



T.C.

**KIRIKKALE ÜNİVERSİTESİ
FEN BİLİMLER ENSTİTÜSÜ**

**SOSYAL MEDYA MAKİNA ÖĞRENME TEKNİKLERİNİ
KULLANARAK YAZAR KİMLİK DOĞRULAMASI**

SULEYMAN ALTERKAVI

BİLGİSAYAR MÜHENDİSLİĞİ ANABİLİM DALI

DOKTORA TEZİ

DANIŞMAN

Prof. Dr. Hasan ERBAY

KIRIKKALE-2021



T.C.

KIRIKKALE ÜNİVERSİTESİ

FEN BİLİMLER ENSTİTÜSÜ

**SOSYAL MEDYA MAKİNA ÖĞRENME TEKNİKLERİNİ
KULLANARAK YAZAR KİMLİK DOĞRULAMASI**

SULEYMAN ALTERKAVI

BİLGİSAYAR MÜHENDİSLİĞİ ANABİLİM DALI

DOKTORA TEZİ

DANIŞMAN

Prof. Dr. Hasan ERBAY

KIRIKKALE-2021

Aileme

ETİK BEYANI

Kırıkkale Üniversitesi Fen Bilimleri Enstitüsü Tez Yazım Kurallarına uygun olarak hazırladığım bu tez çalışmada;

- Tez içinde sunduğum verileri, bilgileri ve dokümanları akademik ve etik kurallar çerçevesinde elde ettiğimi,
- Tüm bilgi, belge, değerlendirme ve sonuçları bilimsel etik ve ahlak kurallarına uygun olarak sunduğumu,
- Tez çalışmada yararlandığım eserlerin tümüne uygun atıfta bulunarak kaynak gösterdiğimi,
- Kullanılan verilerde herhangi bir değişiklik yapmadığımı,
- Bu tezde sunduğum çalışmanın özgün olduğunu,

bildirir, aksi bir durumda aleyhime doğabilecek tüm hak kayıplarını kabullendiğimi beyan ederim.

Suleyman ALTERKAVI

15/09/2021

ÖZET

SOSYAL MEDYA MAKİNA ÖĞRENME TEKNİKLERİNİ KULLANARAK YAZAR KİMLİK DOĞRULAMASI

ALTERKAVI, Suleyman¹

Kırıkkale Üniversitesi

Fen Bilimleri Enstitüsü

Bilgisayar Mühendisliği Anabilim Dalı, Doktora tezi

Danışman: Prof. Dr. Hasan ERBAY

Ağustos 2021, 71 sayfa

Sosyal medya ağlarının kullanımı giderek yaygınlaşmaktadır, bununla beraber kötü niyetli mesajlar yaymak amacıyla kullanıcı hesaplarının saldırganlar tarafından ele geçirildiği sosyal mühendislik saldırılarının yeni türleri ortaya çıkmaktadır. Bu saldırıların üstesinden gelmek için yazarlık doğrulaması bilinen bir metnin kullanıcıya ait olup olmadığını tespit eden sınıflandırma yöntemine ihtiyaç vardır. Ayrıca, bu doğrulama doğru ve hızlı olmalıdır.

Burada, bir kişi tarafından ele geçirilen sosyal medya hesapları için yeni bir yazarlık doğrulama yöntemi önerilmektedir. Twitter tabanlı bir veri kümesinden önemli metinsel özellikler türetilmiştir. Twitter veri kümesi 280 karakterli ve yazarlık bilgileriyle manuel olarak açıklama eklenmiş 16124 tweet'ten oluşur. Daha sonra XGBoost algoritması veri kümesindeki her metinsel özelliğin önemini vurgulamak için kullanılmıştır. Ayrıca, özellik seçimi için üç MCDM yöntemini (ANP, VIKOR ve ELECTRE) sıralama sonuçları üçgenlenir ve özellik ağırlıklandırma için Sıra Üs ve Sıra Toplamı ağırlık yöntemleri uygulanır.

Birçok sınıflandırıcı ile değerlendirilen azaltılmış veri kümesi ve elde edilen F1-skoru sonucu %94,5'tir.

Anahtar kelimeleri: Yazarlık doğrulama, Doğal Dil İşleme, Makine Öğrenimi, Çok Ölçütlü Karar Verme (MCDM).

¹ Sleman Alterkawi: Yazarın çifte vatandaşlığından dolayı ismi iki farklı şekilde yazılabilir.

ABSTRACT

SOCIAL MEDIA AUTHORSHIP VERIFICATION USING MACHINE LEARNING TECHNIQUES

ALTERKAVI, Suleyman²

Kırıkkale University

Graduate School of Natural and Applied Sciences

Department of Computer Engineering, Ph. D. Thesis

Supervisor: Prof. Dr. Hasan ERBAY

August 2021, 71 sayfa

Social media networks' usage is spreading but accompanied by a new shape of the social engineering attacks in which users' accounts are compromised by attackers to spread malicious messages for different purposes. To overcome these attacks, authorship verification, a classification problem for classifying a text, whether it belongs to a user or not, is needed. Moreover, the verification must be accurate and fast.

Herein, a novel authorship verification framework for hijacked social media accounts, compromised by a human, is proposed. Significant textual-features are derived from a Twitter-based dataset. The dataset is composed of 16124 tweets with 280 characters crawled and manually annotated with the authorship information. XGBoost algorithm is used, then, to highlight the significance of each textual-features in the dataset. Furthermore, the ranking results of three MCDM approaches (ANP, VIKOR and ELECTRE) are triangulated for feature selection, and the Rank Exponent and Rank Sum weight methods are applied for feature weighting. The reduced dataset evaluated with many classifiers, and the achieved result of the F1-score is 94.5%.

Keywords: Authorship verification· Natural Language Processing· Machine Learning· Multi-Criteria Decision Making (MCDM).

² Sleman Alterkawi: Yazarın çifte vatandaşlığından dolayı ismi iki farklı şekilde yazılabilir.

TEŞEKKÜR

İlk olarak, bu tez ve araştırmayı tamamlamak için bana güç ve cesaret veren yüce Allah'ıma binlerce kez şükrediyorum.

Tezimin hazırlanması esnasında hiçbir yardımı esirgemeyen Sayın Prof. Dr. Hasan ERBAY'a, tez çalışmalarım esnasında, bilimsel konularda yardımlarını aldığım Sayın Doç. Dr. Cenker BİÇER'e ve Sayın Doç. Dr. Bülent Gürsel EMİROĞLU'na teşekkür ederim.

Bu doktora tez çalışmamı hayatımdaki her aşamada bana destek veren ve bu günlere gelmemde üzerimde çok emeği olan ve maddi manevi desteğini üzerimden hiç eksik etmeyen rahmetli sevgili babam Mohamad ALTERKAWI'nin aziz hatırasına ithaf ederim.

Maddi manevi desteklerini üzerimden hiç eksik etmeyen her konuda beni destekleyen Sevgili annem Ghada ALABYAD'a, kıymetli kardeşime ve ablama sonsuz teşekkür ederim.

Bu araştırmayı kıymetli eşime, oğlum Muhammet'e, kızım Aline'e ve gelecek kardeşlerine ithaf ediyorum. Umarım bu araştırma en yüksek bilimsel dereceleri almaları için bir motivasyon olur.

Bu günlere gelmemde üzerimde emeği olan bütün Kırıkkale Üniversitesi öğretmenlerime ve hocalarıma teşekkür ederim.

Suleyman ALTERKAWI (Sleman ALTERKAWİ)

İÇİNDEKİLER DİZİNİ

ÖZET	iv
ABSTRACT	v
TEŞEKKÜR	vi
İÇİNDEKİLER DİZİNİ.....	vii
ÇİZELGELER DİZİNİ.....	xi
ŞEKİLLER DİZİNİ	xii
ALGORİTMALAR DİZİNİ	xiii
KISALTMALAR DİZİNİ	xiv
1. GİRİŞ	1
1.1. Temel Bilgi ve Literatür Taraması	3
1.1.1. Ele Geçirme (Hijacking) Türleri.....	3
1.2. Literatür Taraması	4
1.3. Araştırma Fizibilitesi Ve Katkısı.....	7
2. MATERYAL VE YÖNTEM.....	10
2.1. Yazarlık Analizi.....	10
2.1.1. Yazarlık Doğrulaması	10
2.1.2. Yazarlık Atfetme.....	12
2.1.3. Yazarlık Profilleme	12
2.1.4. Yazarlık Kümeleme	12
2.2. Çevrimiçi Sosyal Ağlar (ÇSA'lar)	13
2.2.1. Twitter'da Yazarlık Doğrulamasının Zorluğu	14
2.3. Veri kaynakları	15
2.3.1. Twitter Platformundan Ücretsiz Tweet Alma	15

2.3.1.1. API'yi Geliştirici Tarafından Oluşturma.....	16
2.3.1.2. Mevcut Bir Aracı Veya Hizmetleri Kullanma	18
2.3.2. Veri Kümesi Satın Alma	18
2.3.3. Veri Kümesi Sahibi İle İşbirliği.....	20
2.3.4. Twitter Kazıma	21
2.4. Doğal Dil İşleme.....	22
2.4.1. Belirsizlik	22
2.4.1.1 Sözcük Belirsizliği	22
2.4.1.2 Sözdizimsel Belirsizlik	22
2.4.1.3 Anlamsal Belirsizlik	23
2.4.1.4 Söylem Belirsizliği	23
2.4.1.5 Pragmatik Belirsizlik	23
2.4.2. Doğal Dil İşleme Seviyeleri	23
2.4.2.1 Karakter Seviyesi	23
2.4.2.2 Sözcük Seviyesi.....	23
2.4.2.3 Sözdizimsel Seviyesi	24
2.4.2.4 Anlamsal Seviyesi	24
2.4.2.5 İçeriğe Özel Seviye.....	24
2.4.2.6 Yapısal Seviye.....	24
2.4.2.7 Kendine Özgü Seviye	25
2.4.3. Doğal Dil İşleme Uygulamaları.....	25
2.4.4. Doğal Dil İşleme Teknikleri.....	25
2.5. Makine Öğrenme Türleri	26
2.5.1 Denetimli öğrenme.....	26
2.5.2. Denetimsiz öğrenme	27
2.5.3. Takviye öğrenimi.....	27

2.6. Denetimli Öğrenme Algoritmaları.....	28
2.6.1. Destek Vektör Makinesi (SVM).....	28
2.6.2. Lojistik Regresyon.....	29
2.6.3. Random Forest.....	30
2.6.4. XGBoost.....	30
3. ARAŞTIRMA BULGULARI.....	32
3.1. Veri Toplama.....	35
3.1.1. Yazarlıklarını Doğrulamak İçin Twitter Hesaplarını Belirlenmesi	35
3.1.2. Nesnelerin Seçilmesi.....	36
3.1.3. Tarayıcıyı Geliştirmek	38
3.2. Verileri ön işleme	39
3.3. Özellik Çıkarma.....	41
3.3.1. Sözlüksel Özellik	41
3.3.2. Sözdizimsel Özellik	42
3.3.3. Anlamsal Özellik	42
3.4. Ön İşlemci Makine Öğrenimi Algoritmasını Seçme	45
3.5. Özellik Seçimi	47
3.5.1. Alternatifleri Tanımlama.....	48
3.5.2. Kriterleri Tanımlama.....	48
3.5.3. Kriterlerin Ağırlıklarını Tanımlama	49
3.5.4. Alternatifleri Değerlendirmek İçin Yöntemi Seçme.....	50
3.5.5. Problem İfadesine Karşı Çözümleri Doğrulama.....	54
3.6. Özelliklere Ağırlıklarına Göre Ağırlık Atama	54
4. SONUÇ VE ÖNERİLER.....	56
4.1. Deneyler Ve Tartışma	56
4.2. Conclusion.....	Error! Bookmark not defined.

5. KAYNAKLAR.....	64
6. ÖZGEÇMİŞ.....	71



ÇİZELGELER DİZİNİ

<u>Çizelge</u>	<u>Sayfa</u>
Çizelge 2-1. Netlytic.org müşterilerinin planları (netlytic.org, 2021).....	18
Çizelge 3-1. Gömme kelimesinin örneği	43
Çizelge 3-2. BoW örneği	44
Çizelge 3-3. tf örneği	45
Çizelge 3-4. idf örneği	45
Çizelge 3-5. Performans ölçüm parametreleri	46
Çizelge 3-6. Bazı makine öğrenimi algoritmalarının performansı.....	47
Çizelge 3-7. Çıkarılan metinsel özelliklerin performansı arasındaki karşılaştırma	48
Çizelge 3-8. Saaty ölçeği	49
Çizelge 3-9. Her bir kriterin diğer kriterlere göre önemi.....	49
Çizelge 3-10. Normalleştirilmiş Matris	49
Çizelge 3-11. Ağırlık matrisi.....	50
Çizelge 3-12. ANP yöntemine dayalı metinsel özellikler sıralamasını göstermektedir.	51
Çizelge 3-13. VIKOR yöntemi için sıralama tablosu.....	52
Çizelge 3-14. ELECTRE yöntemi için sıralama tablosu	54
Çizelge 4-1. Her bir algoritma için durma kelimelerini özelliği olmadan ve Her bir algoritma için stop kelimeleri özelliği olmadan ve “Sıra Toplamı Ağırlık” yöntemi kullanılarak ağırlıklandırma sonrası tahmin ölçüleri.	56
Çizelge 4-2. Her bir algoritma için stop kelimeleri özelliği olmadan ve“Sıra Üs ağırlığı” yöntemi kullanılarak ağırlıklandırma sonrası tahmin ölçüleri.....	57

ŞEKİLLER DİZİNİ

ŞEKİL	Sayfa
Şekil 2-1. Başvuru formu oluşturma.....	16
Şekil 2-2. Uygulama detayları.....	17
Şekil 2-3. "Consumer_key" and "Access keys" alma	17
Şekil 2-4. Farklı düzeylerde veri erişim planları (developer.twitter.com: Pricing policy, 2018).....	20
Şekil 2-5. Makine Öğrenimi Türleri (Sultan ve Ark., 2018)	26
Şekil 2-6. SVM algoritmasını kullanarak iki veri sınıfını ayırt etme (Andrew, 2021).	29
Şekil 2-7. Lojistik regresyon konsepti için grafik gösterimi (Andrew, 2021)	29
Şekil 2-8. Makine öğrenimi algoritmaları arasında karşılaştırma (Blondel, 2021)....	31
Şekil 3-1. Modelin ayrıntılı mimarisini	32
Şekil 3-2. Akış şemasını	33
Şekil 3-3. CBOW ve Skip-gram modelleri mimarisi (Mikolov Chen, 2013).....	44
Şekil 4-1. Önerilen modeli uygulamadan önce ve sonra sınıflandırıcıların karşılaştırılması.....	58
Şekil 4-2. Lojistik regresyon sınıflandırıcısının performans değerlendirmesi.....	59
Şekil 4-3. Önerilen modelde lojistik regresyon sınıflandırıcısı için eğitim süresi karşılaştırması.....	60
Şekil 4-4. Önerilen model performansının ilgili çalışmalarla karşılaştırılması	61

ALGORİTMALAR DİZİNİ

<u>Algoritma</u>	<u>Sayfa</u>
Algoritma 3-1. Önerilen modelin temel aşamaları	34
Algoritma 3-2. Veri tarayıcısı için Pseudo-code	39
Algoritma 3-3. ANP için sözde kod	50
Algoritma 3-4. VIKOR için sözde kod.....	52
Algoritma 3-5. ELECTRE için sözde kod	54



KISALTMALAR DİZİNİ

ÇSA (OSN)	Çevrimiçi Sosyal Ağlar (Online Social Network)
ABD (USA)	Amerika Birleşik Devletleri (United States Of America)
AP	Associated Press
AMO (SMO)	Ardışık Minimal Optimizasyon (Sequential Minimal Optimization)
k-NN	k-en Yakın Komşu (k-nearest neighbors)
BoW	Kelime Torbası (Bag of Words)
tf-idf	Term Frequency-Inverse Document Frequency
LSTM	Uzun Kısa Süreli Bellek (Long short-term memory)
EER	Eşit Hata Oranı (Equal Error Rate)
KL-diverjansının (KL-divergence)	Kullback-Leibler Diverjansının (Kullback–Leibler Divergence)
UbCadet	Kullanıcı Davranışı Analizi Tabanlı Ele Geçirilmiş Hesap Algılama (User Behaviour Analytics Based Compromised Account Detection)
ML	Makine Öğrenimi (Machine Learning)
MCDM	Çok Ölçütlü Karar Verme (Multi Criteria Decision Making)
API	Uygulama Programlama Arayüzleri (Application Programming Interface)
NLP	Doğal Dil İşleme (Natural Language Processing)
POS	Konuşma Parçası Etiketleme (Part-Of-Speech)
GDPR	Genel Veri Koruma Yönetmeliği'nin (General Data Protection Regulation)
ANP	Analitik Ağ Süreci (Analytic network process)
VIKOR	VIseKriterijumsa Optimizacija I Kompromisno Resenje

1. GİRİŞ

Çevrimiçi Sosyal Ağlar (ÇSA), dünya genelinde milyonlarca insanı bir araya getiren sanal platformlardır. 2025 yılına kadar 4,41 milyar insanın, yani dünya nüfusunun %53,88'inin ÇSA kullanıcısı olacağı tahmin edilmektedir (Statista, 2020; Worldometers, 2020). ÇSA'lar, hızlı büyüyen, günlük yaşantımızı tanımlayıcı teknolojiler ve önemli bilgi kaynaklarıdır. Örneğin, her gün ortalama 500 milyon tweet atılmaktadır (live stats, 2020). Kullanıcılar toplamda günlük 10,5 milyon dakika Facebook'ta zaman geçirmektedirler (Singh ve Banerjee, 2019).

Yüksek kullanım oranı, farklı kullanıcı türleri ve kullanıcıların öz güveni ÇSA'ları saldırganlara karşı korumasız bırakmaktadır. Bunlara hijack saldırılarının da eşlik etmesi, ÇSA problemlerini araştırmacıların çözmesi gereken önemli bir husus haline getirir (Trang ve Ark., 2015).

Günümüzde sosyal medya hesaplarını ele geçirmek, siber suçlular için önemli bir kazanç kaynağıdır. Saldırganlar, ÇSA hesaplarını ele geçirerek, büyük kullanıcı kitlesine zararlı mesajları veya sahte içerikleri yayabilmektedirler. 2016 yılında Time dergisinin yayınladığı rapora göre, Rus korsanlar 2016 yılı ABD seçimlerine, birçok Twitter hesabını ele geçirip yanlış bilgileri paylaşarak müdahale etmiştir (Calabresi, 2017).

Yüksek profilli kullanıcı hesaplarının sonucu daha yıkıcı olabilir. Örneğin, Associated Press'in (AP) 2013 yılında beyaz sarayda bir bombanın patladığı hakkında yayınladığı hikâye, Standart & Poor'un indeksinin %1'lik bir azalmasına ve geçici olarak 136 milyar ABD doları kaybetmesine sebep olmuştur (Lee, 2013). Güvenlik ihlali mesajın tespit edilip kaldırılmasına kadar, AP mesajı 3.000'den fazla kullanıcı tarafından paylaşılmıştır. Bu tür olaylar bize, ÇSA'da bir hesabın yasal sahibi tarafından yazılmamış mesajların güvenilir yollarla algılanıp engellemesinin çok önemli bir husus olduğunu göstermektedir. Diğer taraftan, stilometri biliminden gelen yazarlık doğrulaması, metindeki özellikleri çıkartıp dikkate alarak yazarın kişiliğini tespit etmektedir (Stamatatos, 2009). Yazar doğrulamasının amacı iki ayrı metnin aynı yazar tarafından yazılıp yazılmadığını belirlemek (Boenninghoff ve Ark., 2019) iken Roche ve Ark. yazarlık doğrulamasını, "iki tweet verildiğinde aynı yazara ait olup olmadığını tahmin etme işlemi" olarak tanımlamaktadırlar (Rocha ve

Ark., 2016). Her insan, yazma veya konuşmada idiolect adı verilen benzersiz bir ize sahip olduğundan (Al-Khatib ve Al-qaoud, 2020), Boenninghoff ve Ark. idiolecti tespit etmek için önerdikleri paradigmlar ya profil tabanlıya karşılık örnek tabanlı ya dışsal karşılık içsel ya da ikiliye karşılık birli sınıflandırmalar bulunmaktadır (Boenninghoff ve Ark., 2019).

Twitter veri kümelerinde yazar doğrulamada yanlış oranların azaltılmasının, ele geçirilen sosyal medya hesaplarının erken tespitini sağladığı ve hasarı azalttığı sadece birkaç çalışmada ele alınmıştır (Suman ve Ark., 2021). Bu çalışma kapsamında, ÇSA'larda bulunan kısa mesajların içindeki metinsel özellikler aracılığıyla kullanıcıların profillerini tespit etmek için bir yöntem bulmanın yanı sıra, araştırmacılar yüksek doğrulukla çalışan, özelliklerin sınıflandırmaya katkısına göre sıralayabilen bir yazarlık doğrulaması modeli önermektedirler.

Çalışmanın ana adımları şöyle özetlenebilir. Twitter tabanlı bir veri kümesi oluşturulup yazarlık bilgileri önerilen modeli geliştirmek ve değerlendirmek için manual olarak etiketlenmiştir. Model, tweet'lerden bazı stilometri özelliklerini çıkartmaktadır.

Tweet'ten çıkarılan özelliklerin yazar doğrulamada katkısını ölçmek için ön işlemci olarak XGBOOST algoritması kullanılmıştır. Sınıflandırma sonuçlarını daha da iyileştirmek için en önemli özellikleri seçen ve sınıflandırıcının performansına olumsuz etki eden özellikleri eleyen Çok Ölçütlü Karar Verme (MCDM) yöntemi kullanılarak özellikler sıralanıp ayıklanmıştır. Geriye kalan özelliklere bir ağırlıklandırma yöntemi vasıtasıyla sırasına göre bir ağırlık atanmıştır. Farklı sınıflandırma algoritmaları test edilmiş ve Logistic Regression algoritmasının en iyi F1-skoruna ve doğruluğuna sahip olduğu tespit edilmiştir. Daha sonra, önerilen modelin bu problem için kullanılan diğer modeller ile doğruluk, duyarlılık, F1-skor ve hassasiyet performans metriklerini kullanarak karşılaştırması yapılmıştır.

Çalışmanın geriye kalan kısmı 3 ana başlıkta toplanmıştır.

1.1. Temel Bilgi ve Literatür Taraması

Tez çalışmasının bu kısmında, Sybil hesaplara, bot tarafından ve insanlar tarafından ele geçirilmiş hesaplara yer verilmiştir.

1.1.1. Ele Geçirme (Hijacking) Türleri

Yukarıda anlatıldığı üzere, ÇSA'ların saldırıları 3 ana gruba ayrılmıştır. Bunlar, Sybil hesaplar, bot tarafından ele geçirilmiş hesaplar ve insan tarafından ele geçirilmiş sosyal medya hesapları. Aşağıda her biri kısaca açıklanmaktadır.

Sybil hesapları: Bunlar, sahte haberleri yaymak, sistemleri tehlikeye atmak veya Sybil saldırıları gerçekleştirmek için ÇSA'lara enjekte edilen sahte hesaplardır. İnsanlar ya da robotlar tarafından kontrol edilirler ve çoğu zaman kötü amaçlar doğrultusunda kullanılırlar. Bu hesaplar aynı mesajları, örneğin reklamlar ya da zararlı yazılımları, sürekli göndererek kullanıcıların çalışmalarını etkilemek ya da bağlantılarını ekleyip dostça davranarak bencil hedefler için o kullanıcıların bilgilerini elde etmeye çalışırlar. Sybil hesaplar normal kullanıcıların hesaplarından farklı olarak anormal davranış göstermektedir. Günümüzde ÇSA'ların popülaritesi artmasıyla birlikte sahte hesaplar artarak Facebook hesaplarında %3'ten %4'e kadar bir orana sahip olmaktadır (Singh ve Banerjee, 2019). Bu tür saldırıların tespit edilmesi diğer iki ÇSA saldırı türleriyle kıyasla daha kolaydır. Sosyal medya ekiplerinin bu tür saldırıları tespit etme çabalarına ek olarak, birçok araştırmacı Sybil hesaplarını sadece tespit etmekle kalmayıp aynı zamanda analiz etmek ve karakterize etmek için de çalışmaktadır (Benevenuto ve Ark., 2010; Gong ve Ark., 2014; Gao ve Ark., 2018; Singh ve Banerjee, 2019).

Bot tarafından ele geçirilmiş hesaplar: Bilgisayar korsanları, kimlik avı ve spam eylemleri için kullanıcının kimlik bilgilerini çalarak gerçek kullanıcıların hesaplarını kontrol ederler (Singh ve Ark., 2018; VanDam ve Ark., 2017). Çoğunlukla bu tür hesaplar, ÇSA platformlarında kötü amaçlı içeriği yaymak için kimlik avı saldırıları kullanarak kullanıcı hesaplarını ele geçirilmektedir (Barbon ve Ark., 2017). Zangerle ve Ark. bu tür saldırıları "Hesabın asıl sahibinin bilgisi olmadan hesabın üçüncü kişi tarafından hacklenip ve ardından tweetleri paylaşmak için direkt mesajları göndermek veya takip etmek ve takipten çıkarmak gibi eylemleri gerçekleştirmek için kullanılması" olarak ele almışlardır (Zangerle ve Specht, 2014). Örneğin, 15

Temmuz 2020'de bilgisayar korsanları, Bitcoin'de bir bağış istek mesajı yaymak ve takipçilere kripto dolandırıcılığı mesajları göndermek için Twitter'da büyük takipçisi olan hesapların kontrolünü ele geçirmişlerdir (Okereafor ve Adelaiye, 2020).

İnsan tarafından ele geçirilmiş sosyal medya hesapları: Bunlar da yasal kullanıcıların hesaplarıdır. Bu tür saldırılarda ana hedef, yüksek profilli çok sayıda takipçisi olan bireysel hesaplardır. Bilgisayar korsanları, ÇSA'nın hesaplarına, kimlik bilgileri yoluyla ya da kullanıcılar tarafından hesap güvenlik özelliklerini kötüye kullanarak erişim elde etmektedirler. İzinsiz giren kişiler, sahte haberler yayınlamak veya kullanıcının ağından özel bilgiler elde etmek için gerçek kullanıcı gibi davranır ve bunun sonucu gerçek kullanıcının itibar ve ekonomik kayba yol açar (Savyan ve Bhanu, 2020).

Örneğin, yukarıda bahsedilen 2013 yılındaki Associated Press Twiter hesapları ele geçirip “Acil haber: beyaz saray'da iki patlama olmuş ve Barack Obama yaralanmış” tweetleyerek finansal piyasada paniğin yayılmasına sebep olmuştur (Trang ve Ark., 2015). Trang ve Ark. (2015) bu tür saldırıları açıklayan birkaç örnekten bahsetmişlerdir (Trang ve Ark. 2015), Bu tür saldırılar Li ve Ark. (2014) ve Suman ve Ark. (2021) tarafından da ele alınmıştır.

Sahte hesapların ve bot tarafından ele geçirilen hesapların gösterdiği anormal davranışlar, araştırmacıların ortaya çıkarması gereken temel problemlerdendir (Viswanath ve Ark., 2014; Okereafor ve Adelaiye, 2020). Twitter, sahte hesapların %77'ini hesabın oluşturmamasından bir gün, geriye kalanın %92'ini ise üç gün sonra yakalayabilmektedir (Trang ve Ark., 2015). Twitter 2018 yılında bu hesaplardan 70 milyon anesini kapatmıştır. İnsanlar tarafından ele geçirilmiş hesapların tespiti zorluk düzeyi yüksek problemlerdendir (Okereafor ve Adelaiye, 2020) çünkü ÇSA'larda kullanıcının yazma stiline göre yazarlık doğrulaması yapılamamaktadır.

1.2. Literatür Taraması

Egele ve Ark. (2013) ele geçirilmiş hesapların gösterdiği tutarsız davranışları yakalamak için profil-tabanlı paradigmaları izleyerek COMPA adlı bir sistem geliştirmiştir. Bu sistem, kullanıcı idiolectini tanımlayan davranışsal özellikleri kümeleyerek ve Ardışık Minimal Optimizasyon (AMO) sınıflandırıcıyı kullanarak

tahmin yapmaktadır. Bu yöntem, Twitter ve Facebook veri setleri üzerine test edildiğinde %3,6 yanlış pozitif, küçük boyutlu hackleme saldırıları üzerine ise düşük bir hassasiyet (precision) oranına sahip olmuştur. Burada dikkat edilmesi gereken husus COMPA sistemi bir zaafiyeti avcılık ataklarında (phishing campaigns) yayınlanan benzer mesajları arayıp onların kümeleri oluşturmak için tasarlanmış olmasından dolayı ele geçirilmiş hesapları tek mesaj ile tespit edememesidir.

Trang ve Ark. (Trang ve Ark., 2015) nesne tabanlı paradigma izleyerek COMPA'yı iyileştirmiştir. Trang ve Ark.'a göre (Trang ve Ark., 2015), COMPA, gerçek dünyadaki uygulamalar için yüksek yanlışlık oranına sahiptir. Bu nedenle, insanlar tarafından ele geçirilen ÇSA hesaplarını tespit etmede COMPA'yı iyileştirmek için, hesapların kümelenmesi adımı, benzer bir davranışı olan, ele geçirilmiş veya geçirilmemiş değerlendirme adımıyla değiştirilmiştir. Yazarlar, tek bir hesap özelliğini kullanarak iyi sonuçlar elde etmenin zor olduğunu belirtmişlerdir. Ayrıca, gelecekteki çalışmalarda, stil ve zaman gibi ek özellikler dahil ederek, hesap için değil, tek mesaj için anormallik puanını dahil edip COMPA'nın metodolojisini değiştirmeyi önermektedirler.

Barbon ve Ark. (Barbon ve Ark., 2017), kullanıcının profilini tanımlayan metinsel özellikleri birleştirerek bir bot tarafından ele geçirilmiş hesapları tespit etmek için bir model önermişlerdir. Önerilen modelde K-en yakın komşu (K-NN) algoritması sınıflandırıcı olarak kullanılmıştır. Modelin 10000 kullanıcı içeren Twitter data kümesi üzerine değerlendirilmiş ve %93 doğruluk (accuracy) elde etmişler. Yalnızca düzenli kullanılmayan ve gönderilerinde kullanıcı ile bahsedilen konular arasında tekdüzelik olmayan Twitter hesaplarında doğruluk azalmaktadır. Gelecekteki çalışmalar için, yazarlar doğruluğu artırma amacıyla dilsel olmayan (non-linguistic) özellikler eklemeyi önermektedirler.

Filip Lagerholm (2017), iyi ve kötü niyetli kullanıcı tweetleri arasındaki benzerliği ölçmeyi önermiştir. Önerilen modelde temel özellik kümesi olarak, N-gram, terim frekansı – ters metin frekansı (term frequency-inverse document frequency (tf-idf)), ve Kelime Torbası (Bag of Words (BoW)) kullanılmıştır. Daha sonra, Uzun Kısa Süreli Bellek (LSTM) Sinir ağı sınıflandırıcı olarak kullanılmaktadır. Deneysel değerlendirme sonucunda farklı konularla sekiz farklı kullanıcının tweet'lerini kullanarak % 93.32'lik doğruluğa ulaşmışlardır. Öte yandan, Barbon ve Ark. (Barbon

ve Ark., 2017), modeli gerek hayatta daha uygulanabilir hale getirmek iin benzer konuları olan bir veri kmesinin kullanılmasını nermiřlerdir.

Brocardo ve Ark. (2017), tarafından yapılan bařka bir arařtırmada, srekli kimlik doęrulaması iin derin inan aęları ve sınıflandırma amacıyla tanıtılan Gauss birimlerinin kombinasyonu olan bir sistem geliřtirilmiřtir. nerilen yaklařım, Enron e-postalarına ve Twitter beslemelerine dayalı olarak 140, 280 ve 500 karakterlik metin bloklarından oluřan kısa mesajlar kullanılarak deęerlendirilmiř ve farklı konfigrasyonlarda % 8,21 ila % 16,73 arasında deęiřen bir Eřit Hata Oranı (EER) elde edilmiřtir.

Yakın zamanlarda, Kaur ve Ark. tarafından verimli bir denetimli (supervised) yazarlık doęrulama modeli nerilmiřtir (Kaur ve Ark., 2018). Bilinen ve bilinmeyen kullanıcılar tarafından yazılmıř metindeki farklılıęı len bir model sunulmuřtur. Bilinen veri rneęinden bazı zellikleri seerek kullanıcının davranıřsal profilini oluřturulmuřtur. Ardından, aynı zellikleri bilinmeyen rneklerden ıkartılmıřtır. Bylece, iki zellik grupları karřılařtırma sonucunda yazarlık doęrulaması yapılmıřtır. Kaur ve Ark. tweetleri, her belgenin zelliklerini S1, S2, S3 ve SR adlı gruplara koyarak drt belgeye ayırmıřlardır. S1, "İlk Giren İlk ıkar " řeklinde gncellenen temel profili temsil eder, S2 eęitim kmesini belirler, S3 ve SR test kmesidir. K-means algoritması sınıflandırma amaları iin kullanılmaktadır. Onların zellik kmesi, BoW, N-gram, folksonomi ve stilometri zellikleri iermektedir. Onların modeli Twitter'ın aık veri kmesi zerine deęerlendirilmiř ve %89.24 sınıflandırma doęruluęu elde edilmiřtir.

Benzer bir model, Seyler ve Ark. (2018) tarafından, aık kaynak Twitter veri kmelerinden ıkarılan istatistiksel lmleri ieren zellik kmesi kullanılarak geliřtirilmiřtir. Modelin, kelimeleri kullanan gerek kullanıcının spam gnderenden farklı bir olasılık daęılımına sahip olduęunu ve iki daęılım arasındaki benzerlięi tespit etmek iin Kullback Leibler diverjansının (KL-diverjansının) kullanıldığını varsayarak, ele geirilmiř hesap tespitine uygulandıęı gsterilmiřtir. Yazarlar θ -kullanıcı ve θ -spam olmak zere iki grup $t?$ ve t^* olarak iki zaman noktası tanımlamıřlardır. Zaman noktaları, spam tarafından hesabı kontrol etmenin bařlangı ve bitiř noktasını temsil etmektedir. Bunların iindeki kelime spam gndericinin stilini temsil ederken bunların dıřındakiler ise gerek kullanıcının stilini temsil

etmektedir. Yazarlar, KL-diverjans puanlarının en büyük, en küçük, varyans ve ortalama değerlerini lojistik regresyon sınıflandırıcısı için özellikler olarak kullanmışlardır ve sınıflandırıcının doğruluğu 600 tweet'te doğrulandı ve % 85'i yapay veriler için doğrulamışlardır.

Son zamanlarda, PV ve Bhanu (Savyan ve Bhanu, 2020), Kullanıcı Davranışı Analizi Tabanlı Ele Geçirilmiş Hesap Algılama (UbCadet) adlı yazar doğrulaması için denetimsiz bir sistem önermişlerdir. Bu sistem, tweet metni, hashtag, zaman ve coğrafi konum arasındaki benzerliği ölçerek ÇSA kullanıcısının anormal davranışını analiz etmektedir. UbCadet sistemi, Yelp ve Twitter veri kümeleri kullanılarak değerlendirilip beş kullanıcıdan oluşan özellik kümesini analiz ederken % 83.1 genel doğruluk elde edilmiştir.

Suman ve Ark. (Suman ve Ark., 2021), tweet çiftinin benzerliğini tespit etmek için semantik temsillerini karşılaştırarak Siamese LSTM sinir ağı mimarisi kullanmıştır. Ağ girişi, tweetlerden çıkarılan metin ve emojiler iken çıktı ise hesaplanan benzerlik sonucudur. Yazarlar, farklı ilgi alanlarına sahip 108 kullanıcı için özel tweet veri kümesini toplamış ve manuel olarak etiketlemişlerdir, sonuçta %61.56 doğruluğu elde etmişlerdir.

1.3. Araştırma Fizibilitesi ve Katkısı

Sosyal ağ sitelerinin çeşitlenerek artmasıyla birlikte ÇSA'lerin önemi hayatımızda gittikçe artmaktadır. Sosyal ağ siteleri, kişiselleştirilmiş web profilleri aracılığıyla bireylerin ağlarını çevrimiçi topluluklara bağlayan web siteleridir (Robbins, 2012). Ira P. Robbins, "Konu çevrimiçi gönderilerine gelince insanlar aptaldır. İnsanlar çevrimiçi olarak sergilenen aptalca davranışlar gösteriyor" demiştir (Robbins, 2012). İnsanlar, kendi konumları, yaşamlarıyla ilgili ayrıntıları, ilişki durumları ve doğum tarihleri gibi bilgi paylaşmaktadırlar. Üstelik, kişiler belirli bir konu hakkında fikir ve düşüncelerini yayınlar, kendileriyle aynı fikirde oldukları diğer konuları başkalarının hesaplarından paylaşır, diğer mesajlara yorum yapar, arkadaşlarını etiketler. Tüm bu bilgiler bilgisayar korsanları tarafından aynı kişisel bilgileri kopyalayan orijinaline benzer sahte hesaplar yapmak için kullanılabilir. Birçok ÇSA'nin, kullanıcıları için gerçek adlar gerektirmeyen MySpace gibi bunu etkinleştirdiği ve her biri bir takma

adla kaydolabilir. Veya gerçek bir hesabı hacklemek ve hesabı mağdurları veya başka amaçlarla kötüye kullanmak için kullanmaktadır.

Birçok ÇSA için hizmet süresi sözleşmesinde kullanıcılarına sahte kimlikler kullanmaması için uyarır ve bazıları kimlik doğrulaması için bir iyileştirme ister. Ancak diğer ele geçirilmiş hesaplar durumunda, kullanıcılar hesabı ele geçirilmiş olarak bildirmesi veya hesabın farkedebilecek anormal davranış gösterene kadar ÇSA'nın yapacağı bir şey bulunmamaktadır.

Bugünlerde ÇSA gönderileri mahkemede birçok dava ve soruşturma gönderi yazarına bir delil olarak kullanılmaktadır. Gönderiler birçok dava ve soruşturma da kullanılmaktadır (Robbins, 2012). Ayrıca ifade özgürlüğünü kısıtlandığı ülkelerde de, Twitter'da tweet attığı için kişiler tutuklanabilmekte ve sosyal medya ağların izlenmesi, terörist ağları tespit etmek için kullanılmaktadır. Ama ya birisi başka bir hesabın güvenliğini ihlal ederse ve onun hesabından gönderi yayınlarsa; bu durumda yazarı doğrulamak için bir sistem gerekmektedir.

Mevcut çalışmaların çoğu Sybil hesapları ve botlar tarafından ele geçirilen sosyal medya hesapların çeşitli makine öğrenimi modelleri kullanarak yazar doğrulaması üzerine odaklanmıştır. Ancak bir insan tarafından ele geçirilen sosyal medya hesapları için yazar doğrulama doğruluğunu iyileştirmek için en önemli stilometrik özelliklerin çıkarılmasına ve seçilmesine çok az vurgu yapılmıştır.

Öte yandan, bu çalışma, herhangi bir kısa metin mesajlaşma servisinde etkili bir şekilde kullanılabileceği, sözcüksel, sözdizimsel ve anlamsal özellikleri birleştirerek bir tweet'in yazarlığını doğrulamak için bir model önerilmektedir. Model, en etkili özellikleri seçmek için özelliklerin azaltılmasını (features reduction) dikkate almakta ve en alakalı özelliklerin çıkarılmasına ve uygun boyut azaltma yönteminin uygulanmasına bağlı olarak Makine Öğrenimi ML modelinin yüksek tahmin doğruluğuna vurgu yapmaktadır. Benzer yaklaşımı Gadekallu ve Ark. Da önermektedir (Gadekallu ve Ark., 2020).

Çalışma kapsamında bir insan tarafından ele geçirilmiş sosyal medya hesaplarının yazar doğrulama doğruluğunu, mesajları sınıflandırmak için ML algoritmalarının izlediği bir boyut azaltma yaklaşımı önererek iyileştirmeye hedeflemektedir. Çalışmanın ana katkıları aşağıda listelenmiştir:

- Önerilen sistemi değerlendirmek için aynı türden veya aynı konuya ait, yazarlık bilgileri ile birlikte Twitter tabanlı veri kümesi oluşturmak ve manuel olarak etiketlemek,
- Bir tweet'in yazarlığını doğrulamak ve herhangi bir kısa mesajlaşma hizmetinde etkili bir şekilde kullanılabilecek sözcüksel, sözdizimsel ve anlamsal özellikleri birleştirmek,
- Geleneksel makine öğrenimi sınıflandırıcılarını, meta öğrenme algoritmaları olarak yürütüp özellik seçim sürecinde ön işlemci olarak kullanmak,
- Aşağıdakileri içeren boyut azaltma yaklaşımının uygulanması: Tweetin yazarlığını doğrulamada her özelliğin katkısını ölçmek için aşılın meta-öğrenme algoritmasının bir ön işlemci olarak kullanılması. Ardından, en az etkili özellik kümesinin göz ardı edildiği MCDM yaklaşımlarını kullanarak bu özellikleri sıralayın. Kalan özellikler, ağırlık yöntemini kullanarak derecelerine göre bir ağırlık atmak.
- En yüksek performansı elde etmek için mesaj sınıflandırması için farklı ML algoritmaları uygulamak. Şekil 1'de gösterilen ve önerilen model, Bölüm 3'te detaylandırılmıştır.

2. MATERYAL VE YÖNTEM

Tarihsel olarak bakıldığında yazarlık analizi teknikleri, Mendenhall'ın yazarın kişiliğini belirtmek için özellikleri sayma fikrini ilk icat ettiği 1887'de ortaya çıkmıştır (Lagerholm, 2017). Bunu Yule (1938) ve Morton'un (1965) yazarın kimliğini yargılamak için cümle uzunluklarını kullanan çalışmaları izlemiştir (Tekkalmaz ve Dedeoğlu, 2006). Daha sonra bu alandaki çalışmalar birçok dilde devam etmiştir.

Yazarlık analizi süreci, metnin kompozisyonu hakkında sonuçlar çıkarmak için metnin özelliklerini incelemeye odaklanır, esasında yazarlık analizi, istatistiksel ölçümleri edebi yönetime uygulayan doğal dil işlemenin bir dalı olan stilometri biliminden gelmektedir. Birçok alanda kullanılan yazarlık analizi, örneğin, anonim veya çelişkili kitaplarda yazarın gerçek yazar olup olmadığının anlaşılması için benzerliğin ifşa edilmesi, aynı zamanda e-posta ve haber gruplarının yazarlarını tespit etmeye yönelik hukuki soruşturmalardır (Stamatatos, 2008). Son yıllarda istihbarat, ceza ve medeni hukuk gibi birçok alanda uygulama alanları artmaktadır. Bu, adli yazarlık olarak anılır ve bu, bir metnin kaynağının bilinen bir yazar kümesinden bir (veya muhtemelen birkaç) yazar tarafından ele alındığı fakat yazarların araştırmacılar tarafından bilinmediği varsayımı yapan problemdir (Rocha ve Ark., 2016).

Yapay zekadaki hızlı gelişmelerle birlikte, bilgisayar inanılmaz sayıda yazarın metinsel özelliklerini işleyebilmekte ve özellikle yazarlık analizini daha zor hale getiren belirli özelliklere sahip ÇSA'larda gerçek yazarı yüksek doğrulukla sınıflandırabilmektedir. Bu bölümde, yazarlık analizi türleri, ÇSA'lar, doğal dil işleme teknikleri ve makine öğrenimi gibi çalışmadaki ana kavramları tanıttacaktır.

2.1. Yazarlık Analizi

Rao ve Ark. araştırmalarında dört tür yazarlık analizi tanımlamışlardır (Rao ve Ark., 2017). Bunlar:

2.1.1. Yazarlık Doğrulaması

Yazarlık atfetme, bilinen kişilerden oluşan bir liste arasından bir hedef belgenin en olası yazarının belirlenmesi problemidir. Buna karşın yazarlık doğrulaması, hedef belgenin belirli bir kişi tarafından yazılıp yazılmadığının tespiti problemidir. Koppel'e göre doğrulama, atfetmeden önemli ölçüde daha zordur (Koppel ve Schler, 2004) ki bu, bir belgenin veya kitabın yazarını doğrulamak içinde geçerlidir, yani yazarlık niteliklerini bir dizi yazar arasında yapmak kolaydır, ancak yazarlık doğrulaması yazarın yazım tarzındaki değişiklikler nedeniyle zordur. Koppel ve Schler, 10 yazar tarafından yazılmış yirmi bir 19. yüzyıl İngilizce kitabın üzerinde çalışması bu yöndedir (Koppel ve Schler, 2004). Ancak, ikili sınıflandırma problemi olan yazarlık doğrulaması çoklu sınıflandırma problemi olan yazarlık niteliklendirmeden daha kolaydır. Rocha ve ark. kullanıcı başına 500 tweet olan 50 kullanıcıyı analiz etmek için 3.30GHz ve 96GB RAM belleğe sahip 5820K CPU bilgisayar kullanmışlardır (Rocha ve Ark., 2016).

Rocha ve ark., yazarlık doğrulamasını “İki tweet verildiğinde, aynı yazara ait olup olmadıklarına dair bir tahmin yapma” problemi olarak tanımlarken (Rocha ve Ark., 2016), Brocardo ve ark. ise verilen bir hedef belgenin belirli bir kişi tarafından yazılıp yazılmadığını tespit etme problemi olarak tanımlamıştır (Brocardo ve Ark., 2017).

Koppel ve Schler yazarlık doğrulamasını “verilen tek yazarlı yazı örneklerinden, verilen metinlerin bu yazar tarafından yazılıp yazılmadığının belirlenmesi” problemi olarak tanımlar. Bir kategorileştirme problemi olarak yazarlık doğrulamasını tanımlamanın daha iyi bir yolu, verilen iki örnek kümesinin tek bir üretici süreç (yazar) tarafından mı yoksa iki farklı süreç tarafından mı üretildiğinin sorulmasıdır (Koppel ve Schler, 2004).

Rao ve Ark. yazarlık doğrulaması, “Tek bir yazara ait bir dizi belge ve sorgulamanın bir belge göz önüne alındığında, sorgulanan belgenin söz konusu yazar tarafından yazılıp yazılmadığını belirleme” problemidir (Rao ve Ark., 2017).

Buna göre yazarlık doğrulaması, bir metni bir kullanıcıya ait olsun ya da olmasın sınıflandırmak için bir sınıflandırma problemidir. Bu araştırmada, tek bir yazar için verilen ve sorgulanan gönderi için bir dizi ÇSA gönderisi olacaktır. Yazarlık doğrulaması, söz konusu yazar tarafından yazılıp yazılmadığını belirleyecektir.

Araştırmacılar, yazarlık doğrulama problemini, tek bir yazardan gelen tüm metinleri kümülatif olarak ele alan Profil Tabanlı Paradigma gibi çeşitli perspektiflerden ele almışlardır. Her metin örneğini ayrı ayrı ele alan örnek tabanlı paradigma en çok kullanılan paradigmadır (Potha, 2019).

Diğer araştırmacılar, ilkinin bilinen bir tweetin bilinmeyen bir tweet ile benzerliğine dayanarak yazarı tahmin ettiği, birkaç dilde etkili olan içsel ve dışsal yöntemler arasında ayrım yaparken, ikinci yöntem, sorunu örneklediği ikili bir sınıflandırmaya aktarır. Şüpheli yazarları negatif sınıfa dönüştürür ve bilinmeyen metinle bilinen metin ve negatif sınıftan metin arasındaki benzerliğin sonucunu karşılaştırır; bu ilk yöntemden daha doğrudur (Suman ve Ark., 2021).

2.1.2. Yazarlık Atfetme

Brocardo ve Ark. yazarlık atıfetmeyi “bilinen bireyler listesi arasından bir hedef belgenin en olası yazarını belirlemek” olarak tanımlamıştır (Brocardo ve Ark., 2017).

Rao ve Ark. İse yazarlık kimliğini “tartışmasız yazarlık metinlerinin bulunduğu bir dizi aday yazar göz önüne alındığında, amaç doğru yazarı bulma” işlemidir (Rao ve Ark., 2017).

2.1.3. Yazarlık Profilleme

Yazarlık profillemesi veya nitelendirmesi, anonim bir belgede yazarın özelliklerini (örneğin cinsiyet, yaş ve ırk) belirleme ile ilgilidir (Brocardo ve Ark., 2017).

2.1.4. Yazarlık Kümeleme

Rao ve Ark. göre yazar kümeleme problemi, klasik yazarlık atıfetme probleminden daha zordur. Bir belge koleksiyonu verildiğinde, amaç, aynı yazar tarafından yazılmış belgeleri, her bir küme farklı bir yazara karşılık gelecek şekilde gruplamaktır. Belgeleri dahil edilen farklı yazarların sayısı verilmemiştir (Rao ve Ark., 2017).

2.2. Çevrimiçi Sosyal Ağlar (ÇSA'lar)

Çevrimiçi Sosyal Ağlar, bireylerin kendilerini tanımlamalarına ve sosyal ilişkiler kurdukları sosyal ağlara katılmalarına izin veren elektronik ağlar olarak da bilinir (Bai ve Yao, 2010).

Bu ağlar, arkadaşlıklar, ortak çalışma veya bilgi alışverişi gibi belirli ilişkiler içinde birbirleriyle iletişim kuran farklı insan gruplarından oluşur. Bu ağların varlığını koruması, kullanıcılar arasındaki etkileşimin devam etmesine bağlıdır (Chou ve Chou, 2009). Günümüzde ÇSA'lar dünyanın her yerindeki diğer insanlarla tanışmak, bir şeyi tartışmak, haberleri yaymak ve düşünceleri özgürce ve sınırsız olarak paylaşmak için büyük platformlardır.

Sosyal ağlar, belli temel ortak özelliklere sahipken, ağın ve kullanıcıların doğası gereği empoze ettiği özelliklerle birbirinden farklılaşır. Ağlarda bulunan en belirgin ortak özellikler şunlardır:

- Profil sayfaları: Bu sayfalarda kullanıcı, cinsiyet, doğum tarihi, ülke, ilgi alanları ve resimler gibi kişisel bilgilerini ve diğer bilgileri sunar.
- Arkadaşlar veya bağlantılar: Kullanıcılar, belirli bir amaç için birbirine bağlanırken, ÇSA, bağlantıda olan kişiyi bir arkadaş olarak tanıtır.
- Kullanıcı mesajı: Mesaj aracılığıyla kullanıcı belirli bir konu hakkında düşüncelerini yazabilir ve bunları arkadaşlarıyla veya başkalarıyla paylaşabilir. Bu mesaj twitter'da Tweet, Facebook'ta yayın gibi ÇSA'ya göre farklı bir adı vardır.
- Hashtag: Hashtag, belirli bir konu hakkındaki tüm kullanıcıların mesajlarını, diğer kullanıcıların bu konu hakkında konuşulanları kolayca bulmasına yardımcı olan ve fikirlerine katkıda bulunabilecek bir etiket altında birleştirir (Costa ve Ark., 2017). Sosyal medyada “#” karakteri ile kullanılır, konu etiketlenmek istenirse hashtag kelimeler arasında boşluk olmaz, örneğin: #SocialMedia veya alt çizgi arasında ayırım yapmak için kullanılabilir, yani #social_media.
- Bahsetme: Bu, özellikle mesajın potansiyel aday kullanıcılarının dikkati çekmede kullanılır. Sosyal medyada, kullanıcı mesaja dikkatini çekmek için diğer kullanıcılardan bahseder.

En önemli ÇSA'ların bazıları; Facebook, Twitter, YouTube ve MySpace'dir. Bu çalışmada, önemli ÇSA'lardan biri olan, kısa mesajlaşmalara izin veren, tweet olarak anılan mesajların en fazla 140 karakter olduğu ancak Eylül 2017'de bu sayının 280 karaktere genişlettiği (phys.org, 2017) Twitter'da yazarlık doğrulaması ele alınmaktadır. Kullanıcılar diğer kullanıcıların tweet'lerine abone olabilir, bu "takip etme" olarak bilinir ve aboneler "takipçiler" veya "tweetler" olarak bilinir. Bireysel tweet'ler diğer kullanıcılar tarafından "retweet" olarak bilinen bir işlemle kendi özet akışlarına iletilir. Bir retweet düğmesine tıklayarak, Kullanıcılar ayrıca bireysel tweetleri "beğenebilir" (önceden "favori"). “@” karakterini ve ardından kullanıcı adını kullanarak bir kullanıcıdan bahsedebilir veya ona cevap verebilirsiniz (Espinhera ve Albuquerque, 2013).

Yazarlık Doğrulaması yapmak için, metin veya tweet'in kendisiyle ilgili olmayan, ancak platform tarafından sağlanan, her platformun yazara veya gönderiye tanımlaması için belirli özellikler verdiği ve platform politikasına göre bir platformdan diğerine farklı olduğu özelliklerden yararlanılabilir. Twitter'da retweet, bahsetme, duygular, zaman damgası vb. gibi birçok özellik vardır.

2.2.1. Twitter'da Yazarlık Doğrulamasının Zorluğu

Yazarlık doğrulaması, yeni bir bilim değildir, onlarca yıldan beri uygulanmaktadır ve ilk başlangıcı kitapların, şiirlerin yazarlarını veya diğer edebi yenilikleri doğrulamakla başlamıştır. Marcelo ve arkadaşlarına göre, daha önce yapılmış çalışmaların çoğu büyük metin içeriğine sahip olan metin belgelerine odaklanılmıştır (Murphy, 2012). Bu metinler büyük kelime içeriğine ve yazarın stilini tespit etmek için kullanılan birçok stilometrik özelliğin çıkarılmasına yardımcı olan belirli bir yapıya sahipti. Ancak, ÇSA'larda ve özellikle Twitter'da Yazarlık doğrulaması daha zordur çünkü mesajlar sınırlı sayıda karaktere sahip olmanın yanı sıra kullanıcının mesajlarındaki dil gramerlerini takip etmesi gerekmektedir.

ÇSA'larda kısa mesajlarda, yine Rocha ve Ark. belirttiği gibi net bir bağlam yoktur ve yazarların yaptığı bir araştırma, internet kültüründe yaygın olan ve geleneksel olmayan noktalama işaretleri, kısaltmalar ve karakter tabanlı göstergeler kullanılmakta (Rocha ve Ark., 2016) bunun sonucu olarak yazarlık doğrulama problemi daha karmaşık bir hal almaktadır.

ÇSA’larda kullanılan yeni ortak dilin yanı sıra örneğin duygular (😊 ..), kelime ve cümlelerin kısaltması (brb: be right back, thx: teşekkürler vb.) gibi geleneksel kitaplarda ve makalelerde bulunmayan stilometrik özellikleri temsil eden özellikler ÇSA’larda bulunur.

Kısa mesajların yazarlık doğrulamadaki diğer zorluğu ise ÇSA’larda çıkarılabilen küçük stilometrik özellikler kümesi olmasıdır. Bu problemlerin üstesinden gelmek için birçok öneri vardır, örneğin mesajları tek bir belge yapmak için birleştirmek (Rocha ve Ark., 2016). Diğer bir yaklaşım ise kısa mesaj yazarlık doğrulamasının sınırlamalarını genişletmek için dilsel olmayan başka özellikler de kullanmaktır.

2.3. Veri kaynakları

Twitter’da yazar doğrulaması için genel bir veri kümesi olmadığından, çalışma amacına ulaşmak için özel veri kümesi oluşturuldu. Twitter platformundan tweet elde etmenin birkaç yöntemi vardır. Bu yöntemler aşağıda verilmiştir.

2.3.1. Twitter Platformundan Ücretsiz Tweet Alma

Twitter, kullanıcıların, araştırmacıların ve geliştiricilerin uygulama programlama arayüzleri (API’ler) aracılığıyla sınırlı miktarda tweet’i ücretsiz olarak almalarına olanak tanır. Kullanıcılar tarafından oluşturulan yazılımın Twitter ile bütünleşmesini sağlayan bir erişim noktası gibi davranan API, ya anahtar kelime tabanlı arama Stream API ya da özel hesap adı tabanlı bir arama REST API ile genel verileri alma olanağı sağlar (Twitter.com, 2021). Stream API, “api.search” uç noktası aracılığıyla, veri akışı belirli anahtar kelimelere, hashtag’lere, kullanıcı konumuna, kullanıcı kimliklerine veya kullanılan dile göre filtreleme özelliği ile ancak gizlilik endişeleri nedeniyle örneklerin nasıl seçildiği ve ilişkisi hakkında bir açıklama yapılmaksızın, her gün yaklaşık 5 milyon tweet ve ilgili tüm meta verilerden oluşan 12 gigabayt ham veri olan tüm tweet’lerin %1’inin alınmasına izin verir (Steinert-Threlkeld, 2018). Bu yöntem, belirli bir olayla ilgili verilerle güncel olmak istediğinizde iyidir.

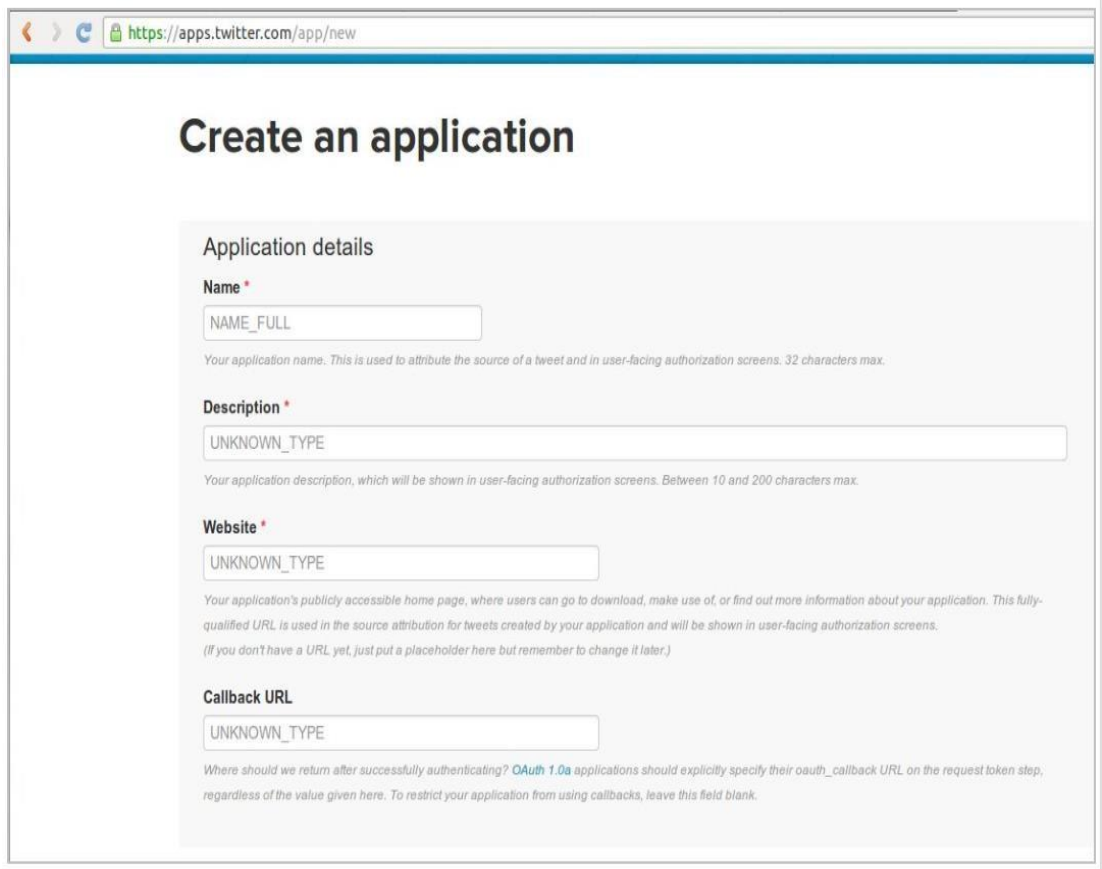
“api.user_timeline” uç noktası aracılığıyla, REST API belirli bir hesap için yaklaşık 3200 son tweet’in alınmasına izin verir, ancak her istekte 15 dakika içinde 180 istek 200 tweet alınabilir, bu nedenle çeşitli meta veri özelliklerine sahip 3200 tweet’i toplamak için 16 istek gerekir (Steinert-Threlkeld, 2018). Bu yöntemle, verileri

almak için önerilen iki yaklaşım vardır: API'yi Geliştirici Tarafından Oluşturma ve Mevcut Bir Aracı Veya Hizmetleri Kullanma.

2.3.1.1. API'yi Geliştirici Tarafından Oluşturma

API'yi oluşturmak için şu anahtarlar edinilmelidir: consumer_key, consumer_secret, access_token, access_token_secret, yeni bir Twitter uygulaması oluşturduktan ve aşağıdaki bağlantıya giriş bilgilerini girdikten sonra elde edilen link: <http://apps.twitter.com> olacaktır.

Mevcut twitter politikasına göre, hesap sahibinin birkaç adımda oluşturulan bir geliştirici hesabına sahip olması gerekir. İlk adım, gerçek bir telefon numarası ve e-posta hesabı tarafından yetkilendirilmiş ayrıntılı kişisel bilgilerin girilmesidir. İkinci adım uygulama oluşturma amacını beyan etme ve üçüncü adım geliştiricilerin koşullarını kabul etmektir, son aşama ise twitter ekibinin doğrulama adımıdır. (Bu aşama birkaç gün sürebilir) (Developer.twitter.com, 2021). İlgili şema aşağıdaki gibidir:



<https://apps.twitter.com/app/new>

Create an application

Application details

Name *

Your application name. This is used to attribute the source of a tweet and in user-facing authorization screens. 32 characters max.

Description *

Your application description, which will be shown in user-facing authorization screens. Between 10 and 200 characters max.

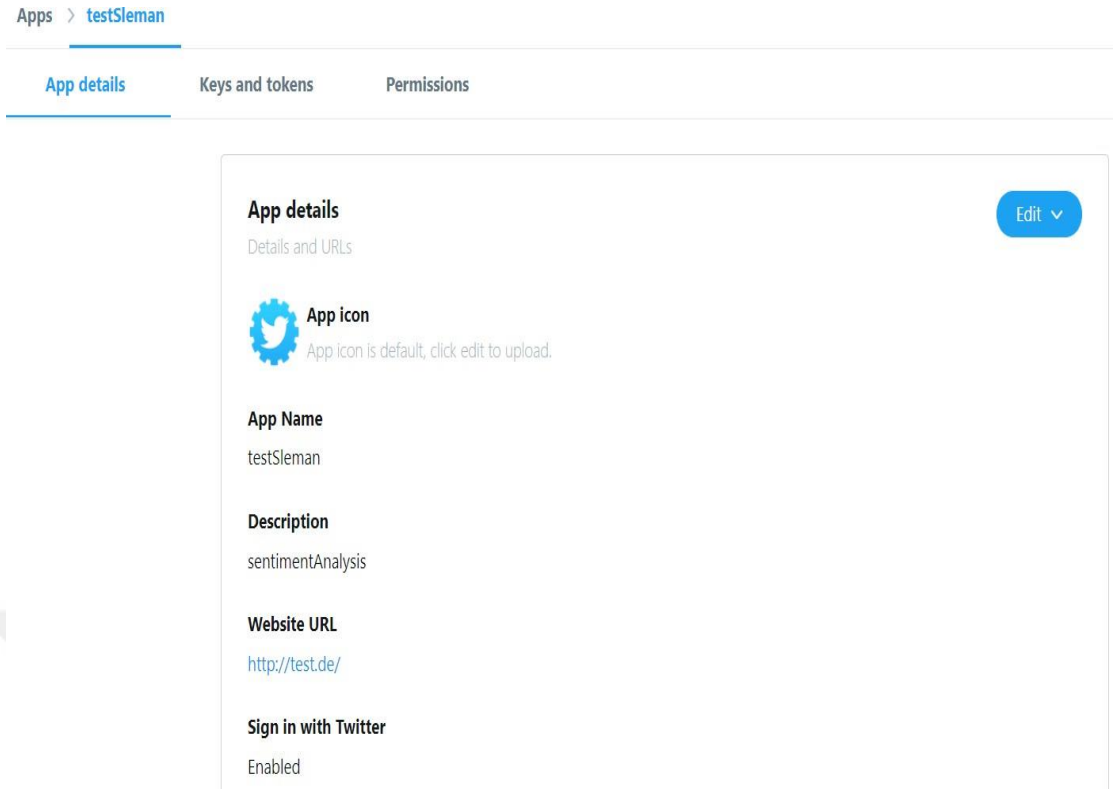
Website *

Your application's publicly accessible home page, where users can go to download, make use of, or find out more information about your application. This fully-qualified URL is used in the source attribution for tweets created by your application and will be shown in user-facing authorization screens. (If you don't have a URL yet, just put a placeholder here but remember to change it later.)

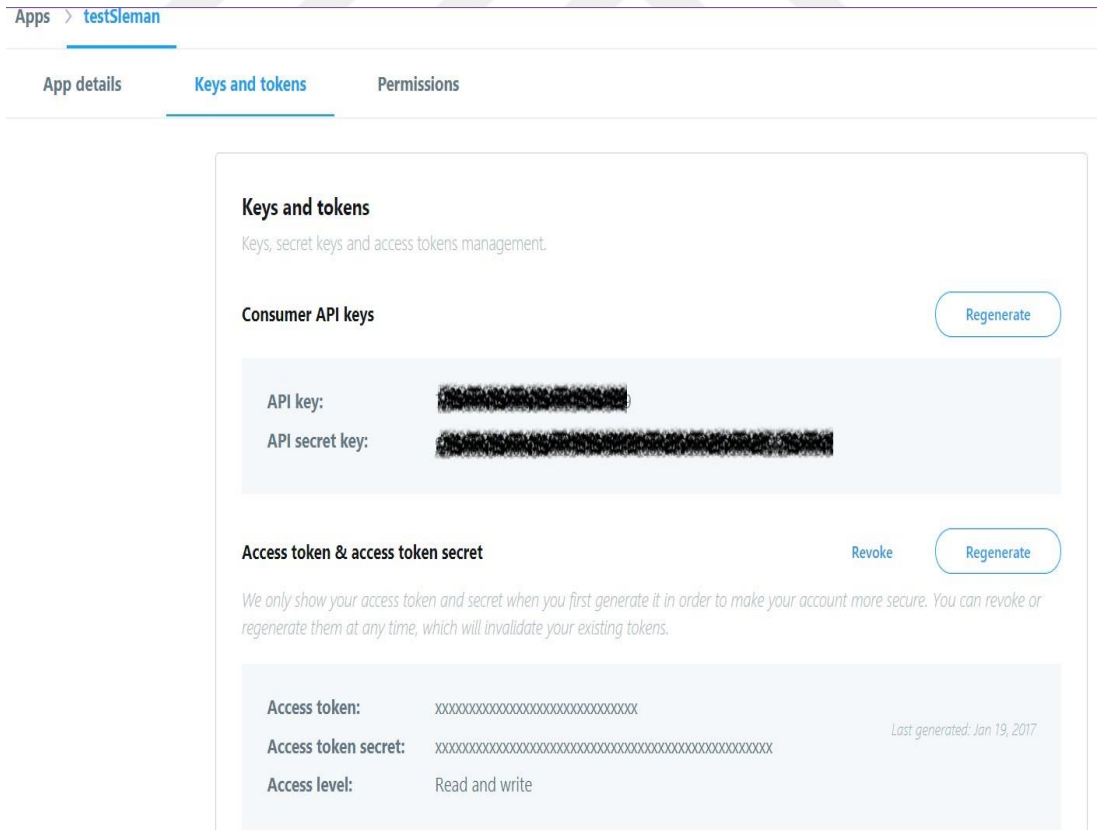
Callback URL

Where should we return after successfully authenticating? OAuth 1.0a applications should explicitly specify their oauth_callback URL on the request token step, regardless of the value given here. To restrict your application from using callbacks, leave this field blank.

Şekil 2-1. Başvuru formu oluşturma



Şekil 2-2. Uygulama detayları



Şekil 2.3. "Consumer_key" and "Access keys" alma

2.3.1.2. Mevcut Bir Aracı Veya Hizmetleri Kullanma

Herhangi bir programlama dilinde kod yazmadan tweet almak için alternatif yaklaşım, hazır araçlar veya servisler kullanmaktır. Kullanıcı adı, şifre ve gerekli tweetler ile ilgili parametreleri sağlayarak; bu araçlara aşağıdaki bazı örneklerden ulaşılabilir:

Netlytic.org: Belirli bir sorgu için tweet'leri alır, toplama işlemi her 15 dakikada bir çalışır ve çalıştırma başına 1000 tweet'e kadar alır, ancak ücretsiz planda 2 haftadan eski tweet'leri sağlar (netlytic.org, 2021).

Web sitesi, ünlü stand-up komedyeni Edward John Izzard'a geri dönen @eddieizzard hesabı aranarak test edilmiştir. 1000 tweet alındı, çoğu hesap tarafından değil hesap hakkında tweetlendi, arama kutusunu "eddieizzard'dan" olacak şekilde özelleştirirken, kullanıcı son iki haftada bir tweet attığından bir tweet aldı.

Web sitesi, daha fazla özellikli veri almak için Prime planlar sunar. Bunlar aşağıdaki tabloda gösterilmektedir.

Çizelge 2.1. Netlytic.org müşterilerinin planları (Netlytic.org, 2021)

	Tier 1 (Free)	Tier 2 (Free)	Tier 3 (Community-Supported)
Max # Of Datasets	3	5	300
Max # Of Records/Dataset	2500	10000	100000

Greptweet.com: web sitesi, aşağıdakiler gibi çeşitli nesnelerle veri sağlıyordu: tweet kimlik numarası, zaman damgası ve metin. Web sitesinin Twitter'daki resmi hesabında bildirildiği üzere, hizmetler Twitter'ın I.B.3 ve I.B.6 (twitter.com/greptweet, 2021) şartlarını ihlal etmesi nedeniyle 2016-08-27'de askıya alındı.

DMI-TCAT: "Digital Methods Initiative Twitter Capture and Analysis Toolset" in kısaltması, Bu, çok fazla programlama deneyimi olmadan Akış API'sini kullanarak tweet'leri almak için bir araçtır (Borra ve Rieder, 2013), maalesef hizmetler artık genel kullanıma açık değildir.

2.3.2. Veri Kümesi Satın Alma

Veri kümeleri için birçok sağlayıcı talebi bulunduğundan veri kümesi ticareti ünlüdür, bu yöntemin örnekleri aşağıdaki gibidir:

Gnip: Paket için 500 \$ 'dan başlayan plan sağlama, kısa bağlantılar için genişleyen bağlantılar gibi özel meta veriler, dili algılamak gibi daha akıllı meta veriler..., Gnip için Twitter alımı, 1 milyon tweet için 1500 \$ 'dan başlayan aylık veya yıllık plan ile kullanıcılara tweetlere sınırsız erişim, ayrıca eğitim kurumları için bazı indirimler sağlanmıştır (Steinert-Threlkeld, 2018).

DataSift: Fiyatlandırma planları geliştirici gereksinimlerine dayalıdır; son zamanlarda Twitter onlarla olan sözleşmeyi iptal etti ve hizmet artık mevcut değil (Steinert-Threlkeld, 2018).

Texifter: Aynı DataSift pazarlama planına sahip, sorguya göre fiyatı hesaplamak için çevrimiçi hesaplayıcı ile web sitesi, geliştiricinin Twitter'ın hizmet süresi nedeniyle bir veya daha fazla sorgu aracılığıyla günde 50000 tweet almasını kısıtladı (Steinert-Threlkeld, 2018).

Twitter: Twitter web sitesine göre, aşağıda açıklandığı gibi, her plan için farklı özelliklerle alınan tweet'ler göre üç plan vardır (developer.twitter.com: Pricing policy, 2018):

- Ücretsiz olan standart API (önceki bölümde açıklanmıştır).
- Uç nokta için ücretsiz ve ücretli özelliklere sahip premium API.
- Uç nokta için tamamen ücretli özellikler olan kurumsal API, tamamen twitter verilerine bağlı olanlar için sağlanmıştır.

Tweets	Standard (free)	Premium	Enterprise
Publish and engage	✓		
Search Tweets: 7-days	✓		
Search Tweets: 30-days		✓	✓
Search Tweets: Full-archive		✓	✓
Filter Tweets	✓		✓
Sample Tweets	✓		✓
Batch Tweets			✓
Direct Messages	✓		
Account and users	✓	✓	✓
Metrics			✓
Ads API	✓		
Publisher tools and SDKs	✓		

Şekil 2-4. Farklı düzeylerde veri erişim planları (developer.twitter.com: Pricing policy, 2018)

2.3.3. Veri Kümesi Sahibi İle İşbirliği

Günümüzde sosyal medyadaki metinleri analiz etmek bir trend; dünya çapında pek çok araştırmacı, geliştirici ve şirket twitter verilerini topluyor, ön işleme tabi tutuyor ve analiz ediyor, bunların çoğu bir araştırmada ortaklık yapmaktan çekinmiyor. Bu yöntem en uygun yöntem olmaya devam etmektedir, ancak hiçbir maliyeti yoktur, zamandan tasarruf sağlar ve araştırmacının çabalarını araştırmının çekirdeğine aktarır (Steinert-Threlkeld, 2018).

Burada, Twitter'ın hizmet şartlarının araştırmacıların günde elli binden fazla tweet paylaşmasına izin vermediğini, ancak kimlik numaralarının paylaşılmasında bir sınır olmadığını belirtmekte fayda var (Steinert-Threlkeld, 2018).

Deneye göre, araştırmacıların çoğu, ülkelerindeki araştırma merkezleri veya hükümet partileri tarafından kendilerine sunulan veriler üzerinde çalışıyor, bu nedenle satır verilerini paylaşma izinleri yok; diğer taraftan bazıları, yalnızca araştırmada kullanılan ve çoğu durumda diğerlerinin araştırmasına uygun olmayan önceden işlenmiş verileri paylaşmayı kabul eder; bu araştırmada diğerleri tarafından sunulan verilerin çoğu, özelliklerin çoğundan yoksundu.

2.3.4. Twitter Kazıma

Web siteleri verileri yapılandırmak için HTML'yi kullandığından, web kazıma verileri almak için bu yapıdan yararlanmaktadır (Cording ve Lyngby, 2011). Kazımanın arkasındaki mantık, siteyi belirli bir hiyerarşiye ayırtmak ve ardından bu hiyerarşideki belirli yoldan verileri almaktır.

Web sitelerinin çoğu, verileri kullanıcılara sunmak için bir API sağlar, ancak API olmadığı durumlarda, gerekli veriler otomatik olarak kazınabilir ve saklanabilir. Twitter'da sağlanan API'ler sınırlıdır ve bazı araştırmaların gereksinimlerini karşılamamaktadır. Bu nedenle, gerekli verilerin kazımanın mevcut çözümlerden biri olabilir. Ancak kazımanın iki dezavantajı var.

- Tweetlerin kazınmasının yasallığı: Twitter, API'lerini ticari amaçlarla kısıtladı, Twitter'ın kullanım şartlarında "NOT, robots.txt dosyası hükümlerine uygun olarak yapılması halinde hizmetlerin taranmasına izin verilir, ancak hizmetler'in Twitter'dan önceden izin alınmadan kazınması kesinlikle yasaktır" şeklinde ifade edilmiştir" (Twitter: Term of Services, 2021). Bu nedenle, tweetleri kazımak yasadışı bir işlemdir.
- Kazıma kullanılan tekniklerden bazıları uygulanabilir değildir, çünkü web sitesine göre dinamik olarak değişen işaretlemeye veya görünüm güncellemesine bağlı olabilir (Cording ve Lyngby, 2011).

Yukarıdaki yöntemleri araştırdıktan ve test ettikten sonra ve araştırmanın özel gereksinimleri nedeniyle, tweetleri almak için Rest API seçildi.

2.4. Doğal Dil İşleme

Doğal dil işleme (Natural Language Processing - NLP), yapılandırılmamış metin analizi için bir dizi tekniktir, metin oluşturmak ve ondan bilgi çıkarmak için kullanılmaktadır (Espinhera ve Albuquerque, 2013).

Doğal diller, Arapça, İngilizce, Fransızca ve diğerleri gibi insan dilleri için kullanılan terimlerdir, bu diller doğal olarak gelişmiş ve yıllar içinde bilinçli bir planlama olmaksızın evrimleşmişlerdir ve dil bilgisinin çıkarılması dilin varlığından sonra gerçekleşmiştir. Öte yandan, programlama dilleri gibi dil bilgisi dilden önce var olan doğal olmayan diller de vardır. Ayrıca, Lojban dili gibi belirsizliğin olmadığı için insan iletişimini geliştirmeyi amaçlayan, tamamen insan iletişimine yönelik kurgulanmış diller de bulunmaktadır.

2.4.1. Belirsizlik

Belirsizlik, dilsel anlamda netlik eksikliği durumudur. Belirsizlik; sözcüksel, sözdizimsel, anlamsal, söyleysel ve pragmatik olabilir (Anjali ve Babu, 2014).

2.4.1.1. Sözcük Belirsizliği

İçinde görüldüğü bağlam veya cümle ne olursa olsun, tek tek kelime düzeyindedir. Örneğin, “bar” kelimesinin birçok anlamı vardır. Mutfağın bir parçası olan bir masa veya başka bir şey anlamına gelebilir. Alkol servisi yaptığınız yer anlamına gelebilir, başka bir bağlamda çikolata anlamına da gelebilir. Ayrıca müzikte (müzik notalarının yazılmasında) zamanı ölçmeye yarayan bir başka anlam da vardır. Bu tür belirsizlikte, belirsiz kelimenin anlamını bulunduğu bağlamdan çıkarmak mümkün olduğundan, kelimenin anlamını bilmek daha kolaydır (Anjali ve Babu, 2014).

2.4.1.2. Sözdizimsel Belirsizlik

Belirli bir cümlenin anlamsal olarak birden fazla olası analizine sahip olması durumudur, örneğin, "The woman saw the man with the Eyeglasses". Anlamsal belirsizlik kadının gözlük takan bir adamı mı yoksa gözlüklerinin arkasından mı gördüğü belirsiz. Buradaki anlam, “with” edatının erkeğe mi yoksa kadına mı bağlı olduğuna bağlıdır (Anjali ve Babu, 2014).

2.4.1.3. Anlamsal Belirsizlik

Anlamsal belirsizlik, cümlenin anlamının tam olarak açık olmadığı halidir. Örneğin, "I met a colleague in the workplace" ifadesi iki anlam taşır. İlk olarak, konuşmacı kendisiyle aynı yerde çalışan biriyle tanışır. Diğer anlam ise tanışmış olduğu meslektaşıyla aynı işyerinde çalışıyor zorunluluğunun olmaması; belki bir sınıf arkadaşısıdır ve onunla iş yerinde karşılaşmıştır. Bu durum atasözlerinde de sıklıkla karşımıza çıkmaktadır (Anjali ve Babu, 2014).

2.4.1.4. Söylem Belirsizliği

Anjali ve Babu (2014), bunu şu şekilde tanımlamıştır: Söylem düzeyinde işlemenin, paylaşılan bir dünyaya veya paylaşılan bilgiye ihtiyaç duyması olarak tanımlamıştır ve yorumlama bu bağlam kullanılarak gerçekleştirilir.

2.4.1.5. Pragmatik Belirsizlik

Pragmatik belirsizlik, cümlenin, onu netleştirecek çok fazla bilgiye sahip olmadığına olur ve bu, konuşmacının niyetine bağlıdır. Örneğin "Ben de seni seviyorum" cümlesi "Tıpkı senin beni sevdiğin gibi", "bir başkasının yaptığı gibi", "Başkasını seviyorum" ve "senin gibi" gibi pek çok anlam içermektedir (Anjali ve Babu, 2014).

2.4.2. Doğal Dil İşleme Seviyeleri

Doğal dil işleme alanının birkaç seviyesi vardır:

2.4.2.1. Karakter Seviyesi

Karakter seviyesinde metindeki karakter özellikler doğal dil işlemede kullanma hedeflenir. Bu özelliklerden arasında harfin sıklığı, büyük harf kullanımı, kullanılan karakterlerin toplamı (Maria, 2016), ayrıca Arapçada olduğu gibi karakterin metindeki durumu da özellikler arasında bulunmaktadır.

2.4.2.2. Sözcük Seviyesi

Sözcüksel veya dilsel nitelikler, tümceden çıkarılan tüm karakterleri ve kelime tabanlı istatistiksel ölçümleri içerir (Barbon ve Ark., 2017; Kaur ve Ark., 2018). Sözcüksel özellikler, sözcük düzeyine veya karakter düzeyine göre çıkarılabilir

(Brocardo ve Ark., 2017) öyle ki belirli bir sözcüğün metindeki sıklığı, sözcük sayısı, özgün sözcük toplam sayısı, sözcük bakımından zengin bir metin, hapax legomenon sayısı metinde tek kelime anlamına gelen dis legomenon, metinde iki kelime anlamına gelen dis legomenon, metinde üçlü olarak geçen kelime anlamına gelen tris legomenon. Bu seviye dilden bağımsızdır, yani hemen hemen tüm dillerde bir belirteç kullanılarak uygulanır (Maria, 2016). Belirteçleştirme genellikle NLP işlemede birincil bir adımdır, sözcük bağlamında cümleleri ayrılmış sözcüklere döndürmek için kullanılır (Schütze ve Ark., 2018).

2.4.2.3. Sözdizimsel Seviyesi

Sözdizimsel özellikler, bağlamdan bağımsız olarak yazar özelliğini tanımlar (Brocardo ve Ark., 2017) ve cümleleri tanımlamak için belge sınırlarını vurgulayan noktalama işaretleri aracılığıyla belirtilebilir (Brocardo ve Ark., 2017). Durdurma sözcükleri ve Konuşma Parçası Etiketleme (Part-Of-Speech - POS) ölçüleri, sözcüğün bağlam içindeki işlevini tanımlar (Maria, 2016). POS, tümcedeki kelimeler arasındaki ilişkiyi dilbilgisel olarak tanımlayan zamirler, edatlar, bağlaçlar, yardımcı fiiller olarak kategorize edilebilir (Juola, 2008).

2.4.2.4. Anlamsal Seviyesi

Bu seviye anlamı eş anlamlılar olarak işlemek için kullanılır ve anlamsal seviyede eş seslik veya eş anlamlılık gibi anlam dikkate alınır (Castro ve Lindauer, 2012) . Aslında anlamı anlamak makine için zor sorunlardan biridir, bir kelimenin bağlam içindeki anlamı, anlamsal ilişkiler gibi birçok faktörü içerir (Maria, 2016).

2.4.2.5. İçeriğe Özel Seviye

Aynı yaş seviyesindeki yazarlar, iletişim kurmalarını sağlayan belirli kelime türlerini kullanırlar, aynı konuda çalışan yazarlar da bunu yapar. Siyasette kullanılan kelime dağarcığı sanattan farklıdır. Ayrıca aynı milliyetten olan insanlar benzer davranışı gösterirler. Bu düzey metindeki bu kelime dağarcığını ölçmekle ilgilenir (Maria, 2016).

2.4.2.6. Yapısal Seviye

Bu seviye, bir metnin biçimini ve organizasyonunu analiz etmekle ilgilenir. Bloglar ve e-posta gibi çevrimiçi belgelerde daha fazla esneklik sağlar (Maria, 2016).

2.4.2.7. Kendine Özgü Seviye

Bu seviye, metindeki dilbilgisi veya heceleme hataları gibi yazım hatalarını ölçmekle ilgilenir. Aslında kullanılabilecek tüm hataları içeren bir liste yoktur (Maria, 2016).

2.4.3. Doğal Dil İşleme Uygulamaları

Doğal dil işlemenin birçok uygulama alanı vardır. Bu uygulama alanlarından bazıları:

- Yazarlık analizi
- Bilgi Çıkarma
- Soru Cevaplama
- Makine Çevirisi
- Öneri sistemi
- Duygu analizi
- Metin sınıflandırması
- Metin oluşturma
- Metin özetleme.

2.4.4. Doğal Dil İşleme Teknikleri

NLP'de kullanılan özellikler çıkarma bölümünde daha detaylı anlatacağımız birçok teknik vardır, bunların tümü alt seviyelerde çalışır. Bazı teknikler:

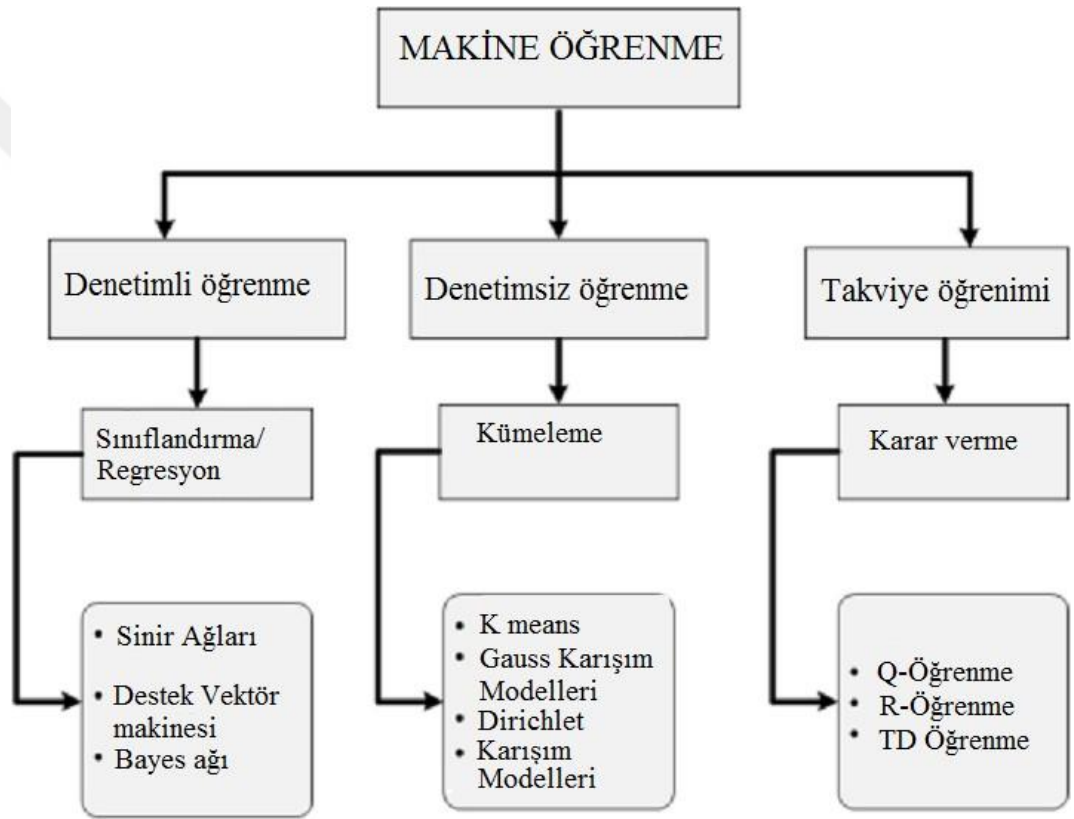
- Karakter N-gramı
- Kelime N-gramı
- POS etiketleri
- İşlev sözcüklerini ve noktalama işaretlerini kullanma
- Dilbilgisi hataları.

Bu çalışma, birkaç doğal dil işleme tekniğini birleştirerek metinden metinsel özellikleri çıkarmaktadır.

2.5. Makine Öğrenme Türleri

Makine öğrenme, yapay zeka dallarından biridir. Bilgisayarın önceden programlama yapmadan bir eylemi yapmasını sağlayan, başka bir deyişle belirli eylemlere programcı tarafından öğretmeden kendi kendine doğru şekilde nasıl tepki vereceğini öğrenen bilim olarak tanımlanabilir (Murphy, 2012).

Makine öğrenimi, aşağıdaki şekilde gösterildiği gibi denetimli, denetimsiz ve takviye öğrenme olarak sınıflandırılır (Sultan ve Ark., 2018):



Şekil 2-5. Makine Öğrenimi Türleri (Sultan ve Ark., 2018)

2.5.1. Denetimli öğrenme

Aynı zamanda tahmine dayalı öğrenme olarak da adlandırılır, bu türde makine, çıktısının ne olduğunu önceden bilen etiketli girdiden öğrenir (önceden spam veya spam olmayan olarak etiketlenmiş bir dizi e-posta) ve gerekli olan, makinenin gelecekte herhangi bir yeni girdiden çıktıyı tahmin edebilmesi için gereken girdi ile

çıktı arasında bağlantı öğrenmesidir. Denetimli öğrenmede en önemli modeller Sınıflandırma ve Regresyondur (Murphy, 2012).

- Sınıflandırma: Sınıflandırma, makine öğreniminde en çok kullanılan modeldir. Bu modelde girdi daha önce sınıflara ayrılmıştır ve öğrenmenin amacı, yeni girdiyi bu sınıflara ayırmaktır. Sınıflandırma örnekleri arasında e-posta sınıflandırması, yüz tanıma, çiçek görüntü sınıflandırması bulunmaktadır. Sınıflandırma modellerinden ikili sınıflandırma günlük yaşantımıza sıklıkla karşımıza çıkmaktadır.
- İki sınıflandırma modeli: Girdinin iki sınıftan birine sınıflandırıldığı modeldir (Robbins, 2012). Örneğin, tweet'in belirli bir yazara ait olup olmadığına göre, yazarın özelliğine göre sınıflandırıldığı ikili sınıflandırma modeli. Bu sınıflandırma türü bizim vaka çalışmamız için uygundur.
- Regresyon: Bu model sınıflandırmaya benzer, ancak işlevi ayrılmış sınıflar yerine değerler atamasıdır; Bu modelin döviz fiyatlarını tahmin etme, bir kişinin ömrünü tahmin etme gibi pek çok uygulaması vardır.

2.5.2. Denetimsiz öğrenme

Tanımlayıcı öğrenme de denir, önceki türün aksine, bu türde makine, bilinen herhangi bir çıktı olmadan sadece girdi verisi üzerinde öğrenir. Burada amaç, veriler arasında yeni bir desen ve gizli ilişkiler çıkarmaktır (Murphy, 2012). Denetimsiz öğrenmede en önemli yaklaşımlarından biri kümelemedir. Bu yaklaşımda girdiler önceden bilinmeyen kümelerle ayrılır. Bu yaklaşımda, birinin hareketlerini öğrenmesi gibi makinenin daha sonra bu hareketleri algılayabildiği ve uygun reaksiyonlarla ilişkilendirebildiği pek çok uygulama vardır. Ayrıca E-ticarete, kullanıcıları satın almalarına ve taramalarına göre kümeler halinde kümeleme sürecinden sonra bunları kullanarak her kümeyi ilgilendiren reklam mesajları göndermek için kullanılır.

2.5.3. Takviye öğrenimi

Bu tipte makine, eyleme ceza veya ödül işaretleri vererek belirli bir durumda hareket etmeyi öğrenir (Murphy, 2012).

2.6. Denetimli Öğrenme Algoritmaları

Makine öğrenimi, NLP problemlerini çözmek için kullanılabilecek çeşitli algoritmalar sağlamaktadır. Denetimli, denetimsiz veya takviyeli makine öğrenimi türü, ikili veya çoklu sınıflandırma türü, döndürülen doğruluk gibi uygun algoritmayı belirleyen çeşitli faktörler vardır.

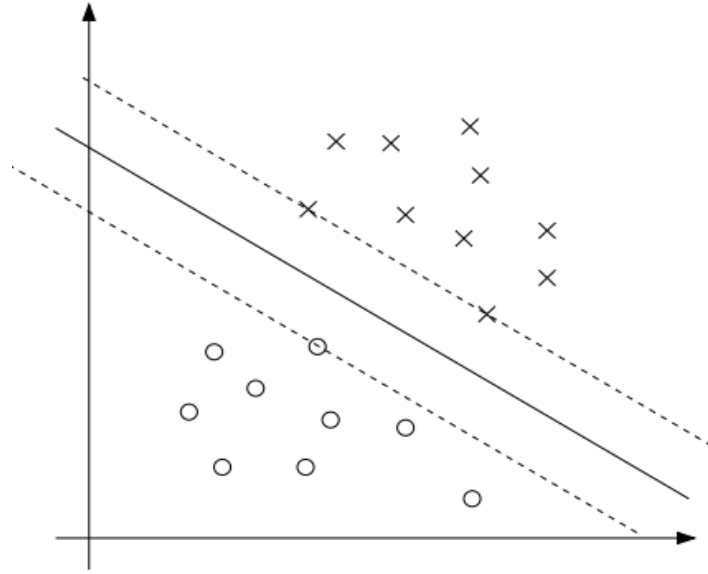
Araştırma vakası denetimli öğrenme ve ikili sınıflandırma altında ele alınacaktır ve tahmin yapmak için literatürde kullanılan birçok algoritma vardır. Bu algoritmalar, tahmin doğruluğu arasında karşılaştırma yapmak, etkili bir model oluşturmak için test edilecektir. Aşağıda deneylerde kullanılacak genel denetimli öğrenme algoritmaları için kısa bir açıklama bulunmaktadır.

2.6.1. Destek Vektör Makinesi (SVM)

SVM, sıklıkla kullanılan denetimli öğrenme algoritmalarından biridir (Benevenuto ve Ark., 2010), (Zangerle ve Specht, 2014) (Li ve Ark., 2014). İki tür veri arasında, bunları net bir şekilde sınıflandırmak için Hiperdüzlem olarak adlandırılan bir ayırma alanı bulmayı amaçlar, marj ne kadar genişse o kadar iyidir, hiperdüzlemi tanımlamak için kullanılan noktalara destek vektörleri denir (Andrew, 2021).

SVM, büyük verilerle eğitimde uzun zaman alır, çünkü Hiperdüzlemin en iyi yerini aramak için her seferinde bir özellik eklemeye bağlıdır, bu da onu az ve temiz veri için daha uygun hale getirir.

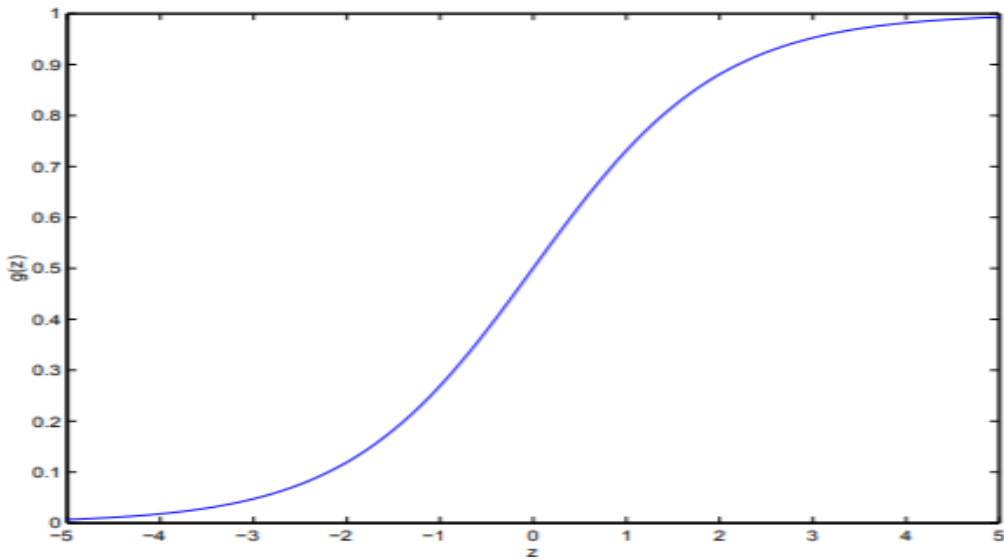
Aşağıdaki şekilde Hiperdüzlem iki veri sınıfı arasında ayrım yapmaktadır ve çizgiler üzerindeki üç nokta, karar sınırından en küçük marjları temsil eden destek vektörleridir, fakat bunları silmek karar sınırını yeniden tahsis etmektedir (Andrew, 2021).



Şekil 2-6. SVM algoritmasını kullanarak iki veri sınıfını ayırt etme (Andrew, 2021).

2.6.2. Lojistik Regresyon

Lojistik model veya logit model, kategorik bağımlı değişken ile çoklu bağımsız değişkenler arasındaki ilişkiyi analiz eder ve bir olayın meydana gelme olasılığını tahmin etmek için verileri lojistik bir eğriye sığdırır ve meydana gelme olasılığını öngördüğü için değerleri $[0,1]$ arasındadır (Park, 2013). Daha önce Seyler ve Ark. sosyal medya hesabının ele geçirilip geçirilmediğini tahmin etmek için kullanılmıştır (Seyler ve Ark., 2018). Lojistik regresyon kavramı için Şekil 2.7 gösterilmiştir.



Şekil 2-7. Lojistik regresyon konsepti için grafik gösterimi (Andrew, 2021)

2.6.3. Random Forest

Her biri, aşırı uyumu azaltmak için rastgele verilerle oluşturulmuş birkaç ağaçla temsil edilen ormanın ağaç sayısı arttıkça doğruluk ta artar; bu nedenle, aşağıdaki algoritma sözde kodunda en uygun ağaç sayısına ve derinliğine karar vermek için optimizasyon algoritmasına ihtiyaç vardır (Kumar, 2019).

- Eğitim kümesindeki vaka sayısının N olduğunu varsayın, bu durumda bu N vakanın numunesi rasgele değiştirme ile alınır.
- Özniteliklerin M girdi değişkeni varsa, her düğümde M 'den rastgele m değişken seçilecek şekilde bir $m < M$ sayısı belirtilir. Bu m üzerindeki en iyi bölme, düğümü bölmek için kullanılır. Ormanı büyütürken m değeri sabit tutulmaktadır.
- Her ağaç mümkün olduğu kadar büyütülür ve budama yapılmaz.
- n ağacın tahminlerini toplayarak yeni verileri tahmin edin (yani, sınıflandırma için çoğunluk oyları, regresyon için ortalama)”, örneğin, sınıflandırma durumunda, sınıfla olan tahmin sonucu daha yüksek oyu aldı ve regresyonda tahmin sonucu, sonuçların ortalamasıdır.

Random Forest, sınıflandırmada regresyondan daha iyi olduğu için, büyük veri setleri ve boyutları ile özellik seçiminde de kullanılmaktadır. Geliştiricilerin sadece farklı parametreleri ve rastgele tohumları denemesine izin veren bir kara kutu algoritmasıdır (Kumar, 2019).

2.6.4. XGBoost

Aşırı gradyan artırma “XGBoost” algoritması, "GBM" Gradient Boosting Machine'in geliştirilmiş bir versiyonudur. Gradyan artırma, çoklu katkılı regresyon ağaçları, stokastik gradyan artırma veya gradyan artırma makineleri gibi farklı isimlere sahiptir (Yozgatlıgil, 2017).

2016 yılında geliştirilen ve Kaggle yarışmasını kazanan ağaç algoritması, bir ağaç dizisi oluşturarak ve her ağacı bir önceki ağaçtaki eğitime göre geliştirerek, diğer güçlendirme algoritmalarından daha hızlıdır (Chen ve Guestrin, 2016). Boosting fikri, gradyan artırma denilen kaybı en aza indirmede gradyan iniş algoritmasını kullanarak hatalarını önlemek ve en iyi performans sonucunu elde etmek için mevcut modellere yeni modeller eklemeye devam eder (Yozgatlıgil, 2017). XGBoost,

Random Forest algoritmasının aksine, ağacı büyük ölçüde önlemek için torbalama tekniği olarak adlandırılan budama tekniğini kullanarak ağaçları olabildiğince az böler (Yozgatlıgil, 2017). Modelleri sıralı değil paralel olarak eğitmek, farklı modelleri manipüle ederek aşırı uyumu ele almak ve hata oranını azaltarak yanlılığı azaltmak için boosting'de hesaplamalı torbalama tercih edilmektedir (Yozgatlıgil, 2017).

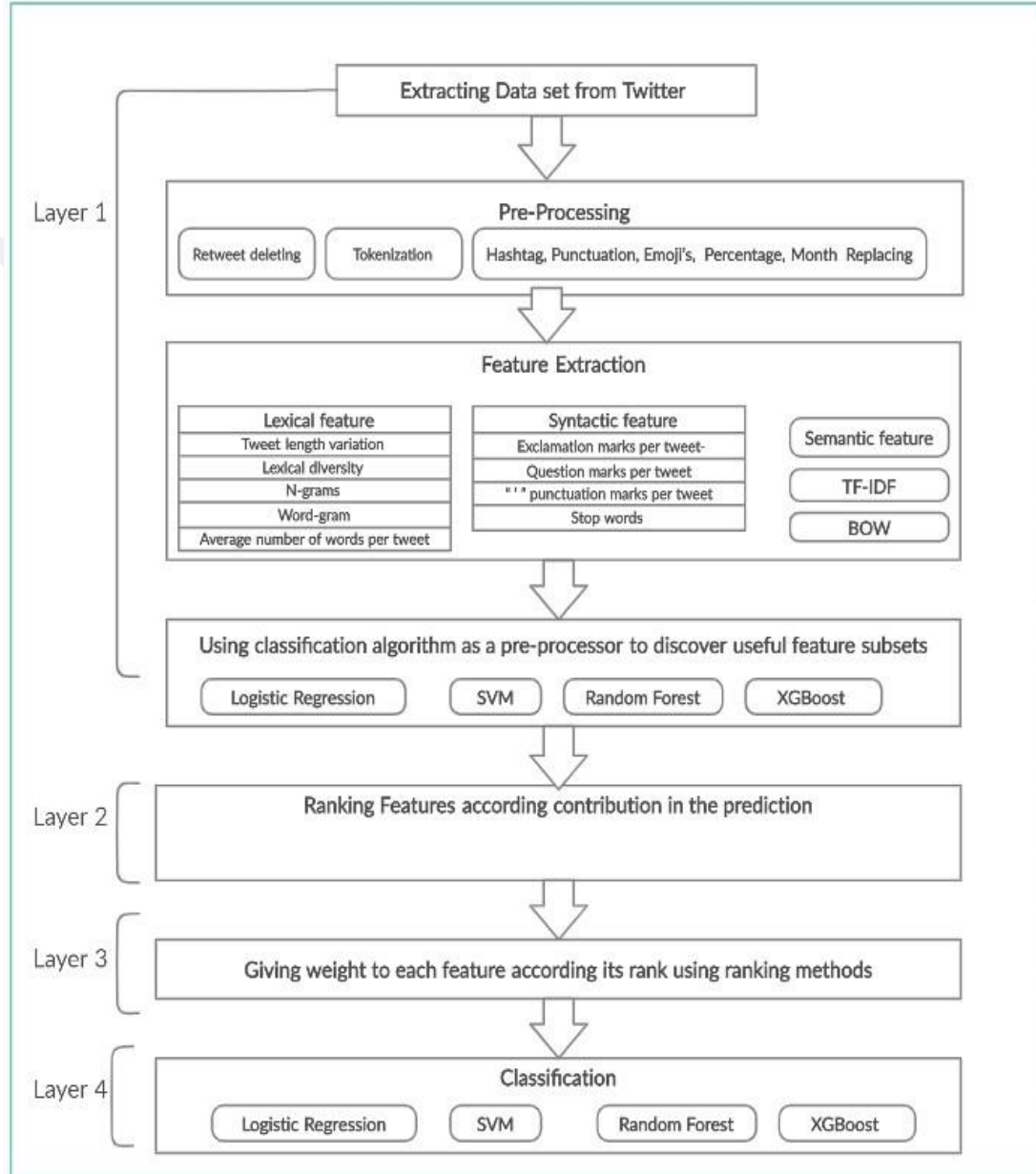
Aşağıdaki tablo, farklı makine öğrenimi algoritmaları ile yetenekleri arasındaki karşılaştırmayı göstermektedir (Blondel, 2021).

Algo	Type	Tolerance number features	Parametrization	Memory size	Minimal required quantity	Com m	Overfitting Tendency	Difficulty	Time for Learning	Time for predicting
Linear Regression	R	Weak	Weak	Small	Small	++	Low	Weak	Weak	Weak
Logistic Regression	C	Weak	Simple	Small	Small	++	Low	Weak	Weak	Weak
Decision Tree	R & C	Strong	Simple / intuitive	Large	Small	+++	Very high	Weak	Weak	Weak
Random Forest	R & C	Strong	Simple / intuitive	Very Large	Large	++	Average	Average	Costly	Costly
Boosting	R & C	Strong	Simple / intuitive	Very Large	Large	+	Average	Average	Costly	Weak
Naive Bayes	C	Weak	No params.	Small	Small	++	Low	Weak	Weak	Weak
SVM	C	Very strong	Not intuitive	Small	Large	--	Average	High	Costly	Weak
Neural Network (NN)	C	Very strong **	Not intuitive	Inter	Large	---	Average	Very high	Costly	Weak
Deep Neural Network	C	Very strong **	Not intuitive	Very Large	Very Large	---	High	Very high	Very costly	Weak
K-Means	CL*	Strong	Simple / Intuitive		Small	+		High	Weak	
One class SVM	A	Very strong	Not intuitive	Weak	Large	--	Average	High	Costly	Weak

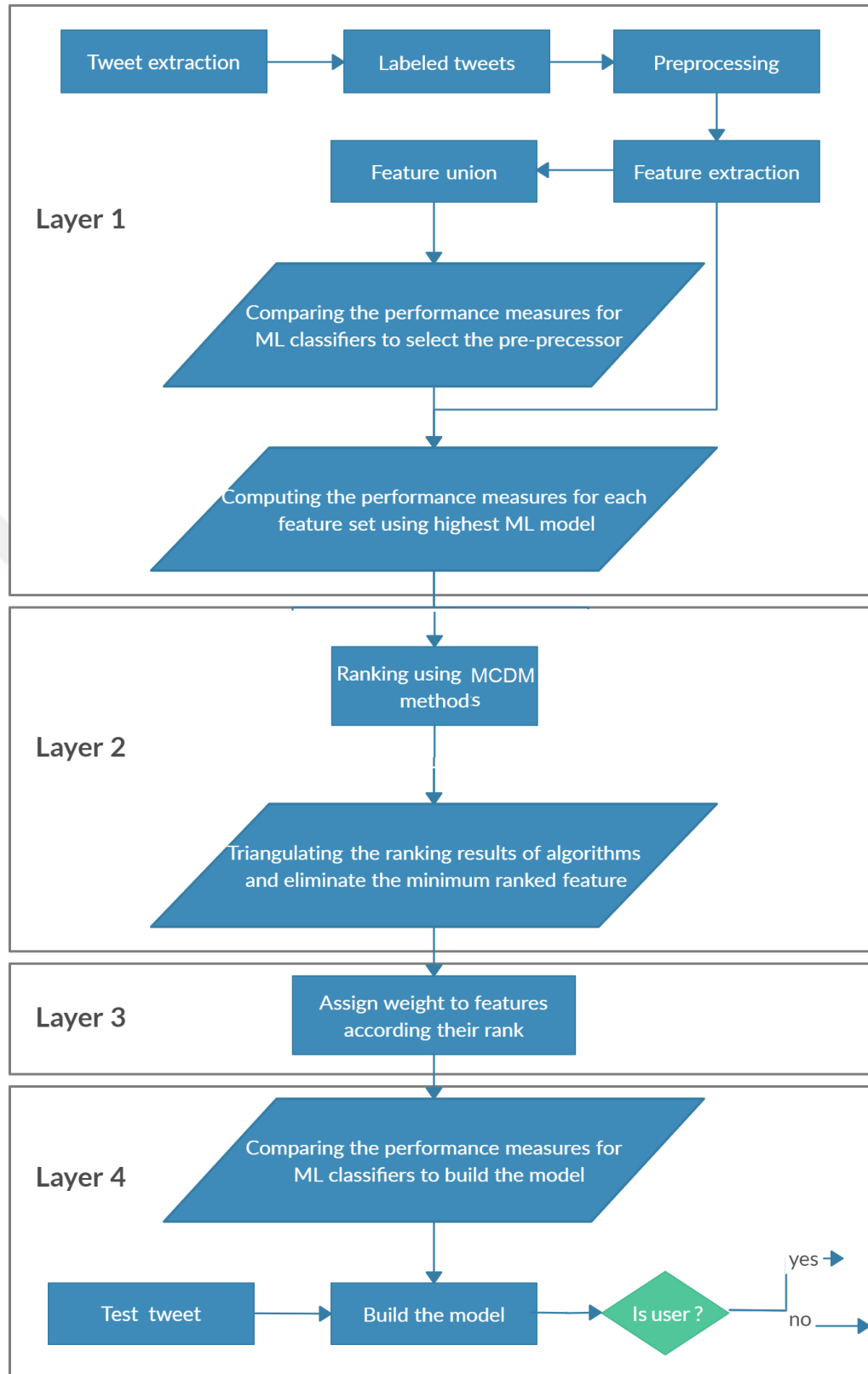
Şekil 2-8. Makine öğrenimi algoritmaları arasında karşılaştırma (Blondel, 2021)

3. ARAŞTIRMA BULGULARI

Bu bölüm önerilen modelle ilgili ayrıntıları sunmaktadır. Şekil 3-1. Modelin ayrıntılı mimarisi, Şekil 3-2. Akış şemasını sözde kodları ve yedi ana alt süreçlerini içeren akış şemasını göstermektedir.



Şekil 3-1. Modelin ayrıntılı mimarisi



Şekil 3-2. Akış şemasını

Girdi:	DatasetPath.
Çıktı:	Classifiers prediction
Adım 1	Tweetleri ön işleme /* Retweeti silinmiştir; noktalama işaretlerini, emojileri, hashtag'leri, yüzdeyi, ay adlarını değiştirilir ve tokenizasyon uygulandı.*/
Adım 2	Özellikler çıkarma ve birleştirme /* Sekiz özellik kümesi çıkarılmıştır BoW, tf-idf, kelime 2-gram, karakter 2-gram, Durma kelimeleri, Sözlüksel, Sözdizimsel, Anlamsal.*/
Adım 3	Yararlı özellik alt kümelerini keşfetmek için bir ön işlemci olarak bir birincil öğrenme algoritması kullanmak Foreach model in {XG Boost, Lojistik regresyon, Random Forest, SVM} Calculate Çalışma süresi, Doğruluk, Duyarlılık, Kesinlik, F1-skoru Return En yüksek performans gösteren model
Adım 4	Ön işlemci makine öğrenimi algoritmasını kullanarak tahmindeki katkıya göre özellik kümelerini listeleme Foreach kelimelerle ayarlanan özellik Calculate Çalışma süresi, Doğruluk, Duyarlılık, Kesinlik, F1-skoru Return tüm özellik kümesi performans ölçüleri
Adım 5	MCDM yaklaşımlarını kullanan özellik kümelerini sıralama
Adım 6	Algoritmaların sıralama sonuçlarını üçgenleme ve minimum dereceli özelliği ortadan kaldırma
Adım 7	Sıralama yöntemlerini kullanarak sıralamaya ağırlık atama
Adım 8	Ağırlıklı özellik kümeleri için hesaplama sınıflandırıcıları tahmini

Algoritma 3-1. Önerilen modelin temel aşamaları

Yazarlık doğrulama süreci modeli, Şekil 3-1. Modelin ayrıntılı mimariside gösterildiği gibi sekiz ana alt adımdan oluşmaktadır. İlk adım, kullanıcıların tweet'lerini, tweet'lerinin geçmişini, mevcut ve ilgili tüm özellikleri toplamaktır. İkinci adımda, toplanan bu tweet'ler kullanılmayan öz nitelikler arındırılarak ve tweet'ler tek bir formatta standartlaştırılmıştır. Üçüncü adımda, metinden farklı metinsel desenlere ait özellik vektörleri çıkarılıp sonrasında tweet'ler yığını temsil etmek için tek bir matriste birleştirilir. Dördüncü adımda, her bir özelliğin tweet'in yazarlığını doğrulama yeteneğini ölçmek için dört farklı sınıflandırıcı arasından en iyisi sınıflandırma algoritması bir ön işlemci olarak seçilir. Daha sonra, beşinci adımda MCDM yaklaşımları kullanılarak özellik setleri sıralanarak özellik indirgeme yapılır. Altıncı adımda, algoritmaların sıralama sonuçlarının üçgenlenmesi ve

minimum dereceli özelliğın ortadan kaldırılması. Yedi adımda, hatırlatma özellik setleri sıralama yöntemleri kullanılarak sıralarına göre ağırlıklar atanır. Son olarak, mesaj yazarlığını doğrulamada en yüksek performansı elde etmek için ağırlıklı özellik setlerine dört farklı sınıflandırıcı yöntemi uygulanır. Aşağıdaki alt bölümlerde, model adımlarının ayrıntılı bir açıklaması bulunmaktadır.

3.1. Veri Toplama

Veri toplama aşağıdaki adımlardan oluşmaktadır:

3.1.1. Yazarlıklarını Doğrulamak İçin Twitter Hesaplarını Belirlenme

320 milyondan fazla aktif kullanıcısı olan Twitter, en popüler mikroblog ÇSA'lerden biri haline gelmiştir (Singh ve Banerjee, 2019). Ne yazık ki, yazarlık doğrulama çalışmaları için herkese açık standart bir Twitter veri kümesi bulunmamaktadır (Boenninghoff ve Ark., 2019). Bu nedenle, yazarlık doğrulama modelleri geliştirmek için bir Twitter veri kümesi oluşturmak büyük önem arz etmektedir. Twitter verilerini elde etmenin en popüler yollarından biri onun API'sidir (Singh ve Banerjee, 2019). Ancak Twitter, API üzerinden, 280 karakter bloklu yaklaşık 3200 tweet ve bunlarla ilgili özelliklerin çeşitli istekler yoluyla sınırlı miktarda veri alınmasına izin vermektedir (Steinert-Threlkeld, 2018). Bu nedenle, veri kümesini oluşturmak için Python programlama dili ve tweepy kütüphanesi kullanılarak bir crawler yazılımı geliştirilmiştir. Veriler, aynı konu ve türe sahip farklı kullanıcılardan toplanmıştır. Bilinmelidir ki, gerçek hayatta, yazarlar konu ve tür bakımından farklılık gösterir (örneğin, belgeler, e-posta, tweet, vb), bu yüzden önemli olan yazarın stilometrisine odaklanmaktır (Usha ve Thampi, 2017; Barlas ve Stamatatos, 2020). Bu arada, çapraz konulardan veya çapraz türlerden gelen bilgiler modeli yanlış yönlendirebilir (Barlas ve Stamatatos, 2020), bu da yazarlık doğrulamasını zorlaştırmaktadır (Kestemont ve Ark., 2012).

100 kullanıcıdan oluşan bir örneklem seçilmiştir ve her biri kendini teknoloji alanında yazar olarak tanıtmaktadır ve Twitter'daki içerikleri sadece teknolojiye odaklanmaktadır. Bu da araştırma sonucunu daha güçlü kılmakta, makinenin konuların farklılaşmasına değil sadece yazarın tarzına odaklanmayı sağlayacaktır; Ancak, güvenliği ihlal edilmiş veya ele geçirilmiş mesajın siyaset veya edebiyat gibi

farklı bir alandan geldiğini tahmin etmek zor değildir. Ancak model, konudaki değişimi tespit edebilecektir. Model doğruluğunu artırmak için kullanıcıların aynı etki alanından seçilmiş olması, Usha ve Thampi'nin belirttiği gibi “Bunun nedeni, yazarlar arasında muazzam derecede değişen konuların, yanıltıcı bir şekilde doğru sınıflandırma modellerine ve güvenilir eğitim örneklerine yol açabilmesidir” (Usha ve Thampi, 2017).

businessinsider.com'dan, hesap örneğini seçmek için twitter'da en çok etkilenen yüz teknoloji uzmanının referans olarak kullandığı bir rapor mevcuttur (Borison, 2021). Hesaplar kişisel hesap olup, bu da sadece hesap sahibinin hesap üzerinden tweet attığı anlamına gelmektedir. Ancak çalışmamız 280 karakterli tweetleri kullandığından 3200 tweet'in Eylül 2017'den sonra tweet'lenmesi ön koşulu vardır. Yazarların tamamı sık tweet atmadığından, bu koşul tüm örnek hesaplar için mevcut olmayıp, 5 kullanıcının çalışma koşulunu sağlayan hesaplardan rastgele seçildiğinden ne zaman daha fazla tweet atılması modelin eğitiminde etkili olmuştur.

Son dipnot olarak veriler, kullanıcının ekran adı üzerinden alınabilir, ayrıca kullanıcı tarafından değiştirilebilir ve twitter buna izin verir; bu nedenle, verileri user_id üzerinden almak daha doğrudur, kullanıcıya kimlik numarasını temsil eden ve değiştirilemez (Steinert-Threlkeld, 2018).

3.1.2. Nesnelerin Seçilmesi

Twitter, tweet meta verilerini temsil eden tweet ile birçok nesnenin alınabilmesini sağlar ve bu nesneler bir kök olabilir ve içinde {} içine koyan birçok alt dal olabilir, aşağıdaki tüm mevcut nesnelerde:

- Kullanıcıyla İlgili Nesneler (Developer.twitter.com: user objects, 2021)
"user" : {"id" , "id_str" , "name" , "screen_name" , "location" , "description" ,
"url" , "followers_count" , "friends_count" , "listed_count" , "verified" ,
"favourites_count" , "statuses_count" , "created_at" , "utc_offset" ,
"time_zone" , "geo_enabled" , "lang" ,
"profile_background_image_url_https" , "entities" : {"url" : {"urls" ,
"expanded_url" , "display_url" , "indices" , "description" , "protected" ,
"contributors_enabled" , "is_translator" , "is_translation_enabled" ,
"profile_background_color" , "profile_background_image_url"

"profile_background_tile" , "profile_image_url" , "profile_image_url_https" ,
"profile_banner_url" , "profile_link_color" , "profile_sidebar_border_color" ,
"profile_sidebar_fill_color", "profile_text_color",
"profile_use_background_image" , "has_extended_profile" , "default_profile"
 , "default_profile_image" , "following" , "follow_request_sent" ,
"notifications" , "translator_type" }, "geo", "contributors" }

- Tweetle ilgili nesneler (Developer.twitter.com: tweet objects, 2021)
{ "created_at", "id", "id_str" , "text", "source", "truncated" ,
"in_reply_to_status_id"
 , "in_reply_to_status_id_str", "in_reply_to_user_id", "in_reply_to_user_id_str",
"in_reply_to_screen_name", "coordinates": { "coordinates", "type" }, "place": { "i
d", "url", "place_type", "name", "full_name", "country_code", "country",
"attributes", "bounding_box" : { "type", "coordinates" }, "quoted_status_id",
"quoted_status_id_str" "is_quote_status", "quoted_status",
"retweeted_status", "quote_count", "reply_count", "retweet_count" ,
"favorite_count" , "entities": { "hashtags" , "symbols" ,
"user_mentions", "urls", "media", "polls" }, "favorited", "retweeted",
possibly_sensitive" , "filter_level", "lang", "matching_rules" { "tag", "id",
"id_str" } }.

Tüm bu nesneleri araştırıp almaya çalıştıktan sonra, onları dört gruba ayırabiliriz:

- REST API aracılığıyla alınan nesneler
Bu nesneler farklı python sözdizimi tarafından alınır, örneğin “dil”,
“tweet.lang” sözdizimi kullanılarak alınır, "hashtag'ler",
"tweet.entities['hashtags']" sözdizimini kullanarak alınırken, “medya”,
“tweet.entities.get('media',[])" sözdizimini kullanarak alıyor, ve kullanıcı adı
“tweet.user.name” sözdizimi kullanılarak alınıyor.
- STREAM API aracılığıyla alınan nesneler
REST API yerine STREAM API aracılığıyla birkaç nesne alınır, örneğin:
kaynak tweet'in kimliğini temsil eden "quoted_status_id" ve
"quoted_status_id_str", alıntılanan tweet olan “quoted_status”, ve tweet
hakkında iyi veya kötü bir Boole değeri alan "possibly_sensitive".
- Premium plan aracılığıyla alınan nesneler

Tweet ile ilgili birkaç nesne, yalnızca "quote_count" ve "reply_count" gibi kurumsal ve premium planlar için sağlanmıştır (Developer.twitter.com: user objects, 2021).

➤ Devlet kısıtlaması nedeniyle yasak nesneler

Genel Veri Koruma Yönetmeliği'nin (GDPR) kısıtlaması nedeniyle ve kullanıcı gizliliğini korumak için 23 Mayıs 2018'de birkaç nesneyi çağırmak, "tweet.user.utc_offset" ve "tweet.user.time_zone" gibi boş değerler alıyor. Aynı kısıtlama nedeniyle diğer nesneleri çağırmak, ilgili geçmiş olmadan varsayılan değeri alıyor, bu tür "profile_background_image_url" ve "profile_background_image_url_https" nesneleri (Piper, 2021).

3.1.3. Tarayıcıyı Geliştirmek

Algoritma 3.2 twit API adımlarını göstermektedir. Girdi olarak erişim verileri ve tarama koşulları bulunmaktadır. Çıktı ise tweetlerdir.

Girdi:	Access Tokens, Retriving Conditions
Çıktı:	Datasets
Adım 1	Gerekli kütüphaneleri içe aktarma
Adım 2	Gerekli kimlik bilgilerinin atanması /* Verileri almak için uygulamaya gerekli kimlik bilgilerini atayan API anahtarlarından bahsetme */
Adım 3	Fonksiyonu tanımlama
Adım 4	Twitter yetkilendirme
Adım 5	2018'in başlangıcından sonra verilerin alınması /* Daha önce de belirtildiği gibi twitter, tweet'te izin verilen harfleri 140 harften 280 harfe çıkarmaya başladı Eylül 2017'den sonra */
Adım 6	Istek gönderme /* Her istekte 200 tweet'in alınmasına izin verilir, bu nedenle her kullanıcı için "Başlangıç Tarihi"nden sonra yazılan mevcut tüm tweet'leri almak için bir döngü gereklidir new_tweets = api.user_timeline(user_id = user_id,count=200) alltweets.extend(new_tweets) oldest = alltweets[-1].id - 1 while (len(new_tweets) > 0) and (new_tweets[-1].created_at >

```

        startDate):
            new_tweets = api.user_timeline(user_id =
            user_id,count=200,max_id=oldest )
            alltweets.extend(new_tweets)
            oldest = alltweets[-1].id - 1

    */
Adım 7   Gerekli nesnelerin belirlenmesi
    /*
        Kullanıcı veya tweet ile ilgili gerekli nesneleri belirleme
    */
Adım 8   Alınan verileri .csv dosyasına kaydetme
Adım 9   Her kullanıcı için alınan tweetleri sayma
Adım 10  Metin dosyasından user_id okuma
    /*
        Fonksiyon bu dosyadan ID'leri okuyacak ve her user_id için veriler
        toplanacak ve .csv dosyasına kaydedilecektir
    */
Adım 11  Zamanlanmış Görevi fonksiyona ekleme
    /*
        API'yi kapatıp 15 dakika sonra tetiklemek için fonksiyona bir
        zamanlanmış görev eklendi; twitter ise bağlantıyı kapatır veri sınırına
        erişildiyse ve 15 dakikada 180 istek yapılmasına izin verir, her istekte
        200 tweet alınabilir.
    */

```

Algoritma 3-2. Veri tarayıcısı için Pseudo-code

Twitter veri kümesi, farklı kullanıcılardan alınan 16124 örneği içerir. Burada tweet ile ilgili alınan nesneler sırasıyla yazarın "kimliğini" ve mesajı temsil eden "kimlik" ve "metin"dir. Yazarlık doğrulamasının tek sınıflı bir sınıflandırma problemi olduğu düşünüldüğünde (Halvani ve Ark., 2016), örnekler bilinen yazara ait tweet kümesi için 1 ve bilinmeyen yazara ait tweet kümesi için 0 ile etiketlenmiştir.

3.2. Verileri ön işleme

Verileri ön işleme, verileri analize hazırlamak için çok önemli bir süreçtir. Benzer çalışmalarda kullanılan ön işlemlerden bazılarında bahsetmek gerekirse, örneğin; Barbon ve ark. önemli yazarlık doğrulamada etkisinin olmadığı için metinden durma kelimelerini çıkarmıştır (Barbon ve Ark., 2017).

Mikros ve Perifanos çalışmalarında @replies, hashtag'leri ve üç kelimeden küçük tweetleri kaldırmış, ayrıca her yirmi üç tweet'i tek bir korpusta birleştirmiş, yirmi üç

tweet'i tek bir korpusta birleştirmenin yazar kimliğinin doğruluğunu artırdığını belirtmişlerdir (Mikros ve Perifanos, 2013).

McCollister, her tweet'i bir belge olarak değerlendirip ve neredeyse birbiriyle aynı tweet'leri ile birlikte bilinmeyen kaynak kullanıcıları, ayrı kelimeleri, kısaltılmış URL'leri, nadir ve yaygın kelimeleri kaldırmış ve tüm harfleri küçük harfe dönüştürmüştür (McCollister, 2016).

Maitra ve Ark. tüm harfleri küçük harfe dönüştürmüş ve kelimeler için kök ve POS etiketi kelimelerini çıkarmak için dosyayı kelimelerine ayıştırmıştır (Maitra ve Ark., 2016).

Zangerle ve Specht noktalama işaretlerini kaldırmış, tüm harfleri küçük harfe dönüştürmüş ve İngilizce olmayan kelimeleri kaldırmıştır. Araştırmacılara göre URL'lerin kaldırılması sınıflandırma doğruluğunu artırdığı, ancak hashtag'lerin, durma kelimelerini, kullanıcıların bahsettiği veya kelimelerin kökünün çıkarılması doğruluğu artırmamaktadır (Zangerle ve Specht, 2014).

Brocardo ve Ark. e-posta veri kümesi üzerindeki deneyleri aracılığıyla aşağıdaki ön işleme adımlarını gerçekleştirmişlerdir (Brocardo ve Ark., 2017):

- E-posta adresini “@” karakterleriyle değiştirmek
- HTTP adresini “http” kelimesiyle değiştirmek
- Para biriminin “XX\$” ile değiştirilmesi
- Yüzdeyi “%XX” ile değiştirme
- Farklı sayıların “0” ile değiştirilmesi
- Tüm harfleri küçük harfe dönüştürdü
- Boşlukları ve noktalama işaretlerini silin
- Tüm kullanıcı mesajlarını tek bir korpusta toplayın

Rao ve Ark. metni UTF-8 kodlamasına dönüştürüp, durma kelimelerini, virgülleri, tam durak sayılarını ve özel karakterleri metinden temizlemiştir (Rao ve Ark., 2017).

Usha ve Thampi kısa tweetleri birleştirip, birçok yazarın ilgili konudaki tweetlerini tutarken ve alıntı yapılan tweet'leri, retweet'leri ve bağlantıları çıkarmıştır (Usha ve Thampi, 2017).

Rocha ve Ark. ingilizce olmayan ve kısa (dört kelimedenden az) mesajları çıkarmışlardır. Araştırmacılara göre bu tür kısa mesajların yazar stilini temsil

etmediği ve sınıflandırma sürecinde gürültü yaptığını belirtmektedirler. Ayrıca yazar tarafından yazılmadığı için retweetleri temizlemişler ve tarihler, sayılar, URL'ler ve saatler gibi seyrek özellikler ile birlikte harf olmayan tüm karakterleri bir sembolle değiştirmişlerdir (Rocha ve Ark., 2016). Toplanan veriler, yazarın stilini tam olarak ifade etmek ve gürültüyü takip etmek için simple-regex kullanılarak ön işleme tutulmuştur. Rocha ve Ark. (Rocha ve Ark., 2016), yazarlık doğrulamasındaki veri kümesi dikkatli bir şekilde ön işlemden geçirilmesi gerektiğini belirtmektedir. Aynı zamanda, veriyi kaldırmak veya yeniden şekillendirmek, yazarın kendine özgü özelliklerini etkileyebildiğini vurgulamaktadırlar.

Sonuç olarak verileri hazırlamak için birkaç farklı adımı takip edebiliriz ancak, bunların tamamı mevcut çalışma için uygun olmayabilir, yani bazı temizlenmiş kelimeler veya karakterler, yazarlık doğrulamasında önemli özellikleri temsil edebilir. Bu nedenle, birkaç yazardan 3000 tweet örneği manuel olarak incelenmiş ve bu aşağıdaki ön işleme adımlarını önermiştir:

- Retweetler, yazarın stilini yansıtmadığı için silinmiştir.
- Noktalama işaretleri “!” ile değiştirilir.
- Emojiler “?” ile değiştirilir.
- “#” işaretinden sonra bir kelime ile temsil edilen tüm hashtag'ler "hash" ile değiştirilir.
- Yüzde işaretleri "%", "/p" ile değiştirilir.
- Ay adları "m" ile değiştirilir.
- Veri kümesine tokenizasyon uygulandı.

3.3. Özellik Çıkarma

Makine öğrenimi algoritmaları, farklı boyutlarda diziler içeren metin verilerinden değil, önceden belirlenmiş boyuttaki sayısal vektörlerden öğrenmek için tasarlanmıştır. Bu nedenle, devam etmeden önce metin verileri sayısal vektörlere çevrilmelidir. Burada, aşağıdaki sayısal özellikler çıkarılmıştır.

3.3.1. Sözlüksel Özellik

Aşağıdaki sözcüksel özellikler kullanılmıştır:

- Tweet başına ortalama kelime sayısı.

- Tweet uzunluğu varyasyonu.
- Tweetlerin ne kadar zengin kelime dağarcığı içerdiğini gösteren sözcük çeşitliliği, başka bir deyişle kullanıcının aynı kelime dağarcığını aynı veya farklı bir anlamı ifade etmek için kullanması; örneğin metin boyunca her seferinde "referans" kelimesi farklı anlamda kullanması veya "source", "resource", "source material", "weld".
- N-gram: 2-gram veya bigram karakterlerinin kullanılması, örneğin "love" kelimesinde 2-gram karakteri: "lo, ov, ve".
- Word-gram: 2-gram kelimelerinin kullanılması, örneğin "oğlan sütü sever" cümlesinde 2-gram veya bi-gram kelimesi "oğlan sütü, sütü sever"dir.

3.3.2. Sözdizimsel Özellik

Aşağıdaki sözdizimsel özellikler kullanılmıştır:

- Tweet başına ünlem işareti.
- Tweet başına soru işaretleri.
- Tweet başına " ' " noktalama işaretleri.
- Durma kelimeleri.

3.3.3. Anlamsal Özellik

Anlamsal özellikler, bir kelimenin veya cümlenin dilsel bağlamdaki anlamını ve diğer dilsel birimlerle olan ilişkilerini anlamaya çalışmak için kullanılır (Maria, 2016). Çalışmada çıkarılan aşağıdaki anlamsal bileşenler kelimeyi gömme, BoW, ve term frequency-inverse document frequency (tf-idf).

Kelime gömme, Kelime gömmede metni sayısal vektörlerde temsil ederken, aynı anlama sahip kelimeler tam temsile sahiptir. Kelimeler vektörleştirilip, birleştirilerek gömme kelimesi oluşturulur. Vektör uzunluğu, kelimeyi tanımlayan özellik numarası kullanılarak hesaplanır (örneğin, 200 özellik olduğunu varsayalım, sonra vektör uzunluğunun 200 olduğunu varsayalım), bu durumda özellik sayısı, toplam kelime sayısından küçüktür, her bir özellik değeri $[-1, 1]$ arasındadır ve değer 1'e yakın olduğunda kelimeyi temsil eder.

Aşağıdaki örnek, gömme kelimesini açıklamaktadır:

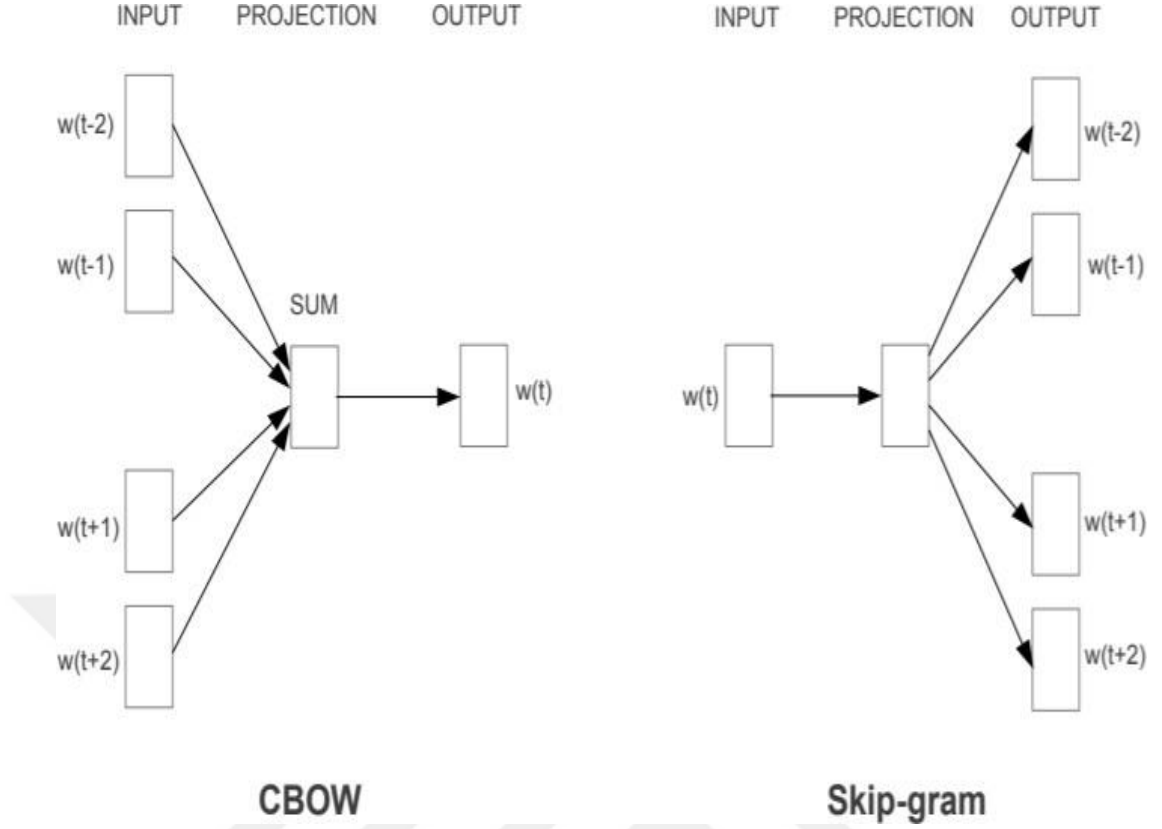
3 kelime (erkek, kız, portakal) ve 3 özellik (cinsiyet, uzun, meyve) olduğunu varsayarsak, bunlar her kelime vektörünü ifade eden aşağıdaki tabloyu verir:

Çizelge 3-1. Gömme kelimesinin örneği

	Cinsiyet	Uzun	Meyve
Erkek	-1	0.2	0
Kız	1	0.3	0
Portakal	0	0	0.95

Cinsiyetin sadece erkek, kadın ile ilgili olduğu ve portakalla ilgili olmadığı, bunun için portakalın cinsiyet değerinin 0 olduğuna dikkate ediniz, ayrıca, erkek için değer -1, kadın için ise 1'dir, burada iki kelime oldukça zıttır. Modelin algoritması Word2vec olarak adlandırılan 2013 yılında Mikolov tarafından geliştirilen önceden eğitilmiş bir kelime gömme işlemidir (Mikolov ve Ark., 2013a).

Mikolov ve ark. (Mikolov ve Ark., 2013b), veri kümelerinden kelimelerin vektörünü temsil etmek için kullanılan iki ilişkili model önermiştir. Birincisi, sürekli kelime torbası modelidir ve geleneksel kelime torbası'ndan farklı olarak, kelimeyi bağlama göre tahmin eder. İkincisi, çevreleyen kelimeleri tahmin etmek için giriş kelimesini kullanan sürekli "Skip-gram"dır (Şekil 3-3).



Şekil 3-3. CBOW ve Skip-gram modelleri mimarisi (Mikolov ve Ark., 2013b)

Kelime Torbası: BoW, bir metnin, dilbilgisi kurallarına ve kelimelerin metin içindeki sırasına bakılmaksızın çok sayıda kelime kümesinde temsil edilmesidir. BoW, çoklu kümede sadece kelimenin frekansına odaklanır (Sivic ve Zisserman, 2008).

Aşağıdaki örnek, modeli açıklar:

Cümle 1: I like programming

Cümle 2: programming is interesting

Çizelge 3-2. BoW örneği

Kelime	I	Like	Programming	Is	interesting
Belge Frekansı	1	1	2	1	1

İlk kayıt, iki cümleli dökümandaki tüm benzersiz kelimeleri içerirken, ikinci satır, dökümandaki her kelimenin sıklık sayısını temsil eden Döküman Frekansı içerir.

Term frequency-inverse document frequency: tf-idf bir metnin, tüm belgelerdeki görünümüne kıyasla bir belgenin içinde ne sıklıkta görüldüğünü hesaplayarak, kelimenin önemini belge aracılığıyla yansıtan istatistiksel bir temsildir (Ramos, 2003).

(tf-idf) modeli için kısa terim iki bölümden birleştirilmiştir:

Terim frekansı (tf): kelimenin bir belge içinde ne sıklıkla geçtiği ile ilgilidir.

Ters belge frekansı (idf): belgelerin içinde sıkça geçen sözcüğün hangi sınıfa indirilmesidir.

Yukarıdaki örnek için, iki cümlede tf aşağıdaki gibi olacaktır:

Çizelge 3-3. tf örneği

Kelime	I	Like	Programming	Is	interesting
Cümle 1	1	1	1	0	0
Cümle 2	0	0	1	1	1

Ve idf aşağıdaki gibi olacak:

Çizelge 3-4. idf örneği

Kelime	I	Like	Programming	Is	interesting
Cümle 1	1	1	½	0	0
Cümle 2	0	0	½	1	1

3.4. Ön İşlemci Makine Öğrenimi Algoritmasını Seçme

Farklı dilsel özelliklerin ayıklanması, her özelliğin meta-öğrenme aşaması olarak adlandırılan tweet'in yazarlığını doğrulama yeteneğini ölçmenin önemini artırmaktadır. Hospedales ve arkadaşlarına göre, meta-öğrenme, birçok aşamada makine öğrenimi modeline deneyim (Hospedales ve Ark., 2020) transferi yapar. Bununla birlikte, son on yılda, bilgi işlem gücündeki ve mevcut veri setlerindeki gelişmeler, makine öğrenimi algoritmalarının çeşitliliğinde bir artış sağlamıştır. Ayrıca, meta-öğrenme, algoritma seçimini ve parametrelerinin ayarlanmasını daha az karmaşık hale getirdi (Lemke ve Ark., 2015). Öte yandan, Lemke ve arkadaşlarına göre SVM, k-en yakın komşu gibi geleneksel ML algoritmaları, meta-öğrenme

algoritması seçiminde oldukça başarılıdır (Lemke ve Ark., 2015). Schönfeld ve ark (Schoenfeld ve Ark., 2018), pratik makine öğrenimi çözümlerinin, tek bir algoritma yerine makine öğrenimi operatörlerinin ardışık düzenlerine sahip olması gerektiğini savunurken, ihtiyaç sadece bir sınıflandırma algoritması seçmek değil, bununla birlikte bir önışlemci algoritması içeren işlem hattı oluşturmaktır.

Özellik boyutunu azaltmak ve dolayısıyla sınıflandırıcı doğruluğunu artırmak için, çıkarılan veri kümesi SVM, lojistik regresyon, random forest ve XGBoost algoritmaları gibi birçok meta-öğrenme algoritması kullanılarak değerlendirilmiştir.

Öte yandan, önışlemcilerin performansını ölçmek için doğruluk, duyarlılık, kesinlik ve F1-skoru gibi farklı metrikler kullanılmıştır, aşağıda tanımlanmıştır (Gadekallu ve Ark., 2020) (Kaur ve Ark., 2018):

Çizelge 3-5. Performans ölçüm parametreleri

T	Test edilen yazar tarafından yazılan tweetler
N	Diğer yazarlar tarafından yazılan tweetler
TP	Tweet, test edilen yazara aittir ve sistem bunu doğru tahmin etmiştir.
FP	Tweet, test edilen yazara ait değil ve sistem onları kendisine ait olarak tahmin etti
FN	Tweet, test edilen yazara aittir ve sistem onları yanlış tahmin etmiştir.
TN	Tweet, test edilen yazara ait değil ve sistem onları doğru tahmin etti

- Doğruluk: aşağıdaki gibi tanımlanabilir

$$\text{Doğruluk} = \frac{\text{doğru tanınan tweetler}}{\text{toplam tweet sayısı}} \quad \text{Denklem 3-1}$$

- Duyarlılık: "Doğru Pozitif"+"Yanlış Negatif" olayların toplamı üzerindeki "Doğru Olumlu" olayların yüzdesi

$$\text{Duyarlık} = \frac{TP}{TP+FN} \quad \text{Denklem 3-2}$$

- Kesinlik Skoru: "Doğru Pozitif"+"Yanlış Pozitif" oluşumların toplamı üzerindeki "Doğru Pozitif" oluşumların yüzdesi

$$\text{Kesinlik} = \frac{TP}{TP+FP} \quad \text{Denklem 3-3}$$

- F1-skoru : " Kesinlik" ve " Duyarlık" ölçülerine bağlı bir ölçüdür

$$\text{F1-skoru} = 2 * \left(\frac{\text{Kesinlik} * \text{Duyarlık}}{\text{Kesinlik} + \text{Duyarlık}} \right) \quad \text{Denklem 3-4}$$

Özellik seçim algoritması olarak kullanılacak en iyi performans gösteren önışlemciyi seçmek için, birkaç ML algoritması ile bazı deneyler yapıldı ve çalışma zamanı dahil performansları ölçülmüştür. Çizelge 3-6 sonuçları göstermektedir.

Çizelge 3-6. Bazı makine öğrenimi algoritmalarının performansı

	Çalışma süresi	Doğruluk	Duyarlılık	Kesinlik	F1-skoru
Lojistik regresyon	1.06 Min	0.9029	0.962	0.90	0.939
SVM	0.99 Min	0.9029	0.947	0.90	0.938
Random forest	0.63 Min	0.7903	0.965	0.760	0.877
XGBoost("Maks Derinlik"=7)	21.71 Min	0.8932	0.9599	0.8896	0.933
XGBoost("Maks Derinlik"=9)	27.23 Min	0.9049	0.967	0.902	0.94

XGBoost algoritma açıklamasına göre, booster türü gibi daha iyi sonuçlar verecek şekilde birçok parametre ayarlanabilir. Booster parametreleri, öğrenme süreci parametreleriyle ilgilidir ve aşırı uyum parametrelerinden kaçınılır. Ayrıca, ağaç derinliğine maksimum değeri temsil etmek için maksimum derinlik gibi ağaçla ilgili parametreler vardır ve ağacın içindeki al dal ağırlığı için minimum değeri temsil eden Min dal Ağırlığı, eğitim sürecine giren örneklerin yüzdesini temsil eden alt örnek, ağacı oluşturmak için kullanılan sütunların oranını temsil eden ağaca göre örnek olarak ifade edilebilir.

Parametreler hakkında daha fazla bilgi için bakınız (Castro ve Lindauer, 2012) . "Maks Derinlik" = 9, birçok turdan sonra en yüksek doğruluğu veren bu parametre için optimize edilmiş değerdı. Çizelge 3-6. Bazı makine öğrenimi algoritmalarının performansı, parametreleri ayarlandıktan sonra XGBoost performansındaki değışikliği yansıtır.

3.5. Özellik Seçimi

Kotsiantis ve ark. araştırmalarında, yararlı özellik alt kümelerini keşfetmek için bir önışlemci olarak bir birincil öğrenme algoritması kullanmasının mümkün olduğunu belirtmişlerdir (Kotsiantis ve Ark., 2006). Çizelge 3-7. Çıkarılan metinsel özelliklerin performansı arasındaki karşılaştırma, XGBoost algoritmasını kullanarak bir tweet'in yetkisini tanımda çıkarılan her bir özelliğin katkısını yansıtmaktadır.

Çizelge 3-7. Çıkarılan metinsel özelliklerin performansı arasındaki karşılaştırma

	Çalışma süresi	Doğruluk	Duyarlılık	Kesinlik	F1-skoru
BOW	10.32 Min	0.7867	0.91359	0.7789	0.8556
tf-idf	10.51 Min	0.7844	0.9168	0.7766	0.8548
kelime 2-gram	20.96 Min	0.7257	0.9347	0.7070	0.8251
karakter 2-gram	18.61 Min	0.87363	0.9380	0.8718	0.9112
Durma kelimelerini	73.64 Min	0.7867	0.9135	0.7789	0.8556
Sözlüksel	27.88 Min	0.7748	0.8997	0.7520	0.8609
Sözdizimsel	28.82 Min	0.7748	0.8997	0.7520	0.8609
Anlamsal	0.66 Min	0.7848	0.9006	0.7653	0.8674

Özellik seçiminde kullanılan birçok algoritma vardır, örneğin: χ^2 (Chi-Squared) yöntemi (Benevenuto ve Ark., 2010), bilgi kazanımı (IG) (Criado ve Ark., 2016), Grey İlişkisel Analiz “GRA” teorisi, çok amaçlı optimizasyon yöntemi (Ma ve Ark., 2018), ve Çok Kriterli Karar Verme “MCDM” yaklaşımlarından biri olan TOPSIS (Singh ve Ark., 2014; Kaur ve Ark., 2018). Aşağıdaki bölüm, birçok MCDM yaklaşımının sonuçlarını üçgenlemenin, özelliklerin bireysel analizini optimize ettiğini iddia etmektedir.

MCDM, birçok alternatifi sıralamak için birden fazla kriter kullanan karar problemleriyle açıkça ilgilenen bir operasyonel araştırma alt disiplini (Jahan ve Ark., 2016), birbirini takip eden aşamalar olarak: alternatifleri tanımlama, kriterleri tanımlama, kriterlerin ağırlıklarını tanımlama, alternatifleri değerlendirmek için yöntemi seçme ve problem ifadesine karşı çözümleri doğrulama (Füülop, 2005). Aşağıda bu aşamaların açıklaması yer almaktadır.

3.5.1. Alternatifleri Tanımlama

Alternatifler, çalışma probleminde önemlerine göre sıralanacak metinsel özellikler arasında BOW, tf-idf , kelime 2-gram, karakter 2-gram, durma kelimelerini, sözlüksel, sözdizimsel ve anlamsal bulunmaktadır.

3.5.2. Kriterleri Tanımlama

Kriterler, bunlara göre alternatifin sıralanacağı ölçülerdir: çalışma süresi, doğruluk, duyarlılık, kesinlik, F1-skor.

3.5.3. Kriterlerin Ağırlıklarını Tanımlama

Her bir kriter, aşağıdaki tabloda gösterildiği gibi Saaty ölçeğinden (Saaty, 2008) elde edilen ağırlıklar olan ikili karşılaştırma matrisini oluşturmak için karar vermedeki etkisinden dolayı bir ağırlık belirler.

Çizelge 3-8. Saaty ölçeği

Anlamı	Ölçek
Eşit derecede önemli	1
biraz önemli	2
orta derecede önemli	3
Orta seviyenin üzerinde	4
Kesinlikle önemli	5
Güçlünün üstünde	6
Çok güçlü	7
Çok önemli	8
Son derece önemli	9

Kaur ve Ark. tarafından yapılan deneylerde (Kaur ve Ark., 2018), daha fazla güvenilirliği yansıtmak için çalışma zamanı kriteri eklenmiştir. Bu da Çizelge 3-9 'da gösterilen önem kriterleri tablosunun oluşturulmasına çok önemli ağırlıkla katkıda bulunmuştur.

Çizelge 3-9. Her bir kriterin diğer kriterlere göre önemi

	F1-skoru	Kesinlik	Duyarlılık	Doğruluk	Çalışma süresi
F1-skoru	1	2	3	4	5
Kesinlik	½	1	2	3	4
Duyarlılık	1/3	1/2	1	2	4
Doğruluk	1/4	1/3	1/2	1	5
Çalışma süresi	1/5	1/4	1/4	1/5	1

Kriter önem matrisinin normalleştirilmesi, Çizelge 3-10 'teki sonuç matrisini meydana getirmiştir.

Çizelge 3-10. Normalleştirilmiş Matris

	F1-skoru	Kesinlik	Duyarlılık	Doğruluk	Çalışma süresi
F1-skoru	0.43796	0.489795918	0.44444444	0.39215686	0.26316
Kesinlik	0.21898	0.244897959	0.2962963	0.29411765	0.21053
Duyarlılık	0.14599	0.12244898	0.14814815	0.19607843	0.21053
Doğruluk	0.10949	0.081632653	0.07407407	0.09803922	0.26316
Çalışma süresi	0.08759	0.06122449	0.03703704	0.01960784	0.05263

Aynı satırdaki diğer kriterlere göre normalleştirilmiş kriter önem ortalaması Çizelge 3-11'deki ağırlık matrisini vermektedir.

Çizelge 3-11. Ağırlık matrisi

F1-skoru	0.405502
Kesinlik	0.252963
Duyarlılık	0.164637
Doğruluk	0.125279
Çalışma süresi	0.051618

3.5.4. Alternatifleri Değerlendirmek İçin Yöntemi Seçme

Bu alt bölüm, metinsel özellikleri sıralamak için Analitik Ağ Süreci (ANP), VIKOR ve ELECTRE yöntemlerinin bazı deneylerini içerir.

Analitik Ağ Süreci (ANP) 1996 yılında Thomas L. Saaty tarafından geliştirilmiş ve birçok araştırmadan yararlanılmıştır. ANP yönteminin temel amacı, sınırlı sayıda karar kriteri açısından sınırlı sayıda alternatifin sıralamasını elde etmektir. İkili karşılaştırmaları kullanarak kriterler ve alternatifler arasında tercihler oluşturmanın temelidir. Bir karar verici için en uygun alternatif her zaman en üst sırada yer alan alternatiftir (Saaty, 2005). ANP'nin adım adım uygulaması aşağıdadır (Saaty, 2001).

Girdi:	Karar verme matrisi
Çıktı:	Alternatifler Sıralaması
Adım 1	Karar verme problemini belirleme
Adım 2	Kriterlerin birbirleriyle olan ilişkilerinin belirlenmesi
Adım 3	Kriterler arasında ikili karşılaştırmalar yapmak
Adım 4	Ağırlıksız süper matris oluşturma /* Kriterler arası değerlendirmeler sonucunda elde edilen değerler */
Adım 5	Ağırlıklı süper matrisin oluşturulması /* Ağırlıksız süper matrisin değerleri, içindeki ilgili kümelerin ağırlıklı değerleri ile çarpılarak oluşturulur */
Adım 6	En iyi alternatifin belirlenmesi ve seçilmesi /* Alternatifler, ideal değere yakınlıklarından dolayı sıralanmıştır. */

Algoritma 3-3. ANP için sözde kod

Çizelge 3-12. ANP yöntemine dayalı metinsel özellikler sıralamasını göstermektedir.

Alternatifler	Sıralama
BOW	2
tf-idf	3
kelime 2-gram	5
karakter 2-gram	4
Durma kelimelerini	8
Sözlüksel	6
Sözdizimsel	7
Anlamsal	1

VlseKriterijumsa Optimizacija I Kompromisno Resenje (VIKOR) 1998 yılında Serafim Opricovic tarafından çok karmaşık karar problemlerini çözmek için önerilmiştir, birçok deney yoluyla başarılı bir uygulama elde edilmiştir. VIKOR'un adım uygulaması aşağıdadır (Brestovac ve Grgurina, 2013).

Girdi: Karar verme matrisi
Çıktı: Alternatifler Sıralaması

Adım 1 Her bir kriter için “Çıkarılan metinsel özelliklerin performansı karşılaştırma karşılaştırması” tablosundaki maksimum ve min değerlerini seçme

$$f_i^* = \max_j f_{ij} \quad f_i^- = \min_j f_{ij} \quad \text{Denklem 3-5}$$

Adım 2 S_j ve R_j , formüller kullanılarak her alternatif için hesaplanır.

$$S_j = \sum_{i=1}^n \omega_i \cdot \frac{(f_i^* - f_{ij})}{(f_i^* - f_i^-)} R_j = \max_i [\omega_i \cdot \frac{(f_i^* - f_{ij})}{(f_i^* - f_i^-)}]$$

where $j = 1, \dots, m$
and $i = 1, \dots, n$

Denklem 3-6

ω_i : Ölçütlerin göreceli ağırlığını ifade edin

Adım 3 Her alternatif için Q_j değerleri şu şekilde hesaplanmıştır:

$$Q_j = v \frac{(S_j - S^*)}{(S^- - S^*)} + (1 - v) \frac{(R_j - R^*)}{(R^- - R^*)} \quad \text{Denklem 3-7}$$

$$R^- = \max_j R_j \quad R^* = \min_j R_j \quad S^- = \max_j S_j \quad S^* = \min_j S_j \quad \text{Denklem 3-8}$$

where $j = 1, \dots, m$ and $i = 1, \dots, n$

$v = 0.5$ kullanılır

Adım 4 Q_j değerleri listelenir. Alternatifler, Q_j ideal değerine yakın olmaları nedeniyle sıralanmıştır.

Algoritma 3-4. VIKOR için sözde kod

Çizelge 3-13, VIKOR yöntemine yönelik metinsel anlatıları içermektedir.

Çizelge 3-13. VIKOR yöntemi için sıralama tablosu

Alternatifler	Sıralama
BOW	2
tf-idf	3
kelime 2-gram	4
karakter 2-gram	7
Durma kelimelerini	8
Sözlüksel	5
Sözdizimsel	6
Anlamsal	1

ELECTRE yöntemi, Çizelge 3-9'ten en güçlü özellikleri seçmek için kullanılır. ELECTRE, çok kriterli karar problemlerinde birçok alternatifi sıralamak için MCDM yaklaşımıdır (Jahan ve Ark., 2016). ELECTRE yöntemi, gerçeği yansıtan eleme ve seçim anlamına gelir (Mesran ve Ark., 2017). Karmaşık karar verme problemlerini birçok alternatif ve az kriterle çözmek için bir felsefe olarak geliştirilmiştir (Yanie ve Ark., 2018). Uygun kriterlere göre alternatifler arasında ikili üstünlük karşılaştırmalarına dayanır (Botti ve Peypoch, 2013). ELECTRE'nin adım uygulaması aşağıdaki şekilde gösterilmiştir (Yücel ve Görener, 2016):

Girdi: Karar verme matrisi
Çıktı: Alternatifler sıralaması

Adım 1 Ağırlıklı normalleştirilmiş karar matrisinin hesaplanması
/*

Çizelge 3-11, $W = (w_1; w_2; \dots; w_n)$ kriterlerin ağırlık vektörlerini içerir,

where $w_j > 0, \sum_{j=1}^n w_j = 1$ Denklem 3-9

Çizelge 3-10'te elde edilen normalize karar matrisinin kriter ağırlıkları ile çarpılmasıyla hesaplanan ağırlıklı normalize karar matrisi
*/

Adım 2 Uyumsuzluk ve uyum kümesini belirleyin

/*

Ağırlıklı normalleştirilmiş karar matrisindeki her çiftin öğeleri karşılaştırılır ve uyum kümesi, aşağıdaki ilişki ile belirlenen her alternatif çiftinin en iyi veya eşit öğelerini içerecektir:

$$C(p; q) = \{j, v_{pj} \geq v_{qj}\} \quad \text{Denklem 3-10}$$

Uyumsuzluk kümesi, aşağıdaki ilişki ile belirlenen her alternatif çiftinin en kötü öğelerini içerecektir:

$$D(p; q) = \{j, v_{pj} < v_{qj}\} \quad \text{Denklem 3-11}$$

*/

Adım 3 Uyum ve uyumsuzluk matrisinin hesaplanması

/*

Uyum matrisi, uyum kümesinin eleman ağırlıkları toplanarak hesaplanır

$$C_{pq} = \sum_{j \in C(p; q)} w_j \quad \text{Denklem 3-12}$$

Uyumsuzluk matrisi, uyumsuzluk kümesi öğelerinin toplam farkının, ölçüt öğelerinin toplam farkına bölünmesiyle hesaplanır

$$D_{pq} = \sum_{j \in D(p; q)} |v_{pj} - v_{qj}| / \sum_{j \in C(p; q)} |v_{pj} - v_{qj}| \quad \text{Denklem 3-13}$$

*/

Adım 4 Avantajların hesaplanması

/*

Uyumsuzluk ve uyum ortalamaları hesaplanır. Uyum indeks matrisini elde eden uyum ortalamasından büyük veya eşit olan uyum matrisindeki değerler için evet denmesi, uyumsuzluk matrisindeki değerler için uyum ortalamasından küçük veya eşit değerler için 'hayır' derken uyumsuzluk indeksi matrisini elde edin

*/

Adım 5 Net üstün ve düşük değerlerin hesaplanması

/*

Alternatifler net üstün ve alt değerlerine göre sıralanır. Bunları hesaplamak için aşağıdaki formüller kullanılır

$$C_p = \sum_{k=1}^m C_{pk} - \sum_{k \neq p}^m C_{kp} \quad \text{Denklem 3-14}$$

$$D_p = \sum_{k=1}^m D_{pk} - \sum_{k \neq p}^m D_{kp} \quad \text{Denklem 3-15}$$

*/

Adım 6 ELECTRE yönteminde metinsel özelliklerin üstün sıralamasını gösterir

Algoritma 3-5. ELECTRE için sözde kod

Çizelge 3-14, ELECTRE yöntemine dayalı metinsel özellikler sıralamasını göstermektedir.

Çizelge 3-14. ELECTRE yöntemi için sıralama tablosu

Alternatifler	Sıralama
BOW	5
tf-idf	4
kelime 2-gram	2
karakter 2-gram	7
Durma kelimelerini	8
Sözlüksel	3
Sözdizimsel	6
Anlamsal	1

3.5.5. Problem İfadesine Karşı Çözümleri Doğrulama

Üçgenleme ANP, VIKOR ve ELECTRE sıralama sonuçları, üzerinde çalışılan metinsel özelliklerden bu özelliğin vazgeçilebilir kararını destekleyen durma kelimelerini en kötü idealliği üzerinde anlaşmayı göstermektedir; ayrıca, diğerleri tweet'in yazarlığını doğrulamak için önceliklerine göre yakın sıralama bazında elde ederken anlamsal özellik en ideal olanıydı, özelliklere ağırlıklarına göre ağırlık atama.

3.6. Özelliklere Ağırlıklarına Göre Ağırlık Atama

Özellik seçimi, tahminde kullanılan öznitelikleri belirlerken, özellik ağırlıklandırma, seçilen özelliklere önceliklerine göre farklı ağırlıklar atayarak önemini göstermektedir (Hall, 1999). Kullanılan özellik ağırlıklandırması (Alazab ve Ark., 2014; Savyan ve Bhanu, 2020).

Özellik ağırlıklandırmanın yararlı olduğu bazı durumlar olsa da araştırmacılar arasında hala tartışma konusu olduğunun bilinmesi önemlidir (Hall, 1999). Örneğin Alazab, özelliklere çıktıya katkılarına göre ağırlık verilmesine dayanan bir yaklaşım geliştirmiştir (Alazab ve Ark., 2013). Bu çalışmada, sıra toplam ağırlık yöntemi (Roszkowska, 2013) formülü kullanılarak her bir özelliğe sıralarına göre bir ağırlık atanmıştır:

$$w_j(RS) = (n - r_j + 1) / \sum_{k=1}^n (n - r_k + 1) = 2(n+1-r_j) / (n(n+1)) \quad \text{Denklem 3-16}$$

r_j özellik sıralaması olduğunda, n özellik sayısı, $j = 1; 2; \dots; n$.

Ve ‘‘Sıra Üs ağırlığı’’ yöntemi (Roszkowska, 2013):

$$w_j(RS) = (n - r_j + 1)^P / \sum_{k=1}^n (n - r_k + 1)^P \quad \text{Denklem 3-17}$$

r_j özellik sıralaması, n özellik sayısı, P en önemli kriter ağırlığı belirtmektedir.



4. SONUÇ VE ÖNERİLER

4.1. Deneyler ve Tartışma

Veri kümesini oluşturmak için öncelikle yirmi farklı Twitter kullanıcısına ait 16124 tweet taranmış, ardından her bir tweet temizleme işlemine tabi tutulmuştur. Daha sonra, temizleme işlemini sonucu elde edilen veri kümesindeki tweet'lerin özellikleri ilgili temizlenmiş tweet'lerden çıkarılmıştır. Bu işlemler Python 3 programlama dili ortamında gerçekleştirildi. Tüm deneyler 16 GB RAM'e sahip bir bilgisayar kullanılarak yürütülmüştür. K-katlama, her bir kat yinelemesinde eğitim ve test kümelerini indekslemek için bir çapraz doğrulama türü olarak kullanılmıştır.

ANP, VIKOR ve ELECTRE sıralama sonuçlarının üçgenlenmesinden ve çalışılan metinsel özelliklerden kelimeleri durdur özelliğinin çıkarılmasından sonra; "Sıra Toplamı Ağırlık" yöntemi kullanılarak her bir özelliğe sıralarına göre bir ağırlık atanmıştır.

Öznitelikler için ağırlıklı sıralanmış değerler, sınıflandırıcıları eğitmek için birleştirilmiştir. Çizelge 4-1, tweet'lerin yazarlık doğrulamasına ilişkin tahminleri göstermektedir.

Çizelge 4-1. Her bir algoritma için durma kelimelerini özelliği olmadan ve her bir algoritma için stop kelimeleri özelliği olmadan ve “Sıra Toplamı Ağırlık” yöntemi kullanılarak ağırlıklandırma sonrası tahmin ölçüleri.

	Çalışma süresi	Doğruluk	Duyarlılık	Kesinlik	F1-skoru
Lojistik regresyon	0.41 Min	0.909	0.967	0.907	0.943
SVM	0.42 Min	0.897	0.94	0.896	0.934
Random forest	0.63 Min	0.796	0.977	0.777	0.881
XGBoost	29.22 Min	0.887	0.955	0.883	0.929

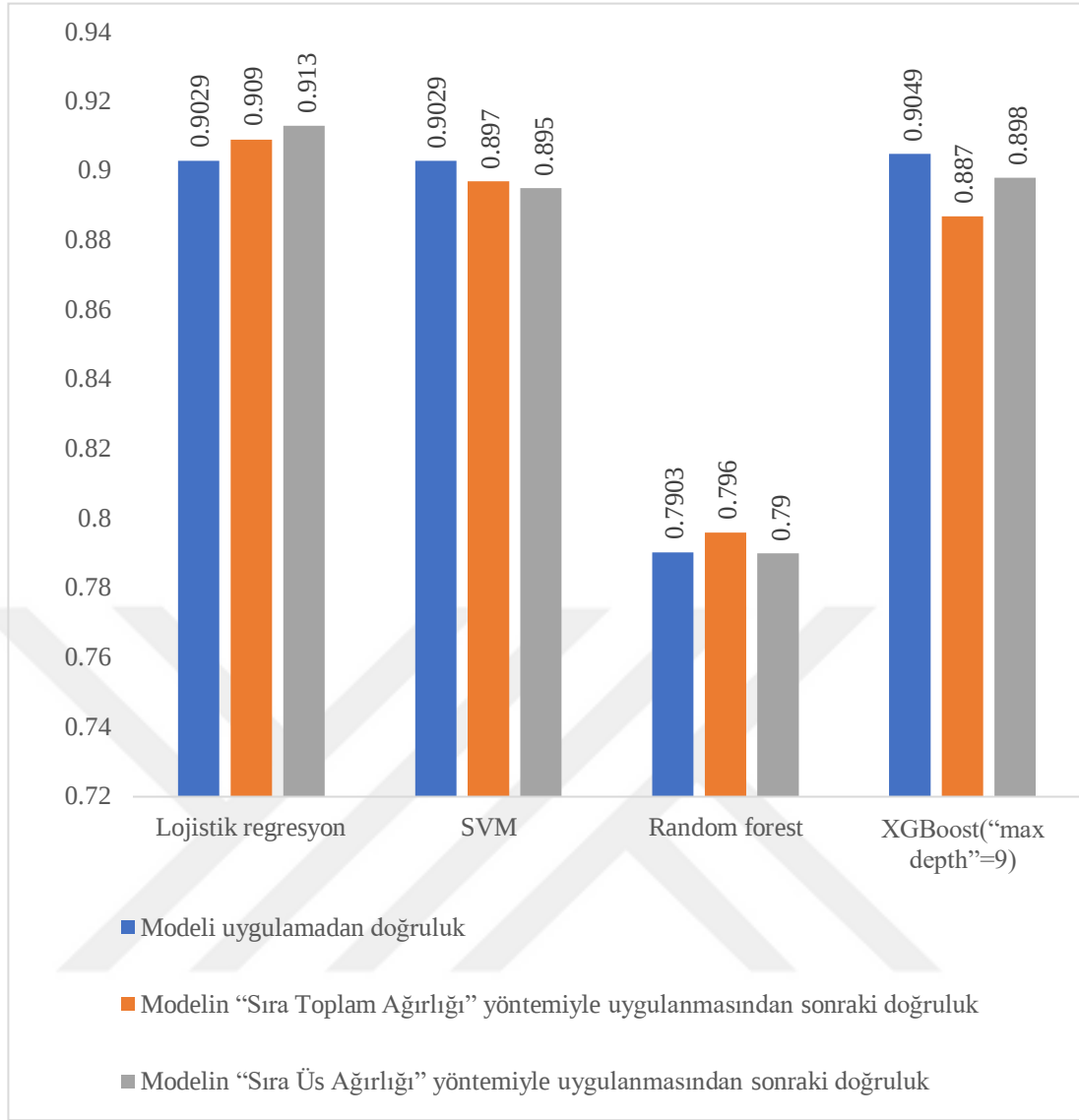
Çizelge 4-2, lojistik regresyon, SVM, Random Forest ve XGBoost algoritmalarının, çalışma süresi, doğruluk, duyarlılık, kesinlik ve F1-skoru ölçümlerine dayalı olarak tweet'lerin yazarlığını doğrulama konusundaki performansını göstermektedir. MCDM yaklaşımları ve “Sıra Üs Ağırlığı” yöntemi kullanılarak özellikleri seçimi ve ağırlıklandırmadan sonra sınıflandırıcıların karşılaştırmalı analizi, lojistik regresyon algoritmasının performansının diğer ML sınıflandırıcılarından daha iyi performans

gösterdiği gerçeğini vurgulamaktadır. Bu nedenle, modelimizde sınıflandırma adımında lojistik regresyon kullanılması, tweet'in yazarlığını doğrulamasında daha etkin güvenilir bir yaklaşımdır.

Çizelge 4-2. Her bir algoritma için stop kelimeleri özelliği olmadan ve “Sıra Üs ağırlığı” yöntemi kullanılarak ağırlıklandırma sonrası tahmin ölçüleri.

	Çalışma süresi	Doğruluk	Duyarlılık	Keskinlik	F1-skoru
Lojistik regresyon	0.43 Min	0.913	0.97	0.913	0.945
SVM	0.38 Min	0.895	0.94	0.895	0.933
Random forest	0.57Min	0.79	0.97	0.794	0.88
XGBoost	3.39 Min	0.898	0.97	0.897	0.936

Araştırmada deneyler için dört farklı sınıflandırıcı kullanılırken, lojistik regresyon sınıflandırıcı diğerlerinden daha iyi performans göstermiştir. Şekil 4-1, önerilen model aracılığıyla “Sıra Üs Ağırlığı” ve “Sıra Toplam Ağırlığı” yöntemlerinin uygulanmasından önce ve sonra bu sınıflandırıcılar arasındaki Tweet yazarlık doğrulamasının doğruluğunu karşılaştırmaktadır. Önerilen model yalnızca doğruluğu %90.29'dan %91.3'e çıkarmakla kalmamakta, aynı zamanda yazarlığın doğrulanmasında yanlış oranı da azaltmaktadır.

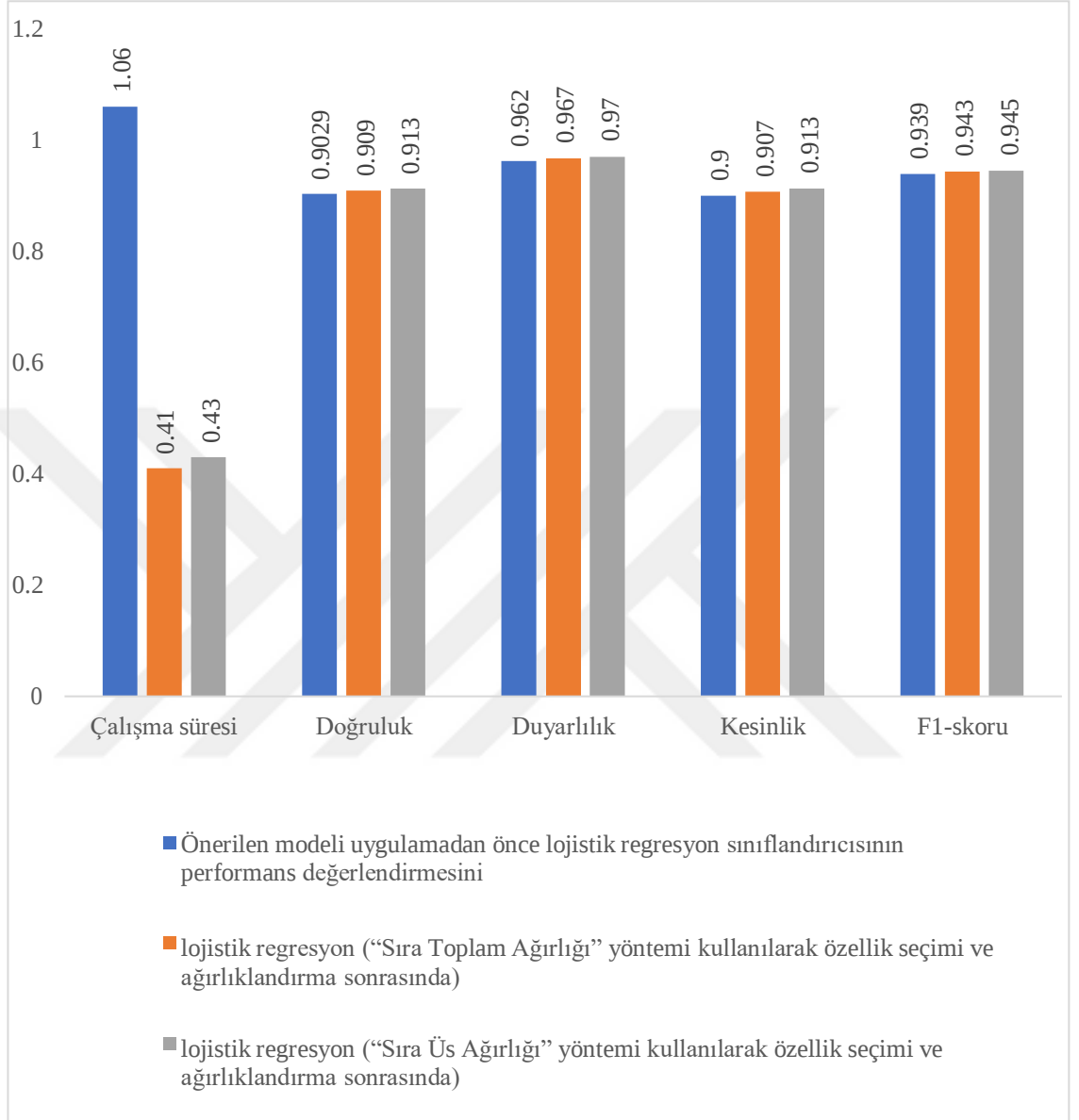


Şekil 4-1. Önerilen modeli uygulamadan önce ve sonra sınıflandırıcıların karşılaştırılması

Şekil 5, önerilen modeli uygulamadan önce lojistik regresyon sınıflandırıcısının performans değerlendirmesini, “Sıra Toplam ağırlığı” yöntemi kullanılarak özellik seçimi ve ağırlıklandırma sonrasında ve “Sıra Üs ağırlığı” yöntemi kullanılarak özellik seçimi ve ağırlıklandırma sonrasında göstermektedir.

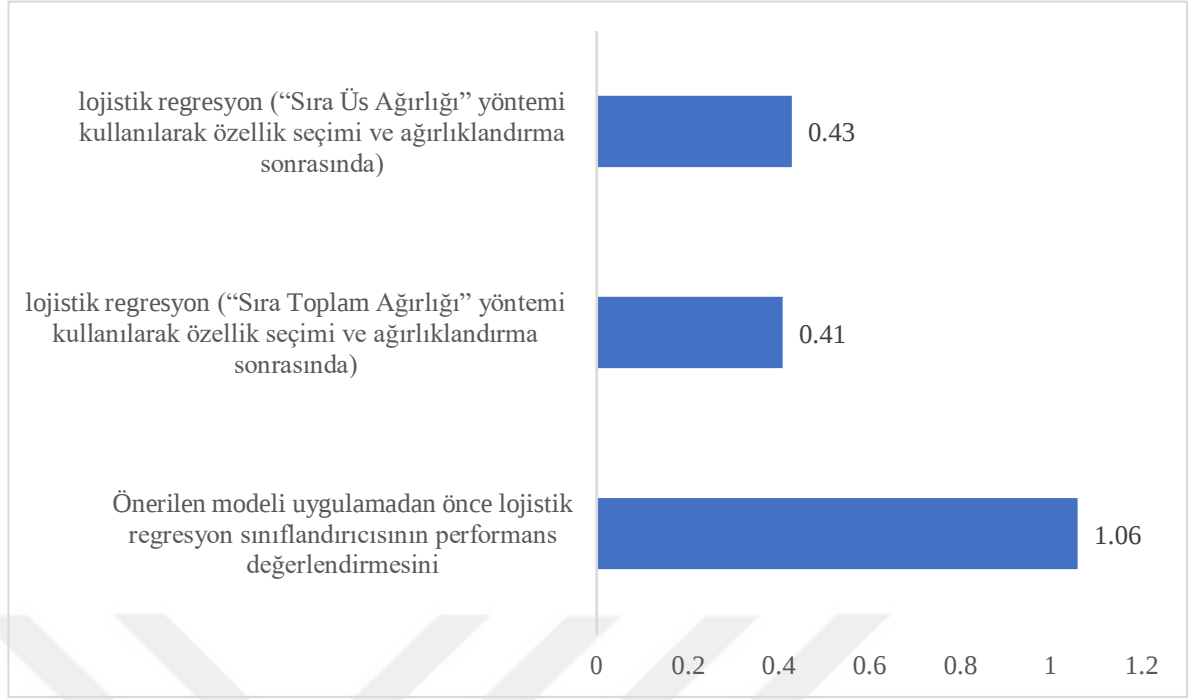
Özellik seçimi ve ağırlıklandırmayı içeren önerilen modelin “Sıra Üs Ağırlığı” yöntemi kullanılarak uygulanmasının, model uygulanmadan lojistik regresyona göre ve “Sıra Toplam Ağırlığı” yöntemi kullanılarak özellik seçimi ve ağırlıklandırması yapılan modelden daha iyi performans sağladığı şekilden anlaşılmaktadır. Diğerlerine göre art arda daha yüksek doğruluk, duyarlılık, kesinlik ve F1-skoru

değeri ile önerilen modelin üstünlüğünü doğrulamaya yardımcı olurken, önerilen model doğruluğu %90,29'dan %91,3'e yükseltir, bu da daha fazla tweet'in doğru tanındığı anlamına gelmektedir.



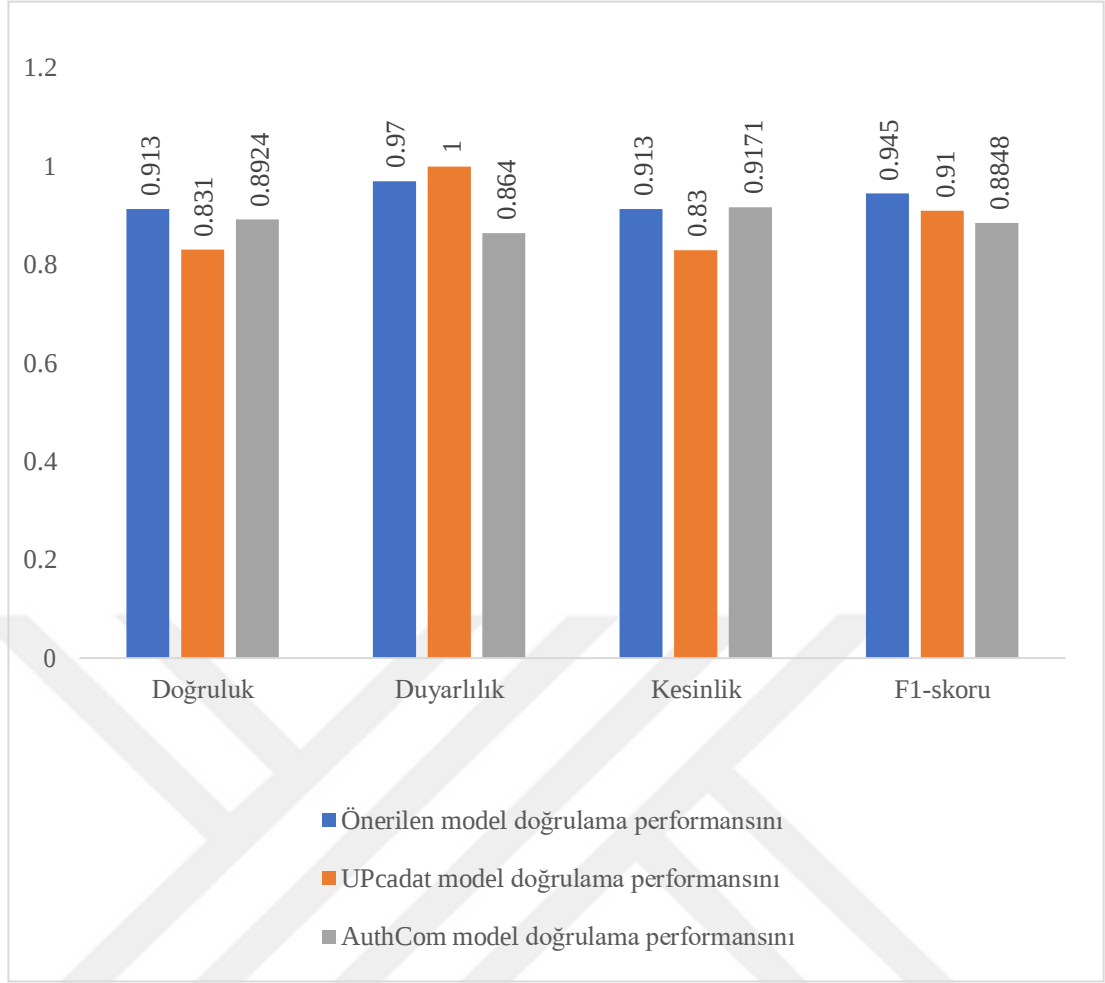
Şekil 4-2. Lojistik regresyon sınıflandırıcısının performans değerlendirmesi

Şekil 4-3. Önerilen modelde lojistik regresyon sınıflandırıcısı için eğitim süresi karşılaştırması, tweet'in yazarlığını doğrulamak için lojistik regresyon sınıflandırıcısının eğitim süresini azaltan MCDM yaklaşımlarını ve “Sıra Üs ağırlığı” yöntemini kullanan özellik seçimi uygulamasını göstermektedir.



Şekil 4-3. Önerilen modelde lojistik regresyon sınıflandırıcısı için eğitim süresi karşılaştırması

Şekil 4-4, bu çalışmada önerilen model ile (Savyan ve Bhanu, 2020)'deki UbCadet modeli ve (Kaur ve Ark., 2018)'deki AuthCom modeli arasındaki tweet'in yazarlığını doğrulama performansını karşılaştırmaktadır. Bu çalışmada önerilen modelin UbCadet ve AuthCom modellerinden daha yüksek doğruluk sağladığı gözlemlenmiştir. Ancak UbCadet modelinin duyarlılığı, önerilen modelde %97, AuthCom modelinde ise %86,4 ile karşılaştırıldığında en yüksek değeri 1 almıştır. UbCadet tweetlerin %83,1'ini doğru olarak tanırken, AuthCom %89,24'ü doğru olarak tanır ve önerilen model %91,3'ünü doğru olarak tanır; Tweetlerin yazarlık doğrulaması konusunda karar vermede önerilen modelin güvenilirliğini yansıtmaktadır. Şekil 4-4, bu çalışma boyunca önerilen çerçevenin, literatürdeki ilgili mevcut çalışmaları (Savyan ve Bhanu, 2020; Kaur ve Ark, 2018) karşılaştırarak güvenliği ihlal edilmiş hesapları doğrulamada daha doğru olduğunu göstermektedir.



Şekil 4-4. Önerilen model performansının ilgili çalışmalarla karşılaştırılması

Önerilen model ile ilgili aşağıdaki sonuçları vurgulamaktadır:

- “Sıra Üs Ağırlık” yöntemine ek olarak MCDM yaklaşımlarının uygulanması, lojistik regresyon algoritmasının performansını artırmakta, ayrıca Şekil 4-1’te gösterildiği gibi SVM, XGBoost ve random forest algoritmaları için performans ölçütlerini biraz bozmaktadır.
- Model uygulaması, makine öğrenimi sınıflandırıcıları için eğitim süresini kısaltır.
- “Sıra Toplamı Ağırlık” yöntemine ek olarak MCDM yaklaşımlarının uygulanması, en önemli özellikleri vurgulamak, lojistik regresyonun performansını, Şekil 4-1’te gösterildiği gibi önerilen modeli uygulamadan önce biraz arttırmaktır.

MCDM yaklaşımları belirli ML sınıflandırıcıları için yüksek performans oranı sağlarken, karar verici tarafından kriterlerin önemini belirlemek için atanan ağırlığın

karar verici tarafından öznel girdilere dayanan çalışmamızın sınırlılığı olarak kabul edilmektedir (Botti ve Peypoch, 2013).

(Viswanath ve Ark., 2014)'de bahsedilen teşvikli gizli anlaşma ağları, kullanıcıların hesaplarında promosyon mesajları yayınlamak için ödeme aldığı, Kişisel hesaplar aracılığıyla yayınlanan toplu uygulamaların yanı sıra, yazarın kişiliğini belirten metindeki metinsel özellikleri değiştiren yazarlık doğrulama sürecini engelledikleri için bir sınırlama olarak kabul edilir.

Araştırma, yalnızca Eylül 2017'den sonra tweetlenen mesajlar için geçerli olan kişisel hesaplardan Twitter beslemelerine dayalı olarak 280 karakterlik metin bloklarından oluşan kısa mesajlar kullanılarak değerlendirildi; ancak Twitter'da yazarlık doğrulaması için genel bir veri kümesi yoktur ve Twitter her hesap için 3200'e kadar tweet alınmasına izin verir. Bu koşullara göre veri kümesi toplanmış olup, sık sık tweet atan kullanıcılardan; ayrıca bu koşullar tüm örnek hesaplar için mevcut değildi. Özellikle çoğu tweet, veri temizleme işlemi sırasında silinen retweetler, yeniden paylaşım bağlantıları veya emojilerdir.

4.2. Sonuç

Bu çalışma, ÇSA'larda yazarlık doğrulaması için hibrit bir model önermektedir. Yazarlık doğrulama amacıyla standartlaştırılmış Twitter veri kümesinin herkese açık olmaması nedeniyle, aynı tür ve konu Twitter tabanlı bir veri kümesi taranmış ve yazarlık bilgileriyle manuel olarak açıklama eklenmiştir. Veri kümesini özellik çıkarımına hazırlamak için ardışık ön işleme adımları uygulanmıştır.

Bu nedenle dört katmanlı Yazarlık doğrulama yaklaşımı başlatılmıştır. Başlangıçta, XGBoost algoritması, toplanan, açıklamalı ve önceden işlenmiş tweet veri kümesinden çıkarılan her bir kriterde tweet yazarlığının doğrulanmasında metinsel özelliğin performansını hesaplamak için bir ön işlemci olarak seçilmiştir.

Özellik seçimi, bir MCDM problemi olarak ele alınarak araştırılmıştır, bu nedenle üç MCDM yaklaşımının (ANP, VIKOR ve ELECTRE) sonuçları, özelliklerin önceliklerini sıralamak için üçgenleştirilmiştir; üç yaklaşım, tweet'in yazarlığını doğrulamada, durma kelimelerini önemsiz, önemli ve anlamsal özelliğin diğer özelliklerle karşılaştırıldığında üstünlüğü konusunda hem fikirdi; incelenen metinsel

özelliklerden bu özelliğin vazgeçilebilir kararını desteklemektedir. Metodolojik düzeyde, benzer problemler için faydalı olması beklenen, bu alana üç MCDM yaklaşımının (ANP, VIKOR ve ELECTRE) sonuçlarının üçgenleştirilmesinin ilk uygulamasıdır. Kaur ve Ark. tarafından (Kaur ve Ark., 2018) çoğu kriter ve bunların nispi ağırlıkları elde edilmiştir. MCDM, birçok alternatifi sıralamak için birden fazla kriter kullanan karar problemleriyle açıkça ilgilenen bir operasyonel araştırma alt disiplini, MCDM, tweet'lerden çıkarılan sekiz metinsel özelliği sıralamak için ikili karşılaştırmaları kullanmaktadır.

Ayrıca, seçilen önceliklerin ağırlıklandırılması için Sıra Üs ve Sıra Toplamı Ağırlık yöntemleri kullanılmıştır. Azaltılmış veri kümesi, dört makine öğrenimi sınıflandırıcısı ile değerlendirilmiştir; bunlardan ikisi, çalışma süresi, doğruluk, duyarlılık, kesinlik ve F1-skoru açısından geleneksel makine öğrenimi ile karşılaştırıldığında artan performansı göstermiştir. Lojistik regresyon sınıflandırıcı ile önerilen model için deney analizi, tweet'in yazarlığını doğrulamada 280 karakterlik blok boyutları için %94,5'e ulaşan yüksek bir F1-skoru sonucunu yansıtmaktadır. Bu, gelecekteki çalışma olarak yüksek özellik sayısına sahip başka bir sınıflandırma probleminde model uygulamasını genişletebilir.

5. KAYNAKLAR

1. Statista (2020). Number of social network users worldwide from 2017 to 2025. <https://www.statista.com/statistics/278414/number-of-worldwide-socialnetwork-users/> İndirilme Tarihi: 22.05.2020.
2. Worldometers (2020). World population projections. <https://www.worldometers.info/world-population/world-population-projections/> İndirilme Tarihi: 15.06.2020.
3. live stats (2020). Twitter usage statistics. <https://www.internetlivestats.com/twitter-statistics/> İndirilme Tarihi: 13.02.2020.
4. Singh, Y. ve Banerjee, S. (2019). Fake (sybil) account detection using machine learning. In Proceedings of International Conference on Advancements in Computing & Management (ICACM) , USA, October.
5. Trång, D., Johansson, F., & Rosell, M. (2015). Evaluating algorithms for detection of compromised social media user accounts. In 2015 Second European Network Intelligence Conference (pp. 75-82). IEEE, USA, September.
6. Calabresi, M. (2017, May 18). Inside russia's social media war on america. Retrieved from times Blog <https://time.com/4783932/inside-russia-social-media-war-america/>
7. Lee, E. (2013, May 26). Associated press twitter account hacked in marketmoving attack. Retrieved from Bloomberg Technology. <https://www.bloomberg.com/middleeast/>
8. Stamatatos, E., A survey of modern authorship attribution methods. Journal of the American Society for information Science and Technology 60(3), 538-556,2009.
9. Boenninghoff, B., Hessler, S., Kolossa, D., & Nickel, R. M. (2019). Explainable authorship verification in social media via attention-based similarity learning. In 2019 IEEE International Conference on Big Data (Big Data) (pp. 36-45). IEEE, USA, December.
10. Rocha, A., Scheirer, W. J., Forstall, C. W., Cavalcante, T., Theophilo, A., Shen, B., & Stamatatos, E. (2016). Authorship attribution for social media forensics. Journal of IEEE transactions on information forensics and security, 12(1), 5-33.
11. Al-Khatib, M. A., & Al-qaoud, J. K. (2020). Authorship verification of opinion articles in online newspapers using the idiolect of author: a comparative study. Journal of Information, Communication & Society, 1-19.
12. Suman, C., Saha, S., Bhattacharyya, P., & Chaudhari, R. S. (2021). Emoji helps! a multi-modal siamese architecture for tweet user verification. Journal of Cognitive Computation, 13(2), 261-276.
13. Benevenuto, F., Magno, G., Rodrigues, T., & Almeida, V. (2010). Detecting spammers on twitter. In Collaboration, electronic messaging, anti-abuse and spam conference (CEAS). July.
14. Gong, N. Z., Frank, M., & Mittal, P. (2014). Sybilbelief: A semi-supervised learning approach for structure-based sybil detection. IEEE Transactions on Information Forensics and Security, 9(6), 976-987.

15. Gao, P., Wang, B., Gong, N. Z., Kulkarni, S. R., Thomas, K., & Mittal, P. (2018). Sybilfuse: Combining local attributes with global structure to perform robust sybil detection. In 2018 IEEE conference on communications and network security (CNS) (pp. 1-9). IEEE, USA, May.
16. Singh, M., Bansal, D., & Sofat, S. (2018). Who is who on twitter—spammer, fake or compromised account? a tool to reveal true identity in real-time. *Cybernetics and Systems*, 49(1), 1-25.
17. VanDam, C., Tang, J., & Tan, P. N. (2017). Understanding compromised accounts on twitter. In *Proceedings of the International Conference on Web Intelligence*, August.
18. Barbon, S., Igawa, R. A., & Zarpelao, B. B. (2017). Authorship verification applied to detection of compromised accounts on online social networks. *Multimedia Tools and Applications*, 76(3), 3213-3233.
19. Zangerle, E., & Specht, G. (2014). " Sorry, I was hacked" a classification of compromised twitter accounts. In *Proceedings of the 29th annual acm symposium on applied computing*, March.
20. Okerefor, K., & Adelaiye, O. (2020). Randomized cyber attack simulation model: a cybersecurity mitigation proposal for post covid-19 digital era. *International Journal of Recent Engineering Research and Development (IJRERD)*, 5(07), 61-72.
21. Nauta, M. (2016). Detecting hacked twitter accounts by examining behavioural change using twttr metadata. In *Proceedings of the 25th Twente Student Conference on IT*, Hollanda, November.
22. Seyler, D., Li, L., & Zhai, C. (2018). Identifying compromised accounts on social media using statistical text analysis. *arXiv e-prints*, arXiv-1804.
23. PV, S., & Bhanu, S. (2020). UbCadet: detection of compromised accounts in twitter based on user behavioural profiling. *Multimedia Tools & Applications*, pp. 1-37.
24. Viswanath, B., Bashir, M. A., Crovella, M., Guha, S., Gummadi, K. P., Krishnamurthy, B., & Mislove, A. (2014). Towards detecting anomalous user behavior in online social networks. In *23rd {USENIX} Security Symposium ({USENIX} Security 14)*, USA, 20-24 August.
25. Egele, M., Stringhini, G., Kruegel, C., & Vigna, G. (2013). Detecting compromised accounts on social networks. *NDSS, The Internet Society journal*, University of California. Retrieved from https://sites.cs.ucsb.edu/~vigna/publications/2013_NDSS_compa.pdf
26. Li, J. S., Monaco, J. V., Chen, L. C., & Tappert, C. C. (2014). Authorship authentication using short messages from social networking sites. In *2014 IEEE 11th International Conference on e-Business Engineering* (pp. 314-319). IEEE, USA, November.
27. F. Lagerholm,. (2017). Using artificial intelligence to verify authorship of anonymous social media posts, Independent thesis Basic level, Mälardalen University, Sweden.
28. Brocardo, M. L., Traore, I., Woungang, I., & Obaidat, M. S. (2017). Authorship verification using deep belief network systems. *International Journal of Communication Systems*, 30(12), e3259.
29. Kaur, R., Singh, S., & Kumar, H. (2018). AuthCom: Authorship verification and compromised account detection in online social networks using AHP-

- TOPSIS embedded profiling based technique. *Expert Systems with Applications*, 113, 397-414.
30. Robbins, I. P. (2012). Writings on the Wall: The Need for an Authorship-Centric Approach to the Authentication of Social-Networking Evidence. *Jurnal of Minn. JL Sci. & Tech.*, 13, 1.
 31. Gadekallu, T. R., Khare, N., Bhattacharya, S., Singh, S., Maddikunta, P. K. R., Ra, I. H., & Alazab, M. (2020). Early detection of diabetic retinopathy using PCA-firefly based deep learning model. *Journal of Electronics*, 9(2), 274.
 32. Tekkalmaz, M, Dedeoğlu, Y. (2020, May 22). Authorship Attribution CS533 – Information Retrieval Systems. Retrieved from slideplayer <https://slideplayer.com/slide/3980244/>
 33. Rao, O. Srinivasa, NV Ganapathi Raju, Y. Srilalitha, and Mrs P. Bharathi. (2017). Authorship Attribution using Unsupervised Clustering Algorithms on English C50 News. *international Advanced Research Journal in Science, Engineering and Technology*. ISO 3297:2007 Certified Vol. 4, Issue 7.
 34. Koppel, M., & Schler, J. (2004). Authorship verification as a one-class classification problem. In *Proceedings of the twenty-first international conference on Machine learning* (p. 62), USA, 4-8 July.
 35. Brocardo, M. L., Traore, I., Saad, S., & Woungang, I. (2013). Authorship verification for short messages using stylometry. In *2013 International Conference on Computer, Information and Telecommunication Systems (CITS)* (pp. 1-6). IEEE, USA, 1-6 May.
 36. N. Potha.(2019). Authorship Verification. Doktora Tezi. AEGEAN Üniversitesi, Greece.
 37. Bai, X & Yao, Y.(2010). Facebook on Campus: The Use and Friend Formation in Online Social Networks. *SSRN Electronic Journal*. 10.2139/ssrn.1535141.
 38. Chou, A. Y., & Chou, D. C. (2009). Information system characteristics and social network software adoption. In *Proceedings of the SWDSI conference* (pp. 335-343), Italy, March.
 39. Costa, J., Silva, C., Antunes, M., & Ribeiro, B. (2013). Defining semantic meta-hashtags for twitter classification. In *International Conference on Adaptive and Natural Computing Algorithms* (pp. 226-235). Springer, Berlin, Heidelberg, 4-6 April.
 40. phys.org (2021). Twitter to double tweet limit to 280 characters (update). <https://phys.org/news/2017-11-twitter-character-limit.html> İndirilme Tarihi: 12.05.2021.
 41. Espinhara, J., & Albuquerque, U. (2013). Using online activity as digital fingerprints to create a better spear phisher. Technical report, Trustwave SpiderLabs, Rapor No: HITB, 2013.
 42. Murphy, K. P. (2012). *Machine learning: a probabilistic perspective*. London: MIT press.
 43. Twitter (2021). About Twitter's APIs. <https://help.twitter.com/en/rules-and-policies/twitter-api> İndirilme Tarihi: 12.05.2021.
 44. Steinert-Threlkeld, Z. (2018). *Twitter as data*. London: Cambridge University Press.
 45. Developer.twitter.com (2021). reviewing applications. <https://developer.twitter.com/en/review> İndirilme Tarihi: 12.05.2021.

46. Netlytic.org (2021). Data Collection. https://netlytic.org/home/?page_id=10834 İndirilme Tarihi: 12.05.2021.
47. Twitter.com/greptweet (2021). <https://twitter.com/greptweet?lang=en>. İndirilme Tarihi: 12.05.2021.
48. Borra, E., & Rieder, B. (2014). Programmed method: Developing a toolset for capturing and analyzing tweets. *Aslib journal of information management*, Vol. 66 No. 3, pp. 262-278.
49. Developer.twitter.com (2018). Pricing policy. <https://developer.twitter.com/en/pricing.html> İndirilme Tarihi: 12.05.2018.
50. Cording, P. H., & Lyngby, K. (2011). Algorithms for web scraping. Lyngby: Technical University of Denmark
51. Twitter: Term of Services (2021). <https://twitter.com/en/tos> İndirilme Tarihi: 12.05.2021.
52. lojban.org (2021). What is lojban?. https://mw.lojban.org/papri/Welcome!/en#Learning_Lojban İndirilme Tarihi: 12.05.2021.
53. Anjali, M. K., & Babu, A. P., Ambiguities in natural language processing. *International Journal of Innovative Research in Computer and Communication Engineering*, 2(5), 392-394. 2014.
54. K. Maria,. (2016). Authorship Attribution Forensics: Feature selection methods in authorship identification using a small e-mail dataset. yüksek lisans Tezi. Technoglossia University, Athens, Greece.
55. Schütze, H., Manning, C. D., & Raghavan, P. (2008). Introduction to information retrieval. Cambridge:Cambridge University Press.
56. Juola, P. (2008). Authorship attribution. Boston:Now publishers inc.
57. A. Castro, B.Lindauer. (2012).Author identification on twitter. Project. Stanford university, California.
58. Sultan, K., Ali, H., & Zhang, Z. (2018). Big data perspective and challenges in next generation networks. *Future Internet journal*, 10(7), 56.
59. Andrew, Ng. (2021, April 12). CS229 Lecture notes (Web log message). Retrieved from stanford univerisyt website <https://see.stanford.edu/materials/aimlcs229/cs229-notes1.pdf>
60. Park, H. (2013). An introduction to logistic regression: from basic concepts to interpretation with particular attention to nursing domain. *Journal of Korean Academy of Nursing*, 43(2), 154-164.
61. Kumar,G. (2019). Machine learning techniques for improved business analytics. Germany: Hershey, PA IGI Global, Business Science Reference.
62. Yozgatligil, C. (2021, April 12). learning material for middle east technical university, (Web log message). Retrieved from middle east technical university.users.metu.edu.tr/ceylan/INTRODUCTION%20TO%20BOOSTIN G.ppt
63. Chen, T., Guestrin, C. (2016). Xgboost: A scalable tree boosting system. In: *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, pp. 785-794. ACM, USA, 13 August.

64. Blondel, P. (2021, March 13). Which Machine Learning Algorithm To Choose For My Problem? (Web log message). Retrieved from SAP company. <https://cai.tools.sap/blog/machine-learning-algorithms>
65. Usha, A., & Thampi, S. M. (2017). Authorship Analysis of Social Media Contents Using Tone and Personality Features. In International Conference on Security, Privacy and Anonymity in Computation, Communication and Storage (pp. 212-228). Springer, China, 12-15 Decenber.
66. Barlas, G., Stamatatos, E. (2020). Cross-domain authorship attribution using pre-trained language models. In: IFIP International Conference on Artificial Intelligence Applications and Innovations, pp. 255-266. Springer, Greece, 5-7 June.
67. Kestemont, M., Luyckx, K., Daelemans, W., Crombez, T. (2012). Cross-genre authorship verification using unmasking. *English Studies journal* **93**(3), 340-356.
68. Borison, R. (2021, March 15). PRESENTING: The 100 Most Influential Tech People On Twitter (Web log message). Retrieved from businessinsider.com, <https://www.businessinsider.com/100-influential-tech-people-on-twitter-2014-4>
69. developer.twitter.com (2021). user objects . <https://developer.twitter.com/en/docs/tweets/data-dictionary/overview/user-object.html> İndirilme Tarihi: 12.05.2021.
70. developer.twitter.com (2021). tweet objects. <https://developer.twitter.com/en/docs/tweets/data-dictionary/overview/tweet-object> İndirilme Tarihi: 12.05.2021.
71. Piper, A. (2021, May 15). Upcoming changes to the developer platform (electronic mailing list message). Retrieved from twittercommunity mailing list; <https://twittercommunity.com/t/upcoming-changes-to-the-developer-platform/104603>
72. Halvani, O., Winter, C., Pflug, A. (2016). Authorship verification for different languages, genres and topics. *Digital Investigation journal* **16**, S33-S43.
73. Mikros, G. K., & Perifanos, K.(2013).Authorship Attribution in Greek Tweets Using Author's Multilevel N-Gram Profiles. In AAAI Spring Symposium: Analyzing Microtext. USA, 25-27 March.
74. C. McCollister. (2016). Predicting author traits through topic modeling of multilingual social media text. Doktora Tezi. Kansas Üniversitesi, USA.
75. Maitra, P., Ghosh, S., & Das, D.(2016).Authorship Verification-An Approach based on Random Forest. arXiv preprint arXiv:1607.08885.

76. Mikolov, T., Sutskever, I., Chen, K., Corrado, G.S., Dean, J. (2013). Distributed representations of words and phrases and their compositionality. NIPS'13: Proceedings of the 26th International Conference on Neural Information Processing Systems - Volume 2, pp. 3111-3119, USA, 5-10 December.
77. Mikolov, T., Chen, K., Corrado, G., Dean, J.(2013). Efficient estimation of word representations in vector space. arXiv preprint arXiv:1301.3781.
78. Sivic, J., Zisserman, A. (2008). Efficient visual search of videos cast as text retrieval. IEEE transactions on pattern analysis and machine intelligence **31**(4), 591-606.
79. Ramos, J., ve Ark. (2003). Using tf-idf to determine word relevance in document queries. In: Proceedings of the first instructional conference on machine learning, vol. 242, pp. 133-142. Piscataway, 5-7 December.
80. Hospedales, T., Antoniou, A., Micaelli, P., Storkey, A. (2020). Meta-learning in neural networks: A survey. IEEE Transactions on Pattern Analysis and Machine Intelligence journal pp. 1-1.
81. Lemke, C., Budka, M., Gabrys, B. (2015). Metalearning: a survey of trends and technologies. Artificial intelligence review journal **44**(1), 117-130,.
82. Schoenfeld, B., Giraud-Carrier, C., Poggemann, M., Christensen, J., Seppi, K. (2018). Preprocessor selection for machine learning pipelines. ICML 2018 AutoML Workshop, Austria , 25-31 July.
83. Kotsiantis, S., Kanellopoulos, D., Pintelas, P. (2006). Data preprocessing for supervised learning. International Journal of Computer Science **1**(2), 111-117.
84. Criado, N., Rashid, A., & Leite, L.(2016). Flash mobs, Arab Spring and protest movements: Can we analyse group identities in online conversations?. Expert Systems with Applications journal, 62, 212-224.
85. Ma, X., Feng, Y., Qu, Y., & Yu, Y. (2018). Attribute Selection Method based on Objective Data and Subjective Preferences in MCDM. international journal of computers communications & control, 13(3), 391-407.
86. Singh, R., Kumar, H., & Singla, R. K. (2014). TOPSIS based multi-criteria decision making of feature selection techniques for network traffic dataset. International Journal of Engineering and Technology, 5(6), 4598-4604,.
87. Jahan, A. Edwards, K.L. Bahraminasab, M.(2016). Multi-criteria decision analysis formsupporting the selection of engineering materials in product design. Boston: Butterworth-Heinemann press.
88. Fulop, J. (2005).Introduction to decision making methods. In: BDEI-3 workshop pp. 1-15, Washington, December 13-15.

89. Saaty, T.L.,(2008). Decision making with the analytic hierarchy process. *International journal of services sciences* **1**(1), 83-98.
90. Saaty, T. L. (2001). Decision making with the analytic network process (ANP) and its super decisions software: The national missile defense (NMD) example. *ISAHP 2001 proceedings Symposium*, Switzerland, 2-4 Augstos.
91. Saaty, T.L.,(2005). Theory and applications of the analytic network process: decision making with benefits, opportunities, costs, and risks. Pittsburgh: RWS publications.
92. G. Brestovac, R. Grgurina. (2013). Applying multi-criteria decision analysis methods in embedded systems design. *yüksek lisans Tezi. Mälardalen University, Sweden.*
93. Mesran, G. Ginting, G.Suginam, r. (2017). Implementation of elimination and choice expressing reality (electre) method in selecting the best lecturer (case study stmik budi darma). *Int. J. Eng. Res. Technol.(IJERT*, vol. 6, no. 2, pp. 141-144.
94. Yanie, A., Hasibuan, A., Ishak, I., Marsono, M., Lubis, S., Nurmawati, N., ... & Ahmar, A. S. (2018). Web based application for decision support system with ELECTRE method. In *Journal of Physics: Conference Series* (Vol. 1028, No. 1, p. 012054). IOP Publishing, Indonesia, 9–10 October.
95. Botti, L. Peypoch, N. (2013). Multi-criteria electre method and destination competitiveness. *Tourism Management Perspectives* **6**, 108-113.
96. Yücel, M. Görener, A. (2016). Decision making for company acquisition by electre method. *International Journal of Supply Chain Management* **5**(1), 75-83.
97. A. Hall. (1999). Correlation-based feature selection for machine learning. *Doktora Tezi. Waikato Üniversitesi, NewZeland.*
98. Alazab, M., Huda, S., Abawajy, J., Islam, R., Yearwood, J., Venkatraman, S., Broadhurst, R.,(2014). A hybrid wrapper-filter approach for malware detection. *Journal of networks* **9**(11), 2878-2891.
99. Alazab, M., Layton, R., Broadhurst, R., Bouhours, B.(2013). Malicious spam emails developments and authorship attribution. In: *2013 Fourth Cybercrime and Trustworthy Computing Workshop IEEE*, pp. 58-68, Australia, 21-22 November.
100. Roszkowska, E., (2013). Rank ordering criteria weighting methods-a comparative overview, *Polish National Science Center journal. Nr 5 (65)*, s. 14-33.

6. ÖZGEÇMİŞ

Adı Soyadı : Suleyman ALTERKAVI (Sleman ALTERKAWI)

Yabancı Dil : İngilizce

Eğitim Durumu :

lisans : Haelp Üniversitesi, 2008

Yüksek lisans : IUL Üniversitesi, 2014

Çalıştığı Kurum/Kurumlar ve Yıl/Yıllar: MIS danışmanı, Serbest çalışan,
(07/09/2020- devam).

Yayımları (SCI) :

SCI-Expanded İndekslerinde Yer Alan Dergilerde Yayımlanan Makaleler :

Makale	Dergi
Novel authorship verification model for social media accounts compromised by a human	Multimedia Tools and Applications, Springer
Design and Analysis of a Novel Authorship Verification Framework for Hijacked Social Media Accounts Compromised by a Human	Security and Communication Networks, Hindawi