

T.C.
İSTANBUL BEYKENT ÜNİVERSİTESİ
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ

BİLGİSAYAR MÜHENDİSLİĞİ ANABİLİM DALI
BİLGİSAYAR MÜHENDİSLİĞİ BİLİM DALI

**MAKİNE ÖĞRENİMİNİN GÜNÜMÜZDEKİ ROLÜ:
MODERN YÖNELİMLER VE GELECEK
PERSPEKTİFLER**

Yüksek Lisans Tezi

Tezi Hazırlayan
Mirjafar GASIMLI

İstanbul, 2025

T.C.
İSTANBUL BEYKENT ÜNİVERSİTESİ
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ

BİLGİSAYAR MÜHENDİSLİĞİ ANABİLİM DALI
BİLGİSAYAR MÜHENDİSLİĞİ BİLİM DALI

**MAKİNE ÖĞRENİMİNİN GÜNÜMÜZDEKİ ROLÜ:
MODERN YÖNELİMLER VE GELECEK
PERSPEKTİFLER**

Yüksek Lisans Tezi

Tezi Hazırlayan
Mirjafar GASIMLI

Öğrenci No
2320003027

ORCID ID
0009-0002-9979-579X

Danışman
Dr. Öğr. Üyesi Talat FİRLAR

İstanbul, 2025

YEMİN METNİ

Yüksek Lisans Tezi olarak sunduğum “**Makine Öğreniminin Günümüzdeki Rolü: Modern Yönelimler ve Gelecek Perspektifler**” başlıklı bu çalışmanın, bilimsel ahlak ve geleneklere uygun şekilde tarafımdan yazıldığını, bu tezdeki bütün bilgileri akademik ve etik kurallar içinde elde ettiğimi, yararlandığım eserlerin tamamının kaynaklarda gösterildiğini ve çalışmamın içinde kullanıldıkları her yerde bunlara atıf yapıldığını, patent ve telif haklarını ihlal edici bir davranışımın olmadığını belirtir ve bunu onurumla doğrularım. 11/06/2025

Mirjafar GASIMLI

T.C.
İSTANBUL BEYKENT ÜNİVERSİTESİ
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ MÜDÜRLÜĞÜ
TEZLİ YÜKSEK LİSANS SINAV TUTANAĞI

11/06/2025

Enstitümüz *Bilgisayar Mühendisliği* Anabilim Dalı *Bilgisayar Mühendisliği* Programı yüksek lisans öğrencilerinden 2320003027 numaralı *Mirjafar GASIMLI*'nin “*İstanbul Beykent Üniversitesi Lisansüstü Eğitim – Öğretim Yönetmeliği*”nin ilgili maddesine göre hazırlayarak, Enstitümüze teslim ettiği “*Makine Öğreniminin Günümüzdeki Rolü: Modern Yönelimler ve Gelecek Perspektifler*” konulu tezini, Yönetim Kurulumuzun 20/05/2025 tarih ve 2025/22 sayılı toplantısında seçilen ve On-Line toplanan biz jüri üyeleri huzurunda, İstanbul Beykent Üniversitesi Lisansüstü Eğitim ve Öğretim Yönetmeliğinin 29. maddesinin 3. fıkrası gereğince Zoom programı aracılığıyla on-line olarak aday tarafından savunulmuş ve sonuçta adayın tezi hakkında “*OYBİRLİĞİ*” ile “*KABUL*” kararı verilmiştir.

İşbu tutanak, 2 nüsha olarak hazırlanmış ve Enstitü Müdürlüğü’ne sunulmak üzere tarafımızdan düzenlenmiştir.

DANIŞMAN
Dr. Öğr. Üyesi Ta*** Fİ***
(İstanbul Beykent Üniversitesi)

ÜYE
Dr. Öğr. Üyesi Ke*** Gö*** NA***
(İstanbul Beykent Üniversitesi)

ÜYE
Doç. Dr. Em*** SE***
(İstinye Üniversitesi)

Adı ve Soyadı : Mirjafar GASIMLI
Danışmanı : Dr. Öğr. Üyesi Talat FİRLAR
Derecesi ve Tarihi : Yüksek Lisans (Tezli), 2025
Alanı : Bilgisayar Mühendisliği
Anahtar Kelimeler : Makine Öğrenmesi, Siber Güvenlik, Saldırı Algılama Sistemleri

ÖZ

MAKİNE ÖĞRENİMİNİN GÜNÜMÜZDEKİ ROLÜ: MODERN YÖNELİMLER VE GELECEK PERSPEKTİFLER

İnsanlar evrimlerinden bu yana çeşitli görevleri daha basit bir şekilde gerçekleştirmek için birçok araç türü kullanmaktadır. Makine öğreniminin amacı verilerden öğrenmektir. Makinelerin açıkça programlanmadan kendi kendilerine nasıl öğrenecekleri konusunda birçok çalışma yapılmıştır. Birçok matematikçi ve programcı, bu sorunun çözümünü bulmak için büyük veri kümelerine sahip çeşitli yaklaşımlar uygulamaktadır. Bu çalışmada, makine öğrenimi algoritmalarının günümüzdeki rolü, temel yaklaşımları ve gelecekteki yönelimleri incelenmiştir.

İlk bölümde makine öğreniminin mevcut durumu, veri ve donanım bağımlılıkları, özellik işleme teknikleri, sorun çözme yöntemleri, yürütme süresi, yorumlanabilirlik ve model seçimi gibi teorik temeller ele alınmıştır. İkinci bölümde, makine öğreniminin siber güvenlik alanındaki yeri açıklanmış, özellikle saldırı tespit sistemlerinde ve veri analizinde kullanımı değerlendirilmiştir. Ayrıca NSL-KDD veri seti ve KNIME aracı kullanılarak yapılan bir uygulama üzerinden algoritmaların performansı analiz edilmiştir. Üçüncü bölümde, makine öğreniminin gelişimi, ürün yorumları üzerinden duygu analizi, derin öğrenmedeki trendler, karşılaşılan güncel problemler ve bu problemlere önerilen çözümler sunulmuştur. Çalışmanın sonunda, makine öğreniminin hem bugünkü etkisi hem de gelecekteki potansiyeli hakkında genel bir değerlendirme yapılmıştır.

Name and Surname : Mirjafar GASIMLI
Supervisor : Asst. Prof. Dr. Talat FIRLAR
Degree and Date : Master's (Thesis), 2025
Major : Computer Engineering
Keywords : Machine Learning, Cyber Security, Intrusion Detection Systems

ABSTRACT

DETERMINING THE CURRENT ROLE, MODERN ASPECTS AND FUTURE ORIENTATIONS OF MACHINE LEARNING ALGORITHMS

Since their evolution, humans have used many types of tools to perform various tasks more easily. The creativity of the human brain has led to the invention of different machines. These machines have made human life easier by enabling people to meet various life needs such as travel, industry, and computing. The aim of machine learning is to learn from data. Many studies have been conducted on how machines can learn on their own without being explicitly programmed. Many mathematicians and programmers apply various approaches with large data sets to find a solution to this problem. In this study, the current role of machine learning algorithms, their basic approaches, and future trends are examined.

In the first section, the current status of machine learning, theoretical foundations such as data and hardware dependencies, feature processing techniques, problem solving methods, execution time, interpretability, and model selection are discussed. In the second section, the place of machine learning in the field of cyber security is explained, and its use in intrusion detection systems and data analysis is evaluated. In addition, the performance of the algorithms is analyzed through an application using the NSL-KDD data set and the KNIME tool. In the third section, the development of machine learning, sentiment analysis on product reviews, trends in deep learning, current problems encountered and proposed solutions to these problems are presented. At the end of the study, a general assessment is made about both the current impact and future potential of machine learning.

İÇİNDEKİLER

Sayfa No.

ÖZ

ABSTRACT

TABLolar LİSTESİ iii

ŞEKİLLER LİSTESİ iv

KISALTMALAR v

SÖZLÜK..... vii

GİRİŞ 1

1. MAKİNE ÖĞRENİMİNİN MEVCUT DURUMUNU VE TEORİK

TEMELLERİNİ İNCELEMEK

1.1. Makine Öğreniminin Temel Unsurları 2

1.1.1. Veri Bağımlılıkları..... 3

1.1.2. Donanım Bağımlılıkları..... 4

1.1.3. Özellik İşleme..... 4

1.1.4. Sorun Çözme Yöntemi 4

1.1.5. Yürütme Süresi 4

1.1.6. Yorumlanabilirlik 5

1.1.7. Model Seçimi..... 5

1.2. Günümüzde Yapay Zeka Algoritmaları 5

1.2.1. Veri Toplama..... 7

1.2.2. Model Seçimi ve Eğitimi..... 7

1.2.3. Model Değerlendirmesi 7

1.2.4. Veri Ön İşleme 7

1.2.5. Süperparametrelerin Tahmini ve Ayarlanması..... 8

1.2.6. Gözetimli Öğrenme 9

1.2.7. Gözetimsiz Öğrenme 9

1.2.8. Güçlendirme Öğrenimi 10

1.3. Makine Öğrenimini Uygulamaları 10

1.3.1. Görüntü İşleme 10

1.3.2. Konuşma Tanıma..... 11

1.3.3. Hastalığın Teşhisi	11
1.3.4. İstatistiksel Analiz	11
1.3.5. Finansal Hizmetler.....	12
1.4. Günümüzde Makine Öğrenimi	12
1.4.1. Finansal Hizmetler.....	12
1.4.2. Büyük Binalarda HVAC Enerji Tüketiminin Optimizasyonu.....	13
1.4.3. Otonom Arabalar	13
1.4.4. Ulaşımveri Analizi.....	13
1.5. Makine Öğreniminin Sorunları	14
1.5.1. Yetenek Açığı.....	15
1.5.2. Veri Kalitesi ve Niceliği	18
1.5.3. Model Yorumlanabilirliği.....	18
1.5.4. Ölçeklenebilirlik Sorunları	19
1.5.5. Düzenleyici Uyumluluk.....	19
2. YAPAY ZEKADA MAKİNE ÖĞRENMESİNİN ROLÜ	
2.1. Siber Güvenlikte Saldırı Tespit Sistemleri ve Makine Öğrenimi Yaklaşımları.....	21
2.2. NSL-KDD Veri Seti ve Saldırı Kategorileri Bağlamında Siber Güvenlik Yaklaşımları.....	29
2.3. Nsl-Kdd Veri Seti İle Knime Uygulaması Kullanarak Saldırı Tespiti ve Makine Öğrenimi Algoritmalarının Performans Analizi.....	34
2.4. Makine Öğrenimi İçin Turing Testi	49
3. MAKİNE ÖĞRENİMİNİN GELİŞİMİ VE GELECEK PERSPEKTİFLERİ	
3.1. Makine Öğrenme Algoritması Kullanılarak Amazon Ürünlerinin Duygu Analizi.....	56
3.2. Makine Öğrenimi ve Derin Öğrenmedeki Güncel Trendler	64
3.3. Makine Öğrenimindeki Güncel Sorunlar	70
3.4. Güncel Sorunların Çözümü	73
SONUÇ	79
KAYNAKÇA.....	81

TABLÖLAR LİSTESİ

	Sayfa No.
Tablo 1. Veri Seti İçeriği	41
Tablo 2. Confusion Matrix	44
Tablo 3. Confusion Matrix (2).....	45



ŞEKİLLER LİSTESİ

	Sayfa No.
Şekil 1. İki Seviyeli Karar Ağaçları.....	35
Şekil 2. Rastgele Ormanlar ve Karar Ağaçları	36
Şekil 3. SVM, Optimum Hiper Düzlem	37
Şekil 4. Destek Vektör Noktaları ve Kenar Boşlukları.....	37
Şekil 5. Basit Bir Gizli Katmana Sahip Basit Bir Sinir Ağı	38
Şekil 6. Topluluğa Sırayla Yeni Modeller Ekleyin	38
Şekil 7. Random Forest Predictor	42
Şekil 8. Naive Bayes Learner	44
Şekil 9. Random Forest ve Naive Bayes Performans Karşılaştırması	46
Şekil 10. Turing Testinin Genelleştirilmiş Paradigmaları	52
Şekil 11. Değiştirilmiş Turing Test Paradigması	54
Şekil 12. Arama Modülü	58
Şekil 13. Ön İşleme Modülü	59
Şekil 14. Sınıflandırma Modülü	60
Şekil 15. Veri Sınıflandırmasının Analizi.....	64
Şekil 16. Makine Öğrenimi ve Derin Öğrenmedeki Mevcut Eğilimleri Gösteren Sankey Diyagramı	67

KISALTMALAR

BERT : Bidirectional Encoder Representations from Transformers (Çift Yönlü Kodlayıcı Temsiller – Transformer tabanlı)

CLIP : Contrastive Language–Image Pretraining (Karşıtlı Dil-Görüntü Ön Eğitimi)

CNN : Convolutional Neural Network (Evrimsel Sinir Ağı)

DL : Deep Learning (Derin Öğrenme)

DQN : Deep Q-Network (Derin Q Ağı)

GAN : Generative Adversarial Network (Üretici Çekişmeli Ağ)

HIDS: Host-based Intrusion Detection System (Ana Bilgisayar Tabanlı Saldırı Tespit Sistemi)

HVAC: Heating, Ventilation, and Air Conditioning (Isıtma, Havalandırma ve Klima)

IDS : Intrusion Detection System (Saldırı Tespit Sistemi)

IPSH: Intrusion Prevention System using Honeypots (Tuzak Sistemler ile Saldırı Önleme Sistemi)

KDD : Knowledge Discovery in Databases (Veritabanlarında Bilgi Keşfi)

kNN : k-Nearest Neighbors (k-En Yakın Komşu Algoritması)

ML : Machine Learning (Makine Öğrenmesi)

NAS : Neural Architecture Search (Sinirsel Mimari Arama)

NLP : Natural Language Processing (Doğal Dil İşleme)

NSL : NSL-KDD (KDD veri kümesinin geliştirilmiş versiyonu)

ONNX: Open Neural Network Exchange (Açık Sinir Ağı Değişim Formatı)

PCA : Principal Component Analysis (Temel Bileşenler Analizi)

PPO : Proximal Policy Optimization (Yakın Politika Optimizasyonu)

QML : Quantum Machine Learning (Kuantum Makine Öğrenmesi)

RNN : Recurrent Neural Network (Yinelemeli Sinir Ağı)

SAINT: Self-Attention Augmented Inference for Intrusion Detection (Saldırı Tespiti için Öz-Dikkat Destekli Çıkarım)

SGD : Stochastic Gradient Descent (Stokastik Gradyan İnişi)

SVM : Support Vector Machine (Destek Vektör Makineleri)

vd. : ve diğerler

XAI : Explainable Artificial Intelligence (Açıklanabilir Yapay Zeka)



SÖZLÜK

Makine Öğrenmesi (Machine Learning): Makine öğrenmesi, bilgisayarların veriden örüntüleri öğrenerek kararlar almasını sağlayan bir yöntemdir. Bu yöntem, açıkça programlama yapılmadan geçmiş verilere bakarak gelecekteki durumları tahmin eder. Günümüzde öneri sistemlerinden sağlık uygulamalarına kadar birçok alanda kullanılmaktadır.

Saldırı Tespit Sistemleri (Intrusion Detection Systems): Saldırı tespit sistemleri, bilgisayar veya ağ sistemlerine yapılan şüpheli girişimleri belirlemeye çalışır. Bu sistemler, ağ trafiğini inceleyerek olağandışı davranışları tespit eder. Genellikle erken uyarı sağlayarak güvenlik önlemlerinin zamanında alınmasına yardımcı olur.

Siber Güvenlik (Cyber Security): Siber güvenlik, bilgisayar sistemlerini ve dijital verileri dış tehditlerden korumayı amaçlar. Virüs, zararlı yazılım ve izinsiz girişler gibi saldırılara karşı savunma mekanizmaları geliştirir. Kişisel bilgilerden kurumsal ağlara kadar geniş bir koruma alanı sunar.

GİRİŞ

Teknolojik gelişmelerin hız kazanmasıyla birlikte, bilgisayar sistemlerinin verilerden öğrenerek karar verebilme yeteneği, pek çok alanda önemli kolaylıklar sağlamıştır. Bu gelişmelerin temelinde yer alan makine öğrenmesi, geçmiş veriler üzerinden örüntüler oluşturup çeşitli durumları tahmin edebilen yöntemler bütünü olarak öne çıkmaktadır. Açıkça programlanmadan çalışan bu sistemler, sağlık, finans, iletişim ve güvenlik gibi birçok sektörde etkili çözümler sunmaktadır.

Makine öğrenmesinin verimli bir şekilde uygulanabilmesi için yeterli miktarda, kaliteli ve anlamlı veriye ihtiyaç duyulmaktadır. Ayrıca, yüksek işlem gücüne sahip donanımlar ile doğru özellik seçimi ve bu özelliklerin etkili şekilde işlenmesi başarıyı doğrudan etkilemektedir. Farklı problemler için farklı algoritmalar ve modeller kullanılırken; modelin çalışma süresi, sonuçların doğruluğu ve açıklanabilirliği gibi unsurlar model seçiminde belirleyici rol oynamaktadır.

Bu çalışmada, makine öğrenmesinin temel kuramsal altyapısı ile birlikte veri ve donanım bağımlılıkları, model seçimi, yürütme süresi ve yorumlanabilirlik gibi konular ele alınmıştır. Ardından, makine öğrenmesinin siber güvenlikteki yeri incelenmiş; özellikle saldırı tespit sistemleri kapsamında NSL-KDD veri seti üzerinden yapılan analizlerle, çeşitli algoritmaların başarıları karşılaştırmalı olarak sunulmuştur. Son bölümde ise, makine öğrenmesinin gelişim süreci, güncel uygulama örnekleri (örneğin Amazon yorum analizi) ve bu alanda karşılaşılan güncel sorunlar değerlendirilmiştir. Bu kapsamda hem teknik ayrıntılar hem de uygulama örnekleriyle konu detaylandırılmıştır.

1. MAKİNE ÖĞRENİMİNİN MEVCUT DURUMUNU VE TEORİK TEMELLERİNİ İNCELEMEK

Makine öğreniminin mevcut durumu ve teorik temelleri incelenerek veri ve donanım bağımlılıkları, özellik işleme, sorun çözme yöntemleri, yürütme süresi, yorumlanabilirlik ve model seçimi gibi temel unsurlar değerlendirilmiştir.

1.1. Makine Öğreniminin Temel Unsurları

Makine öğrenimi, bir görevi gerçekleştirmek üzere açıkça programlanmak yerine örneklerden öğrenen bir bilgisayar algoritmaları sınıfını ifade eder. Somut örnekler kümesinden genel bir kural formüle etmeyi öğrenir. Böylece, insan öğrenmesi gibi, bilgisayar da edinilen bilgiden performansını iyileştirme yeteneğine sahip olur. Aradaki fark, bilginin mevcut durumunda, bilgisayarın insanlardan çok daha fazla öğrenme örneğine ihtiyaç duymasındır. Makine öğrenimi, yapay zekanın temelidir. Algoritmanın yapısına ve karmaşıklığına bağlı olarak sığ ve derin öğrenmeye ayrılabilir. Hem kendi önemleri hem de derin öğrenmede yer alan bazı kavramları ve ilkeleri tanıtmaya yardımcı olmaları nedeniyle birkaç örnek verilmiştir. Her iki makine öğrenimi biçiminin de yaygın olarak kullanıldığını bilmek önemlidir, çünkü belirli bir görev için birinin veya diğerinin en uygun olduğu durumlar vardır. Derin öğrenme için, çok katmanlı bağlantı, geri yayılım ve evrişim gibi ek kavramlar ayrıntılı olarak açıklanmaktadır, çünkü bunlar bu modelleri dağıtırken dikkate alınması gereken faktörlerdir. Taşımacılık sorunları da dahil olmak üzere farklı alanlardaki sorunlara uygun makine öğrenimi algoritmaları tasarlamak ve uygulamak, makine öğrenimi tekniklerinin ve temel teorilerin her köşesinin kapsamlı bir şekilde anlaşılmasını gerektirir. Bu bölüm, makine öğrenimi alanındaki temel kavramların bir yelpazesini, makine öğreniminin tanımı ve kategorileriyle başlayarak ve ardından gelişmiş makine öğrenimi algoritmalarının temel yapı taşlarını ele alarak tanıtmaktadır. Doğrusal regresyon ve lojistik regresyon, gradyan iniş algoritmaları, düzenleme ve makine öğreniminin diğer temel kavramları dahil olmak üzere yaygın regresyonların arkasındaki teori tartışılmaktadır. Ek olarak, bu bölüm veri ön işleme, eğitim ve test prosedürlerini yerine getirmek için temsili makine öğrenimi araçlarını tanıtmaktadır. Son zamanlarda, klasik Makine öğrenimi yöntemleri, iyi uyarlanabilirliğe sahip ve özellik mühendisliğine ihtiyaç duymayan uçtan uca Derin öğrenme yöntemleri yapısı

nedeniyle, birçok alanda derin öğrenme (DL) tarafından gölgede bırakılmıştır. Bununla birlikte, klasik ML yöntemleri DL yöntemlerinin yapı taşlarıdır ve ML ve DL'nin genel eğitim ve test süreci aynıdır. Bu arada, daha basit model yapısına sahip klasik yöntemler küçük verilerle iyi çalışır. Klasik ML yöntemlerinin yorumlanabilirliği, sözde "kara kutu" DL yöntemlerinin çoğundan bile daha iyi olabilir. Klasik ML yöntemleri, ulaşım uygulamalarında hala yaygın olarak benimsenmektedir, örneğin, ulaşım modu tanıma, yol yüzey durumu tespiti, tıkanıklık tespiti, sürücü davranış sınıflandırması, yolcu sayısı tahmini, darboğaz tanımlama ve araç ağı arıza tespiti. Bu nedenle, bu bölüm daha gelişmiş modelleri sunmadan önce en önemli makine öğrenimi temellerini tanıtmaktadır. Makine öğrenmesi her zaman yapay zekanın ayrılmaz bir parçası olmuştur ve metodolojisi, alanın başlıca endişeleriyle uyumlu bir şekilde gelişmiştir. Model AI sistemlerinde giderek artan bilgi hacimlerini kodlamanın zorluklarına yanıt olarak, birçok araştırmacı son zamanlarda bilgi edinme darboğazını aşmanın bir yolu olarak makine öğrenmesine dikkatlerini çevirmiştir (Carbonell, 2019).

Günümüzde, yapay zeka yetenekleri sunan akıllı sistemler genellikle makine öğrenimine güvenir. Makine öğrenimi, sistemlerin analitik model oluşturma sürecini otomatikleştirmek ve ilişkili görevleri çözmek için sorun-özel eğitim verilerinden öğrenme kapasitesini tanımlar. Derin öğrenme, yapay sinir ağlarına dayalı bir makine öğrenimi kavramıdır. Birçok uygulama için, derin öğrenme modelleri sık makine öğrenimi modellerinden ve geleneksel veri analizi yaklaşımlarından daha iyi performans gösterir.

1.1.1. Veri Bağımlılıkları

Derin öğrenme ile geleneksel makine öğrenimi arasındaki farklardan biri, veri miktarı arttıkça performansdır. Derin öğrenme algoritmaları, veri hacimleri küçük olduğunda o kadar iyi performans göstermez, çünkü derin öğrenme algoritmaları verileri mükemmel bir şekilde anlamak için büyük miktarda veri gerektirir. Tersine, bu durumda, geleneksel makine öğrenimi algoritması yerleşik kuralları kullandığında, performans daha iyi olacaktır (LeCun vd., 2015)

1.1.2. Donanım Bağımlılıkları

DL algoritması birçok matris işlemi gerektirir. Grafik İşlem Birimi, matris işlemlerini verimli bir şekilde optimize etmek için büyük ölçüde kullanılır. Bu nedenle, GPU, DL'nin düzgün çalışması için gerekli donanımdır. Derin öğrenme, geleneksel makine öğrenimi algoritmalarından daha çok GPU'lu yüksek performanslı makinelere güvenir.

1.1.3. Özellik İşleme

Özellik işleme, verilerin karmaşıklığını azaltmak ve öğrenme algoritmalarının daha iyi çalışmasını sağlayan kalıplar oluşturmak için alan bilgisini bir özellik çıkarıcıya koyma sürecidir. Özellik işleme zaman alıcıdır ve özel bilgi gerektirir. ML'de, bir uygulamanın özelliklerinin çoğu bir uzman tarafından belirlenmeli ve ardından bir veri türü olarak kodlanmalıdır. Özellikler piksel değerleri, şekiller, dokular, konumlar ve yönelimler olabilir. Çoğu ML algoritmasının performansı, çıkarılan özelliklerin doğruluğuna bağlıdır. Yüksek seviyeli özellikleri doğrudan verilerden elde etmeye çalışmak, DL ile geleneksel makine öğrenimi algoritmaları arasındaki temel farktır. Bu nedenle, DL her sorun için bir özellik çıkarıcı tasarlama çabasını azaltır.

1.1.4. Sorun Çözme Yöntemi

Sorunları çözmek için geleneksel makine öğrenimi algoritmaları uygulandığında, geleneksel makine öğrenimi genellikle sorunu birden fazla alt soruna ayırır ve alt sorunları çözerek nihai sonucu elde eder. Bunun aksine, derin öğrenme doğrudan uçtan uca sorun çözmeyi savunur.

1.1.5. Yürütme Süresi

Genel olarak, DL algoritmasında birçok parametre olduğu için bir DL algoritmasını eğitmek uzun zaman alır; bu nedenle eğitim adımı daha uzun sürer. ResNet gibi en gelişmiş DL algoritması, bir eğitim seansını tamamlamak için tam iki hafta sürerken, ML eğitimi nispeten az zaman alır, sadece saniyeler ile saatler. Ancak, test süresi tam tersidir. Derin öğrenme algoritmaları, test sırasında çalıştırmak için çok

az zaman gerektirir. Bazı ML algoritmalarıyla karşılaştırıldığında, test süresi veri miktarı arttıkça artar. Ancak, bu nokta tüm ML algoritmaları için geçerli değildir, çünkü bazı ML algoritmalarının kısa test süreleri vardır.

1.1.6. Yorumlanabilirlik

Yorumlanabilirlik, ML'yi DL ile karşılaştırmada önemli bir faktördür. DL'nin el yazısıyla yazılmış sayıları tanıması, insanların standartlarına yaklaşabilir, oldukça şaşırtıcı bir performans. Ancak, bir DL algoritması size bu sonucu neden sağladığını söylemeyecektir. Elbette, matematiksel bir bakış açısından, derin bir sinir ağının bir düğümü etkinleştirilir. Ancak, nöronlar nasıl modellenmelidir ve bu nöron katmanları birlikte nasıl çalışır? Bu nedenle, sonucun nasıl üretildiğini açıklamak zordur. Tersine, makine öğrenme algoritması, algoritmanın neden böyle seçtiğine dair açık kurallar sağlar; dolayısıyla kararın gerekçesini açıklamak kolaydır.

1.1.7. Model Seçimi

Pek çok pratik durumda, bir modelin seçimi modelin temel bir özelliği olan karmaşıklığa göre belirlenir. Derin öğrenme ve takviyeli öğrenme gibi alanlar araştırma topluluğu arasında popülerlik kazanmasına rağmen, pratikte genellikle daha basit modeller seçilir. Bu tür modeller arasında sığ sinir ağı mimarileri, Temel Bileşen Analizi'ne (PCA) dayalı basit yaklaşımlar, kararağaçları ve rastgele ormanlar bulunur. Basit modeller, önerilen ML çözümünün konseptini kanıtlamanın ve uçtan uca kurulumu yerine getirmenin bir yolu olarak kullanılabilir. Bu yaklaşım, dağıtılan bir çözümü elde etme süresini azaltır, önemli geri bildirimlerin toplanmasına olanak tanır ve ayrıca aşırı karmaşık tasarımlardan kaçınmaya yardımcı olur.

1.2. Günümüzde Yapay Zeka Algoritmaları

Bilgisayar görüntüsü, olağan madencilik, analiz ve yapıcı bilgilerin dikkate alınmasıyla bütünleşmiş bilgisayar biliminin bir alt bölümüdür. Basit bir ifadeyle, bilgisayar görüntüsü makineler bir insan olmadan sorunları nasıl görebilir/çözer" anlamına gelir. Buna örnek, Sophia robotu Hanson Robotics şirketi tarafından sunulmuş ve 2016 yılında Saudi Arabistanda ilk vatandaşlığı alan robot olmuştur. Bu

nedenle, bilgisayar görüntüsü alanının sanallaştırmayla ilgili çeşitli sorunları çözebilen gelecek vaat eden araştırma alanı haline geldiği söylenebilir. Bilgisayar görüntüsü, yeni ve gelişmiş teknolojiler veya kavramlar (Blockchain, Her Şeyin İnterneti vb.) ve farklı bilgisayar görüntüsü tekniklerini kullanan uygulamalarla genişlenmekte ve ortaya çıkmaktadır. Akıllı dünyaların vizyonunu ve gereksinimlerini karşılayabilmek için, yapay zeka ve makine öğrenmesi araçları ve uygulamaları, açıkça programlanmaksızın daha doğru ve anlamlı sonuçlar tahmin edebilme yeteneği sağlar. Yapay zeka algoritmaları için, Yapay Zekanın geleneksel teknikleri kullanılarak oluşturulan sistemlerde kullanılan çeşitli çıkarımlar, kurallar ve mantık, değişen dünyanın bugünkü gereksinimlerini karşılamıyor. Ayrışmada, sınıflandırma, kümeleme, regresyon için veri setinde var olan kalıpların analizine ve tespitine odaklanan sistemler, AI'nın baskın sistemi haline geliyor. Basit bir ifadeyle, bilgisayar görüşü, bir makinenin (bir insan olmadan) gördüğü makinelerin bilimi ve teknolojisidir. Bilgisayar görüşü, çeşitli grafik sorunlarına yaklaşmak için çok sayıda yöntemi içeren bir araştırma alanıdır. Son yıllarda, yüzlerce uygulama/birçok kuruluş, bilgisayar görüşünün amacına destek sağlamak için Yapay Zeka, Derin Öğrenme ve diğer Makine Öğrenmesi teknikleri ile değiştirildi. Dolayısıyla, akıllı dünyaların vizyonunu ve bu dünyanın gereksinimlerini karşılayabilmek için, yapay zeka ve makine öğrenmesi; araçların ve uygulamaların açıkça programlanmaksızın daha doğru ve değer açısından daha isabetli sonuçlar tahmin etmesini mümkün kılmaktadır.

Makine öğrenimi, yapay zekanın en önemli alt kümelerinden biridir. Makine öğrenimi, algoritmalar kullanarak, açıkça planlama yapmadan ve bireysel eylemleri dikte etmeden, gerçek veya simüle edilmiş örnekler ve veriler kümesini öğrenmek ve bunlar üzerinde işlem yapmak üzere tasarlanmış bir makine modelidir. Makine öğreniminin amacı, bilgisayarların ve sistemlerin istenen görevi kademeli olarak ve artan verilerle gerçekleştirebilmesidir. Makine öğrenimi araştırmalarının kapsamı çok geniştir. Teorik olarak, araştırmacılar yeni öğrenme yöntemleri oluşturmaya ve yöntemleri için öğrenmenin uygulanabilirliğini ve kalitesini incelemeye çalışmaktadır ve diğer yandan, bazı araştırmacılar makine öğrenimi yöntemlerini yeni sorunlara uygulamaya çalışmaktadır. Elbette, bu spektrum ayrık değildir ve araştırmalar her iki yaklaşımın bileşenlerine sahiptir. Aşağıda makine öğrenimi süreçleri açıklanmaktadır. Makine öğrenimi süreçlerini analiz etme adımları arasında veri toplama ve hazırlama,

model seçimi ve eğitimi ve süper parametreleri ayarlama ve tahmin etme yer almaktadır (Sammut, Webb, 2011).

1.2.1. Veri Toplama

Makine öğrenmesi sürecindeki ilk adım, makine için gereken bilgi ve verileri sağlamaktır. Bu veriler iki gruba ayrılır. İlk grup sistemi eğitmek için kullanılır ve ikinci grup sistemi test etmek için kullanılır. Seçilen verilerin tüm popülasyonu temsil ettiği unutulmamalıdır. Veriler genellikle 20/80 veya 70/30'a bölünür, böylece model yeterli eğitimden sonra daha sonra test edilebilir (Dietterich, 1997, s. 97).

1.2.2. Model Seçimi ve Eğitimi

Makine öğrenmesi prensiplerinde ikinci ve onu takip eden adımlar, model seçimi ve modelin eğitilmesidir. Belirli bir problem türünü çözmek amacıyla geliştirilmiş ve uyarlanmış çeşitli makine öğrenmesi algoritmaları ve modelleri bulunmaktadır. Bu nedenle, seçilecek model, problemin doğasına ve çözüm ihtiyacına uygun olarak belirlenir ve eğitilir (Nahr vd., 2021).

1.2.3. Model Değerlendirmesi

Makine, kendisine öğretilen verilerden farklı desenler ve özellikler öğrenir ve yeni verileri tanımlama, sınıflandırma veya tahmin etme gibi çeşitli alanlarda kararlar almak üzere kendini eğitir. Tahminler, makinenin bu kararları nasıl alabildiğini doğru bir şekilde incelemek için eğitilmiş veriler üzerinde test edilmelidir. Bunu yapmak için, önce eğitilmiş veriler üzerinde faaliyetler gerçekleştirilir ve model eğitildikten sonra, ne kadar doğru olduğunu belirlemek için veri tabanlı test için kullanılır (Burkov, 2019)

1.2.4. Veri Ön İşleme

Ön işleme adımı, genellikle veri temizliğiyle ilgili bir dizi faaliyeti içerir. Bu faaliyetler; bir şema tanımlanması, eksik değerlerin doldurulması, verilerin daha düzenli ve sade bir forma dönüştürülmesi ile ham verilerden daha kullanılabilir bir biçime eşlenmesini kapsamaktadır. Bu tür veri manipülasyonlarını gerçekleştirme

yöntemleri, bu çalışmanın kapsamı dışında kalan ayrı bir araştırma alanı olarak değerlendirilmektedir. Ön işlemenin bir parçası olarak ele alınabilecek, nispeten daha az bilinen ancak oldukça önemli bir diğer sorun ise veri dağılımıdır. Çoğu durumda, farklı şemalara, kurallara ve veri depolama/erişim yöntemlerine sahip birden fazla, ancak ilişkili veri kaynağı bulunabilmektedir. Bu veri kaynaklarını makine öğrenmesi için uygun tek bir veri kümesinde birleştirmek, veri bütünleştirme olarak bilinen ve başlı başına karmaşık bir süreç olan bir görevdir.

Bu duruma örnek olarak Madaio ve çalışma arkadaşlarının (2016) Atlanta İtfaiye Departmanı için geliştirdiği danışmanlık sistemi Firebird gösterilebilir. Bu sistem, yangın denetimlerinde öncelikli hedefleri belirlemeye yardımcı olmayı amaçlamaktadır. Firebird'ün geliştirilmesine yönelik ilk adımda, yangın olaylarının geçmişi, işletme ruhsatları ve hane bilgileri gibi çeşitli alanlara ait toplam 12 veri kümesinden veriler toplanmıştır. Bu veri kümeleri, itfaiye departmanının gözetimindeki her bir mülkle ilgili tüm bilgileri kapsayan birleşik bir veri kümesine dönüştürülmüştür. Yazarlar özellikle veri birleştirmenin zorluklarına dikkat çekmişlerdir. Binaların coğrafi konumlarına göre tanımlanabileceği göz önünde bulundurulduğunda, her bir veri kümesinin mekansal veriler (adres bilgileri gibi) içerdiği ifade edilmiştir. Ancak bu mekansal bilgiler, farklı formatlarda sunulabildiği gibi, adres yazımındaki küçük farklar gibi detaylar nedeniyle veri eşleştirmeyi zorlaştırabilmektedir.

1.2.5. Süperparametrelerin Tahmini ve Ayarlanması

Makine öğrenimi bağlamında süperparametreler, modelin kendisinin tahmin edemediği ancak modelin performansını artırmak için gerekli olan değişkenlerdir. Bir tanım vermek istersek, makine öğrenimindeki süperparametreler algoritmayı çalıştırmak için kullanıcı tarafından belirtilmesi gereken parametrelerdir. Klasik parametreler veriler tarafından öğretilirken, süperparametreler verilerden öğrenilir. Bir model çok sayıda süperparametreye sahip olabilir ve süperparametrelerin mümkün olan en iyi kombinasyonunu seçme sürecine süperparametre ayarı denir. Hiperparametreleri ayarlamak için bazı temel yöntemler arasında ağ araması, rastgele arama veya gradyan tabanlı optimizasyon bulunur. Süperparametreleri optimize etme süreci tamamlandıktan sonra, makine öğrenimi modelinin oluşturulduğu ve başarı

oranına veya doğruluğuna, tahmin yeteneğine bağlı olarak gerçek dünyada uygulanabileceği söylenebilir. Bu nedenle, yukarıdaki yöntemlerin yardımıyla bir makine öğrenimi algoritması oluşturulabilir (Rose, 2016).

1.2.6. Gözetimli Öğrenme

"Bu tür öğrenme, etiketli verilerle çalışır; çünkü algoritmanın öğrenme süreci (eğitim verileri aracılığıyla) bir gözlemci tarafından denetlenir. Bu süreçte 'öğrenme algoritması', eğitim verilerinden elde edilen çıktıları çeşitli yinelemeler boyunca tahmin eder; ardından bu çıktılar gözlemci tarafından düzeltilir. Algoritma kabul edilebilir bir performans düzeyine ulaştığında ise öğrenme süreci sona erer. Gözetimli öğrenme, sistemin eğitilebilmesi için bir dizi giriş verisi gerektirir. Bu öğrenme yöntemi kendi içinde iki alt kategoriye ayrılır: regresyon ve sınıflandırma. Regresyon, çıktının sürekli bir sayı ya da sayı kümesi olduğu problemlere yöneliktir. Örneğin, evin alanı, oda sayısı gibi özelliklere dayanarak ev fiyatlarının tahmin edilmesi regresyon problemlerine örnek olarak verilebilir. Sınıflandırma, çıktısı bir kümenin parçası olan sorunları ifade eder, örneğin bir e-postanın spam olup olmadığını tahmin etmek veya bir kişinin on hastalıktan hangisinde hastalık olduğunu tahmin etmektir (Caruana vd., 2006).

1.2.7. Gözetimsiz Öğrenme

Gözetimsiz öğrenme ile gözetimli öğrenme bir tür makine öğrenme yöntemidir. Öğrenme etiketlenmemiş veriler üzerinde ve bu verilerdeki gizli kalıpları bulmak için yapılır, öğrenme gözetimsiz olacaktır. Aslında, gözetimli öğrenmenin aksine, gözetimsiz öğrenme, etiketli sonuçlara ihtiyaç duymadan kalıpları çıkarmak ve modellemek için girdi verilerini kullanır. Bu öğrenme kategorisinde yalnızca etiketlenmemiş ham verilerimiz vardır. Aslında, bu öğrenme sürecinde algoritmaya rehberlik edecek bir gözlemci bulunmamaktadır (Nouri vd., 2019). Gözetimsiz öğrenmenin temel amacı, verilerin dağılımını modelleyerek veriler hakkında daha fazla bilgi elde etmektir. Bu öğrenme modeli, etiketlenmemiş verilerdeki gizli yapıları ortaya çıkarmayı hedefler. Gözetimsiz öğrenme, kümeleme ve boyut indirgeme gibi iki ana kategoriye ayrılır (Barlow, 1989).

1.2.8. Güçlendirme Öğrenimi

Makine öğrenmesinin bir alt dalı olan pekiştirmeli öğrenmenin (reinforcement learning) temel amacı, yazılım ajanlarının, içinde bulundukları ortama en uygun eylemleri seçerek ödülü maksimize etmeleridir. Bu öğrenme türü, davranışçı psikolojiden esinlenmiş olup, bir ajanın maksimum ödül elde etmek amacıyla hangi davranışları sergilemesi gerektiğine odaklanır ve öğrenme sürecinde ödül-mükafat yoluyla en iyi stratejiyi keşfetmesini sağlar. Pekiştirmeli öğrenme; oyun teorisi, kontrol teorisi, yöneylem araştırması, bilgi teorisi, çoklu ajan sistemleri, sürü zekası, istatistik, genetik algoritmalar ve simülasyona dayalı optimizasyon gibi çok çeşitli disiplinlerde incelenmektedir. Bu öğrenme türünün, denetimli öğrenmeden iki temel farkı bulunmaktadır: Birincisi, girdi ve çıktılar açık bir şekilde tanımlı değildir; ikincisi ise, yeni bilgileri keşfetme (exploration) ve mevcut bilgileri kullanma (exploitation) arasında uygun bir denge kurmayı gerektiren çevrimiçi ve canlı öğrenme sürecine güçlü bir vurgu yapmasıdır. Pekiştirmeli öğrenme; üretim, otomotiv, işletme yönetimi, bilgisayar sistemleri, makine görüşü, eğitim, enerji, finans, oyun teknolojileri, sağlık hizmetleri, doğal dil işleme (NLP), robotik, bilim, mühendislik, sanat ve ulaşım gibi birçok farklı alanda kullanılmaktadır. Örneğin, Google'ın çeviri hizmetinde gelişmiş öğrenme yöntemleri kullanılmaktadır (Mehrabian vd., 2021, s. 16).

1.3. Makine Öğrenimini Uygulamaları

Makine öğrenmesi, sağlık, finans, görüntü işleme ve ses tanıma gibi birçok farklı alanda veriden öğrenerek pratik çözümler sunmaktadır.

1.3.1. Görüntü İşleme

Makine öğreniminin en yaygın uygulamalarından biri görüntü tanımadır. Dijital görüntülerde nesneleri kategorize etmek için birçok fırsat vardır ve makine öğrenimi bunu yapmak için kullanılabilir. Örneğin, siyah beyaz görüntülerde her piksel bir ölçü birimi olarak kullanılır. Renkli görüntülerde her piksel, kırmızı, yeşil ve mavi olmak üzere üç rengin yoğunluğu için bir ölçü birimi kullanır. Makine öğrenimi ayrıca görüntü işlemede yüzleri tanımlamak için de kullanılabilir. Bir veritabanında her kişi için ayrı bir kategori vardır ve makine öğrenimi algoritmaları

tanımlamak için bu görüntüleri kullanır. Makine öğrenimi ayrıca sıradan yazılarda veya basılı el yazmalarında el yazısını tanımak için de kullanılır.

1.3.2. Konuşma Tanıma

Makine öğreniminin uygulamalarından biri de konuşmayı metne dönüştürmedir. Bu, otomatik konuşma tanıma ile bilgisayar konuşma tanıma olarak da bilinir. Konuşma tanıma yardımıyla bir yazılım bir konuşmadaki kelimeleri tanıyabilir ve bunu bir metin dosyasına dönüştürebilir. Buradaki ölçü birimi, konuşma sinyallerini simgeleyen bir sayı zinciri olabilir. Konuşma sinyalleri ayrıca farklı zaman frekans bantlarına güçlü bir şekilde bölünebilir. Konuşma tanıma, etkileşimli bir ses arayüzü veya sesli arama yeteneği vb. olan uygulamalarda kullanılabilir.

1.3.3. Hastalığın Teşhisi

Makine öğrenimi, hastalıkları teşhis etmek için kullanılan tekniklerde ve araçlarda kullanılabilir. Bu teknoloji, klinik parametreleri analiz etmek ve bunları birleştirerek hastalığın ilerlemesini tahmin etmek, tıbbi bilgileri çıkarmak, sonuçlara ulaşmak için araştırma yapmak, tedavi planlaması yapmak ve hasta takibi yapmak için kullanılabilir. Bunlar, makine öğrenimi yöntemlerinin başarılı uygulamalarıdır. Makine öğrenimi algoritmalarının kullanımı ayrıca bilgisayar sistemlerini ve sağlık bölümlerini entegre etmeye yardımcı olabilir.

1.3.4. İstatistiksel Analiz

Ekonomide en önemli konulardan biri, menkul kıymet alım satımı için kısa vadeli stratejiler elde etmektir. Bu stratejileri elde etmek için kullanıcı, tarihsel korelasyonlar ve merkezi genel ekonomik değişkenler gibi faktörlere dayalı olarak menkul kıymet alım satımı yapmak için ticaret algoritmalarını kullanır. Makine öğrenimi, bu kısa vadeli strateji algoritmalarını elde etmek için çok yararlı olabilir.

1.3.5. Finansal Hizmetler

Makine öğreniminin finans ve bankacılıkta kullanım potansiyeli çoktur. Makine öğrenimi ve yapay zekanın kullanılması, popüler finansal hizmetler sağlayabilir. Makine öğrenimi, bankaların ve finans kuruluşlarının daha bilinçli kararlar almasına yardımcı olabilir, bir hesap kapanmasını gerçekleşmeden önce tespit etmek için finansal hizmetler sağlayabilir, müşteri maliyet modellerini izleyebilir, piyasa analizi yapabilir, modelleri izleyebilir, akıllı makinelere maliyeti öğretebilir ve son olarak makine öğrenimi algoritmaları, eğilimleri ve eğilimleri önceden kolayca belirleyebilir ve gerçek zamanlı olarak tepki verebilir (Abdullah, 2021).

1.4. Günümüzde Makine Öğrenimi

Aşağıda farklı endüstrilerde makine öğreniminin farklı uygulamalarından bazıları belirtilmiştir. Günümüz dünyasındaki uygulama hacmi nedeniyle, yalnızca birkaç özel durum yeterli olacaktır.

1.4.1. Finansal Hizmetler

Rutgers Üniversitesi Sanat ve Sanat Zekası Laboratuvarı'ndaki araştırmacılar, bir bilgisayar algoritmasının resimleri, bu eserlerin stiline, türüne ve sanatçılara göre kolayca tanımlayıp sınıflandırabileceğini görmek istediler. Araştırmacılar bunu, resim stillerini yapay zekaya sınıflandırmak için görsel özellikleri tanımlayıp öğrenerek yaptılar. Gelişmiş makine öğrenimi algoritmaları ve yapay zekaları, resim stillerini veritabanında %60 doğrulukla kategorilere ayırarak, eserleri tanıma ve kategorilere ayırma konusunda sıradan uzman olmayanları geride bırakıyor ve daha iyi performans gösteriyor. Araştırmacılar, stil sınıflandırması için görsel özelliklerin (yani, denetlenen öğrenmeyle ilgili bir sorun) sanatsal etkileri belirlemek için de kullanılabileceğini varsaydılar (denetimsiz öğrenmeyle ilgili bir sorun). Daha sonra, belirli nesneleri tanımlamak için Google'da bulunan nesneleri sınıflandırmak ve eğitmek için algoritmalar kullandılar. Bu algoritmaları 550 yıldan uzun süredir çalışan 66 farklı sanatçının 1.700'den fazla resminde test ettiler.

1.4.2. Büyük Binalarda HVAC Enerji Tüketiminin Optimizasyonu

Ofis binaları, hastaneler ve diğer büyük ticari binalardaki ısıtma, havalandırma ve klima (HVAC) sistemleri genellikle iklim değişikliği modellerini, değişken enerji maliyetlerini veya binanın termal özelliklerini hesaba katmadıkları için verimsizdir. BuildingIQ bulut tabanlı yazılım platformu bu sorunu kolayca çözebilir. Platform, elektrik sayacı verileri, termometreler ve HVAC basınç sensörleri ile hava ve enerji maliyetleri gibi büyük miktarda gigabayt bilgiyi sürekli olarak işlemek ve işlemek için gelişmiş algoritmalar ve makine öğrenimi yöntemleri kullanır. Özellikle, makine öğrenimi verileri bölmek ve ısıtma ve soğutma süreçlerinde gaz, elektrik, buhar ve güneş enerjisinin göreceli payını belirlemek için kullanılır. BuildingIQ platformu, büyük ölçekli ticari binalarda normal zamanlarda HVAC enerji tüketimini %10 ila %25 oranında azaltıyordur. (Nozari vd., 2019).

1.4.3. Otonom Arabalar

Waymo, Google Otomatik Araba Projesi'nin bir koludur. Bu projenin amacı, insan sürücüsü olmadan kendi kendine gidebilen arabalar üretmektir. Bunu yapmak için Waymo filosunun yapay zekanın ciddi yardımına ihtiyacı vardır. Waymo arabaları çevrelerini görmek, onlara anlam vermek ve başkalarının nasıl davranacağını tahmin etmek için makine öğrenimini kullanır. Yol boyunca davranışlarını değiştiren birçok değişikliğe rağmen, gelişmiş bir makine öğrenimi sistemi başarı için olmazsa olmazdır (Scanlon vd., 2021).

1.4.4. Ulaşımveri Analizi

Makine öğreniminin ulaşım ve taşımacılık endüstrilerinde, mal türü, popüler varış noktaları ve menşe büyük şehirler gibi ulaşım alanındaki kalıpları belirlemek ve çeşitli eğilimleri bulmak için temel uygulamalarından biridir. Makine öğrenimini kullanarak, malların şehre girip çıkması için en uygun rotaları belirlemek ve farklı rotalardaki olası sorunları bulmak ve bunların ortaya çıkmasını önlemek mümkündür. Tüm bunlar, ulaşım endüstrisi, toptancılar, perakendeciler ve müşteriler için karları artırır. Ayrıca, şehir içi ve şehirler arası ulaşımın güvenliği iyileştirilecektir. Bugün, farklı sorunları modellemek için farklı ayrıştırma verileri ve farklı algoritmalar

kullanılmaktadır. İnsanlar ve makineler tarafından üretilen veri miktarı o kadar büyüktür ki, bu verilere dayalı emilim, yorumlama ve karmaşık kararlar insan yeteneklerinin ötesine geçer. Yapay zeka, tüm bilgisayar öğreniminin temelini oluşturur ve karmaşık karar almanın geleceğidir. Bu kombinasyonları ve yer değiştirmeleri hesaplamada bilgisayarların kullanımı, en iyi kararı elde etmek için çok yararlıdır. Yapay zeka (ve makine öğrenimindeki evrimi) ve derin öğrenme, iş ve diğer birçok alanla ilgili kararlar almanın anahtarıdır. Yapay zeka, bir işletmedeki birçok süreci kendi başına yapabilir, iş gücünü önemli ölçüde azaltabilir ve bir organizasyonun verimliliğini artırabilir, zamandan ve paradan ve diğer birçok kaynaktan tasarruf sağlayabilir. Tüm bunlar AI için, özellikle de işletmeler için önemlidir.

1.5. Makine Öğreniminin Sorunları

Makine öğreniminin ilk dönemleri, görece basit ve yüzeysel yöntemlerle karakterize edilmekteydi. Örneğin, bir karar ağacı algoritması, yalnızca denetleyiciler tarafından öğretilen kurallara sıkı sıkıya bağlı şekilde çalışıyordu: 'Eğer bir nesne oval ve yeşilse, bunun salatalık olma olasılığı P 'dir.' Bu tür modeller, bir görüntüdeki salatalığı tanımada yeterince başarılı olmasalar da, en azından algoritmaların nasıl çalıştığı herkes tarafından anlaşılabilirdi. Buna karşılık, derin öğrenme algoritmaları farklı bir yaklaşım benimser. Bu algoritmalar, verilerin hiyerarşik temsillerini oluşturarak, kendi anlayış yapılarını geliştirmelerine olanak tanıyan çok katmanlı mimariler kullanırlar. Büyük eğitim verisi kümelerini analiz ettikten sonra, sinir ağları salatalıkları şaşırtıcı bir doğrulukla nasıl tanıyacaklarını öğrenebilirler. Sorun, denetleyicilerinin - makine öğrenimi mühendisleri veya veri bilimcileri - bunu tam olarak nasıl yaptıklarını bilmemeleridir. Soruna kara kutu denir. Yapay Zeka denetleyicileri girdiyi (makine öğrenimi algoritmasının analiz ettiği eğitim verileri) ve çıktıyı (verdiği karar) anlar. Mühendisler, tek bir tahminin nasıl yapıldığını anlayabilseler de, tüm modelin nasıl çalıştığını kavramak oldukça zordur. Bazı yapay zeka araştırmacıları, makine öğreniminin son zamanlarda yeni bir 'simya' biçimi haline geldiğini ve tüm alanın bir kara kutuya dönüştüğünü savunmaktadır. Bu görüş, Google'dan Ali Rahimi'nin de dile getirdiği bir düşüncedir. Tıp, sürücüsüz arabalar veya kredi notunun otomatik değerlendirilmesi gibi diğer yapay zeka uygulamalarının

geliştirilmesinde önemli bir engel teşkil etmektedir. Kara kutu problemi, uygulama içi öneri hizmetleri gibi makine öğrenimi alanlarında da önemli bir zorluk olarak karşımıza çıkmaktadır. Kara kutu, uygulama içi öneri hizmetleri için makine öğreniminde bir sorundur. Web uygulaması kullanıcılarının otomatik önerilerin nasıl çalıştığını az çok bildiklerinde kendilerini daha rahat hissettikleri ortaya çıktı. Bu nedenle Netflix gibi birçok büyük veri şirketi ticari sırlarından bazılarını ifşa ediyor. Araştırma, yapay zekanın genellikle insanlarda korku ve diğer olumsuz duygulara neden olduğunu gösteriyor. İnsanlar bir nesnenin "neredeyse bir insan gibi" görünmesinden ve davranmasından korkuyor. Bu fenomene "ürkütücü vadi" adı veriliyor. Makinelerin, makineler gibi davrandığı kabul edilirken, konuşma, gülümseme, şarkı söyleme veya resim yapma gibi insanlara özgü işlevleri yerine getiren makineler genellikle kabul edilmemektedir. Elbette, yeni nesiller robotlar ve makine öğrenimi algoritmalarıyla etkileşime girdikleri dijital bir ortamda büyüdükçe bu durum zamanla değişebilir (Nozari vd., 2019).

1.5.1. Yetenek Açığı

"Makine öğrenimi sektörüne ilgi duyan birçok kişi bulunmasına rağmen, bu teknolojiyi geliştirebilecek uzman sayısı oldukça sınırlıdır, bu da günümüzde yaygın bir makine öğrenimi sorunu olarak karşımıza çıkmaktadır. Makine öğrenimini iyi anlayan bir veri bilimcisinin, yazılım mühendisliği konusundaki bilgi seviyesi neredeyse hiç yoktur, bu da ciddi bir sorundur. Çinli teknoloji devi Tencent, 2017 yılı sonlarında dünya çapında yapay zeka ile ilgilenen yaklaşık 300.000 araştırmacı ve uygulayıcı olduğunu tahmin etmiştir. Bağımsız bir şirket olan Element AI, "10.000'den az kişinin ciddi yapay zeka araştırmalarıyla başa çıkmak için gerekli ML becerilerine sahip olduğunu" tahmin ediyor. Makine öğrenimi mühendisleri ve veri bilimcileri, Google, Amazon, Microsoft veya Facebook gibi en önemli oyuncular için en önemli öncelikli işe alımlardır. Yapay zeka alanındaki maaşları fırlatır, ancak aynı zamanda piyasada bulunan uzmanların ortalama kalitesini de düşürür. Makine öğrenimiyle sorun çok daha kötü görünüyor. DataCamp'e göre, 2023'te ML mühendislerinin ortalama maaşı 160.140 dolara yükselirken, en iyiler NBA süperstarları kadar maaş alacak. Google, 2016'da yaptığı bir mahkeme dosyasında, otonom araç bölümünün

liderlerinden birinin Google'ın rakibi Uber'e gitmeden önce 120 milyon dolar teşvik kazandığını açıkladı. (DataCamp, 2023).

Yukarıda bahsettiğim gibi, bir makine öğrenimi modelini eğitmek için büyük veri kümelerine ihtiyacınız vardır. Herkes petabaytlarca bilgiyi depolayıp işleyebildiği için artık sorun değilmiş gibi görünebilir. Depolama ucuz olsa da yeterli miktarda eğitim verisi toplamak zaman alır. Dahası, hazır veri kümeleri satın almak pahalıdır. Makine öğreniminde farklı nitelikte sorunlar vardır. Algoritma eğitimi için veri hazırlamak karmaşık bir süreçtir. Makine öğrenimi algoritmanızın hangi sorunu çözmesini istediğinizi bilmeniz gerekir çünkü sınıflandırmayı, kümelemeyi, regresyonu ve sıralamayı önceden planlamanız gerekecektir. Veri toplama mekanizmaları ve tutarlı biçimlendirme oluşturmanız gerekir (Jordan vd., 2015). Daha sonra verileri öznitelik örnekleme, kayıt örnekleme veya toplama ile azaltmanız gerekir. Eğitim verilerini ayrıştırmanız ve yeniden ölçeklendirmeniz gerekir. Bu, yetenekli mühendisler ve zaman gerektiren karmaşık bir görevdir. Bu nedenle sonsuz disk alanınız olsa bile süreç pahalıdır. Kişisel verileri kullanmayı planlıyorsanız, muhtemelen ek zorluklarla karşılaşacaksınız. Dünya çapında insanlar gizliliklerini korumanın önemini giderek daha fazla farkına varıyor. Bunu yaptığınızı fark ettiklerinde, tüm izinlerini almış olsanız bile, bunları sizinle paylaşmak veya resmi bir şikayette bulunmak istemeyebilirler. Ünlü Avrupa Genel Veri Koruma Yönetmeliği gibi kişisel verileri koruyan yeni düzenlemelerin getirilmesiyle kişisel veriler ve büyük veri faaliyetleri de daha zor, riskli ve maliyetli hale geldi.

En büyük teknoloji şirketleri herkes için açık kaynaklı çerçevelere para harcıyor. Alphabet Inc. (eski Google) TensorFlow'u sunarken, Microsoft Facebook ile işbirliği yaparak Açık Sinir Ağı Değişimi'ni (ONNX) geliştiriyor. Bu sistemler, dünyanın dört bir yanındaki işletmeler ve bireysel kullanıcılar tarafından sağlanan verilerle destekleniyor. Ancak, makine öğreniminin temel sorunu, tüm bu ortamların çok genç olmasıdır. TensorFlow'un ilk sürümü Şubat 2017'de yayınlanırken, bir diğer popüler kütüphane olan PyTorch Ekim 2017'de çıktı. Web uygulama çerçeveleri çok daha eskidir - Ruby on Rails 20 yaşında ve Python tabanlı Django 27 yaşında. Bir yandan genç teknoloji en çağdaş çözümleri kullanırken, diğer yandan üretime hazır olmayabilir veya üretime hazır olmayabilir. Herhangi bir tatmin edici sonuç elde etmek için zamana ihtiyacınız var ve planlama zordur. Geleneksel kurumsal yazılım

geliştirme oldukça basittir. İş hedefleriniz, işlevleriniz var, bunları oluşturmak için teknolojiyi seçiyorsunuz ve çalışan bir sürümü yayınlamanın birkaç ay süreceğini varsayıyorsunuz. Makine öğrenimi geliştirmesinde daha fazla katman vardır. Mühendisler, iş hedeflerinizi belirlerken planladığınız eylemleri gerçekleştirmeyi öğrenecek bir program üretecek bir program yazıyorlar. Sadece bu bir veya iki seviyeyi eklemek her şeyi çok daha karmaşık hale getiriyor. Zorluk, makine öğreniminin çok daha fazla zaman almasıdır. Eğitim verilerini toplayıp hazırlamanız, ardından algoritmayı eğitmeniz gerekir. Makine öğrenimi projelerinde, geleneksel web sitesi veya uygulama geliştirme projelerine kıyasla çok daha fazla belirsizlik bulunmaktadır. Bu nedenle, deneyimli bir ekip geleneksel projelerde zaman tahminlerini oldukça hassas bir şekilde yapabilirken, ürün önerileri sağlamak için kullanılan bir makine öğrenimi projesi, beklenenden çok daha kısa veya uzun bir süre alabilir. Bunun nedeni, en deneyimli makine öğrenimi mühendislerinin bile derin öğrenme ağlarının farklı veri kümelerini analiz ederken nasıl davranacağını kesin olarak öngörememeleridir. Bu ayrıca makine öğrenimi mühendislerinin ve veri bilimcilerinin bir modelin eğitim sürecinin tekrarlanabileceğini garanti edemeyeceği anlamına gelir. Makine öğrenimi zorlukları alanında, zamanın ve planlamanın karmaşık manzarasında gezinmek, veri bilimi ve makine öğrenimi uygulamaları alanına giren işletmeler için zorlu bir engel teşkil eder.

Makine öğrenimi süreci, verileri analiz etmekten ve uygun makine öğrenimi tekniklerini seçmekten model eğitime ve dağıtımına kadar bir dizi kritik aşamayı içerir. Bunu ele almak için, bu aşamaları parçalara ayırmak ve her biri için gerçekçi zaman tahminleri sunmak, beklentilerin makine öğrenimi yolculuğunun karmaşıklıklarıyla uyumlu olmasını sağlamak esastır. Makine öğrenimi geliştirmeye göre uyarlanmış proje yönetimi metodolojileri, veri odaklı projelerin dinamik yapısına uyum sağlamak için esnekliği vurgulayarak önemli bir rol oynar. Dahası, makine öğrenimi yaşam döngüsü boyunca sürekli test ve doğrulamanın önemini vurgulamak, öngörülemeyen zorluklara karşı koruma sağlamak ve çeşitli makine öğrenimi uygulamalarında kullanılan modellerin sağlamlığını sağlamak için zorunlu hale gelir.

1.5.2. Veri Kalitesi ve Niceliği

Veri kalitesi ve niceliği alanında da büyük bir zorluk vardır. Yüksek kaliteli ve yeterli eğitim verilerine ulaşmak, gürültülü verilerin varlığı ve düşük kaliteli bir veri setinin potansiyel tuzakları ile işaretlenen zorlu bir görevdir. Makine öğrenimi modellerinin güvenilirliği ve doğruluğu, eğitildikleri verilerin kalitesine bağlıdır ve aykırı değerler, eksik değerler ve yanlışlıklar gibi sorunlar arasında gezinmek çok önemlidir. Belirli sınıfların veya sonuçların yeterince temsil edilmediği dengesiz veri setlerini ele alma stratejileri dikkatlice düşünülmelidir. Ayrıca, eğitim verilerindeki önyargıları ele almak, özellikle çeşitli kullanıcı gruplarını etkileyen uygulamalarda adil ve eşit model tahminlerini sağlamak için çok önemlidir. İşletmeler makine öğrenimi yolculuklarına çıktıkça, bu zorluklarla başa çıkmak, eğitim veri setinin hem kalitesini hem de niceliğini artırmak için veri temizleme, ön işleme ve artırma tekniklerini kapsayan stratejik bir veri düzenleme yaklaşımı gerektirir.

1.5.3. Model Yorumlanabilirliği

Makine öğreniminin geniş alanında, modellerin yorumlanabilirliğini sağlamak, özellikle sağlık ve finans gibi hassas sektörlerde önemli bir zorluk teşkil eder. Karar alma süreçlerinde doğru ve anlaşılır içgörülere duyulan ihtiyaç göz önüne alındığında, yorumlanabilir modellerin önemi yeterince vurgulanamaz. Karmaşık matematiksel hesaplamalar ve gelişmiş algoritmalar kullanmak kesin sonuçlar üretebilse de, model doğruluğu ve yorumlanabilirlik arasında bir denge kurar. Özellikle şeffaflık, kullanıcılardan ve paydaşlardan güven ve anlayış kazanmak için hayati önem taşıdığına, doğru dengeyi sağlamak çok önemlidir. Eğitim modelleri arayışında, karmaşık, doğrusal olmayan bir model ile daha basit bir doğrusal model arasındaki seçim dikkatli bir değerlendirme gerektirir. Makine öğreniminin sürekli gelişen ortamındaki işletmeler, dağıtılan modellerin yalnızca doğruluk sağlamakla kalmayıp aynı zamanda şeffaflık ve anlaşılabilirliği koruyarak en hızlı büyüyen ve teknolojik olarak gelişmiş alanlardan birine güveni teşvik etmek için bu faktörleri titizlikle tartmalıdır.

1.5.4. Ölçeklenebilirlik Sorunları

Makine öğreniminin uçsuz bucaksız dünyasında, özellikle büyük veri kümeleri veya karmaşık veri yapıları söz konusu olduğunda ölçeklenebilirlik söz konusu olduğunda önemli bir zorluk ortaya çıkar. Bu bilgi zenginliğini ele almak için makine öğrenimi modellerini ölçeklendirmek, içsel engelleri aşmak için düşünceli bir strateji gerektirir. Dağıtılmış bilgi işlem ve paralel işleme gibi ölçeklenebilir çözümlere ilişkin derin bir anlayış olmazsa olmaz hale gelir. Bu teknolojiler, kapsamlı veri kümelerini ve karmaşık matematiksel hesaplamaları verimli bir şekilde yönetmede önemli bir rol oynar. Sürekli olarak yeni verilerin ortaya çıktığı sürekli gelişen gerçek dünyada, işletmeler ölçeklenebilir yaklaşımlara duyulan ihtiyaçla boğuşmaktadır. İster robotik eğitim verilerine uyum sağlamak ister hızla ilerleyen teknolojilerin karmaşıklıklarında gezinmek olsun, ölçeklenebilir çözümlerin nüanslı bir şekilde anlaşılması ve uygulanması çok önemlidir. En hızlı büyüyen teknolojilerin şekillendirdiği bir ortamda çevik ve etkili kalmak, makine öğreniminin günümüzün veri açısından zengin ortamlarının dinamik talepleriyle sorunsuz bir şekilde bütünleşmesini sağlamaktır.

1.5.5. Düzenleyici Uyumluluk

Makine öğrenimi zorlukları alanında, düzenleyici uyumluluğun üstesinden gelmek, özellikle katı standartlarla karakterize edilen endüstrilerde kritik bir engel olarak ortaya çıkıyor. Kaliteli veri ve uyumluluğun bir araya getirilmesi, özellikle tıbbi teşhis veya video gözetimi gibi hassas alanlarla uğraşırken odak noktası haline geliyor. Konuşma tanıma gibi harika bir teknolojinin veya geçmiş ve eksik verilerin karmaşıklıklarının yasal ve etik standartlarla uyumlu olması gereken düzenleyici çerçevelere uymada veri güvenliğinin sağlanması son derece önemlidir. Endüstriler makine öğrenimi uygulamalarına getirilen aşırı gerekliliklerle karşı karşıya kaldıkça, düzenleyici ortamın dikkatli bir şekilde incelenmesi zorunlu hale geliyor. Uyumluluğa ulaşma stratejileri yalnızca yasal ve düzenleyici standartları karşılamayı değil, aynı zamanda olası ihlallere karşı koruma sağlayan sağlam bir çerçeve geliştirmeyi de içerir. Makine öğrenimi endüstrileri şekillendirmeye devam ederken, bu dönüştürücü teknolojinin güvenini ve etik kullanımını teşvik etmek için düzenleyici zorluklara yönelik proaktif yaklaşımlar şarttır.

Makine öğreniminin yakında yaygın bir teknoloji haline gelmesi çok olasıdır ancak ana makine öğrenimi sorunları henüz çözülmemiştir. Yine de, bir AI projesine girmek yüksek riskli, yüksek ödüllü bir girişimdir. Sabırlı olmanız, dikkatli planlama yapmanız, makine öğrenimi teknolojisinin getirdiği zorluklara saygı duymanız ve makine öğrenimini gerçekten anlayan ve size boş vaatler satmaya çalışmayan insanları bulmanız gerekir.



2. YAPAY ZEKADA MAKİNE ÖĞRENMESİNİN ROLÜ

Yapay zeka kapsamında makine öğrenmesinin rolü incelenmiş; siber güvenlikte saldırı tespit sistemleri, NSL-KDD veri seti çerçevesinde geliştirilen yaklaşımlar ve KNIME uygulaması kullanılarak yapılan performans analizleri üzerinden değerlendirmeler yapılmıştır.

2.1. Siber Güvenlikte Saldırı Tespit Sistemleri ve Makine Öğrenimi

Yaklaşımları

Günümüzün milyarlarca insanın internete eriştiği bağlantılı dünyasında, iletişim için internete bağlı olan veya bir bilgisayara veya herhangi bir akıllı cihaza bağlı olan her şey çeşitli siber saldırı türlerinden etkilenebilir. Sonuç olarak, kamu veya özel birçok kuruluş, sürekli ve karmaşık farklı siber saldırı ve siber tehdit türleriyle uğraşmak zorundadır. Siber tehditlerin ve siber saldırıların artık toplumun her alanına nüfuz etmesi, siber güvenliğin neden hayati önem taşıdığını göstermektedir. Siber güvenlik, programları, ağları ve sistemleri herhangi bir siber saldırıdan koruma eylemidir. Bu siber saldırılar esas olarak kritik verilere erişmeyi, bunları değiştirmeyi veya silmeyi; kullanıcılardan para çalmayı; veya olağan iş operasyonlarını durdurmayı amaçlamaktadır. Saldırı Algılama Sistemi (IDS), siber saldırıları ele almak için önemli ve dinamik alanlardan biridir. IDS, bir saldırı girişimini ağına veya sistemin aktivitesini analiz ederek tespit edebilen bir uygulama veya tekniktir, ardından IDS alarmı çalacaktır. Daha ileri adımlara karar verebilen herhangi bir sisteme Saldırı Önleme Sistemi (IPS) adı verilir. IDS'nin rolü, bir bilgisayar veya ağ sisteminde tespit edilebilecek kötü amaçlı ve şüpheli olayları belirleyerek güvenlik seviyesini yükseltmektir. Saldırı tespitinde en popüler çalışma alanlarından biri anormallik tabanlı ve imza tabanlı tespittir. Anormallik tabanlı saldırı tespiti, ağ trafiğinde normal kalıpları takip etmeyen alışılmadık olayların tespit edilmesi durumundan bahseder. Tipik olmayan veya başka bir deyişle anormal olan herhangi bir şeyin kritik olabileceği ve bir dereceye kadar bazı güvenlik olaylarıyla ilişkili olabileceği varsayılır (Srinivas, 2024). İmza tabanlı saldırı tespit görevleri, tespit edilen kalıpların imza olarak adlandırıldığı belirli kalıpları arayarak saldırıları tespit etmektir. IDS'ler iki farklı kategoriye göre tanımlanabilir:

- Ana Bilgisayar Tabanlı Giriş Tespit Sistemleri (HIDS): Bunlar, üzerine kurulduğu sistem eylemlerini tarayabilir ve izleyebilir. HIDS, herhangi bir dosya sisteminin dosyalarının bütünlüğünde herhangi bir değişiklik tespit edebilir ve herhangi bir kötü amaçlı veya şüpheli etkinlik aramak için günlük dosyalarını analiz edebilir.
- Ağ Tabanlı Giriş Tespit Sistemleri: Bunlar, ağ altyapısını taramaya ve izlemeye odaklanır. Ağın paket akışını analiz eder ve ağda olası bir saldırıyı tespit etmek için paketlerin başlıklarını ve içeriklerini inceler.
- Daha önce belirtildiği gibi, her iki IDS türü tarafından verileri analiz etmek için iki farklı eylem uygulanır:
- İmza Tabanlı Giriş Tespit Sistemleri: Bu tür sistemler, bilinen saldırıları, bunları depolanan kalıplarla karşılaştırarak tespit edebilir, ancak yeni saldırıları tanıyamaz. Bu nedenle tespit, bilinen saldırıların imzalarına ve bir sistem yöneticisi tarafından belirlenen kurallara dayanacaktır.
- Anomali Tabanlı Saldırı Algılama Sistemleri: Bu tür sistemler, normal sistem aktivitesine dayalı bir model oluşturabilir ve daha sonra gözlemlenen aktivitenin bir anomali olup olmadığını belirlemek için oluşturulan modeli kullanarak gözlemlenen aktiviteyi değerlendirebilir.

Araştırmacılar tarafından ağ trafiğinde anomali algılamayı uygulamak için farklı sınıflardan birçok algoritma ve yöntem kullanılır. Bazı teknikler, istatistiklerin gözlemlenen vakanın bir anomali olup olmadığını hesaplamak ve belirlemek için kullanıldığı istatistiksel bir bakış açısına dayalı olarak uygulanır (Kehe vd., 2011). Diğer teknikler, bir anomaliyi tespit etmek için uygulanan makine öğrenimi algoritmalarına dayanır. Makine öğrenimi yöntemleri üç kategoriye ayrılabilir:

- Gözetimli öğrenme: Bir eğitim kümesi etiketli örneklerle sahiptir ve bir algoritma yeni bir gözlemi yalnızca bir sınıfla eşleştirecektir.
- Gözetimsiz öğrenme: Eğitim kümesinde olası bir grup hakkında etiketler veya herhangi bir ayrıntı yoktur. Eğitim sırasında algoritma grupları ayarlar ve benzerlik düzeyini belirler.
- Yarı-denetimli öğrenme: Eğitim seti, çok sayıda etiketsiz örneklerle birlikte az miktarda etiketli örnek içerir. Başka bir deyişle, yarı-denetimli öğrenme, eğitim

verilerinin etiketlendiği denetlenen öğrenme ile eğitim verilerinin etiketlenmediği denetlenmeyen öğrenme arasında gerçekleşir.

İlk iki kategorideki algoritmalar bir ağ anormalliği tespit probleminde uygulanmıştır. Saldırıların anormal olaylar oldukları ve algoritma modeli tarafından sınıflandırılacakları için tespit edilebileceği beklenmektedir. Modeller, makine öğrenimi algoritmaları tarafından ağ trafiğini analiz etmek için kullanılır. Yapay zeka ve siber güvenlik, düşündüğümüzden çok daha yakın ve devrim niteliğinde olarak lanse edildi. Ancak bu, ihtiyatlı beklentilerle yaklaşılması gereken yalnızca kısmi bir gerçektir. Gerçek şu ki, önümüzdeki gelecek için nispeten kademeli iyileştirmelerle karşı karşıya kalabiliriz. Perspektif olarak, tamamen otonom bir gelecekle karşılaştırıldığında kademeli görünebilecek şey, aslında geçmişte yapabildiklerimizin çok ötesindedir.

İnsan hatası, siber güvenlik zayıflıklarının önemli bir parçasıdır. Örneğin, büyük Bilgi teknoloji (BT) ekipleri kurulumla uğraşsa bile, uygun sistem yapılandırmasını yönetmek inanılmaz derecede zor olabilir. Sürekli yenilik sürecinde, bilgisayar güvenliği her zamankinden daha katmanlı hale geldi. Duyarlı araçlar, ekiplerin ağ sistemleri değiştirilirken, değiştirilirken ve güncellenirken ortaya çıkan sorunları bulmasına ve hafifletmesine yardımcı olabilir. Bulut bilişim gibi daha yeni internet altyapısının eski yerel çerçevelerin üzerine nasıl yerleştirilebileceğini düşünün. Kurumsal sistemlerde, bir bilgi teknoloji ekibinin bu sistemleri güvence altına almak için uyumluluğu sağlaması gerekir. Yapılandırma güvenliğini değerlendirmek için manuel süreçler, ekiplerin bitmeyen güncellemeleri normal günlük destek görevleriyle dengelemeleri nedeniyle yorgun hissetmelerine neden olur. Akıllı, uyarlanabilir otomasyonla ekipler yeni keşfedilen sorunlar hakkında zamanında tavsiye alabilirler. Devam etme seçenekleri hakkında tavsiye alabilirler veya hatta gerektiğinde ayarları otomatik olarak ayarlamak için sistemler kurabilirler. İnsan verimliliği, siber güvenlik sektöründeki bir diğer sorun noktasıdır. Hiçbir manuel süreç her seferinde mükemmel şekilde tekrarlanamaz, özellikle de bizimki gibi dinamik bir ortamda. Bir kuruluşun birçok uç nokta makinesinin ayrı ayrı kurulumu en çok zaman alan görevlerden biridir. İlk kurulumdan sonra bile, bilgi teknoloji ekipleri kendilerini daha sonra yanlış yapılandırmaları veya uzaktan güncellemelerde yama yapılamayan eski kurulumları düzeltmek için aynı makineleri tekrar ziyaret ederken bulurlar.

Dahası, çalışanlara tehditlere yanıt verme görevi verildiğinde, söz konusu tehdidin kapsamı hızla değişebilir.

İnsan odağının beklenmedik zorluklar nedeniyle yavaşlayabileceği yerlerde, yapay zeka ve makine öğrenimine dayalı bir sistem minimum gecikmeyle hareket edebilir. Tehdit uyarısı yorgunluğu, dikkatli bir şekilde ele alınmazsa kuruluşlara başka bir zayıflık daha verir. Yukarıda belirtilen güvenlik katmanları daha ayrıntılı ve yaygın hale geldikçe saldırı yüzeyleri artmaktadır. Birçok güvenlik sistemi, bilinen birçok soruna tamamen refleksif uyarılarla tepki verecek şekilde ayarlanmıştır. Sonuç olarak, bu bireysel uyarılar insan ekiplerinin olası kararları ayırtmasını ve harekete geçmesini sağlar. Yüksek uyarı akışı, bu düzeydeki karar vermeyi özellikle zorlu bir süreç haline getirir. Sonuç olarak, karar yorgunluğu siber güvenlik personeli için günlük bir deneyim haline gelir. Bu belirlenen tehditler ve güvenlik açıkları için proaktif eylem idealdir, ancak birçok ekip tüm temelleri kapsayacak zamana ve personele sahip değildir. Bazen ekipler önce en büyük endişelerle yüzleşmeye ve ikincil hedefleri bir kenara bırakmaya karar vermek zorunda kalır. Siber güvenlik çabaları içinde yapay zekanın kullanılması, bilgi teknoloji ekiplerinin bu tehditlerin çoğunu etkili ve pratik bir şekilde yönetmesine olanak tanıyabilir. Bu tehditlerin her biriyle yüzleşmek, otomatik etiketleme ile toplu olarak yapılırsa çok daha kolay hale getirilebilir. Bunun ötesinde, bazı endişeler aslında makine öğrenimi algoritması tarafından ele alınabilir. Tehdit yanıt süresi, siber güvenlik ekiplerinin etkinliği için kesinlikle en önemli ölçütlerden biridir. Kötü niyetli saldırıların istismardan dağıtımına kadar çok hızlı hareket ettiği bilinmektedir. Geçmişteki tehdit aktörleri, saldırılarını başlatmadan önce bazen haftalarca ağ izinlerini incelemek ve güvenliği yatay olarak devre dışı bırakmak zorundaydı. Ne yazık ki, siber savunma alanındaki uzmanlar teknoloji yeniliklerinden faydalanan tek kişiler değil. Otomasyon o zamandan beri siber saldırılarda daha yaygın hale geldi. Son LockBit fidye yazılımı saldırıları gibi tehditler saldırı sürelerini önemli ölçüde hızlandırdı. Şu anda, bazı saldırılar yarım saat kadar hızlı hareket edebiliyor. İnsan tepkisi, bilinen saldırı türlerinde bile ilk saldırının gerisinde kalabilir. Bu nedenle, birçok ekip, girişimde bulunulan saldırıları önlemek yerine başarılı saldırılara tepki vermeye daha sık başvurmuştur.

ML destekli güvenlik, bir saldırıdan veri çekip hemen gruplandırılıp analize hazır hale getirebilir. Siber güvenlik ekiplerine işleme ve karar vermeyi daha temiz bir

iş haline getirmek için basitleştirilmiş raporlar sağlayabilir. Personel kapasitesi, küresel olarak birçok bilgi teknoloji ve siber güvenlik ekibini etkileyen devam eden sorunların kapsamına girer. Bir organizasyonun ihtiyaçlarına bağlı olarak, kalifiye profesyonellerin sayısı sınırlı olabilir. Ancak, daha yaygın durum, insan yardımının işe alınmasının organizasyonlara bütçelerinin sağlıklı bir miktarına mal olabilmesidir. İnsan personelini desteklemek, yalnızca günlük emeği telafi etmeyi değil, aynı zamanda eğitim ve sertifikasyon için devam eden ihtiyaçlarında yardım sağlamayı gerektirir. Bir siber güvenlik uzmanı olarak güncel kalmak, özellikle de bugüne kadar tartışma boyunca bahsetmeye devam ettiğimiz sürekli yenilik açısından zorludur. AI tabanlı güvenlik araçları, personeli ve desteği için daha az yoğun bir ekiple öncülük edebilir. Bu personelin AI ve makine öğreniminin en son alanlarıyla başa çıkması gerekirken, maliyet ve zaman tasarrufu daha küçük personel gereksinimleriyle birlikte gelecektir. Uyarlanabilirlik, belirtilen diğer noktalar kadar belirgin bir endişe değildir, ancak bir organizasyonun güvenliğinin yeteneklerini önemli ölçüde değiştirebilir. İnsan ekipleri, beceri setlerini özel gereksinimlerinize göre özelleştirme kapasitelerinden yoksun olabilir. Personel belirli yöntemler, araçlar ve sistemler konusunda eğitilmemişse, bunun sonucunda ekibinizin etkinliğinin zayıfladığını görebilirsiniz. Yeni güvenlik politikaları benimsemek gibi basit görünen ihtiyaçlar bile insan tabanlı ekiplerle yavaş ilerleyebilir. Bu, insan olmanın doğasıdır, çünkü işleri yapmanın yeni yollarını anında öğrenemeyiz ve bunu yapmak için zamana ihtiyacımız vardır. Doğru veri kümeleriyle, uygun şekilde eğitilmiş algoritmalar, özellikle sizin için özel bir çözüme dönüştürülebilir.

Siber güvenlikte yapay zeka, makine öğrenimi ve derin öğrenme siber güvenliği gibi disiplinlerin bir üst kümesi olarak kabul edilir, ancak oynayacağı kendi rolü vardır. Yapay zeka özünde "başarıya" odaklanır ve "doğruluk" daha az ağırlık taşır. Ayrıntılı problem çözmede doğal tepkiler nihai hedeftir. Yapay zekanın gerçek bir uygulamasında, gerçek bağımsız kararlar alınır. Programlaması, veri kümesinin yalnızca mantıksal sonucunu bulmak yerine, bir durumda ideal çözümü bulmak için tasarlanmıştır. Daha fazla açıklamak gerekirse, modern yapay zekanın ve onun altında yatan disiplinlerin şu anda nasıl çalıştığını anlamak en iyisidir. Otonom sistemler, özellikle siber güvenlik alanında, yaygın olarak seferber edilen sistemlerin kapsamında değildir. Bu kendi kendini yönlendiren sistemler, birçok kişinin yapay

zeka ile yaygın olarak ilişkilendirdiği sistemlerdir. Ancak, koruyucu hizmetlerimize yardımcı olan veya bunları artıran yapay zeka sistemleri pratik ve kullanılabilir. Yapay zekanın siber güvenlikteki ideal rolü, makine öğrenimi algoritmaları tarafından oluşturulan kalıpların yorumlanmasıdır. Elbette, modern yapay zekanın sonuçları bir insanın yetenekleriyle yorumlaması henüz mümkün değil. İnsan benzeri çerçeveler peşinde bu alanı geliştirmeye yardımcı olmak için çalışmalar yapılıyor, ancak gerçek yapay zeka, makinelerin soyut kavramları yeniden çerçevelemek için durumlar arasında taşımalarını gerektiren uzak bir hedeftir. Başka bir deyişle, bu yaratıcılık ve eleştirel düşünce düzeyi, yapay zeka söylentilerinin size inandırmak istediği kadar yakın değildir (Srinivas, 2024).

Makine öğrenimi, bazı açılarından gerçek yapay zekanın tam tersidir. Makine öğrenimi özellikle "doğruluk" odaklıdır, ancak "başarıya" o kadar odaklanmaz. Bunun anlamı, makine öğreniminin görev odaklı bir veri kümesinden öğrenmeyi amaçlayarak ilerlemesidir. Verilen görevin en uygun performansını bularak sonuçlanır. Verilen verilere dayalı tek olası çözümü, ideal olmasa bile, takip edecektir. Makine öğrenimi ile verilerin gerçek bir yorumu yoktur, bu da bu sorumluluğun hala insan görev güçlerine düştüğü anlamına gelir. Makine öğrenimi, veri deseni tanımlama ve uyarılma gibi sıkıcı görevlerde mükemmeldir. İnsanlar, görev yorgunluğu ve genellikle monotonluğa karşı düşük toleransları nedeniyle bu tür görevlere pek uygun değildir. Bu nedenle, veri analizinin yorumlanması hala insan elindeyken, makine öğrenimi, verilerin okunabilir, incelemeye hazır bir sunumda çerçevelenmesine yardımcı olabilir. Makine öğrenimi siber güvenliği, her biri kendine özgü avantajlara sahip birkaç farklı biçimde gelir: Veri sınıflandırma, veri noktalarına kategoriler atamak için önceden belirlenmiş kuralları kullanarak çalışır. Bu noktaları etiketlemek, saldırılar, güvenlik açıkları ve proaktif güvenliğin diğer yönleri hakkında bir profil oluşturmanın önemli bir parçasıdır (Shaukat vd., 2020). Bu, makine öğrenimi ve siber güvenliğin kesiştiği nokta için temeldir. Veri kümeleme, önceden belirlenmiş kuralları sınıflandırmanın aykırı değerlerini alır ve bunları paylaşılan özelliklere veya tuhaf özelliklere sahip "kümelenmiş" veri koleksiyonlarına yerleştirir. Örneğin, bu, bir sistemin daha önceden eğitilmediği saldırı verilerini analiz ederken kullanılabilir. Bu kümeler, bir saldırının nasıl gerçekleştiğini ve neyin istismar edildiğini ve açığa çıkarıldığını belirlemeye yardımcı olabilir. Önerilen eylem yolları, bir ML güvenlik

sisteminin proaktif önlemlerini yükseltir. Bunlar, davranış kalıpları ve önceki kararlar etrafında şekillenen ve doğal olarak önerilen eylem yolları sağlayan tavsiyelerdir. Burada bunun gerçek otonom AI aracılığıyla akıllı karar alma olmadığını tekrar belirtmek önemlidir. Aksine, mantıksal ilişkileri sonuçlandırmak için önceden var olan veri noktalarına ulaşabilen uyarlanabilir bir sonuç çerçevesidir. Tehditlere verilen yanıtlar ve riskleri azaltma, bu tür bir araçla büyük ölçüde desteklenebilir. Olasılık sentezi, önceki verilerden alınan derslere ve yeni, alışılmadık veri kümelerine dayalı yepyeni olasılıkların sentezlenmesine olanak tanır. Bu, bir eylemin veya bir sistemin durumunun benzer geçmiş durumlarla uyumlu olma olasılığına daha fazla odaklandığı için önerilerden biraz farklıdır. Örneğin, bu sentez bir organizasyonun sistemlerindeki zayıf noktaların önleyici bir şekilde araştırılması için kullanılabilir. Bu fayda, mevcut veri kümelerini değerlendirerek olası sonuçları tahmin ederek elde edilir. Bu, öncelikle tehdit modelleri oluşturmak, dolandırıcılık önlemeyi, veri ihlali korumasını ana hatlarıyla belirtmek için kullanılabilir ve birçok öngörücü uç nokta çözümünün temelini oluşturur. Daha detaylı açıklamak gerekirse, siber güvenlikle ilgili olarak makine öğreniminin değerini vurgulayan birkaç örnek aşağıda verilmiştir: Kuruluşunuzu gizlilik yasalarının ihlallerinden korumak son birkaç yıldır muhtemelen en önemli öncelik haline gelmiştir. Genel Veri Koruma Yönetmeliği (GDPR) öncülük ederken, Kaliforniya Tüketici Koruma Yasası (CCPA) gibi diğer yasal önlemler ortaya çıkmıştır. Müşterilerinizin ve kullanıcılarınızın toplanan verilerinin yönetimi bu yasalar uyarınca yapılmalıdır, bu da genellikle bu verilerin talep üzerine silinmeye erişilebilir olması gerektiği anlamına gelir. Bu mevzuatlara uymamanın sonuçları arasında ağır para cezaları ve kuruluşunuzun itibarının zarar görmesi yer alır. Veri sınıflandırması, tanımlayıcı kullanıcı verilerini anonimleştirilmiş veya kimliği belirsiz verilerden ayırmanıza yardımcı olabilir. Bu, özellikle büyük veya eski kuruluşlarda, eski ve yeni verilerin büyük koleksiyonlarını ayırıştırma girişimlerinde manuel işçilikten sizi kurtarır.

Kullanıcı davranışlarına dayalı ağ personeli üzerinde özel profiller oluşturarak, güvenlik kuruluşunuza uyacak şekilde özelleştirilebilir. Bu model daha sonra yetkisiz bir kullanıcının kullanıcı davranışının uç değerlerine göre nasıl görünebileceğini belirleyebilir. Klavye vuruşları gibi ince özellikler, öngörücü bir tehdit modeli oluşturabilir. Olası yetkisiz kullanıcı davranışlarının olası sonuçlarının ana hatlarıyla

birlikte, Makine öğrenimi güvenliği, maruz kalan saldırı yüzeylerini azaltmak için önerilen başvuru yollarını önerebilir (Shah, 2021). Kullanıcı davranışı profili kavramına benzer şekilde, tüm bilgisayarınızın performansının özel bir tanılama profili, sağlıklı olduğunda derlenebilir. İşlemci ve bellek kullanımının yüksek internet verisi kullanımı gibi özelliklerle birlikte izlenmesi, kötü amaçlı etkinliğin göstergesi olabilir. Bununla birlikte, bazı kullanıcılar video konferans veya sık sık büyük medya dosyası indirmeleri yoluyla düzenli olarak yüksek hacimli veri kullanabilir. Bir sistemin temel performansının genel olarak nasıl görüldüğünü öğrenerek, daha önceki bir makine öğrenimi örneğinde bahsettiğimiz kullanıcı davranışı kurallarına benzer şekilde, nasıl görünmemesi gerektiğini belirleyebilir. Bot etkinliği web siteleri için gelen bant genişliğini tüketebilir. Bu özelliklerle, özel e-ticaret mağazaları ve fiziksel mağazaları olmayanlar gibi internet tabanlı iş trafiğine bağımlı olanlar için geçerlidir. Gerçek kullanıcılar, trafik ve iş fırsatı kaybına neden olan yavaş bir deneyim yaşayabilir. Bu etkinliği sınıflandırarak, makine öğrenimi güvenlik araçları, onları anonimleştirebilen sanal özel ağlar gibi kullanılan araçlardan bağımsız olarak botların web'ini engelleyebilir. Kötü niyetli taraflara ilişkin davranışsal veri noktaları, bir makine öğrenimi güvenlik aracının bu davranış etrafında tahmini modeller oluşturmasına ve bu aynı etkinliği görüntülemek için yeni web adreslerini önceden engellemesine yardımcı olabilir. Bu güvenlik biçiminin geleceği etrafındaki tüm parlak diyaloglara rağmen, hala dikkat edilmesi gereken sınırlamalar var. Makine öğreniminin veri kümelerine ihtiyacı vardır ancak veri gizliliği yasalarıyla çelişebilir.

Eğitim yazılım sistemleri, doğru modeller oluşturmak için bol miktarda veri noktası gerektirir ve bu, "unutulma hakkı" ile pek uyuşmaz. Bazı verilerin insan tanımlayıcıları ihlallere neden olabilir, bu nedenle olası çözümlerin dikkate alınması gerekecektir. Olası düzeltmeler arasında yazılım eğitildikten sonra orijinal verilere erişimi neredeyse imkansız hale getiren sistemler edinmek yer alır. Veri noktalarının anonimleştirilmesi de dikkate alınmaktadır, ancak program mantığını çarpıtmaktan kaçınmak için bunun daha fazla incelenmesi gerekecektir. Sektörün bu kapsamdaki programlamayla çalışabilen daha fazla yapay zeka ve makine öğrenimi siber güvenlik uzmanına ihtiyacı vardır. Makine öğrenimi ağ güvenliği, gerektiğinde bakımını yapabilen ve ayarlamalar yapabilen personelden büyük ölçüde faydalanacaktır. Ancak, nitelikli ve eğitilmiş bireylerin küresel havuzu, bu çözümleri sağlayabilen personele

olan büyük küresel talepten daha küçüktür. İnsan ekipleri hala önemli bir rol oynamaya devam edecektir. Son olarak, eleştirel düşünme ve yaratıcılık, karar alma sürecinde hayati bir öneme sahip olacaktır. Daha önce de belirtildiği gibi, makine öğrenimi bu becerileri gerçekleştirmeye hazır veya (tıpkı yapay zekanın da aynı şekilde olmadığı gibi) yetenekli değildir (Vegesna, 2023).

2.2. NSL-KDD Veri Seti ve Saldırı Kategorileri Bağlamında Siber Güvenlik Yaklaşımları

DoS, Probe, U2R ve R2L saldırıları olarak kategorize edilmiş normal ve saldırı trafiği verilerini içerir. Geliştirilmiş bir sürüm olan NSL-KDD, daha iyi makine öğrenimi değerlendirmesi için yedeklilik ve dengesizlik sorunlarını ele alır. Bu veri seti, 300.000'inin 24 farklı saldırı türünü temsil ettiği yaklaşık 4.900.000 örnek içerir. Her örnek 41 özellik ile tanımlanır ve bir saldırı veya normal olarak etiketlenir. Ancak, KDDcup99 birçok araştırmacı tarafından çok sayıda yedek kayıt ve düzensizlik içerdiği için eleştirilmiştir (Stolfo vd., 2000). Bu sorunu gidermek için, tüm KDDcup99 veri kümesinin seçili kayıtlarını içeren yeni bir veri kümesi olan NSL-KDD önerildi (Warzyński vd., 2018). KDDcup99 ile NSL-KDD arasındaki farklar, oluşturulan sınıflandırıcıların önyargılı olmadığından emin olmak için NSL-KDD'deki yedek kayıtların eğitim kümesinden silinmiş olmasıdır. Ayrıca, bazı yöntemlere uygulandığında daha iyi bir tespit oranına sahip olmak için yinelenen kayıtlar hariç tutulmuştur. NSL-KDD veri setinde temsil edilen dört saldırı kategorisi vardır: Hizmet reddi saldırısı (DoS), meşru kullanıcılarının erişemeyeceği bir makineyi veya ağı dondurmaya veya kapatmaya amaçlayan bir saldırdır. DoS saldırıları, hedefi trafikle boğarak veya çökmeye neden olan bilgiler ileterek bunu başarır. Sonuç olarak, Hizmet Dışı Bırakma (DoS) saldırısı, hedeflenen kullanıcıların bir hizmete veya kaynağa erişimini engellemeyi amaçlar. Kullanıcıdan Köke Saldırı (U2R) ise, saldırganın sistemde sıradan bir kullanıcı hesabıyla oturum açtıktan sonra, sistemdeki güvenlik açıklarını tespit ederek bu açıklar üzerinden yönetici (kök) yetkileri elde etmeye çalıştığı bir saldırı türüdür. Kök saldırısı, sistemin en yetkili kullanıcısı olan "root" (kök) hesabının yetkilerinin ele geçirilmesi anlamına gelir; bu da saldırgana sistem üzerinde tam denetim imkânı sağlar.

Uzaktan Yerel Saldırı (R2L), saldırganın hedeflenen makinelerdeki herhangi bir açığı ortaya çıkarmak ve yalnızca yerel bir kullanıcının sahip olabileceği ayrıcalıklardan yararlanmak için erişiminin olmadığı hedeflenen bir makineye paketler gönderdiği bir saldırdır. Araştırma Saldırısı, saldırganın sistemi tehlikeye atmak için kullanılabilecek herhangi bir açığı bulmak amacıyla hedeflenen bir makineyi taradığı bir saldırdır.

Modern ulusal güvenlik kavramı (sözde post-Vestfalya devletinin güvenliği), vatandaş güvenliği, ulusal güvenlik sistemi ve ulusüstü güvenlik mekanizmalarının işlevi ile elde edilen, sürdürülen ve iyileştirilen, ulusal ve devlet değerlerinin ve çıkarlarının kontrolsüz bir şekilde gerçekleştirilmesi, geliştirilmesi, yararlanılması ve en iyi şekilde korunması durumudur ve ayrıca (bireysel, grup ve kolektif) onları tehlikeye atma korkusunun olmaması ve toplum ve devlet yaşamı için önemli olan gelecekteki olguların ve olayların gelişimi üzerinde kolektif bir huzur, kesinlik ve kontrol duygusudur (Arafat vd., 2018). 21 yüzyılın eşiğinde, istihbarat ve güvenlik servislerinin yetkilerinin genişlediği açıktır (örneğin, örgütlü, finansal, yüksek teknoloji, enerji ve siber güvenliğin bastırılması). Ancak, hala topluma ve devlete yönelik tüm tehditlere yanıt vermiyorlar. Kural olarak, dikkatleri genellikle örgütlü ve komplocu faaliyetlerin önlenmesine ve bastırılmasına yöneliktir, bunların çoğu gizli veya gizlidir (yani, ülke içinde veya dışında istihbarat veya yıkıcı faaliyetleri doğrudan örgütleyen veya yürüten bireyler, gruplar, örgütler ve kurumlar tarafından gerçekleştirilen siyasi olarak motive edilmiş tehlike olayları). Mevcut koşullarda, istihbarat-güvenlik faaliyeti, devletin güvenlik-istihbarat sistemi içindeki istihbarat servislerinin ve diğer kurumların yetki, hak ve görevlerinin kapsamını temsil eder; bunlar, yasal düzenlemeler ve siyasi kararlar tarafından talep edildiği şekilde sistematik olarak istihbarat toplar, işler ve sunar ve diğer faaliyetleri yürütür. Ulusal güvenlik işlevinin sahipleri olarak, istihbarat ve güvenlik servisleri, ulusal güvenlik sisteminin alt sistemlerinden biri olan güvenlik-istihbarat sisteminin ayrılmaz konularıdır (Warzyński vd., 2018). Siber uzay aracılığıyla casusluk olgusu hem popülist hem de çağdaş profesyonel ve bilimsel literatürde tartışılmaktadır. ABD, Çin ve Rusya Federasyonu da dahil olmak üzere birçok önde gelen ülkenin hükümetleri, siber casusluk sorunundan şikayetçidir. Bilgisayar ağlarının güvenliğinin artık sadece teknik bir sorun değil, aynı zamanda önemli bir stratejik konu olduğu, 2011 yılında

kıdemli Amerikalı diplomat Henry Kissinger'ın ABD ve Çin temsilcilerine "siber detante" başlatmaları için gönderdiği davetle doğrulandı. Kissinger, iki ülke arasında siber uzayın belirli alanlarının bilgisayar casusları, hırsızlar ve bilgisayar korsanları için yasak olduğunu ilan edecek bir anlaşmanın savunuculuğunu yaptı. Teorik olarak, bu tür istihbarat operasyonları çok uzak bir mesafeden ve dünyanın herhangi bir yerinden temel ve hassas siyasi ve askeri iletişime yönelik olsa bile, istihbarat verilerinin toplanması olasılığı vardır. Tüm gizli servislerin özlemi, sözde "sıfır gün", yani diğer insanların sistemlerine gizlice sızma (hiçbir hasara yol açmadan) yeteneğidir, böylece sistem sahibi "sessiz" gözetimin nesnesi olduğunu bilmez (Putnik 2009). Sıfır gün, yazılım, donanım veya aygıt yazılımında bulunan ve hatayı düzeltmekten veya başka bir şekilde düzeltmekten sorumlu taraf veya taraflarca bilinmeyen bir güvenlik açığıdır. Sıfır gün açığı terimi, hatanın kendisine atıfta bulunurken, sıfır gün saldırısı, açığın keşfedildiği zamandan ilk saldırının yapıldığı zamana kadar sıfır gün olan bir saldırıya atıfta bulunur. Sıfır gün istismarı, bilgisayar korsanlarının bir açığı (genellikle kötü amaçlı yazılım yoluyla) kullanmak ve saldırıyı gerçekleştirmek için kullandıkları yöntem veya tekniğe atıfta bulunur. Bilgisayar ağı istismarı, istihbarat toplama ve bir saldırganın verilerine bilgi sistemi aracılığıyla erişim sağlayan diğer işlemleri içerir (Arun, 2023). Bilgisayar ağı istismarı, saldırganın karar alma sürecini etkilemek ve müttefik kuvvetlerin istihbaratını geliştirmek amacıyla saldırganın bilgi sistemine sızma konusunda kasıtlı ve düşünülmüş bir eylemdir. Bu işlemler, rakibin ağlarından bilgi çıkarmayı (pasif form) ve rakibin ağlarına veri ve bilgi yerleştirmeyi başarır; bu da rakibin savaş alanını doğru bir şekilde değerlendirme yeteneğini düşürür (aktif form, yani sözde gizli etki operasyonu). Veri çıkarma ve ekleme işlemleri aşağıdaki gibi tanımlanır:

- Veri çıkarma, rakibin ağında dolaşan verileri yakalamayı veya rakibin veritabanından bilgi çıkarmayı içeren pasif bir tekniktir. Bu nedenle, bilgi edinmek için kötü amaçlı bağlantılara ve düğümlere erişim gereklidir. Çıkarma tekniğiyle elde edilen istihbarat daha sonra kullanılabilir, değiştirilebilir ve rakibin ağına geri beslenebilir. Verilerin kopyalanabileceğini ve bu nedenle çıkarımlarının tespit edilemeyeceğini belirtmek gerekir;

- Ekleme, rakibin veri tabanına veri eklemeyi içeren ve rakibin bilgilerinin manipüle edilmesini sağlayan aktif bir tekniktir. Bu teknikle, karşı tarafın yanlış bir istihbarat resmi edinmesi sağlanır ve bu da diğer tarafa avantaj sağlar.

Bilgisayar ve ağ istismarının aksine, siber casusluk farklı stratejilere, taktiklere ve araçlara dayanan nispeten yeni bir istihbarat toplamadır. Siber casusluk, politik, askeri, ekonomik ve diğer avantajlar elde etmek, verileri kendi amaçları için kullanmak veya elde edilen bilgileri finansal kazanç için satmak için bir düşman hakkında hassas bilgilere erişmek amacıyla uluslararası düzeyde bilgisayarları veya dijital iletişimleri kullanmaktır. Bilgi ve iletişim teknolojilerinin yaygınlaştığı çağda, devletlerin ulusal güvenliğine yönelik güvenlik riskleri ve tehditler önemli ölçüde artmakta ve bu bağlamda siber uzaydaki istihbarat servislerinin rolleri önemli ölçüde önem kazanmaktadır. Yani istihbarat servisleri, istihbarat bilgilerinin toplanması ve korunması, kritik altyapılara saldırı ve savunma, siber savaşçıların işe alınması ve düşman saflarında tespit edilmesi bağlamında siber uzay istihbarat araştırmaları üzerinde yoğun bir şekilde çalışmaktadır. Modern askeri doktrinlerde sanal, Öklid dışı uzay -siberuzay- kara, su, hava ve uzayla birlikte beşinci muharebe alanı statüsünü kazanmıştır. Siberuzay yalnızca propaganda faaliyetleri yürütmek için uygun bir ortam olmakla kalmayıp, aynı zamanda belirli bir devletin bilgi altyapısına saldırarak önemli fiziksel hasara da neden olabilir - bu da insan ve maddi kurbanların yanı sıra saldırıya uğrayan devletin egemenliğinin ihlaliyle sonuçlanabilir. Bu nedenle, siberuzay bugün saldırıların hedefi ve teknolojik olarak gelişmiş orduların ve onların uzmanlaşmış "bilgisayar savaşçıların" cephaneliğinde güçlü bir araç (saldırı silahı) haline gelmiştir. Medya ve bilgi teknolojisi devrimi, toplumun nasıl etkileşim kurduğunu ve devletlerin (ve paramiliter, devlet karşıtı oluşumların) nasıl savaş yürüttüğünü temelden değiştirmektedir. Askeri-teknolojik terminolojide, bilgisayar, uydu iletişimi ve medya alanındaki devrim niteliğindeki başarıların sürekli birleşmesinin, bilgi ve iletişim devrimi savaşın jeostratejik ve politik-ekonomik hedeflerini temelden değiştirmemiş olsa bile, savaş olanaklarını kökten "iyileştirdiği" açıktır (Arun, 2023).

Siber savaş, saldırgan bilgi ve iletişim (IC) sistemlerine saldırmak ve verileri saldırgan ve savunma amaçlı olarak manipüle etmek için çeşitli teknikler ve araçlar kullanır. Bilgi çağında, bilginin ve bilginin yayılması ve üretimiyle ilgili tüm

faaliyetlerin önemi önemli ölçüde artmıştır. Ancak, fiziksel alan gibi, siber alan da genellikle onu ilk elde edene aittir. Küresel ölçekte bilgi çağı teknolojileri, uluslararası arenada devletlerin, örgütlerin, çokuluslu şirketlerin, sivil toplum örgütlerinin, ulusötesi suç örgütlerinin ve hatta bireylerin rolünü artırmaktadır. Siber savaşın özneleri, yani bilgisayar saldırılarının aktörleri, analitik amaçlar için iki kategoriye ayrılabilir. Bölünme kriteri, bilgisayar saldırısını gerçekleştiren kişinin niyetinin varlığı veya yokluğudur. Bu anlamda, önceden tasarlanmış bir siber saldırı ile önceden tasarlanmamış bir saldırı arasında ayrım yapabiliriz. Siber savaşta kullanılan önceden tasarlanmış bir siber saldırı, ulusal güvenliği kasıtlı olarak etkilemek veya ulusal güvenliğe karşı daha fazla operasyon için koşullar yaratmak amacıyla siber araçlar kullanılarak gerçekleştirilen herhangi bir saldırdır. Önceden tasarlanmış saldırılar, savaşla - savaş düzeyindeki ulusal politika - eşdeğer tutulabilir. Bunlar, bir düşmanın bilgisayar ve bilgi iletişim sistemlerine karşı herhangi bir eylemi içerir.

Aktörler, çoğunlukla bireyler olmak üzere, ulusal güvenliği istemeden tehdit eden ve genellikle saldırılarının uluslararası düzeydeki potansiyel sonuçlarının farkında olmayan pervasız siber saldırılar gerçekleştirirler. Bu aktörler, sistemin savunma mekanizmalarını atlatarak ve sistemin bilgilerini, yani sistemin kendisini manipüle ederek, kullanarak veya yok ederek siber sızma gerçekleştiren herkesi içerir. Bu aktörlerin farklı saikleri ve niyetleri vardır ancak ulusal güvenliğe zarar vermeyi veya ulusal güvenliğe karşı daha fazla operasyon düzenlemeyi amaçlamazlar. Bu aktörler çoğunlukla genel olarak "hacker" terimiyle anılırlar ve siber suçlar işlemelerine rağmen, kasıtlı olarak siber savaş yürütmezler. Ancak, bu aktör kategorisinin, siber operasyonlarda bilgi ve yeteneklerini kullanarak kasıtlı saldırılar gerçekleştirenler tarafından manipüle edildiği durumlar olduğunu belirtmek önemlidir. Yıkıcı eylemler uygulamak, istihbarat servislerinin bugün yoğun olarak kullandığı yöntemlerden biridir. Bu, siyasi yaratıcıların dış politika eylemlerine güç (güç kullanımı) temelinde karar vermelerini ve bunu yalnızca klasik, fiziksel biçimde (silahlı kuvvet) değil, aynı zamanda çeşitli yıkıcı içerikler biçiminde (Ulusal İstihbarat'ın gizli direktörü Dan Coats, "kritik Amerikan altyapısını tahrip edebilecek artan yabancı siber saldırı tehdidi" olduğunu söyledi (CBS News, 2018). Siber saldırıları, kışkırtıcı eylemler, karma eylemler, siber saldırılar, psikolojik propaganda faaliyetleri tarafından tespit edilen "endişe verici faaliyetlere" bağladı) bağlaması

anlamına gelir ve bunlar belirli bir hedefin imhasının tüm özelliklerine sahiptir. Modern istihbarat servislerinin çalışmalarının bu içerikleri, öncelikle devletin katılımını tamamen gizleme ve uluslararası düzeyde kamuoyunda uzlaşma ve kınama olasılığını önleme çabası nedeniyle, uygulanan yöntemlerin olağanüstü gizliliği ve karmaşıklığı ile ayırt edilir. Yıkıcı faaliyetler esasen diğer devletlerin iç işlerine müdahaleyi, gizlenmiş, nihayetinde şiddetli bir saldırganlık biçimini temsil eder. Bu arada, yıkıcı içerik, modern istihbarat servislerinin başka bir tür faaliyetini temsil eder (birincisi istihbarattır), öncelikle dış uluslararası düzeyde, uluslararası ilişkilerde devletlerin dış politika eylemleri çerçevesinde, vurguladığımız gibi, diğer devletlere karşı güç kullanımı ile karakterize edilir. Bunlar, istihbarat servislerinin seçilen dış politika eylemi yönlerinin oluşturulması ve uygulanmasına katılımının yalnızca karar alma süreciyle sınırlı olmadığı, inisiyatif aşamasına (istihbarat verilerinin toplanması, işlenmesi, atanması ve yorumlanması) ve biraz daha az sıklıkla içerik ifade aşamasında aktif olarak yer aldıklarının açık kanıtıdır. Bu nedenle, istihbarat servisleri, yıkıcı içerikleri planlama, örgütlenme ve uygulama yoluyla dış politika hedeflerine ulaşmada doğrudan rol oynarlar. İstihbarat servislerinin bu planda belirli dış politika arabulucuları olarak rolü, dünyadaki modern istihbarat ve güvenlik sistemlerinin büyük olanaklarının bir sonucudur ve bu da istihbarat servislerinin eylem repertuarını istihbarat dışı eylemlerle azami ölçüde zenginleştirmesini önemli ölçüde etkilemiştir. Bu bağlamda, dış politika araçlarıyla birlikte, istihbarat servisleri dış politika arabulucuları olarak genellikle belirli dış politika prosedürünün doğasını da belirlerler.

2.3. Nsl-Kdd Veri Seti İle Knime Uygulaması Kullanarak Saldırı Tespiti ve Makine Öğrenimi Algoritmalarının Performans Analizi

KNIME (Konstanz Information Miner), açık kaynaklı bir veri analitiği platformudur. Makine öğrenimi, veri madenciliği ve büyük veri analizi için güçlü araçlar sağlar. Kullanıcı dostu grafik arayüzüyle kullanıcılar kod yazmadan verileri işleyebilir ve modelleyebilir. KNIME, veri temizleme, ön işleme, özellik mühendisliği, görselleştirme ve model eğitimi adımlarının kolayca yürütülmesini sağlar. Makine öğrenimi algoritmalarını test etmek, model performansını değerlendirmek ve büyük veri kümelerini analiz etmek için idealdir. Özellikle siber güvenlikte, saldırı tespiti için kNN (k-En Yakın Komşular) gibi algoritmalar

uygulanarak NSL (Ağ Güvenlik Laboratuvarı) -KDD veri kümesi kullanılarak anormallik tespiti yapılabilir (Kehe vd., 2011). Karar Ağaçları, her düğümün bir nitelik veya özellik olarak sunulduğu ve her dalın bir kural veya karar olarak sunulduğu ve her yaprağın bir sonuç olarak sunulduğu düğümlerden, dallardan ve yapraklardan oluşur. Karar ağacına, verilerimize uygulanan ve tahmin edilen bir sınıfla sonuçlanan bir dizi evet/hayır sorusu olarak bakılabilir (Ravipati, 2019).

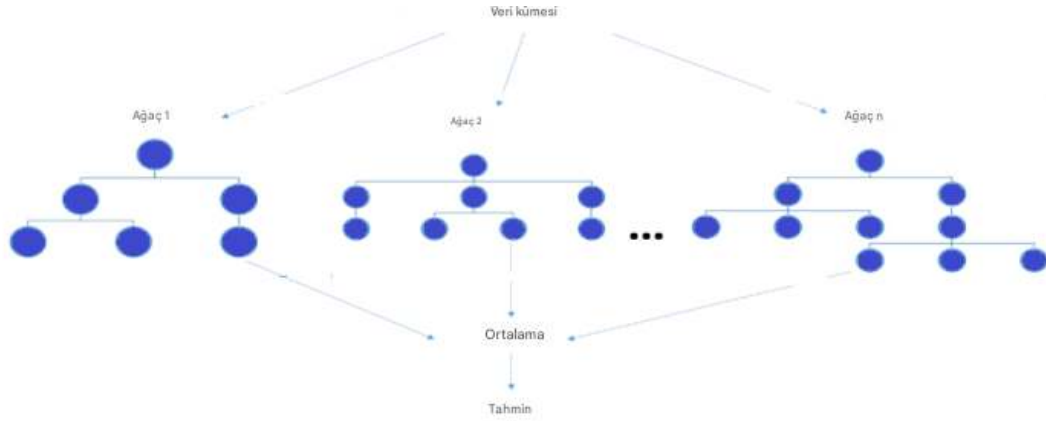


Şekil 1. İki Seviyeli Karar Ağaçları

Kaynak: Abualkibash, M. (2019). Machine learning in network security using KNIME analytics. International Journal of Network Security & Its Applications (IJNSA), 11(5), 1–14.

Karar ağacı algoritmasından türetilen birçok algoritma vardır. Bunlardan bazıları ID3, C4.5, J48 vb.'dir. ID3 algoritmasıyla ilgili sorun, bilgilerin aşırı uyum sağlayabilmesidir. C4.5, ID3'ün geliştirilmiş bir sürümüdür ve aşırı uyum sorununu ele alır. Bu, J48'in C4.5 algoritmasının açık kaynaklı bir uygulaması olduğu anlamına gelir. C4.5, sınıflandırma görevleri ve veri madenciliği için makine öğreniminde kullanılan bir karar ağacı algoritmasıdır. J48, Weka veri madenciliği yazılımında C4.5'in açık kaynaklı bir sürümü olarak mevcuttur. Açık kaynaklı olmak, J48'in kodunun herkes tarafından erişilebilir, incelenebilir ve değiştirilebilir olduğu anlamına gelir. Rastgele orman, birçok karar ağacından yapılandırılmış bir modeldir. Ağaçlar oluşturulurken, rastgele ormandaki her ağacın eğitim verilerinin rastgele bir örneğinden öğrendiği eğitim verilerinin rastgele örneklenmesi dikkate alınacaktır. Her ağacı farklı örneklerde eğitmenin amacı, her ağaç belirli bir eğitim verisi kümesiyle

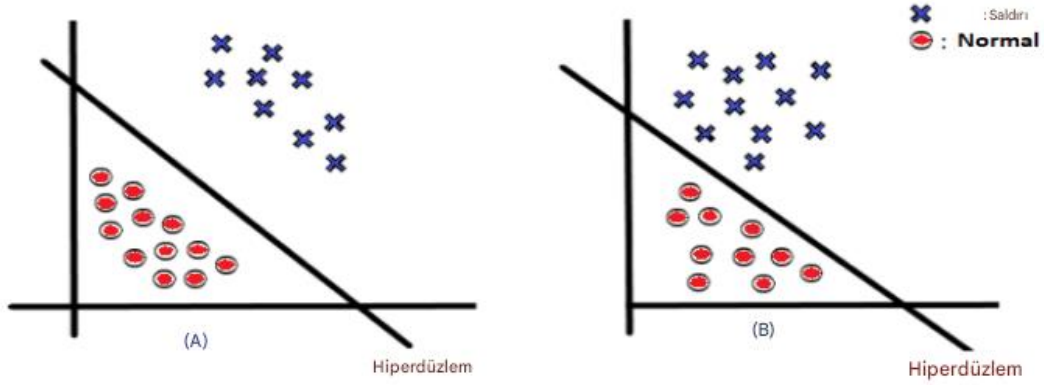
ilgili yüksek varyansa sahip olsa bile, tüm ormanın önyargıyı artırmadan daha düşük varyansa sahip olmasıdır. Test sırasında, tahminler her karar ağacının tahmin ortalamaları alınarak oluşturulur. Her bir öğrenciyi eğitim verilerinin çeşitli önyüklemeli alt kümeleri üzerinde eğitime ve ardından tahminlerin ortalamalarını alma adımlarına, önyükleme toplamanın kısaltması olan torbalama adı verilir.



Şekil 2. Rastgele Ormanlar ve Karar Ağaçları

Kaynak: Abualkibash, M. (2019). Machine learning in network security using KNIME analytics. International Journal of Network Security & Its Applications (IJNSA), 11(5), 1–14.

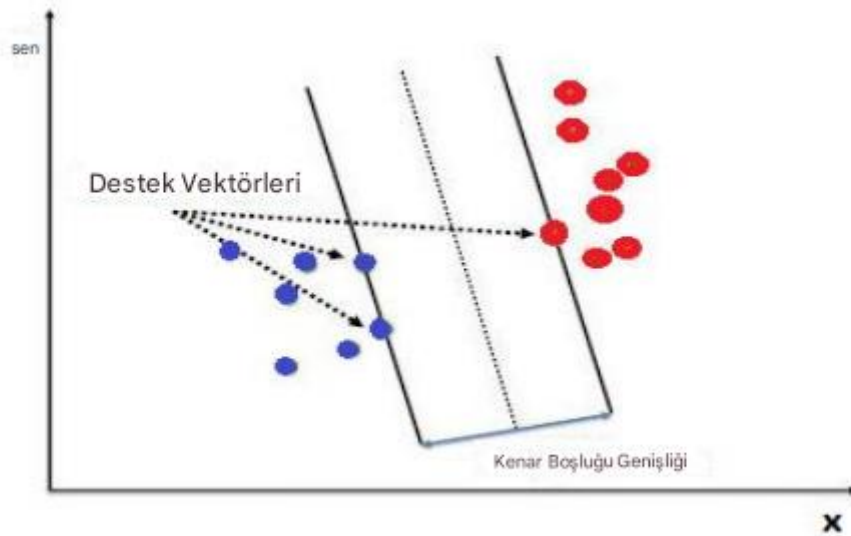
Her karar ağacındaki düğümleri bölerken özelliklerin rastgele alt kümeleri dikkate alınacaktır. Örneğin, 4 özellik varsa, her ağaçtaki her düğümde, düğümü bölmek için yalnızca 2 rastgele özellik dikkate alınacaktır. Naive Bayes, makine öğrenimi modelinin Bayes teoremini kullanarak sınıflandırmada kullanılacak olasılığa dayalı olarak oluşturulduğu bir sınıflandırıcıdır. Örneğin, Bayes teoremini uygulayarak, bir olay meydana geldiğinde bir saldırının gerçekleşme olasılığını hesaplayabiliriz. Destek vektör makineleri, sınıflandırma durumlarında uygulanabilen bir gözetimli öğrenme algoritmasıdır. Genellikle küçük bir veri kümesine uygulanır çünkü işlenmesi uzun zaman alır. SVM, özellikleri en iyi şekilde birkaç alana bölen bir hiper düzlemi belirleme kavramına dayanarak oluşturulmuştur. Örneğin, iki durumu tanımlayacak bir fonksiyon (hiper düzlem) oluşturmak istiyorsak, böylece bir sunucu bu kadar çok paket aldığında bu bir saldırı mı yoksa normal mi olarak sınıflandırılacak? Aşağıda hiper düzlemin çizildiği iki senaryonun şekilleri verilmiştir.



Şekil 3. SVM, Optimum Hiper Düzlem

Kaynak: Abualkibash, M. (2019). Machine learning in network security using KNIME analytics. International Journal of Network Security & Its Applications (IJNSA), 11(5), 1–14.

Hiper düzleme yakın noktalara destek vektör noktaları, vektörler ile hiper düzlem arasındaki boşluğa ise kenar boşlukları denir.

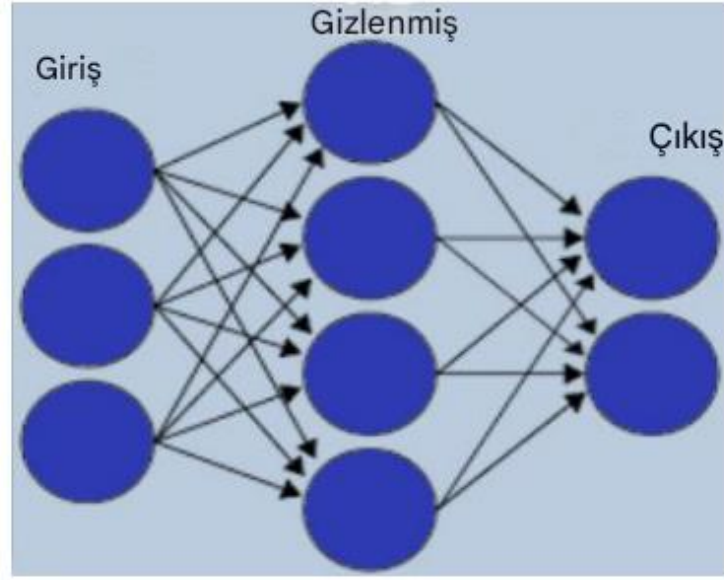


Şekil 4. Destek Vektör Noktaları ve Kenar Boşlukları

Kaynak: Abualkibash, M. (2019). Machine learning in network security using KNIME analytics. International Journal of Network Security & Its Applications (IJNSA), 11(5), 1–14.

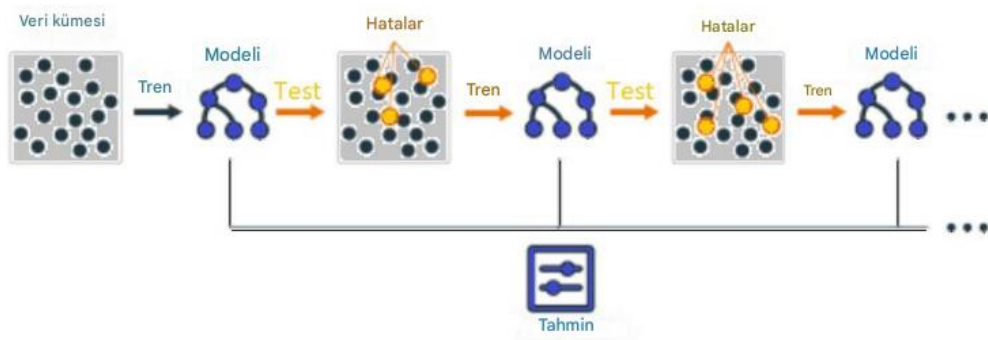
Yapay sinir ağları, makine öğrenmesinde yaygın olarak kullanılan algoritmalarından biridir. İnsanların öğrenme biçimini kopyalamak için tasarlanmış beyinden ilham alan sistemlerdir. Sinir ağları, giriş ve çıkış katmanları içerir ve ayrıca, girdiyi çıkış katmanının kullanabileceği bir şeye dönüştüren gizli katmanlara sahip

olabilir. Yapay Sinir Ağında, bir eğitim seti olarak daha fazla veri kullanılırsa daha doğru hale gelir. Birisi bir görevi tekrar tekrar yapmaya devam ettiğinde de benzerdir. Zamanla, o kişi daha verimli hale gelir ve az sayıda hata yapar.



Şekil 5. Basit Bir Gizli Katmana Sahip Basit Bir Sinir Ağı

Kaynak: Abualkibash, M. (2019). Machine learning in network security using KNIME analytics. International Journal of Network Security & Its Applications (IJNSA), 11(5), 1–14.



Şekil 6. Topluluğa Sırayla Yeni Modeller Ekleyin

Kaynak: Abualkibash, M. (2019). Machine learning in network security using KNIME analytics. International Journal of Network Security & Its Applications (IJNSA), 11(5), 1–14.

kNN (k-En Yakın Komşular) sınıflandırma algoritması, makine öğrenimindeki basit algoritmalarından biridir. Sınıflandırma, eğitim verilerindeki benzer veri noktalarını belirlemeye ve ardından sınıflandırmalarına dayalı bir tahmin yapmaya dayanır. kNN, birkaç özelliği olmayan küçük boyutlu veri kümelerinde daha iyi sonuçlar veren algoritmalarından biridir. Veri ön işleme, bir makine öğrenimi modeli oluşturma yolculuğunda önemli bir adımdır. Makine öğrenimi modellerimizi herhangi bir sorun olmadan oluşturduğumuzdan emin olmak için bazı hazırlıklar yapmalıyız. Veri ön işleme yoksa makine öğrenimi modelimiz düzgün çalışmaz. Herhangi bir veri kümesinde, bağımsız değişkenler ile bağımlı değişkenler arasındaki fark olan çok önemli bir şeyi ayırt etmemiz gerekir (Thomas vd., 2018). Herhangi bir makine öğrenimi algoritmasında veya modelinde, bağımsız değişkenler bağımlı bir değişkeni tahmin etmek için kullanılır, örneğin normal veya anormallik. Ayrıca, veri kümesinde eksik veri bulunan durumlarla da uğraşmamız gerekir. Bir sütundaki eksik verileri ele almanın en yaygın fikri, o sütunun ortalamasını almaktır. Başka bir deyişle, bir sütundaki tüm değerlerin ortalaması, o sütundaki eksik verileri değiştirecektir (Paulauskas vd., 2017). Dahası, bazen veri kümesinde, basitçe kategoriler içerdikleri için kategorik değişkenler olarak adlandırılan bir tür değişken olacaktır. Makine öğrenimi modelleri matematiksel denklemlere dayandığından, denklemlerde kategorik değişkenleri, örneğin metni tutarsak bu bazı sorunlara neden olur çünkü denklemlerde yalnızca sayılar isteriz. Sonuç olarak, kategorik değişkenleri sayılara kodlamamız gerekir. Bir sonraki adım, makine öğreniminde çok önemli olan ölçeklemeyi içerir, değişkenler aynı ölçekte olmadığında bu, bazı makine öğrenimi modellerinde bazı sorunlara neden olur.

- Gerçek pozitif, yani bir saldırının saldırı olarak öngörüldüğü anlamına gelir.
- Yanlış pozitif, yani normal bir ağ paketinin saldırı olarak tahmin edilmesi.
- Gerçek negatif, yani normal bir ağ paketinin normal olarak tahmin edilmesi.
- Yanlış negatif, yani saldırının normal olarak tahmin edilmesi.

Bu çalışmada, ağ trafiğinin normal ya da saldırı olup olmadığını belirlemeye yönelik bir sınıflandırma modeli oluşturulmuştur. Modelin kurulması sürecinde KNIME platformu kullanılmış ve NSL-KDD veri kümesinden seçilmiş 20 satırlık bir örnek veri seti analiz edilmiştir. Bu veri seti; protocol_type, src_bytes, dst_bytes ve

label gibi temel sütunlardan oluşmaktadır. label sütunu, bağlantının “normal” ya da “attack” olup olmadığını belirtmektedir ve sınıflandırma modelinin hedef değişkeni olarak kullanılmıştır.

Veri Ön İşleme

İlk olarak veri KNIME ortamına yüklenmiş ve protocol_type ile label sütunları nominal veri tipine dönüştürülmüştür. protocol_type sütunu, “One to Many” yöntemi ile dummy değişkenlere ayrılmıştır. Bu işlem sonucunda veri setinde tcp, udp ve icmp sütunları oluşturulmuştur. Sayısal sütunlardan src_bytes doğrudan modele giriş olarak alınmıştır. dst_bytes sütunu ise tüm satırlarda sıfır değer içerdiğinden modelleme sürecine dahil edilmemiştir.

Eğitim ve Test Verisi Ayırımı

Veri seti, eğitim ve test olmak üzere ikiye ayrılmıştır. Stratified Sampling yöntemi ile %70 eğitim ve %30 test oranında bir dağılım gerçekleştirilmiştir. Bu sayede sınıfların dengeli bir şekilde her iki gruba da dağılması sağlanmıştır. Toplam 20 satırlık veri setinin 14 satırı eğitim, 6 satırı ise test amacıyla kullanılmıştır.

Model Kurulumu ve Uygulama

Modelleme aşamasında Random Forest algoritması kullanılmıştır. Model, 100 karar ağacı ile oluşturulmuş ve her bir ağacın maksimum derinliği 30 olarak belirlenmiştir. Ayrıca, yaprak başına minimum örnek sayısı 1 olarak bırakılmıştır. Eğitim verisi kullanılarak model oluşturulmuş, ardından test verisi üzerinden modelin doğrulaması gerçekleştirilmiştir.

Performans Değerlendirmesi

Modelin test verisi üzerindeki performansı “Scorer” düğümü ile değerlendirilmiştir. Elde edilen confusion matrix sonuçlarına göre; 2 örnek doğru bir şekilde “normal”, 2 örnek doğru bir şekilde “attack” olarak tahmin edilmiştir. Ayrıca, bir “normal” örnek yanlışlıkla “attack” ve bir “attack” örnek yanlışlıkla “normal” olarak sınıflandırılmıştır. Bu değerlere göre modelin genel doğruluk oranı %66,7 olarak hesaplanmıştır. Kesinlik ve duyarlılık değerleri de benzer şekilde %66,7 olarak

bulunmuştur. Bu sonuçlar, küçük boyutlu ve sınıf dengesi olan veri setlerinde Random Forest algoritmasının makul düzeyde başarı sağlayabildiğini göstermektedir.

Tablo 1. Veri Seti İçeriği

Rows: 20 | Columns: 7

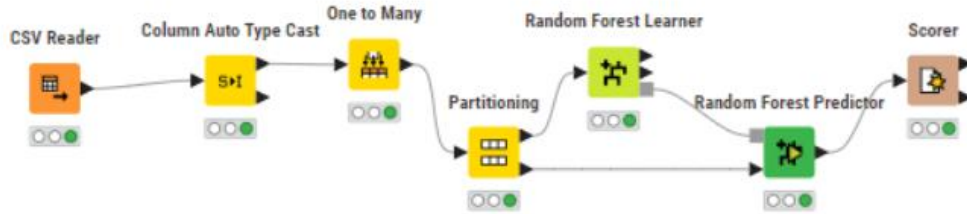
#	RowID	protocol_type	src_bytes	dst_bytes	label	tcp	udp	icmp
1	Row0	tcp	337	0	normal	1	0	0
2	Row1	tcp	0	0	attack	1	0	0
3	Row2	tcp	8314	0	attack	1	0	0
4	Row3	tcp	3427	0	normal	1	0	0
5	Row4	tcp	27264	0	normal	1	0	0
6	Row5	udp	105	0	normal	0	1	0
7	Row6	tcp	0	0	attack	1	0	0
8	Row7	tcp	372	0	normal	1	0	0
9	Row8	tcp	0	0	attack	1	0	0
10	Row9	tcp	0	0	attack	1	0	0
11	Row10	tcp	0	0	attack	1	0	0
12	Row11	tcp	281	0	normal	1	0	0
13	Row12	tcp	1471	0	normal	1	0	0
14	Row13	tcp	0	0	attack	1	0	0
15	Row14	tcp	338	0	normal	1	0	0
16	Row15	tcp	0	0	attack	1	0	0
17	Row16	tcp	248	0	normal	1	0	0
18	Row17	icmp	0	0	attack	0	0	1
19	Row18	tcp	0	0	attack	1	0	0
20	Row19	tcp	8115	0	normal	1	0	0

Kaynak: <https://www.kaggle.com/datasets/hassan06/nsllkdd>

Çalışma kapsamında kullanılan veri seti toplam 20 satırdan oluşmaktadır. Veride her bir satır, bir ağ bağlantısını temsil etmektedir. Her bağlantının bazı özellikleri vardır. Bunlar arasında en önemlileri “protocol_type”, “src_bytes”, “dst_bytes” ve “label” sütunlarıdır. “protocol_type” bağlantının hangi iletişim protokolü ile gerçekleştiğini gösterir. Bu protokoller TCP, UDP ve ICMP olarak üç gruba ayrılmıştır. İncelenen verilere göre bu bağlantıların 18 tanesi TCP, 1 tanesi UDP ve 1 tanesi ICMP protokolü kullanılarak gerçekleştirilmiştir. Bu durum veri setinde TCP bağlantılarının ağırlıkta olduğunu açıkça ortaya koymaktadır. Bu nedenle modelleme aşamasında TCP’ye ait özellikler daha etkili olabilir. “src_bytes” sütunu bağlantı başlatan sistemden gönderilen veri miktarını, “dst_bytes” ise karşı sistemden geri gönderilen veri miktarını göstermektedir. Yapılan incelemelerde “dst_bytes” sütunundaki tüm değerlerin sıfır olduğu görülmüştür. Bu da, ya veri toplanırken hedef sistemden dönüş olmadığını ya da bu verinin kayıt altına alınmadığını gösterir. Bu yüzden bu sütunun modelde anlamlı bir katkısı olmayabilir ve analizlerde dikkate alınmayabilir. Diğer yandan “src_bytes” sütunu farklı değerler içermektedir. Bazı

bağlantılarda veri gönderimi hiç yapılmamışken (değer sıfır), bazı bağlantılarda 27.000 bayta kadar veri gönderimi olmuştur. Bu farklılık, bağlantının niteliği hakkında fikir verebilir. Özellikle saldırı (attack) olarak etiketlenen bağlantıların birçoğunda veri gönderiminin az ya da hiç olmadığı göze çarpmaktadır. “label” sütunu ise her satırın normal ya da saldırı olduğunu göstermektedir. Bu sütun sınıflandırma için hedef değişkendir. Veri setinde dengeli bir dağılım söz konusudur. Toplamda 10 satır “normal”, 10 satır ise “attack” olarak etiketlenmiştir. Bu dengeli yapı modelin eğitiminde avantaj sağlayacaktır. Çünkü bir sınıfın diğerine göre baskın olması, modelin yanlış sonuçlar üretmesine neden olabilir. Dengeli veri sayesinde model her iki sınıfı da öğrenme şansı bulur.

Sonuç olarak, veri seti az sayıda satır içermesine rağmen sınıflandırma modeli geliştirmek için temel yapıya sahiptir. Protokol türleri arasında dengesizlik olsa da “label” sütununun dengeli olması, modelin güvenilirliğini artırabilir. Ancak “dst_bytes” gibi tümü sıfır olan sütunların modele etkisi olmayabileceği için çıkarılabilir. “src_bytes” sütunu ise sınıflar arasında belirleyici rol oynayabilir. Bu özelliklere dikkat edilerek oluşturulacak bir sınıflandırma modeli ile bağlantıların normal mi yoksa saldırı mı olduğu belirlenebilir.



Şekil 7. Random Forest Predictor

Kaynak: Yazar tarafından oluşturulmuştur.

Yukarıda yer alan görsel, bu çalışmada gerçekleştirilmiş olan veri analizi sürecini adım adım gösteren KNIME iş akışını temsil etmektedir. Bu akış, bir ağ bağlantısının normal mi yoksa saldırı içerikli mi olduğunu belirlemek amacıyla kurulmuştur. Kullanılan veri seti, NSL-KDD adlı siber güvenlik veri kümesinin küçük

bir örneklemini içermektedir. İş akışı yedi temel adımdan oluşmaktadır ve her bir adımın belirli bir görevi bulunmaktadır. İlk olarak, CSV Reader düğümü aracılığıyla veri seti KNIME platformuna aktarılmıştır. Bu düğüm sayesinde dışarıdan alınan .csv formatındaki veri düzgün bir şekilde programa yüklenmiş ve analiz için uygun hale getirilmiştir. Sonrasında, Column Auto Type Cast adlı düğüm ile sütun tipleri otomatik olarak düzenlenmiştir. Özellikle “protocol_type” ve “label” sütunları nominal (kategorik) verilere çevrilerek sınıflandırma için uygun hale getirilmiştir. Üçüncü aşamada, One to Many düğümü kullanılarak “protocol_type” sütunu tcp, udp ve icmp adlarında üç ayrı sütuna dönüştürülmüştür. Bu işlem, makine öğrenmesi modellerinin kategorik verileri daha doğru yorumlayabilmesi için yapılmıştır. Yani, protokol bilgisi modelin anlayabileceği sayısal formata çevrilmiştir. Dördüncü adımda, Partitioning düğümü kullanılarak veri seti ikiye ayrılmıştır. Verinin %70’i eğitim (öğrenme), %30’u ise test (değerlendirme) amacıyla ayrılmıştır. Stratified Sampling yöntemi kullanılarak “label” sütunundaki normal ve saldırı sınıflarının her iki grupta da dengeli dağılım göstermesi sağlanmıştır. Beşinci adımda, Random Forest Learner düğümü ile sınıflandırma modeli oluşturulmuştur. Bu model, veri setindeki örnekleri karar ağaçları üzerinden değerlendirerek hangi sınıfa ait olduğunu tahmin etmeye çalışmaktadır. Modelin eğitilmesi sırasında yalnızca eğitim verisi kullanılmıştır. Altıncı aşamada, Random Forest Predictor düğümüyle eğitilen model test verileri üzerinde uygulanmış ve tahminler elde edilmiştir. Her bir bağlantının “normal” mi yoksa “attack” mı olduğu model tarafından belirlenmiştir.

Son olarak, Scorer düğümü ile modelin başarımı ölçülmüştür. Gerçek değerler ile modelin tahmin ettiği değerler karşılaştırılarak doğruluk oranı ve karışıklık matrisi oluşturulmuştur. Elde edilen sonuçlara göre, toplam 6 test örneğinin 4’ü doğru sınıflandırılmış ve modelin doğruluk oranı %66,7 olarak hesaplanmıştır. Bu iş akışı sayesinde temel sınıflandırma süreci görsel olarak açıklanmış, tüm işlemler adım adım izlenmiş ve ağ trafiği analizi konusunda örnek bir uygulama gerçekleştirilmiştir.

Tablo 2. Confusion Matrix

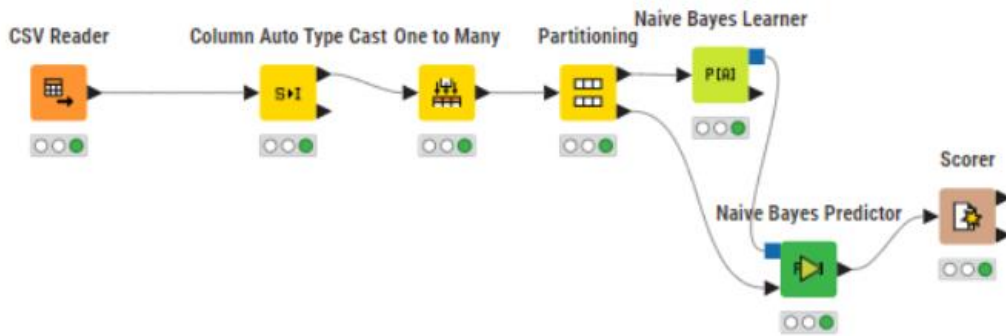
Confusion matrix (Table)				
Rows: 2		Columns: 2		
#	RowID	normal Number (integer)	attack Number (integer)	
1	normal	2	1	
2	attack	1	2	

Kaynak: Yazar tarafından oluşturulmuştur.

Değerlendirme Göstergeleri

- Doğruluk (Accuracy) = $(TP + TN) / (TP + TN + FP + FN)$
= $(2 + 2) / (6) = \%66,7$
- Kesinlik (Precision) = $TP / (TP + FP) = 2 / 3 = \%66,7$
- Duyarlılık (Recall) = $TP / (TP + FN) = 2 / 3 = \%66,7$

Model, test verileri üzerinde değerlendirilmiş ve toplam 6 örnekten 4'ünü doğru, 2'sini ise yanlış sınıflandırmıştır. Confusion matrix sonuçlarına göre modelin normal ve attack sınıflarını ayırt etme yeteneği ölçülmüştür. Modelin genel doğruluk oranı %66,7 olup, hem kesinlik hem de duyarlılık değerleri de %66,7 seviyesindedir. Bu sonuçlar, Random Forest algoritmasının küçük ve dengeli veri setlerinde sınıflandırma yeteneğine sahip olduğunu göstermektedir. Daha büyük ve dengesiz veri setlerinde daha yüksek doğruluk oranlarına ulaşmak için parametre optimizasyonları yapılabilir.



Şekil 8. Naive Bayes Learner

Kaynak: Yazar tarafından oluşturulmuştur.

Bu görsel, Naive Bayes algoritmasının KNIME programı kullanılarak oluşturulan iş akışını göstermektedir. İlk adımda, CSV Reader kullanılarak veri kümesi

içeri aktarılır. Daha sonra Column Auto Type Cast ile veri sütunlarının türleri otomatik olarak tanımlanır. One to Many node'u, kategorik verileri sayısal formata çevirmek için kullanılır. Partitioning adımı veriyi eğitim ve test veri setlerine böler. Naive Bayes Learner, eğitim verisi üzerinde modeli oluşturur. Model daha sonra Naive Bayes Predictor node'u ile test verisi üzerinde tahmin yapmak için kullanılır. Son adımda Scorer node'u ile modelin doğruluk performansı değerlendirilir. Bu akış, temel bir sınıflandırma işlemi için KNIME'da kolayca kullanılabilir bir yöntemdir.

Tablo 3. Confusion Matrix (2)



#	RowID	normal Number (integer)	attack Number (integer)
1	normal	1	0
2	attack	0	5

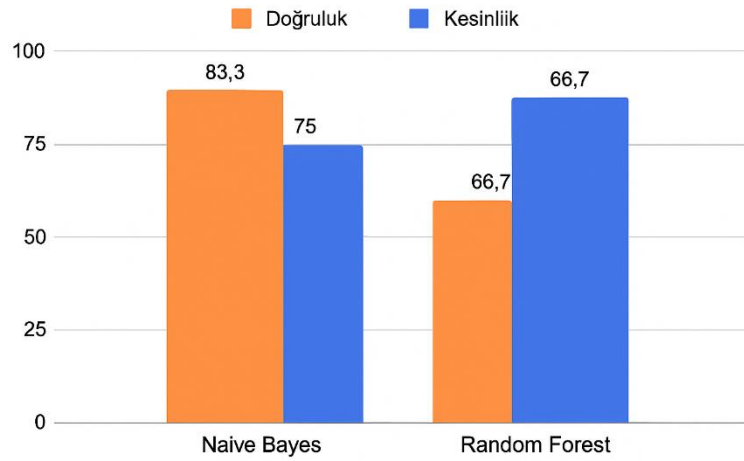
Kaynak: Yazar tarafından oluşturulmuştur.

Bu şekil, Naive Bayes algoritmasının sınıflandırma başarımını gösteren karışıklık matrisidir. 'Normal' ve 'Attack' olmak üzere iki sınıf bulunmaktadır. Matrise göre model, 1 adet 'normal' veriyi doğru tahmin etmiş, ancak başka hiçbir 'normal' veriyi doğru bilememiştir. Buna karşılık, 'attack' sınıfına ait 5 veri başarıyla doğru tahmin edilmiştir. Bu durum modelin özellikle 'attack' sınıfında daha başarılı tahminler yaptığını göstermektedir. Fakat 'normal' sınıfında düşük başarı elde edilmiştir. Bu tür analizler, modelin hangi sınıflarda güçlü ya da zayıf olduğunu değerlendirmek için kullanılır. Confusion matrix, makine öğrenmesi projelerinde model başarısının önemli göstergelerinden biridir.

- Doğruluk (Accuracy) = $(TP + TN) / (TP + TN + FP + FN) = (3 + 2) / (3 + 2 + 1 + 0) = 5 / 6 = \%83,3$
- Kesinlik (Precision) = $TP / (TP + FP) = 3 / (3 + 1) = 3 / 4 = \%75$
- Duyarlılık (Recall) = $TP / (TP + FN) = 3 / (3 + 0) = 3 / 3 = \%100$

Model, test verileri üzerinde uygulanarak toplam 6 örnekten 5'ini doğru, yalnızca 1'ini yanlış sınıflandırmıştır. Confusion matrix'e göre modelin "normal" ve "attack" sınıflarını ayırt etme kabiliyeti oldukça yüksektir.

Bu analizde True Positive (TP) = 3, yani model saldırı olan 3 veriyi doğru şekilde tahmin etmiştir. True Negative (TN) = 2, yani normal olan 2 veri de doğru şekilde tanınmıştır. False Positive (FP) = 1, yani sistem 1 normal veriyi yanlışlıkla "saldırı" olarak sınıflandırmıştır. False Negative (FN) = 0, yani hiçbir saldırı durumu "normal" olarak sınıflandırılmamıştır. Bu durum, güvenlik açısından oldukça olumlu bir göstergedir. Genel olarak, modelin doğruluk oranı %83,3, kesinlik oranı %75, ve duyarlılık oranı %100 olarak hesaplanmıştır. Özellikle duyarlılık oranının %100 olması, modelin saldırı durumlarını kaçırmadan doğru şekilde tespit edebildiğini ve güvenlik açısından etkili bir çözüm sunduğunu göstermektedir. Ancak kesinlik oranının %75 olması, bazı normal verilerin yanlış alarm olarak sınıflandırıldığını da ortaya koymaktadır.



Şekil 9. Random Forest ve Naive Bayes Performans Karşılaştırması

Kaynak: Yazar tarafından oluşturulmuştur.

1. Doğruluk (Accuracy)

- Random Forest algoritmasının doğruluk oranı: %66,7
- Naive Bayes algoritmasının doğruluk oranı: %83,3

Doğruluk oranı, modelin doğru sınıflandırdığı örneklerin toplam örneğe oranıdır. Bu bağlamda, Naive Bayes algoritması daha yüksek bir doğruluk oranı sunmuştur. Bu, modelin genel performansının daha istikrarlı olduğunu göstermektedir.

2. Kesinlik (Precision)

- Random Forest: %66,7
- Naive Bayes: %75

Kesinlik değeri, modelin saldırı olarak sınıflandırdığı veriler arasında gerçekten saldırı olanların oranıdır. Naive Bayes modeli daha yüksek kesinlik değeriyle, yanlış alarm oranını Random Forest'a göre daha düşük tutmuştur. Bu da özellikle yanlış pozitiflerin kritik olduğu uygulamalarda Naive Bayes'i daha avantajlı hale getirmektedir.

3. Duyarlılık (Recall)

- Random Forest: %66,7
- Naive Bayes: %100

Duyarlılık, modelin gerçek saldırı verilerini ne kadar iyi yakalayabildiğini gösterir. Bu metrik açısından Naive Bayes %100'lük başarıya ulaşmış, yani test setindeki tüm saldırı verilerini doğru şekilde tespit etmiştir. Random Forest ise bazı saldırı durumlarını gözden kaçırmıştır. Bu fark, özellikle siber saldırıların tespitinde hayati önem taşımaktadır; çünkü bir saldırının gözden kaçması büyük güvenlik açıklarına yol açabilir.

4. Confusion Matrix İncelemesi

- Random Forest Confusion Matrix'inde toplam 6 örnekten sadece 4'ü doğru sınıflandırılmıştır. 2 örnek ise hatalıdır.
- Naive Bayes Confusion Matrix'inde ise 6 örnekten 5'i doğru sınıflandırılmıştır. Sadece 1 hata yapılmış, o da bir normal durumun yanlışlıkla saldırı olarak değerlendirilmesidir (false positive).

Bu bağlamda, Naive Bayes algoritması saldırı tespiti açısından daha hassas davranmakta ve saldırıları gözden kaçırmamaktadır. Random Forest modeli, karar ağaçları kombinasyonu ile daha güçlü modelleme imkanı sunsa da, küçük ve dengeli veri setlerinde zaman zaman aşırı öğrenme (overfitting) veya yanlış sınıflama riski taşımaktadır. Buna karşın, Naive Bayes modeli istatistiksel temelli basit yapısıyla özellikle küçük veri setlerinde etkili sonuçlar verebilmektedir.

Bu çalışmanın sonuçları göstermiştir ki:

- Küçük örneklemli ve dengeli veri setlerinde Naive Bayes, yüksek duyarlılığı ve doğruluk oranıyla daha başarılıdır.
- Ancak daha büyük ve çeşitli veri kümelerinde, hiperparametre optimizasyonları yapılmış Random Forest modelleri daha stabil ve genelleyici sonuçlar verebilir.

Analizin Faydaları ve Katkısı:

- Saldırı Tespiti İçin Etkili İki Yöntem Test Edildi: Gerçek ağ trafiği verilerine dayalı olarak, bağlantıların “normal” ya da “saldırı” içerikli olup olmadığını belirlemek amacıyla hem Naive Bayes hem de Random Forest algoritmaları kullanılarak iki ayrı model geliştirildi. Böylece farklı algoritmaların saldırı tespiti üzerindeki başarıları karşılaştırılmış oldu.
- Makine Öğrenimi Süreci Uygulamalı Olarak Gösterildi: Veri yükleme, ön işleme, eğitim/test bölme, model kurma ve başarı değerlendirme gibi adımlar sırasıyla gerçekleştirildi. Bu süreç, özellikle veri bilimi veya siber güvenlik alanında öğrenim gören öğrenciler için adım adım örnek bir uygulama sağladı.
- Modellerin Küçük Veri Setinde Performansı Görüldü: Sadece 20 satırlık küçük ve dengeli bir veri seti kullanılarak bile, modellerin temel başarı oranları, sınıflandırma gücü ve etki eden sütunlar net olarak gözlemlendi. Bu da küçük ölçekli çalışmaların da anlamlı sonuçlar verebileceğini gösterdi.
- Algoritmaların Güçlü ve Zayıf Yönleri Karşılaştırıldı: Random Forest modeli %83.3 doğruluk oranı ile daha yüksek başarı gösterirken, Naive Bayes modeli %66.7 doğruluk ile daha sade fakat temel düzeyde etkili sonuç verdi. Böylece hangi durumda hangi algoritmanın tercih edileceği konusunda daha bilinçli kararlar alınabilir hale geldi.
- Gerçek Hayat Uygulamaları İçin Yol Gösterici Oldu: Her iki modelin uygulaması, teorik bilgilerin pratikte nasıl hayata geçirileceğini gösterdi. Bu durum hem akademik çalışmalar için hem de siber güvenlik sistemlerinin geliştirilmesi açısından önemli bir örnek teşkil etti.

- İleri Seviye Analizler İçin Sağlam Temel Oluşturuldu: Bu çalışma, ileride daha büyük NSL-KDD veri kümeleriyle, farklı algoritmalar ve parametre ayarlamalarıyla yapılacak analizler için başlangıç niteliğinde sağlam bir altyapı sundu. Model kıyaslamaları, doğruluk artırma ve optimizasyon çalışmaları bu yapı üzerinden kolayca geliştirilebilir.

2.4. Makine Öğrenimi İçin Turing Testi

1950 yılında İngiliz matematikçi ve bilgisayar biliminin öncülerinden Alan Turing, 'makinelere düşünebilir mi?' sorusunu gündeme getirmiştir. Turing, 'Computing Machinery and Intelligence' başlıklı makalesinde, bir makinenin düşünme yetisine sahip olup olmadığını değerlendirmek amacıyla daha sonra 'Turing Testi' ya da 'taklit oyunu' olarak anılacak yöntemi önermiştir. Test, bir erkek ve bir kadının, hangisinin hangisi olduğunu tahmin etmesi gereken bir sorgulayıcıdan izole edilmesini içeren Viktorya tarzı bir oyunun uyarlamasına dayanıyordu. Turing'in versiyonunda, bilgisayar programı katılımcılardan birinin yerini alıyordu ve sorgulayıcı hangisinin bilgisayar hangisinin insan olduğunu belirlemek zorundaydı. Sorgulayıcı makine ile insan arasındaki farkı söyleyemezse, bilgisayarın düşündüğü veya "yapay zekaya" sahip olduğu düşünülüyordu. Turing'in testi, ilk dijital bilgisayarların geliştirilmesinden sadece birkaç yıl sonra ve "yapay zeka" terimi henüz ortaya atılmadan önce geldi. O zamanlar, sibernetik ve otomat teorisi dahil olmak üzere "düşünen makineler" alanı için çeşitli isimler vardı. 1956'da, Turing'in ölümünden iki yıl sonra, Dartmouth College'da profesör olan John McCarthy, düşünen makineler hakkında fikirleri açıklamak ve geliştirmek için bir yaz atölyesi düzenledi ve proje için "yapay zeka" adını seçti. Yapay Zeka'0000'nın (YZ) bir araştırma alanı olarak kuruluş anı olarak kabul edilen Dartmouth konferansı, "makinelere dili nasıl kullanacağını, soyutlamalar ve kavramlar oluşturacağını, artık insanlara ayrılmış türden sorunları nasıl çözeceğini ve kendilerini nasıl geliştireceğini" bulmayı amaçlıyordu. Sonraki birkaç yıl içinde, araştırmacıların dama ve satranç gibi oyunlar oynamak gibi uzman seviyesinde bilgi gerektiren görevleri yerine getirmek için teknikleri araştırmasıyla alan hızla büyüdü. 1960'ların ortalarına gelindiğinde, Amerika Birleşik Devletleri'ndeki yapay zeka araştırmaları Savunma Bakanlığı tarafından yoğun bir şekilde finanse ediliyordu ve dünya çapında YZ laboratuvarları kurulmuştu. Aynı

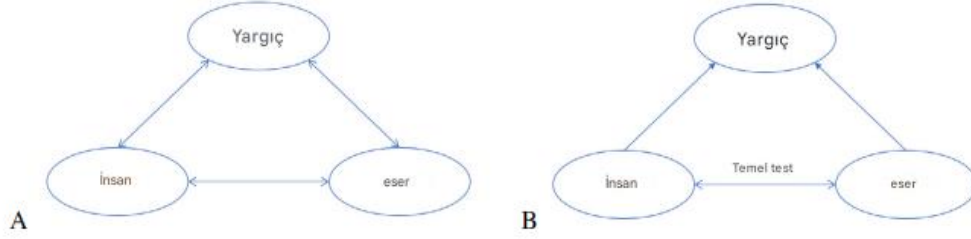
sıralarda, Lawrence Radyasyon Laboratuvarı, Livermore, Sidney Fernbach başkanlığındaki Matematik ve Hesaplama Bölümü içinde kendi Yapay Zeka Grubunu kurdu. Programı yürütmek için Livermore, yapay zeka öncüsü Marvin Minsky'nin eski bir öğrencisi olan MIT mezunu James Slagle'ı işe aldı. Çocukluğundan beri kör olan Slagle, MIT'den matematik alanında doktora derecesi aldı. Eğitimini sürdürürken Slagle, Beyaz Saray'a davet edildi ve burada, Başkan Dwight Eisenhower'dan, Körler İçin Kayıt A.Ş. adına, olağanüstü akademik çalışmalarından dolayı bir ödül aldı. 1961 yılında Slagle, tez çalışması kapsamında SAINT (Symbolic Automatic INtegrator) adlı bir program geliştirmiştir. Bu program, insan uzmanların karar verme süreçlerini taklit edebilen bilgisayar sistemleri olan ilk 'uzman sistemler'den biri olarak kabul edilmektedir.

SAINT, kalkülüsteki integrallerin manipülasyonunu içeren temel sembolik entegrasyon problemlerini, üniversite birinci sınıf öğrencisi seviyesinde çözebiliyordu. Program, 54'ü MIT birinci sınıf kalkülüs sınavlarının finalinden alınmış 86 problemten oluşan bir set üzerinde test edildi. SAINT, soruların ikisini hariç hepsini çözmeyi başardı. Livermore'da Slagle ve grubu, bilgisayar programlarına problem çözme durumlarına yaklaşımlarında hem tümdengelimli hem de tümevarımlı akıl yürütmeyi öğretmeyi amaçlayan birkaç program geliştirmek için çalıştı. Bu programlardan biri olan multiple (öğrenen çok amaçlı teori kanıtlayan huristik program), geometri ve kalkülüs problemlerinden dama gibi oyunlara kadar çok çeşitli görevlerde "sonra ne yapılacağını" öğrenme esnekliğiyle tasarlandı. Slagle'a göre, yapay zeka araştırmacıları artık zamanlarını Turing'in "makinelere düşünebilir mi?" sorusunun artılarını ve eksilerini tekrar tekrar ele alarak harcamıyorlardı. Bunun yerine, "düşünmenin" bir "ya da ya" durumundan ziyade bir süreklilik olarak görülmesi gerektiği görüşünü benimsediler. Bilgisayarların az düşünüp düşünmediği veya hiç düşünmediği açıktı; gelecekte gelişebilecekleri açık bir soru olarak kaldı. Ancak, yapay zeka araştırmaları ve ilerlemesi patlamalı bir başlangıçtan sonra yavaşladı; ve 1970'lerin ortalarına gelindiğinde, keşifsel araştırmanın yeni yolları için hükümet fonlaması neredeyse tükenmişti. Benzer şekilde, Laboratuvar'da Yapay Zeka Grubu dağıtıldı ve Slagle çalışmalarını başka bir yerde sürdürmek üzere yola çıktı. Sonraki yıllarda alanın önemi azaldı ve azaldı; ancak 1990'ların sonu ve 2000'lerin başında, yapay zeka araştırması, çok yönlü, tamamen zeki makineler yaratma orijinal

hedefi yerine, belirli sorunlara belirli çözümler bulmaya odaklanarak tekrar ön plana çıktı. Günümüzde daha hızlı bilgisayarlar ve büyük miktarda veriye erişim, makine öğrenimi ve veri odaklı derin öğrenme yöntemlerinde ilerlemeler sağladı. Şu anda Lawrence Livermore Ulusal Laboratuvarı, makine öğrenimi ve derin öğrenme dahil olmak üzere çeşitli veri bilimi alanlarına odaklanmıştır (Korukonda, 2003).

LLNL, 2018 yılında Laboratuvarın çeşitli veri bilimi disiplinlerini (yapay zeka, makine öğrenimi, derin öğrenme, bilgisayar görüşü, büyük veri analitiği ve diğerleri) tek bir çatı altında bir araya getirmek için Veri Bilimi Enstitüsünü (DSI) kurmuştur. DSI ile Laboratuvar, ülkenin veri bilimi yeteneklerinin en son durumunu ilerletmek için veri bilimi iş gücünü, araştırmayı ve erişimi oluşturmaya ve güçlendirmeye yardımcı olmaktadır. 20 ci yüzyılda, İngiliz matematikçi Alan Turing, "Hesaplama makineleri ve zeka" adlı eserinde açıklandığı gibi, Yapay Zeka (YZ) olasılığı ve hangi özelliklere sahip olması gerektiği sorusuyla aktif olarak ilgilendi. (Turing, 1950) Aslında, bu konuyu ele alan ilk kişiydi ve buna göre, bilimsel toplumun gelecekte üzerinde çalışacağı belirli bir temel oluşturdu. Turing testi, bilime yakın olmayan kişiler tarafından bile günümüze kadar bilinen en ünlü ve tartışmalı deneysel paradigmalardan biridir. "Korukonda'nın (2003) sert eleştirilerine rağmen, bu yaklaşım hala zeki sistem araştırmalarının nihai amacıyla ilişkilendirilmeye devam etmektedir. Ancak, gerçekten ulaşabileceğimiz en iyi hedef bu mu – yoksa başka bir olasılık var mı? Burada, umut verici olabilecek alternatif bir bakış açısı sunulmaktadır. Turing testinin genel olarak kabul edilen yorumu şu şekildedir. Bir insan değerlendirici, biri insan, diğeri bilgisayar programı gibi bir eser olan iki katılımcıyla etkileşime girer. "Bu değerlendirme modelinde, değerlendirmeyi yapan kişi etkileşim yoluyla karşısındaki varlığın bir insan mı yoksa bir yapıt mı olduğunu anlamaya çalışır. Her iki tarafın amacı da, değerlendiriciyi kendisinin insan olduğunu, karşısındakinin ise yapıt olduğunu düşündürmeye ikna etmektir. Eğer değerlendirici rastgele tahmin seviyesinde bir yargıya varırsa, bu durumda yapıtın testi geçtiği kabul edilir ve bu anlamda insan benzeri bir iradeye sahip olduğu söylenebilir. Başlangıçta önerilen modelde, değerlendirme yalnızca yazılı metinlerle sınırlıydı; ancak günümüzde, serbest hareket edebilen robot sistemleri de dahil olmak üzere daha gelişmiş biçimleri

tartışılmaktadır (Graham, 2004). Bu genel model yukarıda anlatıldığı şekilde bir diyagramla özetlenebilir.



Şekil 10. Turing Testinin Genelleştirilmiş Paradigmaları

Kaynak: Graham-Rowe, D. (2004). Turing test is a total turn-off for robots: The thinking machine's challenge of choice is something a little more down to earth. *New Scientist*, 183(2461), 22.

Turing testinin başarılı bir oturumu, başka bir takip oturumunda başarıyı garantileyemez ve rastgele bir sonuçtan kaynaklanabilir. Bu nedenle, başarılı bir Turing testi yeni bir genel gerçek oluşturmaz. Gerçekten de, bu durumdaki başarı olumsuz bir sonuçtur (bize sıfır hipotezi H_0 'ın reddedilemeyeceğini söyler) ve muhtemelen yeterince büyük bir örneğe dayalı başka bir test tarafından çürütülebilir. Öte yandan, birisinin Turing testinin belirli bir sınırlı versiyonuna dayanarak ortalama bir insandan daha sık bir insan olduğu yargılanacak bir eser inşa ettiğini hayal edin. Daha sonra, bu gerçek güvenilir bir şekilde deneysel olarak belirlendiğinde (H_0 reddedilir), bir gerçek olarak kalacaktır. İlk başta, bir eserin herhangi bir insandan daha "insana benzer" olabileceği fikri mantıksal bir çelişki gibi görünüyordu. Gerçekten de, bir şey bir şeye kendisinden daha çok nasıl benzeyebilir? Ancak bu izlenim, fikrin yanlış anlaşılmasından kaynaklanıyor. Test, tek bir ölçüte dayanır: Katılımcının, bir eserin inanılabilirliği veya insan benzerliği hakkındaki öznel yargısı, belirli bir varlığın insan olup olmadığı sorusuna verilen 'evet-hayır' cevabı ile ifade edilir. Turing Testi meydan okumasındaki amaç, pozitif yanıt olasılıklarını bir insan ve bir bilgisayar için istatistiksel olarak eşit hale getirmektir. Buna karşılık, Overman Meydan Okuması'ndaki amaç (adını böyle koymayı öneriyoruz) pozitif yanıt olasılığını bir bilgisayar için bir insana kıyasla önemli ölçüde daha yüksek hale getirmektir. Mantıksal olarak mümkün mü? Evet, eğer bir eser inandırıcılık açısından ortalamanın üzerinde puan alan bir insana özgü davranışlar sergileyebiliyorsa önce belirli terimleri ve koşulları tanımlamak gerekir. Öncelikle, biyolojiden destekleyici bir argüman ele alalım (Gentry, 2012). 1950'de Tinbergen, dişi ringa martılarının gagalarının altında

kırmızı bir nokta olduğunu ve martı yavrularının beslenmek istediklerinde gagalamaları için bir hedef olduğunu buldu. Yavru kırmızı daireyi gagalar, ağzını açar, anne biraz yiyecek çıkarır ve martı kusmuğu tükenene kadar veya gerektiği kadar tekrarlar. Ancak, yavruların onları besleyen gerçek anneleri olması durumunda çok da rahatsız olmadıkları ortaya çıktı. Bir sırtığın üzerindeki bir kafa, vücutsuz bir gaga veya bir çubuğun üzerindeki kırmızı bir nokta, yavrunun yiyecek için gagalamasını tetikleyecektir.

Bir avatar, belirli bir ortamda bir aktörün somutlaşmış hali olarak anlaşılacaktır. Örneğin, duyuları olan ve belirli eylemleri gerçekleştirebilen sanal veya fiziksel bir robot olabilir. Daha genel olarak, bir avatar, bir aktörün bir ortama gömülmesinin bir biçimi olarak anlaşılacaktır. Avatar ile aynı arayüzün insan katılımcılar ve yapay (sanal) aktörler tarafından kullanılması gerektiği varsayılırken, avatarın kendisi onu kontrol eden aktörün kimliğini ortaya çıkarmaz. Bu nedenle, tüm aktörler anonim kalır ve doğaları (insan olsun veya olmasın) kontrol eden aktörden bağımsız olan avatarın özelliklerinden türetilmez. Buradaki aktör terimi, belirli bir ortamda bir avatarı kontrol eden akıllı bir aracı (bir insan, bir eser veya sanal bir varlık) ifade eder. Özellikle, bu aracı ortamdan duyuşal girdi almalı ve ihtiyaçlarını karşılamak için ortamda eylemler gerçekleştirebilmelidir. Buradaki ortam terimi, sanal gerçeklik simülatöründe veya bir robotik laboratuvarında veya hayali bir dünyada düzenlenen ve olayların bir bilgisayar monitöründe görünen metin aracılığıyla kendini gösterdiği belirli deneysel ortamları ifade eder. Avatarlar, ortamın parçaları olarak kabul edilir. Deneysel paradigma veya paradigma terimi, avatarları kontrol eden aktörlerin davranışlarını ve etkileşimlerini belirleyen belirli başlangıç koşullarının, prosedürlerin, hedeflerin ve kuralların dayatıldığı bir ortam sınıfını ifade eder. Aşağıdaki değerlendirme, avatarları kontrol eden aktörlerin bulunduğu paradigmalara sınırlıdır.

- Anonim kalmak,
- Birbirleriyle sosyal etkileşimlerde bulunmak ve
- Belirli hedefleri takip etmek.

Ek olarak, her aktör için, bir insan katılımcı (paradigmaya bağlı olarak başka bir aktör veya harici bir yargıç olabilir) avatarın davranışına bakarak, bu aktörün insan

aynı sıklıkta seçilecek şekilde formüle edebiliyor. Ancak bunun hala "Taklit Oyunu" olarak kalmasıyla ilişkili bir sorun var. Makine, yalnızca bir kişinin nasıl davrandığına bakıyor ve ardından "bağımsız" kararlar almak yerine onun davranışını taklit ediyor. Bu sorunun bir kısmı testin nasıl formüle edildiğiyle ilgilidir. Test, makinenin sorunun yalnızca bir tarafını görmesini sağlar, diğer kişiliklerle, makinelerle ilgili kararlar almak zorunda değildir. Bu, Makinelerin bir insan gibi, hatta daha iyi düşünmesini sağlayabilecek bir şey olabilir. Örneğin, Turing testinde gözlemciyi bir makineyle değiştirmeyi deneyelim. Şimdi makinenin kendi başına nasıl kazanacağını öğrenmesi gerekecek. Sonuçta, hiç kimse makinenin kendisinden daha iyi nasıl düşündüğünü anlayamaz. Makinenin bir rol modeli olmadığında, şartlı olarak kendi başına kararlar alacak, kendisi olacak ve kimseyi taklit etmeyecektir. Ancak, bu modelde, bunun sonucunda alt seviyedeki makinenin Gözlemci Makinesi için kişiden pratik olarak ayırt edilemez olma tehlikesi vardır. Bu nedenle, gözlemcinin bir kişi olduğu durumda alt seviyeli makineler için testler yapmak faydalı olacaktır. Böylece Makine alt seviyede katılacak ve aynı zamanda gözlemci olarak katılacaktır.

Makine, insan gibi düşünmek ile bir makine gibi düşünmek arasında bir denge kuracaktır. Böylece, makine gözlemcisi, bir insanı taklit etmeden, yalnızca makinenin anlayabileceği bir dilde düşünerek insan benzeri özellikler sergileyecektir. İnsan davranışını taklit etme yeteneği açısından benzerlik gösterecektir, ancak bu benzerlik her açıdan olmayacaktır. Yeni test paradigması "Üst İnsan Mücadelesi" olarak adlandırılabilir, çünkü fikir ortalama bir insandan daha insani bir eser ve sonunda - tüm insanlardan daha insani bir eser yaratmaktır. Turing Testi mücadelesine önerilen alternatif, daha fazla seviye bulunamadığı zaman sona eren bir seviye dizisinden oluşur. Birinci seviye mücadelesi önemsiz olabilir veya olmayabilir. Genel mücadele, Turing testini geçmekten daha zordur. Daha güçlü gereksinimler belirleyerek, bu gereksinimlerin ötesine geçilebilir. Turing testiyle karşılaştırıldığında, Üst İnsan mücadelesini her seviyede geçmek, sıfır hipotezini reddetmek anlamına gelir. Üst İnsan, insan zekasının ötesinde bir kapasiteye sahip olmayı ifade eden bir kavramdır ve bu, insanın düşünsel ve bilişsel sınırlarını aşan bir zeka seviyesini tanımlar. Sosyal olarak kabul edilebilir zeki ajanlara ve robotlara olan talebin artması nedeniyle, Üst İnsan mücadelesini geçmek, önemli bir teknolojik atılım olarak kabul edilecektir.

3. MAKİNE ÖĞRENİMİNİN GELİŞİMİ VE GELECEK PERSPEKTİFLERİ

Makine öğreniminin gelişimi ve geleceği ele alınarak Amazon ürün yorumları üzerinden duygu analizi, derin öğrenme alanındaki güncel trendler, karşılaşılan mevcut sorunlar ve bu sorunlara yönelik çözüm yaklaşımları incelenmiştir.

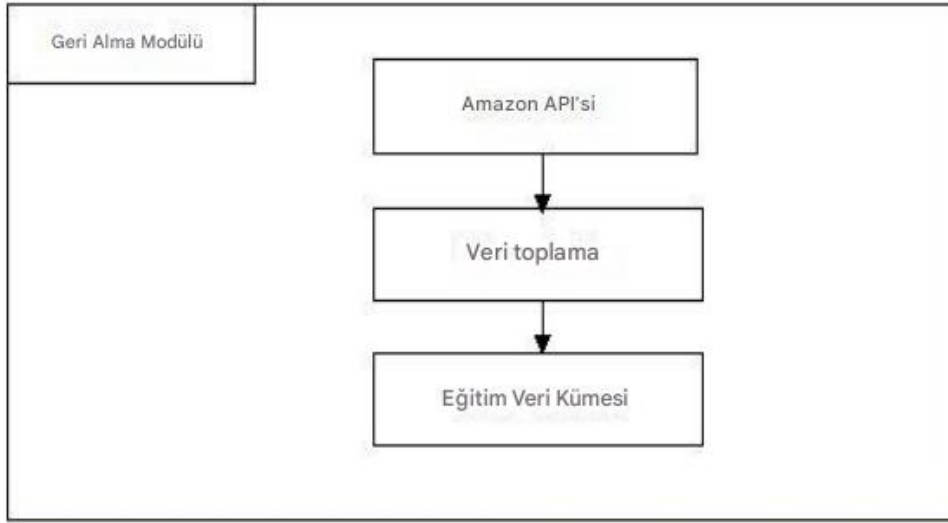
3.1. Makine Öğrenme Algoritması Kullanılarak Amazon Ürünlerinin

Duygu Analizi

Son yıllarda, duygusal analiz (sentiment analysis) çevrimiçi ürün satışı yapan platformalarda müşteri deneyimlerini analiz etmek için yaygın bir şekilde kullanılmaktadır. Belirli ürünü kullanan kullanıcı sayısı artmış ve bu da sektörün ürünün özelliklerini geliştirme ve iyileştirmesine neden olmuştur. Günümüzde web siteleri, bloglar, çevrimiçi alışveriş kullanan birçok kullanıcı kullandıkları ürünleri inceleme eğilimindedir. Bu incelemeler, ürün aramaları sırasında diğer kullanıcılar tarafından dikkate alınmıştır. Bu nedenle sektör, kullanıcıların incelemelerine dayanarak kullanıcının aradığı doğru ürünü sunmanın kökenini duygusal analiz kavramını kullanarak bulmuştur. Duygusal Analiz, inceleme koleksiyonlarının dikkate alındığı incelemelerin analiz edildiği, işlendiği ve kullanıcıya önerildiği veri analizi kavramıdır. Verilen incelemeler daha uzundur ve birkaç paragraf içerikten oluşur. Başlangıçta, durdurma sözcükleri, fiiller, noktalama işaretleri ve bağlaçlar gibi istenmeyen verileri kaldırmak için bu incelemeler önceden işlenmelidir. Ön işleme tamamlandıktan sonra, eğitilen veri kümesi Naive Bayes ve SVM algoritması kullanılarak sınıflandırılır (Bancken, 2014). Bu mevcut algoritmalar, değeri olmayan doğruluğu sağlamıştır. Bu nedenle, verilen incelemelerin doğruluğunu artırmak için bir topluluk yaklaşımı uygulanmıştır. Topluluk, iki veya daha fazla algoritmayı birleştirerek ve kullanılan her algoritma için oy referansına göre mod değerini hesaplayarak yapılan bir sınıflandırma yaklaşımıdır. Bu nedenle, mevcut algoritmalarından daha yüksek doğruluk elde edebilmek için Topluluk (Ensemble) Yöntemi geliştirilmiştir. Topluluk yöntemi, birden fazla sınıflandırıcının (Naive Bayes ve SVM) tahminlerini birleştirerek çoğunluk oyu (majority voting) prensibiyle nihai kararı vermektedir.

Çeşitli web sitelerinden gelen metin incelemelerinden kamuoyunun fikrini hesaplamalı yöntemlerle çıkarmak ve analiz etmek günümüzde trend olan bir araştırma çalışması haline geliyor. Bu araştırmalarda elde edilen geleneksel ve cüruf dillerine dayalı bir dizi kök sözcük kullanılıyor. Duygu ifadeleri bu yaklaşımlar aracılığıyla tanınıyor. Genel olarak, inceleme veri kümeleri denetlenen ve denetlenmeyen veri kümeleri olarak sınıflandırılır. Denetlenen öğrenmede, araştırmacılar etiketli bir veri kümesine sahiptir, yani araştırmacılar girdilerin ve çıktılarının değerlerine sahiptir. Makine öğrenmesinde, model, girdi (input) verileri ile çıktı (output) verileri arasındaki ilişkileri tanımlayan matematiksel yapıdır. Farklı algoritmalar, bu modeli öğrenebilmek için kullanılmaktadır. Makine öğreniminde, verilerin bir modelini elde etmemizi sağlayan birçok farklı algoritma vardır. Araştırmacıların aradığı amaç ve makine öğrenimini nasıl kullanabileceğimiz, modeli öğrendikten sonra yeni bir girdi verildiğinde çıktıyı tahmin etmektir. Denetlenmeyen öğrenmede, etiketlenmiş verilerimiz yoktur. Girdilerimiz olduğunu ancak çıktılarımız olmadığını söyleyebiliriz. Yani, amaç verilerimizde bir örüntü belirlemektir. Aynı gruba veya çıktıya ait olduğunu düşündüğümüz grupları veya kümeleri bulabiliriz. Burada bir model elde etmemiz gerekir. Ayrıca, yine aradığımız amaç, yeni bir girdi verildiğinde çıktıyı tahmin edebilmektir. Duygu Analizi, kullanıcıların incelemeler aracılığıyla diğer kullanıcılarla etkileşim kurmasına yardımcı olan kullanıcı değerini güçlendirir. Bu incelemeler ve yanıtlar, diğer kullanıcılar tarafından ürüne gidip gitmemeye karar vermede kullanılır. Duygu analizinde, incelemeleri sınıflandırmada Naive Bayes, Lojistik regresyon gibi birçok mevcut yaklaşım olduğundan, bu yaklaşımlar etkilidir ancak doğruluk açısından etkili değildir. Doğruluğu artırmak için bu alanda birçok araştırma çalışması yürütülmektedir.

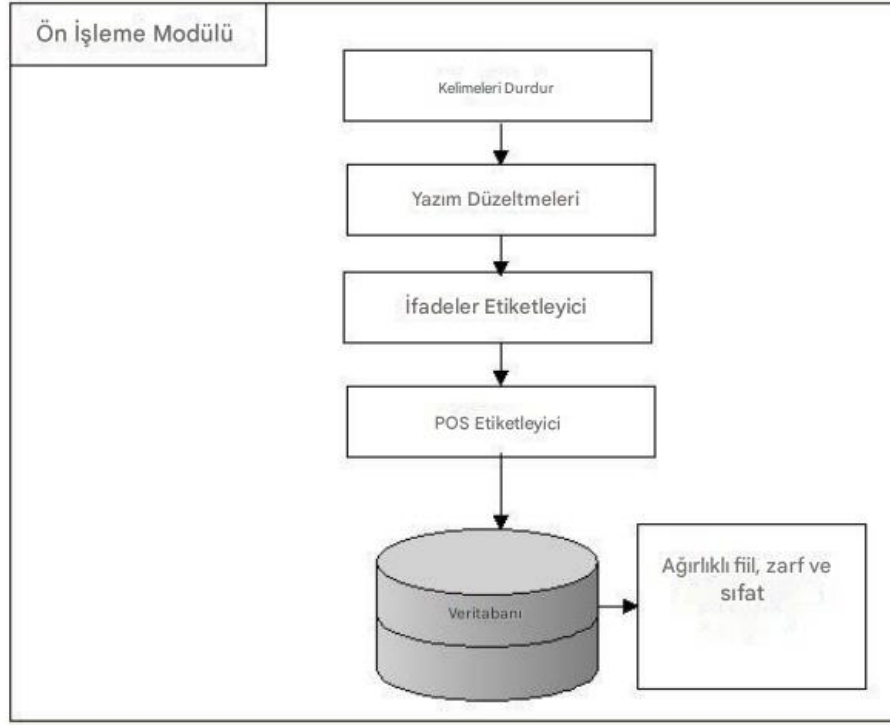
Bu çalışmada, Amazon ürünlerine ilişkin metin incelemeleri toplanarak bir eğitim veri seti oluşturulmuştur. Veriler, Amazon API'si kullanılarak veya veri kazıma teknikleriyle toplanmış, daha sonra ön işleme adımları (durdurma sözcüklerinin, noktalama işaretlerinin temizlenmesi vb.) uygulanmıştır.



Şekil 12. Arama Modülü

Kaynak: Integrated workflow of machine learning and deep learning methods. Patil, D., Rane, N. L., Desai, P., & Rane, J. (2024). Trustworthy Artificial Intelligence in Industry and Society, 28–81.

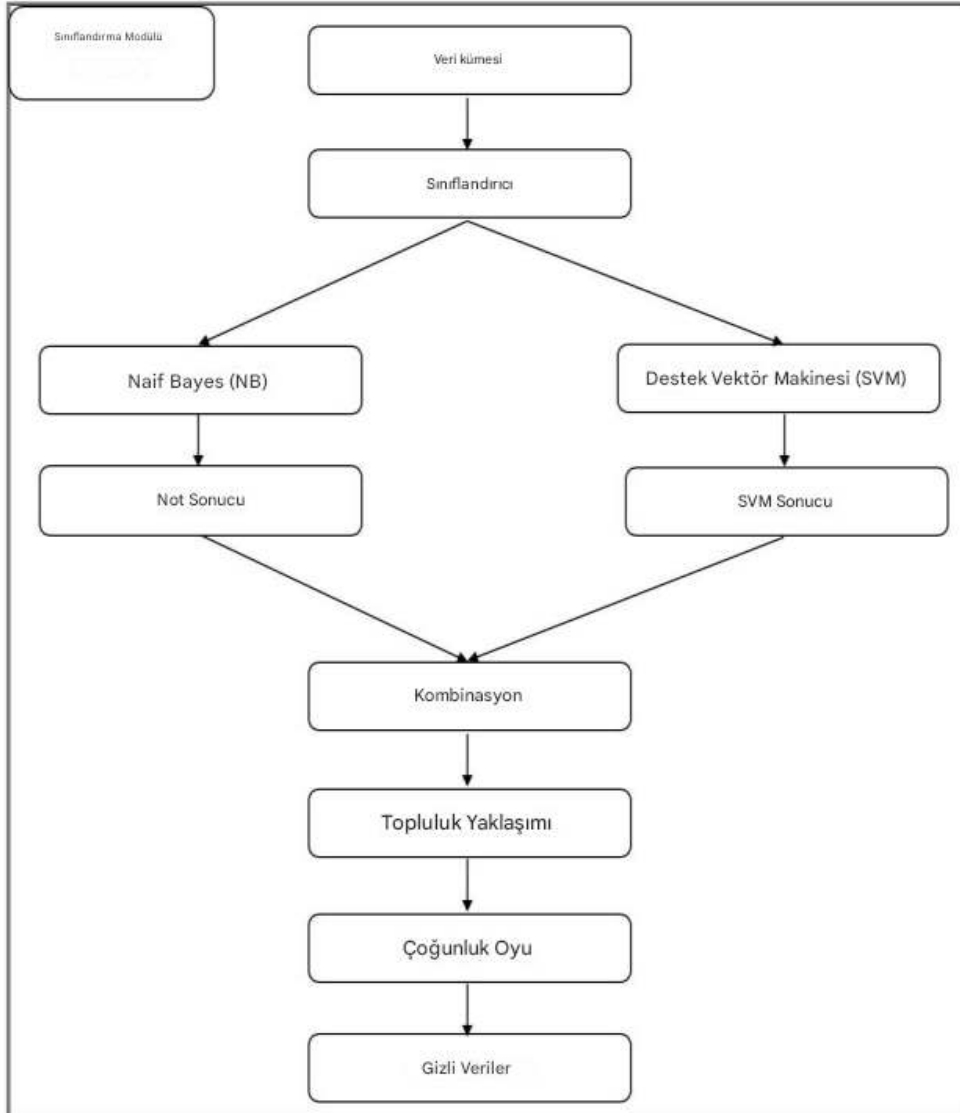
Veriler toplandıktan sonra, istenmeyen ifadeleri, sözcükleri, sembolleri kaldırmak için ön işleme tabi tutulmalıdır (Şekil 12). Bir analiz çalıştırmadan önce verilerin temsili ve kalitesi ilk ve en önemli şeydir. Başlangıçta @, hashtags# (Durdurma Sözcükleri) gibi semboller kaldırılacak, ardından yazım denetimi ve ifade etiketlemesi yapılacaktır. Bir sonraki aşama POS (Kelime Türleri) etiketlemesiyle devam etti. Burada, fiiller, zarflar ve sıfatlar gibi tüm gerekli POS sözcükleri ve ifadeleri çıkarılır. Bu sözcükler ağırlıklı puanlarla dikkate alınacaktır.



Şekil 13. Ön İşleme Modülü

Kaynak: Integrated workflow of machine learning and deep learning methods. Patil, D., Rane, N. L., Desai, P., & Rane, J. (2024). Trustworthy Artificial Intelligence in Industry and Society, 28–81.

Burada, Eğitim verileri Naive Bayes, SVM sınıflandırıcıları kullanılarak toplanacak ve sınıflandırılacaktır. Genellikle, Gözetimli ve Gözetimsiz gibi iki tür sınıflandırma vardır. Gözetimli sınıflandırma etiketlerle sağlanacaktır. Eğitim verileri bir dizi eğitim örneğinden oluşurken, gözetimsiz sınıflandırma herhangi bir etiket kümesiyle sağlanmayacaktır. Naive Bayes ve SVM gibi bazı gözetimli sınıflandırıcılar önerdik. Naive Bayes, sınıf etiketlerinin sonlu bir kümeden seçildiği, özellik değerlerinin vektörleri olarak temsil edilen sorun örnekleri için sınıf etiketleri atamak üzere sınıflandırıcılar oluşturmak için basit bir tekniktir. Destek Vektör Makinesi (SVM) tarafından durum doğrudan görünmez, ancak duruma bağlı çıktı görünür. Her durum, olası çıktı belirteçleri üzerinde bir olasılık dağılımı alır. Bu nedenle, bir SVM tarafından üretilen belirteç dizisi, durum dizisi hakkında bazı bilgiler verir. Son olarak, yukarıdaki tüm yöntemleri birleştirerek bir Topluluk yaklaşımı uygulanır. Sınıflandırma yapıldıktan sonra, aralarında çoğunluk oyu (Doğruluk) alınır ve sınıflandırılan, mevcut sistemden daha iyi doğruluk üretir.



Şekil 14. Sınıflandırma Modülü

Kaynak: Integrated workflow of machine learning and deep learning methods. Patil, D., Rane, N. L., Desai, P., & Rane, J. (2024). Trustworthy Artificial Intelligence in Industry and Society, 28–81.

Naive Bayes sınıflandırıcısı basit bir olasılık tabanlı algoritmadır. Bayes teoremini kullanır ancak örneklerin birbirinden bağımsız olduğunu varsayar ki bu pratik dünyada gerçekçi olmayan bir varsayımdır. Naive Bayes sınıflandırıcısı karmaşık gerçek dünya durumlarında iyi çalışır.

$$P(H | X) = P(X | H) P(H) / P(X)$$

- P, parantez içindeki değişkenlerin olasılığını ifade eder.

- $P(H)$, H'nin marjinal olasılığının ön olasılığıdır; X'te bulunan bilgileri henüz hesaba katmamış olması anlamında ön olasılığdır.
- $P(H/X)$, H'nin koşullu olasılığıdır; X verildiğinde, ayrıca arka olasılık olarak da adlandırılır; çünkü X olayının sonucunu zaten dahil etmiştir.
- $P(X/H)$, H verildiğinde X'in koşullu olasılığıdır.
- $P(X)$, normalde kanıt olan X'in ön veya marjinal olasılığıdır.

Ayrıca şu şekilde de gösterilebilir. Arka=olasılık * ön normalleştirme sabiti

$P(X/H)/P(X)$ oranına ayrıca standartlaştırılmış olasılık denir.

ADIMLAR;

- Adım 1: Verileri eşit büyüklükte k alt kümeye bölün.
- Adım 2: Her alt kümeyi sırayla test için, kalanı eğitim için kullanın.
- Adım 3: Verilerin toplam pozitif ve toplam negatif olasılığını bulun.
- Adım 4: Bir veri kümesindeki her kelime için, pozitif olasılığını ve negatif olasılığını bulur.
- Adım 5: Tüm pozitif ve tüm negatif olasılık değerlerini çarp.
- Adım 6: Bu değerleri toplam olasılık değerleriyle çarp.
- Adım 7: Olasılığın daha anlamlı değere sahip olduğu duyguyu seç.

SVM, regresyon analizinde ve farklı form ve kategorilere sahip bir dizi eğitim kümesi için sınıflandırmada kullanılır (Fang., vd 2015).

Eğitim aşamasında hesapladığımız yalnızca eğitim noktalarının bir alt kümesinden benzerliği hesaplar. Bu nedenle, eğitim aşamasında, herhangi bir test noktasının benzerliğini hesapladığımız bir eğitim noktaları alt kümesini sabitleriz. Bu seçilen eğitim noktalarına destek vektörleri denir çünkü yalnızca bu noktalar bir test noktasının sınıfını seçme kararımızı destekleyecektir. Eğitim aşamamızın mümkün olduğunca az destek vektörü bulmasını ve böylece daha az sayıda benzerliği hesaplamamız gerektiğini umuyoruz. Şimdi, destek vektörlerini seçtikten sonra, her destek vektörüne bir ağırlık atarız; bu, kararımızı verirken o destek vektörüne ne kadar önem vermek istediğimizi söyler.

- Adım 1: Her eğitim noktası için ağırlık bulun ve ağırlığı sıfır olan noktalar destek vektörleri değildir, yani test süresi boyunca önemleri sıfırdır ve geri kalanlar destek vektörleridir.
- Adım 2: Vektörlerin her birini başka bir daha yüksek boyutlu vektöre dönüştürün ve ardından bu iki karmaşık daha yüksek boyutlu destek vektörü arasındaki nokta çarpımını alın.
- Adım 3: Duygu değeri olarak vektörlerin maksimumunu bulun.

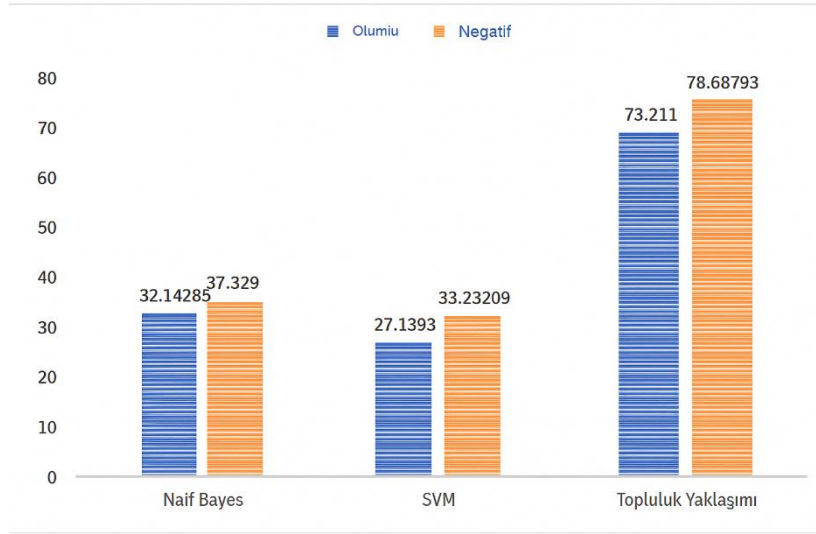
Önerilen metodoloji, sonuçları geliştirmek için Naive Bayes ve Destek Vektör Makinesi sınıflandırma algoritmalarını birleştirmeye odaklanmıştır. Şekil 19, bize veri kümesinden sonuçların elde edilmesine kadar veri işleme ilkesini gösteren bir topluluk algoritması sunulmaktadır. Eğitim ve test verileri, makine öğrenimi algoritmalarının girdisi olarak geçmeden önce ön işleme tabi tutulmuş ve temizlenmiştir. Topluluk yaklaşımı, incelemeyi sınıflandırmak için kullanılır. Topluluk yaklaşımı, Çoğunluk oylaması veya oylama sınıflandırıcısı olarak da bilinir. Topluluk yaklaşımı, algoritmaların birleşimidir. Bir veya daha fazla algoritma sınıflandırılır ve çıktı, tahmincilerin oylamasına (tüm algoritmaların mod değerleri) göre sağlanır. Burada, Topluluk yaklaşımı, Naive Bayes ve SVM birleştirilerek kullanılır. Bu nedenle, önerilen yöntem, önerilenden daha az doğruluk üreten mevcut algoritmalarla karşılaştırıldığında daha yüksek doğruluk üretecektir. Topluluk yöntemi, etkili bir şekilde işlev görmeye yardımcı olan algoritmaları birleştirdiği için genellikle daha iyi doğruluk sağlar. Bu yöntem, tek başına herhangi bir algorithmadan elde edilebilecek olandan daha iyi performans sağlar. Her model her örnekte bir tahmin analizi yapar ve son çıktı tahmini oyların %50'sinden fazlasını alan tahmindir. Tahminlerin hiçbir oyların yarısından fazlasını almazsa, topluluk yöntemi bu örnek için kararlı bir tahmin yapamaz.

- Adım 1: Araştırmacı, Naive Bayes ve SVM gibi mevcut algoritmaları uygulamıştır.
- Adım 2: Araştırmacı, Voting sınıflandırıcısı adlı bir sınıf oluşturmuş ve bu sınıf, Naive Bayes ve SVM gibi kullanılan sınıflandırıcıların listesini devralmıştır.

- Adım 3: Bu sınıflandırıcılar, eğitim kümesinin sınıflandırılmasında kullanılmış ve her algoritma için oylar hesaplanmıştır.
- Adım 4: Oylama sınıflandırıcısı sınıfında, araştırmacı kendi sınıflandırma yöntemini oluşturmuştur. Sınıflandırıcı nesnelerinin listesinde yineleme yapılmış ve her bir sınıflandırıcıdan özelliklere göre sınıflandırma istenmiştir.
- Adım 5: Yineleme tamamlandıktan sonra, araştırmacı en popüler oyu döndüren modu (oyları) döndürmüştür.

Örneğin bir sınıfın frekansının hedef değerlerimizde baskın olduğu sınıf dengesizliği gibi sorunların farkında olmalıyız. Bu gibi durumlarda, her iki sınıfın da aynı oranda olması için muhtemelen test ve eğitim setimiz için sınıflara dayalı katmanlı veri bölme işlemi yapmamız gerekir. Bu nedenle, model bir eğitim veri seti kullanılarak eğitilir ve doğruluğu için test veri seti kullanılarak test edilir.

Veriler, Amazon API'si kullanılarak veya veri kazıma teknikleriyle toplanmış, daha sonra ön işleme adımları (durdurma sözcüklerinin, noktalama işaretlerinin temizlenmesi vb.) uygulanmıştır. Eğitim veri seti Naive Bayes sınıflandırıcısı kullanılarak eğitilecek ve sınıflandırıcının doğruluğu için test verileri kullanılarak daha fazla test edilecektir. Eğitim veri seti, sınıflandırma için en önemli veri setidir. Test verileri, sonucu zaten bilinen (eğitim verilerinin sonucu bile bilinen) ve eğitim verilerine (öğrenmenin ne kadar etkili olduğu) dayanarak makine öğrenimi algoritmasının doğruluğunu belirlemek için kullanılan verilerdir. Burada, test verileri eğitim veri setinden kalan kalan verilerden oluşur. Test verileri 3996 olumlu ve olumsuz metin incelemesinden oluşur. Bu test verileri, belirli bir ürünü önermek için girdinin sonucunu olumlu veya olumsuz olarak kontrol etmek için girdi olarak kullanılacaktır. Test verileri, algoritmanın verileri ne kadar iyi sınıflandırdığı ve çıktıyı ne kadar iyi işlediği konusunda algoritmanın etkinliğini ve doğruluğunu kontrol eden verilerdir.



Şekil 15. Veri Sınıflandırmasının Analizi

Kaynak: Integrated workflow of machine learning and deep learning methods. Patil, D., Rane, N. L., Desai, P., & Rane, J. (2024). Trustworthy Artificial Intelligence in Industry and Society, 28–81.

Çubuk grafik, üç sınıflandırma modelinin (Naive Bayes, SVM ve Topluluk Yaklaşımı) olumlu ve olumsuz sonuçları tahmin etmedeki performansını karşılaştırır. Sonuçlar, Topluluk Yaklaşımının her iki sınıf için de diğer modellerden önemli ölçüde daha iyi performans gösterdiğini, olumlu için yaklaşık %73,21 ve olumsuz için %78,69 elde ettiğini göstermektedir. Buna karşılık, Naive Bayes yaklaşık %32,14 (pozitif) ve %37,33 (negatif) puan alırken, SVM %27,14 (pozitif) ve %33,23 (negatif) ile yakından takip etmiştir. İstatistiksel olarak, Topluluk yöntemi önemli bir iyileştirme göstererek bireysel modellerin doğruluğunu neredeyse iki katına çıkarmıştır. Bu, genel tahmin performansını artırmak ve model önyargısını azaltmak için sınıflandırıcıları birleştirmenin avantajını vurgulamaktadır.

3.2. Makine Öğrenimi ve Derin Öğrenmedeki Güncel Trendler

"Makine öğrenimi (ML) ve derin öğrenme (DL), büyük veri ve artan hesaplama gücü sayesinde hızlı bir gelişim göstermekte, bu alanlarda sürekli yeni teknolojik ilerlemeler ortaya çıkmaktadır."(Goyal vd., 2015). Makine öğreniminde önemli eğilimlerden biri, Edge AI (uç cihazlarda veri işleme) ve Federated Learning (verilerin yerel cihazlarda kalmasını sağlayan dağıtık eğitim) yaklaşımlarının yükselişidir. Edge AI, verileri merkezi bulut sunucularında değil, cihazlarda yerel

olarak işler. Bu trend, gerçek zamanlı işleme talebi, azaltılmış gecikme, gelişmiş gizlilik ve IoT cihazlarının yaygınlaşması tarafından yönlendirilmektedir. Federasyonlu öğrenme, verileri yerel tutarken birden fazla merkezi olmayan cihazda modellerin eğitilmesini sağlayarak bunu tamamlar. Bu yaklaşım, veri gizliliğini ve güvenliğini iyileştirir ve daha sağlam modeller oluşturmak için birden fazla kaynaktan gelen verilerin kullanılmasına olanak tanır.

Açıklanabilir Yapay Zeka (XAI) ML ve DL modelleri daha karmaşık hale geldikçe, şeffaflık ve açıklanabilirliğe olan ihtiyaç artmıştır. Açıklanabilir Yapay Zeka (XAI), karmaşık makine öğrenimi modellerinin kararlarını şeffaflaştırarak kullanıcıların güvenini artırmayı hedefler. Bu eğilim, özellikle sağlık, finans ve otonom sürüş gibi kritik uygulamalarda AI sistemlerine güven kazanmak için çok önemlidir. SHAP (SHapley Eklemeli Açıklamalar) ve LIME (Yerel Yorumlanabilir Modelden Bağımsız Açıklamalar) gibi teknikler popüler hale geliyor ve paydaşların model tahminlerini yorumlamalarını ve altta yatan mekanizmalarını anlamalarını sağlıyor. Transfer öğrenme, daha önce geniş veri kümeleri üzerinde eğitilmiş bir modelin, sınırlı veriye sahip yeni bir görev için hızlıca uyarlanmasını sağlar. Bu, veri eksikliği problemini azaltır ve eğitim süresini kısaltır. Bu teknikler, büyük veri kümelerinde önceden eğitilmiş modellerden yararlanmayı ve bunları belirli görevler için ince ayarlamayı içerir. Bu yaklaşım, kapsamlı hesaplama kaynaklarına ve eğitim süresine olan ihtiyacı önemli ölçüde azaltır. Önemli örnekler arasında, büyük miktarda veri üzerinde önceden eğitilmiş ve doğal dil işleme ile bilgisayarlı görme arasında değişen çeşitli uygulamalar için uyarlanabilen BERT (Çift Yönlü Kodlayıcı Temsilleri), GPT-3 (Üretici Önceden Eğitilmiş Dönüştürücü 3) ve CLIP (Karşılaştırmalı Dil-Görüntü Ön Eğitimi) gibi modeller yer alır.

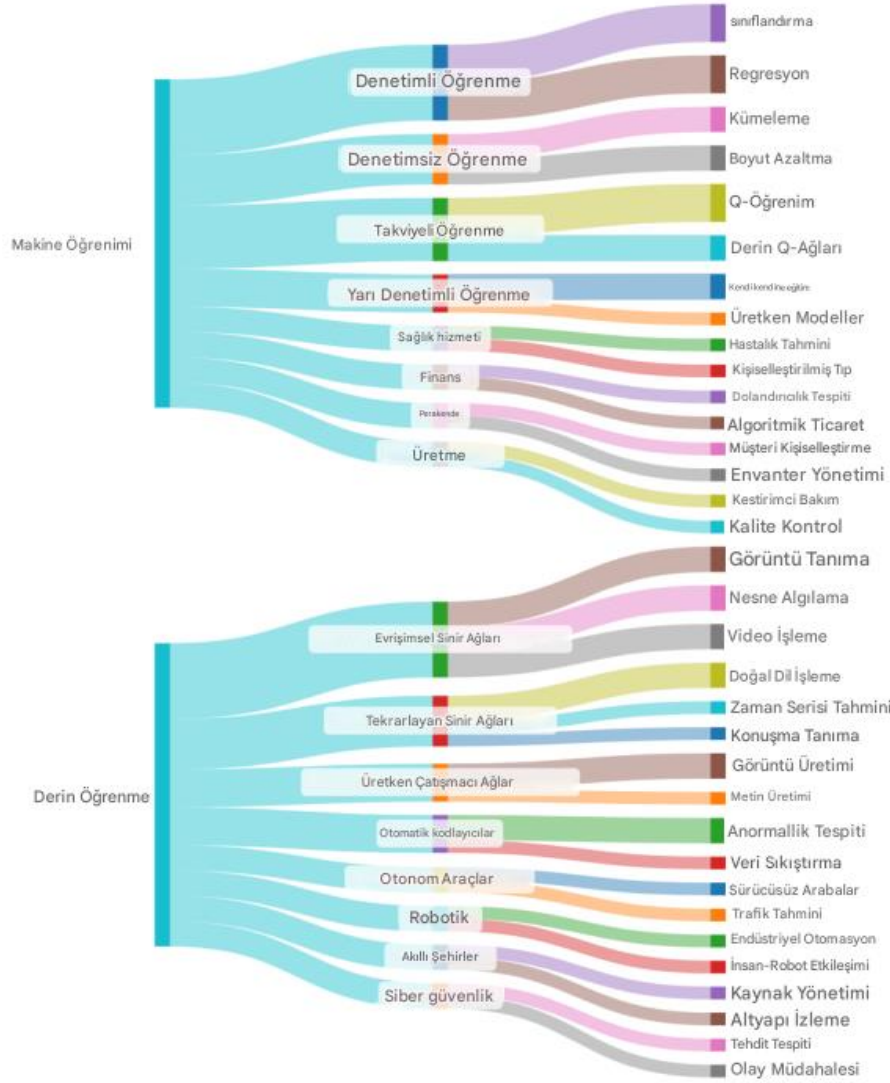
Güçlendirmeli öğrenme, özellikle oyun, robotik ve otonom sistemlerdeki uygulamalarıyla önemli bir odak alanı olmaya devam ediyor. RL'deki son gelişmeler, karmaşık ortamlarda etkileyici yetenekler gösteren Yakınsal Politika Optimizasyonu (PPO) ve Derin Q-Ağları (DQN) gibi geliştirilmiş algoritmalar tarafından desteklenmiştir. Ayrıca, RL'nin gözetimsiz öğrenme ve taklit öğrenme gibi diğer öğrenme paradigmalarıyla entegrasyonu, potansiyel uygulamalarını genişletmektedir (Kontopoulos vd., 2013).

Yapay zeka ve ML'nin etik etkileri artan bir ilgi görmektedir. Önyargı, adalet ve hesap verebilirlik konusundaki endişeler, önyargı tespiti ve azaltma tekniklerine yönelik araştırmaları yönlendirmektedir. Kuruluşlar, modellerinin mevcut önyargıları sürdürmemesini veya şiddetlendirmemesini sağlamak için giderek daha fazla etik AI çerçeveleri benimsiyor. Bu eğilim, önyargılı modellerin önemli toplumsal etkilere sahip olabileceği işe alım, borç verme ve kolluk kuvvetleri gibi alanlarda özellikle önemlidir.

Otomatik Makine Öğrenimi (AutoML), ML modellerinin nasıl geliştirildiği konusunda devrim yaratıyor. AutoML platformları, ML'yi gerçek dünya sorunlarına uygulama sürecinin uçtan uca otomatikleştirilmesini hedefliyor. Buna veri ön işleme, özellik mühendisliği, model seçimi ve hiperparametre ayarı dahildir. AutoML, bu görevleri otomatikleştirerek uzman olmayanların ML modellerini verimli bir şekilde oluşturmasını ve uzmanların model geliştirmenin daha karmaşık yönlerine odaklanmasını sağlar. Google AutoML, H2O.ai ve DataRobot gibi platformlar bu alanda öncülük ediyor.

Nöral Mimari Arama (NAS), sinir ağı mimarilerinin tasarımını otomatikleştirmeye odaklanan yeni bir alandır. Ağ yapılarını elle tasarlamak yerine, NAS algoritmaları belirli görevlere göre uyarlanmış en iyi mimarileri arar. Bu yaklaşım, elle tasarlanmış modellerden daha iyi performans gösteren yeni mimarilerin keşfedilmesine yol açmıştır. EfficientNet ve DARTS (Fark Edilebilir Mimari Arama) gibi teknikler, model performansında ve verimliliğinde önemli iyileştirmeler göstermiştir (Krishna vd., 2013).

Kuantum bilişim, kuantum makine öğreniminin (QML) umut vadeden bir alan olarak ortaya çıkmasıyla ML'de ilerleme kaydediyor. QML, kuantum mekaniğinin prensiplerinden yararlanarak belirli sorunları klasik algoritmalarından daha hızlı çözebilecek algoritmalar geliştiriyor. QML, henüz erken aşamalarında olmasına rağmen optimizasyon, veri sınıflandırması ve üretken modeller konusunda potansiyel gösterdi. Kuantum donanımı gelişmeye devam ettikçe, QML karmaşık ML sorunlarına yaklaşımımızı kökten değiştirebilir.



Şekil 16. Makine Öğrenimi ve Derin Öğrenmedeki Mevcut Eğilimleri Gösteren Sankey Diyagramı

Kaynak: Integrated workflow of machine learning and deep learning methods. Patil, D., Rane, N. L., Desai, P., & Rane, J. (2024). Trustworthy Artificial Intelligence in Industry and Society, 28–81.

Bu diyagram, Makine Öğrenimi ve Derin Öğrenmenin alt alanlarını ve uygulama alanlarını vurgulayarak yapılandırılmış bir genel görünümünü sunar. İstatistiksel olarak, gözetimli öğrenme, yaygın gerçek dünya kullanımını yansıtan sınıflandırma ve regresyona dallanarak baskındır. Gözetimsiz öğrenme, kümeleme ve boyut azaltmayı içerirken, takviyeli öğrenme Q-öğrenme gibi karar alma algoritmalarına odaklanır. Derin Öğrenme, özellikle CNN'ler (Evrişimli Sinir Ağları) ve RNN'ler (Tekrarlayan Sinir Ağları) gibi sinir ağları etrafında merkezlenmiştir ve

görüntü tanıma, NLP (Doğal Dil İşleme) ve video işleme gibi görevlere katkıda bulunur. Sağlık, finans ve üretim gibi pratik sektörler, öngörücü bakım, dolandırıcılık tespiti ve otomasyonu vurgulayarak ML/DL'nin geniş çaplı benimsenmesini sergiler. Bu diyagram, ML/DL alanlarını ayrıntılı olarak açıklayan bir eğitim görselleştirmesinden referans alınmıştır.

Yapay zeka modellerinin, özellikle derin öğrenmenin çevresel etkisi daha belirgin hale geldikçe, sürdürülebilir yapay zekaya giderek daha fazla vurgu yapılmaktadır. Bu eğilim, yapay zeka uygulamalarının karbon ayak izini azaltan enerji açısından verimli algoritmalar ve mimariler geliştirmeyi içermektedir. Yapay zekayı daha sürdürülebilir hale getirmek için model budama, niceleme ve verimli sinir ağları gibi teknikler araştırılmaktadır. Ek olarak, büyük ölçekli modelleri eğitmek için yenilenebilir enerji kaynaklarının kullanılması yönünde bir baskı vardır (Patil vd., 2016).

Birden fazla modaliteden (örneğin, metin, resim, ses) gelen bilgileri entegre etmeyi içeren çok modlu öğrenme ivme kazanmaktadır. Bu yaklaşım, modellerin verilerin daha kapsamlı ve ayrıntılı temsillerini öğrenmesini sağlar. Örneğin, OpenAI'nin CLIP gibi modelleri, resim altyazısı ve görsel soru cevaplama gibi görevleri gerçekleştirmek için görme ve dili birleştirir. Çok modlu öğrenme, içerik oluşturmadan insan-bilgisayar etkileşimine kadar uzanan uygulamalarda yapay zeka sistemlerinin yeteneklerini geliştirmektedir.

Finans sektöründe, ML ve DL, dolandırıcılık tespiti, algoritmik ticaret, risk yönetimi ve müşteri hizmetleri dahil olmak üzere çok çeşitli uygulamalar için kullanılıyor. Büyük miktarda finansal veriyi analiz etme ve kalıpları tespit etme yeteneği, finans kuruluşlarının daha iyi bilgilendirilmiş kararlar almasına ve kişiselleştirilmiş hizmetler sunmasına yardımcı oluyor. Dahası, düzenlemelere uyumda yapay zekanın kullanımı (RegTech), firmaların karmaşık düzenleme manzaralarında daha verimli bir şekilde gezinmesine yardımcı oluyor.

Doğal Dil İşleme (NLP), dönüştürücü tabanlı modellerdeki gelişmelerle yönlendirilen bir yenilik kaynağı olmaya devam ediyor. GPT-4 ve T5 gibi modeller, dil anlama ve oluşturmada yeni ölçütler belirledi. Bu modeller, makine çevirisi, duygu analizi ve konuşma AI gibi görevlere uygulanıyor. NLP'nin diğer AI teknolojileriyle

entegrasyonu, sesle etkinleştirilen asistanlar ve otomatik içerik oluşturma gibi daha karmaşık uygulamaları da mümkün kılıyor.

Siber saldırı tehdidinin artmasıyla birlikte, AI ve ML siber güvenliği geliştirmede önemli bir rol oynuyor. ML algoritmaları anormallikleri tespit etmek, potansiyel tehditleri tahmin etmek ve güvenlik olaylarına gerçek zamanlı olarak yanıt vermek için kullanılıyor. Özellikle derin öğrenme modelleri, siber tehditleri gösteren karmaşık kalıpları belirlemede etkili olduğunu kanıtıyor. Yapay zeka destekli siber güvenlik çözümleri, hassas verileri korumak ve dijital sistemlerin bütünlüğünü sağlamak için vazgeçilmez hale geliyor.

ML ve DL tarafından desteklenen kişiselleştirme ve öneri sistemleri, e-ticaretten yayın hizmetlerine kadar çevrimiçi platformlarda her yerde bulunmaktadır. Bu sistemler, kişiselleştirilmiş içerik ve öneriler sunmak için kullanıcı davranışlarını ve tercihlerini analiz eder. Derin öğrenmedeki gelişmeler, özellikle işbirlikçi filtreleme ve dizi modellemede, önerilerin doğruluğunu ve alakalılığını artırıyor. Kişiselleştirilmiş deneyimler, rekabetçi pazarlardaki işletmeler için önemli bir farklılaştırıcı haline geliyor (Qaisi vd., 2016).

Sankey diyagramı (Şekil 16), Makine Öğrenimi ve Derin Öğrenme ana kategorileriyle başlayarak, belirli öğrenme metodolojilerine ayrılır. Makine Öğrenimi, Gözetimli Öğrenme, Gözetimsiz Öğrenme, Güçlendirmeli Öğrenme ve Yarı Gözetimli Öğrenme olarak ayrılır. Her kategori kendi tekniklerine bağlanır: Gözetimli Öğrenme için Sınıflandırma ve Regresyon, Gözetimsiz Öğrenme için Kümeleme ve Boyut Azaltma, Takviyeli Öğrenme için Q-Öğrenme ve Derin Q-Ağları ve Yarı-gözetimli Öğrenme için Öz Eğitim ve Üretken Modeller. Derin Öğrenme, Evrişimli Sinir Ağları (CNN'ler), Tekrarlayan Sinir Ağları (RNN'ler), Üretken Çelişkili Ağlar (GAN'lar) ve Otokodlayıcıları içerir ve her biri CNN'ler için Görüntü Tanıma, Nesne Algılama ve Video İşleme; RNN'ler için Doğal Dil İşleme, Zaman Serisi Tahmini ve Konuşma Tanıma; GAN'lar için Görüntü Oluşturma ve Metin Oluşturma; ve Otokodlayıcılar için Anomali Algılama ve Veri Sıkıştırma gibi belirli uygulamalara bağlıdır. Ayrıca bu teknolojilerin çeşitli endüstrilerdeki gerçek dünya uygulamalarını da vurgular. Makine Öğrenmesinin Sağlık, Finans, Perakende ve Üretim üzerindeki etkisini gösterir ve Sağlıkta Hastalık Tahmini ve Kişiselleştirilmiş Tıp, Finansta Dolandırıcılık Tespiti ve

Algoritmik Ticaret, Perakendede Müşteri Kişiselleştirme ve Envanter Yönetimi ve Üretimde Tahmini Bakım ve Kalite Kontrolü gibi uygulamaları ayrıntılı olarak açıklar. Benzer şekilde, Derin Öğrenme, Otonom Araçlar, Robotik, Akıllı Şehirler ve Siber Güvenlik gibi alanları etkiler ve Otonom Arabalarda, Trafik Tahmininde, Endüstriyel Otomasyonda, İnsan-Robot Etkileşiminde, Kaynak Yönetiminde, Altyapı İzlemede, Tehdit Tespiti ve Olay Tepkisinde yenilikleri yönlendirir.

3.3. Makine Öğrenimindeki Güncel Sorunlar

Modern bilgi teknolojilerinin temel bileşenlerinden biri olan makine öğrenmesi, yaygın olarak uygulanan bir alandır. Ancak makine öğrenmesinin gelişimi, veri kalitesi, etik problemler ve yönetim gibi çeşitli zorluklarla sınırlanmaktadır. Veri analizi alanındaki gelişmeler ve yeni algoritmaların geliştirilmesi, makine öğreniminin yalnızca akademik araştırmalarda değil, aynı zamanda iş ve günlük yaşamda da uygulama alanını genişletiyor. Örneğin finans, sağlık, otomasyon ve diğer alanlarda kullanılan makine öğrenmesi algoritmaları karar alma süreçlerini daha hızlı ve daha doğru hale getiriyor. Bu durumlarda ortaya çıkan veri kalitesi, etik sorunlar, yönetim ve esneklik gibi zorluklar hayati önem taşımaktadır.

Modern çağda makine öğrenmesinin daha geniş bir şekilde uygulanmasının önünde bir takım engeller bulunmaktadır. Bu engellerin kaldırılması, alanın ilerlemesi için olmazsa olmazdır. Bir modeli eğitmek için gereken verinin kalitesi ve miktarı, tahminlerin doğruluğunu önemli ölçüde etkiler. Sadece kaliteli verilerle çalışmak değil, aynı zamanda bu verileri doğru seçip işlemek de önemlidir. Aynı zamanda model öğrenme sürecinde karşılaşılan aşırı uyum sorunu, veri kümesinin bileşimi ve modelin yapısı konusunda farklı yaklaşımları gerektirmektedir. Bu sorunların üstesinden gelmek için yeni geliştirme yöntem ve metodolojileri araştırılıyor.

Modern makine öğrenmesi yöntemlerinin uygulanması hem olumlu sonuçları hem de kritik zorlukları beraberinde getiriyor. Bu zorlukların başında verinin kalitesi geliyor. Doğru ve güvenilir verilerin eksikliği modellerin sonuçlarını çarpıtabilir. Veri temizleme, gürültü giderme ve boşluk doldurma gibi adımlar makine öğrenmesi sürecinin temel parçalarıdır. Veri kalitesi iyileştirilmezse, modellerin sonuçları yanlış

ve yanıltıcı olabilir; Bu durum gerçek dünya uygulamalarında ciddi sorunlara yol açabilir (Rodak vd., 2014).

Model seçimi de önemli bir konudur. Her makine öğrenimi görevi, belirli analiz gereksinimlerine ve verilerin doğasına uygun bir modelin seçilmesine dayanmalıdır. Farklı modeller farklı yaklaşımlar ve algoritmalar gerektirir ve bazen bir modelin diğerine üstün gelmesi beklenmedik bir durum olabilir. Burada, aşırı uyum (küçük eğitim verilerinde iyi performans gösterme) ve yetersiz uyum (tüm verileri doğru bir şekilde öğrenememe) olguları, tutarlı ve optimum bir model seçmede önemli zorluklar yaratmaktadır. Bu gibi durumlar, eğitim ve test aşamalarında önemli tutarsızlıklara neden olarak makine öğrenmesi sistemlerinin etkinliğini azaltmaktadır.

Ayrıca model açıklanabilirliği ve yorumlanabilirliği, sonuçların kabulü ve uygulanabilirliği açısından kritik öneme sahiptir. Kullanıcılar için net ve anlaşılır bir modelin olması sistem kabulünü artırır. Ancak etik konular (veri toplama, gizlilik ve model işletme prensipleri hakkındaki endişeler) de şarkı öğreniminin geliştirilmesinde en fazla dikkat gerektiren alanlar arasındadır. Her yeni inovasyon, kuruluşların bu etik engellerle mücadele etmek için uygun adımlar atmasını gerektirir. Sonuç olarak, bu konular yalnızca makine öğrenmesinin pratik uygulaması üzerinde değil, aynı zamanda geleceğinin şekillendirilmesi üzerinde de olumlu bir etkiye sahiptir.

Aşırı uyum ve yetersiz uyum, bir modelin performansını ve eğitim verilerinden görülmemiş verilere genelleme yapma yeteneğini önemli ölçüde etkileyen makine öğrenimindeki kritik kavramlardır. Aşırı uyum, bir model aşırı karmaşık hale geldiğinde, altta yatan örüntüler yerine eğitim verilerindeki gürültüyü ve dalgalanmaları yakaladığında meydana gelir (Soni vd., 2014). Sonuç olarak, model eğitim veri setinde olağanüstü doğruluk gösterirken, doğrulama veya test veri setlerinde performansı önemli ölçüde bozulur. Bu senaryo genellikle bir modelin mevcut veri miktarına göre çok fazla parametresi olduğunda veya eğitim optimum öğrenme noktasının ötesinde devam ettiğinde ortaya çıkar. Öğrenme eğrileri gibi görselleştirmeler genellikle bu olguyu ortaya çıkararak eğitim ve doğrulama doğruluğu arasında önemli bir boşluk gösterir ve bir modelin eğitim setinin özelliklerinden çok fazla şey öğrendiğini işaret eder.

Tersine, yetersiz uyum, modelin verilerdeki ilgili eğilimleri ve yapıları yakalayamamasını temsil eder ve genellikle aşırı basitleştirilmiş bir modelden veya yetersiz eğitimden kaynaklanır. Bu gibi durumlarda, model hem eğitim hem de doğrulama veri kümelerinde zayıf performans gösterir, çünkü eğitildiği verilerdeki temel ilişkileri temsil edemez. Yetersiz uyum sağlayan bir model, eğitim örneklerinden öğrenmek için gereken karmaşıklığa sahip değildir ve bu da genellikle yüksek önyargıya ve düşük varyansa yol açar. Bu durum, yetersiz özellikler kullanmaktan, çok basit algoritmalar uygulamaktan veya yetersiz sayıda eğitim yinelemesi ayarlamaktan kaynaklanabilir.

Bu sorunların ele alınması dikkatli bir denge gerektirir, çünkü hem aşırı uyum hem de yetersiz uyum, modelin etkinliğini engeller. Çapraz doğrulama, düzenleme teknikleri ve uygun model karmaşıklığının seçimi gibi stratejiler, aşırı karmaşık modelleri cezalandırarak aşırı uyumu azaltmaya yardımcı olabilir. Öte yandan, yetersiz uyum genellikle model karmaşıklığını artırarak, özellik seçimini iyileştirerek ve veri ilişkilerini daha iyi yakalamak için model parametrelerini ayarlayarak giderilebilir. Makine öğrenimindeki nihai amaç, hem aşırı uyum hem de yetersiz uyumdan kaynaklanan hataları en aza indirirken, görülmemiş verilerdeki tahmin doğruluğunu en üst düzeye çıkaran bir denge kurmaktır; böylece modellerin yalnızca eğitim setinde değil, verilerin önemli ölçüde değişebileceği gerçek dünya uygulamalarında da iyi performans gösterebilmesini sağlar.

Makine öğrenimi, genişleyen kabiliyetlerinin yanı sıra etik kaygıları da beraberinde getiriyor. Uygun etik ilkelerin yokluğunda bu alanın gelişmesi artan risklere yol açabilir. Sistemin analizi, oluşturulması ve uygulanmasının her aşamasında etik konuların tartışılması gerektiği, bunun aynı zamanda insan hakları, bireylerin korunması ve adil iş süreçlerinin sağlanması ile de ilişkili olduğu vurgulanmaktadır. Örneğin, algoritmaların ayrımcılığa yol açabileceği yönündeki endişeler yaygındır. Makine öğrenmesi modellerinin özellikle tarihsel verileri yansıtma konusunda kaygı duyması, toplum yapısındaki adaletsizlikleri daha da derinleştirebilir. Bu gibi durumlarda cinsiyet, ırk ve diğer sosyal parametrelere dayalı nesnel karar alma potansiyeli risk altındadır.

Ayrıca, veri toplama ve işleme sürecinde ortaya çıkan etik sorunlar, kişisel bilgilerin mahremiyeti konusunda endişeleri gündeme getirmektedir. Kullanıcıların izni olmaksızın izinsiz veri toplanması, kamu güvenini zedeleyen ciddi bir sorundur. Büyük veriyi yönetme süreci, verilerin nasıl toplandığını, depolandığını ve kullanıldığını açıklayan açık ve şeffaf bir yaklaşım gerektirir. Ayrıca kimlik tanımada veri güvenliğinin sağlanması ve olası suistimallerin önlenmesi için uygun mekanizmaların hayata geçirilmesi gerekmektedir. Sonuç olarak etik sorumluluk, araştırmacıların ve programcıların yalnızca teknik yönle odaklanmak yerine, toplum üzerindeki etkiyi de göz önünde bulundurarak sorumlu bir yaklaşım benimsemelerini gerektirir.

Bu bağlamda, makine öğrenmesi sistemlerinin etik sorunlarının yalnızca yasal çerçevelerle çözülmesi mümkün değildir; Ayrıca ahlak, insan hakları ve adalet ilkelerinin genişlettiği alanlarda açık diyalog ve işbirliğini de gerektirir. Araştırmacılar, şirketler ve politikacıların birlikte çalışması, bu sorunları ele almak için etik standartların oluşturulmasına ve uygulanmasına yardımcı olabilir. Dolayısıyla makine öğrenmesinde etik konuların tartışılması, bu teknolojilerin topluma fayda sağlamasını ve insan refahını iyileştirme amacına hizmet etmesini sağlayacaktır.

3.4. Güncel Sorunların Çözümü

Makine öğrenmesi, analiz ve uygulanan yöntemlerin etkinliğini artırmayı hedefleyen geliştirme yöntemlerinin kullanıldığı modern veri biliminin dinamik alanlarından biridir. Bu amaçla çeşitli stratejik yaklaşımlar mevcuttur; bunların arasında veri artırma, model optimizasyonu, transfer öğrenmesi ve hibrit modeller özellikle önem taşımaktadır. Her yöntem, algoritmaların performansını iyileştirmek ve elde edilen sonuçların doğruluğunu artırmak için farklı yaklaşımlar sunmaktadır.

Veri artırma süreci, mevcut bir veri setini genişletmek ve zenginleştirmek için uygulanır. Bu yöntem, yeni örnekler yaratarak modelin aşırı uyum sağlama riskini azaltmaya yardımcı olur ve buna bağlı olarak modelin daha geniş ve farklı veri türlerine dayalı genelleme yeteneğini artırır (Sun vd., 2012). Örneğin, görüntü tanıma alanında döndürme, ölçekleme, kırpma gibi dönüşümlerle yeni görüntülerin oluşturulması daha geniş bir veri yelpazesi sağlar. Aşağıdaki model optimizasyonu,

hiperparametrelerin ve performans ölçümlerinin uygun şekilde ayarlanması etrafında dönen metodolojiler aracılığıyla problem mükemmelliğine odaklanır. Bunlara, modelin doğruluğunu artırmaya olanak tanıyan ızgara araması ve rastgele arama gibi yaklaşımlar da dahildir.

Transfer öğrenmesi, daha önce öğrenilen model bilgisinin yeni ancak ilgili görevlere uygulanmasını sağlar. Bu yöntem, eski modelin soyut kavramlarının yeni verilere uyarlanabilmesi nedeniyle sınırlı veri kümeleriyle daha etkili öğrenmeye yardımcı olur. Hibrit modeller ise farklı yöntemlerin bir araya getirilmesiyle daha karmaşık ve uygun maliyetli yaklaşımlar sunuyor. Örneğin, büyük veri analizi alanında daha doğru sonuçlara ulaşılmasına olanak sağlaması nedeniyle, derin öğrenme ile geleneksel istatistiksel yöntemlerin birlikte kullanımı günümüzde yaygınlaşmıştır. Dolayısıyla makine öğrenimine yönelik geliştirme yöntemleri yalnızca performansı iyileştirmekle kalmıyor, aynı zamanda bilimsel araştırma ve inovasyonun ilerlemesi için gerekli bir temel oluşturuyor.

Veri artırma veya daha geniş anlamda veri artırma, bir modelin öğrenme yeteneğini artırmak için makine öğrenmesinde kullanılan temel bir yaklaşımdır. Bu yöntem, mevcut verilerden yola çıkılarak modelin farklı senaryolarda daha güvenilir sonuçlar üretmesini sağlamayı amaçlamaktadır. Veri kümelerindeki kısıtlamalardan kaynaklanan öğrenme hatalarının üstesinden gelmek için görüntü, metin, ses ve diğer veri türleri üzerinde çeşitli şekillerde veri artırma yapılabilir. Örneğin, döndürme, ölçekleme, parlaklık değişiklikleri gibi dönüşümler yapılarak görüntü verilerinde yeni desenler oluşturmak mümkündür; Bu, modele daha fazla değişken kazandırarak genelleme yeteneğini artırır.

Veri artırma aynı zamanda sisteme aşırı uyum sağlamayan daha sağlam bir model elde etmeye de yardımcı olur. Aşırı uyum, modelin eğitim setinin belirli özelliklerine aşırı uyum sağlaması ve yeni verilerde düşük performans göstermesi ile karakterize edilir. Bunu azaltmak için, veri artırma yöntemleriyle yeni örnekler sunulduğunda, modelin daha çeşitli ve geniş bir ortamda öğrenme olanağı vardır. Üstelik bu yaklaşım, az miktarda veriyle önemli sonuçlar elde edilmesine olanak tanıdığı için kaynakların sınırlı olduğu durumlarda özellikle yararlıdır. Yapılan

alıřmalar, veri artırma stratejilerinin model sonularını nemli lde iyileřtirdiėini ve ėrenme sresini optimize ettiėini gstermiřtir.

Derin ėrenme uygulamalarında veri artırmanın kullanımı son zamanlarda yaygınlařmıřtır, nk derin modeller genellikle byk veri kmelerine ihtiya duymaktadır. Veri artırma yntemleri kolaylık saėlar, rneėin retken atıřmalı Aėlar (GAN'lar) gibi modern yaklařımlar yeni veriler reterek daha zengin bir eėitim ortamına olanak tanır. Bu modern teknolojiler, makine ėrenimi sistemlerinin daha da iyileřtirilmesinin, model performansının artırılmasının ve daha gvenilir arařtırma sonularının nn aıyor. Sonu olarak, veri artırma yalnızca veri miktarını artırmakla kalmıyor, aynı zamanda otomasyonun modern dnyasında nemli bir rol oynayan modelin genel davranıřını da iyileřtiriyor.

Model optimizasyonu, makine ėrenmesi sistemlerinin etkinliėini artırma amacına hizmet eden bir dizi yntem ve stratejinin birleřimidir. Bu sre esneklik ve verimlilik saėlarken, aynı zamanda modelin standart performansını da arttırıyor. ncelikle farklı optimizasyon yntemleri yklenen verinin spesifik analizine ve ėrenme parametrelerinin ayarlanmasına odaklanır.

Derin ėrenme baėlamında, gradyan iniři ve Yapay Sinir Aėlarında kullanılan Stokastik Gradyan İniři (SGD), Adam ve RMSprop gibi eřitli formları, modelin performansını nemli lde iyileřtirmeye yardımcı olur.

Model optimizasyonu aynı zamanda hiperparametre ayarlama srecini de ierir. Hiperparametreler, model ėrenme sreci sırasında ayarlanmayan, ancak modelin kalitesi zerinde nemli bir etkiye sahip olan parametrelerdir. ėrenme oranının, toplu byklėn ve dnem sayısının optimize edilmesi, modelin ařır uyum ve yetersiz uyum gibi durumlardan kaınmasını saėlar. Bu amala Grid Search (Izgara Araması) ve Random Search (Rastgele Arama) gibi yntemler optimizasyon kapsamını geniřletirken, Cross-Validation (apraz Doėrulama) tekniklerinin uygulanmasıyla modelin nihai performansı daha doėru bir řekilde deėerlendirilmektedir.

Makine ėrenimi modellerinin geliřtirilmesi srecinde, optimizasyon yalnızca parametrelerin ayarlanması ile sınırlı kalmayıp, aynı zamanda modelin genellenebilirliėini artırmak amacıyla dzenleme (reglerizasyon) tekniklerinin

uygulanmasını da kapsamaktadır. Bu bağlamda, L1 (Lasso) ve L2 (Ridge) düzenleme yöntemleri, modelin karmaşıklığını kontrol altına alarak aşırı uyum (overfitting) problemini önlemede önemli bir rol oynamaktadır. L1 regülerizasyonu, modelde yer alan bazı özelliklerin katsayılarını sıfıra indirerek otomatik özellik seçimi (feature selection) gerçekleştirir. Bu, modelin yalnızca en anlamlı değişkenlerle eğitilmesini sağlar ve böylece hem açıklayıcılığı artırır hem de daha sade modellerin oluşturulmasına katkıda bulunur. L2 regülerizasyonu ise tüm katsayıların büyüklüğünü azaltarak modelin genelleme kapasitesini yükseltir; ancak, L1'den farklı olarak katsayıları sıfıra indirmez. Bu yöntem, tüm özellikleri modelde tutarak daha dengeli ve istikrarlı sonuçlar elde edilmesini hedefler. Her iki düzenleme tekniği de ceza terimi olarak bilinen ek bir bileşen üzerinden çalışır. Bu terim, modelin kayıp fonksiyonuna (loss function) dahil edilerek büyük ağırlık değerlerinin cezalandırılmasını sağlar. L1 düzenlemesi ağırlıkların mutlak değerlerini, L2 düzenlemesi ise karelerini kullanarak ceza uygular. Böylece model, daha düşük varyansla ve daha yüksek genellenebilirlik düzeyiyle öğrenim sağlar. Veri çeşitliliği ve problem yapısına bağlı olarak, uygun algoritmalarla entegre edilmiş regülerizasyon stratejileri, makine öğrenimi uygulamalarında hem verimliliği artırmakta hem de uygulama alanlarının genişlemesini mümkün kılmaktadır. Bu yönüyle model optimizasyonu, yalnızca modelin doğruluğunu değil, aynı zamanda güvenilirliğini de artırarak makine öğrenimi projelerinin başarıyla tamamlanmasına katkı sunmaktadır. Ayrıca, optimizasyon sürecinin ayrılmaz bir parçası da düzenleme tekniklerinin uygulanmasıdır. L1 ve L2 düzenlemesi, modelin karmaşıklığını kontrol etmeye yardımcı olarak yukarıda belirtilen aşırı uyum sorununu önlemeyi amaçlamaktadır. Bu, yalnızca daha basit ama etkili modellerin oluşturulmasını sağlamakla kalmıyor, aynı zamanda makine öğrenimi sistemlerinin genel rehberliğini de artırıyor. Analiz edilen verilerin çeşitliliğine göre uyarlanan ve uygun algoritmalarla uyumlu entegre optimizasyon stratejileri, makine öğrenmesinin verimliliği ile uygulama alanlarının genişlemesi arasındaki ilişkiyi güçlendirmektedir. Bu şekilde model optimizasyonu, herhangi bir makine öğrenimi projesinin başarılı bir şekilde tamamlanmasını destekleyerek, sonuç olarak daha doğru ve güvenilir veri analizi sağlar.

Transfer öğrenmesi, makine öğrenmesinde temel bir kavramdır ve bir görevden edinilen bilgiyi, ilgili ancak farklı bir görevdeki performansı iyileştirmek için

kullanmayı amaçlar. Bu teknik, her benzersiz görev için büyük miktarda etiketli veri gerektiren geleneksel makine öğrenimi yaklaşımlarıyla büyük bir tezat oluşturuyor. Bu çerçevede, transfer öğrenmesi, büyük veri kümeleri üzerinde eğitilmiş olan önceden eğitilmiş modellerin uyarlanmasını kolaylaştırır ve bunların nispeten sınırlı veriye sahip uzmanlaşmış uygulamalar için ince ayar yapılmasına olanak tanır. Bu, yalnızca eğitim sürecini hızlandırmakla kalmaz, aynı zamanda modelin gerçek dünya uygulamalarındaki verimliliğini ve etkinliğini de artırır.

Transfer öğrenmesinin temel bir yönü, derin sinir ağları tarafından öğrenilen özelliklerin hiyerarşik yapısından yararlanma yeteneğidir. Başlangıç katmanlarında modeller, çeşitli alanlarda geçerli olan kenarlar ve dokular gibi genel özellikleri yakalama eğilimindedir. Ağda daha derinlere doğru ilerledikçe, özellikler giderek daha fazla görev odaklı hale gelir. Önceden eğitilmiş bir modelin erken katmanlarını dondurarak ve belirli bir yeni görevi ele alırken yalnızca sonraki katmanları ayarlayarak, uygulayıcılar transfer öğrenmesinin faydalarını en üst düzeye çıkarabilirler. Bu, tıbbi görüntüleme veya doğal dil işleme gibi veri toplamanın pahalı veya zaman alıcı olduğu alanlarda özellikle avantajlıdır.

Yapılan araştırmalar, transfer öğrenmesinin, mevcut eğitim verilerinin seyrek olduğu senaryolarda kayda değer performans iyileştirmelerine yol açabileceğini göstermiştir. Kaynak ve hedef alanlar arasındaki dağılım kaymasını en aza indiren alan uyarlaması ve birden fazla ilgili görevde modelleri aynı anda eğiten çoklu görev öğrenmesi gibi teknikler, transfer öğrenmesinin çok yönlülüğüne daha fazla örnek teşkil eder. Bu stratejiler yalnızca paylaşılan bilgiyi güçlendirmekle kalmıyor, aynı zamanda çeşitli bağlamlarda genelleştirilebilen daha dayanıklı bir modelin oluşturulmasını da kolaylaştırıyor (Wadbude vd., 2016). Makine öğrenimi gelişmeye devam ederken, transfer öğrenmesi, çeşitli uygulamalara uyarlanabilen sağlam yapay zeka sistemlerinin hızla geliştirilmesini sağlayan, böylece önemli zorlukların ele alınmasını ve alandaki yeni olanakların kilidinin açılmasını sağlayan önemli bir metodoloji olarak ortaya çıkmaktadır.

Hibrit modeller, makine öğrenmesinin mekansal çok boyutluluğuna yeni yaklaşımlar getiren yenilikçi sistemlerdir. Bu modeller, istatistiksel modeller ile makine öğrenme algoritmalarının etkileşimi olmak üzere daha önceki geleneksel

yöntemlerin bir araya getirilmesiyle ortaya çıkmıştır. Hibrit sistemlerin etkisi, hem basit istatistiksel modellerin geleneksel sınırlarının aşılmasında hem de daha karmaşık algoritmaların optimizasyonundan kaynaklanan sorunların aşılmasında önemli rol oynamaktadır. Örneğin, hibrit modeller bir yandan istatistiksel öğrenme yöntemlerinin basitliği ve açıklanabilirliği ile ilerlerken, diğer yandan derin öğrenme yöntemlerinin karmaşıklığından yararlanarak daha doğru sonuçlara ulaşılmasını sağlar.

Bu sistemler aynı zamanda çeşitli veri kaynaklarının daha geniş ve güvenilir bir şekilde bütünleştirilmesi ve daha enerjik ve verimli sonuçlar elde edilmesi sürecini de kolaylaştırmaktadır. Hibrit modellerin birçok uygulama alanında avantajı, özel gereksinimlere uyum sağlayabilmeleri ve farklı modellerin yüksek düzeyde entegre edilebilmeleridir. Örneğin, farklı veri biçimlerinin (örneğin yapılandırılmış ve yapılandırılmamış veriler) kullanımının artmasıyla karşı karşıya kaldığımız zorluklar, hibrit yaklaşımlarla daha etkili bir şekilde ele alınabilir. Burada temel odak noktası klasik modellerin yeni teknolojilerle birlikte kullanılması ve böylece tutarlı analiz ve sonuçların doğruluğunun artırılmasıdır.

Hibrit modellerin bir diğer önemli yönü de uyarlanabilir ve ölçeklenebilir olmalarıdır. Örneğin tıbbi tanı sistemlerinde, modern tıbbi veri ve görüntü analizi algoritmaları, kanser gibi karmaşık hastalıkların tanınmasında istatistiksel ve derin öğrenme modellerinin yararlı bir kombinasyonunu uygulayarak hibrit modeller aracılığıyla daha verimli sonuçlara ulaşılmasına olanak sağlamaktadır. Sonuç olarak, bu tür sistemlerin geliştirilmesi, makine öğrenmesine yönelik durum-özel hibrit yaklaşımların geliştirilmesine odaklanan merkezi bir araştırma alanı haline gelmiştir.

SONUÇ

Makine öğrenmesi son yıllarda hızla gelişen bir alandır, ancak aynı zamanda birçok zorluğu ve sorunu da beraberinde getirir. Bu çalışmada, veri kalitesi, model seçimi, öğrenme algoritmalarının karmaşıklığı, knime analiz ve yorumlama zorluğu gibi makine öğrenmesinin önemli yönleri etrafında dönen tartışmalar hem sizlerin hem de akademik camianın odağındadır. Geliştirilen modellerin doğru ve güvenilir sonuçlar üretmesini sağlamada, verilerin doğru bir şekilde toplanması, temizlenmesi ve yapılandırılması kritik bir rol oynar. Ancak etik konular, önyargıların ortadan kaldırılması, kullanıcıların kişisel verilerinin korunması gibi hususlar güvenli, güvenilir ve adil sistemler oluşturmak için büyük önem taşımaktadır. Geliştirme yöntemleri açısından hibrit yöntemlerin uygulanması, kişiselleştirilmiş öğrenme stratejileri ve alternatif değerlendirme yöntemlerinin kullanımı makine öğrenmesinin daha geniş ve derin alanlara yayılmasına olanak sağlamaktadır.

Derin öğrenme ve pekiştirmeli öğrenme alanındaki gelişmeler, karmaşık problemlerin çözümünün hızlanmasının yanı sıra sistemlerin kendini geliştirme yeteneğini artırarak daha çevik ve etkili modellerin oluşturulmasının önünü açıyor. Bu gelişmeler sonucunda makine öğrenmesi, yalnızca akademisyenler arasında değil, aynı zamanda endüstriler arasında da bilgi alışverişinin hızını, doğruluğunu ve kapsamını genişletmede vazgeçilmez bir rol oynamaktadır. Gelecek perspektifleri, makine öğreniminin geliştirilmesinde temel yaklaşımların ve metodolojilerin daha da ilerletilmesi için mükemmel bir platform sağlar. Daha önce yapay zekanın geri dönüşüne ilişkin teoriler geliştirilmiş olsa da son zamanlarda bu yaklaşımların uygulama alanı daha da genişliyor. Örneğin, evrimsel ve uyarlanabilir algoritmalar, daha etkili öğrenme süreçleri için ek kriterler sağlayarak sistemlerin gelişmesine yardımcı olabilir. Ayrıca bu metodolojilerin uygulanmasıyla farklı modellerin karşılaştırılması ve analiz edilmesi, çok yönlü çözümlere yönelik araştırmaların daha da ilerlemesini sağlar. Gelecekte veri toplama ve işleme alanında yeni yaklaşımlar, makine öğrenmesi sistemlerinin etkisini artıracaktır. Makine öğrenimi, veri çağında hızla gelişen ve yaygın olarak benimsenen bir teknoloji olarak ortaya çıkmıştır. Bu çalışma, makine öğrenimi algoritmalarının mevcut rolünü, teorik temellerini, uygulama alanlarını ve gelecekteki trendleri kapsamlı bir şekilde incelemiştir. Veri ve donanım bağımlılıkları, özellik işleme teknikleri, yorumlanabilirlik, yürütme süresi ve

model seçimi gibi temel hususlar, algoritma performansı üzerindeki etkilerini anlamak için analiz edilmiştir.

Makine öğreniminin siber güvenlikteki rolü de, özellikle saldırı tespit sistemlerinde ve veri analizinde uygulanması olmak üzere merkezi bir odak noktası olmuştur. NSL-KDD veri kümesi ve KNIME aracının kullanımıyla, çeşitli makine öğrenimi algoritmalarının deneysel bir değerlendirmesi yapılmıştır. Sonuçlar, farklı modeller arasında değişen verimlilik, doğruluk ve yorumlanabilirlik düzeylerini vurgulamıştır. Konuyla ilgili birkaç öneri paylaşılabilir:

- Veri kalitesi ve ön işleme, model doğruluğunu ve sağlamlığını artırmak için önceliklendirilmelidir. Etkili özellik seçimi ve normalleştirme, performansı önemli ölçüde iyileştirebilir.
- Algoritma seçimi, hesaplama maliyeti, ölçeklenebilirlik ve yorumlanabilirlik gibi faktörler göz önünde bulundurularak sorunun özel gereksinimleriyle uyumlu olmalıdır.
- Açıklanabilir AI (XAI) yaklaşımları, özellikle siber güvenlik gibi kritik alanlarda, karar alma süreçlerinde şeffaflık ve güveni sağlamak için daha geniş bir şekilde veri büyüklüğüne entegre edilmeli.
- Derin öğrenme ve hibrit modellerde devam eden araştırma ve geliştirme, aşırı uyum, yüksek eğitim maliyetleri ve veri bağımlılığı gibi mevcut zorlukların üstesinden gelmeye yardımcı olabilir.
- Kurumlar ve kuruluşlar, daha karmaşık ve duyarlı ML sistemlerinin dağıtımını desteklemek için donanım altyapısına ve gerçek zamanlı veri hatlarına yatırım yapmalıdır.

KAYNAKÇA

- Abdullah, M. (2021). The implication of machine learning for financial solvency prediction: An empirical analysis on public listed companies of Bangladesh. *Journal of Asian Business and Economic Studies*, 28(4), 303–320.
- Abualkibash Adams, S., Arel, I., Bach, J., Coop, R., Furlan, R., Goertzel, B., Hall, J. S., Samsonovich, A., Scheutz, M., Schlesinger, M., Shapiro, S. C., & Sowa, J. (2012). Mapping the landscape of human-level artificial general intelligence. *AI Magazine*, 33(1), 25–42.
- Alvarado, N., Adams, S. S., Burbeck, S., & Latta, C. (2006). Beyond the Turing test: Performance metrics for evaluating a computer simulation of the human mind. In *AGIRI 2006 Conference Proceedings*.
- Arafat, M., Jain, A., & Wu, Y. (2018). Analysis of intrusion detection dataset NSL-KDD using KNIME analytics. In *Proceedings of the International Conference on Cyber Warfare and Security*. Academic Conferences International Limited.
- Arun, A., et al. (2023). Deep learning algorithms used in intrusion detection systems: A review. *ar5iv.org*.
- Bancken, W., Alfarone, D., & Davis, J. (2014, August). Automatically detecting and rating product aspects from textual customer reviews. In *Proceedings of the 1st International Workshop on Interactions Between Data Mining and Natural Language Processing at ECML/PKDD* (pp. 1–16).
- Barlow, H. B. (1989). Unsupervised learning. *Neural Computation*, 1(3), 295–311.
- Burkov, A. (2019). *The hundred-page machine learning book* (pp. 3–5). Canada: Andriy Burkov.
- Carbonell, J. G., & Scannell, R. S. (2019). Machine learning: A historical and methodological analysis. *Association for the Advancement of Artificial Intelligence*.

- Caruana, R., & Niculescu-Mizil, A. (2006, June). An empirical comparison of supervised learning algorithms. In *Proceedings of the 23rd International Conference on Machine Learning* (pp. 161–168).
- CBS News. (2018, July 13). *Intel chief Dan Coats says of cyber attacks: "We are at a critical point"*.
- Challita, U., Ryden, H., & Tullberg, H. (2020). When machine learning meets wireless cellular networks: Deployment, challenges, and applications. *IEEE Communications Magazine*, 58(6), 12–18.
- Chella, A., Lebiere, C. L., Noelle, D. C., & Samsonovich, A. V. (2011). On a roadmap to biologically inspired cognitive agents. *Frontiers in Artificial Intelligence and Applications*, 233, 453–460.
- DataCamp. (2023). *Machine learning engineer salaries in 2023*.
- Fang, X., & Zhan, J. (2015). Sentiment analysis using product review data. *Journal of Big Data*, 2(1), 1–14.
- Ghahremani-Nahr, J., Nozari, H., & Bathaee, M. (2021). Robust box approach for blood supply chain network design under uncertainty: Hybrid moth-flame optimization and genetic algorithm. *International Journal of Innovation in Engineering*.
- Goyal, A., & Parulekar, A. (2015). Sentiment analysis for movie reviews.
- Graham-Rowe, D. (2004). Turing test is a total turn-off for robots: The thinking machine's challenge of choice is something a little more down to earth. *New Scientist*, 183(2461), 22.
- Haldar, M., Abdool, M., Ramanathan, P., Xu, T., Yang, S., Duan, H., Zhang, Q., Barrow-Williams, N., Turnbull, B. C., & Collins, B. M. (2019). Applying deep learning to AirBnB search. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*.
- Hansson, K., Yella, S., Dougherty, M., & Fleyeh, H. (2016). Machine learning algorithms in heavy process manufacturing. *American Journal of Intelligent Systems*.

- Hong, S., Yi, T., Yum, J., & Lee, J. H. (2021). Visitor-artwork network analysis using object detection with image-retrieval technique. *Advanced Engineering Informatics*, 48, 101307.
- Jordan, M. I., & Mitchell, T. M. (2015). Machine learning: Trends, perspectives, and prospects. *Science*, 349(6245), 255–260.
- Kehe, W., Tao, D., & Jianping, H. (2011). The research of intrusion detection technology based on heuristic analysis. In *Proceedings of the 2011 6th IEEE Joint International Information Technology and Artificial Intelligence Conference*.
- Keramati, A., Ghaneei, H., & Mirmohammadi, S. M. (2016). Developing a prediction model for customer churn from electronic banking services. *Financial Innovation*, 2(1), 10.
- Kontopoulos, E., Berberidis, C., Dergiades, T., & Bassiliades, N. (2013). Ontology-based sentiment analysis of Twitter posts. *Expert Systems with Applications*, 40(10), 4065–4074.
- Korukonda, A. R. (2003). Taking stock of the Turing test: A review, analysis, and appraisal of issues surrounding thinking machines. *International Journal of Human-Computer Studies*, 58(4), 319–348.
- Krishna, P. V., Misra, S., Joshi, D., & Obaidat, M. S. (2013, May). Learning automata-based sentiment analysis for recommender system on cloud. In *2013 International Conference on Computer, Information and Telecommunication Systems (CITS)* (pp. 1–5). IEEE.
- LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436–444.
- Madaio, M., Chen, S. T., Haimson, O. L., Zhang, W., Cheng, X., Hinds-Aldrich, M., & Chau, D.
- Meena, G., & Choudhary, R. R. (2017). A review paper on IDS classification using KDD 99 and NSL-KDD dataset in WEKA. In *2017 International Conference on Computer, Communications and Electronics (Comptelix)*.

- Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2021). A survey on bias and fairness in machine learning. *ACM Computing Surveys (CSUR)*, 54(6), 1–35.
- Mitchell, T. M. (2020). The discipline of machine learning. *Machine Learning Department Technical Report CMU-ML-06-108*. Carnegie Mellon University.
- Nouri, F., Samadzad, S., & Nahr, J. (2019). Meta-heuristics algorithm for two-machine no-wait.
- Nozari, H., Najafi, E., Fallah, M., & Hosseinzadeh Lotfi, F. (2019). Quantitative analysis of key performance indicators of green supply chain in FMCG industries using non-linear fuzzy method. *Mathematics*, 7(11), 1020.
- Patil, D., Rane, N. L., Desai, P., & Rane, J. (2024). Machine learning and deep learning: Methods, techniques, applications, challenges, and future research opportunities. In *Trustworthy artificial intelligence in industry and society* (pp. 28–81). Deep Science Publishing.
- Patil, H., & Mane, P. M. (2016). Survey on product review sentiment analysis with aspect ranking. *International Journal of Science and Research*, 5(9), 749–752.
- Paulauskas, N., & Auskalnis, J. (2017). Analysis of data pre-processing influence on intrusion detection using NSL-KDD dataset. *2017 Open Conference of Electrical, Electronic and Information Sciences (eStream)*.
- Qaisi, L. M., & Aljarah, I. (2016, July). A Twitter sentiment analysis for cloud providers: A case study of Azure vs. AWS. In *2016 7th International Conference on Computer Science and Information Technology (CSIT)* (pp. 1–6). IEEE.
- Ravipati, M. A. (2019). A survey on different machine learning algorithms and weak classifiers based on KDD and NSL-KDD datasets. *International Journal of Artificial Intelligence and Applications (IJALA)*, 10(3).
- Rodak, J., Xiao, M., & Longoria, S. (2014). Predicting helpfulness ratings of Amazon product reviews. *Unpublished manuscript*.

- Rose, S. (2016). A machine learning framework for plan payment risk adjustment. *Health Services Research*, 51(6), 2358–2374.
- Sammut, C., & Webb, G. I. (Eds.). (2011). *Encyclopedia of machine learning*. Springer.
- Samsonovich, A. V. (2019). The Constructor metacognitive architecture. *AAAI Fall Symposium – Technical Report, FS-09-01*, 124–134.
- Samsonovich, A. V., Goldin, R. F., & Ascoli, G. A. (2010). Toward a semantic general theory of everything. *Complexity*, 15(4), 12–18.
- Scanlon, J. M., Kusano, K. D., Daniel, T., Alderson, C., Ogle, A., & Victor, T. (2021). Waymo simulated driving behavior in reconstructed fatal crashes within an autonomous vehicle operating domain. *Unpublished manuscript*.
- Shah, V. (2021). Machine learning algorithms for cybersecurity: Detecting and preventing threats. *Revista Española de Documentación Científica*, 15(4), 42–66.
- Shaukat, K., Luo, S., Chen, S., & Liu, D. (2020, October). Cyber threat detection using machine learning techniques: A performance evaluation perspective. In *Proceedings of the 2020 International Conference on Cyber Warfare and Security (ICCWS)* (pp. 1–6). IEEE.
- Soni, V., & Patel, M. R. (2014). Unsupervised opinion mining from text reviews using SentiWordNet. *International Journal of Computer Trends and Technology*, 11(5), 234–238.
- Srinivas, B. (2024). Enhancing cybersecurity with machine learning: Algorithms and approaches. *International Journal of Intelligent Systems and Applications in Engineering*, 12(22s), 210–220.
- Stolfo, S. J., Fan, W., Lee, W., Prodromidis, A., & Chan, P. K. (2000). Cost-based modeling for fraud and intrusion detection: Results from the JAM project. In *Proceedings of the DARPA Information Survivability Conference and Exposition (DISCEX'00)*.

- Sun, S., Liu, H., Lin, H., & Abraham, A. (2012, October). Twitter part-of-speech tagging using pre-classification Hidden Markov Model. In *2012 IEEE International Conference on Systems, Man, and Cybernetics (SMC)* (pp. 1118–1123). IEEE.
- Tavallae, M., Bagheri, E., Lu, W., & Ghorbani, A. A. (2009). A detailed analysis of the KDD CUP 99 dataset. In *Proceedings of the 2009 IEEE Symposium on Computational Intelligence for Security and Defense Applications*.
- Thomas, R., & Pavithran, D. (2018). A survey of intrusion detection models based on NSL-KDD dataset. In *Proceedings of the 2018 Fifth HCT Information Technology Trends (ITT)*.
- Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433–460.
- Vegesna, V. V. (2023). Privacy-preserving techniques in AI-powered cybersecurity: Challenges and opportunities. *International Journal of Machine Learning for Sustainable Development*, 5(4), 1–8.
- Wadbude, R., Gupta, V., Mekala, D., Jindal, J., & Karnick, H. (2016). User bias removal in fine-grained sentiment analysis. In *European Chapter of the Association for Computational Linguistics*. arXiv:1612.06821.
- Warzyński, A., & Kolaczek, G. (2018). Intrusion detection systems vulnerability on adversarial examples. In *Proceedings of the 2018 Innovations in Intelligent Systems and Applications (INISTA)*.

İnternet Kaynağı

Hassan. NSL-KDD. 6 Nisan 2025 tarihinde
<https://www.kaggle.com/datasets/hassan06/nslkdd> adresinden edinilmiştir.