



T.C.
İSTANBUL ÜNİVERSİTESİ-CERRAHPAŞA
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ



DOKTORA TEZİ

**SİBER TEHDİT İSTİHBARATI İÇİN YENİ TARAMA MODELİ
GELİŞTİRİLMESİ**

Ebu Yusuf GÜVEN

DANIŞMAN

Doç. Dr. Muhammed Ali AYDIN

II. DANIŞMAN

Dr. Öğr. Üyesi Ali BOYACI

Bilgisayar Mühendisliği Anabilim Dalı

Bilgisayar Mühendisliği Programı

Mart, 2023

Aziz Şehitlerimize ithaf ediyorum.

BÜTÇE DESTEKLERİ

SİBER TEHDİT İSTİHBARATI İÇİN YENİ TARAMA MODELİ GELİŞTİRİLMESİ

Bu tez çalışması için herhangi bir kurumdan bütçe desteği alınmamıştır.

TEŞEKKÜR

Tez sürecinde desteklerinden dolayı aileme ve dostlarıma teşekkür ederim. Çalışmalarına vermiş oldukları destek ve yönlendirmeleriyle katkı sağlayan kıymetli hocalarım Doç. Dr. Muhammed Ali AYDIN ve Dr. Öğr. Üyesi Ali BOYACI' ya teşekkürü borç bilirim. Tez düzenleme ve kontrol sürecinde destek olan Dr. Öğr. Üyesi Zeynep GÜRKAŞ AYDIN ve Araş. Gör. Melike BAŞER'e teşekkür ederim.

Mart 2023

Ebu Yusuf GÜVEN

İÇİNDEKİLER

	Sayfa No
BÜTÇE DESTEKLERİ	v
TEŞEKKÜR	vi
İÇİNDEKİLER	vii
ŞEKİL LİSTESİ	ix
TABLO LİSTESİ	xii
SİMGE VE KISALTMA LİSTESİ	xiv
ÖZET	xvii
ABSTRACT	xix
1. GİRİŞ	1
1.1 Veri İhlali Sınıflandırması	2
1.2 Kullanıcı Bilgileri Analizi	3
1.3 Katkılar	5
1.4 Organizasyon	6
2. KAVRAMSAL ÇERÇEVE	7
1.1. Siber Tehdit İstihbaratı	7
1.1.1. CTI Mevcut Modeller	9
1.1.2. CTI Araçlar ve Platformlar	10
1.2. Veri İhlal İstihbaratı	11
1.2.1. İstihbarat Verilerinin Değerlendirilmesi	12
1.3. Siber Tehdit İstihbaratında Kullanılan Özellik ve Anlam Çıkarma Yöntemleri	15
1.3.1. Düzenli İfadeler	15
1.3.2. İsimlendirilmiş (Adlandırılmış) Varlık Tanıma	15
1.3.3. Doğal Dil İşleme	16
1.3.4. Optik Karakter Tanıma	17
1.4. Veri Sınıflandırma Yöntemleri	17
1.4.1. Rastgele Orman	17
1.4.2. Destek Vektör Makinesi	18

1.4.3.	K-En Yakın Komşu	19
1.4.4.	Naive Bayes	20
1.4.5.	Yinelemeli Sinir Ağı	21
1.4.6.	Modellerin Performans Hesaplaması	21
1.5.	İlgili Çalışmalar	22
3.	YÖNTEM	32
3.1.	Siber Tehdit İstihbaratı için Yeni Tarama Modeli	32
3.1.1.	Sızıntı Verisi Kaynağı	34
3.1.2.	Sızıntı Verisine Erişim ve Çözümleme	37
3.1.3.	Veri Yönetimi	42
3.1.4.	Veri Sınıflandırma	43
3.1.5.	Sunum Katmanı	44
3.2.	AuthInfo Veri Seti	44
3.3.	Siber Tehdit İstihbaratı Veri Seti	47
3.3.1.	REGEX	48
3.3.2.	Varlık İsim Tanıma (Named Entity Recognition - NER)	49
3.4.	Deneysel Araçlar	50
4.	BULGULAR	52
4.1.	Sızıntı Verilerinin Analizi	52
4.2.	Veri Sızıntılarının Sınıflandırılması	61
4.3.	Kullanıcı Adı-Parola Dosyalarının Analizi: Parola Seçim Davranışı Tespiti	82
4.4.	Vaka Çalışması: Parola Seçim Davranışı Kullanarak Kaba Kuvvet Saldırısı	88
4.5.	Parola Davranışı Duyarlı Parola Politikası yöntemi	95
5.	TARTIŞMA	98
6.	SONUÇ VE ÖNERİLER	101
	KAYNAKLAR	104
	İNTİHAL RAPORU İLK SAYFASI	116

ŞEKİL LİSTESİ

	Sayfa No
Şekil 2.1: Hiper düzlem örnekleri ile iki özelliği ayıran kategoriler arasındaki maksimum marjın [72].	19
Şekil 2.2: Üç sınıflı bir sınıflandırma problemi için KNN algoritması örneği [76].	20
Şekil 3.1: Siber Tehdit İstihbaratı için Yeni Tarama Modeli Katmanları.	34
Şekil 3.2: Açık web tarama için mevcut arama motorları üzerinden veri akışı.	35
Şekil 3.3: Veri elde edilmesi akışı.	36
Şekil 3.4: Crawler Modülü'nün veri akışı.	37
Şekil 3.5: API Modülü veri akışı.	38
Şekil 3.6: Siber Tehdit İstihbaratı için Yeni Tarama Modeli Ayırıştırıcı Modülü ve Alt Modülleri.	39
Şekil 3.7: İçerik Ayırıştırıcı Modülü İçerik Ayırıştırıcı alt modülü veri akışı.	40
Şekil 3.8: Veri Ayırıştırıcı Modülü veri akışı.	41
Şekil 3.9: Siber Tehdit İstihbaratı için Yeni Tarama Modeli Veri tabanı tasarımı.	43
Şekil 3.10: AuthInfo veri seti oluşturulması aşamaları.	45
Şekil 3.11: Siber Tehdit İstihbaratı için Yeni Tarama Modeli Veri İşleme Mimarisi.	51
Şekil 4.1: Siber Tehdit İstihbaratı için Yeni Tarama Modeli İçerik Ayırıştırıcı yapısal ve serbest metin dosyaların dağılımı.	53
Şekil 4.2: Model tarafından kural tabanlı etiketlenen verinin operatör tarafından düzenlenmesinin karmaşıklık matrisi.	56
Şekil 4.3: 4 Etiket Model Veri Seti ve 4 Etiket Operatör Veri Seti karşılaştırması.	56
Şekil 4.4: 2 Etiket Model Veri Seti ve 2 Etiket Operatör Veri Seti karşılaştırması.	57
Şekil 4.5: 4 Etiket Operatör Veri Seti karşılaştırması.	58
Şekil 4.6: 2 Etiket Operatör Veri Seti karşılaştırması.	58
Şekil 4.7: Veri Seti sık geçen kelimeler.	60

Şekil 4.8: Etkisiz kelime temizleme operasyonu sonrası veri seti sık geçen kelimeler.	61
Şekil 4.9: KNN Algoritması REGEX Veri Seti Karmaşıklık Matrisi.	62
Şekil 4.10: KNN Algoritması NER Veri Seti Karmaşıklık Matrisi.	63
Şekil 4.11: KNN Algoritması REGEX ve NER Veri Seti Karmaşıklık Matrisi.	63
Şekil 4.12: Naive Bayes Algoritması REGEX Veri Seti Karmaşıklık Matrisi.	64
Şekil 4.13: Naive Bayes Algoritması NER Veri Seti Karmaşıklık Matrisi.	65
Şekil 4.14: Naive Bayes Algoritması REGEX ve NER Veri Seti Karmaşıklık Matrisi.	65
Şekil 4.15: Rastgele Orman Algoritması REGEX Veri Seti Karmaşıklık Matrisi.	66
Şekil 4.16: Rastgele Orman Algoritması NER Veri Seti Karmaşıklık Matrisi.	67
Şekil 4.17: Rastgele Orman Algoritması REGEX ve NER Veri Seti Karmaşıklık Matrisi.	67
Şekil 4.18: SVM Algoritması REGEX Veri Seti Karmaşıklık Matrisi.	68
Şekil 4.19: SVM Algoritması NER Veri Seti Karmaşıklık Matrisi	69
Şekil 4.20: SVM Algoritması REGEX ve NER Veri Seti Karmaşıklık Matrisi.	69
Şekil 4.21: REGEX Veri Seti üzerinde doğruluk değeri karşılaştırması.	70
Şekil 4.22: NER Veri Seti üzerinde doğruluk değeri karşılaştırması.	70
Şekil 4.23: REGEX+NER Veri Seti üzerinde doğruluk değeri karşılaştırması.	71
Şekil 4.24: KNN Algoritması çoklu etiket REGEX Veri Seti Karmaşıklık Matrisi.	72
Şekil 4.25: KNN Algoritması çoklu etiket NER Veri Seti Karmaşıklık Matrisi.	73
Şekil 4.26: KNN Algoritması çoklu etiket REGEX ve NER Veri Seti Karmaşıklık Matrisi.	73
Şekil 4.27: Naive Bayes Algoritması çoklu etiket REGEX Veri Seti Karmaşıklık Matrisi.	74
Şekil 4.28: Naive Bayes Algoritması çoklu etiket NER Veri Seti Karmaşıklık Matrisi.	75
Şekil 4.29: Naive Bayes Algoritması çoklu etiket REGEX ve NER Veri Seti Karmaşıklık Matrisi.	75
Şekil 4.30: Rastgele Orman Algoritması çoklu etiket REGEX Veri Seti Karmaşıklık Matrisi.	76
Şekil 4.31: Rastgele Orman Algoritması çoklu etiket NER Veri Seti Karmaşıklık Matrisi.	77
Şekil 4.32: Rastgele Orman Algoritması çoklu etiket REGEX ve NER Veri Seti Karmaşıklık Matrisi.	77

Şekil 4.33: SVM Algoritması çoklu etiket REGEX Veri Seti Karmaşıklık Matrisi.	78
Şekil 4.34: SVM Algoritması çoklu etiket NER Veri Seti Karmaşıklık Matrisi.	79
Şekil 4.35: SVM Algoritması çoklu etiket REGEX ve NER Veri Seti Karmaşıklık Matrisi.	79
Şekil 4.36: Çoklu etiket REGEX Veri Seti üzerinde doğruluk değeri karşılaştırması.	80
Şekil 4.37: Çoklu etiket NER Veri Seti üzerinde doğruluk değeri karşılaştırması.	80
Şekil 4.38: Çoklu etiket REGEX+NER Veri Seti üzerinde doğruluk değeri karşılaştırması	81
Şekil 4.39: AuthInfo veri seti, farklı parola uzunluklarının dağılımı	83
Şekil 4.40: Karakter uzunlukları içindeki karakter türlerinin yüzde dağılımı	84
Şekil 4.41: Sızıntı kaynaklarının parola uzunluk dağılımı yüzdeleri	88
Şekil 4.42: Çoklu sızıntıların dağılımı	89
Şekil 4.43: Birden fazla sızıntıya maruz kalan kullanıcıların dağılımı	89
Şekil 4.44: Kurban kullanıcının dokuz çözülmüş parolalarının karakter tipi dağılımı	95
Şekil 4.45: Güvenli parola seçimi kontrol adımları	96

TABLO LİSTESİ

	Sayfa No
Tablo 2.1: KVKK kapsamında kişisel veri kabul edilen bazı veri türleri.	13
Tablo 3.1: : Anahtar kelimeleri ve sızıntı kaynakları.	35
Tablo 3.2: Ham dosyaların tutulduğu klasörlerin dağılımı.	42
Tablo 3.3: Parola içindeki karakter tiplerinin gösterimi.	46
Tablo 3.4: AuthInfo veri setinde bulunan özellikler.	47
Tablo 3.5: Serbest metin veri ayrıştırmak için kullanılan REGEX kuralları ve özellikler.	48
Tablo 3.6: NER yöntemi ile çıkarılan özellikler.	49
Tablo 4.1: Dosya içerdiği bilgi türüne göre ayrıştırılabilen klasör ağacı	53
Tablo 4.2: Siber Tehdit İstihbaratı için Yeni Tarama Modeli dosya etiketleme kuralları.	54
Tablo 4.3: Dosyalardan model sonucu ve operatör yardımı ile etiketlenen veri setleri.	55
Tablo 4.4: NLP Ön işleme Öncesi Veri Seti sık geçen tokenler.	59
Tablo 4.5: KNN algoritması REGEX ve NER veri setleri model performans metrikleri.	61
Tablo 4.6: Naive Bayes algoritması REGEX ve NER veri setleri model performans metrikleri	64
Tablo 4.7: Rastgele Orman algoritması REGEX ve NER veri setleri model performans metrikleri	66
Tablo 4.8: SVM algoritması REGEX ve NER veri setleri model performans metrikleri	68
Tablo 4.9: KNN algoritması NER veri setleri model performans metrikleri	71
Tablo 4.10: Naive Bayes algoritması çoklu etiket NER veri setleri model performans metrikleri	74
Tablo 4.11: Rastgele Orman algoritması çoklu etiket NER veri setleri model performans metrikleri	76
Tablo 4.12: SVM algoritması çoklu etiket NER veri setleri model performans metrikleri	78

Tablo 4.13: RNN algoritması çoklu etiket Operatör Veri Setleri için model performans metrikleri.	81
Tablo 4.14: Parolaların uzunluk ve karakter türlerine göre dağılımı	84
Tablo 4.15: Parola içeriği karakter türü dağılımı	85
Tablo 4.16: Parola ilk karakter dağılımı	86
Tablo 4.17: Parola maskesinin dağıtımı	87
Tablo 4.18: dce9fef8b5c3ba9d576ee9ad37259632 AnonID'ye sahip bir kullanıcının sızmış paroları	90
Tablo 4.19: AuthInfo istatistikleri kullanılarak kırılan MD5 özetleri parolalarının özellikleri	92
Tablo 4.20: Farklı parola uzunluğuna sahip saldırı süreleri	93

SİMGE VE KISALTMA LİSTESİ

Kisaltmalar	Açıklama
APT	: Advanced Persistent Threat - Gelişmiş Sürekli Tehdit
API	: Application Programing Interface - Uygulama Programlama Arayüzü
AuthInfo	: Authentication Information - Kimlik Doğrulama Bilgileri
AnonID	: Anonymous ID - Anonim Kimlik
BT	: Bilgi Teknolojileri
BERT	: Bidirectional Encoder Representations from Transformers – Transformers'tan Çift Yönlü Kodlayıcı Temsilleri
C2	: Command and Control - Komuta ve Kontrol
CIRCL	: The Computer Incident Response Center Luxembourg – Lüksemburg Bilgisayar Olay Müdahale Merkezi
CTI	: Cyber Threat Intelligence - Siber Tehdit İstihbaratı
CNN	: Convolutional Neural Network - Evrişimli Sinir Ağları
DLP	: Data Leak Prevention - Veri Sızıntı Önleme
DLD	: Data Leak Detection - Veri Sızıntı Tespiti
DML	: Detection Maturity Level - Tespit Olgunluk Düzeyi
GDPR	: General Data Protection Regulation - Genel Veri Koruma Yönetmeliği
GloVe	: Global Vectors for Word Representation – Kelime Temsili için Global Vektörler
HUMINT	: Human Intelligence - İnsan İstihbaratı
HTML	: Hyper Text Markup Language - Hiper Metin İşaretleme Dili
IFC	: Information Flow Control - Bilgi Akış Kontrolü
IoT	: Internet of Things - Nesnelerin İnterneti
IOC	: Indicator of Compromise - Uzlaşma Göstergesi
VPN	: Virtual Private Network - Sanal Özel Ağ

IP	: Internet Protocol - İnternet Protokolü
ICS	: Industrial Control System - Endüstriyel Kontrol Sistemi
KVKK	: Kişisel Verileri Koruma Kanunu
KNN	: K-Nearest Neighbor - K En Yakın Komşu
LDA	: Latent Dirichlet Allocation - Gizli Dirichlet Tahsisi
MISP	: Malware Information Sharing Platform – Zararlı Yazılım Bilgi Paylaşım Platformu
MD5	: Message-Digest 5 - Mesaj Özeti 5
NER	: Named Entity Recognition - Adlandırılmış Varlık Tanıma
NFV	: Network Function Virtualization - Ağ İşlevi Sanallaştırması
NLP	: Natural Language Processing - Doğal Dil İşleme
NLTK	: Natural Language Toolkit - Doğal Dil Araç Seti
PII	: Personally Identifiable Information - Kişisel Tanımlanabilir Bilgiler
UserID	: User ID - Kullanıcı Kimliği
OCR	: Optical Character Recognition - Optik Karakter Tanıma
OSINT	: Open Source Intelligence - Açık Kaynak İstihbaratı
REGEX	: Regular Expression - Düzenli İfadeler
RF	: Random Forest - Rastgele Orman
RNN	: Recurrent Neural Networks - Yinelemeli Sinir Ağı
SDN	: Software Defined Networking - Yazılım Tanımlı Ağlar
SOCMINT	: Social Media Intelligence - Sosyal Medya İstihbaratı
SVM	: Support Vector Machine - Destek Vektör Makinesi
SFTP	: SSH File Transfer Protocol - SSH Dosya Aktarım Protokolü
SSH	: Secure Shell Host - Güvenli Kabuk
TOR	: The Onion Router - Anonim Ağ
TTP	: Tactics, Techniques, and Procedures - Teknik, Taktik ve Prosedür

ÖZET

DOKTORA TEZİ

SİBER TEHDİT İSTİHBARATI için YENİ TARAMA MODELİ GELİŞTİRİLMESİ

Ebu Yusuf GÜVEN

İstanbul Üniversitesi-Cerrahpaşa

Lisansüstü Eğitim Enstitüsü

Bilgisayar Mühendisliği Anabilim Dalı

Bilgisayar Mühendisliği Programı

Danışman : Doç. Dr. Muhammed Ali AYDIN

II. Danışman : Dr. Öğr. Üyesi Ali BOYACI

Kıymet ithaf edilen varlık ve olguların gizlenerek korunması insanlığın varoluşundan beri süregelen bir durumdur. İçinde bulunduğumuz birbirine bağlı nesnelerin dijital çağında, güvenlik ve gizlilik ihtiyaçları çözülmemiş en büyük olgu ise bilgidir. Bilgi güvenliği gizlilik, bütünlük ve erişilebilirlik temel kavramlarının sağlanması ile ifade edilmektedir. Birbirine bağlı dijital sistemlere kaydedilen bilginin güvenliğini sağlamak için bu sistemlerin güvenlik risklerini ve kullanan insanlardan kaynaklanan riskleri birlikte yönetmek gerekmektedir. Riskleri belirlemek, mümkün güvenlik sıkılaştırmalarını yapmak ve olası veri ihlallerini takip etmek için siber tehdit istihbaratı kavramı öne çıkmaktadır. Siber tehdit istihbaratı sistemlere içeriden ve dışarıdan oluşabilecek tehditlere karşı bilgi edinme faaliyetidir. Oluşan tehditlerin anlaşılıp tehdit kaynağına ait olan bilgilerin ve teknik, taktik ve prosedürlerin toplanması faaliyetlerinin bütünüdür. İstihbarat verisi her zaman tehditlere yönelik olmaz, saldırı veya kötü yapılandırmalardan kaynaklı veri sızıntıları da bulunabilmektedir. Siber suçlar, bilgisayar korsanlığı, siber casusluk, güvenlik zafiyetleri ve kötü yapılandırma gibi sebeplerle ortaya çıkan veri sızıntıları genellikle yapısal olmayan formlarda internet üzerinden paylaşılmaktadır. Sızıntı verileri anonim dosya paylaşım platformları, metin ve resim paylaşım platformları,

sosyal medya uygulamaları ve mesajlaşma uygulamaları gibi çeşitli kaynaklarda bulunabilmektedir. Bu kaynaklara erişim arama motorları veya veri sızıntısı ve siber saldırganların kullandığı forum sayfaları üzerinden yapılabilmektedir. Bu veri ihlali istihbaratı analizi çalışmaları genellikle operatör destekli ve kural tabanlı sistemler ve organizasyonlar tarafından yürütülmektedir. Bu tez çalışmasında, yapay zeka ve doğal dil işleme kullanarak veri sızıntıları ve ihlalleri için minimum operatör müdahalesiyle tarama, analiz ve değerlendirme çalışmaları yapılmıştır. Değerlendirilen veri kişisel, kurumsal, potansiyel bilgi ve anonim olarak sınıflandırılmıştır. Geliştirilen kural tabanlı yöntem ile %83 doğruluğa ulaşılırken makine öğrenimi modelleri ile en iyi %89 doğruluk oranlarıyla sızıntı verileri sınıflandırılmıştır. Siber tehdit istihbaratı ve veri ihlali istihbaratı çalışmalarında doğal dil işleme ve yapay zeka kullanılarak verilerin değerlendirilebileceği gösterilmiştir.

Mart 2023 , 139 sayfa.

Anahtar kelimeler: Siber Güvenlik, Siber Tehdit İstihbaratı, Sızmış Veri, Veri İhlali

ABSTRACT

Ph.D. THESIS

DEVELOPMENT OF A NEW SCANNING MODEL IN CYBER THREAT INTELLIGENCE

Ebu Yusuf GÜVEN

İstanbul University-Cerrahpaşa

Institute of Graduate Studies

Department of Computer Engineering

Computer Engineering

Supervisor : Assoc. Prof. Dr. Muhammed Ali AYDIN

Co-Supervisor: Assist. Prof. Dr. Ali BOYACI

The concealment and protection of the assets and facts that are valued have been going on since the existence of humanity. In the digital age of connected objects, we live in, information is the biggest unresolved issue of security and privacy needs. Information security provides the basic concepts of confidentiality, integrity, and accessibility. To ensure the information security recorded in interconnected digital systems, it is necessary to manage the security risks of these systems and the risks arising from the people using them together. Cyber threat intelligence comes to the fore to identify and track risks, make possible security tightening, and monitor potential data breaches. Cyber threat intelligence is obtaining information against threats that may arise inside and outside the systems. It is the whole activity of understanding the threats and collecting the information and techniques, tactics, and procedures belonging to the source of the threat. Intelligence data is not always for threats. There may also be data leaks due to attacks or wrong configurations. Data leaks due to cybercrime, hacktivism, cyber espionage, security vulnerabilities, and incorrect

configuration are generally shared over the internet in unstructured forms. Data breach sources can be found in various sources such as anonymous file-sharing platforms, text and image-sharing platforms, social media applications, and messaging applications. Access to these resources can be done through search engines or forum pages. Operator -assisted and rule-based systems and organizations conduct these data breach intelligence analysis studies. In this thesis, scanning, analysis, and evaluation studies were carried out for data leaks and violations with minimal operator intervention using artificial intelligence and natural language processing. Evaluated data is classified as personal, institutional, and public. As a result of the developed models, it has been shown that data can be estimated using natural language processing and artificial intelligence in cyber threat intelligence and data breach intelligence studies with 89% accuracy.

March 2023, 139 pages.

Keywords: Cyber Security, Cyber Threat Intelligence, Leaked Data, Data Breach

1. GİRİŞ

İnternet ağında çalışmak, alışveriş yapmak ve sosyal ortamlar gibi birçok amaçla kullanıcılar kimlik numarası, adı-soyadı, e-posta, telefon numarası gibi kişisel bilgilerini hizmet sağlayıcılar ile paylaşırlar. İnternet ortamında hizmet sağlayıcıların sık karşılaştıkları veri ihlalleri insanların veri gizliliği endişelerini artırmıştır [1]. Veri Sızıntıları, güvenlik açıkları, güncel olmayan veya doğru yapılandırılmamış bir servisler, sosyal mühendislik saldırıları kullanılarak hizmet sağlayıcı veya son kullanıcının bilgileri ele geçirilmektedir [2]. Tehdit aktörleri ve vektörleri açısından farklılık gösterse de veri sızıntıları bilgisayar korsanlığı, propaganda [3], şantaj, siber casusluk [4], finansal kayıplar [5] ve itibar kayıplarına yol açmak gibi ortak amaçlarla ortaya çıkmaktadır.

Veri ihlali denildiğinde dijital sistemler ve İnternet akla gelse de fiziksel kopyalar ve çevrimdışı diskler üzerinden de veri sızıntılarının gerçekleşmesi mümkündür. İmha ve fiziksel güvenlik prosedürleri fiziksel sistemlerde daha kolay işletildiği için bu tür veri sızıntıları nadir gerçekleşmektedir. Siber saldırgan tarafından yasal olmayan bir yolla elde edilen veriler, verinin niteliği, önemi, gizlilik derecesi ve finansal değeri gibi kriterlere göre değişiklik gösterse de saldırgan, veri güncelliğini kaybetmeden maddi ya da manevi fayda elde etmeye çalışır [6]. Saldırgan ele geçirdiği veri sızıntısını doğrudan satarak ya da şantaj yaparak veri ihlaline uğrayan veri sorumlunun üzerinden maddi fayda elde etmeye çalışır. Yeraltı forumları olarak nitelendirilen yasadışı forum sitelerinde veri sızıntılarının kripto paralar cinsinden satılması oldukça yaygındır [7]. Güncel veri sızıntıları ya da önceden paylaşılmış veri sızıntıları, propaganda veya suçun yaygınlaştırılması amacıyla paylaşılmaktadır. Veri ihlali duyuruları genellikle veri sızıntısı bağlantı yolunun yer aldığı yasadışı forumlardan [8], sosyal ağlar veya grup mesajlaşma uygulamaları üzerinden paylaşılmaktadır. Paylaşılan bağlantı yolları genellikle anonim metin, dosya veya resim paylaşım siteleri olmaktadır. Genellikle anonfiles.com, pastebin.com veya anonim bir e-posta hesabı ile açılmış dosya paylaşım (drive) uygulamaları kullanılmaktadır.

Farklı tehdit aktörleri ve vektörleri olduğu gibi ortaya çıkan veri sızıntısının saklanma biçimi de farklılık göstermektedir. Bazı veri sızıntıları serbest metin olarak karşımıza çıkarken, yapısal olarak virgülle, iki noktayla veya tab karakteri ile ayrılmış kendi içinde bir kuralı olan

dosyalar da bulunmaktadır. Ancak bu yapısal olarak nitelendirilen dosyaların da içerisinde karakter seti problemleri, sütunların yer değiştirmesi veya eksik olması gibi format hataları bulunmaktadır. Veri sızıntılarının genelleştirilebilir bir bilgi çıkarma yöntemi bulunmaması doğrudan sızıntı dosyaları üzerinde analiz yapmayı zorlaştırmakta ve bu nedenle veri sızıntılarının ön işlemden geçirilmesi gerekmektedir. Bu durum tehdit aktörü ve tehdit vektörü bazında ön işleme ihtiyaçları ortaya çıkarmaktadır ve çalışmaların operatör kontrolünde devam etmesini zorunlu kılmaktadır.

1.1 Veri İhlali Sınıflandırması

Günümüz, siber tehditlerin hızla geliştiği ve çeşitlendiği, bireysel ve gruplardan oluşan siber saldırı ekiplerinin ülke destekli gelişmiş kalıcı tehdit gruplarına dönüştüğü bir siber savaş devridir. Siber tehditler; güvenlik açıkları, hatalı yapılandırmalar ve kötü niyetli kullanıcılar gibi birçok durum veri sızıntılarına ve veri ihlallerine sebep olmaktadır. Veri ihlali (data breach), hassas, korunan veya gizli verilerin yetkisiz bir kişi tarafından kopyalandığı, iletildiği, görüntülendiği, çalındığı veya kullanıldığı bir güvenlik olayıdır [9]. Sen ve Borle [10] hassas, korunan veya gizli verilerin, kişisel sağlık bilgilerini, kişisel tanımlanabilir bilgileri, ticari sırları veya fikri mülkiyeti ve kişisel mali verileri içerebileceğini ortaya koymuşlardır.

Artan tehditler ile hükûmet, sağlık, finansal hizmetler, sigorta, sosyal medya ve daha fazlası dahil olmak üzere çeşitli sektörlerde veri ihlalleri meydana gelmektedir [11]. Veri ihlalleri, yalnızca güvenlik uzmanları için bir endişe kaynağı değildir; ayrıca müşteriler, paydaşlar, kuruluşlar ve işletmeler de doğrudan veya dolaylı olarak etkilenirler. Veri ihlallerinin, erişim şekli, kaynağı ve içeriği farklı türlerde olsa da etkileri genellikle aynıdır. Veri ihlalleri bireylere ve kuruluşlara maddi ve manevi zararlar verebilir. Veri sızıntıları kurum ve kuruluşların imajını zedeleyerek itibarlarını ve marka değerlerini zedeler. Maschler ve arkadaşları tarafından sızıntı dosyalarının ve bu sızıntı dosyalarıyla ilişkili kayıtların gerçekliğini doğrulayan bir yazılım geliştirilmiştir [12]. Veri ihlalleri genellikle dahili (internal) ve harici (external) olmak üzere iki ana kategoriye ayrılır: Dahili veri ihlalleri, dahili bir ajanın yardımıyla meydana gelen olaylardan oluşur. Bunlar ayrıcalığın kötüye kullanılması, yetkisiz erişim/ifşa, gereksiz ancak hassas verilerin uygunsuz şekilde imha edilmesi, kayıp, hırsızlık veya gizli verilerin yetkisiz bir tarafla kasıtsız olarak paylaşılması olabilir. Harici veri ihlalleri, herhangi bir harici ajan veya kaynağın neden olduğu olaylardır.

Bunlar, kötü amaçlı yazılım saldırısı, fidye yazılımı saldırısı, kimlik avı, casus yazılım veya çalınan kredi kartlar vb. biçimindeki dolandırıcılık gibi herhangi bir bilgisayar korsanlığı ve bilgi teknolojileri (BT) olayını içerir [13].

Tez çalışmamız kapsamında açık ve özel internet üzerinden erişilebilir kaynaklardan veri sızıntılarının taranması, analiz edilmesi, yapısal hale getirilmesi, değerlendirilmesi ve sınıflandırılması süreçlerini kapsayan yeni bir model önerilmektedir. Sahada genellikle operatör destekli kural tabanlı sistemlerin kullanılmasının aksine, bu tez çalışmasında minimum operatör desteği ile verilerin elde edilerek sınıflandırılması amaçlanmıştır. Veri ihlalleri, arama motorları, TOR ağı, Telegram kanalları, veri ihlali paylaşan uygulama programlama arayüzleri (Application Programming Interface-API) üzerinden toplanmaktadır. Geliştirilen ayrıştırıcı modülleri ile dosyalar; türleri, dosya içerik yapıları ve dosyaların içerdikleri bilgilere göre ayrıştırılarak depolanmaktadır. Ayrıştırma işlemi sırasında ve veri kazıma aşamalarında REGEX, NER ve Doğal Dil İşleme operasyonları ile oluşturulan özellikler kullanılarak kural tabanlı ve yapay zeka destekli sınıflandırma işlemleri gerçekleştirilmiştir. Çalışma kapsamında altı farklı veri seti oluşturularak eğitilen makine öğrenmesi ve yapay zeka modellerinin performansları karşılaştırılmıştır.

1.2 Kullanıcı Bilgileri Analizi

Kullanıcılar, hizmetlere erişim sağlamak için kullanıcı doğrulaması yapmak, veri gizliliğini ve güvenliğini sağlamak amacıyla hala metin tabanlı parolaları kullanmaya devam etmektedir [14]. Çevrimiçi hizmetlerdeki metin tabanlı parolalar, genellikle bir kişiyi tanımlayan bir kullanıcı kimliği (userID - kimlik, takma ad, telefon numarası, kullanıcı adı veya e-posta) ile birlikte kullanılır. Web üzerinden hizmet sağlayan herhangi bir kuruluş veya kişi, kişisel/özel bilgi erişimini doğrulamak için parola kullanır. Bu nedenle, bu hizmet sağlayıcılar, kullanıcıların parolalarının yeterince karmaşık ve tahmin edilmesi zor olduğundan emin olmak için bir tür parola politikasına ihtiyaç duyarlar [15]. Ancak saldırganlar, kullanıcıların güvenliğini zayıflatmak için kaba kuvvet (brute force), sözlük tabanlı (dictionary-based) ve gökkuşağı tablosu (rainbow table) tabanlı saldırılar dahil olmak üzere çeşitli karşı önlemler geliştirmiştir [16]. Ayrıca saldırganlar, kullanıcı oturum bilgileri olmadan sistem güvenlik açıklarından yararlanarak da kullanıcı bilgilerine erişebilir [17].

Kullanıcılar genellikle farklı hizmetler için aynı parolayı kullanma eğilimindedir [18]. Kullanıcıların kimlik bilgilerinden biri (userID veya parola) ele geçirildiğinde, bu bilgiler potansiyel olarak kullanıcı hesaplarının daha fazla kötüye kullanılmasına yol açar. Sonuç olarak, sızan veya çalınan veriler yasa dışı bir şekilde satılabilir ve anonim dosya veya metin paylaşım platformları üzerinden paylaşılabilir. Bu platformlar arasında bulut tabanlı depolama, Github gibi kod barındırma hizmetleri ve dosyaların veya metinlerin anonim olarak paylaşılmasını sağlayan uygulamalar yer alabilir. Bu durum, yasadışı olarak elde edilen kullanıcı verilerinin ticareti veya paylaşımı için izlenemez bir pazar yeri kurulmasının önünü açmaktadır. Öte yandan, insanlar açık kaynaklarda paylaşılan kimlik bilgilerini "haveibeenpwd.com" ve "breachalarm.com" gibi, kullanıcıların kullanıcı kimlikleri/parola sızıntı durumlarını sorgulamalarına olanak tanıyan hizmetler oluşturulmuştur. 11 milyarın üzerinde kullanıcı bilgisine sahip veri sızıntı kontrol sistemleri, güvenlik araştırmacıları için bir çalışma alanı oluşturmaktadır [19]. Ancak, platformlar yalnızca keşfedilen veya herkese açık verilerden oluştuğu için sızan kullanıcı bilgileri ve parolaların toplam sayısı daha fazla olabilir.

Parola politikaları, kullanıcıların tahmin edilmesi kolay veya düzenli olarak kullanılan parolaları seçmesini engellemeyi amaçlamaktadır [20]. Ancak, yeterli parola sızıntısıyla bu politikalar bile kullanıcıların güvenliğini güçlendirmede yetersiz kalabilmektedir [21]. Tüm güvenlik ihlallerini önlemek mümkün olmadığından, kullanıcıların güvenliğini sağlamak için parola politikalarında daha karmaşık tekniklerin gerekli olduğu ortaya çıkmaktadır.

Siber Tehdit İstihbaratı (Cyber Threat Intelligence - CTI) için yeni bir model geliştirilmesi kapsamında, kullanıcı adı-parola bilgileri içeren veri dosyaları ayrıştırılmış ve kullanıcı hesap bilgilerinden parola seçimi davranışı analiz edilmiştir. Veri sızıntıları arasından seçilen sekiz farklı kaynak üzerinden on milyondan fazla kullanıcı kaydı ayrıştırılmıştır. Toplanan verilerde yer alan kullanıcılar ayrıca bir veri sızıntısına maruz bırakılmamak amacıyla kullanıcı adları anonimleştirilmiş ve parola üzerinden özellikler çıkarılarak AuthInfo veri seti oluşturulmuştur. Parola seçim eğilimlerini analiz eden AuthInfo veri seti Github platformu aracılığıyla paylaşılmıştır. Ayrıca istatistiksel yöntemler ve görselleştirme kullanılarak kullanıcı parola seçim davranışları değerlendirilmiştir.

Kullanıcıların birden fazla parolası ele geçirildiğinde aynı parola seçim davranışı ile farklı parolalar seçmesi halinde bir sonraki parolasının tahmininde sızmış parolaların

kullanılabileceği Parola Seçim Davranışı Kullanarak Kaba Kuvvet Saldırısı bölümde gösterilmiştir. Sonuçlar, kullanıcı davranışının istatistiksel analizinin, birden fazla veri ihlali kurbanı olan bir kullanıcıya yönelik kaba kuvvet saldırılarının başarı oranını artırmaya yardımcı olabileceğini göstermiştir. AuthInfo veri seti üzerinde çoklu veri ihlaline maruz kalmış bir kullanıcı üzerinden kaba kuvvet saldırısı gerçekleştirilmiştir. Mevcut işlem gücü ile kaba kuvvet saldırı uzayı dışında olan parolaları AuthInfo veri seti özellikleri ve analiz sonuçları kullanılarak başarılı saldırılar gerçekleştirilmiştir. Parola davranışının, kullanıcı parola üzerinde değişikliğe gitse bile devam ettiği görüldüğü için AuthInfo veri seti üzerinde oluşturduğumuz özellikler kullanılarak Parola Seçim Davranışı Duyarlı Parola Güvenlik Politikası bölümünde önerilmiştir.

1.3 Katkılar

Bu tez çalışması kapsamında veri sızıntılarının ve ihlallerinin açık ve özel web üzerinden taranması, analiz edilmesi, yapısal hale getirip sınıflandırması;

- Siber Tehdit İstihbaratı için Yeni Tarama Modeli sunulmuştur,
- Veri Sızıntılarının Analizi ve Sınıflandırması için;
 - Veri sızıntılarının anlaşılması için analiz ve görselleştirme çalışması yapılmıştır,
 - REGEX, NER ve doğal dil işleme yöntemleri ile özellik çıkarımı yapılarak CTI Veri Seti paylaşılmıştır,
 - Veri sızıntılarını sınıflandırmak için %83 doğruluk ile çalışan kural tabanlı bir sistem geliştirilmiştir,
 - Veri sızıntılarını sınıflandırmak için %88 doğruluk ile çalışan makine öğrenmesi ve yapay zeka modelleri eğitilmiştir,
- Kullanıcı hesap bilgileri veri sızıntılarından ayrıştırılarak ayrıca analiz edilmiştir;
 - Kullanıcı bilgilerinin gizliliğinin ve benzersizliğinin korunarak analizi için anonimleştirme yapılmıştır,
 - Parolalar üzerinde parola seçim davranışı ve hizmet sağlayıcısı parola politikasını analiz için yeni bir özellik çıkarımı uygulanmıştır,
 - Tez kapsamında toplanan veri sızıntıları içerisindeki kullanıcı bilgileri kullanılarak AuthInfo veri seti oluşturulmuştur,

- AuthInfo veri seti kullanılarak kullanıcıların parola seçimi davranışı ve parola seçiminde kullanılan karakterlerin analizi paylaşılmıştır,
- AuthInfo veri seti özellikleri ve analizleri kullanılarak bir vaka çalışması ile kaba kuvvet saldırısı nasıl iyileştirilebileceği açıklanmıştır,
- Kullanıcıların parola seçme davranışlarını saklamak için güvenli bir yöntem önerilmiştir,
- Parola seçim davranışını kullanarak parola seçiminde kullanılmak üzere yeni bir güvenlik politikası adımı önerilmiştir,

Bu katkılar ve sonuçlar ışığında veri sızıntılarının arama, yapılandırma ve sınıflandırılmasında asgari operatör müdahalesi ile çalışan Siber Tehdit İstihbaratı için Tarama Modeli geliştirilmiştir.

1.4 Organizasyon

Tez çalışmasının 2. bölümünde Siber Tehdit İstihbaratı, Mevcut Standartlar Modeller ve Araçlar anlatılmıştır. Kavramsal çerçevede ayrıca istihbarat verilerinin değerlendirilmesi için Kişisel ve Kurumsal veri kavramları açıklanmıştır. Kavramsal çerçevede son olarak Siber Tehdit İstihbaratı için Yeni Model Geliştirilmesi kapsamında kullanılan yöntemler ve ilgili çalışmalar ayrıntılı olarak açıklanmıştır. 3. bölümde serbest ve yapısal veriler içerisinden özellik çıkarmak için kullanılan yöntemler açıklanmıştır. Ayrıca CTI için geliştirilen yeni model ve modülleri ayrıntılı şekilde açıklanmıştır. Veri akış şemaları ve görsellerle model ayrıntıları ve modül ilişkileri verilmiştir. Geliştirilen modelin hem çıktısı hem de modelin yapay zeka modülünün eğitiminde kullandığımız AuthInfo veri seti ve CTI veri setleri açıklamıştır. 4. bölümde toplanan sızıntı verilerinin analizi ve görselleştirilmesi yapılmıştır. Veri sızıntılarının sınıflandırılması bölümünde 6 farklı veri seti ile 5 farklı modelin ikili ve çoklu sınıflandırma performansları sunulmuştur. Bulgular kapsamında son olarak kullanıcı adı- parola dosyaları üzerinden parola seçim davranışı analizi, vaka çalışması olarak parola seçim davranışı kullanarak kaba kuvvet saldırısı gerçekleştirilmesi ve parola davranışı duyarlı parola politikası yönteminin sonuçları incelenmiştir. Tez çalışmasının 5. bölümde elde edilen tüm sonuçlar tartışılırken 6. bölümünde sonuç ve öneriler verilerek tez tamamlanmıştır.

2. KAVRAMSAL ÇERÇEVE

Siber Tehdit İstihbaratı (Cyber Threat Intelligence - CTI) kapsamında veri sızıntıları ve veri ihlalleri tespit, işleme ve sınıflandırma çalışmalarının yapılması siber güvenlik konusu olduğu kadar veri mühendisliği, doğal dil işleme ve yapay zeka çalışmalarının da birlikte yapılmasını gerektirmektedir [22]. Öncelikli olarak CTI'nın kapsamında yapılan saldırı öncesi, saldırı anı ve saldırı sonrası çalışmalar özetlenmiştir. Veri ihlali ve veri sızıntıları genellikle siber tehdit vektörlerinin neticesi olarak karşımıza çıkmaktadır. Tez kapsamında saldırı öncesi ve saldırı anı siber istihbarat çalışmaları kapsam dışı bırakılmış, doğrudan Veri İhlal İstihbaratı üzerine yoğunlaşmıştır. Yapısal ve serbest metin olarak açık ve özel web alanlarında bulunan herkese açık veriler üzerinde çalışılmıştır. Veri sızıntılarının sınıflandırılması için kullanılan alt sınıfların Kişisel Veri ve Kurumsal Veri kavramları açıklanmıştır. Veri kazıma, anlamlandırma ve özellik çıkarımı işlemleri için kullanılan Düzenli İfadeler (REGular EXpressions - REGEX), Doğal Dil İşleme (Natural Language Processing - NLP), Adlandırılmış Varlık Tanıma (Named Entity Recognition - NER) ve Optik Karakter Tanıma (Optical Character Recognition - OCR) yöntemleri açıklanmıştır. Veri sızıntılarının sınıflandırma işlemleri için kullanılan Rastgele Orman, Destek Vektör Makinesi (SVM), Naive Bayes, K-En Yakın Komşu (KNN) ve Yinelemeli Sinir Ağı (RNN) yöntemleri açıklanmıştır. Son olarak diğer araştırmacıların yapmış olduğu Siber Tehdit İstihbaratı ve Veri İhlal İstihbaratı ile ilgili ilgili çalışmalara yer verilmiştir.

1.1. Siber Tehdit İstihbaratı

Siber Tehdit İstihbaratı, siber güvenlikle ilgili bilgilerin toplanması, analiz edilmesi ve bunlara göre hareket edilmesi olarak tanımlanabilir. Bir siber saldırı hakkında farklı bilgi parçalarını toplama, gerçekleştiği bağlamı ve bu bağlamın ne anlama geldiğini belirleme sürecidir [23]. Hayatımızın büyük bir kısmı teknolojiye bağlı olduğundan, siber güvenlik olaylarının sayısı artmakta ve hem bilgi sistemleri hem de siber saldırılar doğası gereği daha karmaşık hale gelmektedir. CTI, bir saldırının zamanı ve yeri, hangi tür kötü amaçlı yazılımın kullanıldığı, özet (hash) değerleri, hangi platformların bu saldırıdan etkilendiği veya bu saldırıya karşı savunmasız olduğu, IP adresleri gibi uzlaşma göstergeleri (Indicator of Compromise - IoC),

kimlik avı e-postaları (phishing e-mails) gibi saldırı vektörleri hakkında bilgiler içerebilir. CTI' nin yapıcı kullanımı, siber saldırıların hem tespit edilmesinde hem de önlenmesinde yardımcı olabilmektedir.

Tehdit istihbaratı düz yazı veya standartlaştırılmış çeşitli formatlardan oluşabilir. CTI' yı paylaşmak, kuruluşların işbirliği yoluyla siber savunmalarını geliştirmelerine, tehdit ortamını daha iyi anlamalarına ve etkilerini azaltmak için yeni tehditlere verilen yanıtları koordine etmelerine yardımcı olabilir [24]. Verimli bilgi paylaşımının önündeki engeller, kuruluşların rekabetlerine yardımcı olma endişesi duymaları veya gizli verileri gizli olmayan verilerden ayırmakta zorlanmaları, misilleme yoluyla daha büyük hedefler haline gelme endişelerinin yanı sıra CTI formatı ve paylaşımı için standartların bulunmamasından kaynaklanmaktadır [25]. Operatör desteği ile genellikle manuel olarak yürütülen süreçlerde yeterli insan, finansal kaynak ve firmalar organizasyon gerektiren bir süreç olarak karşımıza çıkmaktadır.

Açık Kaynak İstihbaratı (Open Source Intelligence - OSINT), halka açık kaynaklardan herkesin erişebileceği ham verilerin toplanmasından oluşan bir istihbarat disiplini [26]. İşlenerek ve analiz edilerek ham veriler eyleme geçirilebilir ve istihbarat bilgisine dönüştürülür veya başka bir deyişle, istihbarat faaliyetinin çok önemli bir parçası olan bir istihbarat ürününe dönüştürülebilir [27]. OSINT terimi ilk olarak ordu ve istihbarat topluluğu tarafından ulusal güvenlik konularında stratejik olarak önemli, kamuya açık bilgileri toplayan istihbarat faaliyetlerini belirtmek için kullanılmıştır [28]. Soğuk savaş döneminde casusluk, insan kaynakları (Human Intelligence - HUMINT) [29] veya elektronik sinyaller (Signal Intelligence SIGINT) [30] yoluyla bilgi elde etmeye odaklanılmış ve 1980'lerde OSINT, ek bir istihbarat toplama yöntemi olarak öne çıkmıştır. İnternetin, sosyal medyanın ve dijital hizmetlerin ortaya çıkışıyla, OSINT, istihbarat toplamak için çok sayıda kaynağa erişim sağlar. Sosyal Medya İstihbaratı (Social Media Intelligence - SOCMINT), OSINT'in bir alt disiplini olarak kabul edilir. SOCMINT, sosyal medya platformlarından bilgi toplanmasını ve analiz edilmesini sağlayan teknikler, teknolojiler ve araçlar olarak tanımlanabilir. SOCMINT, belirli kişiler, gruplar, olaylar veya herhangi bir sayıda başka bir hedef hakkında bilgi edinmek için devlet veya devlet dışı aktörler tarafından kullanılabilir [31]. OSINT ve SOCMINT kaynakları siber saldırganlar tarafından genellikle propaganda ve suçun yaygınlaştırılması için uygun enstrümanlar (ücretsiz anonim dosya ve metin paylaşım platformları gibi) sağlamaktadır. Bu nedenle doğrudan saldırgan kişi ya da grupları izleme ve takip etme imkanı sağlamaktadır. CTI ile ilgili kavramların açıklanmakta ve önemli modeller

ve araçlar kronolojik olarak sunulmaktadır. Açıklanan modeller ve kavramlar CTI toplulukları ve hizmet veren ticari ve devlet kurum ve kuruluşların yaptıkları çalışmalar ve bu tezde sunulan çalışmanın temelini oluşturmaktadır.

1.1.1. CTI Mevcut Modeller

Siber Tehdit İstihbaratı Birleşik Krallık Ulusal Siber Güvenlik Merkezi (The UK National Cyber Security Centre) tarafından stratejik (strategic), taktiksel (tactical), operasyonel (operational) ve teknik (technical) olmak üzere dörde ayırır [32]. Stratejik tehdit istihbaratı, risk ve gerçekleşme sıklığı (olasılık) gibi üst düzey kavramlarla ilgili olmakla birlikte ulusal kuruluşlar ve güvenlik endüstrisi uzmanları gibi üst düzey kaynaklardan gelir. Operasyonel tehdit istihbaratı, bir saldırganın kimliği veya bir saldırının ne zaman gerçekleşeceği gibi belirli saldırılarla ilgili bilgilerdir. Saldırıları tetikleyebilecek olaylar hakkındaki bilgilerden veya çevrimiçi etkinliğin izlenmesinden elde edilebilir. Taktiksel tehdit istihbaratı, tehdit aktörlerinin taktikleri, teknikleri ve prosedürleri (Tactics, Techniques, and Procedures - TTP) hakkında bilgidir ve raporlardan, adli tıp ve kötü amaçlı yazılım analizi yoluyla toplanabilir. Teknik tehdit istihbaratı, bir saldırgana ait bilgilerin ayrıntılarını içerir. Kötü amaçlı yazılım imzaları, IP adresleri, alan adları, dosya ve kayıt defteri gibi etkinlikler olabilir. Yararlı bilgileri analiz etmek çok miktarda ayrıntı içermesi nedeniyle zordur.

Lockheed Martin 2011 yılında, gelişmiş bir kalıcı tehdidin başarılı bir saldırı gerçekleştirmek için geçeceği aşamalar hakkında genel bir bakış sunan Siber Ölüm Zinciri modelini (Cyber Kill Chain) [33] önermiştir. Model, saldırganın başarı kazanmasını önlemek için saldırının "etkisizleştirileceğini (killed)" yedi aşamadan oluşur. Keşif (Reconnaissance), Silahlanma (Weaponization), İletme (Delivery), Sömürme (Exploitation), Yükleme (Installation), Komuta ve Kontrol (Command and Control - C2), Eylem (Actions on Objectives) aşamaları sırasıyla Siber Ölüm Zincirinin aşamalarıdır. Siber Ölüm Zinciri, çeşitli yaygın siber saldırıların aşamalarını ve buna bağlı olarak bilgi güvenliği ekibinin saldırganları önleyebileceği, tespit edebileceği veya durdurabileceği noktaları özetler. Siber Ölüm Zinciri'nin amacı, saldırganların bir saldırıyı gözetlemek ve planlamak için önemli ölçüde zaman harcadıkları, gelişmiş kalıcı tehditler (Advanced Persistent Threat - APT'ler) olarak da bilinen karmaşık siber saldırılara karşı savunma yapmaktır.

Elmas Saldırı Analizi Modeli (Diamond Model) 2013 yılında önerilmiş ve o zamandan beri siber güvenlikte yaygın kullanılan bir model haline gelmiştir [34]. Analistlerin kötü amaçlı

aktiviteyi nasıl değerlendirip anladıklarını açıklar ve izinsiz giriş analizi yapmak için resmi bir yöntem tanımlar. Bu yöntemde "olay" - düşman, altyapı, yetenek ve kurbanın bir bileşimi - izinsiz giriş faaliyetinin temel atomik unsuru olarak kabul edilir.

Emek/Acı Piramidi (Pyramid of Pain), 2014'te David Bianco tarafından bir blog yazısı olarak yayınlanmış ve güvenlik alanındaki çalışmalarda sıkça alıntılanan bir model haline gelmiştir [35]. CTI'nın tehdit algılama operasyonlarında etkili bir şekilde kullanılması için kavramsal bir model olan Emek/Acı Piramidi' ne (Pyramid of Pain) [36] göre özet (hash) bilgisinin değişikliği saldırganların en az çaba ve zaman harcadığı ayak izi (footprint) olarak değerlendirilmektedir [37]. Çok ilgi gören ve sıklıkla atıfta bulunulan bir diğer model, Ryan Stillions tarafından ilk kez 2014'te sunulan Tespit Olgunluk Düzeyi Modeli 'dir (Detection Maturity Level - DML) [38]. DML modeli, siber saldırıların istihbarata dayalı tespitinde dokuz olgunluk seviyesinden oluşur. Seviye ne kadar yüksek olursa, saldırıları tespit etmek için tehdit istihbaratı o kadar iyi uygulanabilmektedir. Düşük seviyeler teknik olarak daha spesifik, yüksek seviyeler ise daha soyuttur.

1.1.2. CTI Araçlar ve Platformlar

Zararlı Yazılım Bilgi Paylaşım Platformu (Malware Information Sharing Platform -MISP) [39], 2011 yılında Lüksemburg Bilgisayar Olay Müdahale Merkezi'nden (Computer Incident Response Center Luxembourg - CIRCL) bir grup geliştirici ve diğer katılımcılar tarafından ücretsiz yazılım/açık kaynak olarak geliştirilmiştir. MISP, özel ve kamu sektörlerinde tehdit göstergelerinin, tehdit istihbaratının paylaşıldığı bir platformdur. MISP, kuruluşların tehdit istihbaratı, göstergeler, tehdit aktörü bilgileri veya MISP 'de yapılandırılabilir her türlü tehdit bilgilerini paylaşmasına olanak tanır. MISP kullanıcıları, mevcut kötü amaçlı yazılım veya tehditler hakkında ortak bilgilerden yararlanır. Bu güvenilir platformun amacı, hedefli saldırılara karşı kullanılan önlemlerin geliştirilmesine yardımcı olmak ve önleyici eylemler ve tespitler oluşturmaktır.

MITRE ATT&CK (Adversarial Tactics, Techniques and Common Knowledge - Saldırgan Taktikler, Teknikler ve Yaygın Bilgi) çerçevesi [40], saldırganların taktikleri ve tekniklerini içeren bir bilgi tabanıdır. Gerçek dünyada meydana gelen ataklar gözlemlenerek oluşturulmuştur. Bu framework ile saldırganların bir sonraki hamleleri tahmin edilebilmektedir. Saldırganların yapabileceği saldırı eylemlerini gösteren, sistemlere nasıl sızdıkları, sızma işlemlerini gerçekleştirirken nasıl bir yol izledikleri konusunda taktik, teknik

ve prosedürleri içeren bir matris mevcuttur. Saldırganların davranışlarını kategorize etmek için matris oluşturulmuştur. Çünkü saldırganlar sürekli gündemde olan güvenlik yöntemleriyle tespit edilmemenin yollarını bulmaktalar ve dolayısıyla defansif tarafta bulunanlar daha önceki defansif yaklaşımlarını değiştirmek zorunda kalmaktadırlar. MITRE ATT&CK'i avantajlı hale getiren şey, tüm taktik, teknik ve prosedürlerin gerçek dünyadaki gerçek saldırı gruplarının gözlemlerine dayanmış olmasıdır. Bu grupların çoğu aynı teknikleri kullanır. Siber saldırı gruplarının sistemlere saldırırken kendi oyun kitapları vardır ve yeni üyeleri hızlı bir şekilde üretken hale getirmek için bu oyun kitabını kullanırlar. Windows, Linux, MacOS işletim sistemleri ile ilgili teknikleri içeren matris “Enterprise” olarak adlandırılmıştır. Saldırganların ağ veya işletim sistemi üzerinden saldırmadan önce uyguladıkları teknik ve taktikler içeren matris “Pre-Att&ck” olarak isimlendirilmiştir. Mobil cihazlar için oluşturulan matris “Mobile Att&ck” ve endüstriyel kontrol sistemlerini (Industrial Control System - ICS) hedefleyen operasyonlarda kullanılabilecek teknikleri içeren matris "ICS" olarak isimlendirilmiştir.

Yarı Otomatik CTI [41] projesi 2017'de başlatılmış ve CTI'nın kullanılmasına, analizine, zenginleştirilmesine ve paylaşılmasına olanak tanımaktadır. 2019 ilkbaharında OpenCTI [42] yayınlanmıştır. Platform, ACT projesi ile benzer fikirlere sahip olup, güçlü sorgulama olanaklarına sahip, kaynak kombinasyonunun işlendiği ve veri modelinin bağlantılı verileri etkinleştirdiği grafik tabanlı bir platform oluşturmaktadır. Seçilen açık kaynak lisansı, ACT platformundan farklıdır. Kullanımda olan ancak kapalı kaynak, abonelik için ödemeye tabi olan SocRadar, CYTHREAT, U.S.T.A. ve BrandDefence gibi bir dizi tehdit istihbarat platformları bulunmaktadır.

1.2. Veri İhlal İstihbaratı

Veri ihlali (data breach), hassas, korunan veya gizli verilerin yetkisiz bir kişi tarafından kopyalandığı, iletildiği, görüntülendiği, çalındığı veya kullanıldığı bir güvenlik olayıdır [43]. Veri ihlalleri genellikle, geniş kapsamlı etkileri ile bilgi güvenliği [44] için önemli olduğu düşünülen hassas verilerin dış taraflara büyük ölçekli olarak yayılmasını içerir. Veri sızıntıları hassas, korunan veya gizli veriler, sağlık bilgileri, doğrudan ya da dolaylı kişisel tanımlanabilir bilgiler, ticari sırlar veya fikri mülkiyet ve kişisel mali veriler içermektedir [45]. Zaman içinde hükümet, sağlık, finansal hizmetler, sigorta, sosyal medya ve daha fazlası dahil olmak üzere çeşitli sektörlerde veri ihlalleri meydana gelmiştir. [46].

Gizli verilerin kuruluş içinde veya dışında yetkisiz bir kişiye ifşa edilmesi, veri sızıntısı olarak adlandırılır [47]. Bir kişi tarafından yanlışlıkla veya kasıtlı olarak yapılabilir. Bilgi teknolojileri (BT) ürünlerinde siber güvenlik ihlalleri yapılmakta ve birçok veri sızıntısı yaşanmaktadır. Veri sızıntılarına, bilgi sistemlerinde bulunan zafiyetler neden olabildiği gibi doğrudan ya da dolaylı olarak insan zafiyetlerini sömüren sosyal mühendislik saldırıları da neden olabilmektedir. Bu veriler, kullanıcı adları ve parolalar olabildiği gibi şirketlerin hatta ülkelerin gizli kalması gereken kıymetli verileri de olabilmektedir. Veri sızıntıları, kuruluşları maddi ve manevi kayıplar vermesine neden olur [47]. Ticari kuruluşlar için ciddi güven ve itibar kaybına sebep olabileceği gibi hizmet verilen ülkeye göre değişen hukuki yaptırımları da olabilmektedir. Bu veri sızıntıları nedeniyle ticari kuruluşlara veri gizliliğinin ihlal edilmesi gerekçesiyle para cezaları da uygulanabilmektedir. Çalışanlar ile yapılan ağır gizlilik anlaşmaları, verilen güvenlik eğitimleri ve bilgi sistemlerinin penetrasyon testleri ile siber risklerin yönetilmesine rağmen güvenliğin tamamen sağlanması mümkün değildir. Bir risk yönetim süreci olarak tarif edilen siber güvenliğin yönetilebilmesi için veri sızıntılarının paylaşıldığı veya satıldığı ortamların taranması önemlidir. Tespit edilen veri sızıntısı durumlarında yasal regülasyonlar (hizmet veren kuruluşun bulunduğu ülke veya hizmet verdiği ülkenin regülasyonları) uyarınca alınması gereken yasal ve teknik önlemler konusunda hukukçulardan ve siber güvenlik uzmanlarından danışmanlık hizmeti alması gerekmektedir. Bu durum birlikte çalışmayı gerektiren ayrı uzmanlık alanları için yüksek maliyetlerle ekip kurulması ya da hizmet alınmasını beraberinde gerektirmektedir.

1.2.1. İstihbarat Verilerinin Değerlendirilmesi

Veri ihlalleri, doğrudan veya dolaylı olarak dijital ve analog sistemlere izni olmayan kişilerin eline geçmesi süreçleridir. Veri ihlallerinin ortaya çıkmasında tespit ve sınıflandırılması problemleri ortaya çıkmaktadır. İşlenen kişisel verilerin büyük hacmi, kişisel verilerin yanlış uygulanması, yanlış kullanılması ve yanlış işlenmesi gibi önemli riskleri beraberinde getirmektedir [48]. Yasal mevzuat içerisinde veri türleri belirlenmiş ve tanımlanmış olsa da belirli bir çerçevesi bulunmayan kurum ve kuruluşlara göre tanımı ve önemi değişen bilgilerde bulunmaktadır. Ayrıca bir veri birden fazla etikete sahip olabilmektedir. Bu durum da etiketlerin önem sırası ya da sınıflandırıcı tarafın belirlediği önceliği değerlendirilmektedir.

1.2.1.1. Kişisel Veri

24 Mart 2016 tarihli ve 6698 sayılı Kişisel Verilerin Korunması Kanunu'na (KVKK) göre kimliği belirli veya belirlenebilir gerçek kişiye ilişkin her türlü bilgi kişisel veridir [49]. Ayrıca, kişisel veriler çerçevesinde yalnızca kişinin kişisel verileri değil, ailesinin kişisel verileri de değerlendirilmektedir. Bu kapsamda kişisel veriler "bir kişinin kişisel, mesleki ve ailevi özelliklerini yansıtan, o kişiyi diğer kişilerden ayıran ve niteliklerinin açıklanmasına olanak sağlayan" doktrinde her türlü bilgi olarak tanımlanabilir [50]. Kişisel veri, gerçek bir kişiye ilişkin doğrudan ya da dolaylı olarak kimliği belirlenebilen nitelikteki her türlü veridir. Kimliği belirlenebilir kişi, fizyolojik, fiziksel, zihinsel, ekonomik, kültürel veya sosyal mensubiyetine özgü bir veya daha fazla faktörün yanı sıra doğrudan veya dolaylı olarak tespit edilebilen gerçek kişidir [51]. Aşağıda kişisel veri olarak nitelendirilen veri türleri ayrıntılı olarak **Tablo 2.1**'de verilmiştir:

Tablo 2.1: KVKK kapsamında kişisel veri kabul edilen bazı veri türleri.

Kişisel Veri Türleri			
Ad Soyad	Doğum Tarihi	Doğum Yeri	Cinsiyet
Kimlik Numarası	Sosyal Güvenlik Numarası	Pasaport Numarası	Cinsel Yönelim
İkamet Adresi	Sabit Telefon Numarası	E-posta Adresi	Parolalar
Rumuz, Kullanıcı Adı	Cep Telefon Numarası	Taşıt Plakası	Pin Numaraları
Sendika, Siyasi görüş	Sosyal Medya Hesapları	IP, MAC adresi	Kredi Kartı Bilgileri
Aile Bilgileri (Anne Adı, Baba Adı)	Sağlık Bilgileri (Tahlil Sonuçları, Kan Grubu, Geçirilen Hastalıklar, Kullanılan İlaçlar)	Yüz ve Bedenini tanımlayan fotoğraflar	Biyometrik verileri (Parmak izi, göz kornesi)

Ad ve soyad, medeni durum, kimlik numarası, doğum yeri, doğum tarihi gibi kimlik bilgileri ve yara izi, ben, dövme gibi vücuttaki diğer ayırt edici işaretler, görüntü, fotoğraf, telefon numarası, banka hesabı/IBAN numarası, kredi kartı numarası, emeklilik, şirket kaydı, sosyal güvenlik numarası, internet protokolü (IP) adresi, e-posta adresi, sosyal medya hesap bilgileri, ilişkili şifreler, güvenlik kodları, el yazısı ve/veya elektronik imzalar, iletişim ve posta veya sosyal ağlar yoluyla yapılan paylaşımlar kişisel veriler olarak nitelendirilmektedir [52, 53]. Ayrıca ırk, etnik köken, siyasi düşünce, inanç, din, mezhep veya diğer inançlar, kılık ve kıyafet, kulüp, vakıf veya dernek üyeliği, sağlık, cinsel hayat, sabıka geçmişi ve güvenlik tedbirleriyle ilgili veriler (kayıt ve bilgiler) , parmak izi, retina görüntüsü, DNA özelliği vb. özel (hassas) kişisel veriler olarak tanımlanmıştır (KVKK.m.6/I) [53, 54].

KVKK’de de kişisel veri ve kişisel veri işlemenin tanımı yapılmıştır. Buna göre, kişisel veri, “kimliği belirli veya belirlenebilir gerçek kişiye ilişkin her türlü bilgiyi”; kişisel verilerin işlenmesi ise, “kişisel verilerin tamamen veya kısmen otomatik olan ya da herhangi bir veri kayıt sisteminin parçası olmak kaydıyla otomatik olmayan yollarla elde edilmesi, kaydedilmesi, depolanması, muhafaza edilmesi, değiştirilmesi, yeniden düzenlenmesi, açıklanması, aktarılması, devralınması, elde edilebilir hâle getirilmesi, sınıflandırılması ya da kullanılmasının engellenmesi gibi veriler üzerinde gerçekleştirilen her türlü işlemi” ifade etmektedir.

1.2.1.2. Kurumsal Veri

Kurumsal veri bir kuruma ait ya da kurum tarafından üretilmiş gizli ya da halka açık verilerin tamamı olarak değerlendirilebilir. Bu verilerin çeşitli yollarla bir kurumdan kasıtlı veya bilgisiz olarak çıkarılmasına kurumsal veri ihlali denir. Örnekler arasında NSA sızıntıları, Sony sızıntıları, HSBC sızıntıları, Panama Belgeleri ve Türk kimlikleri sayılabilir. Bir kurum tarafından oluşturulan ve işlenen verilerdir. Şahıs şirketlerinde kişisel bilgilerin kurumsal amaçlarla kullanılması neticesinde kurumsal nitelikli bilgiler olarak değerlendirmesi gerekmektedir. Kurumsal veri örnekleri **Tablo 2.2**’de verilmiştir. Ayrıca kişisel bilgilerin de kurum ya da kuruluş ilişkilendirmesi durumunda bu tez çalışması kapsamında “Kurumsal Veri” olarak kabul edilmiştir. Örneğin bir hizmet sağlayıcısından sızan kişisel verilerin yasal sorumluluğu (Veri Sorumlusu) hizmet sağlayıcı olduğu için kurumsal veri olarak nitelendirilmiştir. Veri sızıntısı ile Veri Sorumlusu arasında ilişki kurulamayan veri ise “Kişisel Veri” olarak nitelendirilmiştir.

Tablo 2.2: Kurumsal veri olarak kabul edilen veriler.

Kurumsal Veri Türleri		
Firma Ünvanı	Finansal Bilgileri (Cirosu, Karı)	Sosyal Medya Adresleri
Ticari Sicil numarası	Ortakları	Kurumsal e-Posta Adresi
Faks Numarası	Yönetim Kurulu Üyeleri	Çalışanlarının e-Posta Adresi
Vergi Dairesi	Kurum Resmi İnternet Sitesi	Yatırım Bilgileri (İhale, Borsa)
Vergi Adresi	Vergi Numarası	Sabit Telefon Numarası
Kurum Yöneticileri Kişisel e-Posta Adresleri	Çalışanların Kurumsal Cep Telefon Numaraları	Kurum Yöneticileri Sosyal Medya Adresleri

Veri sızıntısı, şirketler ve devlet kurumları gibi kurumsal operasyonlar için ciddi bir tehdittir. Hassas bilgilerin kaybı önemli itibar zedelenmesine ve finansal kayıplara yol açabilir ve hatta bir kuruluşun uzun vadeli istikrarını da etkileyebilir. Sızan bilgilerin yaygın türleri, çalışan/müşteri verilerinden, fikri mülkiyetten tıbbi kayıtlara kadar uzanır [55].

1.3. Siber Tehdit İstihbaratında Kullanılan Özellik ve Anlam Çıkarma Yöntemleri

Veri sızıntıları saldırganların veriyi elde etme yöntemleri ve yararlanma amaçlarına göre değişik formatlarda paylaşılsa da genel geçer bir veri sızıntı formatı bulunmamaktadır. Bu nedenle elde edilen veri sızıntılarını ayrıştırma ve anlamlandırmak için Düzenli İfadeler, Doğal Dil İşleme ve Adlandırılmış Varlık Tanıma gibi yöntemler kullanılmaktadır. Ayrıca sızıntı türü ve saldırı vektörüne göre özel programlama kabiliyetleri kullanılarak çeşitli ayrıştırma yöntemleri uygulanmaktadır.

1.3.1. Düzenli İfadeler

Düzenli İfade (Regular Expression-REGEX), örüntü eşleştirme problemlerini (pattern-matching problems) çözmek için güçlü ve zaman kazandıran bir mekanizmadır. Matematiksel teoriye dayalı olarak yazılır ve metnin bir kısmını açıklayabilir. Bilgisayar bilimi teorisinde, düzenli ifadeler, düzenli bir dilin kompakt bir metinsel temsilidir [56]. Düzenli ifadeler temel olarak kelime arama, metin düzenleme, dosya ayrıştırma, kullanıcı girişi doğrulama ve erişim kontrolleri gibi dize arama ve değiştirme görevlerinde kullanılır. Arama motorlarında [57], veri tabanı sorgulamada [58] ve ağ güvenliğinde [59, 60] daha gelişmiş kullanımlar görülebilir. Düzenli ifade, belirli metin dizelerini eşleştirmek için kullanılan bir modeldir. Python'da düzenli ifadeler şu şekilde kullanılır:

```
eşleşen_ifadeler = re.findall('<.*?>+', metin)
```

```
değiştirilmiş_metin = re.sub('<.*?>+', 'htmlKod', metin)
```

Python programlama dilinde re kütüphanesi findall metodu metin değişkeni içerisinde '<.*?>+' kuralı ile eşleşen tüm ifadeleri liste formatında döndürmektedir. Herhangi bir eşleşme bulunmaması halinde boş liste döndürmektedir. re kütüphanesi sub metodu metin değişkeni içerisinde geçen '<.*?>+' kuralı ile eşleşen tüm ifadeleri htmlKod kelimesiyle değiştirmektedir. Bu metod metni HTML (Hyper Text Markup Language - Hiper Metin İşaretleme Dili) kodlardan temizlemek için kullanılabilir.

1.3.2. İsimlendirilmiş (Adlandırılmış) Varlık Tanıma

İsimlendirilmiş (Adlandırılmış) Varlık Tanıma (Named Entity Recognition - NER), yapılandırılmış veya yapılandırılmamış metinden oluşan bir cümlemin kelimelerini ve bölümlerini önceden tanımlanmış varlıklar halinde sınıflandırmayı amaçlayan Bilgi Çıkarma

ve Toplamanın (Information Extraction and Gathering) bir alt kümesidir [61]. NER, metinden önemli nesneleri (örn. kişi, kuruluş, konum) tanımlamaya yönelik Doğal Dil İşleme (NLP) görevini ifade eder. NER, NLP'de dizi etiketleme adı verilen genel bir problem sınıfına aittir [62]. Özel adla gösterilen herhangi bir şey örneğin, bir kişi, bir yer veya bir kuruluş, adlandırılmış bir varlık olarak kabul edilir. Ayrıca, adlandırılmış varlıklar tarih, saat veya para gibi şeyleri de içerir. İsimlendirilmiş varlıkların işaretlendiği örnek bir metin aşağıda gösterilmiştir:

[*ORG* İstanbul Üniversitesi-Cerrahpaşa] bu [*TIME* Cuma'dan] itibaren [*LOC* Fatih] [*GPE* Cerrahpaşa Yerleşkesi] güzergâhı servislerine kayıt [*PER* Yusuf Selim]'e başvurmaları gerekmektedir.

Bu cümle, *ORGANİZASYON* olarak etiketlenmiş bir kelime, *KONUM* olarak etiketlenmiş bir kelime, *YER* olarak etiketlenmiş bir kelime *ZAMAN* olarak etiketlenmiş bir kelime ve *PER* (*kişi*) olarak etiketlenmiş bir kelime olmak üzere beş adlandırılmış varlık içermektedir.

İsim varlık tanımada, bir metin içinde uygun adlara karşılık gelen dizeler bulunmaya ve bu dizelerle gösterilen varlık türü sınıflandırılmaya çalışılır. Cümle bölümlendirmesindeki belirsizlikten dolayı zor olsa da hangi kelimelerin adlandırılmış bir varlığa ait olduğu ve hangilerinin olmadığı çıkarılmalıdır. Bazı kelimelerin bir kişinin, bir organizasyonun veya bir yerin adı olarak kullanılması gibi durumlarda bazı zorluklar ortaya çıkar. Örneğin, İstanbul bir yer adı olarak kullanılabilir veya “İstanbul Seyahat” ifadesinde bir firmayı veya “İstanbul Üniversitesi-Cerrahpaşa” ifadesinde kuruluşu ifade edebilir.

1.3.3. Doğal Dil İşleme

Doğal Dil İşleme (Natural Language Processing-NLP), insanlar ve bilgisayarlar arasında doğal dilde etkili iletişimi sağlayan teori ve yöntemleri içerir. Bilimsel bir çalışma alanı olarak NLP, insan (veya doğal) dilini bilgisayarlar tarafından yürütülebilecek komutlara çevirmek amacıyla bilgisayar bilimini, dilbilimi ve matematiği özümser [63]. NLP, metinsel verilerden anlamsal anlam çıkarmak için istatistiksel tekniklere dayanan algoritmik yaklaşımları kullanır. NLP'nin ilk uygulamaları, dil çevirisi gibi görevleri tamamlamada basit kalıpları tanımlamak için kullanılan kapsamlı, değişmeyen ve önceden belirlenmiş bir kural kitabı (örneğin, belirli bir ifadenin Fransızca ve İngilizce anlamları arasındaki kuralları eşleme) geliştirmeye dayanmaktadır. NLP'ye yönelik son yaklaşımlar, makine öğrenimini kullanmaları açısından daha dinamiktir [64]. NLP'nin avantajları arasında işlenebilen büyük miktarda metinsel veri,

eğitimli bir model kullanarak analiz hızı (saniyeler içinde yüzlerce sayfalık bilgi işlenebilir) ve ayırt edilemeyecek ince kalıpları tespit etme yeteneği yer alır. Dezavantajlar arasında, bazı NLP modellerini eğitmek için gereken veri kümelerinin çokluğu; günler sürebilen bazı NLP modellerini eğitmek için gereken süre; ve NLP model analizinin sonuçlarını insanlar tarafından okunabilir bir formata çevirmenin zorluğu yer almaktadır [65]. NLP içinde kullanılan yöntemlerden Word2Vec, her kelimelerin anlamsal durumunun sayısal değerlere çevrilmekte, bir yöne ve uzunluğa sahip vektörler olarak temsil etmekte ve bir noktaya göre n-boyutlu uzayda bir konumu temsil etmektedir. Word2Vec yöntemi Mikolov, Chen, Corrado ve Dean (2013) tarafından geliştirilmiştir ve büyük veri kümelerini yüksek kaliteli kelime vektörü temsillerine dönüştürmek için sığ sinir ağlarını kullanır [65].

1.3.4. Optik Karakter Tanıma

Optik Karakter Tanıma (Optical Character Recognition -OCR), otomatik tanımlama yapan ve karakterlerin optik bir mekanizma aracılığıyla otomatik olarak tanınmasını sağlayan bir tekniktir. OCR, basılı, el yazısı veya daktilo ile yazılmış metnin taranan görüntülerinin makine tarafından kodlanmış metne dönüştürülmesidir. Doğal süreçte insan gözü optik bir mekanizmadır; gözün gördüğü görüntü beynin girdisidir. OCR bir insanın okuma yeteneği gibi çalışır. OCR'nin amacı, alfanümerik veya başka bir harfe karşılık gelen görsel kalıpları sınıflandırmaktadır. OCR, bilginin hem bir insan hem de bir makine tarafından okunabilmesi gerektiğinde istenmektedir [66].

1.4. Veri Sınıflandırma Yöntemleri

Veri sızıntıları kullanılarak özellik çıkarımı ve NLP yöntemleri kullanılarak çıkarılan özelliklerin kural tabanlı sınıflandırması yeni tip verilerin öğrenmesine karşı duyarlı değildir. Her yeni veri ve özellik için kuralların güncellenmesi ihtiyacı doğmaktadır. Bu nedenle verilerin sınıflandırılmasında makine öğrenmesi ve yapay zeka yöntemleri kullanılarak hem mevcut sınıflandırma doğruluğu ve hassasiyeti artırılması hem de yeni tip dosyaların sınıflandırılmasına imkan vermesi için modeller eğitilmiştir. Tez çalışması kapsamında Rastgele Orman, Destek Vektör Makinesi, K- En Yakın Komşu, Naive Bayes ve Yinelemeli Sinir Ağı algoritmaları kullanılmıştır.

1.4.1. Rastgele Orman

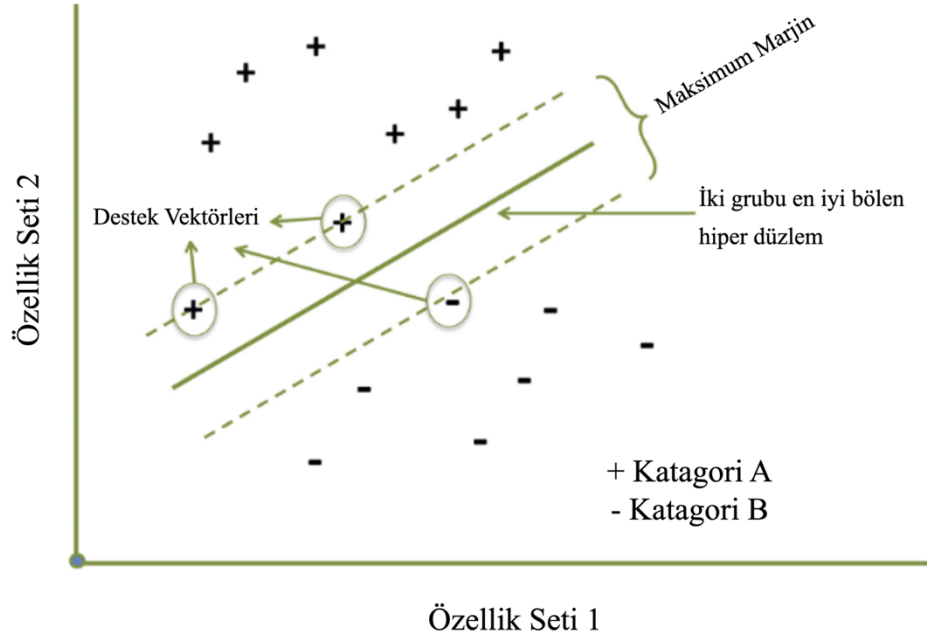
Rastgele Orman (Random Forest), denetimli bir makine öğrenimi algoritmasıdır. Hem sınıflandırmalar hem de regresyonlar için kullanılırlar. Birden fazla karar ağacı oluşturularak gerçekleştirilen sınıflandırma, regresyon ve diğer görevler için kullanılan Rastgele Orman, denetimli öğrenmeye dayalıdır ve bu algoritmanın en büyük avantajı hem sınıflandırma hem de regresyon için kullanılabilmesidir [67]. Rastgele Orman algoritması, mevcut diğer tüm sistemlerle karşılaştırıldığında daha iyi doğruluk sağlar ve en yaygın kullanılan algoritmadır. Karışık, eksik ve gürültülü veri kümeleri ile uğraşırken bile hızlı ve doğru sonuçlar ürettiği için en popüler yöntemlerden biridir [68]. Sonuçların doğruluğu, algoritmadaki ağaç sayısı ile doğru orantılıdır. Ağaç sayısı arttıkça doğruluğun da arttığı öne sürülmüştür [69].

Rastgele Orman modelinin arkasındaki mantık, birden fazla ilişkisiz modelin (bireysel karar ağaçları) bir grup olarak tek başına yaptıklarından çok daha iyi performans göstermesidir. Sınıflandırma için Rastgele Orman kullanıldığında, her ağaç bir sınıflandırma veya bir "oy" verir. Orman, "oyların" çoğunluğu ile sınıflandırmayı seçer. Regresyon için Rastgele Orman kullanıldığında, orman tüm ağaçların çıktılarının ortalamasını alır. Buradaki anahtar, bireysel modeller arasında, yani daha büyük Rastgele Orman modelini oluşturan karar ağaçları arasında düşük (veya hiç) korelasyon olmaması gerçeğinde yatmaktadır. Bireysel karar ağaçları hata üretebilirken, grubun çoğunluğu doğru olacaktır, böylece genel sonuç doğru yönde ilerleyecektir.

1.4.2. Destek Vektör Makinesi

Destek Vektör Makinesi (Support Vector Machine - SVM) tekniğinde, sınıflar arasında tekrarlanabilir bir hiperdüzlem tanımlanır [70]. Öğrenme yöntemleri genellikle her örneği özelliklerine göre doğru bir şekilde sınıflandırmayı amaçlar. Bu durum, özellikle çok uygun eğitim verileri için öğrenme modellerinden ziyade eğitim verilerinin ezberlenmesine ve sınıflandırıcının yeterince genelleme yapamamasına neden olmaktadır. SVM algoritması, eğitim setindeki sınıflar içindeki tüm örnekleri değerlendirerek, aralarındaki marjı en üst düzeye çıkaran bir yüzeyle ayırarak genelleştirme yeteneğini en üst düzeye çıkarır [71]. SVM karar fonksiyonunun eğitim süreci, **Şekil 2.1**'de görüldüğü gibi her iki sınıf etiketinin destek vektörleri arasındaki marjı maksimize etmektir. SVM en çok lineer olarak tercih edilse de, Polynomial, Gaussian (RBF) ve Sigmoid ve Precomputed olarak da uygulanabilir. Doğrusal SVM problemlerinin karmaşıklığı, kullanılan özellik sayısına bağlı olarak değişir. İki özellik için hiperdüzlem basitçe bir çizgi ile temsil edilirken, üç özellik olması durumunda

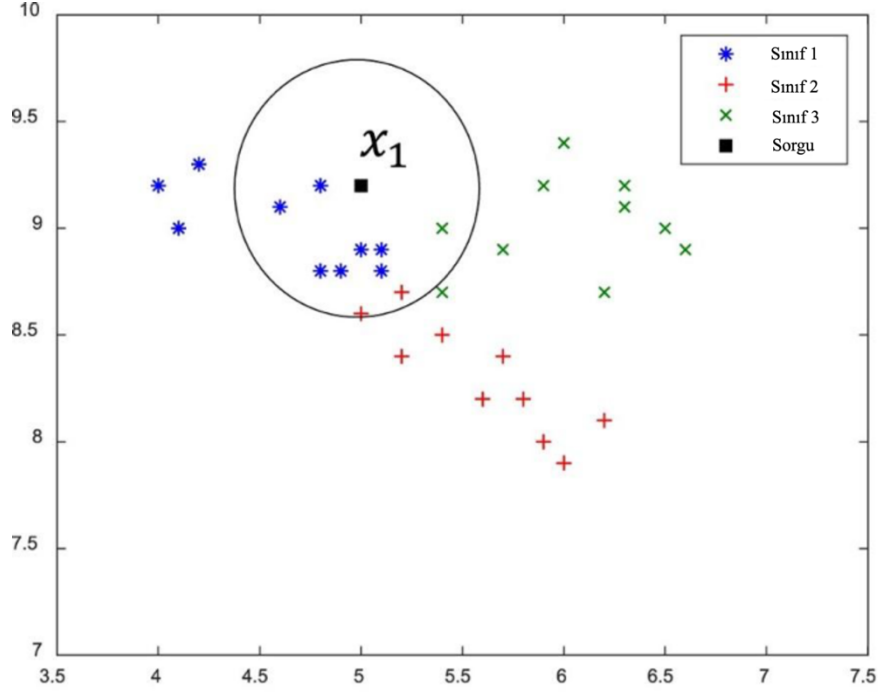
hiperdüzlem iki boyutlu bir düzleme karşılık gelir. SVM için kullanılan özelliklerin bu şekilde doğrusal olarak ayrılabilirliğini varsayarsak, ilgili sınıfın iki etiketini ayıran özelliklerin grafiği üzerinde kolayca düz bir hiper düzlem (doğrusal sınıflandırıcı) çizilebilir.



Şekil 2.1: Hiper düzlem örnekleri ile iki özelliği ayıran kategoriler arasındaki maksimum marjin [72].

1.4.3. K-En Yakın Komşu

K-En Yakın Komşu (K-Nearast Neighbour - KNN), sınıflandırma ve regresyon için kullanılan parametrik olmayan istatistiksel bir yöntemdir [73]. KNN, benzer özelliklere sahip örnekleri sınıflandırmak için bir vektör uzay modeli kullanır. Bilinmeyen sınıf örneklerinin olası sınıflandırmasını değerlendirmek için bilinen kategori örneklerine benzerliğini hesaplayabilir. Basit ve etkili olmasının yanı sıra aşamalı öğrenme için uygundur. Kümeleme, büyük veri ve çok etiketli öğrenme gibi birçok alanda kullanılmaktadır [74]. Klasik KNN algoritması yüksek zaman ve uzay karmaşıklığına sahiptir ve K değerini belirlemek zordur [75]. K değeri çok küçükse girişim hassasiyeti artar ve sınıflandırma doğruluğu düşer. K değeri çok büyükse ve veri seti dengesizse gürültülü örnekler en yakın komşular olarak seçilir ve sınıflandırma performansını olumsuz etkiler [76]. KNN sınıflandırma parametresi **Şekil 2.2**'de gösterildiği gibi öğrenmede seçilen özellikler sınıflandırma performansını doğrudan etkiler. Sınıfların aşırı değerleri için yanlış sınıflandırma riskleri de vardır.



Şekil 2.2: Üç sınıflı bir sınıflandırma problemi için KNN algoritması örneği [76].

1.4.4. Naive Bayes

Naive Bayes (NB), sınıflandırma için sınıf verildiğinde tüm özniteliklerin birbirinden bağımsız olduğu varsayımına dayanarak yeni bir örneğin bir sınıfa ait olma olasılığını çıkarır [77]. Tahmin yapabilmek için her sınıfa ait özelliklerin olasılıkları ile birlikte en yüksek olasılık değerine sahip sınıf seçilir. Çok sınıflı tahminler için Naive Bayes'in kullanımı kolaydır ve doğru sonuçlar üretir. Kategorik veri kümelerinde sayısal veri kümelerinden daha iyi performans gösterir [78].

Bayes teoremi, $P(c)$, $P(x)$ ve $P(x|c)$ 'den sonsal olasılığı, $P(c|x)$ hesaplamının bir yolunu sağlar. Naive Bayes sınıflandırıcısı, bir tahmincinin (x) değerinin belirli bir sınıf (c) üzerindeki etkisinin diğer tahmincilerin değerlerinden bağımsız olduğunu varsayar **Denklem 2.1** de görülmektedir. Bu varsayıma sınıf koşullu bağımsızlığı denir.

$$P(C|X) = P(X|C) * P(C) / P(X) \quad (2.1)$$

- $P(c|x)$ verilen tahminciye göre (*attribute*) sınıfın sonsal olasılığıdır (*target*).
- $P(c)$ sınıfın önceki olasılığıdır.
- $P(x|c)$ verilen tahmincinin olasılığıdır.
- $P(x)$ tahmincinin önceki olasılığıdır.

1.4.5. Yinelemeli Sinir Ağı

Yinelemeli Sinir Ağı (Recurrent Neural Network - RNN), sıralı verileri kullanan bir tür yapay sinir ağıdır. Dil çevirisi, NLP, konuşma tanıma ve resim alt yazısı gibi sıralı veya zamansal problemler için yaygın olarak kullanılır [79]. Yinelemeli Sinir Ağı, derin sinir ağı mimarisinin başka bir biçimidir. Genel olarak, Evrişimli Sinir Ağları (Convolutional Neural Network - CNN) değişmeyen işlevleri temsil etmek için kullanılırken, RNN'ler sıralı bir mimariyi (örneğin, cümleler, paragraflar, vb.) temsil eder. Hiyerarşik CNN'lere kıyasla RNN'lerin NER gibi dizi modelleme görevleri için daha uygun olacağı görülmektedir [80].

Yinelemeli Sinir Ağı (RNN), önceki adımın çıktısının mevcut adıma girdi olarak beslendiği bir Sinir Ağı türüdür. Geleneksel sinir ağlarında tüm girdiler ve çıktılar birbirinden bağımsızdır ancak bir cümlenin bir sonraki kelimesini tahmin etmek gibi durumlarda önceki kelimelere ihtiyaç duyulur ve bu nedenle önceki kelimelerin hatırlanması gerekir. Böylece bu sorunu bir Gizli Katman (Hidden Layer) yardımıyla çözen RNN'nin ana ve en önemli özelliği, bir dizi hakkında bazı bilgileri hatırlayan Gizli Katman (Hidden Layer) durumudur. RNN, hesaplamalarla ilgili tüm bilgileri hatırlayan bir "belleğe" sahiptir. Çıktıyı üretmek için tüm girdilerde veya gizli katmanlarda aynı görevi yerine getirdiği için her girdi için aynı parametreleri kullanır. Bu, diğer sinir ağlarının aksine parametrelerin karmaşıklığını azaltır.

1.4.6. Modellerin Performans Hesaplaması

Yapay zeka ve makine öğrenimi modellerinin başarımlarının ve tahmin performansının karşılaştırılması için doğruluk (accuracy), kesinlik (precision), duyarlılık (recall), F1-Skoru (F1-score) ölçüm yöntem kullanılmaktadır [81]. Bu metriklerin hesaplanırken bir örnek verinin pozitif veya negatif olma durumun sınıflandırılma ihtimalleri üzerinden anlatılmakta ve kullanılan bir semboller aşağıda verilmektedir. Çoklu sınıflandırma için de bu semboller artırılarak kullanılabilir.

Gerçek Pozitif (True Positive-*TP*), örnek pozitif sınıfına aitken sınıflandırıcı tarafından pozitif sınıflandırılması. Yanlış Pozitifleri (False Positive-*FP*), örnek pozitif sınıfına aitken sınıflandırıcı tarafından negatif sınıfta sınıflandırılmasıdır. Gerçek Negatifleri (True Negative-*TN*), örnek negatif sınıfına aitken sınıflandırıcı tarafından negatif sınıfta sınıflandırılmasıdır. Yanlış Negatifleri (False Negative-*FN*), örnek pozitif sınıfına aitken sınıflandırıcı tarafından negatif olarak sınıflandırılmasıdır.

Doğruluk (Accuracy-ACC), sınıflandırıcının sınıflandırma hatası türüne göre ağırlıklandırılmamış genel bir performans ölçüsü türüdür. Örneğin, "Kişisel Veri" türü dataset içerisinde eğitilen model tarafından ne kadarı doğru şekilde sınıflandırıldığını ifade eder. Ancak doğruluk parametresi tek başına kullanılması dengesiz veri setleri için potansiyel risktir [82]. Sınıfların veri seti içinde bulunma oranı dengeli olduğu durumlarda daha anlamlı bir ölçü birimidir. Model doğruluğun hesaplanması **Denklem 2.2**'de gösterilmiştir.

$$ACC = \frac{TP+TN}{TP+FP+TN+FN} \quad (2.2)$$

Kesinlik (Precision -PRC), gerçek pozitif sınıflandırılmış verilerin, model tarafından o sınıf olarak sınıflandırılan toplam veri sayısına oranını tanımlar. Örneğin, "Kurumsal Veri" olarak sınıflandırılan verilerin ne kadarı aslında kurumsal veri olarak etiketlendiğini ifade eder. Kesinlik, eğitilen modelin FP sayısını asgariye düşüğünü göstermek için kullanılır [82]. Model kesinliğinin hesaplanması **Denklem 2.3**'de gösterilmiştir.

$$PRC = \frac{TP}{TP+FP} \quad (2.3)$$

Duyarlılık (Recall-REC), sınıf başına doğru şekilde sınıflandırılan verilerin oranını tanımlamaktadır. Örneğin, "Potansiyel Veri" olarak etiketlenen verilerin ne kadarı doğru sınıflandırıldığını ifade eder. Duyarlılık, sınıflandırıcının FN tahmini başarısını göstermek için hesaplanır. Model duyarlılığı hesaplanması **Denklem 2.3**'de gösterilmiştir.

$$REC = \frac{TP}{TP+FN} \quad (2.4)$$

F1-Skoru ise bir modelin kesinlik ve duyarlılığı oranlarının birleştiren performans metriğidir. Kesinlik ile Duyarlılık metriklerinin harmonik ortalaması alınarak hesaplanmaktadır [83]. Modelin F1-Skoru hesaplanması **Denklem 2.4**'de gösterilmiştir.

$$F1 - Skoru = \frac{2*PRC*REC}{PRC+REC} \quad (2.5)$$

1.5. İlgili Çalışmalar

Çevrim içi hizmetlerin hızla gelişmesi hem bireyler hem de kurumlar için kolaylıklar sağlamaktadır. Ancak çevrimiçi hizmetleri kullanırken paylaşılan kişisel bilgiler artan güvenlik tehditleri nedeniyle kullanıcıların gizlilik endişelerini artırmaktadır. Son yıllarda,

veri ihlalleri sıklaştırmış ve dikkat çekici hale gelmiştir. ızan kişisel veriler, siber saldırı grupları tarafından çarpışma saldırısı (collision attack), kredi kartı dolandırıcılığı (credit card fraud) vb. gibi saldırılarda kullanılmaktadır. Bu nedenle, gün geçtikçe daha fazla araştırmacı veri sızıntısı konusuna odaklanmıştır. Birçok araştırmacı özellikle Veri Sızıntısı Tespiti (Data Leak Detection - DLD) [84, 85] ve Veri Sızıntısı Önleme (Data Leak Prevention - DLP) konularına odaklanmışlardır. Ağ güvenliği analistleri, mevcut ağ veri sızıntısı durumunu keşfedebilir ve anlayabilir. İleti dizilerinin otomatik olarak tanımlanması, hızlı bir şekilde veri kayıp tahmini sağlayabilir.

Yüzey Ağı (Surface Web) ve Karanlık Ağ (Dark Web), siber suç ekosisteminde önemli rol oynayan çok sayıda yeraltı forumları (underground forums) içermektedir. Yeraltı forumları, bilgisayar korsanlarının öğrenme, bilgi için iletişim, güvenlik açığı ifşası, araç değişimi ve siber suçlar için dağıtım merkezi faaliyetleri yürütmesi için popüler oluşumlardır. Birçok forum ayrıca kötü amaçlı yazılım ticareti, bilgi hırsızlığı ve diğer hizmetler için gizli işlemler sağlamaya adanmıştır. Yeraltı forumları sızan kişisel bilgilerin ticaretinde de önemli bir rol oynamaktadır [86]. Son veri ihlallerinin ve/veya zaten sızdırılmış bilgilerin paylaşımının yapılması neredeyse her forumda yaygındır. Hatta birçok bilgisayar korsanı, güvenliği ihlal edilmiş kuruluşlardan veya bireylerden veri alır almaz yeraltı forumlarında duyurular yayınlar. Bu nedenle çoğu zaman veri ihlalleri ilk önce yeraltı forumlarında ifşa edilir. Örneğin, 4 Eylül 2016'da tıklama başına ödeme yapan ClixSense web sitesi bir veri ihlali yaşamıştır. Veriler, saldırganlar tarafından çevrimiçi olarak yayınlanmış ve 28 Ağustos 2018'de "helen250" adlı bir kullanıcı, Çin yeraltı forumunda China Lodging Group'un 130 milyon kullanıcı verisini satmak için bir ileti dizisi yayınlamıştır. Toplam sızan veri sayısı 500 milyona ulaşmıştır [87]. Madhavan ve arkadaşları Deep Web taramasında kullanılacak web arama dizinine eklenmeye uygun URL'ler üreten URL'lerin tanımlanması için olası formdaki giriş kombinasyonlarının arama alanında etkin bir şekilde gezinen bir algoritma sunmuşlardır. Algoritmalarının etkinliğini doğrulayan kapsamlı bir deneysel değerlendirmeye de makalelerinde yer vermişlerdir [88]. Hsinchun ve arkadaşları belirledikleri anahtar kelimeler ile Dark Web üzerinde bulunan web sayfalarını tespit ve sınıflandırma çalışması yapmışlardır. Çalışma kapsamında aşırılık yanlılarının fikirlerini yayan ve yıkıcı faaliyetlerin yürütülmesinde kullanılan çok çeşitli bilgi kaynaklarından gelen tehditleri tespit etmek için, bilgi toplama ve analiz tekniklerini ve insan etki alanı bilgilerini entegre eden yarı otomatik bir metodoloji önermişlerdir [89].

IBM firması tarafından 3600 firma üzerinde yapılan anket çalışmasına göre katılımcıların %83 ü en az bir ihlal yaşadıklarını belirtmişlerdir [90]. Rapora göre veri ihlallerinde 2022 yılı içinde 4,35 milyon Dolar firmaların zarara uğradığı belirtilmiştir [90]. Firmalar sahip olduklarını verileri analiz ve sınıflandırmak için genellikle içeriğe dayalı analiz, bağlama dayalı analiz ve hibrit analiz olmak 3 yöntem kullanmaktadır[91]. İçeriğe dayalı analizde hassas veri taraması veri bekleme, kullanım ve aktarım aşamalarındayken veri içeriğini denetlenmektedir [91, 92, 93]. İçerik taraması, yanlışlıkla veri kaybına karşı etkili bir şekilde koruma sağlasa da, veri gizlemeye operasyonlarına (şifreleme, sıkıştırma, kodlama, karmaşıklıklaştırma gibi) karşı zayıftır [94, 95]. Bağlama dayalı analiz hassas verinin varlığını belirlemeye çalışmak yerine, esas olarak izlenen verilerle ilişkili meta bilgilerini veya verileri bağlamsal analizini (veriden, dosyadan, kaynaktan veya çevreden çıkarılan özellikleri) kullanılarak gerçekleştirir [96, 97, 98]. Hibrit analiz ise hem içeriği hem de bağlamı analiz eden hibrit yaklaşım bulunmaktadır [99]. Verilerin analizi ve sınıflandırılmasında çalışmalarda genellikle içeriğe dayalı analizler çoğunlukta olmasına rağmen son çalışmalarda bağlam ve hibrit analiz çalışmalarına da rastlanmaktadır.

Cheng ve arkadaşları kurumsal veri sızıntı tehditlerini kasıtlı ve kasıtsız, içeriden ve dışarıdan şeklide sınıflara ayırmışlardır [100]. Bir vaka çalışması üzerinde MapReduce aracı üzerinde N-Gram kullanarak email trafiğini sınıflandırmışlardır [100]. Shaikh ve arkadaşı bulut sistemlerinde veri sınıflandırmasını bulut üzerindeki dosya ve klasörlerin kullanım durumu ve erişim kontrolü kısıtlamalarına göre değerler tanımlamakta ve sınıflandırmaktadır [101]. Yayık ve arkadaşları Lojistik Regresyon ve Bert modelleri kullanarak verilerin örneklerinin ilk olarak konusunu çıkarmıştır [102]. Daha sonra konulardan din, adalet veya sağlık olan konuları kişisel veri etiketlenmiştir. Sınıflandırmayı desteklemek için LSTM temelli NER ve kural tabanlı kişisel veri eşleşmesi yöntemlerini de kullanmıştır [102].

Yong Fang ve arkadaşlarının yapmış olduğu çalışmada [103], veri ihlalleri ile ilgili konuları otomatik olarak tespit eden bir sistem sunulmuştur. Çalışma, Gizli Dirichlet Tahsisi (Latent Dirichlet Allocation - LDA) modeline dayalı özellik çıkarım yöntemi uygulanmıştır. Sistemin performansını iyileştirmek için, denetimli sınıflandırma algoritmaları ile karşılaştırılma yapılmış ve sınıflandırıcı için en iyi yöntem seçilmiştir. Çalışmada geliştirilen sistem aracılığıyla, deneysel veri setindeki veri ihlali zincirlerinin %92'sinden fazlası belirlenmiştir. Çalışmada geliştirilen sistemin yapısı üç bölüme ayrılmıştır: toplayıcı ve ön işlemci, özellik seçici ve sınıflandırıcı. Toplayıcı ve ön işlemci, ham verilerin toplanması ve işlenmesi için

kullanılmıştır. açık ve özel web üzerinde web crawler (tarayıcılar) ile forumlardaki gönderileri, yanıtlar ve üyeler (profilleri dahil), web tarayıcıları için hedef olarak belirlenmiştir. Toplanan verilerden özellik seçici, konu dağılımını ve istatistikleri elde etmek için kullanılmakta ve sınıflandırıcıya özellik vektörleri sağlamaktadır. Konu dağıtım özelliğinde LDA modeli ile, bir iş parçacığının veri ihlali ile ilgisi belirlenmiştir. Ayrıca Naive Bayes, Rastgele Orman, Lojistik Regresyon (LR), Destek Vektör Makinesi (SVM), K-En Yakın Komşu (KNN) sınıflandırma algoritmalarının performans metrikleri de karşılaştırılmıştır. Her algoritmanın optimum parametresini bulmak için Grid Search (Izgara Arama) yöntemi kullanılmıştır. Tarayıcıdan gelen gönderinin başlığı ve içeriğinden konu dağılımı ve istatistiksel nitelikten oluşan özellik vektörü çıkarılmakta sınıflandırıcı tarafından bu özellikler kullanılarak etkilenmektedir [103].

Siber veri ihlali vakalarının giderek artmasıyla birlikte, siber saldırıları araştırmak için tonlarca güvenlik günlüğünün manuel olarak taranması hataya açık ve zaman alıcıdır. İmza tabanlı derin arama çözümleri ise, yalnızca tehdit yapıları tam olarak sağlandığında doğru sonuçlar verir. Ortak saldırı modellerine sahip olan ve aynı saldırı araçlarını kullanan karmaşık siber tehditlerin hızla artmasıyla, zamanında soruşturma yapmak neredeyse imkansız hale gelmiştir. Düşmanın TTP'lerini saldırı hedeflerine ve tespit mekanizmalarına eşleyerek tehdit analizi sürecini otomatikleştirmeye ihtiyaç vardır. Umara Noor ve arkadaşları çalışmalarında [104], gözlemlenen saldırı modellerine dayalı olarak siber tehditleri tanımlayan yeni bir makine öğrenimi tabanlı çerçeve önermişlerdir. Çerçeve, iyi bilinen tehdit kaynaklarından çıkarılan tehditleri ve TTP'leri ilişkili algılama mekanizmalarıyla anlamsal olarak ilişkilendirir. Bu ağ daha sonra tehditler ve TTP'ler arasında olasılıksal ilişkiler oluşturarak tehdit oluşumlarını belirlemek için kullanılır. Çerçeve, bir TTP taksonomisi veri kümesi kullanılarak eğitilir ve performans, tehdit raporlarında bildirilen tehdit yapılarıyla değerlendirilir. Çerçeve, kayıp ve sahte TTP'ler durumunda bile saldırıları %92 doğrulukla ve düşük yanlış pozitiflerle verimli bir şekilde tanımlanır. Bir veri ihlali olayının ortalama tespit süresi, 45 tehdit ailesinden 133 TTP ile eğitilmiş bir ağ için 0,15 saniye olarak hesaplanmıştır. Venkatesh Gauri Shankar ve arkadaşları tarafından büyük verilerde ayıklama, sınıflandırma ve toplama için bir yaklaşım olarak "DataSpeak" sunulmuştur [105]. Yaklaşık 2 GB boyutundaki veri seti test edilmiş ve bu veri setine göre doğruluk yaklaşık %98 olarak elde edilmiştir. Yaklaşımda KNN algoritması kullanılmaktadır ve çapraz doğrulama, kesinlik, duyarlılık, karıştırma matrisi ve karşılaştırmalı analiz gibi istatistiksel yaklaşımlarla önerilen

“DataSpeak” çalışmasının doğruluğu ve yeteneği kontrol edilmiştir. Yaklaşım, mevcut KNN ve değiştirilmiş algoritmalar konseptine kıyasla minimum sürede verimli veri çıkarma ve sınıflandırma ile verilere hızlı erişim sağlar. Veri toplama ve sınıflandırma ile ilgili olarak, önerilen yaklaşım birçok araştırmaya kıyasla veri toplama ve sınıflandırmanın hızlı işlenmesini sağlar. Karışıklık matrisi ise, veri çıkarmanın karşılaştırmalı bir analizini oluşturur ve %98 doğruluk sonucu veri sınıflandırma ve toplama şeklinde bir çıktı üretmiştir.

Gizli verilerin kurum içinden veya dışından yetkisiz bir kişiye ifşa edilmesine veri sızıntısı denir. Veri sızıntısı, kuruluşlara büyük mali ve mali olmayan kayıplara neden olur. Veriler bir kuruluş için kritik bir varlık olduğu için tekrarlayan veri sızıntısı olayları giderek artan bir endişeye sebep olmaktadır. Antivirüs yazılımı, saldırı tespiti ve güvenlik duvarları gibi geleneksel bilgi güvenliği koruma yöntemleri bağımsız olarak başa çıkmak için giderek daha zor hale geldiğinden, veri sızıntısı giderek daha ciddi hale gelmektedir Jun Qian ve arkadaşları tarafından “2012 SocialArks veri ihlali” vaka çalışması baz alınarak veri ihlalleri tartışılmıştır [106]. SocialArks veri ihlali, "Tencent"’in sahibi olduğu SocialArks’a ait ElasticSearch veri tabanının yanlış yapılandırılmasından kaynaklanmıştır. Saldırı metodolojisi klasiktir ve beş yaygın Elasticsearch hatasını ve bu sızıntıların olasılıklarını ele almaktadır. ElasticSearch veri tabanının açık kaynak kimliği de birçok etik soruna neden olmakta, bu da herkesin ücretsiz olarak indirip kurabileceği ve neredeyse her yere kurabileceği anlamına geliyor. Bazı şirketler indirip dahili sunucularına kurarken, diğerleri indirip buluta (istedikleri herhangi bir sağlayıcıya) yükler. Ayrıca, Elasticsearch’ün Tencent Şirketi gibi bulut hizmeti şirketleri de vardır. Savunma çözümü, Elasticsearch sunucusunun nasıl optimize edileceğine odaklanmaktadır. Bu sorunu bir bütün olarak çözmek ve endüstri standartlarını iyileştirmek için, veri işleyen şirketlerin yasal standardizasyonu, gelecekte çok sayıda veri sızıntısını çözenin yollarından biri olabilir. Hamza Saleem ve Muhammad Naveed tarafından yapılan çalışmada [107], 10 ünlü veri ihlali için (Sony, Target, Yahoo, Anthem, OPM, RUAG, NSA, Carphone, Equifax, Zomato) adım adım veri ihlali iş akışları ve ardından 50 rastgele veri ihlali çalışması geliştirilerek saldırganların verileri nasıl ihlal ettiğine ilişkin bilgiler sistematik hale getirilmiştir. Verileri korumanın bir yolunun, saldırganların verileri nasıl ihlal ettiğini anlamak, veri ihlallerini bütünsel bir sorun olarak ele almak, gerçek dünyadaki saldırgan davranışını yakalayan tehdit modelleri geliştirmek, insan hatasına dirençli sistemler oluşturmak ve sonrasında hasarı en aza indirecek yaklaşımlar geliştirmek olduğunu söylemişlerdir. Veri ihlali analizine dayanarak, kuruluşların veri ihlallerini önlemek için

karşılması gereken gereksinimler geliştirilmiştir. Mevcut güvenlik teknolojilerinin hangi gereksinimleri karşıladığı açıklanmış ve veri ihlallerini engellemek için araştırma yönergeleri önerilmiştir. Suvendu Kumar Nayak ve Ananta Charan Ojha tarafından yapılan çalışmada [108], veri sızıntısı tespiti ve önleme sistemi tanımlanmıştır. Farklı veri durumları, konuşlandırma noktaları ve sızıntı tespit yaklaşımlarına dayalı olarak karakterize edilmiştir. Ayrıca, on yıllık mevcut araştırmalar dikkate alınarak sistematik bir literatür incelemesi takip edilmiş ve buradaki araştırma boşluklarının belirlenmesi için eleştirel bir analiz yapılmıştır. Veri sızıntısı tespiti ve önlenmesi alanında önemli araştırma yönergeleri önerilmiştir. Veri sızıntısı algılama modeli, verilerini çevrimiçi olarak herhangi bir üçüncü tarafla paylaşma eğiliminde olan çeşitli endüstriler, kuruluşlar, araştırma toplulukları veya kurumlar için oldukça faydalı olabilmektedir. Riya Naik ve Manisha Naik Gaonkar tarafından önerilen bir tasarımda [109], filigran algoritması kullanılarak bulut verilerinde veri sızıntısı ve kurcalamanın tespit edilmesini sağlar. Sunulan yaklaşımda amaç, aktarılan verilerin güvenliğini artırmak için hızlı yanıt koduyla filigran kullanarak verilerde tahribat olup olmadığını tespit eden ve böylece veri sızdıran(lar)ı bulan bir sistem uygulamaktır. Modelin amacı üç ana bölümden oluşmaktadır, birincisi, veri özelliklerini kullanarak aktarılan ve QR kodu oluşturan verilerden bilginin çıkarımını yapmaktır. İkincisi, frekans alanı tekniklerini kullanarak QR filigranını bulut verilerine yerleştirmektir. Üçüncü bölüm, mevcut veri özelliklerini kontrol ederek, filigranı çıkararak ve filigrandaki bilgileri aracının ayrıntılarıyla karşılaştırarak verileri sızdıran suçlu aracıyı tespit etmek ve herhangi birinin verilere müdahale edip etmediğini tahmin etmektir. Ayrıca, bu verilerin sızdırılıp sızdırılmadığını da algılayabilmektedir.

İnternetin hızlı gelişimi birlikte, akıllı uygulamalar, hassas verileri çalmak için rakiplerin birincil hedefi olmuştur. Bu durum kullanıcıların kimlik güvenliği, finansal güvenliği ve hatta can güvenliği için çok önemli bir tehdit oluşturmaktadır. Araştırma toplulukları ve endüstriler, veri sızıntısı tespiti ve önlenmesi için birçok Bilgi Akış Kontrolü (Information Flow Control - IFC) tekniği önermiştir. Ning Xi ve arkadaşları çalışmalarında [110], Mobil ve IoT (Internet of Things - Nesnelerin İnterneti) uygulamalarındaki veri sızıntısı tespiti ve önlenmesi için çeşitli IFC tekniklerini sistematik olarak anlamak için yeni bir çerçeve sağlayan bilgi akışı tabanlı bir savunma zinciri önermişlerdir. Bu alanda yapılan her çalışma teknik, performans ve sınırlama açısından kapsamlı bir şekilde incelenmiş ve araştırma zorlukları savunma zincirinin bütünlüğü dikkate alınarak belirtilmiştir. Android sistemi ve IOS sistemi, mevcut

akıllı telefon sistemi pazarındaki en yaygın iki sistemdir. Mevcut İnternet kaosu karşısında, her iki sistem de özellikle Android sistemi olmak üzere değişen derecelerde sorunlara maruz kalmaktadır. Kullanıcıların mahremiyetini korumak için araştırmacılar da Android sisteminin mahremiyetinin korunmasına odaklanmaya başlamışlardır. Yun Wang ve Limei Wang çalışmalarında [111], makine öğrenimini kullanan bir gizlilik algılama yöntemi önermektedir. İnternetin hızla gelişmesiyle birlikte sanal alemdeki suçlar giderek yaygınlaşmış ve İnternetteki mali suçlar, en yaygın mali suç türü haline geldi. Bu makale, esas olarak Android uygulama gizlilik sızıntısı izleme için İnternet finansal veri güvenliği ve ekonomik risk önlemeyi incelemektedir, ancak gizlilik kaynaklarının sınıflandırılmasında hala bazı eksiklikler bulunmaktadır. Gelecekte, gizlilik sızıntılarının daha doğru ve etkili bir şekilde izlenebilmesi için eksikliklerin optimize edilmesi hedeflenmektedir.

Veri güvenliği, tüm uygulama ve platformlardaki verileri koruyan veri şifreleme, tokenizasyon ve anahtar yönetimi uygulamalarını içerir ancak bunlarla sınırlı değildir. Amir Khaleghimoghaddam tarafından yapılan tez çalışmasında [112], denetimli makine öğrenimi teknikleri kullanılarak şifrelenmiş veriler analiz edildiğinde veri şifrelemede herhangi bir veri sızıntısının olup olmadığını araştırmayı hedeflemiştir. Araştırmacıların çoğu, veri sızıntısını incelemek için şifrelenmiş verilerin tersine mühendislik veya şifreleme algoritmalarına karşı kaba kuvvet saldırıları üzerinde çalışmışlardır. Bu araştırmadaki amaç ise, şifreli metne tersine mühendislik veya kaba kuvvet uygulamak değil, denetimli bir öğrenme algoritmasının şifreli metinde potansiyel olarak veri sızdırabilecek bir model tanımlayıp tanımlayamayacağını araştırmaktır. Bu amaçla, dört farklı veri kümesi üzerinde (Kaggle Sentiment 140, Kaggle Newsgroup, Kaggle Isis, AG News) dört denetimli öğrenme tekniği (Multinomial Naive Bayes, Random Forest, Logistic Regression ve Linear SVM) kullanılmış olup dört şifreleme algoritması analiz edilmiştir. Sonuçlara göre şifreleme algoritmaları güçlendikçe, veri sızıntısının azaldığı görülmektedir. İnternette saklanan hassas bilgilerin sürekli artan bir şekilde toplanmasıyla, veri ihlallerinin sonuçları geçici mali kayıplardan uzun süreli itibar zararlarına kadar değişmektedir. Bu tür olaylarla ilişkili risklerin farkına varılarak, gizlilik sızıntılarını tespit etmek ve önlemek için çok sayıda teknik geliştirilmiştir. Veri yığınlarının taranmasını içeren sızıntı tespiti ile bir çözümün fizibilitesi, yalnızca özel bilgilerin doğru bir şekilde tanımlanmasına değil, aynı zamanda verimli görev yürütmeye de bağlıdır. Sharleen Joy Y. Go ve arkadaşları çalışmalarında [113], Kişisel Tanımlanabilir Bilgiler (Personally Identifiable Information - PII) sızıntı tespiti amacıyla pratik bir mimari tasarlamak için SDN

(Software Defined Networking - Yazılım Tanımlı Ağlar) ve NFV (Network Function Virtualization - Ağ İşlevi Sanallaştırması) ilkelerinden yararlanmışlardır. Sistem, PII dedektörlerini ve yük dengeleyicileri başlatmaktan sorumlu bir sanal makinelerin konuşlandırıldığı bir OpenStack kümesinden ve internet trafiği oluşturmak için birden çok Mininet ana bilgisayarından oluşur. Çift SDN denetleyicileri, trafiğin bulut sunucuları arasında sorunsuz bir şekilde yönlendirilmesini de destekler. İndüklenen ek yükü ölçmek için, sistemin performansı değişen ağ yükleri ve paket işleme yoğunluğu altında analiz edilmiştir. Sonuçlara göre, ağ trafiği oldukça düşük olduğunda yönlendirmenin genel ek yükün daha büyük bir bölümünü oluşturduğu sonucuna varılmış olup bunun tersi, yeterince yüksek trafikte geçerli olur. PII tespitine ayrılmış ağ cihazı olmadığı için, bunun yerine uygulamanın ölçeklendirme özelliğini kapatarak bir tanesini taklit edilmiş ve performansını ölçekleme özellikli sistemle karşılaştırılmıştır. Sonuçlar, SDN/NFV tabanlı mimariyi kullanmanın gecikmeyi ve paket kaybını sırasıyla %85,1 ve %98,3 oranında etkili bir şekilde azalttığını ve verimi %8,41 artırdığını göstermektedir.

Kişisel veri işleme konusundaki kaygıların katlanarak artması nedeniyle, metin içeriklerinin hangi ortamda bulunduklarının yorumlanması için konularının bilinmesi bir zorunluluk haline gelmiştir. Belgenin konusu, içeriğindeki tüm verilerle birlikte sunulan kapsayıcı bir bilgi parçasıdır, bu nedenle KVKK'da tanımlanan kişisel veri kategorileri, işlendiği ortamdan kaynaklanan anlamsal etkiler nedeniyle farklılık gösterebilir. Apdullah Yayık ve arkadaşları çalışmalarında [114], Türkçe metin içeriklerindeki doğrusal olmayan ilişkileri anlayarak makine öğrenmesi modelleri yardımıyla konularını tahmin etmeyi ve modellerin KVKK uyum katkılarını tartışmayı amaçlamışlardır. fasttext, BERT (Bidirectional Encoder Representations from Transformers - Transformatörlerin Çift Yönlü Kodlayıcı Temsilleri) tabanlı ve BERTurk tabanlı modeller, bir konu sınıflandırıcı oluşturmak için genişletilmiş bir veri seti üzerinde tasarlanmış, eğitilmiş ve değerlendirilmiştir. Lojistik regresyon kullanan hızlı metin modeli ve dikkat katmanlarına sahip derin çift yönlü dönüştürücü BERT modelleri ile birçok deney yapılmıştır. Model performansları ve model tahminlerine ilişkin istatistiksel post-hoc test sonuçları analiz edilmiş, ardından modelin endüstriyel kullanımıyla birlikte konuşlandırılması tartışılmıştır. Sonuç olarak, makro ortalama F1-ölçüsü %94,73 metin modeli oluşturulmuş ve verimli bir şekilde üretime entegre edilmiştir. Aitor Garcia-Pablos ve arkadaşları çalışmalarında [115] klinik alanda veri gizliliğinin korunmasına ilişkin sorunları tanıtmış olup NLP alanındaki Transformatör tabanlı derin öğrenme mimarilerinin

gelişiminden faydalanarak Google'ın BERT modeline odaklanmış ve açıklamışlardır. Google tarafından paylaşılan önceden eğitilmiş çok dilli BERT modelini kullanarak BERT tabanlı dizi etiketleme yaklaşımıyla birkaç deney gerçekleştirilmiştir. Deneylerde MEDDOCAN 2019 paylaşılan görev veri seti ve NUBES-PHI adlı yeni bir İspanyol klinik raporları veri seti kullanılmıştır. Deneylerin sonuçları, NUBES-PHI'da, BERT tabanlı modelin, yalnızca sağlanan etiketli veriler üzerinde eğitilerek, herhangi bir uyarılma veya alana özgü özellik mühendisliği gerektirmeden diğer sistemlerden daha iyi performansa sahip olduğunu göstermektedir. BERT tabanlı modelde ise, diğer sistemlerden çok daha yüksek bir duyarlılık (recall) elde edilmiştir. Ayrıca, karşılaştırılan tüm sistemler için eğitim verileri aşamalı olarak azaltılmış ve bu veri kümesi üzerinde ek bir deney gerçekleştirilmiştir [115].

Bilgi açıktır ve bu bilgiler kişisel verileri araştırmak, kurumsal bir dünyada rakipleri yenmek, suçları çözmek ve hatta savaşları kazanmak için kullanılabilir [116]. Jyri ve arkadaşı çalışmalarında [117], açık kaynak istihbaratını (OSINT) ve büyük veri analitiğini (BDA), siber keşif ve kişisel güvenliğin bir parçası olduğunu vurgulayarak analiz etmişlerdir. Aynı zamanda çalışma, kolluk kuvvetlerinin keşif eylemlerinin kamu tarafından onaylanması için nasıl hareket edebileceklerini analiz etmektedir. Ayrıca bu çalışma, suçların önlenmesi, soruşturulması, tespiti ve verilerin serbest dolaşımına ilişkin yasal çerçeve oluşturmaktadır. Açık web içeriğinin işlenmesi çalışmalarından Richard ve arkadaşları Almanya, Avusturya ve İsviçre ülkelerine yönelik yayın yapan Almanca sağlıkla ilgili web içeriğini analiz etmişlerdir [118]. Almanca sağlık ağı graf yapısını, önemli içerik sağlayıcılarını, kamu / özel paydaş oranını belirlemiştir. Destek Vektör Makinesi kullanan sınıflandırıcısı geliştirmiştir [118].

CTI için önemli bir kaynak da sosyal medyadır. Dmytro Lande ve arkadaşları çalışmalarında [119], sosyal medyayı izlemek ve analiz etmek için kurumsal bir sistem oluşturmaya yönelik, sosyal medyadan siber güvenlik konularında büyük miktarda veri işleme, metin dizilerinden bilgi çıkarma konseptine dayanan yaklaşımlar önermişlerdir. Belirli konularda sosyal ağların içerik izleme sisteminin oluşturulması, sosyal ağlardan ilgili bilgilerin seçilmesi, kullanıcılar tarafından iyileştirilmesi için arama motorunun uygulanması, sorguların RSS beslemeleri olarak kaydedilmesi, müşteri uygulamalarında kişisel veri tabanlarının tutulması için bilgi teknolojileri önerilmektedir. Ivan Sousa Fernandes çalışmasında [120], sosyal ağlardan, yani Twitter'dan veri alma ile başa çıkmak için araçları bütünleştiren tam otomatik bir OSINT sistemi geliştirmeye odaklanmıştır. Bu araçlar, tweetler ilişkili verilerin toplanmasını ve toplanan verilerin metin analizine önem verilerek sınıflandırılmasını içerir. Kapsamı, bilgi

güvenliği, sosyal ağlar ve NLP alanlarının kesişimine girmektedir. Bir sosyal ağ ile ilgilenildiği için API'ler ve web crawler veri elde etmek için kullanılmıştır [120]. Andrea ve arkadaşı Siber Tehdit İstihbaratı (CTI) çalışmalarında kullanılmak üzere OSINT kaynaklarının çekilen içeriklerin otomatik olarak değerlendirilmesi için bir yaklaşım sunmuştur. İstihbarata dayalı uyarılar oluşturmak, kaynaklardan veri toplamak ve kullanılmasının karar vermek için yayımlama kaynağı değerlendirmesini çalışması yapmıştır [121]. OSINT kaynakları tehdit istihbaratı alanında çalışan siber güvenlik uzmanları ve akademik araştırmacılar arasında yapılan bir anketle seçilmiştir. OSINT kaynaklarının uygunluğunu ölçmek için puanlama, özellikler ve zamansallık dikkate alınarak özellik seti tanımlanmıştır [121].

Goularas ve arkadaşı [122] Twitter verilerindeki sentimental analiz için NLP ile birlikte RNN ve CNN olmak üzere iki kategorik sinir ağı kullanmışlardır. Uzun-Kısa Süreli Bellek (Long Short-Term Memory - LSTM) ağlarının CNN topluluklarını ve kombinasyonlarını ve bir RNN kategorisini değerlendirip karşılaştırmışlardır. Ayrıca, Word2Vec ve kelime gösterimi GloVe (Global Vectors for Word Representation -Kelime Temsili için Global Vektörler) modelleri için global vektörler gibi farklı kelime gömme sistemlerini karşılaştırmıştır. Abbass ve arkadaşları [123] veri (tweet) ön işleme, model oluşturucu ve tahminlerin sınıflandırılması için üç modülden oluşan bir çerçeve önermişlerdir. Verilen modeli farklı suç sınıflarına ayıran KNN ve SVM tahmin modelini oluşturmak için Multinomial Naïve Bayes (MNB) kullanılmıştır. N'nin en iyi değerini belirlemek ve Unigram, Bigram, Trigram ve 4-gram gibi farklı seviyelerde sistemin doğruluğunu ölçmek için bu makine öğrenme algoritmalarıyla daha fazla N-Gram dil modeli kullanılmıştır.

Bu tez kapsamında ise açık kaynaklar ve çeşitli API'lar üzerinde veri sızıntıları ve ihlallerinin bulunup elde edilmesi, yapısal hale getirilmesi, anlamlandırılması,sızıntı türü kişisel, kurumsal, potansiyel bilgi ve anonim olarak sınıflandırılması için bir model geliştirilmesidir. Başarılı analiz ve sınıflandırma işlemleri sonrasında tez kapsamında toplanan kullanıcı bilgilerinin üzerinde parola seçim davranışı üzerine analizi yapılmıştır. Yapılan parola seçim davranışı analizinin başarısının gösterilmesi için bir vaka çalışmasında kaba kuvvet saldırısının iyileştirilmesinde kullanılmıştır. Kaba kuvvet uzayı dışında kalan uzunluklarda parolaların kırılmasına imkan sağlayan özellikler kullanılarak parola seçim davranışını kontrol eden parola güvenlik politikası önerilmiştir.

3. YÖNTEM

3.1. Siber Tehdit İstihbaratı için Yeni Tarama Modeli

Siber Tehdit İstihbaratı için Yeni Tarama Modeli, sahada birçok araç ve operatör desteği ile gerçekleştirilen işlemlerin birleştirilerek uçtan uca minimum operatör desteği ile çalışmasını hedeflemektedir. Saldırganlar için sızan kullanıcı verileri önemli bir gelir kaynağı olabilmekte ve ekonomik değeri olmayan verilerin paylaşılması suçu yayma yöntemidir. Saldırganların, anonfiles.com ve pastebin.com gibi açık kaynak aracılığıyla sızan verileri paylaştığı gözlemlenmiştir. Herhangi bir kullanıcı bu açık kaynak hizmetler üzerinden bir veri paylaştığında, kullanıcıya rastgele oluşturulmuş bir URL üretilir. Ancak bu hizmetler, paylaşılan verileri aramak için bir izin veya seçenek sağlamaz ve yalnızca bu bağlantıyı bilen kişiler bu verilere erişebilir. Bu nedenle, ya saldırganların bu bağlantıları paylaşmasını beklemek ya da bu servisler aracılığıyla paylaşılan veri ihlallerini bulmak için bu platformlar üzerinden arama yapmak gerekmektedir.

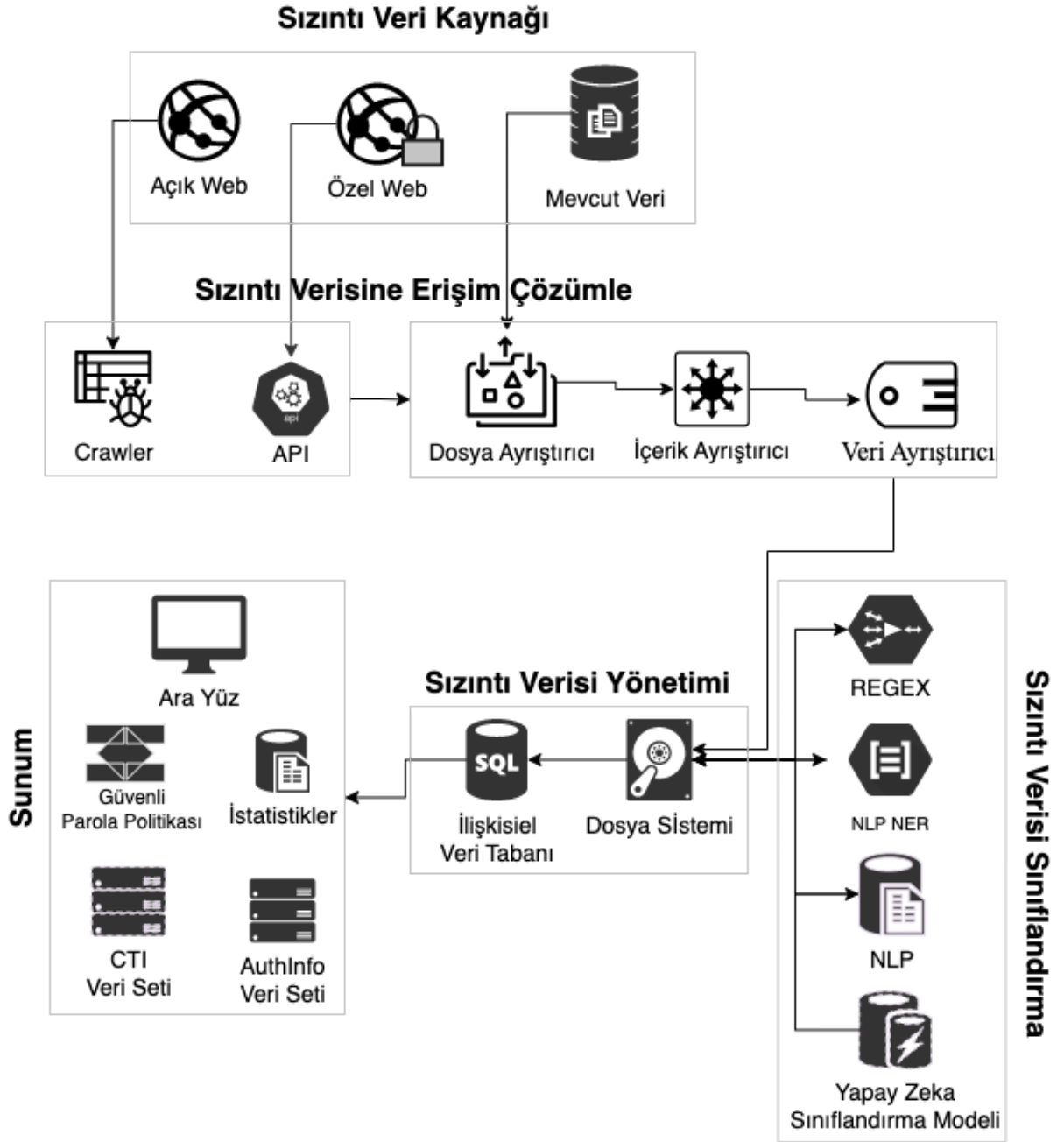
Veri ihlalleri/sızıntıları kişisel ve kurumsal verileri tehdit eden önemli bir siber tehdit vektörüdür. Saldırganlar genellikle ekonomik değer taşıyan verileri genellikle ilk olarak şantaj aracı olarak kullanmayı denerler. Yasadışı web sayfalarında çalınan verileri ilgilenen müşterilere satmaya çalışmaktadırlar. Bir kurumun verisinin çalındığı genellikle bu aşamada kamuoyunda duyulur. Ulaşılabilecek potansiyel müşteri sayısını artırmak için yapılan bu davranış, siber güvenlik uzmanlarına ve araştırmacılara da veri ihlalini araştırmaya yönelik fırsat sağlamaktadır. Modelimizde bu aşamada paylaşılan kanıt görsellerinden ve metinlerin içinde geçen bilgilerden ihlal tanımlamaya çalışılmaktadır. Bir başka aşama verinin ekonomik bir değeri kalmadığında suçun yaygınlaştırılması ve saldırganın izini kaybettirmesini kolaylaştırması amacıyla ücretsiz olarak yapıştırma sitelerinde (pastebin.com), anonim dosya paylaşım sitelerinde (anonfiles.com vb.) veya yasadışı forumlarda (raidforum, TOR siteleri vb.) paylaşılmaktadır. Bu aşamada veri elde edilerek analizi ve içeriğinin değerlendirilmesi yapılmaktadır.

Bu tez çalışmasında, sürekli artan siber istihbarat verilerinin, veri sızıntılarının etkin kullanımı için analiz ve değerlendirme süreçlerinin otomatikleştirilmesi hedeflenmektedir. Minimum

operatör desteği ile, özellikle Siber Tehdit İstihbaratı sürecinde verilerin ayrıştırılması için yapay zeka ve doğal dil işleme tekniklerinin kullanılması ile siber istihbarat verileri değerlendirilmiştir. Açık kaynaklardan ve API'lar üzerinden özel kaynaklardan toplanan siber güvenlik haberleri, veri sızıntıları, veri ihlalleri, forum veya sosyal medya gönderileri gibi farklı yapısal özellikteki metin içeriğini kişisel, kurumsal, potansiyel bilgi veya anonim olarak tanımlamaktadır. Verinin ait olduğu ilgili kişi ve kuruluşu analiz eden ve mevcut tehdit bilgilerinden üretilen özelliklerle karar sistemi gerçekleştirilmektedir.

Bazı durumlarda Google gibi arama motorları bu hizmetleri indeksler ve basit bir anahtar kelime araması ile paylaşılan dosyalara erişmek mümkündür. Ancak, saldırganlar aynı yöntemi diğer kullanıcılara karşı ortalama saldırıları ve kötü amaçlı dosya bulaştırmak için de kullanılabilir. Bu nedenle, araştırmacılar böyle bir görevi güvenli bir sanal alan ortamında gerçekleştirmesi gerekmektedir.

Siber Tehdit İstihbaratı için Yeni Tarama Modeli geliştirilmesini amaçlayan bu tez çalışması kapsamında taranabilir web üzerindeki veri ihlali kaynakları üzerinde sızıntı türü, kaynağı ilgili kişi, kurum ve kuruluş özelliklerine göre sızmış verileri taramak ve analiz etmek için bir model geliştirilmiştir. Tez sürecinde veriye erişim, çözümleme, yönetme ve sınıflandırma gibi çözülmesi gereken bilimsel zorluklar bulunmaktadır. Bu zorlukları Sızıntı Verisi Tespit, Sızıntı Verisine Erişim ve Çözümleme, Sızıntı Verisi Yönetimi, Sızıntı Verisi Sınıflandırma ve Sunum katmanlarıyla beş farklı modül üzerinde çözülmesi **Şekil 3.1**'de görülmektedir.

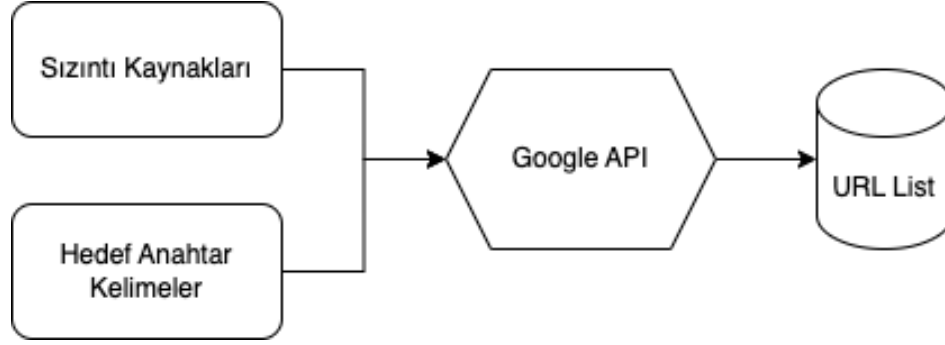


Şekil 3.1: Siber Tehdit İstihbaratı için Yeni Tarama Modeli Katmanları.

3.1.1. Sızıntı Verisi Kaynağı

Ücretsiz ve anonim veri paylaşım platformları aracılığıyla veri ihlallerini toplamak için bilinen standart bir yöntem bulunmamaktadır. Ancak saldırganlar tarafından paylaşılan veri ihlallerine arama motorları ya da saldırganların paylaştıkları forum ve sosyal medya kanalları aracılığıyla erişilebilmektedir. İhlal edilen verilerin sıklıkla ".txt" uzantılı metin dosyaları şeklinde paylaşıldığı gözlemlenmiştir. Arama motorları üzerinde, ".txt" dizisinin ilgili anahtar sözcüklerle birleştirme işlemi yapılarak başlanabilir ve genellikle birden fazla dosyanın

bulunmasıyla sonuçlanır. Tez kapsamında kullanılan arama terimleri **Tablo 3.1'in** içinde verilmiştir. "site:pastebin.com edu.tr" gibi çeşitli siteler ve anahtar kelimeler için arama motoru aramaları yoluyla yüzlerce dosya toplanmıştır. Arama motorları ile Google Dork kullanılarak sızıntı verisi arama yöntemi **Şekil 3.2**'de görülmektedir.



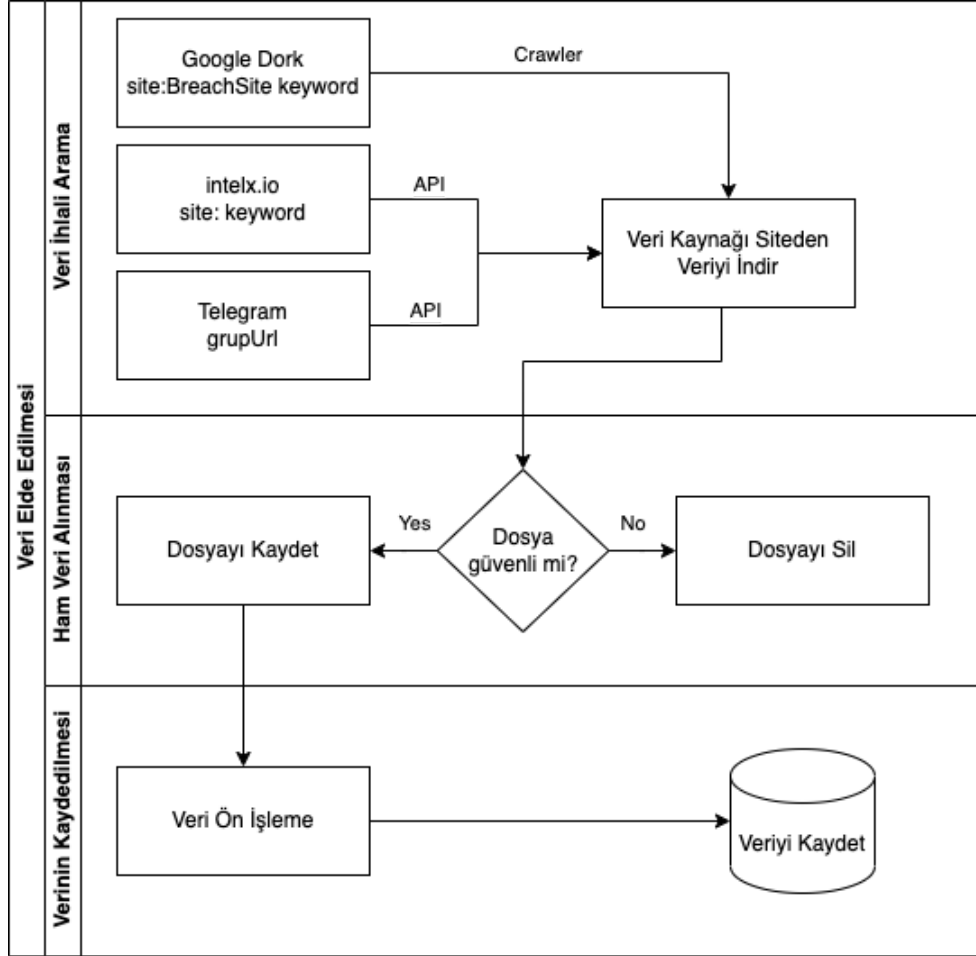
Şekil 3.2: Açık web tarama için mevcut arama motorları üzerinden veri akışı.

Tablo 3.1: Anahtar kelimeleri ve sızıntı kaynakları.

Sızıntı Kaynağı	pastebin, anonfiles, combolist
Arama Kelimesi	edu.tr, data leak, data breach, gmail, hotmail, email, password, VPN, account

Ayrıca veri ihlallerini toplayan ve bunları paylaşan ücretli ve ücretsiz platformlar da bulunmaktadır. Bu platformlardan intelx.io platformu API tez çalışmasına entegre edilerek günlük 200 dosya indirme izni limitiyle sızıntı bilgileri toplanmıştır. Veri ihlallerinin aranması adımı arama motorlarının ve intelx.io sonuçları API ile operatör gözetiminde elde edilir. intelx.io API üzerinden anahtar kelime ve depolama alanı seçerek sorgulama yapılabilmektedir. 24 saat içinde 50 sorgulama, 100 okuma ve 100 görüntüleme izinleri bulunmaktadır. Başka bir sızıntı kaynağı da sıklıkla kullanılan Telegram gruplarıdır. Genellikle suçun yaygınlaştırılması aşamasında gruplar üzerinden kendi topluluklarıyla paylaşmak amacıyla ve herkese açık gruplarda katılan herkesle paylaştıkları belgeleri metin ve dosya formatında paylaşmaktadırlar. Telegram API üzerinden elde edilen metin ve dosyalar da periyodik olarak kaydedilmektedir. Arama süreci **Şekil 3.3**'te gösterildiği gibi veri ihlallerini arama, ham verileri alınması ve verilerin kaydedilmesi olarak üç farklı aşamaya ayrılmıştır. API'ler ve tarayıcılarından elde edilen sızıntı dosyaları ön işlemden geçirilerek kaydedilir. Bulunan dosyalar herhangi bir şekil ve formatta olabilmektedir. Bu nedenle ön işleme aşamasının operatör gözetiminde olması önemlidir. Ayrıca, zararlı yazılımların dahil

olmak üzere yanlış pozitif sonuçların temizlenmesi gereklidir. Bu işlem adımlarının tamamı güvenli bir sanal ortamda yapılmıştır. Olası işleme aşamalarında internetten ve ağdaki diğer bilgisayarlardan izole edilmiştir. Ön işleme aşaması tamamlanan dosyaların sınıflandırılması ve başka analizlerde kullanılmak amacıyla özellik çıkarımı ve anonimleştirme aşamaları tamamlandıktan sonra dosya sistemi ve veri tabanına kaydedilir. Sistemin dosya sisteminde CSV (Comma Separated Values) dosyası, ham dosya ve işlenmiş veri tabanı kaydı oluşturulur.

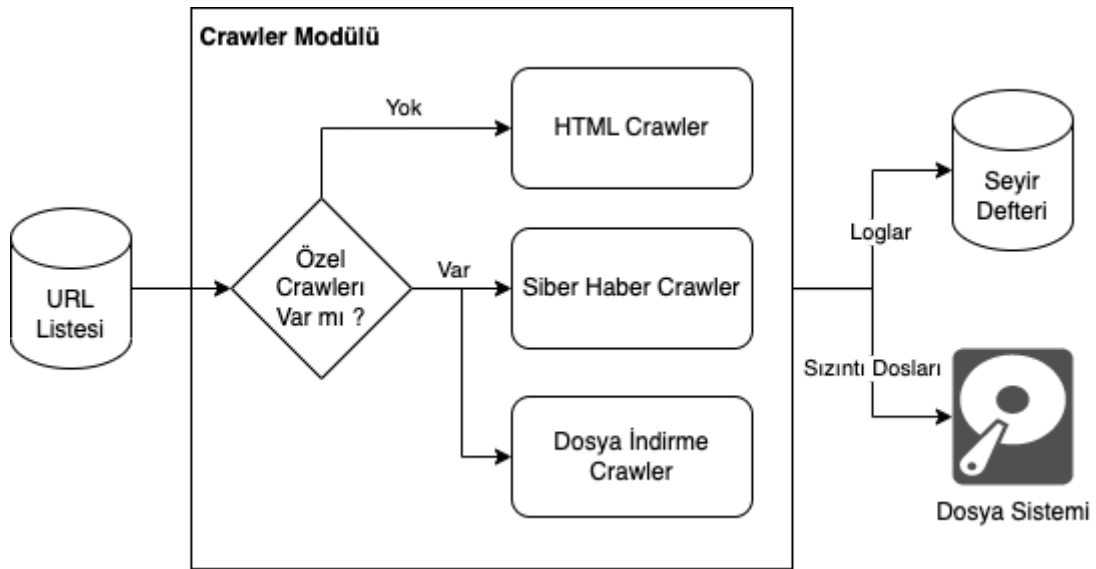


Şekil 3.3: Veri elde edilmesi akışı.

Tez çalışmalarında sızıntı verisi kaynaklarının çeşitliliği yanında bütün açık web (arama motorları, haber siteleri gibi) ve özel web'in (intelx.io, Telegram Grupları, TOR, Sosyal Medya gibi) taranmasıyla tam anlamıyla mümkün olduğu görülmüştür. Tez kapsamında yapılan araştırma ve analiz çalışmaları Google arama motoru, intelx.io ve Telegram gruplarıyla sınırlandırılmıştır. Diğer kaynaklar ve o kaynaklardan elde edilen verilerin sisteme analiz ve sınıflandırma için dahil edilmesi operatör ile yapılan çalışmalar kapsamında mevcut veri olarak sistem ile ilişkilendirilmiştir.

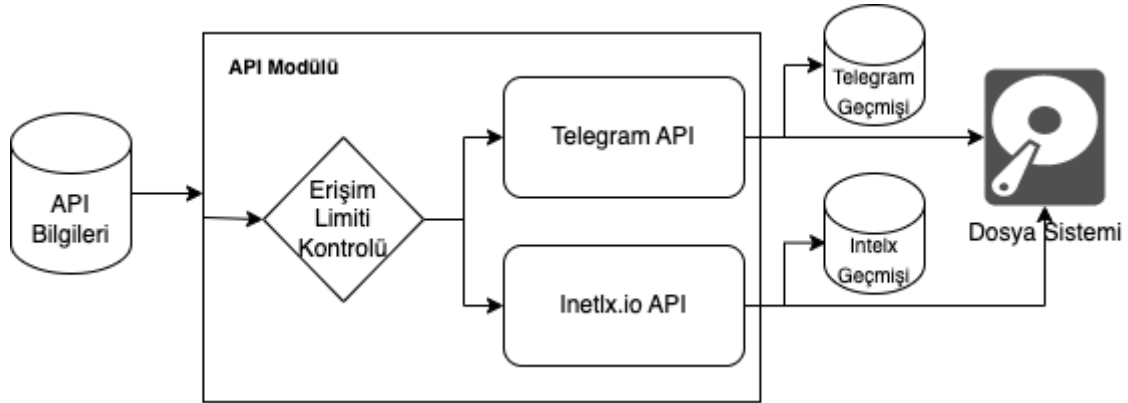
3.1.2. Sızıntı Verisine Erişim ve Çözümleme

Farklı veri kaynaklarında bulunan sızıntı belgelerini elde etmek için farklı yöntem ve prosedürler bulunmaktadır. Bazı kaynaklar doğrudan API sağlarken bazı kaynakların elle kazınması edilmesi gerekmektedir. Tez kapsamında saldırganların özellikle veri paylaşımı için sıklıkla kullandıkları veri paylaşım noktalarına odaklanılmıştır. Sızıntı verilerine erişim ve çözümleme için Veri kazıma (Crawler) module, API Modülü, Dosya Ayırıştırıcı Modülü ve İçerik Ayırıştırıcı modülü bulunmaktadır. Crawler modülü **Şekli 3.4** te görüldüğü gibi URL listesine bir önceki aşamada arama motorlarından eklenen URL'ler ve operatörler tarafından eklenmiş kaynak olası veri sızıntısı bulunabilecek kaynakların URL'leri üzerinde çalışmaktadır. Crawler modülü üzerinde URL listesine eklenen URL'ler için eğer özel bir veri kazıma geliştirilmiş ise ilgili veri kazıma modülü çalıştırılmaktadır. Eğer özelleştirilmiş bir crawler bulunmuyorsa sayfanın arayüz kodu ile birlikte (HTML kodu) kaydedilmektedir. Süreç içerisinde taranan sayfalar, erişim tarihleri ve indirilen dosya isimleri ile birlikte seyir defteri tablosuna yazılmaktadır. Özellikle periyodik olarak kontrol edilen kaynaklar için indirme geçmişi seyir defteri üzerinden takip edilmektedir. Periyodik ve tek sefere mahsus URL listelerinden elde edilen sonuçlar dosyaları (HTML, XML, JSON, txt, PDF, Doc, Docx) dosya sistemi üzerine kaydedilmektedir.



Şekil 3.4: Crawler Modülü'nün veri akışı.

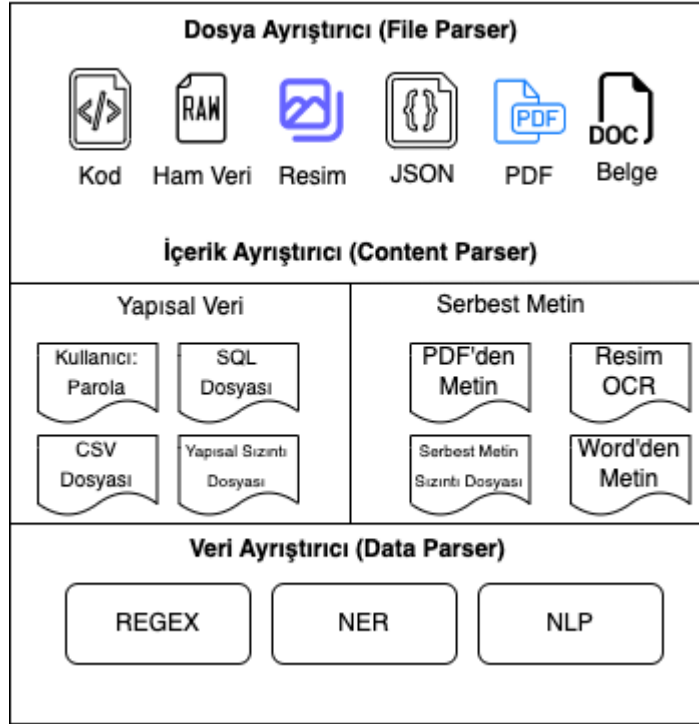
Veri sızıntı kaynaklarının ücretli / limitli kullanım hakkı olan API hizmetleri kullanılarak da sızıntı verisine erişim sağlanmıştır. **Şekil 3.5**'te görüldüğü gibi API modülü öncelikle kullanım limitleri kontrollerini sağlamaktadır. Aksi takdirde ilgili API sağlayıcısı size özel veriler API anahtarınızı banlayabilir veya bir müddet erişiminiz kısıtlanabilmektedir. Veriye erişim için API arayüzü bulunan hizmetlerden alınan ham veri olarak bulunan veriler doğrudan metin belgesi formatında dosya sisteminde depolanmaktadır. API üzerinden sızıntı dosyaları mevcut ise metin belgesi (TXT), PDF, DOC/DOCX Docx, PNG ve JPEG gibi farklı dosya formatlarında sızıntı belgeleri dosya sistemine kaydedilmektedir. Kaydedildiği yer, dosya ismi ve kaynak bilgileri indirilme tarihi gibi bilgiler ise ilişkisel veri tabanı üzerinde kaydedilmektedir. Açık kaynak arama motorunda bulunan seyir defteri gibi Telegram API için Telegram geçmişi, intelx.io için Intelx geçmişi tablolarına çalışma logları yazılmıştır. Sosyal medya gibi çeşitli API ve web kazıyıcılar ile taramaya sıkı kontroller getiren veri kaynakları çalışma kapsamı dışında bırakılmıştır.



Şekil 3.5: API Modülü veri akışı.

Sızıntı verisine erişim sağlandıktan sonra dosya sistemine eklenen verilerin üzerinde asenkron bir şekilde çalışan ve gelen verileri analiz için ayrıştıran Ayrıştırıcı (parser) modülü geliştirilmiştir. Toplanan sızıntı verilerden bilgi çıkarımını artırmak için Ayrıştırıcı modülü veri kazıma aşamasından önce Dosya Ayrıştırıcıda tip belirlenmesi ve buna göre farklı ayrıştırma tekniklerinin uygulanması ihtiyacı doğmuştur. Daha sonra bu dosyaların içerisindeki verilerin yapısal olup olmamasına göre veri kazıma işlemi yapılmıştır. Ayrıştırıcı Modülü dıştan içe doğru (dosya tipinden içindeki veriye kadar) dosyaları ayırarak doğru yöntemle mümkün olan tüm verilerin çıkarılması için **Şekil 3.6**'da görüldüğü üzere Dosya Ayrıştırıcı, İçerik Ayrıştırıcı, Veri Ayrıştırıcı alt modülleri bulunmaktadır. İlk olarak gelen farklı dosya türlerini ayırmak için Dosya Ayrıştırıcı modülü bulunmaktadır. Dosya Ayrıştırıcı

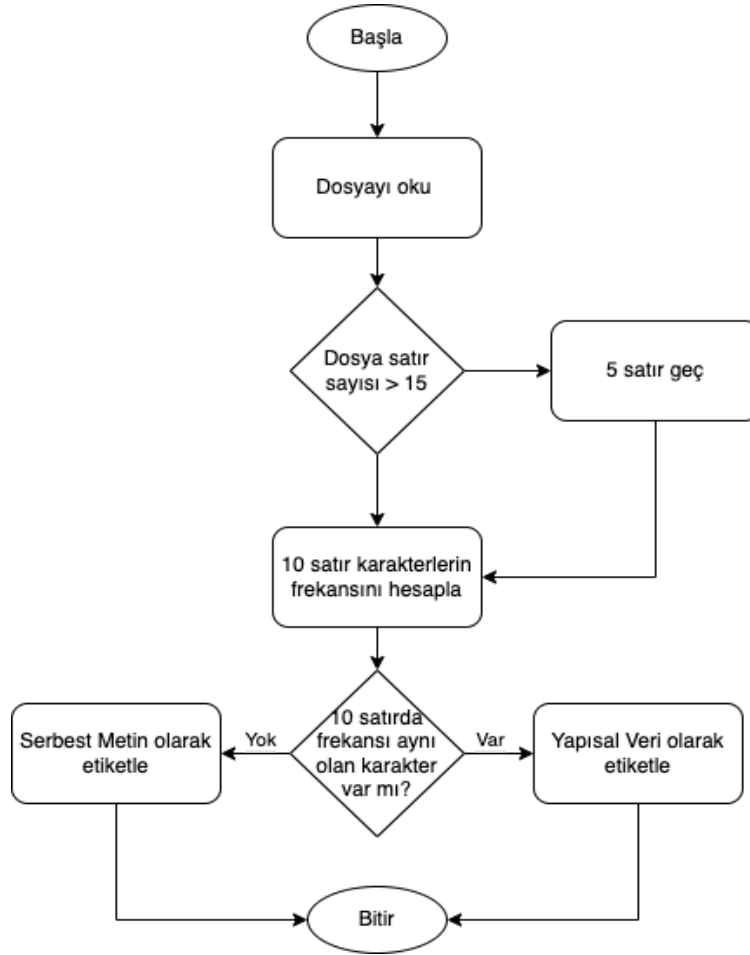
dosyaların uzantılarına göre dosyaları ayırmaktadır. Metin türünde dosyalar (TXT), CSV dosyaları, SQL dosyaları, HTML ve diğer kod dosyaları gibi dosyalar içeriği doğrudan okunabilirken bazı dosya tipleri kendi içerisinde farklı okuma yöntemleri ile içerisindeki veriler okunabilir. Örneğin bazı PDF dosyaları doğrudan PDFtoText kütüphaneleri kullanarak içerisindeki metinlerin okunması mümkün olurken bazı PDF dosyaları ve resim dosyalarının OSR kütüphanesi ile içeriğinin metne dönüştürülmesi gerekmektedir.



Şekil 3.6: Siber Tehdit İstihbaratı için Yeni Tarama Modeli Ayırıştırıcı Modülü ve Alt Modülleri.

Gelen dosyaların içerisindeki verilerin tutulma şekline göre yapısal ve serbest metin olarak ayıran içerik ayırıştırıcı modülü bulunmaktadır. Bir dosyanın içeriği iki nokta, virgöl, tab ve tırnak işareti gibi birçok ayraç kullanılarak ayrılmış olabilir. Bu durumda öncelikle dosyanın içerisinde bir ayraç kullanılıp kullanılmadığının tespit edilmesi gerekmektedir. Eğer bir ayraç kullanılmışsa dosyanın yapısal bir dosya olduğu belirlenebilir. Sızıntı dosyalarının İçerik Ayırıştırıcı alt modülündeki akış diyagramı **Şekil 3.7**'de gösterilmektedir. Sızıntı dosyaları incelendiğinde genellikle ilk beş satırında boşluk sızıntı yapan saldırgan ya da grup bilgisi, propaganda cümleleri gibi bazı mevcut yapısal forma uymayan kısımlar bulunduğu tespit edilmiştir. Bu nedenle yapısal dosyaları daha kapsayıcı bir şekilde analiz etmek için dosya satır sayısı 15'den büyük dosyalar için ilk beş satırı yapılan frekans analizine dahil edilmemiştir. Her bir satırdaki harfleri geçiş sayısı kontrol edilerek satırlarda ortak sayıda

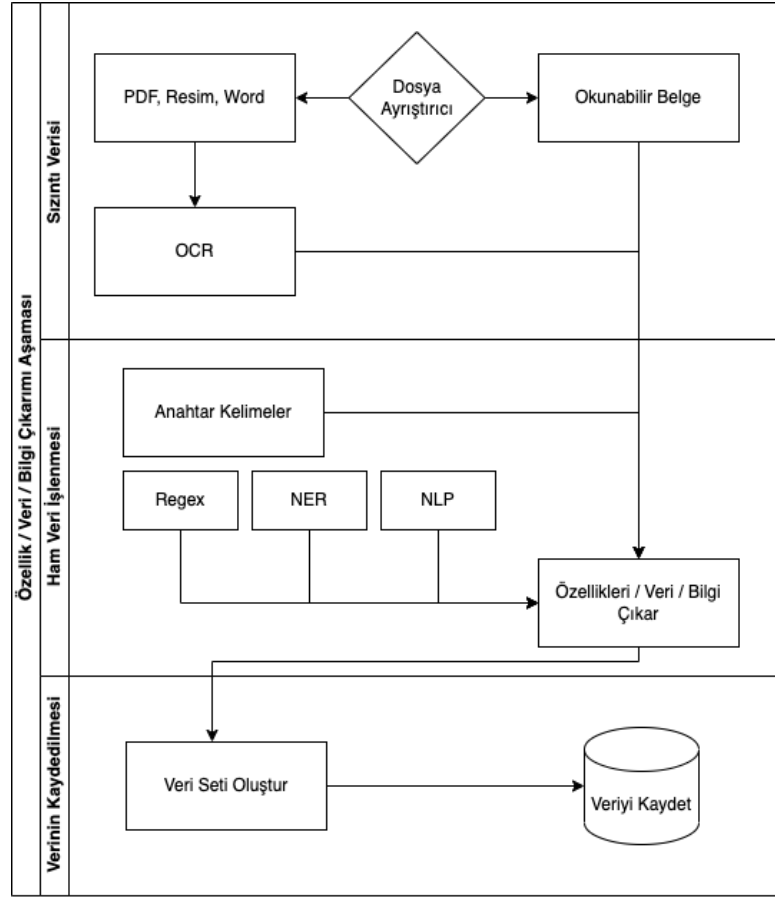
geçen (kolon ayraç) veri ayraçları belirlenmiştir. Bu ayraça sahip metinler yapısal veri dosyaları olarak etiketlenirken bu tür ayraça sahip olmayan dosyalar serbest metin dosyaları olarak etiketlenmiştir. Yapısal olarak işaretlenen dosyalar doğrudan okunması ve veri kazıma işlemine gerek duyulmadan bilgilerin çıkarılması işlemi gerçekleştirilebilmekte veri tanımlama/etiketleme aşamasına geçebilmektedir.



Şekil 3.7: İçerik Ayırıştırıcı Modülü İçerik Ayırıştırıcı alt modülü veri akışı.

Son olarak parser modülünde verinin çıkarılması ve kazıma (data capture) işlemlerini yapmak üzere Veri Ayırıştırıcı alt modülü geliştirilmiştir. Verinin REGEX kural listeleri, NER ve NLP kural ve modelleri yardımı ile verinin tanımlama (identifying) ve anlamlandırılması yapılmaktadır. Verilerin bu aşamada önceki modüllerden gelen veri dosyalarından Veri Ayırıştırıcı ile çıkarılan verileri REGEX, NER ve NLP ile özellikler veri seti oluşturularak CSV uzantılı dosya formatında veri tabanında saklanmaktadır.

Farklı türlerdeki belgelerinin işlenmesi bilgi ve özellik çıkarma çalışmaları aynı veri ayrıştırıcıda gerçekleştirilmektedir. PDF ve Resim belgeleri OCR yöntemi ile metne dönüştürülüp metin içerisinde olası kişisel ve kurumsal veri REGEX, NER ve NLP kullanarak araştırılmıştır. Çalışma dosya sistemi üzerinde çalışmakta ayrıştırılmış klasörler içerisinde dosyalar üzerinde gezerek analiz çalışması gerçekleştirilmektedir. Analizi yapılan dosyalar farklı dizinlere kaldırılarak işlenen dosyalar ayrıştırılmaktadır. Veri Ayrıştırıcı veri akışı **Şekil 3.8**'de görüldüğü gibi anahtar kelime ve REGEX kuralları, NER ve NLP sıralı olarak aynı dosya üzerine uygulanmaktadır. Dosya ayrıştırıcı farklı dosya formatlarına göre gelen veriyi ön işlemeden geçirerek (OCR) dosyalarında işleme yapılmadan da okunabilecek verileri çıkarmaktadır. OCR işlemi Tesseract kütüphanesi kullanılarak gerçekleştirilmektedir. Bu şekilde eğer doğrudan metin formunda olmayan dosyalarda herhangi bir yazı tespit edilirse bir metin içerisine eklenerek özellik / bilgi çıkarma işlemi gerçekleştirilmektedir. Okunabilir formda olan belgeler doğrudan işlenerek özellikler / bilgi çıkarımı yapılmaktadır. Son olarak Veri Seti olarak kullanmak amacıyla bu sonuçları da CSV uzantılı dosya formatında ve veri tabanı üzerinde saklanmaktadır.



Şekil 3.8: Veri Ayrıştırıcı Modülü veri akışı.

3.1.3. Veri Yönetimi

Tez kapsamında çeşitli kaynaklardan operatör tarafından toplanan ve mevcut veri olarak modele katılan sızıntılar, arama motorları üzerinden toplanan veriler ve API'lar üzerinden farklı kaynaklardan elde edilen sızıntı verilerinin ham ve işlenmiş hallerinin depolanması aşamasında dosya sistemi ve ilişkisel veri tabanı kullanılmıştır. Ham dosyaların depolanması için kullanılan klasörlerin dağılımı **Tablo 3.2**'de verilmiştir. Klasör ağacı kural tabanlı ayırtmada farklı türler altında birden aynı klasör tanımına görülebilir. Bu durum farklı veri türlerinin benzer içeriklerden olmasından kaynaklanmaktadır. Kurumsal e-posta adresi ile kişisel e-posta adresi ayırtma için örnek verilebilir.

Tablo 3.2: Ham dosyaların tutulduğu klasörlerin dağılımı.

Klasör Adı	Yapısal Veri Dosyalar	Serbest Metin Dosyalar	Anonim	Kişisel	Kurumsal	Potansiyel Bilgi
domainListesi_dosya	X	X	X			

kullanıcıBilgisi dosya	X	X		X		
ePosta parola	X			X	X	
kullanıcıAdı parola	X	X		X		
pipe dosya	X			X		
2Nokta_dosya	X	X		X	X	
ePosta dosya		X		X	X	
krediKartı bilgi		X		X		
bos dosya		X	X			
url dosya	X				X	X
http dosya	X				X	
tab dosya	X	X			X	
IP dosya		X			X	
html dosya	X		X			X
emailMetin dosya		X			X	
kod dosya		X	X		X	X
db_dump		X			X	

Veri tabanı klasör ağacı, tutulan ham verilerin indirildiği kaynakları, dosya isimleri, indirilme zamanı gibi bilgileri tutmaktadır. Ayrıca ham dosyalara uygulanan ön işleme süreçleri de veri tabanı üzerinden takip edilmektedir. Ön işleme aşamasında özellikle ayrıştırıcı modülü ile gerçekleştirilen veri kazıma ve tanımlama işlemleri sonuçları veri tabanı tablolarında ve dosya sisteminde CSV uzantılı dosya formatında yapısal olarak tutulmaktadır. Asenkron çalışan ayrıştırıcıların veri tabanı üzerinden çalışması sorgulama ve raporlama açısından faydalı olsa da süreç içerisinde veri taşıma ihtiyaçlarından dolayı doğrudan sorgulanmayan işlemler için CSV formatı dosya sisteminde kullanılmaktadır. Kullanılan veri tabanının şeması **Şekil 3.9**'da verilmiştir. Veri seti hazırlama çalışması için kullanılan geçici tablolar **Şekil 3.9**'daki gösterime dahil edilmemiştir.

dosya	seyir_defteri	telegram_seyir_defteri	intelx_seyir_defteri
id	id	id	id
dosya_adi	baslik	dosya_adi	dosya_adi
uzanti	aciklama	kullanici_adi	uzanti
kaynak	url	kaynak	storage_id
analiz_durumu	tarih	telegram_object	bucket
etiket		tarih	intelx_object
tarih			tarih

Şekil 3.9: Siber Tehdit İstihbaratı için Yeni Tarama Modeli veri tabanı tasarımı.

3.1.4. Veri Sınıflandırma

Genellikle farklı formatlarda saklanan kamuya açık kaynaklardan elde edilen verileri analiz etmek ve görselleştirmek için genelleştirilebilen bir yol bulunmamaktadır. Biçim seçimi sınırlı olsa da (virgül ile ayrılmış, iki nokta ile ayrılmış vb.) veri dosyaları üzerine kural yazarak çözümlmek oldukça uzun bir süreçtir. Serbest metin olarak paylaşılmış bir dosyanın analizi ve sınıflandırılması tamamen operatör yardımıyla yapılması mümkün değildir. Bu nedenle Siber Tehdit İstihbaratı için Yeni Tarama Modeli çıkarılan özellikleri kullanarak yapay zeka ve makine öğrenmesi yöntemleri ile dosyaları sınıflandırmaktadır. Veri sınıflandırma işlemi denetimli öğrenme yöntemleri ile gerçekleştirildiği için oluşturulan veri setlerinin geliştirilen arayüz ile elle etiketleme işlemi yapılmıştır. Bu süreç çıkarılan özellikler ve klasörlere ayırma sayesinde 8 bine yakın dosya tek bir operatör ile 30 saat gibi kısa bir zamanda etiketlenmiştir. Çıkarılan özellikler CSV uzantılı dosya da depolarak yapay zeka algoritmalarının eğitim sürecinde kullanılmıştır. Sızıntı verilerinin değerlendirmesi sürekli bir süreç olduğundan eğitim tamamlandıktan sonra elde edilen dosyaların geliştirilen yapay zeka modeli ile etiketlenmesi sağlanmıştır.

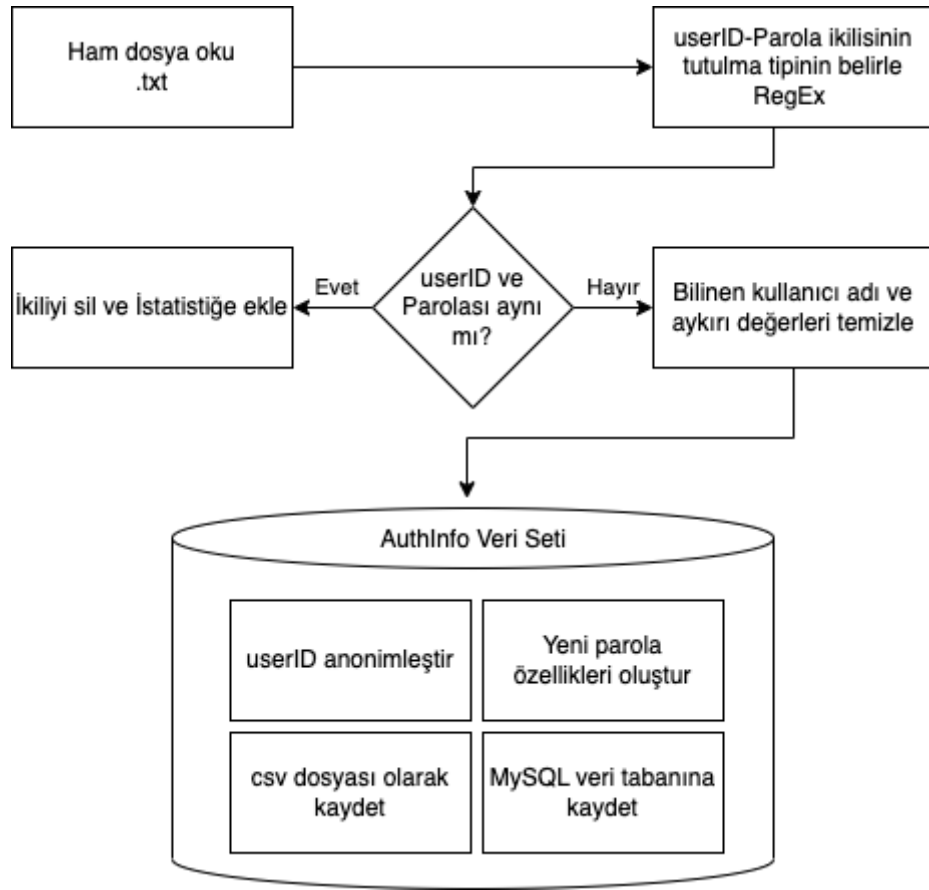
3.1.5. Sunum Katmanı

Siber Tehdit İstihbaratı için Yeni Tarama Modeli sunum katmanında üç farklı araçla sonuçları göstermektedir. Birinci araç bir web arayüzü üzerinden elle etiketleme işlemlerini takip edebildiği dosyaları anonim, kurumsal kişisel ve potansiyel bilgi olarak sınıflandırmaları görebildiği ve düzeltebileceği arayüzlerdir. Yine web arayüzü üzerinde sorgulama ve rapor özellikleri bulunmaktadır. İkinci olarak modelin doğal bir çıktısı olarak oluşan güncel sızıntılardan oluşturulan veri setleridir. Bu veri setleri anonimleştirilerek Github platformu üzerinden diğer kullanıcıların çalışmalarında kullanılması amacıyla paylaşılmaktadır. Üçüncü ve son çıktı olarak da sızıntı verilerinden yaptığımız analizler neticesinde son kullanıcılara faydalı yeni güvenlik yaklaşımlarıdır. Bulgularda sunulacak olan Parola Seçim Davranışı Duyarlı Parola Politikası da bu sunum başlığı altında değerlendirilmektedir.

3.2. AuthInfo Veri Seti

Kullanıcı bilgileri toplanan kaynaklar belli bir disiplin ve kurgu içinde depolanmamış dosyalar olduğundan dosyaların içerisinde sıralama ve format değişikliği olabilmektedir. İlk

olarak userID ve parola alanları dosyada tanımlanmıştır. Genellikle, kaynakların kullanıcı kimliği (userID) ve parola için birer sütun vardır. Ancak, bazı veri kaynaklarında kullanıcı sütunu ile parola sütunları yer değiştirdiği ve bazı kaynaklarda ikiden fazla sütun bulunur. Kullanıcı kimliği ve parola çiftlerini belirlemek için Python programlama dilinde geliştirilmiş bir script kullanılmaktadır. Ardından, bu çiftler analiz için ilişkisel bir veri tabanına aktarılmıştır. AuthInfo veri seti oluşturma adımın iş akışı **Şekil 3.10** 'da verilmiştir. Farklı formatlarda oluşturulmuş veri dosyalarından ortak bir formata dönüştürülen kullanıcı adları ve parolaları anonimleştirilerek ilişkisel bir veri tabanına kaydedilmiştir. Son derece kısa veya uzun ifadelerle sahip parolalar da aykırı değerler olarak kabul edilir ve hatalı kayıtları veya insan olmayan hesapları ortadan kaldırmak için veri kümesinden çıkarılır.



Şekil 3.10: AuthInfo veri seti oluşturulması aşamaları.

Kullanıcı davranışlarını analiz etmek için tek başına parola yeterli değildir. Parolaya ek olarak, kullanıcı kimliğinin analize dahil edilmesi durumunda kişisel ayrıntılarda davranış belirlemek mümkündür. Ancak, kullanıcıyı tanımlayabilecek herhangi bir bilgiyle çalışmak, ülkeye ve düzenlemelerine bağlı olarak yasal zorunluluklar içerir. Örneğin, Türkiye’de

KVKK ve Avrupa Birliğinde GDPR, kullanıcıyı tanımlayan tüm kişisel verilerin anonimleştirilmesini önerir. Bu nedenle, bu çalışmada, bir kişiyi doğrudan veya dolaylı olarak tanımlayabilen userID'ler, anonimlik oluşturmak için özetlenmiştir.

Özet algoritmaları, rastgele bir girdiden sabit boyutlu bir çıktıya bir eşleme oluşturur. Bu eşleme, tek yönlü bir dönüştürme işlevi oluşturur. Ancak, bu algoritmalar kaba kuvvet ve gökkuşağı tablosu saldırılarına karşı savunmasızdır ve yeterli zaman verildiğinde tersine çevrilebilir. Ayrıca, saldırganlar tarafından bilinen kullanıcı adları için karma oluşturmak kolaydır. Bu nedenle, bilinen kullanıcı kimliklerinden farklı çıktılar üretmek için, her kullanıcı için benzersiz bir özet değer oluşturmak üzere kullanıcı kimliğine bir rastgele dize eklenir. Özetleme için yaygın olarak kullanılan MD5 (Message-Digest 5 - Mesaj Özeti 5) algoritması, kullanıcı kimliğini hashlemek için kullanılmıştır ve kullanıcıların kişisel bilgilerini korumak için ek bir koruma olarak bir rastgele değeri eklenmiştir. Anonimleştirme işlemi **Denklem 3.1**'de verilmiştir.

$$AnonID = MD5(RastgeleDizi + userID) \quad (3.1)$$

userID özet algoritması uygulanmış olsa bile, kullanıcı adlarını kısmen veya tamamen parolalarına dahil ettiklerinde bu, kullanıcıların kişisel bilgileri korunmaktadır. Analiz aşamasında kullanıcı kimliklerini parola olarak kullanan 29 kullanıcı ve kullanıcı adının bir kısmını içeren parolaları kullanan 483.869 kullanıcı tespit edilmiştir. Karşı önlem olarak, bu tür sorunları olan bu kullanıcılar veri setinin dışında tutulmuştur. Ayrıca, bir parola olamayacak kadar kısa değerler bulunan (dört harften kısa) parolaları içeren toplam 5.554 kullanıcı kaydı kaldırılmıştır. Son olarak gerçek kullanıcı hesaplarının dahil edilmesini sağlamak için ön işleme işlemi sırasında en sık görülen kullanıcı kimlikleri manuel olarak taranır. Tutarlı sonuçlar elde etmek için "admin" gibi şüpheli veya genel kullanıcı kimlikleri de veri kümesinden kaldırılmıştır. Bu kayıtlar, başlangıçta toplanan kullanıcı bilgisi verilerin %4,83'üne denk gelmektedir.

AuthInfo veri seti, anonim olarak paylaşılan farklı veri kaynaklarının birleştirilmesiyle oluşturulmuştur. Bu kaynaklar, çeşitli hizmetlerden kullanıcı kimliği ve parola çiftlerinin sızıntılarını içerir. Ön işleme adımından sonra veri setini derlemek için 10.174.482 kullanıcı kimliği ve parola çifti kullanılmıştır. Ön işleme aşamasında, parola uzunluğu ve parola maskesi gibi özellikler üretilmiştir. Parola maskesi, parola saldırısında kullanılan hashcat uygulaması gösterimindeki her karakterin türünü gösterir. Hashcat uygulaması saldırganların

ve siber güvenlik araştırmacıların yaygın olarak kullandığı çeşitli parola saldırıları gerçekleştirilebilen bir parola saldırı aracıdır. Hashcat üzerinde kaba kuvvet saldırısı yapılacak parolanın her karakter türü bir harfle temsil edilmiştir. Karakterler türleri ve temsil eden karakterler **Tablo 3.3**'te gösterilmektedir.

Tablo 3.3: Parola içindeki karakter tiplerinin gösterimi.

Temsil Karakteri	Açıklaması	Aldığı Değerler
d	Rakamlar	0,1,2,3,4,5,6,7,8,9
l	Küçük Harfler	a,b,c,...,x,w,z
u	Büyük Harfler	A,B,C,...,Z,W,Z
s	Özel Karakterler	& !'\"%&\/().:;\"=?\@

Ek özellikler, kullanıcı adı ve paroladan türetilmiştir. Hesaplanan parola uzunluğu, paroladaki karakter türleri ve paroladaki her karakterin uygun karakter türüyle değiştirilerek oluşturulan maske dahil olmak üzere çıkarılan özellikleri **Tablo 3.4** açıklamaktadır. Kullanıcı davranışlarını anlamaya yardımcı olmak ve kaba kuvvet saldırıları için parola alanını yoğunlaştırmak için ilk ve son karakter özellikleri eklenmiştir. Bu bilgi, özellikle uzun parolalarda kaba kuvvet saldırıları için arama alanını azaltmaya yardımcı olabilir.

AuthInfo veri kümesi, araştırmacıların kullanması için Github'da CSV uzantılı dosya formatında paylaşılmıştır. Veri kümesi herkese açık olarak paylaşıldığında kullanıcılar için yeni bir ihlale yol açmasını önlemek için kullanıcı adları anonimleştirilmiştir. Ayrıca, her bir userID'nin benzersizliğini korumak için userID'lere özel rastgele bir salt değerle özetlenmiştir. Araştırmacıların bu çalışmayı doğrulamak ve gelecekteki araştırmalar için yeni yollar açmak için veri setini kullanmaları için paylaşılmıştır. Ayrıca, yeni sızıntı kaynakları ortaya çıktıkça veri setinin güncellenmesi için tez projesi kapsamında modül geliştirilmiştir.

Tablo 3.4: AuthInfo veri setinde bulunan özellikler.

Özellik	Açıklama	Örnek
AnonId	Kullanıcı tanımlayıcının bir gizli anahtar eklenerek özetlenmiş hali	ddd5129bba...8470d4a85
Parola	Açık metin parola	Deneme24@
Uzunluk	Parolanın karakter sayısı	9
Karakter Tipi	Parola içerisinde yer alan karakterlerin tipleri	ulds
Kaynak	Verinin elde edildiği kaynak	intelx.io
Maske	Parola karakterlerinin karakter tiplerinin parola karakterleri ile değiştirilmiş hali	u1lll1dds
İlk Karakterin Tipi	Parolanın başladığı ilk karakterinin tipi	u
Son Karakterin tipi	Parolanın başladığı son karakterinin tipi	s
Küçük Karakter Sayısı	Parolanın içerdiği küçük harf karakter sayısı	5

Büyük Karakter Sayısı	Parolanın içerdiği büyük harf karakter sayısı	1
Özel Karakter Sayısı	Parolanın içerdiği özel karakter sayısı	1
Rakam Sayısı	Parolanın içerdiği rakam karakter sayısı	2

3.3. Siber Tehdit İstihbaratı Veri Seti

Veri sızıntıları web robotu (crawler) ve API'lar üzerinden farklı kaynaklardan toplanan 7591 dosya sızıntı dosyasının ayrıştırılması geliştirilen ayrıştırıcı (parser) modülü ile yapılmıştır. Geliştirilen arayüz üzerinden etiketlenmesi tamamlanmış ve sonrasında yapay zeka ile denetimli bir öğrenme yapılabilmesi amacıyla özellik çıkarımı işlemi yapılmıştır. Ön işleme aşamasında özellik üretimi için REGEX kuralları ve NER kullanılmıştır.

3.3.1. REGEX

Sızıntı verilerinin anlamlandırılması süreçlerinde kişisel verilerin format ve yazılım kısıtlarından faydalanılmaktadır. Yazım kısıtları düzenli ifade (Regular Expression - REGEX)'ler kullanılarak oluşturulan kural setleri içerisinde verilerin T.C. kimlik numarasının 11 karakter rakamdan oluşması, vergi numarasının 10 hane rakamdan oluşması gibi kişisel ve kurumsal olabilecek verilerin kuralları çıkarılmıştır. Çıkarılan kurallar **Tablo 3.5**'te özellikler ve uygulanan kurallar verilmiştir. Verileri sızıntıları belirli bir düzen ve disiplin dahilinden paylaşılmadığı için serbest metin veya belli bir yapısal düzeni bulunan metinlerde REGEX kuralları uygularken kuralların dışında kalan, standart dışı ve/veya farklı standartlarda paylaşılan verilerin tespit edilmesi için de aynı özelliğin birden fazla kuralı bulunmaktadır. Bu kurallar sıra ile tüm metne uygulandığı için aynı bilginin birden fazla özellikle temsil edilebilmektedir. Ayrıca sınıflandırılan dosyaların içinde olmasına rağmen tespit edilemeyen veya o kurala uymayıp veri temizleme operasyonları sonra kurala uyduğu için gelen hatalı sonuçlar olabilmektedir. Bu sebeple tez çalışmasında doğrudan REGEX sonuçları kullanmamış olup diğer veri çıkarım yöntemleriyle yapay zeka modellerine Özelliğin bulunma durumu (var, yok) özelliğin kaç kere bulunduğu (sayı) olmak üzere 50 özellik ile modele girdi olarak verilmiştir.

Tablo 3.5: Serbest metin veri ayrıştırmak için kullanılan REGEX kuralları ve özellikler.

Özellik	REGEX Kuralı
T.C. Kimlik Numarası	[1-9]{1}[0-9]{9}[02468]{1}
Vergi Numarası	([0-9]{10})
Telefon	(05 5 905)([0-9]{2})\s?([0-9]{3})\s?([0-9]{2})\s?([0-9]{2})
Telefon 2	(\d{3}[-\.s]?d{3}[-\.s]?d{4} (\d{3})s*d{3}[-\.s]?d{4} \d{3}[-\.s]?d{4})
Sabit Telefon	(0 90?)(\s?)(322 416 272 472 382 358 312 242 478 466 256 266 378 488 458 228 4

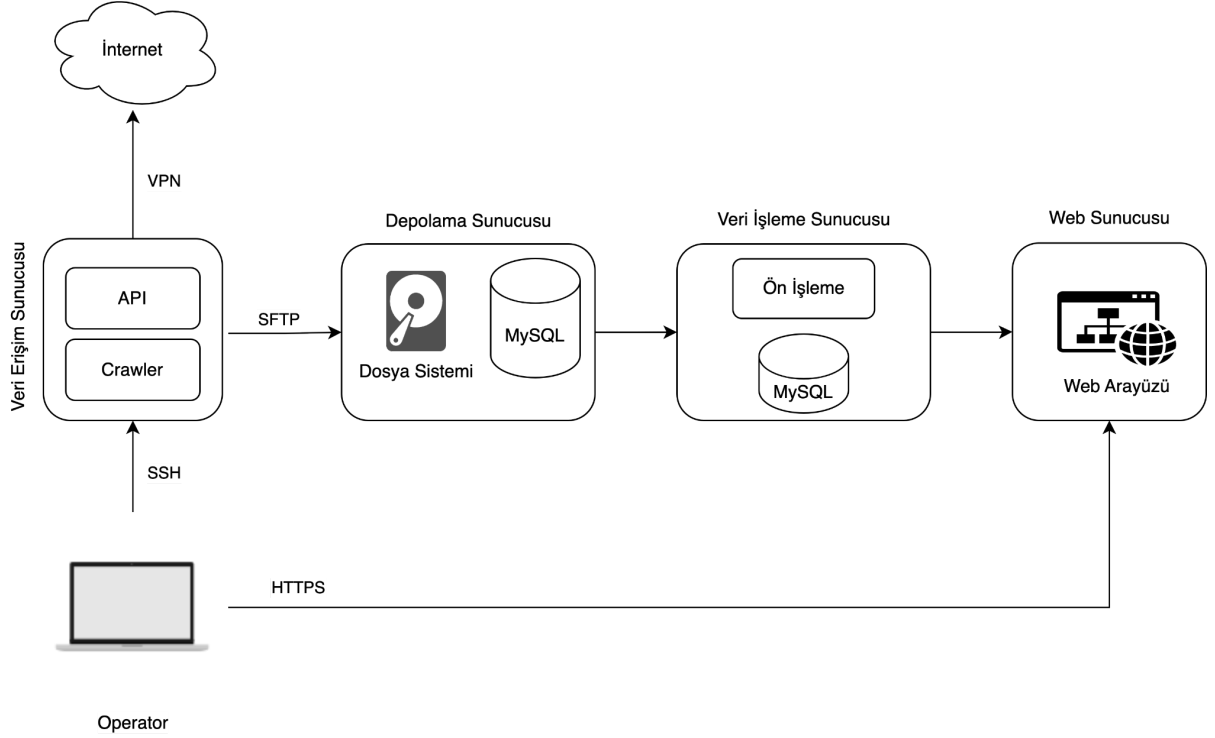
Python programlama dili ve NLTK (Natural Language Toolkit - Doğal Dil Araç Seti) kütüphanesi kullanılarak **Tablo 3.6**'da gösterilen özellikler metin sızıntı verisini işlemek ve analiz etmek için NER kullanılmıştır.

Tablo 3.6: NER yöntemi ile çıkarılan özellikler.

Varlık	Örnek
Organizasyon	İstanbul Üniversitesi - Cerrahpaşa, YÖK
Kişi	Ebu Yusuf GÜVEN
GPE	Türkiye, Erzincan, İstanbul
Konum	Avcılar Yerleşkesi, Bilişim Vadisi
Tarih	6 Mart 2021
Zaman	00:04, Gece Yarısı

3.4. Deneysel Araçlar

Siber Tehdit İstihbaratı için Yeni Tarama Modeli sızıntı verilerinin toplanması, ön işlemesi, analiz edilmesi ve arayüz üzerinden sunulması birden fazla bilgisayarın/sunucunun ve teknolojinin birlikte çalışması gerekmektedir. **Şekil 3.11**'de operatör ve modüllerin çalıştıkları sunucu/bilgisayarların ilişkisi gösterilmiştir. Çalışma kapsamında veri sızıntıları üzerine çalışma yapıldığı için veri toplama kaynaklardan ve indirilen dosyalardan kaynaklı siber güvenlik ve erişim riskleri bulunmaktadır. Erişim sağlanmak istenen veri kaynağı Türkiye'den yasal erişim izni bulunmaması, TOR ağı üzerinde olması, Türkiye IP'lerine kapalı olması veya veri kaynağına erişimde IP adresinin ve MAC adresinin gizlenmesi amacıyla ücretli bir Sanal Özel Ağ (Virtual Private Network - VPN) hizmeti kullanılması tercih edilmiştir. Veri erişim sunucusu üzerinden çekilen veriler SFTP (SSH File Transfer Protocol - SSH Dosya Aktarım Protokolü) uygulaması ile ham dosyalar Depolama Sunucusuna yüklenmektedir. Depolama sunucusu üzerinde dosyalar dosya sisteminde sızıntı kaynağı, erişim bilgileri ise MySQL veri tabanı üzerinde tutulmaktadır. Ayırıştırma modüllerinin çalıştığı ön işleme aşaması Veri İşleme Sunucu üzerinde tamamlanmaktadır. Raporların ve analizlerin tutulduğu MySQL veri tabanı Veri İşleme Sunucusu üzerinde de bulunmaktadır. Web arayüzü üzerinden arama, etiketleme ve rapor gösterme işlemleri için Web Sunucusu bulunmaktadır. Son olarak süreci yönetmek, etiketleme operasyonları için operatör bulunmaktadır. Veri Erişim Sunucusuna SSH üzerinden erişim sağlanırken, diğer sunuculara doğrudan SFTP ya da fiziksel erişimi bulunmaktadır. Ayrıca Operatör Web Arayüzüne HTTPS protokolü üzerinden erişilmektedir.



Şekil 3.11: Siber Tehdit İstihbaratı için Yeni Tarama Modeli Veri İşleme Mimarisi.

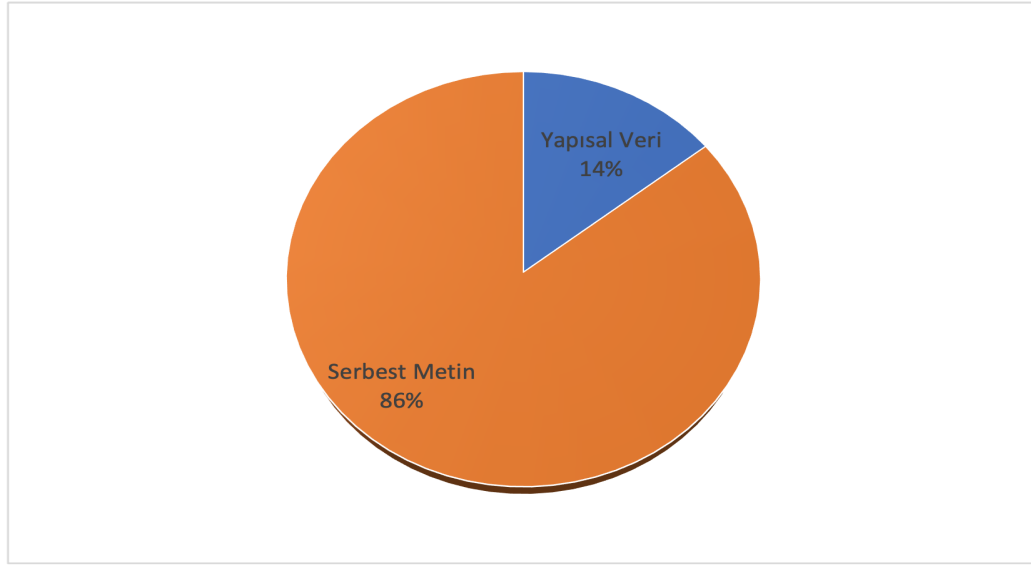
Siber Tehdit İstihbaratı için Yeni Tarama Modeli bilgi işlem ortamı, 250 GB RAM, 2,40 GHz Intel (R) Xeon (R) Gold 6148 CPU ve 14528/15205 MB ve 40MCU ile 4 Tesla T4 GPU donanımına sahip bir Yüksek Performanslı Bilgisayar (High Performance Computing - HPC) üzerine kurulmuş sanal cihazlardan oluşan ortamdır. Ücretsiz olan virtual box uygulaması ile sanallaştırılan kaynaklar Veri Erişim Sunucusu (Linux sunucu işletim sistemi-Ubuntu), Depolama Sunucusu (Linux sunucu işletim sistemi-Ubuntu), Web Sunucusu (Linux sunucu işletim sistemi-Ubuntu) ve Veri İşleme Sunucusu (Windows sunucu işletim sistemi) olmak üzere 4 sanal sunucu ve 1 operatörden oluşmaktadır. Veri erişim sunucusu güvenlik risklerinden dolayı İnternete VPN ile bağlanmaktadır. Sunucular üzerine ya doğrudan işletim sistemi arayüzü veya SSH (Secure Shell Host) ile güvenli bağlantı sağlanmaktadır. Veri Erişim Sunucusu üzerinde indirilen sızıntı dosyalarına SFTP (Secure File Transfer Protocol) ile erişim sağlanmaktadır. Diğer sunucular aynı ağ üzerinde olduğu için birbirleri ile doğrudan ağ üzerinden farklı protokoller ile bağlanmaktadır.

4. BULGULAR

Siber Tehdit İstihbaratı için Yeni Tarama Modeli içerisinde veri sızıntılarını tespit, erişim, analiz ve raporlama aşamalarının ara ve nihai çıktıları başta olmak üzere tez çalışması kapsamında yapılan analiz ve vaka çalışmaları da tez bulgularında paylaşılmıştır. Öncelikli olarak Sızıntı Verilerinin Analizi ve Veri Sızıntılarının Sınıflandırılması çalışmalarının modül çıktı sonuçları sunulmuştur. Daha sonra devamında Kullanıcı Adı-Parola Dosyalarının Analizi, Parola Seçim Davranışı Tespiti, Parola Seçim Davranışı Kullanarak Kaba Kuvvet Saldırısı başlıklı vaka çalışmasının ve Önerilen Parola Davranışı Duyarlı Parola Politikası bulgularına yer verilmiştir.

4.1. Sızıntı Verilerinin Analizi

Siber Tehdit İstihbaratı için Yeni Tarama Modeli veri erişim aşamasında, Google arama motoru üzerinden tespit edilen arama robotu (web crawler) ve operatör tarafından toplanan, Telegram ve intelx.io API'ları üzerinden erişilen verilerin ilk olarak içerik ve yapıları incelenmiştir. İlk olarak dosya türlerine göre sınıflandırılan veriler görsel, PDF ve Word belgelerinde yer alan metinler metin formuna çevrilmiştir. İkinci olarak içeriklerine göre yapısal ve serbest sınıflandırılmıştır. Sızıntı belgelerinin içerisinde başında ve sonunda saldırgan grubun isim ve propaganda verileri bulunduğu için belge ortasından alınan örneklerle yapılan karakter frekans analizi ile belgeler ortak bir karakter(ler) ile ayrılma durumu kontrol edilmiştir. Belgelerin tamamı analize katılmış olmamasına rağmen yapılan analiz sonucunda **Şekil 4.1**'de görüldüğü gibi yüzde 90'a yakın veri sızıntısı beklendiği üzere serbest metin olarak sınıflandırılmıştır. Eğer dosyaların tamamı analize dahil edilmesi durumunda yapısal veri dosyalarının oranı %10'un altına düşmektedir. Ancak yapılan operatör incelemesi sonucunda kullanıcı adı-parola ikilisi gibi genellikle yapısal paylaşılan belgelerin bile düzenli bir formatta tutulmamasından dolayı bu analizde 314 kullanıcı bilgi verisinin serbest metin olarak işaretlendiği görülmektedir.



Şekil 4.1: Siber Tehdit İstihbaratı için Yeni Tarama Modeli İçerik Ayırıştırıcı yapısal ve serbest metin dosyaların dağılımı.

İçerik ayırıştırma aşaması tamamlandıktan sonra Veri Ayırıştırma aşamasında veri dosyalarının içinde bulundurduğu bilgi türüne göre REGEX ve NER ile veri kazıma işlemi uygulanmıştır. Bulunan veri türleri kullanılarak dosya ayırıştırma işlemi yapılmış veriler **Tablo 4.1**'de görülen klasör ağacına bölünmüştür. Dosyaların alt klasörlere bölünmesi kişisel veri, kurumsal veri ve potansiyel veri sırasına göre çalıştırılmıştır. En son kalan verilerin anonim veri olarak değerlendirilmiştir.

Tablo 4.1: Dosya içerdiği bilgi türüne göre ayrıştırılabilen klasör ağacı

Yapısal Veri Dosyalar	Serbest Metin Dosyalar
- YapısalVeri_Dosyalar 1-Anonim - domainListesi_dosya 2-Kisisel - kullanıcıBilgisi_dosya - ePosta_parola - kullanıcıAdı_Parola - pipe_dosya 3-Kurumsal 2nokta_dosya - bos_dosya - url_dosya - http_dosya - tab_dosya - bos_dosya 4-PotansiyelBilgi - html_dosya	- SerbestMetin_Dosyalar 1-Anonim - domainListesi_dosya 2-Kisisel - krediKartı_bilgi - ePosta_dosya 2Nokta_dosya - kullanıcıBilgisi_dosya - kullanıcıAdı_parola 3-Kurumsal - bos_dosya - IP_dosya - emailMetin_dosya - tab_dosya - kod_dosya - db_dump 4-PotansiyelBilgi - html_dosya

Siber Tehdit İstihbaratı için Yeni Tarama Modeli Ayırıştırma modülü sonucunda “Anonim”, “Kişisel Veri”, “Kurumsal Veri” ve “Potansiyel Bilgi” olarak REGEX ve NER yöntemleri kullanılarak kural tabanlı etiketlenmiş verilerden “4 Etiket Model Veri Seti” oluşturulmuştur. Farklı yapay zeka modellerinde ikili sınıflandırmada kullanılmak üzere “Kişisel Veri”, “Kurumsal Veri” ve “Potansiyel Bilgi” başlıkları “Özel Nitelik” etiketinde toplanmıştır. Bu sayede ikili sınıflandırmaya uygun “Anonim” ve “Özel Nitelik” etiketleri bulunan “2 Etiket Model Veri Seti” oluşturulmuştur. “4 Etiket Model Veri Seti” ve “2 Etiket Model Veri Seti” Siber Tehdit İstihbaratı için Yeni Tarama Modeli çıktısı olarak oluşturulmuştur. Sızıntı veri dosyaları ayırıştırıcıdan geçirilirken **Tablo 4.2**'de veri özellikleri kuralları verilen ve tercih edilen etiketler gösterilmiştir. Kişisel verilerin belirlenmesinde KVKK baz alınarak ayırıştırılmıştır. Ancak kurumsal verileri tanımlayan doğrudan bir yönerge olmadığı için kurumla doğrudan ilgili olabilecek bilgiler “Kurumsal” bilgi olarak etiketlenirken firma veya kurumların ilgilenebilecekleri ve takip etmek isteyecekleri bilgiler olduğu düşünülen özelliklerde “Potansiyel Bilgi” olarak etiketlenmiştir. İçeriğinde herhangi veri kazıma operasyonların herhangi bir bilgiye rastlanmayan dosyaların yapısal veri veya serbest metin olmasına bakılmaksızın anonim veri olarak etiketlenmiştir.

Tablo 4.2: Siber Tehdit İstihbaratı için Yeni Tarama Modeli dosya etiketleme kuralları.

Veri Özelliği	Kişisel Veri	Kurumsal Veri	Potansiyel Bilgi	Anonim
Domain Listesi		X		X
Subdomain, alt uzantı veya zafiyetli URL		X		
e-Posta Listesi	X			
Adı Soyadı			X	
T.C. Kimlik Numarası	X			
Vergi Numarası		X		
Firma Ünvanı			X	
Kurumsal e-Posta	X			
Veri tabanı Sızıntısı (SQL veya CSV)	X	X		
Kredi Kartı Bilgisi	X			
Kullanıcı Adı, IP, Domain Listesi		X	X	
Kişi isimleri geçiyor ancak kurum bilgisi yok			X	
Usenet Bilgisi			X	
Kredi Kartı Bilgisi	X			

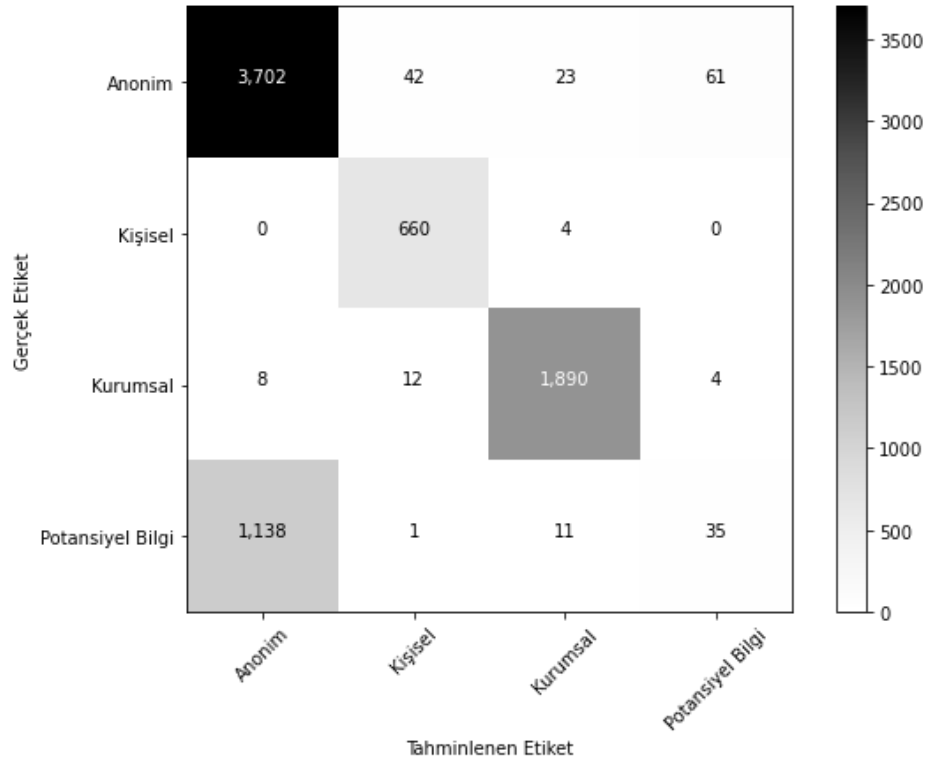
Erişime Açık Yasal Web Sayfaları ve içeriği				X
HTML dosyaları e-posta veya herhangi bilgi yoksa		X		X
EXTM3U Dosyaları			X	
Tab, 2 nokta, noktalı virgül ile ayrılan dosyalar	X	X	X	

Siber Tehdit İstihbaratı için Yeni Tarama Modeli ile geliştirilen bir de arayüz bulunmaktadır. Operatör desteğinin de sızıntı verilerinin sınıflandırılması için kullanılabilmesi amacıyla anahtar kelime arama, etiketleme ve raporlama çalışmaları için arayüz kullanılmaktadır. Yapılan etiketleme çalışması veri ayrıştırma sonrasında arayüz üzerinden gerçekleştirildiği için mevcut kurallar ve etiketler kullanılmıştır. Operatör tarafından yapılan etiketleme çalışması **Tablo 4.3**'te görüldüğü üzere “4 Etiket Operatör Veri Seti” ve “2 Etiket Operatör Veri Seti” olarak isimlendirilmiştir.

Tablo 4.3: Dosyalardan model sonucu ve operatör yardımı ile etiketlenen veri setleri.

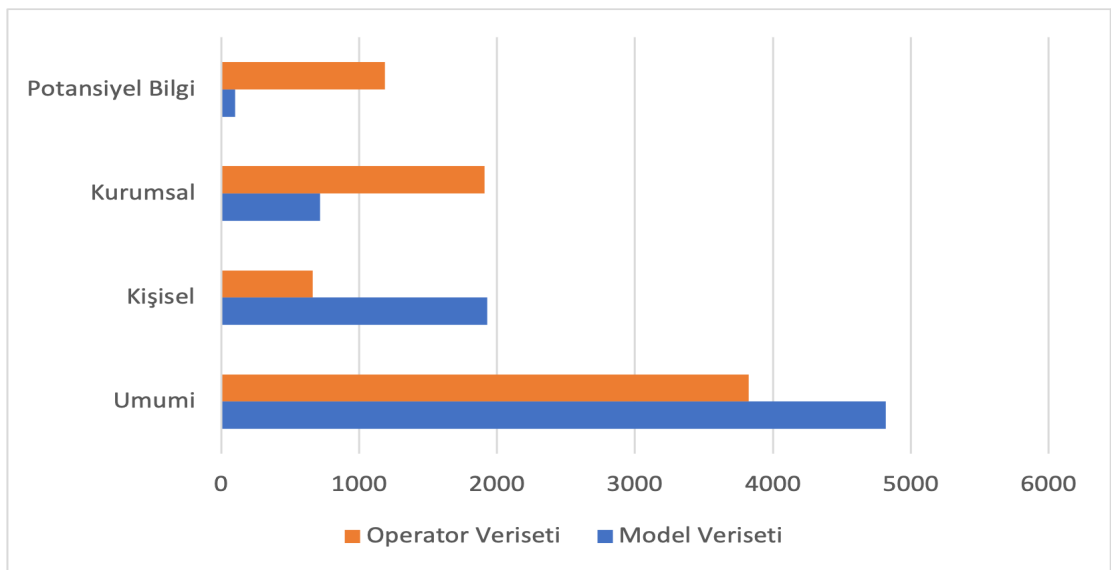
4 Etiket Model Veri Seti		2 Etiket Model Veri Seti		4 Etiket Operatör Veri Seti		2 Etiket Operatör Veri Seti	
İçerik	Etiket	İçerik	Etiket	İçerik	Etiket	İçerik	Etiket
4820	Anonim	4820	Anonim	3826	Anonim	3826	Anonim
1928	Kişisel	2743	Özel Nitelik	665	Kişisel	3762	Özel Nitelik
715	Kurumsal			1909	Kurumsal		
100	Potansiyel Bilgi			1188	Potansiyel Bilgi		

Model sonucu olarak “Anonim” olarak etiketlenen veri sızıntısı dosyalarının %17 si operatör tarafından diğer etiketler tercih edildiği görülmüştür. Ayrıca içerisinde kişisel veri geçmesine rağmen doğrudan veri tabanı sızıntısı olarak görülen verilerin yasal sorumluluğu kurumu ilgilendirdiği için kişisel değil kurumsal olarak işaretlenmiştir. Bu nedenle “Kişisel Veri” etiketinin yaklaşık üçte bir oranında azalırken “Potansiyel Bilgi” etiketinin 11 kat arttığı görülmüştür. Model tarafından etiketlenen verinin operatör tarafından düzeltilmesi sonucu karmaşıklık matrisi **Şekil 4.2**'de görülmektedir. Modelin kural tabanlı olarak doğruluk %83 oranıyla verileri etiketlediği görülmektedir.



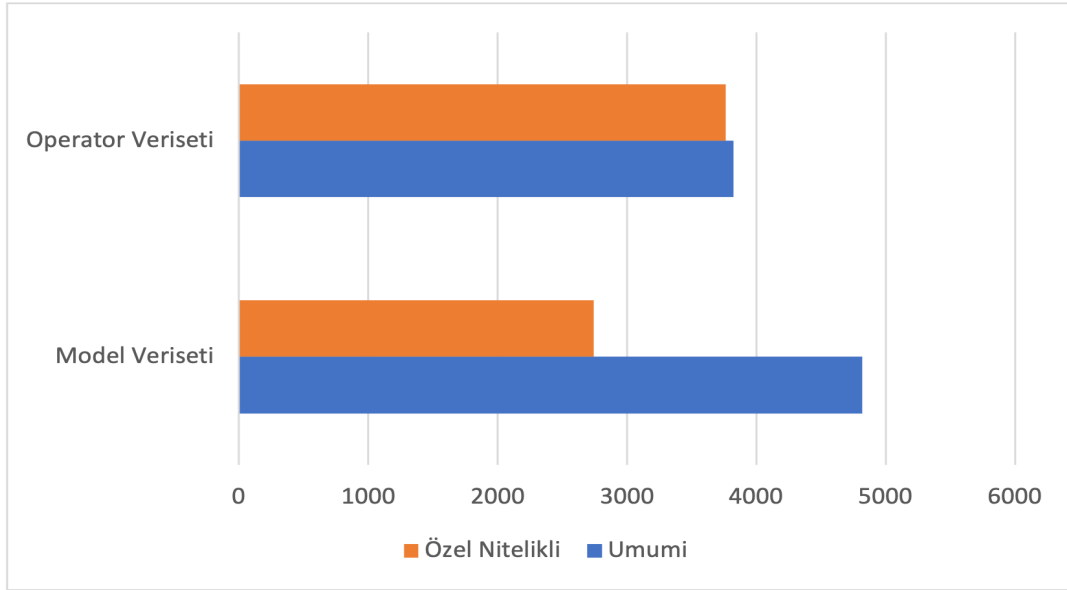
Şekil 4.2: Model tarafından kural tabanlı etiketlenen verinin operatör tarafından düzenlenmesinin karmaşıklık matrisi.

Operatör tarafından yapılan etiketlemede **Şekil 4.3**'te görüldüğü gibi “Anonim”ve “Kişisel Veri” etiketleri azalırken “Kurumsal Veri” ve “Potansiyel Bilgi” etiketlerinin arttığı görülmüştür.



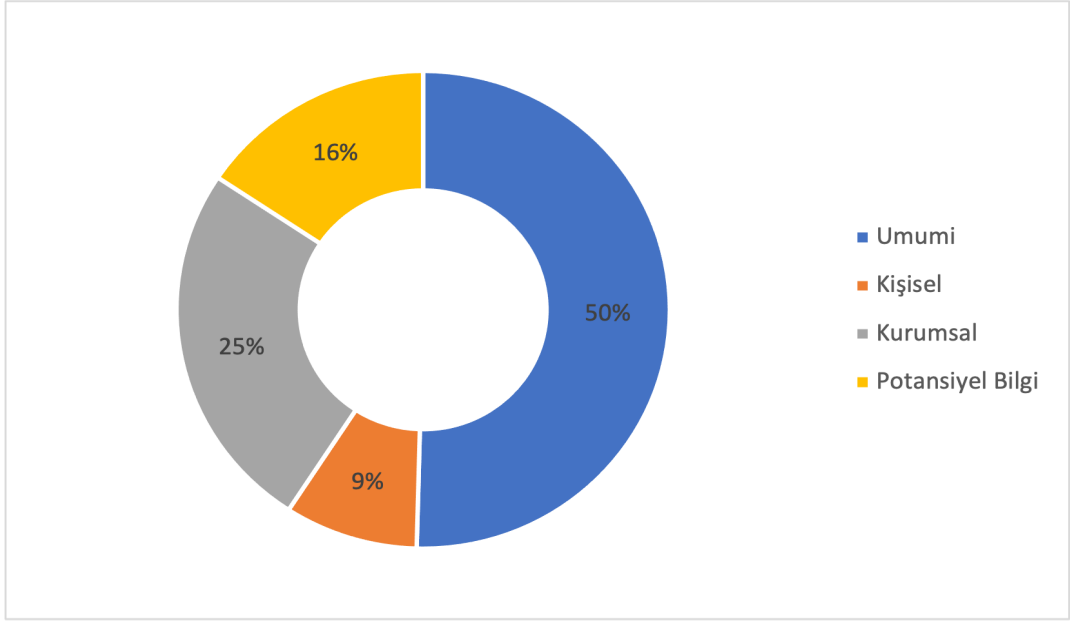
Şekil 4.3: 4 Etiket Model Veri Set seti ve 4 Etiket Operatör Veri Seti karşılaştırması.

İkili etiket dağılımı değerlendirildiğinde operatör etiketlemesi sonrası veri seti içerisinde daha dengeli bir dağılım ortaya çıkmıştır. Buna rağmen **Şekil 4.4**'te görüldüğü üzere veri sızıntısı olarak bulunan dosyaların yarısı ya da daha fazlası “Anonim” veri etiketine sahip olduğu görülmektedir. Bu oran değerlendirici operatör yükünü artırırken farklı dosya türleri ve formlarından yeterince örnek olmamasından dolayı geliştirilen yapay zeka modellerinin başarısını düşürmektedir.



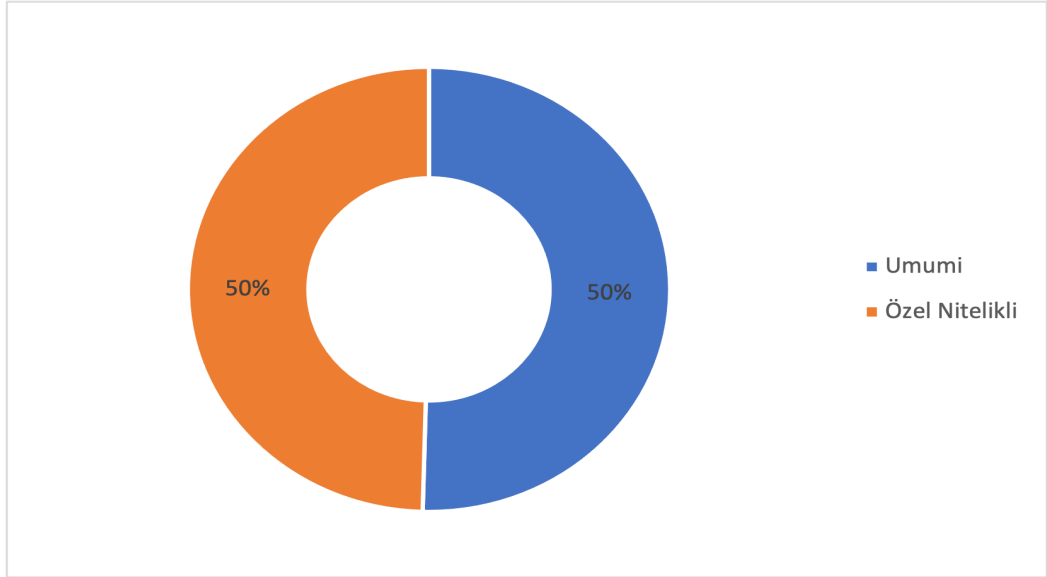
Şekil 4.4: 2 Etiket Model Veriset Seti ve 2 Etiket Operatör Veri Seti karşılaştırması.

Operatör tarafından etiketlenen veri sızıntılarının dört etiket için oranlarının **Şekil 4.5**'te veri seti içindeki dağılım oranları verilmiştir. “Kurumsal Veri” etiketi özel nitelik içeren verilerin yarısını oluştururken “Kişisel Veri” etiketi özel nitelikli verilerinin yaklaşık beşte birini tespit edilebilmektedir. Bu durum operatör tarafından etiketlenen veri setinde kişisel veri tespitinin operatörsüz yapılması için yapay zeka modelini eğitilmesini zorlaştırmaktadır.



Şekil 4.5: 4 Etiket Operatör Veri Seti karşılaştırması.

Operatör tarafından etiketlenen verilerin yarı yarıya çıktığı **Şekil 4.6**'da gösterilmiştir. Bu durum özel nitelikli verilerin tümünün ayrılması için dengeli bir veri seti sağlamaktadır.



Şekil 4.6: 2 Etiket Operatör Veri Seti karşılaştırması.

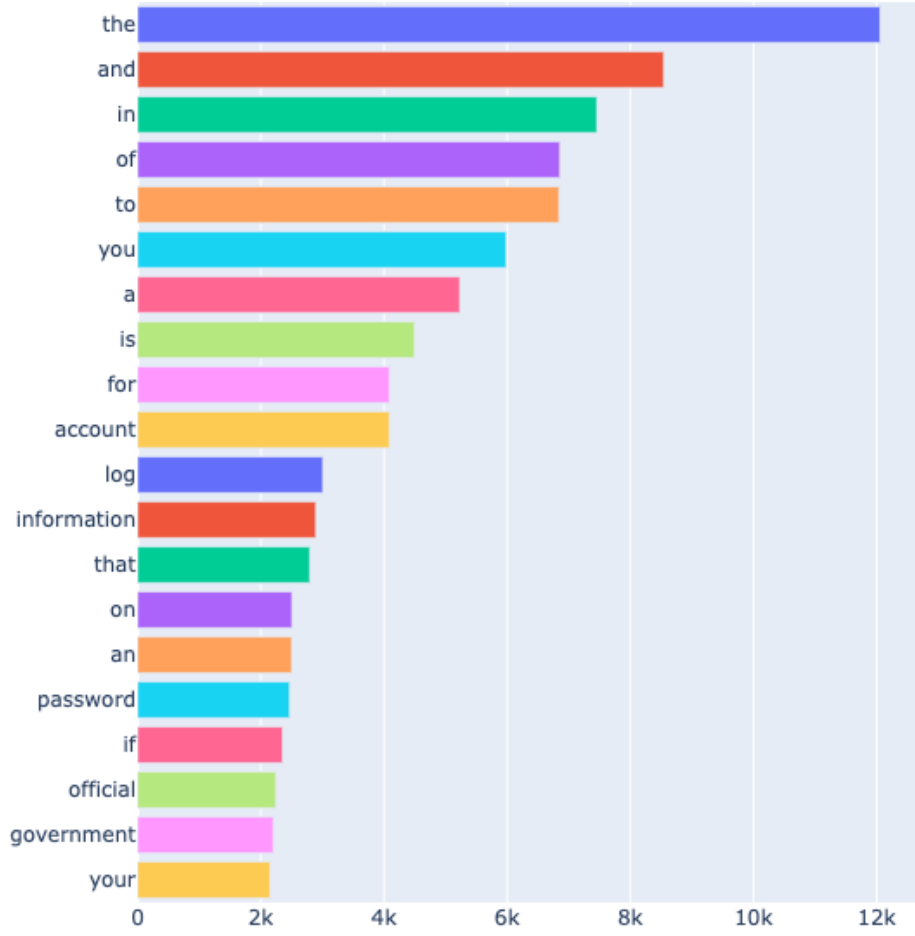
Etkiletime işlemi tamamlandıktan sonra kural tabanlı çıkarılan özelliklerin dışında yapay zeka modellerinde kullanmak üzere NLP operasyonları ile özellik çıkarımı işlemi yapılmıştır. NLP Ön İşleme öncesi boşluk karakterinden yapılan tokenizasyon çalışması sonrası en sık geçen tokenler **Tablo 4.4**'te görülmektedir. Sızıntı veri yapısı kaynaklı özel karakterlerin

ağırlıklı olarak geçtiği görülmüştür. Veri sızıntı dosyalarımız arasında sıklıkla bulunan HTML dosyalara ait karakterin bu listede bulunmamasının sebebi REGEX ve NER operasyonları öncesi sızıntı dosyaları HTML REGEX'i ile temizlendikten sonra operasyonların gerçekleştirilmesinden kaynaklanmaktadır.

Tablo 4.4: NLP Ön işleme Öncesi Veri Seti sık geçen tokenler.

Sıra	Sık Geçen Kelime	Adet
1	(	46860
2	)	37597
3	the	8757
4	and	8318
5	in	7159
6	of	6709
7	to	6665
8	you	5760
9	=	4882
10	a	4704
11	is	4338
12	-	3242
13	{	3094
14	The	3088
15	information	2716
16	that	2715
17	for	2699
18	}	2494
19	on	2378
20	/	2310

Veri kümesi içinde doğrudan bir anlam ifade etmeyen ve tokenizasyona zarar veren karakterlerin temizlik işlemi yapıldıktan sonra **Şekil 4.7**'de tokenlerin geçiş sıklığı görülmektedir. Sık geçen kelimeler arasında “account”, “password”, “information”, “official” ve “government” kelimelerin sıralamaya girdiği görülmektedir. Bu kelimelerin geçtiği dosyalar ile etiketlerin korelasyonu olması halinde sızıntı türü değerlendirmede bu kelimeler kullanılabilir.



Şekil 4.7: Veri Seti sık geçen kelimeler.

Ancak ilk yirmi kelime değerlendirildiğinde içlerinde ağırlığı etkisiz/dolgu kelimelerin (Stop Words) oluşturduğu görülmüştür. Veri dosyaları üzerinde etkisiz kelimelerin temizlenmesi işlemi gerçekleştirilmiştir. Etkisiz kelimelerin temizlenmesi sonrası kelimelerin geçiş sıklığı **Şekil 4.8**'de görülmektedir. İlk 20 kelime değerlendirildiğinde veri etiketleri ile ilişki kurulabileceği düşünülen “log”, “home” ve “search” kelimeleri eklendiği görülmüştür.



Şekil 4.8: Etkisiz kelime temizleme operasyonu sonrası veri seti sık geçen kelimeler.

4.2. Veri Sızıntılarının Sınıflandırılması

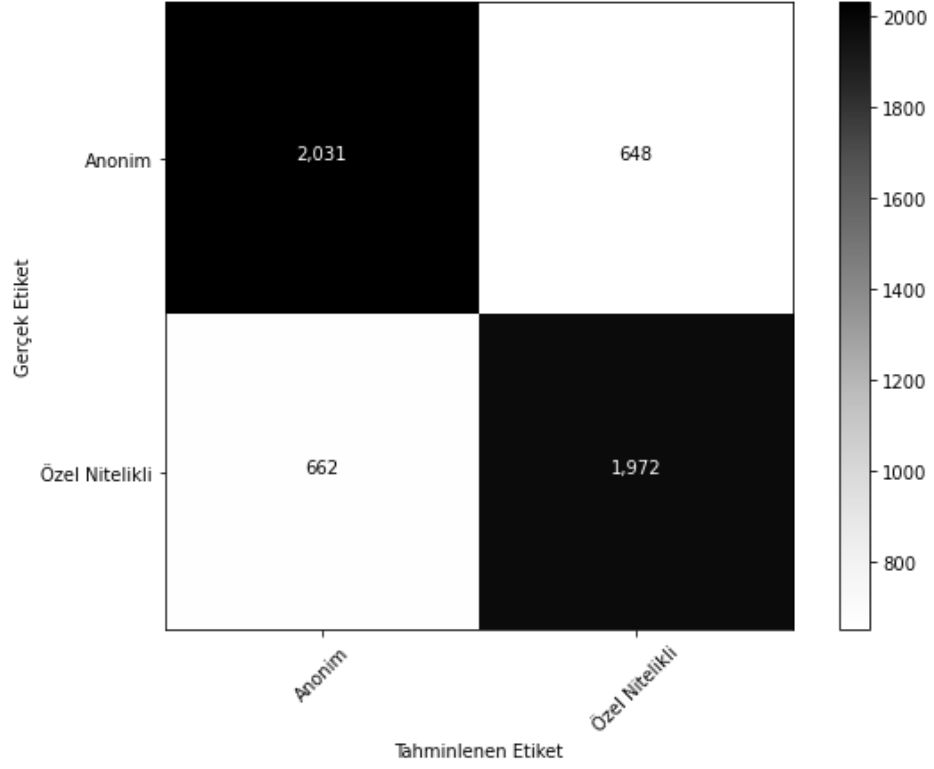
Siber Tehdit İstihbaratı için Yeni Tarama Modeli, sızıntı verilerinin elde edilip, ön işleme operasyonlarından geçirildikten sonra 6 farklı veri seti oluşturmaktadır. İlk ikisi Siber Tehdit İstihbaratı Veri Seti başlığında anlatılan REGEX kurallarından oluşturulan özellikleri içeren REGEX veri seti ve ikincisi NER ile çıkarılan özelliklerden oluşan NER veri setidir. Bu iki özellik kümesi ayrı ve birlikte olarak farklı makine öğrenmesi yöntemlerine sokularak performans metrikleri gözlemlenmiştir. Diğer 4 veri seti ise Sızıntı Verilerinin Analizi kısmında açıklanan “4 Etiket Model Veri Seti”, “2 Etiket Model Veri Seti”, “4 Etiket Operatör Veri Seti” ve “2 Etiket Operatör Veri Seti” olmak üzere doğal dil işleme yöntemleri ile özellik çıkarımı yapılarak yapay zeka yöntemleri ile sınıflandırılmıştır.

İlk olarak ikili sınıflandırma için “Anonim” ve “Özel Nitelikli” (bilgi içeren) verilerin ayrıştırılması yapılmıştır. KNN algoritması üzerinde REGEX ve NER özelliklerinden oluşan veri seti kullanılarak eğitilmiştir. k değeri 4 olup veri seti %70 eğitim %30 test olacak şekilde bölünmüştür. KNN algoritması performans metrikleri **Tablo 4.5**’te gösterilmiştir.

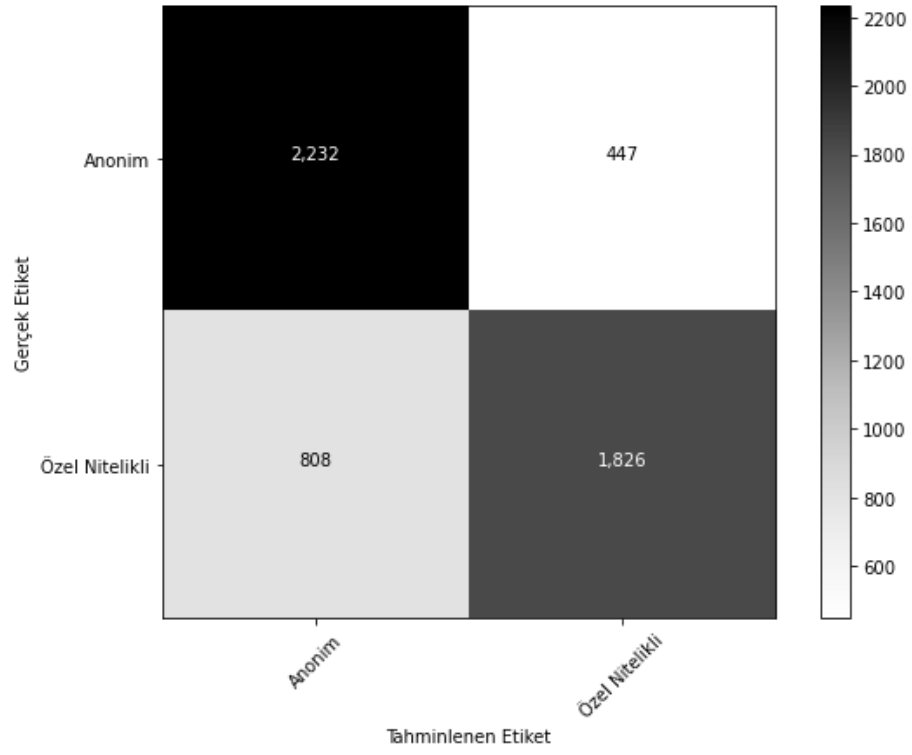
Tablo 4.5: KNN algoritması REGEX ve NER veri setleri model performans metrikleri.

Veri Seti	Doğruluk	Kesinlik	Duyarlılık	F1-Skoru
REGEX Veri Seti	%75	%75	%75	%75
NER Veri Seti	%76	%77	%76	%76
REGEX + NER Veri Seti	%86	%88	%86	%86

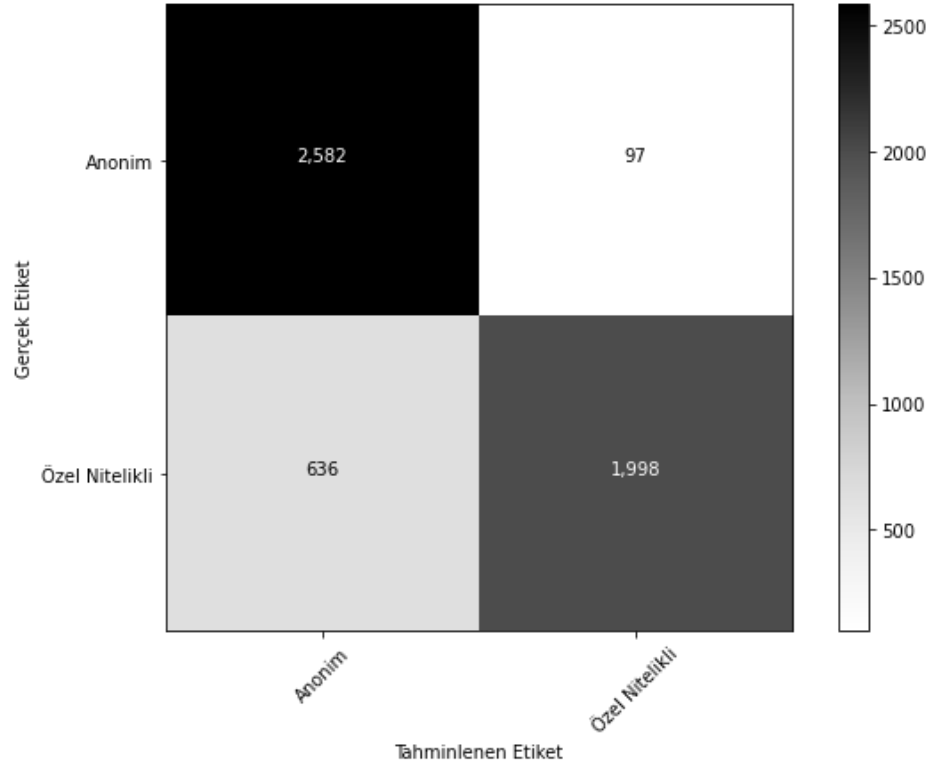
REGEX Veri Seti KNN algoritması ile REGEX Veri Seti için karmaşıklık matrisi **Şekil 4.9**'da, NER Veri Seti için karmaşıklık matrisi **Şekil 4.10**'da ve REGEX + NER Veri Seti için karmaşıklık matrisi **Şekil 4.11**'de görülmektedir.



Şekil 4.9: KNN Algoritması REGEX Veri Seti Karmaşıklık Matrisi.



Şekil 4.10: KNN Algoritması NER Veri Seti Karmaşıklık Matrisi.



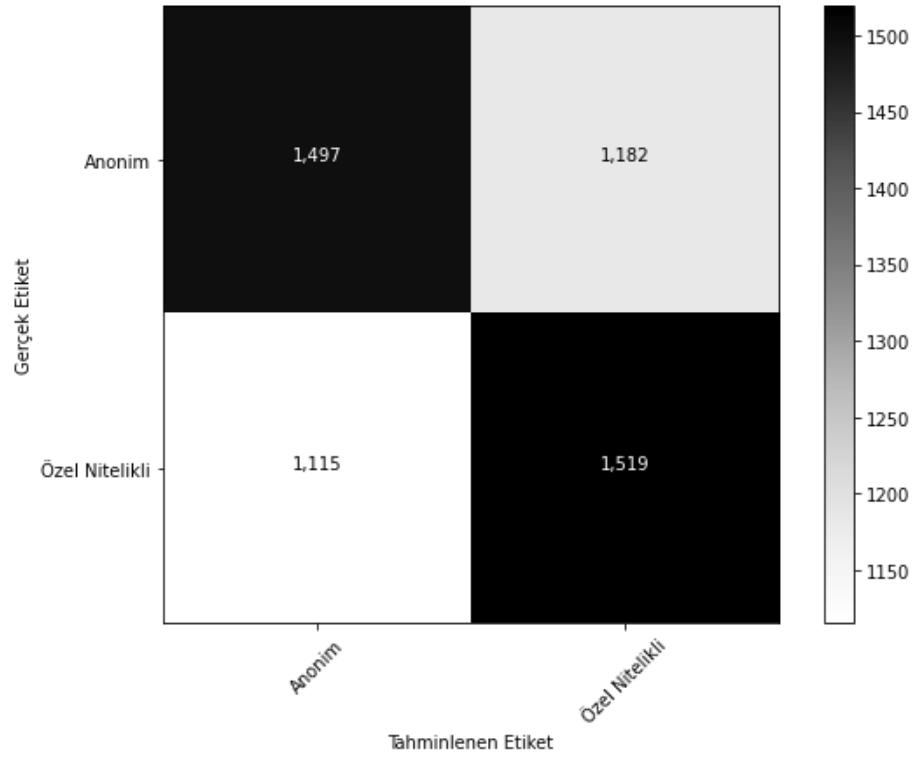
Şekil 4.11: KNN Algoritması REGEX ve NER Veri Seti Karmaşıklık Matrisi.

Naive Bayes algoritması üzerinde REGEX ve NER özelliklerinden oluşan veri seti kullanılarak eğitilmiştir. Veri seti %70 eğitim %30 test olacak şekilde bölünmüştür. Naive Bayes algoritması performans metrikleri **Tablo 4.6**'da gösterilmiştir.

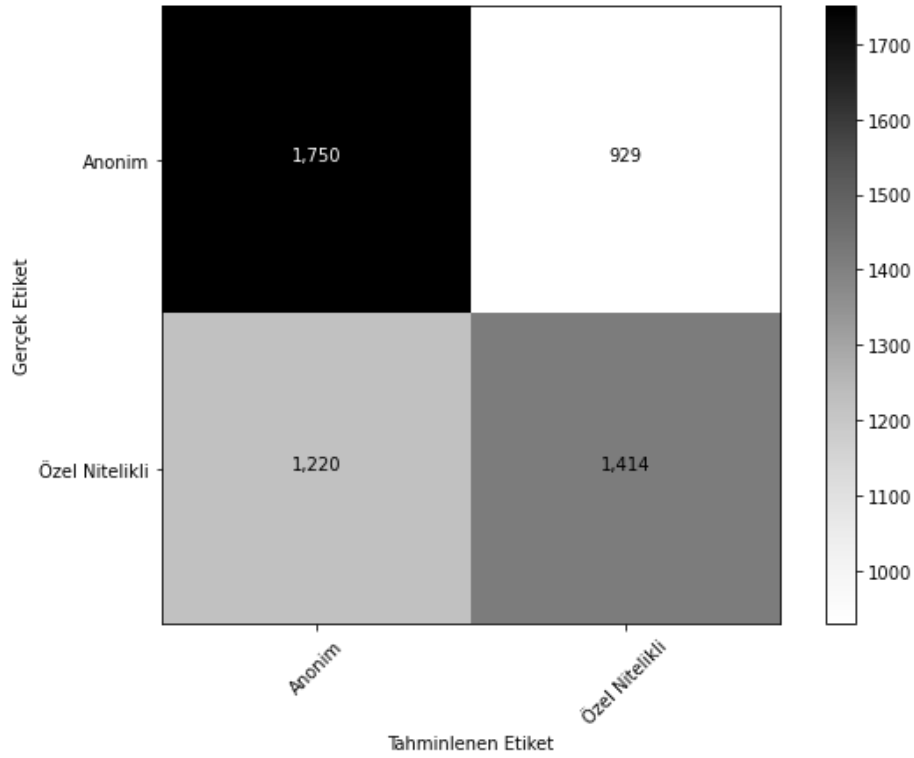
Tablo 4.6: Naive Bayes algoritması REGEX ve NER veri setleri model performans metrikleri

Veri Seti	Doğruluk	Kesinlik	Duyarlılık	F1-Skoru
REGEX Veri Seti	%57	%57	%57	%57
NER Veri Seti	%60	%60	%60	%59
REGEX + NER Veri Seti	%67	%73	%67	%64

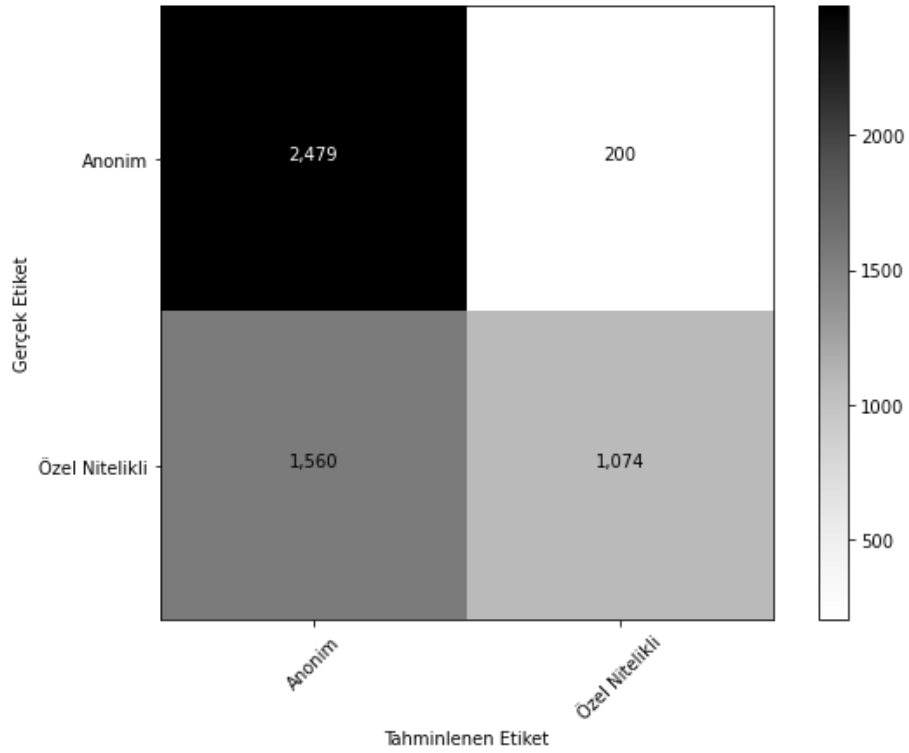
Naive Bayes algoritması ile REGEX Veri Seti için karmaşıklık matrisi **Şekil 4.12**'de, NER Veri Seti için karmaşıklık matrisi **Şekil 4.13**'te ve REGEX + NER Veri Seti için karmaşıklık matrisi **Şekil 4.14**'de görülmektedir.



Şekil 4.12: Naive Bayes Algoritması REGEX Veri Seti Karmaşıklık Matrisi.



Şekil 4.13: Naive Bayes Algoritması NER Veri Seti Karmaşıklık Matrisi.



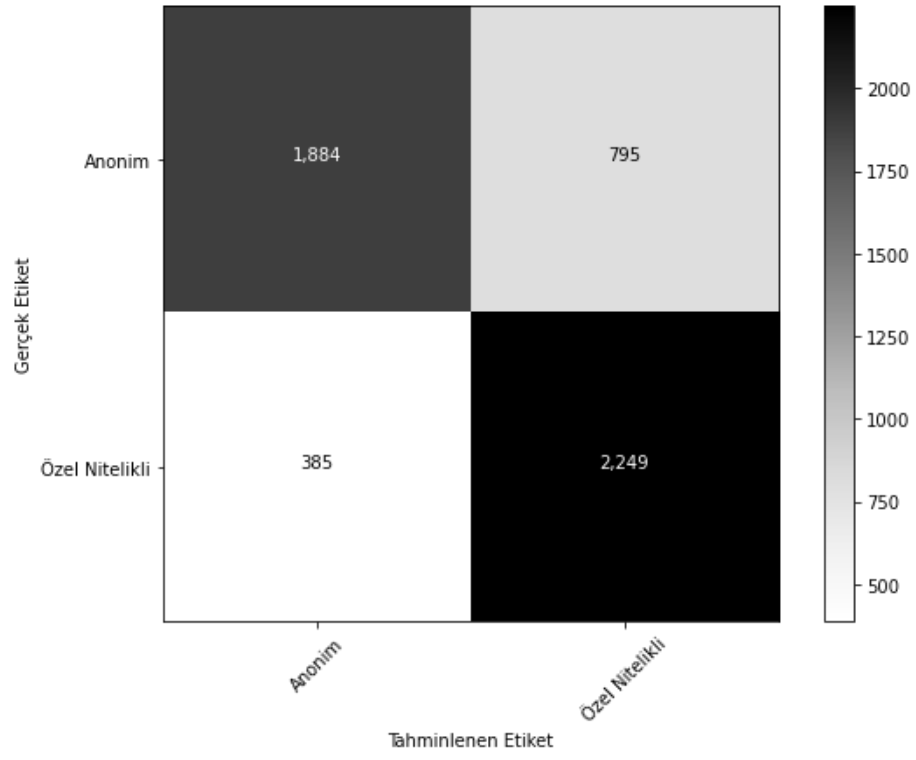
Şekil 4.14: Naive Bayes Algoritması REGEX ve NER Veri Seti Karmaşıklık Matrisi.

Rastgele Orman algoritması üzerinde REGEX ve NER özelliklerinden oluşan veri seti kullanılarak eğitilmiştir. Veri seti %70 eğitim %30 test olacak şekilde bölünmüştür. Rastgele Orman algoritması performans metrikleri **Tablo 4.7**'te gösterilmiştir.

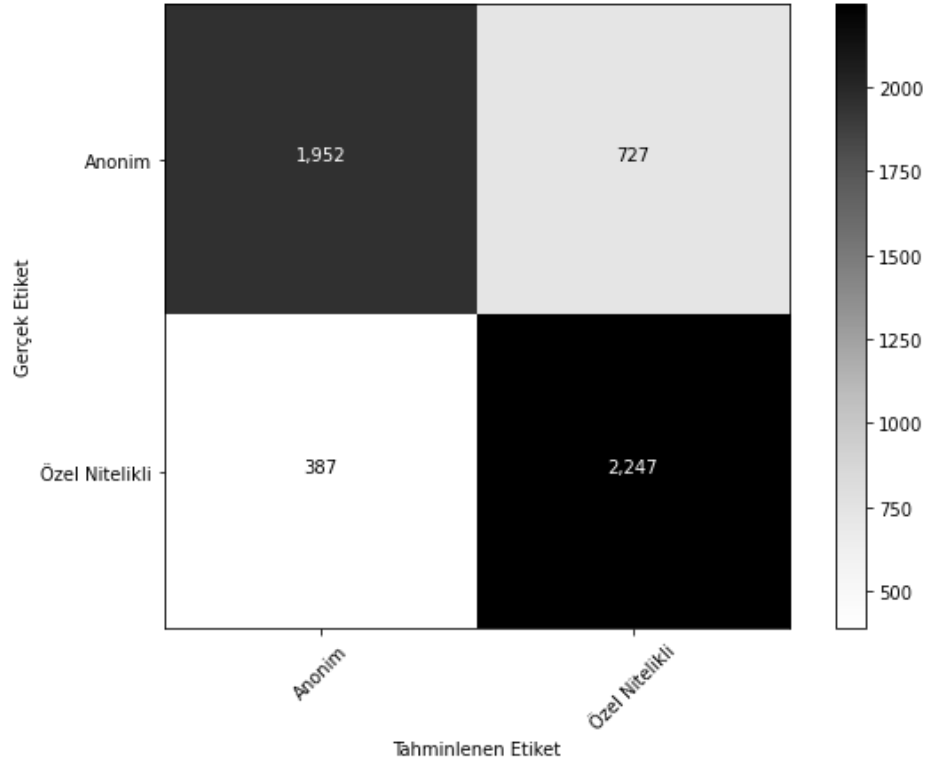
Tablo 4.7: Rastgele Orman algoritması REGEX ve NER veri setleri model performans metrikleri

Veri Seti	Doğruluk	Kesinlik	Duyarlılık	F1-Skoru
REGEX Veri Seti	%78	%78	%78	%78
NER Veri Seti	%79	%80	%79	%79
REGEX + NER Veri Seti	%88	%88	%87	%87

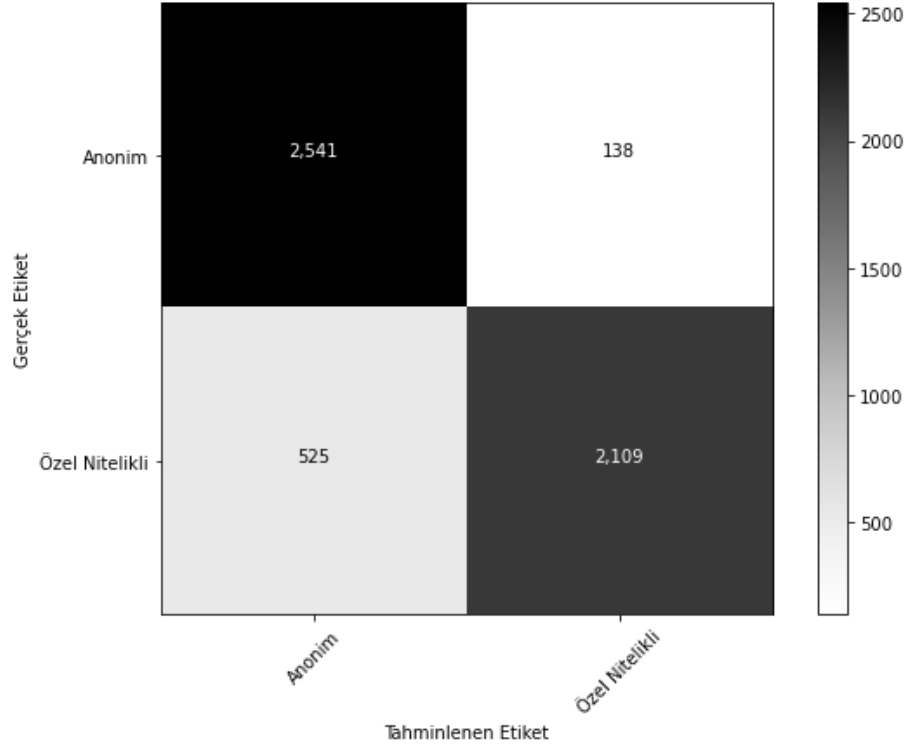
Rastgele Orman algoritması ile REGEX Veri Seti için karmaşıklık matrisi **Şekil 4.15**'de, NER Veri Seti için karmaşıklık matrisi **Şekil 4.16**'de ve REGEX + NER Veri Seti için karmaşıklık matrisi **Şekil 4.17**'de görülmektedir.



Şekil 4.15: Rastgele Orman Algoritması REGEX Veri Seti Karmaşıklık Matrisi.



Şekil 4.16: Rastgele Orman Algoritması NER Veri Seti Karmaşıklık Matrisi.



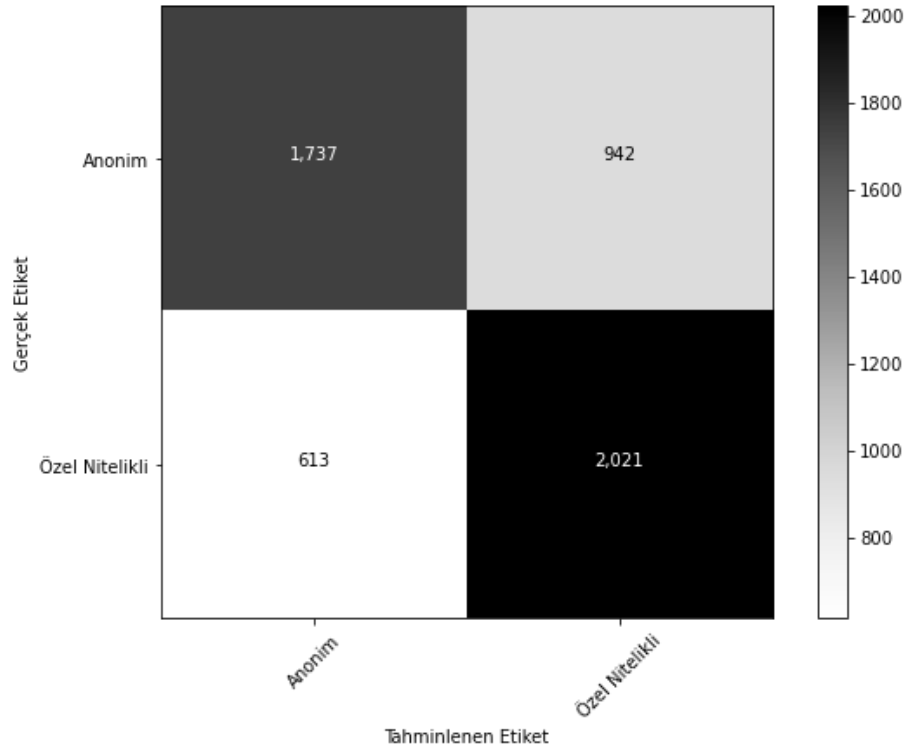
Şekil 4.17: Rastgele Orman Algoritması REGEX ve NER Veri Seti Karmaşıklık Matrisi.

Destek Vektör Makinesi algoritması üzerinde REGEX ve NER özelliklerinden oluşan veri seti kullanılarak eğitilmiştir. Veri seti %70 eğitim %30 test olacak şekilde bölünmüştür. SVM algoritması performans metrikleri **Tablo 4.8**'te gösterilmiştir.

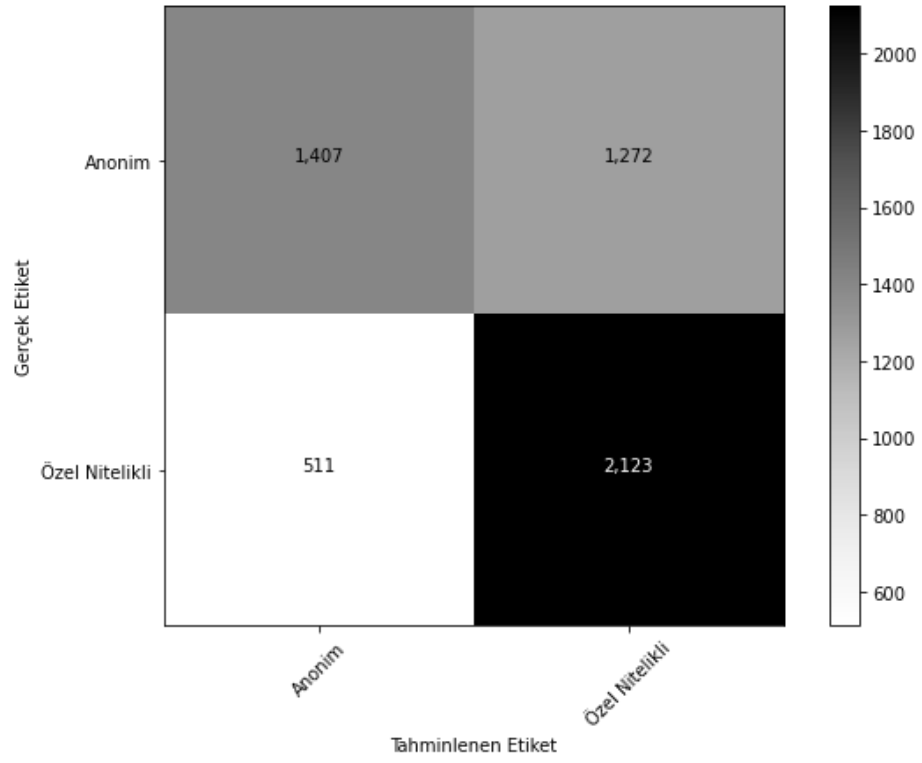
Tablo 4.8: SVM algoritması REGEX ve NER veri setleri model performans metrikleri

Veri Seti	Doğruluk	Kesinlik	Duyarlılık	F1-Skoru
REGEX Veri Seti	%71	%71	%71	%71
NER Veri Seti	%66	%68	%67	%66
REGEX + NER Veri Seti	%84	%82	%82	%82

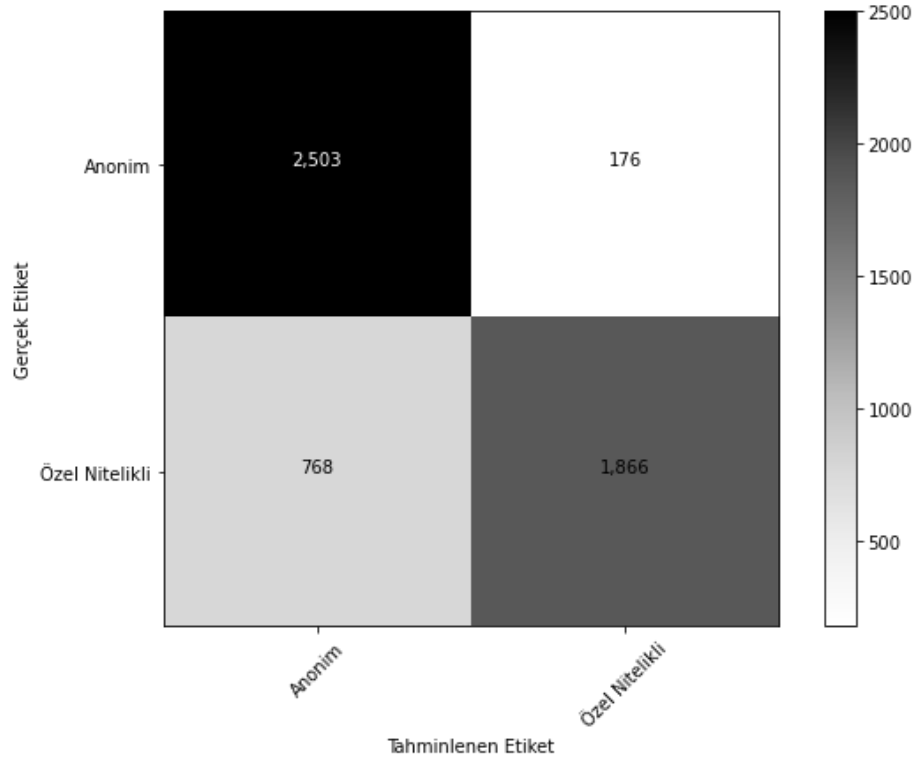
SVM algoritması ile REGEX Veri Seti için karmaşıklık matrisi **Şekil 4.18**'de, NER Veri Seti için karmaşıklık matrisi **Şekil 4.19**'de ve REGEX + NER Veri Seti için karmaşıklık matrisi **Şekil 4.20**'de görülmektedir.



Şekil 4.18: SVM Algoritması REGEX Veri Seti Karmaşıklık Matrisi.

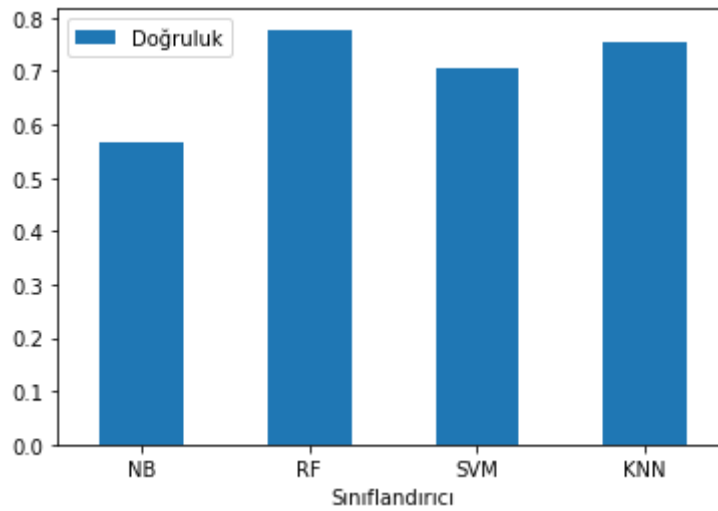


Şekil 4.19: SVM Algoritması NER Veri Seti Karmaşıklık Matrisi



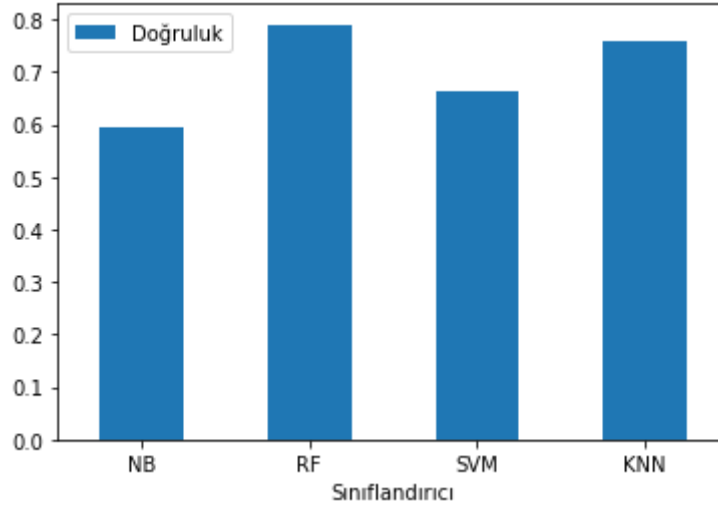
Şekil 4.20: SVM Algoritması REGEX ve NER Veri Seti Karmaşıklık Matrisi.

KNN, Naive Bayes, Rastgele Orman ve SVM algoritmalarından REGEX Veri Seti üzerinde doğruluk değeri en iyi sonuç %78 ile Rastgele Orman Algoritması ile elde edilmiştir. Algoritmaların REGEX Veri Seti üzerinde doğruluk değeri karşılaştırması **Şekil 4.21**'de x ekseninde sınıflandırıcılar y ekseninde 0-1 normalizasyon ile sınıflandırıcı doğrulukları verilmiştir.



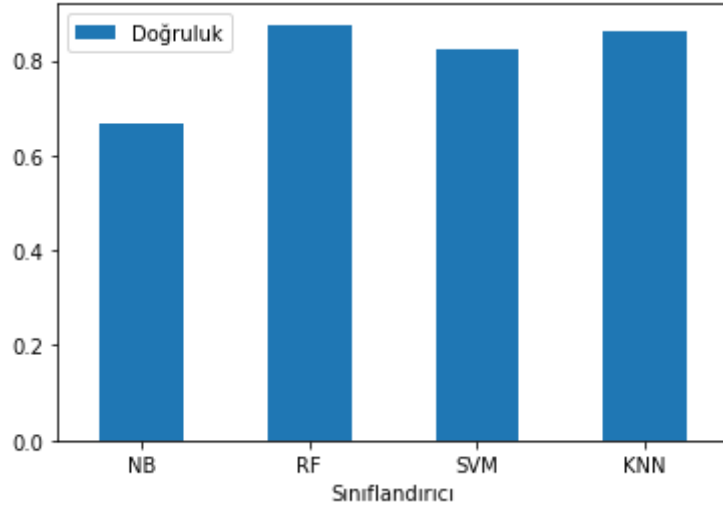
Şekil 4.21: REGEX Veri Seti üzerinde doğruluk değeri karşılaştırması.

KNN, Naive Bayes, Rastgele Orman ve SVM algoritmalarından NER Veri Seti üzerinde doğruluk değeri en iyi sonuç %79 ile Rastgele Orman Algoritması'nda elde edilmiştir. Algoritmaların NER Veri Seti üzerinde doğruluk değeri karşılaştırması **Şekil 4.22**'de x ekseninde sınıflandırıcılar y ekseninde 0-1 normalizasyon ile sınıflandırıcı doğrulukları verilmiştir.



Şekil 4.22: NER Veri Seti üzerinde doğruluk değeri karşılaştırması.

KNN, Naive Bayes, Rastgele Orman ve SVM algoritmalarından REGEX+NER Veri Seti üzerinde doğruluk değeri en iyi sonuç %88 ile Rastgele Orman Algoritmasından elde edilmiştir. Algoritmaların REGEX+NER Veri Seti üzerinde doğruluk değeri karşılaştırması **Şekil 4.23**'de x ekseninde sınıflandırıcılar y ekseninde 0-1 normalizasyon ile sınıflandırıcı doğrulukları verilmiştir.



Şekil 4.23: REGEX+NER Veri Seti üzerinde doğruluk değeri karşılaştırması.

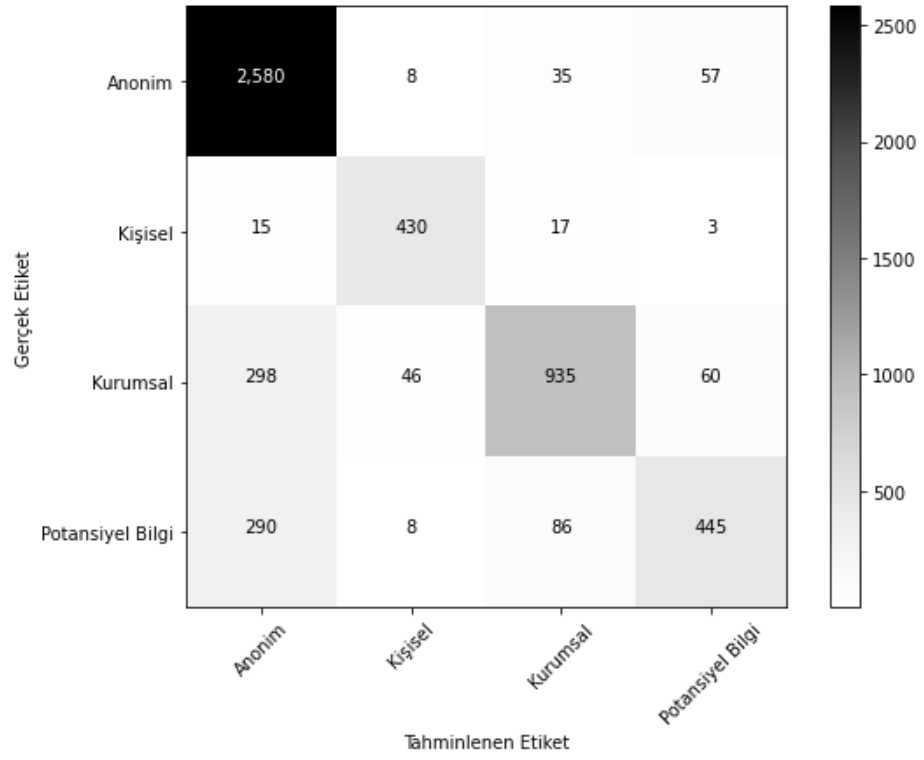
Sınıflandırıcıların performans metrikleri değerlendirildiğinde en iyi performans metriklerine Rastgele Orman algoritması ile ulaşıldığı, sonraki en iyi sonucu ise KNN algoritmasının sağladığı, en başarısız sonuca ise Naive Bayes algoritması ile ulaşıldığı görülmüştür.

İkinci olarak çoklu sınıflandırma için “Anonim”, “Kişisel Veri”, “Kurumsal Veri” ve “Potansiyel Bilgi” verilerin ayrıştırılması yapılmıştır. KNN, Rastgele Orman, Naive Bayes ve SVM algoritmaları üzerinde REGEX ve NER özelliklerinden oluşan veri seti kullanılarak eğitilmiştir. KNN algoritması ile veri seti %70 eğitim %30 test olacak şekilde bölünmüştür. KNN algoritması performans metrikleri **Tablo 4.9**’da gösterilmiştir.

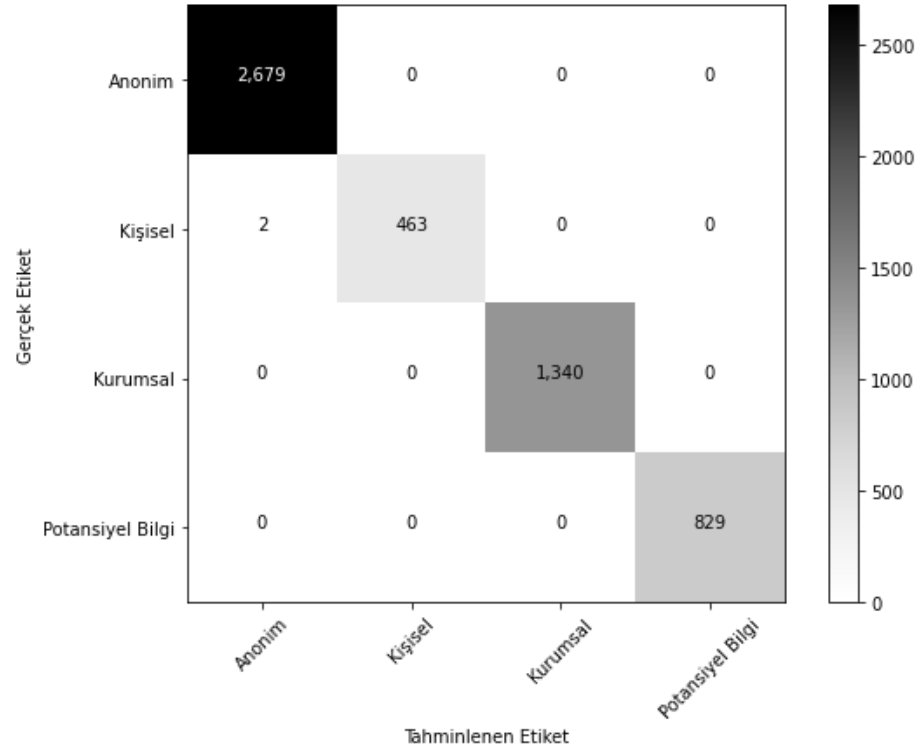
Tablo 4.9: KNN algoritması NER veri setleri model performans metrikleri

Veri Seti	Doğruluk	Kesinlik	Duyarlılık	F1-Skoru
REGEX Veri Seti 4 Etiket	%83	%84	%78	%80
NER Veri Seti 4 Etiket	%100	%100	%100	%100
REGEX + NER Veri Seti 4 Etiket	%87	%86	%84	%85

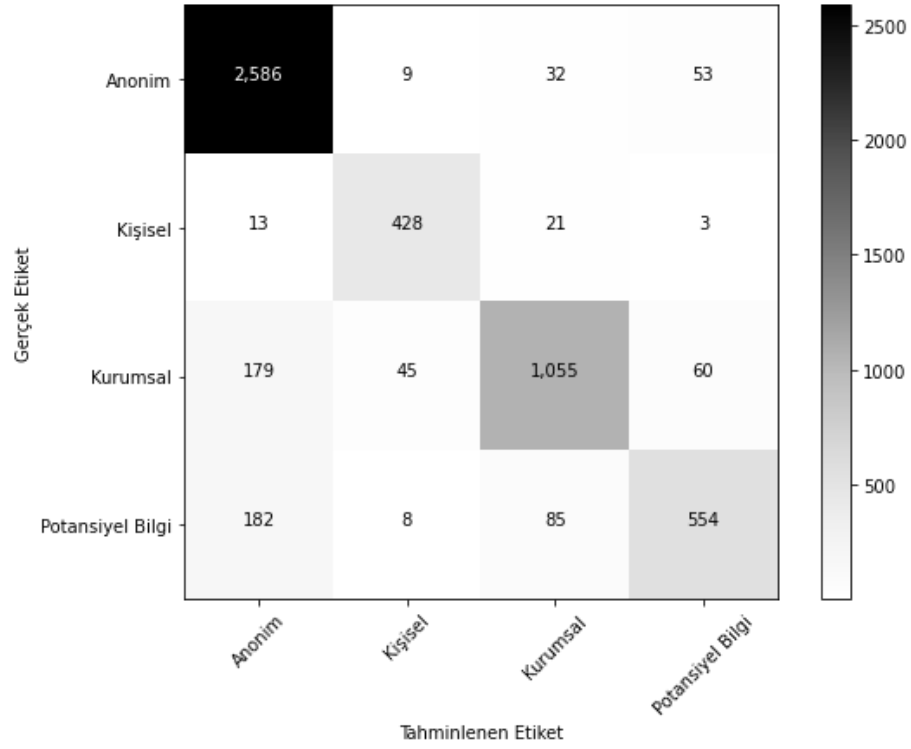
KNN algoritması ile NER Veri Seti için çoklu etiket karmaşıklık matrisi **Şekil 4.24**’de, NER Veri Seti için karmaşıklık matrisi **Şekil 4.25**’de ve NER Veri Seti için karmaşıklık matrisi **Şekil 4.26**’da görülmektedir.



Şekil 4.24: KNN Algoritması çoklu etiket REGEX Veri Seti Karmaşıklık Matrisi.



Şekil 4.25: KNN Algoritması çoklu etiket NER Veri Seti Karmaşıklık Matrisi.



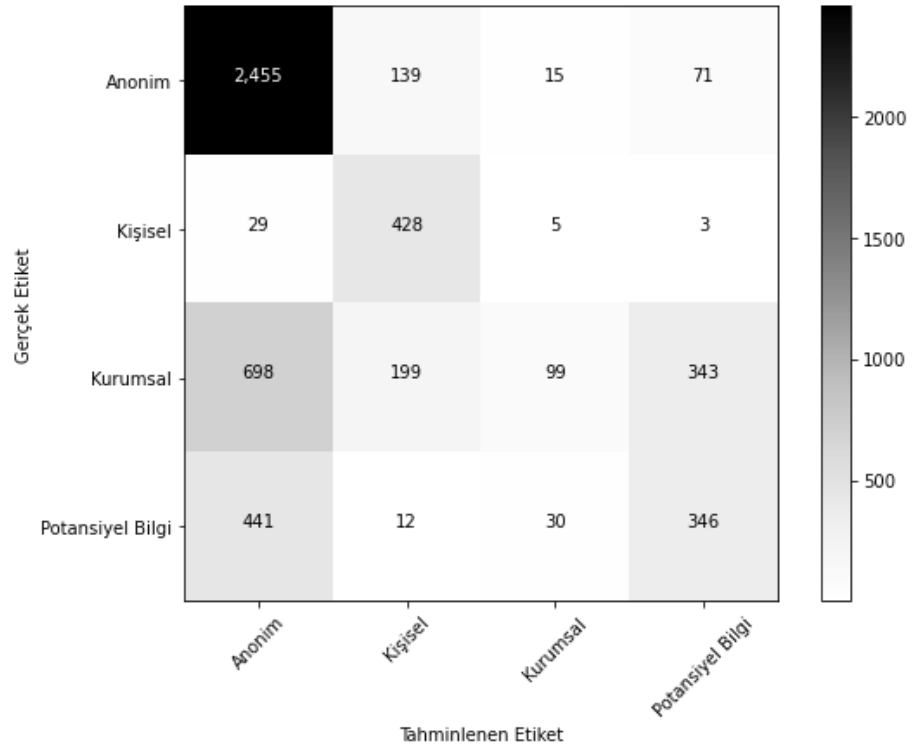
Şekil 4.26: KNN Algoritması çoklu etiket REGEX ve NER Veri Seti Karmaşıklık Matrisi.

Naive Bayes algoritması üzerinde çoklu etiket REGEX ve NER özelliklerinden oluşan veri seti kullanılarak eğitilmiştir. Veri seti %70 eğitim %30 test olacak şekilde bölünmüştür. Naive Bayes algoritması performans metrikleri **Tablo 4.10**'da gösterilmiştir.

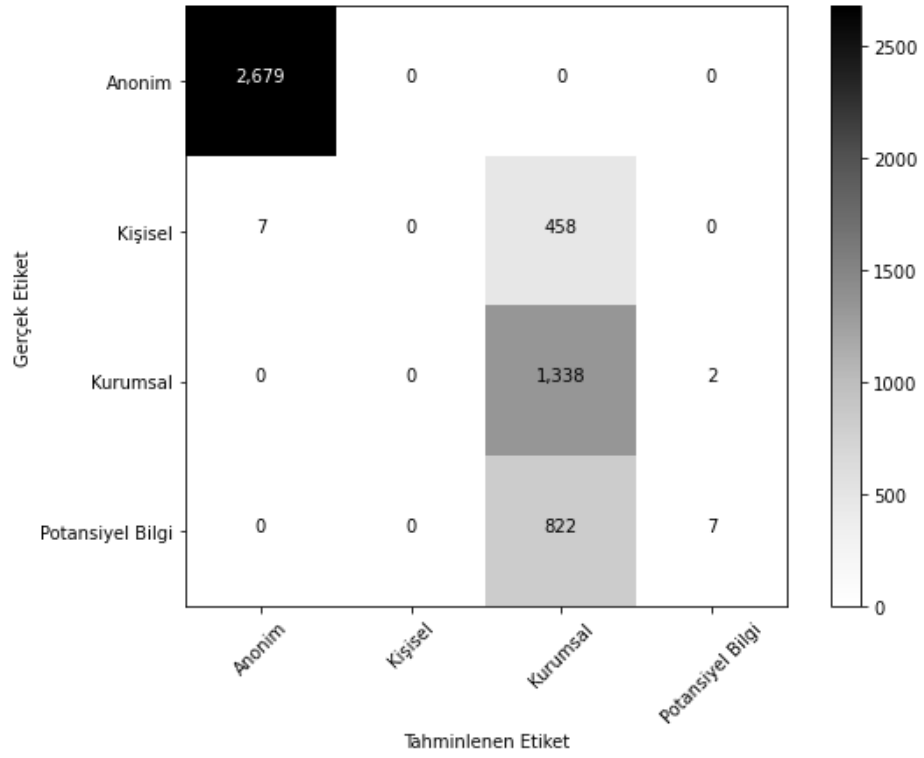
Tablo 4.10: Naive Bayes algoritması çoklu etiket NER veri setleri model performans metrikleri

Veri Seti	Doğruluk	Kesinlik	Duvarlılık	F1-Skoru
REGEX Veri Seti 4 Etiket	%63	%59	%58	%51
NER Veri Seti 4 Etiket	%76	%56	%50	%42
REGEX + NER Veri Seti 4 Etiket	%67	%64	%61	%57

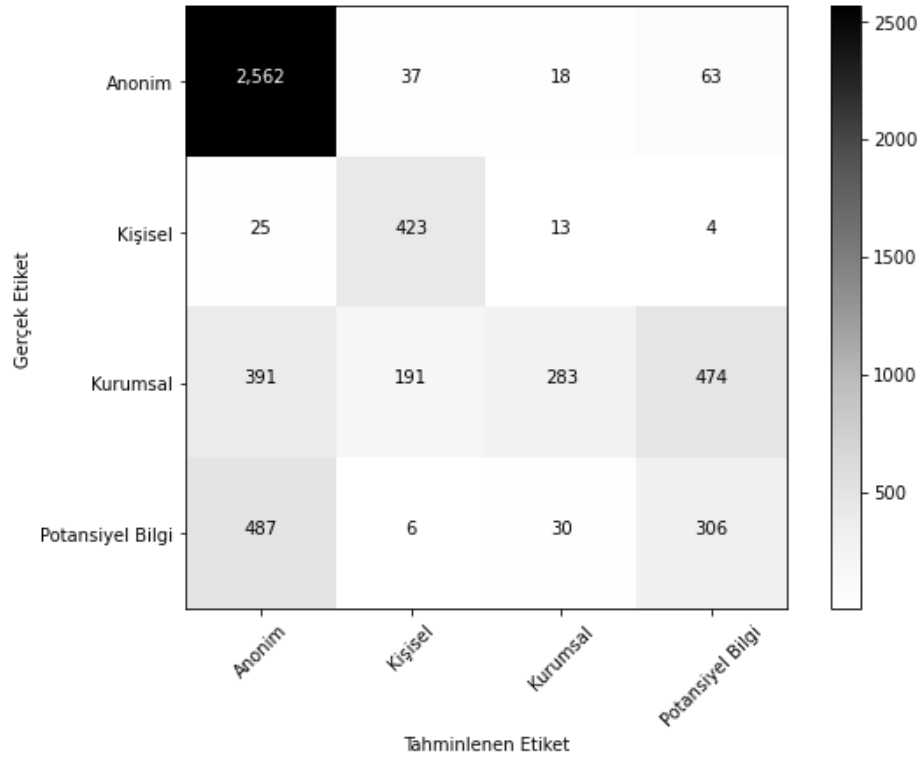
Naive Bayes algoritması ile NER Veri Seti için çoklu etiket karmaşıklık matrisi **Şekil 4.27**'de, NER Veri Seti için karmaşıklık matrisi **Şekil 4.28**'de ve NER Veri Seti için karmaşıklık matrisi **Şekil 4.29**'da görülmektedir.



Şekil 4.27: Naive Bayes Algoritması çoklu etiket REGEX Veri Seti Karmaşıklık Matrisi.



Şekil 4.28: Naive Bayes Algoritması çoklu etiket NER Veri Seti Karmaşıklık Matrisi.



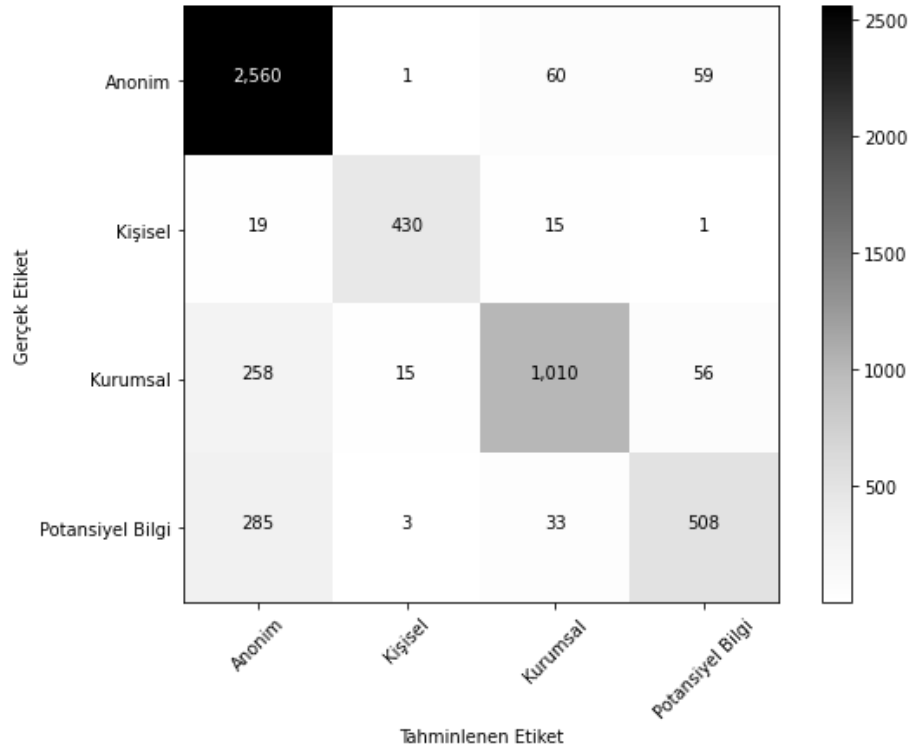
Şekil 4.29: Naive Bayes Algoritması çoklu etiket REGEX ve NER Veri Seti Karmaşıklık Matrisi.

Rastgele Orman algoritması üzerinde çoklu etiket REGEX ve NER özelliklerinden oluşan veri seti kullanılarak eğitilmiştir. Veri seti %70 eğitim %30 test olacak şekilde bölünmüştür. Naive Bayes algoritması performans metrikleri **Tablo 4.11**'de gösterilmiştir.

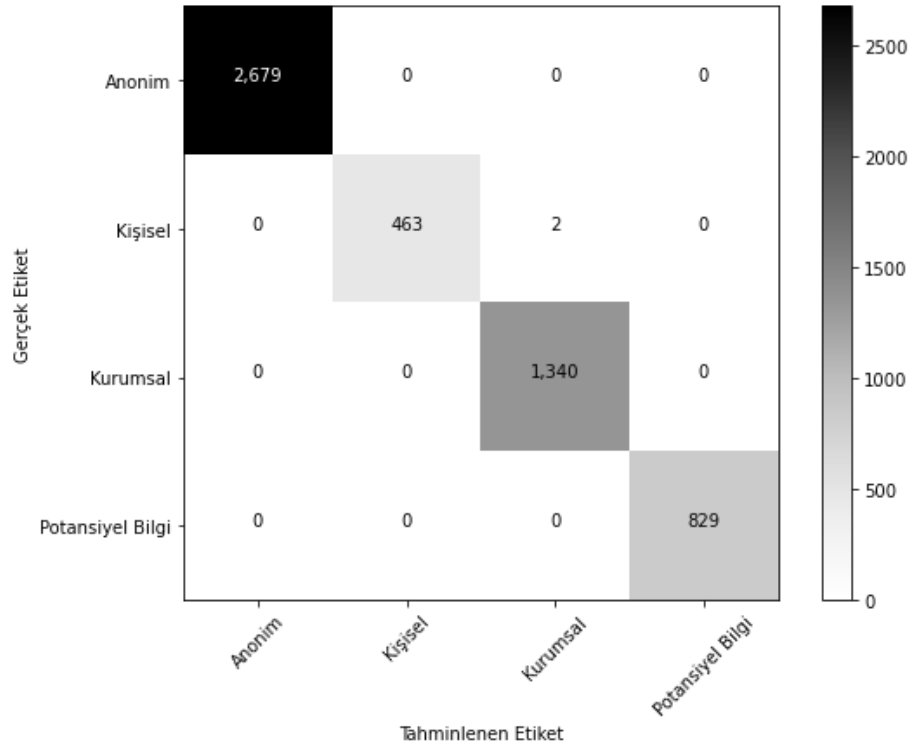
Tablo 4.11: Rastgele Orman algoritması çoklu etiket NER veri setleri model performans metrikleri

Veri Seti	Doğruluk	Kesinlik	Duyarlılık	F1-Skoru
REGEX Veri Seti 4 Etiket	%85	%87	%81	%84
NER Veri Seti 4 Etiket	%100	%100	%100	%100
REGEX + NER Veri Seti 4 Etiket	%89	%91	%85	%88

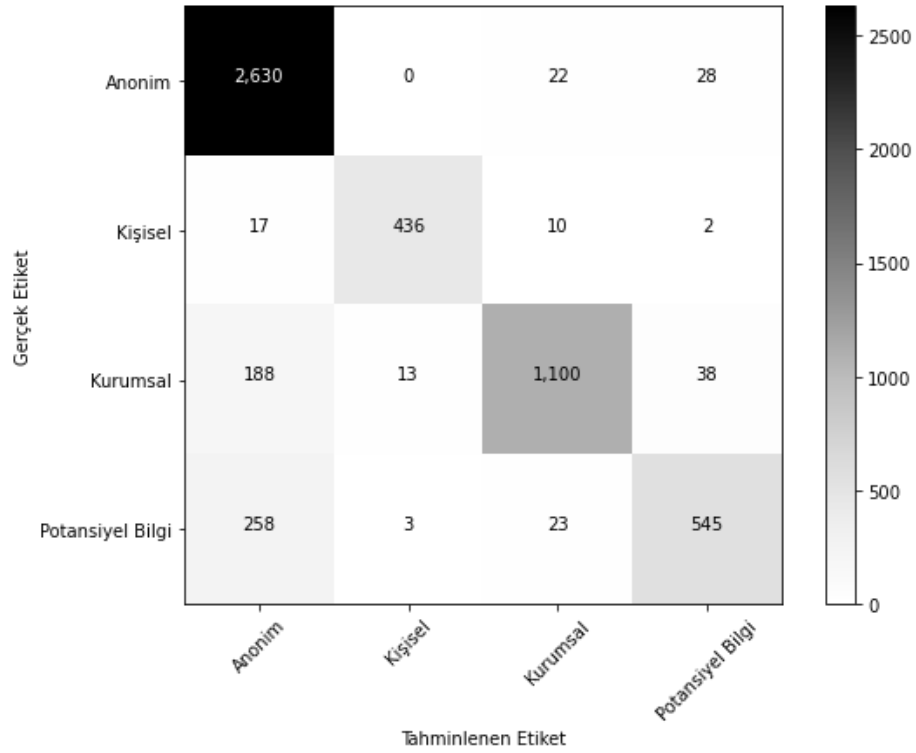
Rastgele Orman algoritması ile NER Veri Seti için çoklu etiket karmaşıklık matrisi **Şekil 4.30**'de, NER Veri Seti için karmaşıklık matrisi **Şekil 4.31**'de ve NER Veri Seti için karmaşıklık matrisi **Şekil 4.32**'de görülmektedir.



Şekil 4.30: Rastgele Orman Algoritması çoklu etiket REGEX Veri Seti Karmaşıklık Matrisi.



Şekil 4.31: Rastgele Orman Algoritması çoklu etiket NER Veri Seti Karmaşıklık Matrisi.



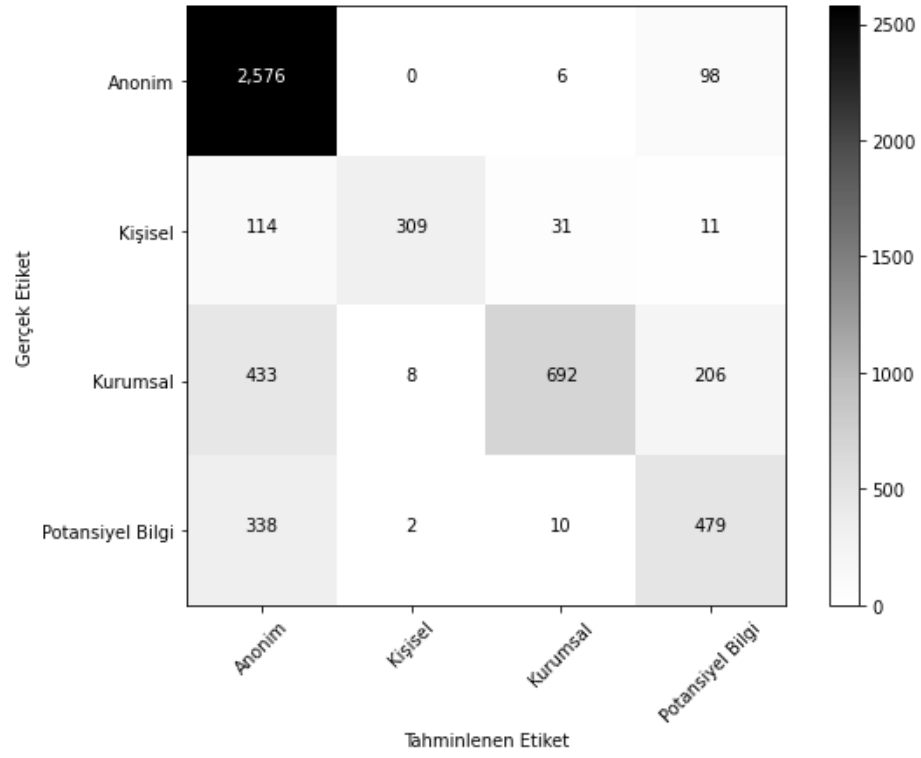
Şekil 4.32: Rastgele Orman Algoritması çoklu etiket REGEX ve NER Veri Seti Karmaşıklık Matrisi.

SVM algoritması üzerinde çoklu etiket REGEX ve NER özelliklerinden oluşan veri seti kullanılarak eğitilmiştir. Veri seti %70 eğitim %30 test olacak şekilde bölünmüştür. Naive Bayes algoritması performans metrikleri **Tablo 4.12**'te gösterilmiştir.

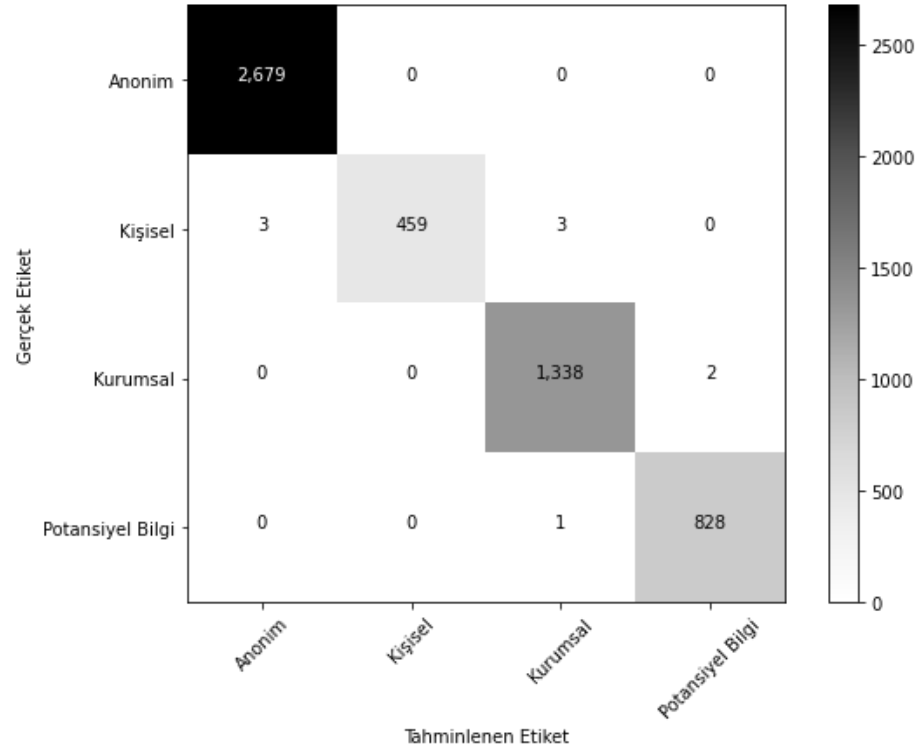
Tablo 4.12: SVM algoritması çoklu etiket NER veri setleri model performans metrikleri

Veri Seti	Doğruluk	Kesinlik	Duyarlılık	F1-Skoru
REGEX Veri Seti 4 Etiket	%76	%81	%68	%72
NER Veri Seti 4 Etiket	%100	%100	%100	%100
REGEX + NER Veri Seti 4 Etiket	%79	%82	%72	%76

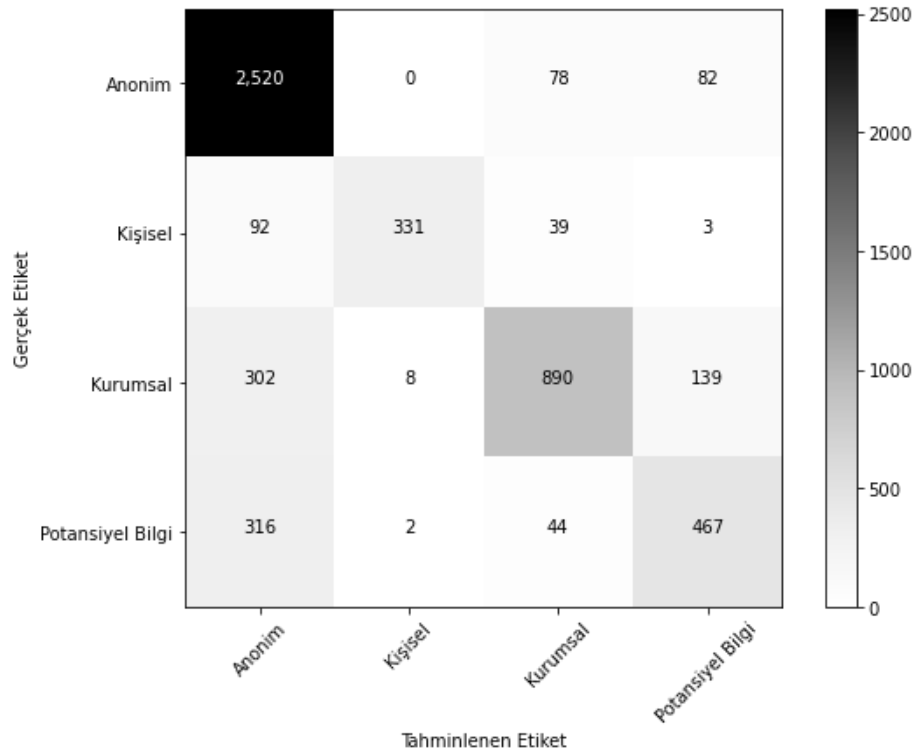
SVM algoritması ile NER Veri Seti için çoklu etiket karmaşıklık matrisi **Şekil 4.33**'de, NER Veri Seti için karmaşıklık matrisi **Şekil 4.34**'de ve NER Veri Seti için karmaşıklık matrisi **Şekil 4.35**'de görülmektedir.



Şekil 4.33: SVM Algoritması çoklu etiket REGEX Veri Seti Karmaşıklık Matrisi.

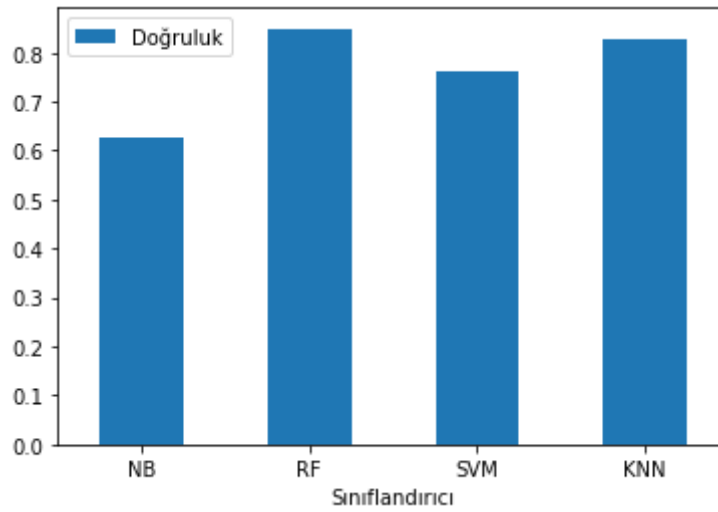


Şekil 4.34: SVM Algoritması çoklu etiket NER Veri Seti Karmaşıklık Matrisi.



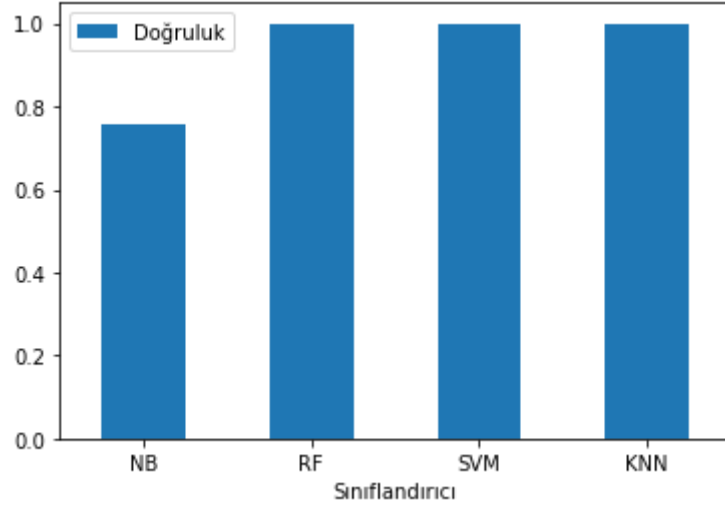
Şekil 4.35: SVM Algoritması çoklu etiket REGEX ve NER Veri Seti Karmaşıklık Matrisi.

KNN, Naive Bayes, Rastgele Orman ve SVM algoritmalarından çoklu etiket REGEX Veri Seti üzerinde doğruluk değeri en iyi sonucu %100 ile KNN ve Rastgele Orman Algoritması vermiştir. Algoritmaların çoklu etiket REGEX Veri Seti üzerinde doğruluk değeri karşılaştırması **Şekil 4.36**'da verilmiştir.



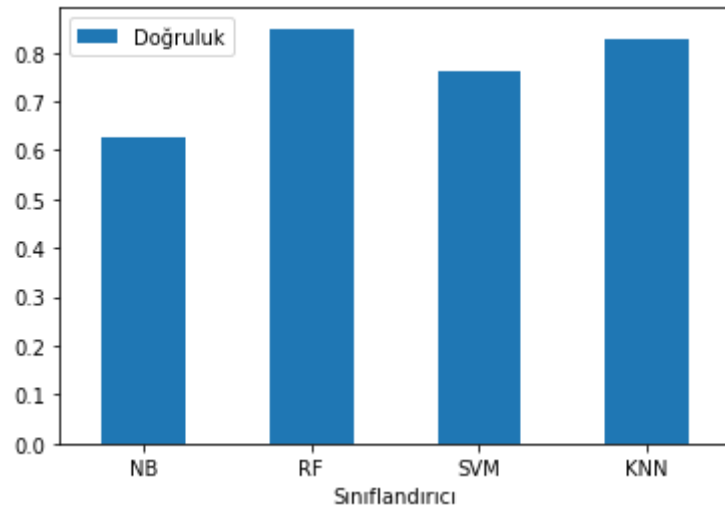
Şekil 4.36: Çoklu etiket REGEX Veri Seti üzerinde doğruluk değeri karşılaştırması.

KNN, Naive Bayes, Rastgele Orman ve SVM algoritmalarından çoklu etiket NER Veri Seti üzerinde doğruluk değeri en iyi sonuç %99 ile KNN SVM ve Rastgele Orman Algoritması vermiştir. Algoritmaların çoklu etiket NER Veri Seti üzerinde doğruluk değeri karşılaştırması **Şekil 4.37**'de verilmiştir.



Şekil 4.37: Çoklu etiket NER Veri Seti üzerinde doğruluk değeri karşılaştırması.

KNN, Naive Bayes, Rastgele Orman ve SVM algoritmalarından çoklu etiket REGEX+NER Veri Seti üzerinde doğruluk değeri en iyi sonuç %84 ile Rastgele Orman Algoritması vermiştir. Algoritmaların çoklu etiket REGEX+NER Veri Seti üzerinde doğruluk değeri karşılaştırması **Şekil 4.38**'de verilmiştir.



Şekil 4.38: Çoklu etiket REGEX+NER Veri Seti üzerinde doğruluk değeri karşılaştırması

Sınıflandırıcıların performans metrikleri değerlendirildiğinde en iyi performans metriklerine Rastgele Orman algoritmasının gösterdiği sonraki en iyi sonucu KNN algoritması sağlarken en başarısız sonucu Naive Bayes algoritması ulaşmıştır.

Üçüncü olarak NLP yöntemleri kullanarak sızıntı verisi içerisinde geçen kelime ve karakterleri RNN ile eğiterek performans metrikleri karşılaştırılmıştır. NLP operasyonu öncesi dosyalarda bulunan anlamsız karakterlerin ve dolgu kelimelerinin temizliği yapılmıştır. Ağ eğitiminde operatör tarafından etiketlenen veri baz alınarak gerçekleştirilmiştir RNN modelinden alınan sonuçlar **Tablo 4.13**'te görülmektedir.

Tablo 4.13: RNN algoritması çoklu etiket Operatör Veri Setleri için model performans metrikleri.

Veri Seti	Doğruluk	Kesinlik	Duyarlılık	F1-Skoru
4 Etiket Operatör Veri Seti	%71	%65	%60	%54
2 Etiket Operatör Veri Seti	%71	%65	%61	%54

4.3. Kullanıcı Adı-Parola Dosyalarının Analizi: Parola Seçim Davranışı Tespiti

Çevrimiçi hizmetler, kullanıcıların kişisel/özel içeriklere erişmesini sağlamak için genellikle parola tabanlı sistemler kullanır. Tez kapsamında toplanan sızıntı belgelerinin içinde ağırlıklı olarak kişisel verilerin kullanıcı adı ve parola çiftlerinden oluştuğu görülmüştür. Parola bilgileri sızmış olma ihtimaline karşın hizmet sağlayıcılar parola güvenliğini artırmak için belirli politikalar kapsamında kullanıcıların parolalarını periyodik olarak değiştirmeye zorlar. Ancak tez kapsamında sızıntı verilerinin analizi, mevcut politikaların kullanıcı parolası seçme alışkanlıklarını dikkate almadığını ve kritik güvenlik ve gizlilik kaygıları oluşturduğunu görülebilir. Parolalar sızıntılarını kullanarak, saldırganlar kullanıcıların parola seçim kümesinin yapısını/modelini inceleme ve muhtemelen tanımlama fırsatına sahip olabilir. Saldırganlar bir sonraki parolayı tahmin edebilir veya kaba kuvvet uzayını önemli ölçüde azaltabilirler. Tez kapsamında vaka çalışmasında, toplanan veri ihlallerinden tespit edilen on milyondan fazla kullanıcı kimliği (e-posta, kullanıcı adı, takma ad vs.) ve parola çiftinin parola seçim davranışlarını incelemek için istatistiksel yöntemler ve görselleştirme teknikleri kullanır. Oluşturulan veri seti, Kullanıcı Kimliği (userID)-Parola çiftlerinin benzersizliği korunarak anonim olarak paylaşıldı. Parolalardan çıkarılan özelliklerle birlikte diğer araştırmacılarla AuthInfo ismi ile Github üzerinden paylaşılmıştır. Analizler neticesinde

kullanıcı parolası seçim davranışı kalıplarının genelleştirilebileceğini ve saldırıların başarı oranını artırmak için kullanılabileceğini göstermektedir.

AuthInfo veri seti üzerinde gerçekleştirilen istatistiksel çıkarım ve veri görselleştirme çalışmaları, özellikleri ve aralarındaki ilişkiyi ortaya koymaktadır. Kaba kuvvet saldırılarında en kritik parametrelerden biri parolanın uzunluğudur. Parola uzunluğunun tahmin edilmesi, uygulanabilir bir kaba kuvvet saldırısı için çok önemlidir, çünkü parola uzunluğunun gereken süre üzerinde üstel bir etkisi vardır. Parolada kullanılan karakter türleri, parola güvenliğini parola uzunluğu kadar etkileyen diğer bir faktördür. Tekrarlanan bir permütasyon, **Denklem 4.1**'de belirli bir karakter kümesinden oluşan olası parolaların toplam sayısını şu şekilde temsil eder:

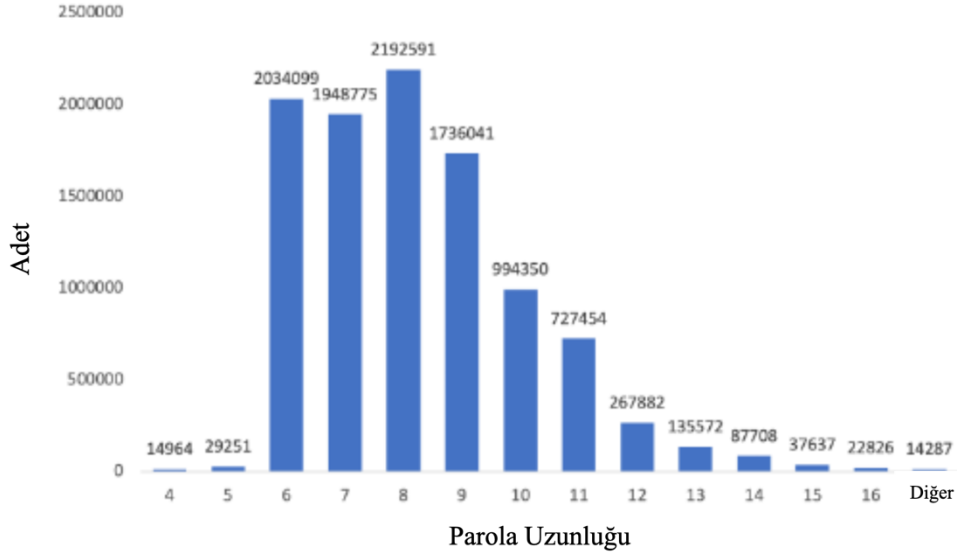
n : karakter seçilecek kümedeki karakter sayısı,

r : parolanın uzunluğu,

$$P(n, r) = n^r \quad (4.1)$$

Denklem 4.1, karakter seti ile parola karakter sayısı uzunluğu arasında üstel bir ilişki gösterir. Örneğin, beş karakterli İngiliz alfabesinden rakamlardan ve 26 küçük harften oluşan bir parola, 60466176 olası parola bulunmaktadır.

AuthInfo veri seti için parola uzunluğu dağılımı **Şekil 4.32**'de göstermektedir. Sonuçlar, 11 karakterden uzun parolaların nispeten düşük olduğunu ve kullanıcıların genellikle altı ila sekiz karakter uzunluğunda bir parola seçme eğiliminde olduğunu gösteriyor. Bu davranış, daha kısa parolaların hatırlanmasının daha kolay olmasıyla açıklanabilir. Ayrıca, çoğu parola politikası için minimum karakter uzunluğu olarak genellikle altı tercih edilmektedir.



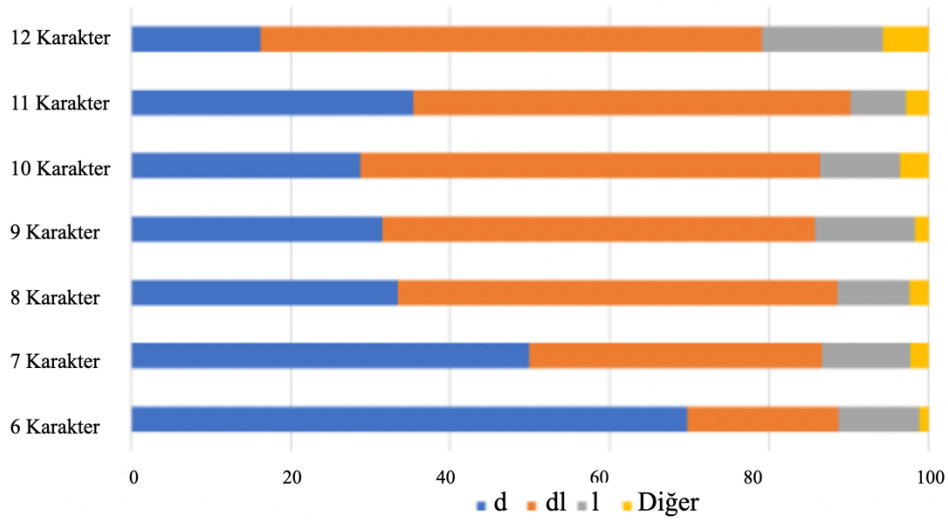
Şekil 4.39: AuthInfo veri seti, farklı parola uzunluklarının dağılımı

Parolaların uzunluk ve karakter türlerine göre dağılımı **Tablo 4.14**'te verilmiştir. Veri seti, altı karakterli bir parola seçen çoğu kullanıcının bunu bir sayı ile yaptığını göstermektedir. Bununla birlikte, parolalarında sekiz karakter kullanmayı tercih eden kullanıcılar, sayıları ve küçük harfleri daha önemli oranda tercih ettiği görülmektedir. Bu durumda, hizmet sağlayıcıların sekiz karakter parola uzunluğu gerektiren parola politikalarının kullanıcıları daha güvenli bir parola seçmeye zorladığı sebebiyle böyle bir yoğunlaşmanın olduğu yorumu yapılabilir. Ayrıca, daha uzun bir parola seçen kullanıcılar, daha farklı karakter setleri kullanmışlardır. Parola seçme uzunluğu arttıkça güvenli bir parolanın seçilme ihtimalini arttığı ve daha bilinçli parola seçimi yapıldığını göstermektedir.

Tablo 4.14: Parolaların uzunluk ve karakter türlerine göre dağılımı

Karakter Uzunluğu	Rakam (d)	Rakam ve Küçük Harf (dl)	Küçük Harf (l)	Diğer
6 Karakter	1.418.316	384.005	207.169	23.597
7 Karakter	972.478	714.900	216.970	43.950
8 Karakter	732.477	1.208.874	198.165	51.933
9 Karakter	538.461	928.638	214.342	29.254
10 Karakter	284.824	573.458	99.160	35.653
11 Karakter	257.041	398.385	50.727	20.447
12 Karakter	43.212	168.040	40.559	15.386

Şekil 4.33'teki değerler, kullanıcı davranışlarını daha anlaşılır kılmak için satır bazında yüzde normalizasyon edilerek gösterilmiştir. Görüldüğü üzere kullanıcıların seçtikleri parola uzunluğu genellikle seçilen karakter tipinin karmaşıklığı ile doğru orantılıdır. Parola uzunluğu arttıkça ikiden fazla karakter türü kullanma olasılığının arttığı gözlemlenmiştir. Veri seti parola kombinasyonları arasında, karma karakter seti (d, dl ve l hariç herhangi bir karakter seti kombinasyonu) kullanımının genellikle parola uzunluğu ile arttığı görülmektedir. Tek istisna, 9 ve 11 karakterli parolalardır. Bu davranış, doğum günleri, telefon numaraları veya resmi kimlik numaraları gibi kalıplara sahip sayıların parola olarak kullanılmasıyla açıklanabilir. 7 karakter uzunluğundan sonra, kullanıcıların yarısından fazlasının dl karakter kombinasyonunu tercih ettiği görülmektedir.



Şekil 4.40: Karakter uzunlukları içindeki karakter türlerinin yüzde dağılımı

AuthInfo veri seti incelendiğinde en sık kullanılan karakterler setleri şu sırayla sıralanmıştır: d, dl, l, ud, udl, sdl, sl, sd, u, ul. d, dl, l ve ud'den sonra kalan sütunlar, alandan tasarruf etmek için "diğer" sütununda birleştirilir. Parolaların çoğu tamamen rakamlardan oluştuğu görülmüştür. Bu dağılım, sızıntı kaynakların parola politikalarının genellikle yeterince katı olmadığını göstermektedir.

Farklı parola uzunlukları için hangi karakter setlerinin kullanıldığını keşfetmek için parola karakter türü dağılımı da hesaplanmıştır. **Tablo 4.15**, en yaygın kullanılan parola uzunluklarını ve bunların tüm kaynaklar üzerindeki dağılımını gösterir. Ayrıca tüm parola uzunlukları için karakter tipi dağılımları **Tablo 4.15**'nin son satırında verilmiştir.

Kullanıcıların mümkün olduğu kadar az sayıda özel karakter kullanmak istedikleri ve genellikle sadece rakam kullandıkları görülmektedir. İlginç bir şekilde, parola uzunluğu arttıkça kullanıcılar özel karakterleri, büyük ve küçük harfleri ve sayıları kullanma eğiliminde artış görülmektedir.

Tablo 4.15: Parola içeriği karakter türü dağılımı

Karakter Uzunluğu	Küçük Harf	Büyük Harf	Özel Karakterler	Rakam
6 Karakter	19,14%	0,45%	0,12%	80,29%
7 Karakter	24,44%	0,57%	0,13%	74,86%
8 Karakter	31,86%	0,52%	0,17%	67,45%
9 Karakter	35,77%	0,67%	0,17%	63,39%
10 Karakter	38%	0,71%	0,23%	61%
Tüm Parola Uzunlukları	32,25%	0,67%	0,20%	66,88%

Bir parolanın olası ilk karakterini belirlemek, kaba kuvvet saldırıları için olası parolaların arama alanını daralttığı için avantajlıdır. Parolanın bir kısmı bilindiğinde, olası kombinasyonların sayısı şu şekilde gösterilebilir:

$$P(n, r) = d^m * n^{r-m} \quad (4.2)$$

Burada n kümedeki karakter sayısıdır, r parola uzunluğunu belirtir ve d bilinen karakter sayısını gösterir. Böylece, parolanın ilk m harflerinin bilindiği varsayılarak, ortaya çıkan olası küme **Denklem 4.2**'de gösterilmektedir.

İngiliz alfabesindeki büyük harf, küçük harf ve ondalık sayıların birleştirilmesi (62 farklı karakter), parolanın ilk harfinin büyük harfle başladığı bilgisi, **Denklem 4.2**'ye göre saldırıyı 36 kat hızlandırır. Çoğu insan cümleye büyük harfle başlama eğiliminde olduğundan, büyük harfli bir parola seçmenin de yaygın olması beklenir. Kullanıcıların parolalarına nasıl başladıklarını anlamak için ilk karakter dağılım yüzdeleri hesaplanmıştır. **Tablo 4.16**'da uzunluk arttıkça rakam kullanımı azalsa da küçük harf kullanımı artmaktadır. Özel karakter ve büyük harf kullanımı diğer sınıflara göre daha az görülmektedir. 11 karakter uzunluğundaki parolalar için daha önce bahsedilen trend dışı davranışın ilk karakter seçiminde de devam ettiği görülmektedir. Bu nedenle 11 haneli parolalar için ayrı bir değerlendirme oluşturulması

saldırının etkinliğini artıracaktır. 12 karakter ve üzeri parola tercih eden kullanıcılar için ilk karakterdeki rakam kullanımının %10'un altına düştüğü gözlemlenmiştir.

Tablo 4.16: Parola ilk karakter dağılımı

Karakter Uzunluğu	Ondalık (D)	Küçük Harf (L)	Özel Karakterler (S)	Büyük Harf (U)
6 Karakter	72,98%	26,26%	0,05%	0,71%
7 Karakter	55,15%	43,57%	0,03%	1,25%
8 Karakter	46,48%	52,37%	0,05%	1,10%
9 Karakter	40,99%	57,12%	0,04%	1,85%
10 Karakter	38,95%	59,48%	0,05%	1,52%
11 Karakter	60,03%	38,96%	0,03%	0,98%
12 Karakter	24,11%	73,91%	0,06%	1,92%
13 Karakter	9,66%	88,41%	0,05%	1,87%
14 Karakter	9,49%	88,69%	0,07%	1,75%
15 Karakter	11,97%	85,13%	0,06%	2,83%

Yine de ilk karakter trendinde rakam kullanımı 15 karakterden sonra tekrar artışa geçmektedir. Parolanın ilk karakterinde küçük harf kullanma eğiliminin genellikle 15 karaktere kadar arttığı, yalnızca 11 karakterden oluşan parolalarda ise %38'e düştüğü görülmektedir. Ayrıca, parolalarında büyük harf kullanan kullanıcıların %90'dan fazlası ilk karakteri tercih etmektedir. Parolanın son karakteri değerlendirildiğinde, özel karakter kullanan kullanıcıların %30'u son karakter olarak özel karakter kümesinden bir eleman kullandığı görülmüştür.

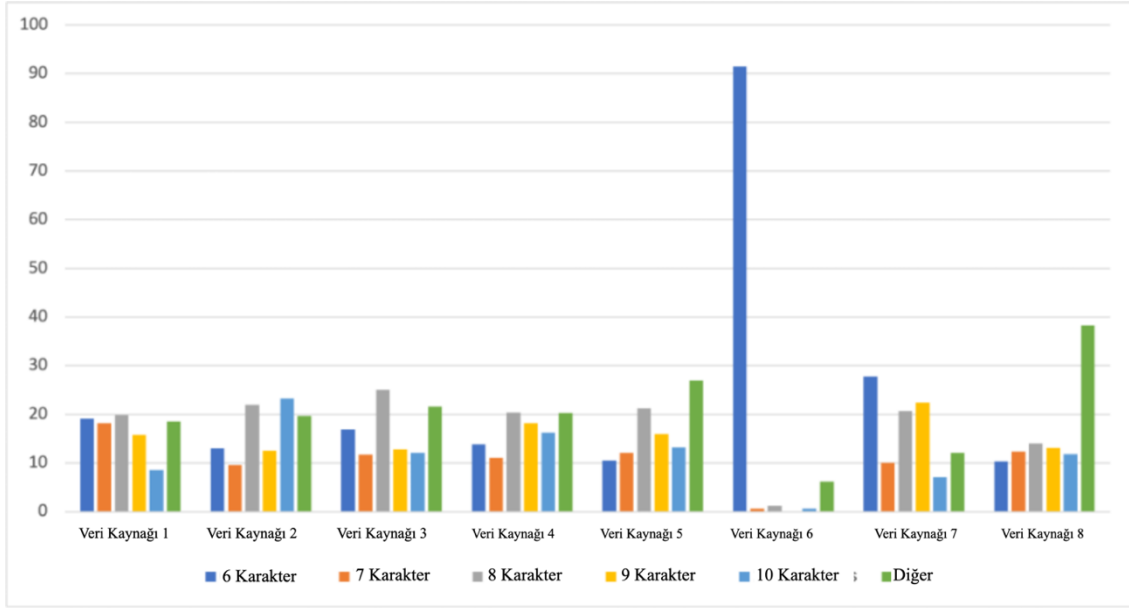
Parolalardaki karakter sayısının dağılımı, parola güvenliğindeki bir diğer kritik parametredir. Yaygın olarak kullanılan bir karakter seti ve uzunluk bilgisi, kaba kuvvet saldırısı deneme alanını önemli ölçüde azaltabilir. **Tablo 4.17**, bağımsız olarak değerlendirilebilen parola karakter türü maskelerinin bir dağıtımını içerir.

Tablo 4.17: Parola maskesinin dağıtımı

Maske	Sayı	Uzunluk
dddddd	1418316	6
dddddd	972478	7
ddddddd	732477	8

ddddddd	538461	9
ddddddddd	398385	11
ldddddd	324169	8
ddddddddd	284825	10
lllllll	216930	7
lllllllll	214308	9
lllllll	207165	6

AuthInfo veri setinde yer alan kaynakların uzunluk dağılımlarını **Şekil 4.41** göstermektedir. Veri kaynaklarının dağılımları genel olarak dengeli olmakla birlikte, veri kaynağı altındaki altı karakterlik uzunlukların dağılımı oldukça yoğundur. Bu anormallik, ilgili kaynak için parola politikası olarak yorumlanabilir.

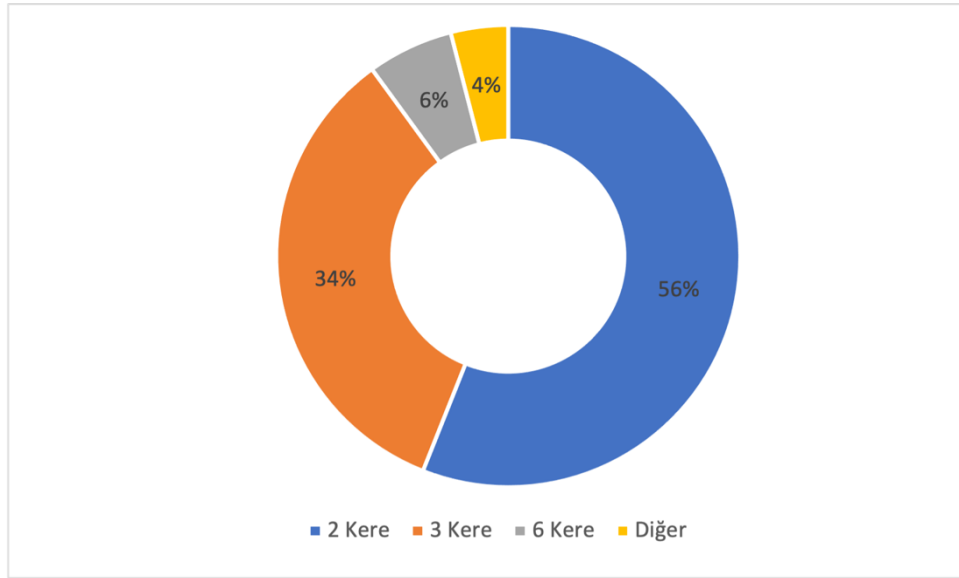


Şekil 4.41: Sızıntı kaynaklarının parola uzunluk dağılımı yüzdeleri

4.4. Vaka Çalışması: Parola Seçim Davranışı Kullanarak Kaba Kuvvet Saldırısı

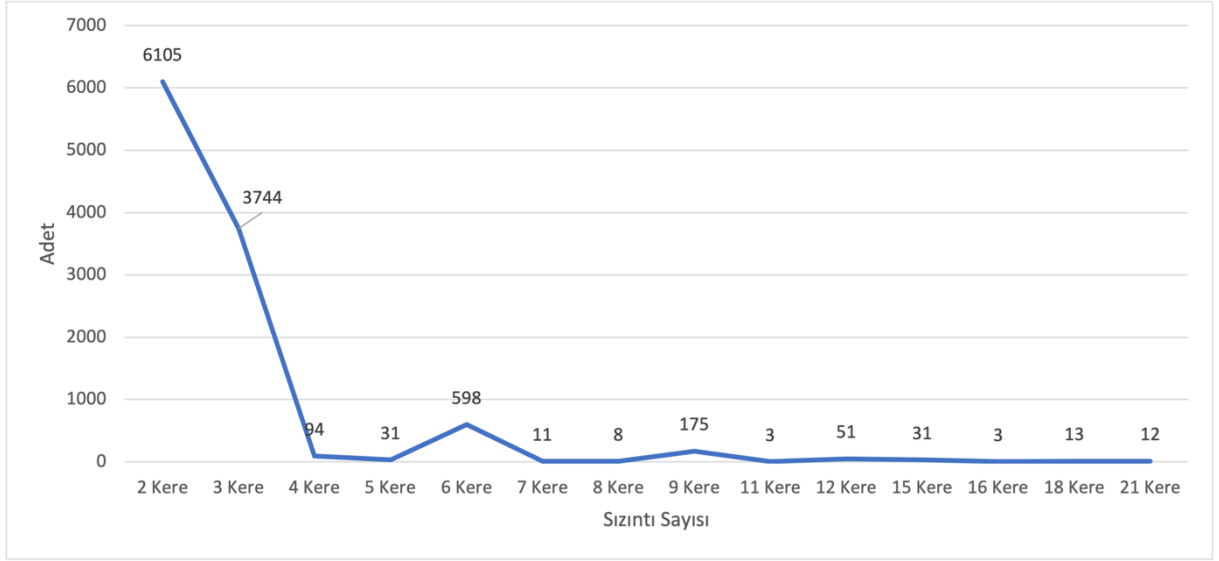
Sızan parolalar, parola seçimi ve genel olarak parola politikası hakkında bilgi sağlar. Sızmış parolalar üzerinden uyarlanmış bir yaklaşım, kullanıcılar için kişiye özgü saldırılara yol açabilir. Birden fazla parola sızıntısı olan kişilerin, saldırganlar tarafından alışkanlıklarını kolayca çıkarabildikleri için özellikle hassastır. NIST raporuna göre [18], insanlar ortalama altı farklı çevrimiçi hizmet kullanmaktadır. Kullanıcın kullandığı altı farklı çevrim içi servisin

birden çok parolanın sızdırılma olasılığını önemli ölçüde artırmaktadır. AuthInfo, birden fazla sızıntıyla karşılaşan kullanıcıların %0,1'e yakın olduğunu göstermektedir. **Şekil 4.42** ve **Şekil 4.43** aynı kullanıcı için veri setindeki çeşitli parola sızıntılarını göstermektedir. Birden fazla parolası sızdırılan kullanıcının büyük çoğunluğunda iki veya üç sızıntı vardır. Bu bilgilerle adım adım yapılan bir kaba kuvvet saldırı deneyi, seçilen bir kullanıcı kimliği için parola seçme davranışını öğrenme durumunun vaka çalışması üzerinde mümkün olduğunu göstermeyi amaçlamaktadır.



Şekil 4.42: Çoklu sızıntıların dağılımı

Bu deney, sızdırılan parolaların salt değerleri olmadan MD5 algoritmasıyla özet haline getirildiğini varsayılmaktadır. Saldırının stratejisi şu şekildedir: başlangıçta, kaba kuvvet saldırısı körü körüne (önceden bilgi olmadan) gerçekleştirilir. Ardından, her başarılı saldırıdan sonra, çözümlenen parola, kullanıcı için bilinen parolalar veri tabanına eklenir. Ayrıca, çözülmüş parolalar, kullanıcı hakkında daha fazla bilgi ile saldırıyı tekrarlamak için kullanılır.



Şekil 4.43: Birden fazla sızıntıya maruz kalan kullanıcıların dağılımı

MD5 kırma süreci, kaba kuvvet saldırısının performansını değerlendirmek için kullanılır. Bilgi işlem ortamı, 250 GB RAM, 2,40 GHz Intel (R) Xeon (R) Gold 6148 CPU ve 14528/15205 MB ve 40MCU ile 4 Tesla T4 GPU donanımına sahip bir Yüksek Performanslı Bilgisayar (HPC) ortamıdır. Ayrıca MD5 çıktılarını kırmak için Hashcat uygulaması kullanılmaktadır.

Parolası on kez sızdırılan bir kullanıcı rastgele seçilir ve kullanıcıların parolaları **Tablo 4.18**'de verilmiştir. Parolalar, tuz değeri olmadan MD5 algoritması ile özetlenir. Ayrıca saldırgan, kurban hakkında bilinen önceki parolalar gibi herhangi bir ön bilgi bilmiyordu. Bir kaba kuvvet saldırısı için ilk parametre genellikle parola uzunluğudur. Yazdırılabilir tüm ASCII değerleri kullanıldığında Hashcat uygulamasının sınırı dokuz karakterdir. Bu nedenle, boyut dokuz karakter olarak ayarlanmıştır. Daha sonra saldırgan, önceden bilgi kullanmadan ilgili kullanıcının parolasının karmasını tahmin etmeye çalışır. Hashcat, küçük ve büyük sayılar ve özel karakterler kullanarak kaba kuvvet saldırısını tamamlamak için 287 gün ve 13 saat hesaplar. Hashcat, zaman kısıtlamaları nedeniyle yaklaşık bir hafta sonra sonlandırıldı. Saldırgan, 4 ila 8 karakter arasındaki tüm karakter setlerini dener ve kaba kuvvet saldırısı sırasında başarısız olur. Mağdur kullanıcının parolasının en az dokuz karakter uzunluğunda olduğu gözlemlenmiştir. Yürütülen komut aşağıda verilmiştir.

Tablo 4.18: dce9fef8b5c3ba9d576ee9ad37259632 AnonID'ye sahip bir kullanıcının sızmış paroları

Parolla	U z u n l u k	Kara kter Tipi	İl k K ar ak te r	L S a y ı s ı	U S a y ı s ı	D S a y ı s ı	S S a y ı s ı	Maske
cambridgl	9	dl	1	8	0	1	0	ld
cambridgedw	11	l	1	11	0	0	0	
cambridgedl	11	dl	1	10	0	1	0	ld
cambridged0w	12	dl	1	11	0	1	0	ldl
cambridged0wq	13	dl	1	12	0	1	0	ldll
cambridgeasdaq	14	l	1	14	0	0	0	
cambridge12177	14	dl	1	9	0	5	0	ldddd
cambridge123qqq	15	dl	1	12	0	3	0	lddddlll
cambridgeazasdld	16	dl	1	15	0	1	0	ldldl
cambridge212311q321qwe	22	dl	1	13	0	9	0	ldddddldd ddlll

Zorluk: ab0ccdfc2bf732ec09b94ec8221e0c39 ve diğer 10 MD5 özet

hashcat -a 3 -m 0 -i --increment-min=4 hash.ebu ?a?a?a?a?a?a?a?a

Çıktı: Yok, Operatör tarafından duruldu.

Tahmini Süre: 287 gün, 13 saat

Bir sonraki deneme, bu tez çalışmanın sonuçları kullanılarak saldırı alanını daraltmak için Authinfo veri seti ile yapılmıştır. Saldırgan, AuthInfo veri kümesinin %80'e yakınının sızdırılmış parolalarının 4 ila 10 karakter uzunluğunda olduğunu bilgisini kullanır. Ek olarak, bu parolaların %85'inden fazlası sayısal ve küçük harfli karakter türleridir. Saldırgan bu bilgileri kullanarak sekiz karakterlik parolaları kolayca kırabilir. Bu işlem için komut ve çıktı aşağıda verilmiştir.

Zorluk: ab0ccdfc2bf732ec09b94ec8221e0c39 ve diğer 10 MD5 Özet

hashcat -a 3 -m 0 -i --increment-min=4 -l ?d?l hash.ebu ?1?1?1?1?1?1?1?1?1?1

Çıktı: cambridgl

Tahmini Süre: 7 dakika, 37 saniye

Saldırgan, küçük harf ve sayısal karakter kümeleri için on karakteri tükettikten sonra, diğer parola uzunluklarının ondan uzun olduğunu belirler. Saldırgan, AuthInfo analizini 11 karakter için incelendiğinde, kullanıcıların parola seçme alışkanlıklarının %52 ondalık ve küçük harf, %36 ondalık ve %9 küçük harften oluşan karakter kümeleri içerdiğini keşfetti. Saldırgan başlangıçta 11 karakterlik ondalık veya küçük harfli karakter kümelerine kaba kuvvet saldırısı girişiminde bulundu; bu, ondalık ve küçük harfli karakter kümeleri için yaklaşık iki ay sürmüştür. Son olarak, saldırı 11 karakter uzunluğundaki ondalık karakter kümesinde başarısız oldu. Saldırgan, küçük harfli bir kaba kuvvet saldırısı kullanarak "cambridgedw" özetini 11 karakter için kırmayı başarmıştır. Bu işlem için komut ve çıktı aşağıda verilmiştir.

Zorluk: 4137b9bce6c0e826abf062cb677a1498 ve diğer 9 MD5 özeti

hashcat -m 0 -a 3 -1 ?d?l hash.ebu ?1?1?1?1?1?1?1?1?1?1

Çıktı: Yok, Operatör İptal Edildi

Tahmini süre: 51 gün, 2 saat

Zorluk: 4137b9bce6c0e826abf062cb677a1498 ve diğer 9 MD5 özeti

hashcat -m 0 -a 3 hash.ebu ?d?d?d?d?d?d?d?d?d?d

Çıktı: Yok

Tahmini süre: 11 saniye

Zorluk: 4137b9bce6c0e826abf062cb677a1498 ve diğer 9 MD5 özeti

hashcat -m 0 -a 3 hash.ebu ?l?l?l?l?l?l?l?l?l?l

Çıktı: cambridgedw

Tahmini süre: 36 dakika 24 saniye


```
hashcat -a 3 -m 0 -i --increment-min=4 -l ?d?l?s hash.ebu ?1?1?1?1?1?1?1?1?1?1
```

Çıktı: Yok

```
hashcat -a 3 -m 0 -i --increment-min=9 hash.ebu ?a?a?a?a?a?a?a?a?a?a
```

Çıktı: Yok

```
hashcat -a 3 -m 0 -i --increment-min=9 hash.ebu cambridge?a?a?a?a?a?a?a?a?a?a
```

Saldırgan, kurbanın 22 karakterlik en uzun parolası için dokuz karakterlik bir ön ek sunsa da hashcat programının kaba kuvvet sınırlarını aştığı için başarılı olamaz. Ancak saldırı en yakın 21 haneli maske ile başladığında hesaplanan süre on yıldan fazladır (Next Bigbang-Bir sonraki Bigbang). Çalıştırılan komutlar ve çıktıları aşağıda verilmiştir.

Zorluk: eeb266e017bab6512e843f96dfb1c6d1

```
hashcat -m 0 -a 3 -l ?d?l hash.ebu cambridge?1?1?1?1?1?1?1?1?1?1?1?1?1
```

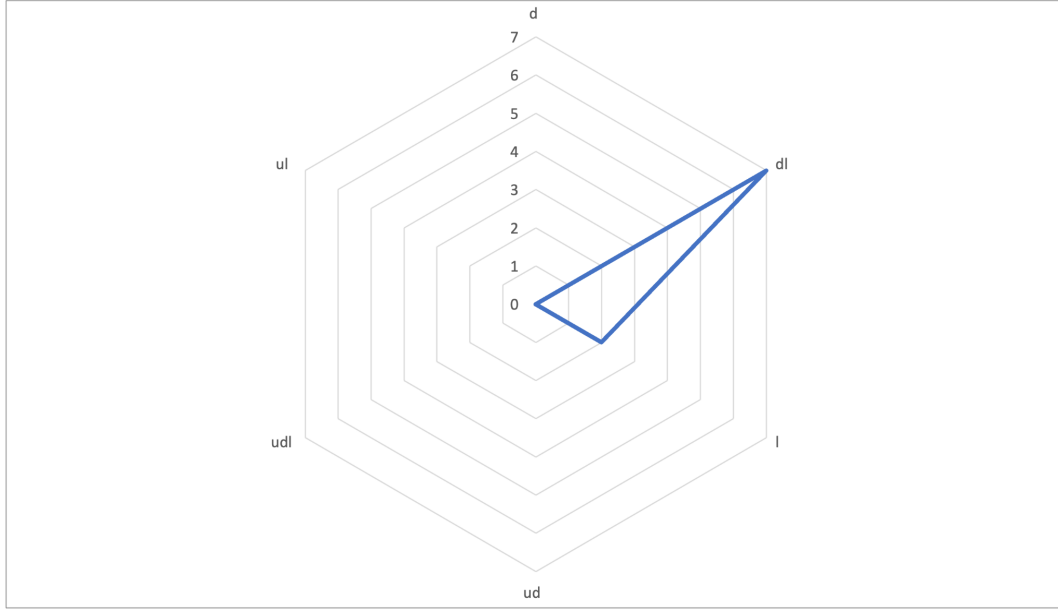
çıktı: yok

```
hashcat -m 0 -a 3 -l ?d?l hash.ebu cambridge?1?1?1?1?1?1?1?1?1?1?1?1?1
```

çıktı: yok

tahmini süre: on yıldan fazla (Next Bigbang)

Kurban kullanıcının radar diyagramı üzerinde **Şekil 4.44** üzerinde görüldüğü gibi ağırlıklı olarak rakam ve küçük harften (dl) oluşan parola seçimi görülmektedir. Kullanıcın 2 adet sadece küçük karakterlerden (l) oluşan 2 parola seçimi yapmasına karşın %90 a yakın kullanıcının rakam küçük harf parola seçme eğiliminde olduğu çıkarılabilir.



Şekil 4.44: Kurban kullanıcının dokuz çözülmüş parolalarının karakter tipi dağılımı

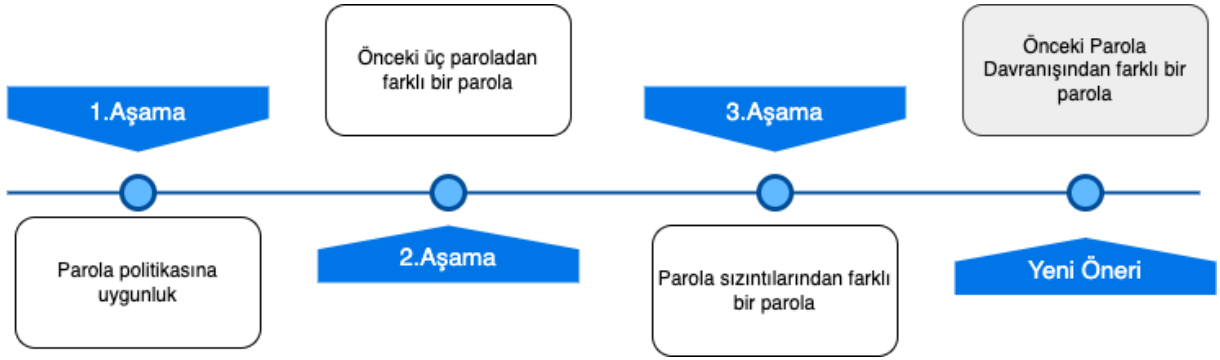
Kullanıcın AuthInfo istatistikleri kullanılarak kırılan parolaları gibi bundan sonra seçeceği parolaları da kırılan/sızan parolaları üzerinden tahmin etmek mümkündür. Ayrıca elde edilen analiz bilgileri kullanılarak, veri setinde belirtilen özellikler kullanılarak kullanıcının bir sonraki parolası tahmin edilerek sözlük oluşturulabilir. Kullanıcı özelinde seçmiş olduğu parolalar benzeyen aynı davranışta farklı kombinasyonlarda parola listeleri de oluşturulabilir. Kullanıcının mevcut havuzdan parola seçmesi durumunda uyarı verilmesi mümkündür. Ancak kullanıcılar 9 ile 22 karakter arasında değişen parola uzunluklarını tercih ettikleri için sözlük oluşturmak işlem gücü ve disk alanı açısından oldukça kaynak ve elektrik tüketimi açısından maliyetli olacaktır.

4.5. Parola Davranışı Duyarlı Parola Politikası yöntemi

Tez çalışması kapsamında incelenen parolalar ve kullanıcı parola seçim davranışı analizi neticesinde, sistemlerin mevcut parola güvenlik seviyesini artırmak ve bir ihlal olması durumunda daha fazla suistimalden kaçınmak için davranışa dayalı yeni bir parola politikası önermektedir. Parola politikaları, kullanıcıların yeterince güvenli bir parola seçmelerini sağlamak için geliştirilmiştir. Bu parola politikaları için karakter uzunluğu ve farklı karakter setlerinin kullanımı veya klavyedeki tuşların mesafesi gibi çeşitli metrikler geliştirilmiştir. Ek olarak, bazı politikalar eski parolaları (muhtemelen özetlenmiş versiyonlar) kaydeder ve kullanıcıları önceki parolalarından farklı bir parola seçmeye zorlar. Ayrıca ek bir önlem olarak

bazı internet servis sağlayıcıları (Github ve Google gibi) mevcut güvenlik politikalarına ek olarak kullanıcı hesapları için parola sızıntılarını kontrol etmektedir. Bununla birlikte, bu çalışmanın vaka çalışması bölümünde, saldırganların parolaların birden çok kez açığa çıkması durumunda parola seçme davranışlarını öğrenebilecekleri gösterilmiştir. Kullanıcının parolası öncekinden farklı olsa bile, kullanıcı davranışının kaba kuvvet saldırısı ile birleştirilmesi gereken süreyi önemli ölçüde kısaltır. Bu nedenle, davranışın parola ile birlikte değiştirilmesi önerilmiştir.

Şekil 4.45'de gösterilen "Parola Politikasıyla Uyumlu", "Önceki Üç Paroladan Farklı" ve "İlgili Parola Sızıntılarından Farklı" adımlarına ek olarak, "Önceki Parola Davranışından Farklı" adımı kapsamlı bir parola politikası oluşturması önerilmiştir.



Şekil 4.45: Güvenli parola seçimi kontrol adımları

"Önceki Parola Davranışından Farklı" özelliğini etkinleştirmek için iki adım gereklidir. İlk adım davranış tespit etmek, ikincisi ise onu güvenli bir şekilde saklamaktır. Parola daha sonra analiz edilemeyeceğinden, parola seçim davranışı parola oluşturma sırasında değerlendirilmelidir. Ancak parola davranışını veri tabanında ham tutmak olası bir sızıntı durumunda da güvenlik riski oluşturacaktır. Bu nedenle, kullanıcıların gizliliğini sağlamak ve önceki davranışlarını izlemek için parola davranışı diğer bilgilerle karıştırılmalıdır. **Denklem 4.3**'de gösterildiği gibi bu özelliklere bir tuz değeri eklenerek farklı hizmetler arasında çıktının benzersizliğini sağlamak da mümkündür.

$$Parola\ Davranış = hash(userID + Maske + \epsilon \times salt) \quad (4.3)$$

Parola davranış bilgisi oluşturma bir yolu, paroladaki karakterleri karakter türleriyle değiştirerek maske verilerini kullanmaktadır. Ancak, maske bilgisi tek başına kullanıcı

parolalarını ortaya çıkarabileceğinden, bu veriler ek kullanıcı bilgileri ile birleştirilmelidir. Alternatif olarak, bu çalışmada kullanılan userID'nin özetlenmiş bir versiyonu olan AnonID de benzer davranışlar sergileyebilir. Son olarak, isteğe bağlı benzersiz salt değeri, aynı kullanıcı verilerinin farklı hizmetlerde çeşitliliğini sağlar.

"Önceki Parola Davranışından Farklı" adımı, diğer aşamalarda zaten kullanılanların ötesinde ek bilgi gerektirmez. Küçük değişikliklerle, mevcut parola politikaları önerilen özelliğe hızla uyum sağlayabilir. Ek olarak, **Denklem 4.3**'deki gibi bir teknik, sızıntılar olsa bile bu bilgilerin gizliliğini sağlayabilir.

5. TARTIŞMA

Veri sızıntısı ve veri ihlali, geleneksel belgelerin saklanmasıyla başlayan ve dijital sistemlerin yaygınlaşmasıyla zirveye ulaşan, içeriden ya da dışarıdan kötü niyetli saldırgan davranışlarıyla veya veri sorumlusu tarafından hatalı prosedür ve yönetimi sonucunda ortaya çıkabilmektedir. Ortaya çıkış şekilleri ve yöntemleri farklı olsa da yayılması ve yaygınlaştırılması için uygulanan yöntemler benzerdir. Yasal olmayan internet forumları, TOR ağı, anonim dosya paylaşım platformları, anonim metin paylaşım platformları üzerinden genellikle paylaşılmaktadır. Arama motorları da tüm taranabilir web sitelerini kapsamak için taramalar gerçekleştirirken veri sızıntı ve veri ihlali için kullanılan kimliği belirsiz kişiler tarafından paylaşılan verileri de indexlemektedirler. Veri ihlallerine ulaşmanın bir diğer yolu da veri sızıntısı üzerine çalışan araştırmacıların kurduğu haveibeenpwn.com ve intelx.io gibi platformlar da ücretli/ücretsiz API'lar ile veri ihlallerini takip etmek için kaynak oluşturmaktadırlar. Ayrıca çeşitli sosyal medya ve mesajlaşma uygulaması grupları üzerinden de veri ihlallerinin doğrudan veya dolaylı olarak (kaynak göstererek) paylaşıldığı görülmektedir.

Siber Tehdit İstihbaratı için Yeni Model Geliştirilmesi başlıklı bu tez kapsamında taranabilir web kaynakları üzerinde veri sızıntıları ve ihlallerinin bulunup elde edilmesi, yapısal hale getirilmesi, anlamlandırılması, sızıntı türü kişisel, kurumsal, potansiyel bilgi ve anonim olarak sınıflandırılması için bir model geliştirilmiştir sınıflandırılma yapılırken ilk olarak “Anonim” ve “Özel Nitelikli” (bilgi içerir) şeklinde ikili sınıflandırma yapılmıştır. Daha sonra özel nitelikli bilgilerin de kendi içinde “Kişisel Veri”, “Kurumsal Veri” ve “Potansiyel Bilgi” olarak ayrılmıştır. Geliştirilen modüllerin çıktıları veri setlerine dönüştürülerek makine öğrenmesi ve yapay zeka modellerinin eğitiminde kullanılmıştır. Oluşturulan veri setleri ve farklı araştırmacıların kullanabilmeleri için anonimleştirilerek Github üzerinde paylaşılmıştır.

Doğrudan sızıntı verilerinin üzerinde çalışan ve topladıkları verileri analiz ve sınıflandırma yapan az da olsa araştırmacılar da bulunmaktadır [9, 84]. Veri sızıntı ve ihlalleri üzerine araştırmacıların genellikle forum gönderileri ve sosyal medya gönderileri ile ağ üzerinde veri sızıntısı tespiti (DLD) ve veri sızıntısı önleme (DLP) üzerinden yapılan analizlerin üzerinde çalışmaların öne çıktığı görülmektedir [10-13, 85]. Veri sızıntıları ve saldırganlar, maddi ve

manevi zarar vermek için sahte veriler de yayınlatabilir, Maschler ve arkadaşları sızıntı doğrulama yazılımı geliştirdi [41]. Tez çalışması verilerin bilgilendirme metninden ziyade veri içeriğinin analizine odaklanılmıştır. Araştırmacıların bir metnin veri ihlalinin bahsetme durumuna odaklanırken tez çalışmasında paylaşılan sızıntı verinin içeriğinin hangi veri ihlali içerme durumu ve hangi tür veri içerdiği (kişisel veya kurumsal gibi) analiz edilmiştir. Siber Tehdit İstihbaratı için Yeni Tarama Modeli değişen regülasyonlar ve kurum ihtiyaçlarına göre esnek olarak veri özelliklerinden bağımsız geliştirilmiştir.

Bu tez kapsamında ise açık kaynaklar ve çeşitli API'lar üzerinde veri sızıntıları ve ihlalleri üzerinde analiz ve sınıflandırma işlemleri sonrasında toplanan kullanıcı bilgilerinin üzerinde parola seçim davranışı üzerine analizi yapılmıştır. Ortalama olarak, bir kullanıcının birden çok internet hizmetinde altı farklı hesabı bulunmaktadır [38]. Kullanıcılar tüm bu hizmetler için farklı bir parola seçseler bile yeterli parola sızıntısı ile saldırganlar bireysel olarak davranışlarını öğrenebilirler. Kullanıcıların genellikle parolalarını oluştururken bir kalıba bağlı kaldıkları görülmektedir. Parola seçim modelinin benzer olduğu varsayılarak, bir saldırı stratejisi tasarlanır. Yapılan parola seçim davranışı analizinin başarısının gösterilmesi için bir vaka çalışmasında kaba kuvvet saldırısının iyileştirilmesinde kullanılmıştır. AuthInfo veri kümesi, bu durumu göstermek için birden çok sızıntıya sahip kullanıcıları içerir. Aşamalı olarak artırılan bir saldırı uzayında yapılan saldırıyı göstermek için bir vaka çalışması olarak, yedi farklı parola sızıntısı olan rastgele seçilmiş bir kullanıcı kullanılır. Bir saldırganın, saldırı uzayı dışında görünen uzunlukta bir parolaya AuthInfo analizleri ve maskeleri kullanılarak bir kaba kuvvet saldırısını yalnızca birkaç saat içinde bitirebileceği kanıtlanmıştır. Bu, saldırganların yeterli veriyle kullanıcıların bir sonraki parolasını nispeten kolay bir şekilde tahmin edildiğini gösterir.

Kullanıcı bilgileri analiz çalışmalarında, araştırmacılar parolaları ve veri ihlallerini incelediler. Veri sızıntılar, yenilikçi kullanıcı davranışı analizi sağlamıştır. breachesalarm.com ve haveibeenpwned.com gibi platformlarda 11 milyar kullanıcı bilgisi kaydı bulunmaktadır. Jayakrishnan ve diğerlerine göre zorunlu parola güncellemeleri, parola oluşumunu yavaşlattı ve güvenliği iyileştirmedi. [42]. Meslek, eğitim, cinsiyet ve milliyet, küçük örneklem büyüklüğü ile sınırlıdır. Murphy'nin tezi, kullanıcı tarafından oluşturulan parola kriptografisi ve zorluğu üzerinedir. Ortalama Levenshtein mesafesi bilinen parolalardan hesaplanmıştır [44]. Honeyword gibi çalışmalar ile parola istilalarını algılama çalışmaları yapılmıştır [35-36]. Bir izinsiz giriş tespit çalışmasında sık kullanılan parolalardan (honeyword) denemeler

yapıldığında uyarı vericek şekilde çalışmaktadır. Bu çalışmalar AuthInfo veri setindeki parolalar ya da kullanıcının daha önce sızmış parolaları da honeyword olarak genişletilebilir. Sızan kullanıcı verileri, kültür ve bölgenin parola seçimini etkilediğini gösteriyor [37-39,41]. Bu araştırmalar, büyük veri kümelerinden yararlı çözümler sağlar, ancak genelleştirilmiş özellik kümelerinden yararlanmaz. Tez 4.3, 4.4, 4.5 içi bölümlerinde parola güvenliği incelemiştir. Literatürden farklı olarak kullanıcı adı parola çiftleri parola üzerinden çıkarılan özellikler ile analiz edilmiş ve yeni bir parola politikası aşaması önerilmiştir.

Parola politikaları, kullanıcıları güvenlikleri için güvenli parolalar seçmeye zorlar. Araştırmacılar tarafından kullanılan parola politikaları genellikle kaba kuvvet çalışmalarını büyütme üzerine odaklanmaktadır. Özellikle kelime Bu tür politikaların çoğu, farklı karakter türleri ve nispeten uzun bir parola gerektirir. Ek olarak, bazı hizmetler, kullanıcının güvenliği ihlal edilmiş veya ortak bir parola kullanıp kullanmadığını kontrol eder. Bununla birlikte, nispeten iyi bir parola politikası bile, kullanıcıların her yeni parolayla birlikte parola seçim modellerini değiştirmelerini gerektirmiyorsa, bir veri ihlali meydana geldiğinde kullanıcılara yardımcı olmaz. Bu nedenle, parola politikalarına yeni bir aşama eklenerek kaba kuvvet uzayı dışında kalan uzunluklarda parolaların kırılmasına imkan sağlayan özellikler içeren parola kalıplarının saklanması ve kullanılması ile parola seçim davranışını kontrol eden yeni bir parola güvenlik politikası yöntemi önerilmiştir. Önerilen yöntem, politikaları daha da güçlendirmek için minimum değişikliklerle mevcut yazılımla kolayca entegre edilebilir. Son olarak, kullanıcı kimliği, parola ve ek özelliklerden oluşan AuthInfo veri seti oluşturularak araştırmacıların kullanımı için Github [42] üzerinde herkese açık olarak paylaşılmaktadır. Kullanıcıların gizliliğini ve mahremiyetini korumak için, kişisel olarak tanımlanabilir unsurlar anonimleştirilir. Ayrıca, ek veri sızıntıları keşfedildikçe depo güncellenecek ve genişletilecektir.

6. SONUÇ VE ÖNERİLER

Sonuç olarak Siber Tehdit İstihbaratı için Yeni Model geliştirilmesi tezi kapsamında veri ihlali ve veri sızıntısı olarak değerlendirilen kaynaklar taranmış, toplanan veriler analiz edilerek kural tabanlı ve yapay zeka destekli olarak sınıflandırılan bir sızıntı verisi işleme modeli geliştirilmiştir. Geliştirilen modelin operatör destekli toplanan verilere, arama motorları sonuçlarına ve programlama arayüzleri (API) üzerinden veri toplamaya imkan sağlamaktadır. Kural tabanlı, dosya yapısı ile genellikle REGEX ve NER ile çıkarılan özellikler kullanılarak ayrıştırılan muhtemel veri sızıntıları dosyaları geliştirilen arayüz ile operatör tarafından etiketlenerek Siber Tehdit İstihbaratı 6 farklı Veri Seti oluşturulmuştur. Tez kapsamında oluşturulan veri seti kullanılarak denetimli öğrenme yöntemleri ile modeller eğitilmiştir.

Tez kapsamında geliştirilen model tarafından sızıntı belgelerinin belgelerin yapısal olarak analiz edildiğinde %90'a yakın veri sızıntısı serbest metin olarak sınıflandırılmıştır. Operatör tarafından Etiketlenen veri seti değerlendirildiğinde 3826 “Anonim” belge, 3762 “Özel Nitelik” bilgi içeren belge bulunduğu görülmüştür. Dengeli bir dağılım olması yapay zeka modelleri açısından faydalı olarak görülmüştür. Ancak özel nitelikli bilgilerin içeriği 665 “Kişisel Veri” 1909 “Kurumsal Veri” ve 1188 “Potansiyel Bilgi” ile alt sınıfların daha dengesiz dağıldığı görülmüştür. Model tarafı kural tabanlı sınıflandırma operatör etiketlemesi ile karşılaştırıldığında belgeleri %83 doğruluk oranıyla veri sızıntı belgelerini etiketlediği görülmektedir.

Sınıflandırıcıların performans metrikleri değerlendirildiğinde en iyi performans metriklerine Rastgele Orman algoritmasının gösterdiği sonraki en iyi sonucu K-En yakın komşu algoritması sağlarken En başarısız sonucu Naive Bayes algoritması ulaşmıştır. K-En Yakın Komşu, Naive Bayes, Rastgele Orman ve SVM algoritmalarından çoklu etiket REGEX+NER Veri Seti üzerinde doğruluk değeri en iyi sonuç %84 ile Rastgele Orman Algoritması vermiştir. Çoklu etiket NER Veri Seti üzerinde doğruluk değeri en iyi sonuç %99 ile KNN, SVM ve Rastgele Orman Algoritması vermiştir. Çoklu etiket REGEX Veri Seti üzerinde doğruluk değeri en iyi sonuç %89 ile Rastgele Orman Algoritması vermiştir. NER Veri Seti 4 Etiket veri seti için SVM, KNN ve Rastgele Orman modeli aşırı öğrenme (overfit) ye maruz kalmış bu nedenle sonuçların değerlendirmesi dışında bırakılmıştır. İkinci olarak çoklu

sınıflandırma için “Anonim”, “Kişisel Veri”, “Kurumsal Veri” ve “Potansiyel Bilgi” verilerin ayrıştırılması yapılmıştır. KNN, Rastgele Orman, Naive Bayes ve SVM algoritmaları üzerinde REGEX ve NER özelliklerinden oluşan veri seti kullanılarak eğitilmiştir. KNN algoritması ile veri seti %70 eğitim %30 test olacak şekilde bölünmüştür. REGEX Veri Seti üzerinde doğruluk değeri en iyi sonuç %78 ile Rastgele Orman Algoritması vermiştir. NER Veri Seti üzerinde doğruluk değeri en iyi sonuç %79 ile Rastgele Orman Algoritması vermiştir. REGEX+NER Veri Seti üzerinde doğruluk değeri en iyi sonuç %88 ile Rastgele Orman Algoritması vermiştir. Rastgele Orman algoritması sınıflandırıcılar arasında en iyi performans metriklerine ile ulaşıldığı görülmüştür. K-En yakın komşu algoritması Rastgele Orman algoritmasından sonra en iyi performans metrik değeri alırken En başarısız sonucu Naive Bayes algoritması ulaşmıştır. Doğrudan metin üzerinden NLP yöntemleri ile özellik çıkarımı yapılarak RNN algoritması ile sınıflandırma yapıldığında 2 etiketli ve 4 etiketli veri setlerinde doğruluk %71 olarak kalmıştır. Maalesef toplanan veri sızıntılarının farklı veri kümelerinden ve tiplerinden olması dolayısıyla kural bazı modellerin başarıları kural tabanlı sınıflandırma başarısının altında kaldığı görülmüştür.

Tez kapsamında toplanan kullanıcı bilgisi içeren verilerden kullanıcı parola seçim davranışlarını anlamak için sızdırılan kullanıcı kimlik bilgilerinin istatistiksel ve görsel analizini kullanır ve ardından bu bilgilere dayalı yeni bir parola politikası önerir. AuthInfo veri setini oluşturmak için 10 milyondan fazla kullanıcı kimliği-parola çifti toplanmıştır. Sonrasında kullanıcı bilgileri anonimleştirilerek analiz edilmiştir. Bu sürecin her adımı, kullanıcı gizliliğini ve güvenliğini korurken kişisel verileri paylaşmak için bir yöntem oluşturmak üzere ayrıntılı olarak tanımlanmış ve akış şemaları ile sağlanmıştır.

AuthInfo veri seti, sızdırılan veri kaynaklarının parola politikalarını anlamanın mümkün olduğunu gösteriyor. Ayrıca, "123456" gibi en sık kullanılan parolalar, geçerli parola politikalarının yeterli uygulanmadığını göstermektedir. Ayrıca, AuthInfo'da toplanan parolaların yaklaşık yarısının altı ila sekiz karakter uzunluğunda olduğu ve parola politikalarındaki minimum gerekliliklerin kullanıcı davranışlarını büyük ölçüde etkilediği gösterilmiştir. Son olarak, parola uzunluğu ile çoklu karakter setlerinin kullanımı arasında doğrudan bir ilişki olduğu görülmüş parola uzunluğu arttıkça farklı karakter kümesinden seçimi de artmaktadır. Ek olarak, kullanıcılar, yaygın olarak kullanılan parolalara benzer şekilde, belirli karakter türleri, karakter konumları ve karakter kullanım sıklığı için bir tercih sergilediği görülmüştür. Ayrıca, AuthInfo veri setindeki kullanıcılar çoğunlukla farklı

uzunluklardaki sayıları ve küçük harfleri, nadiren büyük harfleri ve özel karakterleri kullanmışlardır. Ayrıca, büyük harf kullanan kullanıcıların %90'dan fazlası parolalarında ilk karakteri büyük harf olarak kullanmaktadır. Benzer şekilde, parolada özel bir karakter yer aldığına, parolaların %30'unda son karakter olarak bu karakter kullanılmaktadır. Vaka Çalışmasında gösterildiği gibi, bu istatistikler saldırganların daha başarılı kaba kuvvet saldırıları başlamasına yardımcı olabilir. Bu bulgular, parola politikaları yeterince katı olduğunda bile, benzer parola seçme davranışları sergileyen kullanıcıların güvenliğinin azaldığını göstermektedir.

AuthInfo'nun bu çalışmalar için temel ihtiyaçları karşılaması bekleniyor. Ancak teknolojiyle birlikte güvenlik gereksinimleri ve saldırganların stratejileri hızla değişiyor. Bu nedenle, AuthInfo veri seti bu ihtiyaçları karşılamak ve veri setindeki özellik sayısını artırmak için sık sık güncellenecektir. Ayrıca, araştırmacıların Github deposuna katkıda bulunmaları teşvik edilir çünkü daha büyük bir veri seti, yaygın olarak kullanılan örnekler ve sonekler gibi daha fazla örüntünün algılanmasını ve karmaşık örüntülerin keşfedilmesini sağlar.

Kullanıcının birden fazla sızıntısı, parola seçim davranışı alışkanlıklarını ortaya çıkarmaktadır. Kullanıcının parola davranışının belirlenmesi için gereken minimum sızıntı miktarı, kişisel veya kurumsal düzeyde risk analiziyle belirlenebilir. Parola sızıntıları ile kullanıcı davranışlarının açığa çıkması riski arasındaki ilişkiyi değerlendirmek için olasılıksal bir yöntem geliştirmek mümkündür. Parola politikaları, parolaları daha sağlam ve sızıntılara karşı dayanıklı hale getirmek için desen ve entropi analizi gibi ek adımlarla iyileştirilmelidir. Ayrıca, bu iyileştirilmiş parola politikaları, web siteleri için bulut tabanlı bir hizmet veya bireysel kullanıcılar için yerel uygulamalar olarak sunulabilir.

Son olarak, son kullanıcılar, mahremiyetlerini ve güvenliklerini korumak için daha karmaşık parolalar kullanmanın önemi konusunda uyarılmalı ve eğitilmelidir. Gizliliğini umursamayan bir kullanıcıyı hiçbir politika koruyamaz. Bu nedenle, eğitimler, halkı eğitmeye yardımcı olmak için temel güvenlik ve gizlilik tekniklerini içermelidir.

KAYNAKLAR

- [1] Schlackl F., Link N., and Hoehle H. "Antecedents and consequences of data breaches: A systematic review." *Information & Management* 2022.
- [2] Khan, Freeha, et al. "Data breach management: An integrated risk model." *Information & Management* 58.1 2021.
- [3] McElroy, Sean, and Lisa McKee. "Developing Privacy Incident Responses to Combat Information Warfare." *International Conference on Cyber Warfare and Security*. Vol. 18. No. 1. 2023.
- [4] Dilek, Esma, and Özgür Talih. "Overview of Cyber Espionage Incidents and Analysis of Tackling Methods." *2022 15th International Conference on Information Security and Cryptography (ISCTURKEY)*. IEEE, 2022.
- [5] Foerderer, Jens, and Sebastian W. Schuetz. "Data breach announcements and stock market reactions: a matter of timing?." *Management Science* 68.10 (2022): 7298-7322.
- [6] Busch-Casler, Julia, and Marija Radic. "Personal Data Markets: A Narrative Review on Influence Factors of the Price of Personal Data." *Research Challenges in Information Science: 16th International Conference, RCIS 2022, Barcelona, Spain, May 17–20, 2022*.
- [7] Milunovich, George, and Seung Ah Lee. "Measuring the impact of digital exchange cyberattacks on Bitcoin Returns." *Economics Letters* 221 2022.
- [8] Adnan, Md Akhtaruzzaman, et al. "Identification of Data Breaches from Public Forums." *Innovative Security Solutions for Information Technology and Communications: 14th International Conference, SecITC 2021, Virtual Event, Kasım 25–26, 2021*.
- [9] US Department of Health and Human Services. "State and Tribal Child Welfare Information Systems, Information Security Data Breach Response Plans." *Administration for Children and Families* (2015).

- [10] Sen, Ravi, and Sharad Borle. "Estimating the contextual risk of data breach: An empirical approach." *Journal of Management Information Systems* 32.2 (2015): 314-341.
- [11] Index, Breach Level. *Data Privacy and New Regulations Take Center Stage. Technical Report. Breach Level Index*, 2018.
- [12] F. Maschler, F. Niephaus, J. Risch, Real or fake? large-scale validation of identity leaks, *INFORMATIK* 2017.
- [13] Seh, Adil Hussain, et al. "Healthcare data breaches: insights and implications." *Healthcare*. Vol. 8. No. 2. MDPI, 2020.
- [14] M. H. Barkadehi, M. Nilashi, O. Ibrahim, A. Z. Fardi, S. Samad, Authentication systems: A literature review and classification, *Telematics and Informatics* 35 (5) (2018) 1491–1511.
- [15] D. L. Wheeler, zxcvbn: Low-budget password strength estimation, in: 25th {USENIX} Security Symposium ({USENIX} Security 16), 2016, pp. 157– 173.
- [16] G. Hu, On password strength: a survey and analysis, in: *International Conference on Software Engineering, Artificial Intelligence, Networking and Parallel/Distributed Computing*, Springer, 2017, pp. 165–186.
- [17] T. Malderle, M. Wübbeling, S. Knauer, A. Sykosch, M. Meier, Gathering and analyzing identity leaks for a proactive warning of affected users, in: *Proceedings of the 15th ACM international conference on computing frontiers*, 2018, pp. 208–211.
- [18] Y.-Y. Choong, M. F. Theofanos, H.-K. Liu, United States Federal Employees' Password Management Behaviors: A Department of Commerce Case Study, US Department of Commerce, National Institute of Standards and Technology, 2014.
- [19] Hammouchi, Hicham, et al. "Digging deeper into data breaches: An exploratory data analysis of hacking breaches over time." *Procedia Computer Science* 151 (2019): 1004-1009.
- [20] Murmu, Sanjay, Harsh Kasyap, and Somanath Tripathy. "PassMon: A Technique for Password Generation and Strength Estimation." *Journal of Network and Systems Management* 30 (2022): 1-23.

- [21] Güven, Ebu Yusuf, Ali Boyaci, and Muhammed Ali Aydin. "A novel password policy focusing on altering user password selection habits: a statistical analysis on breached data." *Computers & Security* 113 (2022): 102560.
- [22] Islam, Chadni, et al. "SmartValidator: A framework for automatic identification and classification of cyber threat data." *Journal of Network and Computer Applications* 202 (2022): 103370.
- [23] Roberts, A. (2021). *Cyber Threat Intelligence – What Does It Even Mean?*. In: *Cyber Threat Intelligence*. Apress, Berkeley, CA.
- [24] Zheng, Denise E., and James A. Lewis. "Cyber threat information sharing." *Center for Strategic and International Studies* 62 (2015).
- [25] Haass, Jon C., Gail-Joon Ahn, and Frank Grimmelmann. "Actra: A case study for threat information sharing." *Proceedings of the 2nd ACM workshop on information sharing and collaborative security*. 2015.
- [26] National defense authorization act for fiscal year 2006 [online]. 2006. Available from: <https://www.gpo.gov/fdsys/pkg/PLAW-109publ163/html/PLAW-109publ163>.
- [27] Dokman, Tomislav, and Tomislav Ivanjko. "Open source intelligence (OSINT) issues and trends." *The Future of Information Sciences* 1.2020 (2020): 191.
- [28] Evangelista, João Rafael Gonçalves, et al. "Systematic literature review to investigate the application of open source intelligence (osint) with artificial intelligence." *Journal of Applied Security Research* 16.3 (2021): 345-369.
- [29] Nimier, Vincent, et al. "Challenges in automated HUMINT processing for situational assessment: Experiences from NATO CIMIC Joint Cooperation." *2020 IEEE 23rd International Conference on Information Fusion (FUSION)*. IEEE, 2020.
- [30] Kris, David S. "The NSA's New SIGINT Annex." *Journal of National Security Law & Policy* (2021).
- [31] Dover, Robert. "SOCMINT: a shifting balance of opportunity." *Intelligence and National Security* 35.2 (2020): 216-232.

- [32] National Cyber Security Centre. Threat Intelligence: Collecting, Analysing, Evaluating. Retrieved 31 May 2018. 2015.
- [33] Hutchins, Eric M., Michael J. Cloppert, and Rohan M. Amin. "Intelligence-driven computer network defense informed by analysis of adversary campaigns and intrusion kill chains." *Leading Issues in Information Warfare & Security Research* 1.1 (2011): 80.
- [34] Caltagirone, Sergio, Andrew Pendergast, and Christopher Betz. The diamond model of intrusion analysis. Center For Cyber Intelligence Analysis and Threat Research Hanover Md, 2013.
- [35] Bianco, D. The pyramid of pain.
<http://detectrespond.blogspot.no/2013/03/the-pyramid-of-pain.html>, 2014.
- [36] H. Almohannadi, I. Awan, J. Al Hamar, A. Cullen, J. P. Disso, and L. Armitage, "Cyber threat intelligence from honeypot data using elasticsearch," *Proc. - Int. Conf. Adv. Inf. Netw. Appl. AINA*, vol. 2018-May, pp. 900–906, 2018, doi: 10.1109/AINA.2018.00132.
- [37] K. Oosthoek and C. Doerr, "Cyber Threat Intelligence : A Product Without a Process ? Cyber Threat Intelligence : A Product Without a Process ?," *Int. J. Intell. CounterIntelligence*, vol. 34, no. 2, pp. 300–315, 2021, doi: 10.1080/08850607.2020.1780062.
- [38] Ryan Stillions. The DML model. Retrieved 31 May 2018. 2014. URL:
<http://ryanstillions.blogspot.com/2014/04/on-ttps.html>.
- [39] Wagner, Cynthia, et al. "Misp: The design and implementation of a collaborative threat intelligence sharing platform." *Proceedings of the 2016 ACM on Workshop on Information Sharing and Collaborative Security*. 2016.
- [40] Strom, Blake E., et al. "Mitre att&ck: Design and philosophy." Technical report. The MITRE Corporation, 2018.
- [41] Bromander, Siri, et al. "Investigating sharing of cyber threat intelligence and proposing a new data model for enabling automation in knowledge representation and exchange." *Digital Threats: Research and Practice (DTRAP)* 3.1 (2021): 1-22.
- [42] Sholihah, Intan Maratus, Hermawan Setiawan, and Olga Geby Nabila. *Cyber Threat Intelligence as an Knowledge Sharing Platform*. No. 7077. EasyChair, 2021.

- [43] U.S. Department of Health and Human Services, State and Tribal Child Welfare Information Systems, Information Security Data Breach Response Plans, Administration for Children and Families, 2015.
- [44] S. Goode, H. Hoehle, V. Venkatesh, S.A. Brown, User compensation as a data breach recovery action: An investigation of the Sony PlayStation network breach, *MIS Quarterly* 41 (2017) 703–727.
- [45] R. Sen, S. Borle, Estimating the contextual risk of data breach: An empirical approach, *Journal of Management Information Systems* 32 (2015) 314–341.
- [46] Breach Level Index, Data privacy and new regulations take center stage, 2018. Gemalto, <https://breachlevelindex.com/request-report?thanks=1>.
- [47] Osliak Oleksii, Saracino Andrea, Martinelli Fabio, "A scheme for the sticky policy representation supporting secure cyber-threat intelligence analysis and sharing", *Information & Computer Security*, Emerald Publishing Limited, 687,710, 2019.
- [48] Pamplin, Jeremy, et al. "Improving clinician decisions and communication in critical care using novel information technology." *Military medicine* 185.1-2 (2020): e254-e261.
- [49] “Kişisel Verilerin Korunması Kanunu”, 2016, PP. 1-19
<https://www.mevzuat.gov.tr/mevzuatmetin/1.5.2547.pdf>
- [50] Akkurt, S. Sami. "Kişisel Veri Kavramının Hukuki Niteliğine İlişkin Yaklaşımlara Mukayeseli Bir Bakış." *Kişisel Verileri Koruma Dergisi* 2.1 (2020): 20-32.
- [51] Dokuchaev, Vladimir A., Victoria V. Maklachkova, and V. Yu Statev. "Classification of personal data security threats in information systems." *T-Comm-Телекоммуникации и Транспорт* 14.1 (2020): 56-60.
- [52] Çalışkan, A. "Avrupa İnsan Hakları Mahkemesi Kararları Işığında Kişisel Verilerin Korunması."
- [53] Henkoğlu, Türkay. *Bilgi güvenliği ve kişisel verilerin korunması*. Yetkin Yayınları, 2015.
- [54] Akkurt, Sinan Sami. *Sosyal medyada gerçekleşen ihlaller karşısında kişilik hakkının korunması*. Seçkin Yayıncılık, 2019.

- [55] Cheng, Long, Fang Liu, and Danfeng Yao. "Enterprise data breach: causes, challenges, prevention, and future directions." *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery* 7.5 (2017): e1211.
- [56] Shukla, Neeraj Chandraprakash. *Converting REGEX to Parsing Expression Grammar with Captures*. North Carolina State University, 2021.
- [57] H. Zhao, W. Meng, Z. Wu, V. Raghavan and C. Yu, "Fully automatic wrapper generation for search engines", *Proceedings of the 14th international conference on World Wide Web*, pp. 66-75, 2005.
- [58] A. S. Yeole and B. B. Meshram, "Analysis of different technique for detection of sql injection", *Proceedings of the International Conference & Workshop on Emerging Trends in Technology*, pp. 963-966, 2011.
- [59] Liu, Tingwen, et al. "A prefiltering approach to regular expression matching for network security systems." *Applied Cryptography and Network Security: 10th International Conference, ACNS 2012*.
- [60] D. Ficara, S. Giordano, G. Procissi, F. Vitucci, G. Antichi and A. Di Pietro, "An improved dfa for fast regular expression matching", *ACM SIGCOMM Computer Communication Review*, vol. 38, no. 5, pp. 29-40, 2008.
- [61] Bhardwaj, Bhavya, et al. "Web scraping using summarization and named entity recognition (ner)." *2021 7th international conference on advanced computing and communication systems (ICACCS)*. Vol. 1. IEEE, 2021.
- [62] Erdogan, Hakan. "Sequence labeling: Generative and discriminative approaches." *Proc. 9th Int. Conf. Mach. Learn. Nisan 2010*.
- [63] Kang, Yue, et al. "Natural language processing (NLP) in management research: A literature review." *Journal of Management Analytics* 7.2 (2020): 139-172.
- [64] Leeson, William, et al. "Natural language processing (Nlp) in qualitative public health research: a proof of concept study." *International Journal of Qualitative Methods* 18 (2019).
- [65] Mikolov, Tomas, et al. "Distributed representations of words and phrases and their compositionality." *Advances in neural information processing systems* 26 (2013).

- [66] Hossain, M. Anwar, and Sadia Afrin. "Optical character recognition based on template matching." *Global Journal of Computer Science and Technology* 19.C2 (2019): 31-35.
- [67] Kumar, M. Suresh, et al. "Credit card fraud detection using random forest algorithm." 2019 3rd International Conference on Computing and Communications Technologies (ICCCCT). IEEE, 2019.
- [68] G. Karatas, O. Demir, and O. K. Sahingoz, "Increasing the Performance of Machine Learning-Based IDSs on an Imbalanced and Up-to-Date Dataset," *IEEE Access*, vol. 8, pp. 32150–32162
- [69] A. Makuvaza, D. S. Jat, and A. M. Gamundani, "Deep Neural Network (DNN) Solution for Real-time Detection of Distributed Denial of Service (DDoS) Attacks in Software Defined Networks (SDNs)," *SN Comput. Sci.*, vol. 2, no. 2, p. 107
- [70] Pisner, D. A., & Schnyer, D. M. (2020). Support vector machine. In *Machine Learning* (pp. 101-121). Academic Press.
- [71] Cervantes, J., Garcia-Lamont, F., Rodríguez-Mazahua, L., & Lopez, A. (2020). A comprehensive survey on support vector machine classification: Applications, challenges and trends. *Neurocomputing*, 408, 189-215.
- [72] Schnyer, David M., et al. "Evaluating the diagnostic utility of applying a machine learning algorithm to diffusion tensor MRI measures in individuals with major depressive disorder." *Psychiatry Research: Neuroimaging* 264 (2017): 1-9.
- [73] Freund Y., Schapire R.E. A decision-theoretic generalization of on-line learning and an application to boosting *J. Comput. Syst. Sci.*, 55 (1997), pp. 119-139
- [74] Gallego A.-J., Calvo-Zaragoza J., Valero-Mas J.J., Rico-Juan J.R. Clustering-based k-nearest neighbor classification for large-scale data with neural codes representation *Pattern Recognit.*, 74 (10) (2018), pp. 531-543
- [75] Li, W., Chen, Y., & Song, Y. (2020). Boosted K-nearest neighbor classifiers based on fuzzy granules. *Knowledge-Based Systems*, 195, 105606.
- [76] Pan, Z., Wang, Y., & Pan, Y. (2020). A new locally adaptive k-nearest neighbor algorithm based on discrimination class. *Knowledge-Based Systems*, 204, 106185.

- [77] Chen, Shenglei, et al. "A novel selective naïve Bayes algorithm." *Knowledge-Based Systems* 192 (2020): 105361.
- [78] A. Makuvaza, D. S. Jat, and A. M. Gamundani, "Deep Neural Network (DNN) Solution for Real-time Detection of Distributed Denial of Service (DDoS) Attacks in Software Defined Networks (SDNs)," *SN Comput. Sci.*, vol. 2, no. 2, p. 107
- [79] Jeffrey L Elman. Finding structure in time. *Cognitive science*, 14(2):179–211, 1990.
- [80] Roy, Arya. "Recent trends in named entity recognition (NER)." *arXiv preprint arXiv:2101.11420* (2021).
- [81] Fielding, Alan H., and John F. Bell. "A review of methods for the assessment of prediction errors in conservation presence/absence models." *Environmental conservation* 24.1 (1997): 38-49.
- [82] Shahinfar, Saleh, Paul Meek, and Greg Falzon. "“How many images do I need?” Understanding how sample size per class affects deep learning model performance metrics for balanced designs in autonomous wildlife monitoring." *Ecological Informatics* 57 (2020): 101085.
- [83] Uddin, Muhammad Fahim. "Addressing accuracy paradox using enhanced weighted performance metric in machine learning." *2019 Sixth HCT Information Technology Trends (ITT). IEEE*, 2019.
- [84] P. Papadimitriou and H. Garcia-Molina, "Data leakage detection", *IEEE Trans. Knowl. Data Eng.*, vol. 23, pp. 51-63, Ocak 2011.
- [85] S. A. Kale and S. Kulkarni, "Data leakage detection", *Int. J. Adv. Res. Comput. Commun. Eng.*, vol. 1, no. 9, pp. 668-678, 2012.
- [86] Huang, Cheng, et al. "HackerRank: identifying key hackers in underground forums." *International Journal of Distributed Sensor Networks* 17.5 (2021): 15501477211015145.
- [87] Federal Trade Commission. "Reports in Response to Senate Appropriations Committee Report 116-111 on the FTC’s Use of Its Authorities and the Resources Used and Needed to Protect Consumer Privacy and Security (2020)." (2020).

- [88] Jayant Madhavan, David Ko, Lucja Kot, Vignesh Ganapathy, Alexander Carlin Rasmussen, Alon Y. Halevy, "Google's Deep Web crawl" Proceedings of the VLDB Endowment, Volume 1, Issue 2 August 2008
- [89] Chen, H., Chung, W., Qin, J., Reid, E., Sageman, M. and Weimann, G. (2008), Uncovering the dark Web: A case study of Jihad on the Web. *J. Am. Soc. Inf. Sci.*, 59: 1347-1359.
- [90] Cost of a Data Breach Report 2022. 2022. Available at: <https://www-03.ibm.com/security/data-breach>. (Accessed December 21, 2022).
- [91] Shu X, Yao D, Bertino E. Privacy-preserving detection of sensitive data exposure. *IEEE Trans Inform Forensic Secur* 2015, 10:1092–1103.
- [92] Liu F, Shu X, Yao D, Butt AR. Privacy-preserving scanning of big content for sensitive data exposure with MapReduce. In: Proceedings of the 5th ACM Conference on Data and Application Security and Privacy, CODASPY 2015, San Antonio, TX, 2–4 March, 2015, 195–206.
- [93] Shapira Y, Shapira B, Shabtai A. Content-based data leakage detection using extended fingerprinting. *CoRR abs/1302.2028*. 2013.
- [94] Shu X, Zhang J, Yao D, Feng W. Rapid and parallel content screening for detecting transformed data exposure. In: Proceedings of the Third International Workshop on Security and Privacy in Big Data (BigSecurity), Hongkong, China, April, 2015.
- [95] Shu X, Zhang J, Yao DD, Feng WC. Fast detection of transformed data leaks. *IEEE Trans Inform Forensic Secur* 2016, 11:528–542.
- [96] Costante E, Fauri D, Etalle S, Hartog JD, Zannone N. A hybrid framework for data loss prevention and detection. In: 2016 I.E. Security and Privacy Workshops (SPW), San Jose, USA, 2016, 324–333.
- [97] Katz, Gilad, Yuval Elovici, and Bracha Shapira. "CoBAn: A context based model for data leakage prevention." *Information sciences* 262 (2014): 137-158.
- [98] Ghouse, Mohammed, Manisha J. Nene, and C. Vembuselvi. "Data Leakage Prevention for Data in Transit using Artificial Intelligence and Encryption Techniques." 2019

International Conference on Advances in Computing, Communication and Control (ICAC3). IEEE, 2019.

[99] Carvalho VR, Cohen WW. Preventing Information Leaks in Email. In: Proceedings of the 2007 SIAM International Conference on Data Mining, Minneapolis, USA. Philadelphia: SIAM; 2007, 68–77.

[100] Cheng, Long, Fang Liu, and Danfeng Yao. "Enterprise data breach: causes, challenges, prevention, and future directions." *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery* 7.5 (2017): e1211.

[101] Shaikh, Rizwana, and M. Sasikumar. "Data classification for achieving security in cloud computing." *Procedia computer science* 45 (2015): 493-498.

[102] Yayık, Apdullah, Hasan Apik, and Ayşe Tosun. "Deep learning based topic classification for sensitivity assignment to personal data." *2021 6th International Conference on Computer Science and Engineering (UBMK)*. IEEE, 2021.

[103] Y. Fang, Y. Guo, C. Huang and L. Liu, "Analyzing and Identifying Data Breaches in Underground Forums," in *IEEE Access*, vol. 7, pp. 48770-48777, 2019.

[104] Noor, Umara, et al. "A machine learning framework for investigating data breaches based on semantic analysis of adversary's attack patterns in threat intelligence repositories." *Future Generation Computer Systems* 95 (2019): 467-487.

[105] Shankar, Venkatesh Gauri, Bali Devi, and Sumit Srivastava. "DataSpeak: data extraction, aggregation, and classification using big data novel algorithm." *Computing, Communication and Signal Processing: Proceedings of ICCASP 2018*. Springer Singapore, 2019.

[106] Qian, Jun, et al. "Analyzing SocialArks Data Leak-A Brute Force Web Login Attack." *2022 4th International Conference on Computer Communication and the Internet (ICCCI)*. IEEE, 2022.

[107] Saleem, Hamza, and Muhammad Naveed. "SoK: Anatomy of Data Breaches." *Proc. Priv. Enhancing Technol.* 2020.4 (2020): 153-174.

- [108] Nayak, S.K., Ojha, A.C. (2020). Data Leakage Detection and Prevention: Review and Research Directions. In: Swain, D., Pattnaik, P., Gupta, P. (eds) Machine Learning and Information Processing. Advances in Intelligent Systems and Computing, vol 1101. Springer, Singapore.
- [109] R. Naik and M. N. Gaonkar, "Data Leakage Detection in cloud using Watermarking Technique," 2019 International Conference on Computer Communication and Informatics (ICCCI), Coimbatore, India, 2019, pp. 1-6, doi: 10.1109/ICCCI.2019.8821894.
- [110] Xi, Ning, et al. "Information flow based defensive chain for data leakage detection and prevention: a survey." arXiv preprint arXiv:2106.04951 (2021).
- [111] Wang, Yun, and Limei Wang. "Internet Financial Data Security and Economic Risk Prevention for Android Application Privacy Leakage Detection." Computational Intelligence and Neuroscience 2022 (2022).
- [112] Khaleghimoghaddam, Amir. "Exploring data leakage via supervised learning." (2020).
- [113] Go, Sharleen Joy Y., et al. "An SDN/NFV-enabled architecture for detecting personally identifiable information leaks on network traffic." 2019 Eleventh International Conference on Ubiquitous and Future Networks (ICUFN). IEEE, 2019.
- [114] A. Yayık, H. Apik and A. Tosun, "Deep Learning Based Topic Classification for Sensitivity Assignment to Personal Data," 2021 6th International Conference on Computer Science and Engineering (UBMK), Ankara, Turkey, 2021, pp. 292-297.
- [115] García-Pablos, Aitor, Naiara Perez, and Montse Cuadros. "Sensitive data detection and classification in Spanish clinical text: Experiments with BERT." arXiv preprint arXiv:2003.03106 (2020).
- [116] R Steele. Open source intelligence. Handbook of Intelligence Studies, 42(5):129–147, 2007.
- [117] Rajamäki, Jyri, and Jussi Simola. "„How to apply privacy by design in osint and big data analytics?“. " ECCWS 2019 18th European Conference on Cyber Warfare and Security. Academic Conferences and publishing limited, 2019.

[118] Zowalla R, Wetter T, Pfeifer D Crawling the German Health Web: Exploratory Study and Graph Analysis J Med Internet Res 2020 22(7).

[119] Lande, Dmytro, Igor Subach, and Alexander Puchkov. "A system for analysis of big data from social media." *Information & Security* 47.1 (2020): 44-61.

[120] Rodrigues, Mónica Sofia Amoroso. *Dynamic OSINT System Sourcing from Social Networks*. Diss.

[121] TUNDIS, Andrea; RUPPERT, Samuel; MÜHLHÄUSER, Max. On the Automated Assessment of Open-Source Cyber Threat Intelligence Sources. In: *International Conference on Computational Science*. Springer, Cham, 2020. p. 453-467.

[122] D. Goularas and S. Kamis, "Evaluation of Deep Learning Techniques in Sentiment Analysis from Twitter Data," 2019 International Conference on Deep Learning and Machine Learning in Emerging Applications (Deep-ML), Istanbul, Turkey, 2019, pp. 12-17.

[123] Z. Abbass, Z. Ali, M. Ali, B. Akbar and A. Saleem, "A Framework to Predict Social Crime through Twitter Tweets By Using Machine Learning," 2020 IEEE 14th International Conference on Semantic Computing (ICSC), San Diego, CA, USA, 2020, pp. 363-368.

İNTİHAL RAPORU İLK SAYFASI

Ebu Yusuf Güven

ORJİNALLİK RAPORU

%5

BENZERLİK ENDEKSİ

%4

İNTERNET KAYNAKLARI

%1

YAYINLAR

%1

ÖĞRENCİ ÖDEVLERİ

BİRİNCİL KAYNAKLAR

1

acikbilim.yok.gov.tr

İnternet Kaynağı

%1

2

Ebu Yusuf Güven, Ali Boyaci, Muhammed Ali Aydin. "A Novel Password Policy Focusing on Altering User Password Selection Habits: A Statistical Analysis on Breached Data", Computers & Security, 2021

Yayın

%1

3

Submitted to The Scientific & Technological Research Council of Turkey (TUBITAK)

Öğrenci Ödevi

<%1

4

siberkavramlab.blogspot.com

İnternet Kaynağı

<%1

5

cdn.istanbul.edu.tr

İnternet Kaynağı

<%1

6

Submitted to Mersin Üniversitesi

Öğrenci Ödevi

<%1

7

acikerisim.akdeniz.edu.tr:8080

İnternet Kaynağı

<%1