağılımı tam anlamıyla dengeli değildir. Gelecekte yapılacak çalışmalarda veri dengeleme yöntemleri (ör. SMOTE, ADASYN, sınıf ağırlıklandırma teknikleri) kullanılarak dengesiz sınıflar üzerinde daha tutarlı sonuçlar elde edilebilir.

Beşinci olarak, gelişmiş metin işleme tekniklerinin uygulanması önerilmektedir. Bu çalışmada temel metin madenciliği yöntemleri kullanılmıştır. Gelecekte yapılacak araştırmalarda sözdizimsel ve anlamsal analiz, kelime gömme yöntemleri (Word2Vec, GloVe, FastText), bağlamsal dil modelleri (BERT, RoBERTa vb.) gibi yöntemler kullanılarak daha derinlemesine analizler gerçekleştirilebilir.

Altıncı olarak, zaman boyutunun analize dahil edilmesi yararlı olabilir. Kullanıcı yorumlarının belirli bir zaman diliminde incelenmiş olması, bulguların dönemsel özellikler göstermesine neden olabilir. Gelecekte yapılacak boylamsal çalışmalarda farklı zaman

aralıklarında toplanan veriler incelenerek kullanıcı duygu ve tutumlarındaki değişimlerin izlenmesi mümkündür.

Son olarak, uygulama alanının genişletilmesi önerilmektedir. Bu çalışmada yalnızca deprem uygulamaları incelenmiştir. Gelecekte yapılacak çalışmalarda afet yönetimiyle ilgili diğer mobil uygulamalar (yangın, sel, sağlık acil durum uygulamaları vb.) da kapsam dahiline alınarak, afet yönetiminde kullanıcı deneyimlerine ilişkin daha bütüncül sonuçlara ulaşılabilir.

Kısacası, gelecekte yapılacak çalışmaların odak noktası daha büyük ve çeşitli veri setleri, gelişmiş algoritmalar, zengin öznitelik çıkarım yöntemleri, çok sınıflı ve çok dilli duygu analizi ile hibrit modellerin geliştirilmesi olabilir. Bu gelişmeler, hem akademik açıdan alana katkı sağlayacak hem de afet yönetimi uygulamalarının daha etkili ve kullanıcı odaklı hale getirilmesine yardımcı olacaktır.

KAYNAKÇA

- Agarwal, B., Xie, B., Vovsha, I., Rambow, O., & Passonneau, R. (2011). Sentiment analysis of Twitter data. *Proceedings of the Workshop on Language in Social Media*, 30–38. <https://aclanthology.org/W11-0705/>
- Aggarwal, C. C., & Zhai, C. (2012). *Mining Text Data*. Springer. <https://doi.org/10.1007/978-1-4614-3223-4>
- Aghila, G., & Vidhya, K. A. (2010). A survey of Naïve Bayes machine learning approach in text document classification. *International Journal of Computer Science and Information Security*, 7(2), 206–211. <https://arxiv.org/pdf/1003.1795>
- Allen, R. M., Gasparini, P., Kamigaichi, O., & Bose, M. (2009). *The status of earthquake early warning around the world: An introductory overview*. *Seismological Research Letters*, 80(5), 682–693. <https://doi.org/10.1785/gssrl.80.5.682>
- Allen, R. M., & Melgar, D. (2019). Earthquake early warning: Advances, scientific challenges, and societal needs. *Annual Review of Earth and Planetary Sciences*, 47(1), 361–388. <https://doi.org/10.1146/annurev-earth-053018-060457>
- Amir Asiaee T., Banerjee, A., Tepper, M., & Sapiro, G. (2012). If You are Happy and You Know It... Tweet. In *Proceedings of the 21st ACM International Conference on Information and Knowledge Management*, 29 Ekim, New York, s. 1602–1606. <https://doi.org/10.1145/2396761.239848>
- Arslan, N. (2018). Özelleştirilmiş Naive Bayes algoritması. *ResearchGate*. [Erişim: 17.08.2025, https://www.researchgate.net/publication/326635209_OZELLESTIRILMIS_NAIVE_BAYES_ALGORITMASI]
- Azevedo, A., & Santos, M. F. (2008). KDD, SEMMA and CRISP-DM: A parallel overview. In *Proceedings of the IADIS European Conference on Data Mining*, 24–26 Temmuz, Amsterdam, s. 182–185. [Erişim: 17.08.2025, https://www.researchgate.net/publication/220969845_KDD_semma_and_CRISP-DM_A_parallel_overview#fullTextFileContent]
- Azhar A., Masruroh S.U., Wardhani L. K., & Okfalisa O. (2023). Performance comparison of the Naive Bayes algorithm and the k-NN lexicon approach on Twitter media sentiment analysis. *Sintech Journal of Information and Communication Technology*, 3(2), s. 35–40. <https://doi.org/10.59190/stc.v3i2.229>

- Beckmann, M., Ebecken, N., & Pires de Lima, B.** (2015). A KNN undersampling approach for data balancing. *Journal of Intelligent Learning Systems and Applications*, 7(4), s. 105–116. <https://doi.org/10.4236/jilsa.2015.74010>
- Budak, İ.** (2021). *Veri ve Metin Madenciliği ile Hava Yolu İşletmelerinin Sosyal Medya Yorum ve Skorlarının Değerlendirilmesi*. (Doktora Tezi). Pamukkale Üniversitesi, Sosyal Bilimler Enstitüsü, Denizli.
- Can, F., Koçberber, S., Balçık, E., Kaynak, C., Öcalan, H. C., & Vursavaş, O. M.** (2008). Information retrieval on Turkish texts. *Journal of the American Society for Information Science and Technology*, 59(3), s. 407–421. <https://doi.org/10.1002/asi.20750>
- Chapman, P., Clinton, J., Kerber, R., Khabaza, T., Reinartz, T., Shearer, C., & Wirth, R.** (2000). *CRISP-DM 1.0: Step-by-step data mining guide*. SPSS Inc.
- Chung, J., Gulcehre, C., Cho, K., & Bengio, Y.** (2014). Empirical evaluation of gated recurrent neural networks on sequence modeling. [Erişim: 17.08.2025, *arXiv preprint arXiv:1412.3555*.]
- Cortes, C., & Vapnik, V.** (1995). Support-vector networks. *Machine Learning*, 20(3), s. 273–297. <https://doi.org/10.1007/BF00994018>
- Dasarathy, B. V.** (1991). *Nearest neighbor (NN) norms: NN pattern classification techniques*. IEEE Computer Society Press.
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K.** (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of NAACL-HLT 2019*, s. 4171–4186. <https://doi.org/10.18653/V1/N19-1423>
- Dey, L., Chakraborty, S., Biswas, A., Bose, B., & Tiwari, S.** (2016). Sentiment analysis of review datasets using Naive Bayes and k-NN classifier. *International Journal of Information Engineering and Electronic Business (IJIEEB)*, 8(4), s. 54–62. <https://doi.org/10.5815/ijieeb.2016.04.07>
- Dhiyaulhaq, M. D., & Gunawan P. H.** (2023) Sentiment Analysis of the Jakarta - Bandung Fast Train Project Using the SVM Method. *Jurnal Media Informatika Budidarma*, 7(4). <https://doi.org/10.30865/mib.v7i4.6855>
- Duda, R.O., Hart, P.E. & Stork, D.G.** (2001) *Pattern Classification*. 2nd Edition, John Wiley ve Sons Ltd., New York, s. 202-220.

Eryiğit, G. (2014). ITU Turkish NLP Web Service. In *Proceedings of the Demonstrations at the 14th Conference of the European Chapter of the Association for Computational Linguistics (EACL 2014)* (s. 1–4). Association for Computational Linguistics.

Euclid. (1956). *The thirteen books of Euclid's Elements* (Çev.) T. L. Heath, Dover Publications. (Yaklaşık MÖ 300'de yayınlanan özgün eser).

Farhadloo, M., & Rolland, E. (2016). Fundamentals of sentiment analysis and its applications. In W. Pedrycz ve S.-M. Chen (Eds.), *Sentiment Analysis and Ontology Engineering* (Studies in Computational Intelligence, s. 1–24). Springer. https://doi.org/10.1007/978-3-319-30319-2_1

Feldman, R., & Sanger, J. (2006). *The Text Mining Handbook: Advanced approaches in analyzing unstructured data.* Cambridge University Press. <https://doi.org/10.1017/CBO9780511546914>

Fix, E., & Hodges, J. L. (1989). Discriminatory analysis, nonparametric discrimination: Consistency properties. *International Statistical Review*, 57(3) , s.238-247. [Erişim : 13.09.2025, <https://2024.sci-hub.se/2880/f080d8a194db410a66e8b38e9c740221/fix1989.pdf>]

George, R., & Srividhya, S. R. (2022). Performance evaluation of sentiment analysis on balanced and imbalanced dataset using ensemble approach. *Indian Journal of Science and Technology*, 15(47), s. 2476–2485. <https://doi.org/10.17485/ijst/v15i17.2339>

Go, A., Bhayani, R., & Huang, L. (2009). Twitter Sentiment Classification Using Distant Supervision. *CS224N Project Report, Stanford, I(12)*, 2009. [Erişim : 13.09.2025, <https://www-cs.stanford.edu/people/alecmgo/papers/TwitterDistantSupervision09.pdf>]

Göker, H., & Tekedere, H. (2017). FATİH Projesine Yönelik Görüşlerin Metin Madenciliği Yöntemleri ile Otomatik Değerlendirilmesi. *Bilişim Teknolojileri Dergisi*, 10 (3), s. 291.

Hall, M., Frank, E., Holmes, G., Pfahringer, B., Reutemann, P., & Witten, I. H. (2009). The WEKA data mining software: An update. *ACM SIGKDD Explorations Newsletter*, 11(1), s. 10–18. <https://doi.org/10.1145/1656274.1656278>

Han, J., Kamber, M., & Pei, J. (2012). *Data mining: Concepts and techniques* (3rd ed.). Morgan Kaufmann. [Erişim: 13.09.2025, <https://myweb.sabanciuniv.edu/rdehkharghani/files/2016/02/The-Morgan-Kaufmann-Series-in-Data-Management-Systems-Jiawei-Han-Micheline-Kamber-Jian-Pei-Data-Mining.-Concepts-and-Techniques-3rd-Edition-Morgan-Kaufmann-2011.pdf>]

- Hochreiter, S., & Schmidhuber, J.** (1997). Long short-term memory. *Neural Computation*, 9(8), s. 1735–1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
- Hotho, A., Nürnberger, A., & Paaß, G.** (2005). A brief survey of text mining. *LDV Forum*, 20(1), s. 19–62. <https://doi.org/10.21248/jlcl.20.2005.68>
- Işık, N.** (2019). *Metin Madenciliği Yöntemleri ile E-Ticaret Markalarına Yönelik Sosyal Medya Yorumlarının Analizi*. (Yüksek Lisans Tezi), Marmara Üniversitesi, Sosyal Bilimler Enstitüsü, İstanbul.
- Karamanlı, E.** (2019). *Makine Öğrenmesi Algoritmaları Kullanarak, Metin Madenciliği ve Duygu Analizi ile Müşteri Deneyiminin Geliştirilmesi*. (Yüksek Lisans Tezi), Sosyal Bilimler Enstitüsü, İstanbul Üniversitesi, İstanbul.
- Kaur, G. & Malik, K.** (2021) A Sentiment Analysis of Airline System using Machine Learning Algorithms, *International Journal of Advanced Research in Engineering and Technology*, 12(1), s. 731-742. DOI:[10.34218/IJARET.12.1.2021.066](https://doi.org/10.34218/IJARET.12.1.2021.066)
- Kılıç, S.** (2015). Kappa testi. *Journal of Mood Disorders*, 5(3), s. 142–144.
- Kim, Y.** (2014). Convolutional neural networks for sentence classification. *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, s. 1746–1751. <https://doi.org/10.3115/v1/D14-1181>
- Landis, J. R., & Koch, G. G.** (1977). The measurement of observer agreement for categorical data. *Biometrics*, 33(1), s. 159–174. <https://doi.org/10.2307/2529310>
- Liu, B.** (2012). *Sentiment analysis and opinion mining*. Morgan ve Claypool. [Erişim : 13.09.2025, <https://www.cs.uic.edu/~liub/FBS/SentimentAnalysis-and-OpinionMining.pdf>]
- Mariscal, G., Marbán, Ó., & Fernández, C.** (2010). A survey of data mining and knowledge discovery process models and methodologies. *The Knowledge Engineering Review*, 25(2), s. 137–166. <https://doi.org/10.1017/S0269888910000032>
- Marttin, P. M.** (2025). *KADES Uygulaması Hakkında Kullanıcı Yorumları Üzerinden Web Madenciliği ve Duygu Analizi* (Yüksek lisans tezi). Bilecik Şeyh Edebali Üniversitesi, Lisansüstü Eğitim Enstitüsü, Bilecik.
- Mesri, A.** (2017). *Bir Banka Yazılımı Hakkında Kullanıcı Yorumları Üzerinden Web Madenciliği ve Duygu Analizi* (Tezsiz yüksek lisans dönemi projesi). Hacettepe Üniversitesi, Ankara.

Miner, G. D., Elder, J., Fast, A., Hill, T., Nisbet, R., & Delen, D. (2012). *Practical text mining and statistical analysis for non-structured text data applications*. Academic Press. DOI:[10.1016/C2010-0-66188-8](https://doi.org/10.1016/C2010-0-66188-8)

Munawaroh, K., & Alamsyah, A. (2022). Performance comparison of SVM, Naïve Bayes, and K-NN algorithms for sentiment analysis of public opinion on COVID-19 vaccination on Twitter. *Journal of Advances in Information Systems and Technology*, 4(2), s. 113–125. DOI:[10.15294/jaist.v4i2.59493](https://doi.org/10.15294/jaist.v4i2.59493)

O'Connor, J. J., & Robertson, E. F. (1999). Euclid of Alexandria. *MacTutor History of Mathematics Archive*. University of St Andrews. [Erişim: 17.05.2025, <https://mathshistory.st-andrews.ac.uk/Biographies/Euclid>]

O'Connor, J. J., & Robertson, E. F. (2002). Pafnuty Lvovich Chebyshev. *MacTutor History of Mathematics Archive*. University of St Andrews. [Erişim: 17.05.2025, <https://mathshistory.st-andrews.ac.uk/Biographies/Chebyshev>]

Palen, L., & Liu, S. B. (2007). Citizen communications in crisis: Anticipating a future of ICT-supported public participation. *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, 727–736. <https://doi.org/10.1145/1240624.1240736>

Pang, B., & Lee, L. (2008). Opinion mining and sentiment analysis. *Foundations and Trends® in Information Retrieval*, 2(1–2), s. 1–135. <https://doi.org/10.1561/15000000011>

Paul, J. D., Bee, E., & Budimir, M. (2021). Mobile phone technologies for disaster risk reduction. *Climate Risk Management*, 32, 100296. <https://doi.org/10.1016/j.crm.2021.100296>

Platt, J. (1998). Sequential minimal optimization: A fast algorithm for training support vector machines (MSR-TR-98-14). Microsoft Research. [Erişim: 17.05.2025, <https://www.microsoft.com/en-us/research/wp-content/uploads/2016/02/tr-98-14.pdf>]

Poyraz, O. (2012). *Tıp'da Veri Madenciliği Uygulamaları: Meme Kanseri Veri Seti Analizi* (Yüksek lisans tezi). Trakya Üniversitesi, Fen Bilimleri Enstitüsü, Edirne.

Powers, D. M. W. (2011). Evaluation: From precision, recall and F-measure to ROC, informedness, markedness and correlation. *Journal of Machine Learning Technologies*, 2(1), s. 37–63. [Erişim: 13.09.2025, https://bioinfopublication.org/files/articles/2_1_1_JMLT.pdf]

- Rahadian, I. M., Slamet, C., Andrian, R., Aulawi, H., & Ramdhani, M. A.** (2018). Early warning system in mobile-based impacted areas. *International Journal of Engineering & Technology (UEA)*, 7(3.4), s. 118-121. [Eriřim: 13.09.2025, <https://digilib.uinsgd.ac.id/12268>]
- Raschka, S.** (2014). Naive Bayes and text classification I – Introduction and theory. [Eriřim: 17.05.2025, *arXiv preprint arXiv:1410.5329*.]
- Rogers, R. W.** (1975). A protection motivation theory of fear appeals and attitude change. *The Journal of Psychology*, 91(1), 93–114. <https://doi.org/10.1080/00223980.1975.9915803>
- Sezer, E.** (2018). *Sınıflandırma Sorunu için En Uygun Alt Deęişken Alt Kümesi Seçimi Üzerine Bir Uygulama*. (Yüksek Lisans Tezi). Marmara Üniversitesi, Sosyal Bilimler Enstitüsü, İstanbul. s. 46-47.
- Şeker, S. E.** (2015). Metin Madencilięi (Text Mining). *YBS ansiklopedi*, 2(3), s. 30-32.
- Taboada, M., Brooke, J., Tofiloski, M., Voll, K., & Stede, M.** (2011). Lexicon-based methods for sentiment analysis. *Computational Linguistics*, 37(2), s. 267–307. https://doi.org/10.1162/COLI_a_00049
- Tan, M. L., Prasanna, R., Stock, K., Hudson-Doyle, E. E., Leonard, G. S., & Johnston, D. M.** (2017). Mobile applications in crisis informatics literature: A systematic review. *International Journal of Disaster Risk Reduction*, 24, 297–311. <https://doi.org/10.1016/j.ijdr.2017.06.009>
- Tuzcu, S.** (2020). Çevrimiçi Kullanıcı Yorumlarının Duygu Analizi ile Sınıflandırılması. *ESTUDAM Biliřim Dergisi*, 1(2), s. 1-5.
- Türkoęlu, N.** (2001). Türkiye'nin Yüzölçümü ve Nüfusunun Deprem Bölgelerine Daęılıřı, *Ankara Üniversitesi Türkiye Coęrafiyası Arařtırma ve Uygulama Merkezi Dergisi*, 8, s.133-148. [Eriřim 13.09.2025, https://tucaum.ankara.edu.tr/wp-content/uploads/sites/280/2015/08/tucaum8_7.pdf]
- USGS.** (2021). *Earthquake hazards program*. United States Geological Survey Resmi Sitesi. [Eriřim 13.09.2025, <https://www.usgs.gov/programs/earthquake-hazards>]
- Wirth, R., & Hipp, J.** (2000). CRISP-DM: Towards a standard process model for data mining. In *Proceedings of the 4th International Conference on the Practical Applications of Knowledge Discovery and Data Mining (PKDD)*, s. 29–39. [Eriřim 13.09.2025, <https://www.cs.unibo.it/~montesi/CBD/Beatriz/10.1.1.198.5133.pdf>]

Witten, I. H., Frank, E., Hall, M. A., & Pal, C. J. (2016). *Data mining: Practical machine learning tools and techniques* (4th ed.). Morgan Kaufmann. <https://doi.org/10.1016/C2015-0-02071-8>

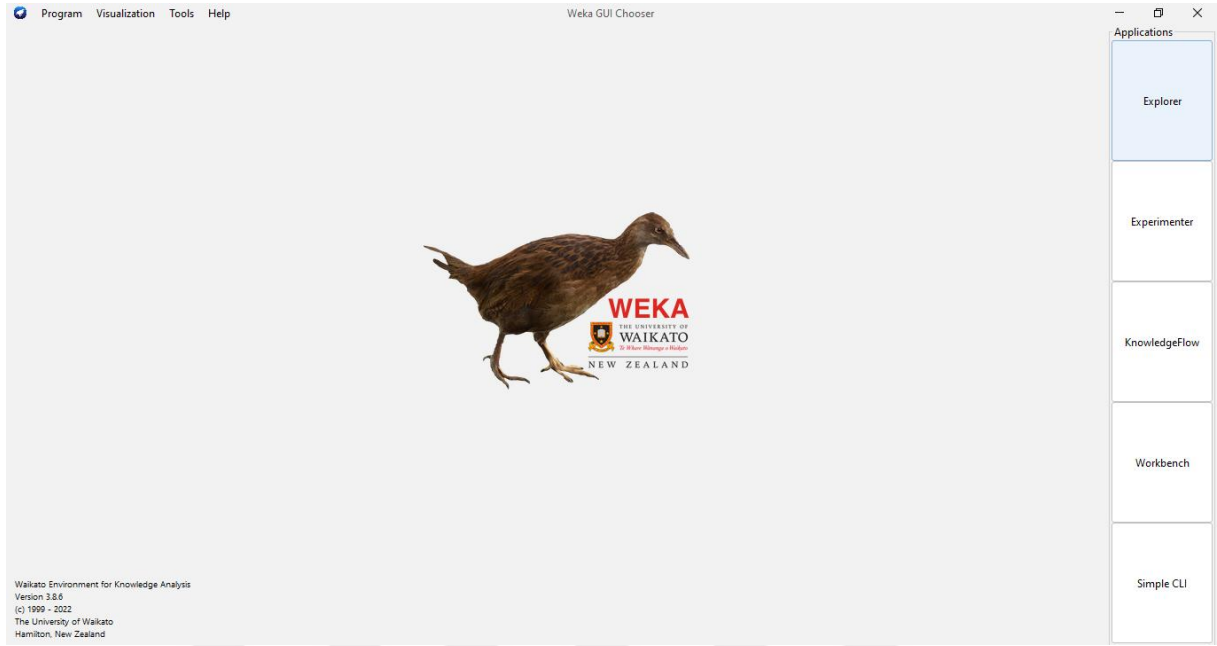
Young, T., Hazarika, D., Poria, S., & Cambria, E. (2018). Recent trends in deep learning based natural language processing. *IEEE Computational Intelligence Magazine*, 13(3), s. 55–75. <https://doi.org/10.1109/MCI.2018.2840738>



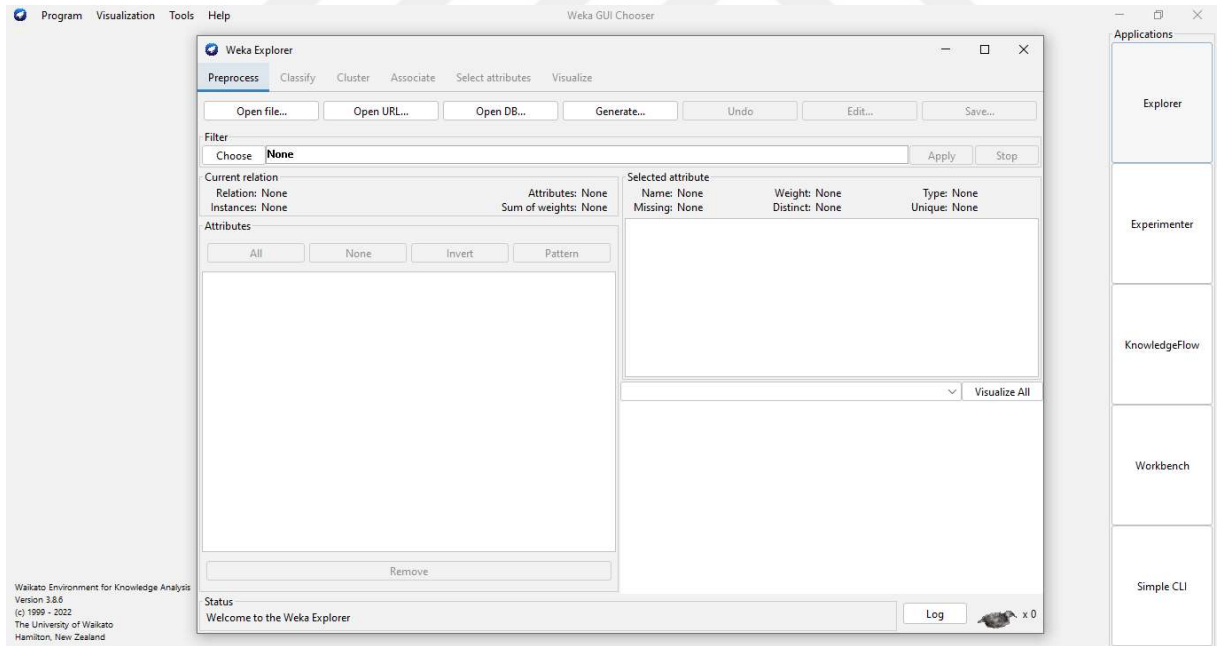


EKLER

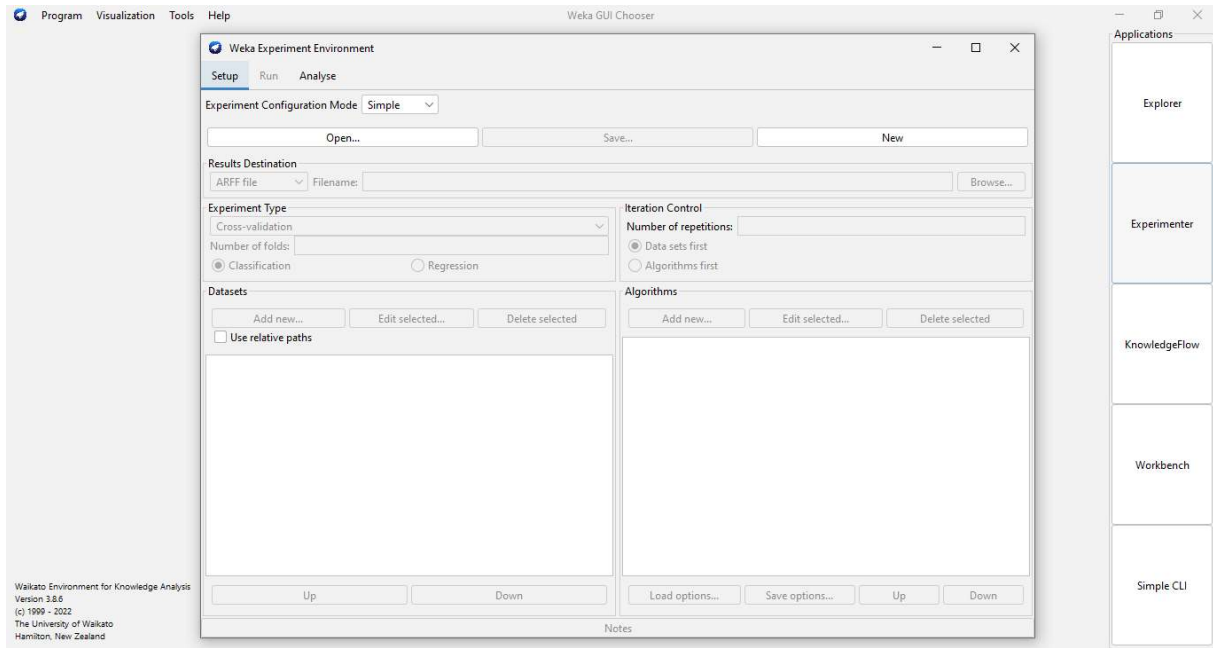
EK-1: WEKA UYGULAMASI



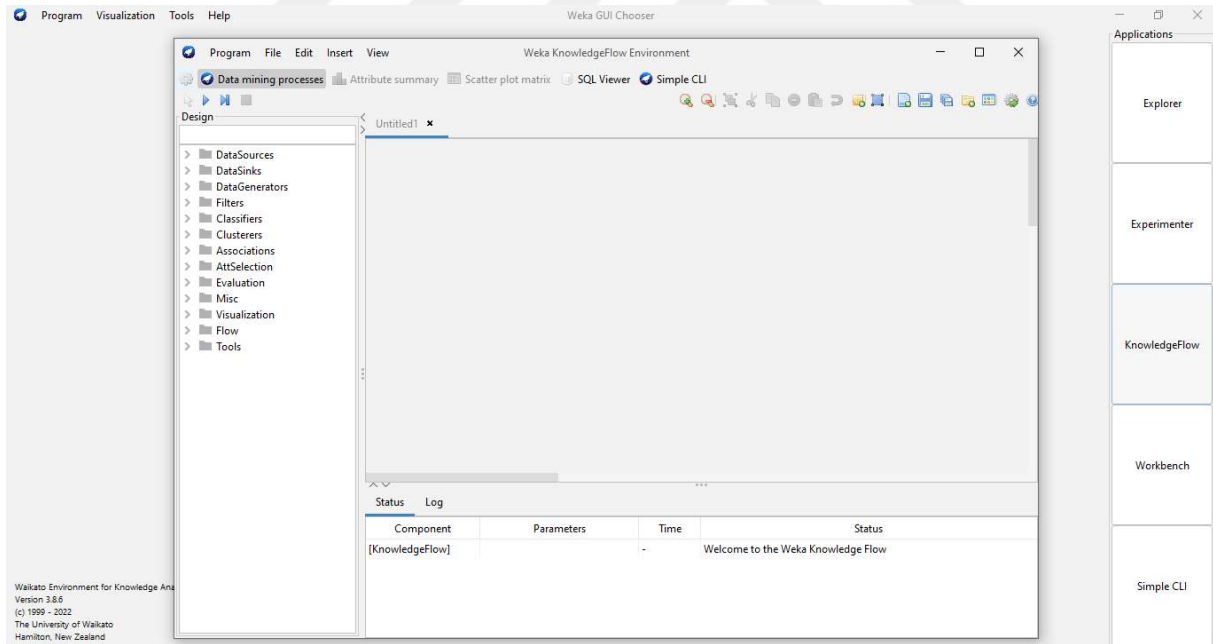
EK-2: WEKA EXPLORER ARAYÜZÜ



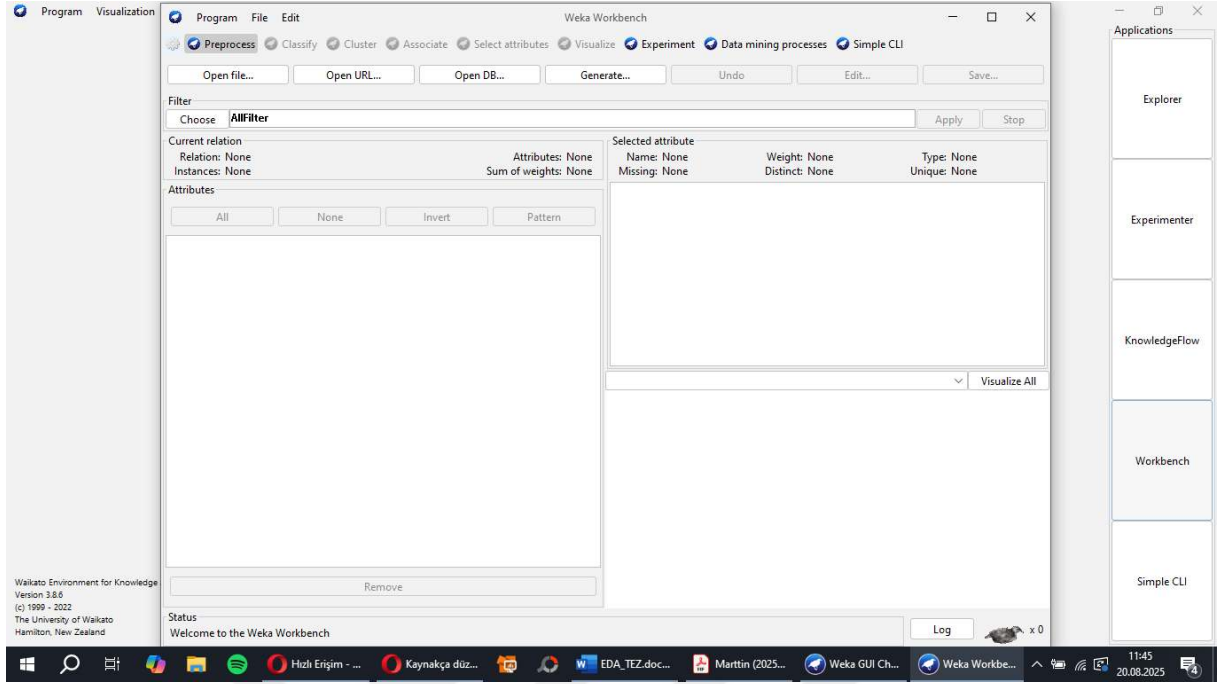
EK-3: WEKA EXPERIMENTER ARAYÜZÜ



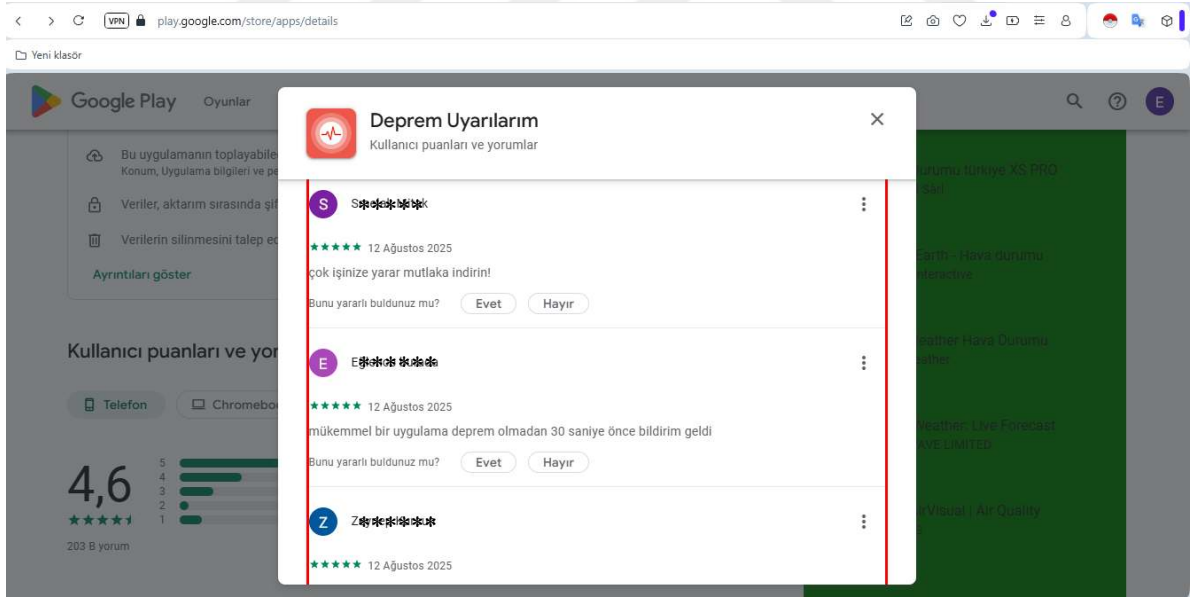
EK-4: WEKA KNOWLEDGE ARAYÜZÜ



EK-5: WEKA WORKBENCH ARAYÜZÜ



EK-6: DEPREM UYGULAMARINA YAPILAN YORUMLAR ÖRNEĞİ



EK-7: INSTANT DATA SCRAPER ARACI İLE VERİLERİN TOPLANMASI ÖRNEĞİ

Instant Data Scraper - Opera

Try another table

Start crawling

☒ Infinite scroll

Min delay: 1 sec

Max delay: 20 sec

Download data or locate "Next" to crawl multiple pages

Kullanıcı	Tarih	Yorum
...	12 Ağustos 2025	çok güzel
...	12 Ağustos 2025	çok işinize yarar mutlaka indirin!
...	12 Ağustos 2025	mükemmel bir uygulama deprem olmadan 30 s
...	12 Ağustos 2025	iyi
...	11 Ağustos 2025	sadece deprem bildirimleri anlık değil geç geliyo
...	11 Ağustos 2025	ayıp olmasın diye atıyor bildirim
...	11 Ağustos 2025	Telefon ayarları açık olmasına rağmen bildirim ç
...	11 Ağustos 2025	reklam var
...	11 Ağustos 2025	olduktan sonra veriyor
...	11 Ağustos 2025	Uygulama çok kötü sürekli bildirim geliyor ama s
...	11 Ağustos 2025	Harika bir uygulama herkes kullanması gerekir
...	11 Ağustos 2025	Kaç saniye önce uyarıyor? Umarım dahada iyi v
...	11 Ağustos 2025	Detaylı gösteriyor nerde kaç büyüklüğünde oldu
...	11 Ağustos 2025	Deprem bekleniyor diye on saniye öncesi alarm
...	11 Ağustos 2025	bildirimleri çok geç geliyor
...	10 Ağustos 2025	deprem uygulamasında bile pro sürüm var arka
...	10 Ağustos 2025	Anında bilgilendirme yapılıyor. Çok güzel bir uyg

Pages scraped: 1
Rows collected: 160
Rows from last page: 160
Working time: 0s

EK-8: VERİLERİN ETİKETLENMESİ ÖRNEĞİ

Yorumlar_Etiketli.xlsx - Excel

Ara

Oturum açın

Paylaş

Dosya Giriş Ekle Sayfa Düzeni Formüller Veri Gözetim Görünüm Yardım

Yapıştır

Pano

Yazı Tipi

Hizalama

Sayı

Genel

Koşullu Biçimlendirme

Tablo Olarak Biçimlendir

Hücre Stilleri

Ekle

Sil

Biçim

Hücreler

Sırala ve Filtre Uygula

Bul ve Seç

Ekleniler

Ekleniler

	A	B	C	D	E	F	G	H	I
8385	GAYET GÜZEL VE ÖZELLİKLİ.	Olumlu							
8386	Güncellenmesi gereken bir uygulama. Android'in yeni sürümlerinde çok ağır çalışıyor ve kasiyor. Daha hafif ve sade bir uygulama	Nötr							
8387	Çok iyi hızlı güvenilir	Olumlu							
8388	Sade, hızlı, Başarılı.	Olumlu							
8389	Harika	Olumlu							
8390	Haberler den önce ,bizlere haberdar etmesi.	Nötr							
8391	Konuya ilgisi olan herkese tavsiye ederim . Mükemmel bir uygulama.	Olumlu							
8392	Zamanında Doğru veriler	Olumlu							
8393	Antalyada deprem oldu fakat uyarı vermedi ? Antalya-kemer 01.03.2022 saat 19.43 4.6 şiddetinde ??	Olumsuz							
8394	Teşekkürler	Olumlu							
8395	Super	Olumlu							
8396	Güzel.	Olumlu							
8397	Çok hızlı ve güvenilir.	Olumlu							
8398	Bildirim gelmiyor	Olumsuz							
8399	Deneyim yaşamadım İnşallah yaşamayız destek amacıyla 5 yıldız verdim	Nötr							
8400	Kesinlikle öneriyorum yıllardır kullandığım harika bir deprem uygulaması	Olumlu							
8401	Harika bir uygulama teşekkürler böyle bir uygulama yaptığınız için	Olumlu							
8402	Super	Nötr							
8403	başarılı	Olumlu							
8404	Widget olsa daha harika olacak. 🤖	Nötr							
8405	Cok iyi program	Olumlu							
8406									

EK-9: VERİYE AİT .csv FORMATININ .arff UZANTILI FORMATA DÖNÜŞTÜRÜLMESİ ÖRNEĞİ

```
*dengeli_verikume.csv - Not Defteri
Dosya Düzen Biçim Görünüm Yardım
İsine yarayan için güvenilir doğru sonuçları veren bir program;olumlu
anlık alarm oluşu muhtesem;olumlu
çok güzel bir uygulama;olumlu
merhaba! uygulamanız çok güzel lakin ve üzeri depremleri göstermiyorsehir olarak da secenegini
çok güzel olmus;olumlu
annem depremi hissettigini soyledi ardından uygulama uyarı verdi çok kaliteli neredeyse anında
ücretsiz ve bildirimleri iyi;olumlu
çok başarılı bir uygulama gayet memnunum ;olumlu
sade ve anlaşılır şekilde uyarı eklerken bildirmesini istediğiniz bölgeyi ve deprem şiddetini b
çok başarılı;olumlu
bu gün sizinle tanıştım geç kalışım emegi gecenlere tesekkurler sagolun muhtesem bir uygulama;
güzel bir uygulama fakat zaman aralığı çok uzun depremi x birebir haber çok önemli tesekur eder
harika uygulama elinize saglık;olumlu
tebrikler en doğru ve tam zamanında bir uygulama;olumlu
rakamsal degerler çok doğru bana anında bildirim geliyor faydalı;olumlu
kullanışlı stabil çalışıyor rabbim kullandırmayı nasip etmesin bundan sonraki gunlerde yıllarda
çok güzel bir uygulama tebrik ediyorum kutluyorum;olumlu
uyarı sesini , ustü olan depremlere ayarlamama ragmen ve tüm ayarlarım doğru olmasına ragmen sız
bildirimler anında gelmiyor dakika sonra gelen bildirimi napayım;olumlu
çok iyi bir uygulama tavsiye ederim herkese;olumlu
anında neden bildirim gelmiyor;olumlu
çok tesekkur ederim deprem olan yerin bildirimini yapıyor lutfen guncel kalsın hep turkiye için
uygulama iphone de çalışmıyor acilis ekranında takılıp kalıyor;olumlu
anında bildirim geliyor çok iyi rasathaneden daha hızlı çalışıyor su gunlerde kesinlikle indirin
herseyiyle harika;olumlu
çok iyi;olumlu
çok faydalı bir uygulama tavsiye ederim;olumlu
allah razı olsun ya elazigda yasıyorum malatyada deprem oluyor dk sonra bildirim geliyor | çok
super tavsiye ederim;olumlu
< St 28, Stn 92 100% Windows (CRLF) BOM ile UTF-8
```

EK-10: WEKA PROGRAMINDA DENGELİ VERİ KÜMESİNİN DAĞILIMI

Weka Explorer

Preprocess Classify Cluster Associate Select attributes Visualize

Open file... Open URL... Open DB... Generate... Undo Edit... Save...

Filter

Choose StringToWordVector -R first-last -W 1000 -prune-rate -1.0 -N 0 -stemmer weka.core.stemmers.NullStemmer -stopwords-handler weka.core.stopw Apply Stop

Current relation

Relation: dengeli_verikume-weka.filters.unsupervised.attribut... Attributes: 655

Instances: 600 Sum of weights: 600

Attributes

All None Invert Pattern

No.	Name
1	☐ class
2	☐ acaba
3	☐ acmamisınız
4	☐ acmanız
5	☐ acık
6	☐ acık
7	☐ acıklanandan
8	☐ acıl
9	☐ acilen
10	☐ acılmıyor
11	☐ acilis
12	☐ actım
13	☐ actık
14	☐ acıpayamda
15	☐ adres
16	☐ afad
17	☐ afet
18	☐ affet
19	☐ aktif

Remove

Selected attribute

Name: class

Missing: 0 (0%) Distinct: 3 Type: Nominal

Unique: 1 (0%)

No.	Label	Count	Weight
1	class	1	1
2	olumlu	300	300
3	olumsuz	300	300

Class: 0Y (Num) Visualize All

300 300

1

EK-11: WEKA PROGRAMINDA DENGESİZ VERİ KÜMESİNİN DAĞILIMI

