

DOKUZ EYLÜL ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ

MAKİNE ÖĞRENMESİ YÖNTEMLERİ İLE
SİBER SALDIRI TESPİTİ

Zehra AKBIYIK

Eylül, 2024

İZMİR

MAKİNE ÖĞRENMESİ YÖNTEMLERİ İLE SİBER SALDIRI TESPİTİ

Dokuz Eylül Üniversitesi Fen Bilimleri Enstitüsü

Yüksek Lisans Tezi

İstatistik Anabilim Dalı, Veri Bilimi Programı

Zehra AKBIYIK

Eylül, 2024

İZMİR

YÜKSEK LİSANS TEZİ SINAV SONUÇ FORMU

ZEHRA AKBIYIK tarafından **DOÇ.DR. SEDAT ÇAPAR** yönetiminde hazırlanan
“MAKİNE ÖĞRENMESİ YÖNTEMLERİ İLE SİBER SALDIRI TESPİTİ”
başlıklı tez tarafımızdan okunmuş, kapsamı ve niteliği açısından bir Yüksek Lisans
tezi olarak kabul edilmiştir.

Doç.Dr.Sedat ÇAPAR

Danışman

Prof.Dr.Abdullah Fırat ÖZDEMİR

Jüri Üyesi

Doç.Dr.Hakan Savaş SAZAK

Jüri Üyesi

Prof. Dr. Okan FISTIKOĞLU

Müdür

Fen Bilimleri Enstitüsü

TEŞEKKÜR

Öncelikle, yüksek lisans çalışmamın geliştirilmesi ve tamamlanmasındaki değerli katkıları için danışman hocam Doç. Dr. Sedat Çapar'a en içten teşekkürlerimi sunarım.

Ayrıca, bu süreçte yanımda olan, desteğini ve sevgisini her daim hissettiren sevgili aileme, özellikle bana güç veren kardeşime teşekkür ederim. Onların varlığının omzuma uzanan bir el oluşu, bu yolculukta hissettiğim en büyük destektir.

Bu vesileyle, akademik hayatımda bana her zaman yol gösterici olan ve sonsuz sabrıyla ilham kaynağım haline gelen Buse Demir'e, iki insanın birbirini anlamasında mükemmel tonu yakaladığım Ayla Altıntaş'a, upuzun ağaçlı yolların eşlikçisi Erdem Engin'e ve bu süreçte beni yüreklendiren, kendime olan inancımı pekiştiren, bu yolculuğu daha güzel hale getiren tüm yol arkadaşlarıma da en içten teşekkürlerimi sunarım.

Zehra AKBIYIK

MAKİNE ÖĞRENMESİ YÖNTEMLERİ İLE SİBER SALDIRI TESPİTİ

ÖZ

Teknolojinin gelişimi ile birlikte siber güvenlik kavramı da hayatımıza girmiştir. Yıllar içinde, mali kayıplar, kişisel veri ihlalleri ve şantaj gibi çeşitli nedenlerle siber saldırılar artmış ve bu saldırılara yönelik çalışmalar hız kazanmıştır. Bu çalışma, siber saldırıların tespitine yönelik ihtiyaçtan doğmuştur.

Çalışmada, siber saldırı tespiti için yaygın olarak kullanılan KDD99 veri seti üzerinde Karar Ağaçları, ADABOOST, Yapay Sinir Ağları (YSA) ve Destek Vektör Makineleri (DVM) gibi çeşitli modeller kurulmuş ve değerlendirilmiştir. Yapılan değerlendirmeler sonucunda, en yüksek başarı oranına sahip modelin Karar Ağaçları olduğu belirlenmiştir. Bununla birlikte, tüm modellerin başarı oranlarının %99 ve üzerinde olduğu gözlemlenmiştir. Bu yüksek başarı oranları, makine öğrenmesi modellerinin siber saldırı tespitinde etkili olduğunu gösterirken, aynı zamanda aşırı uyumlanma (overfitting) riskini de gündeme getirmiştir.

Bu riskin azaltılması amacıyla, altı farklı deney yöntemi uygulanmış ve veri setine ve modellere farklı açılardan yaklaşılmıştır. Ancak, bu yöntemler sonucunda elde edilen modellerin başarı oranlarının yine %99 ve üzerinde olduğu tespit edilmiştir. Bu durum, veri setinin aşırı uyumlanma riskinin tamamen önlenemediği ihtimalini güçlendirmiştir.

Anahtar Kelimeler: Siber güvenlik, siber saldırı, makine öğrenmesi, KDD99 veri seti

CYBER ATTACK DETECTION WITH MACHINE LEARNING METHODS

ABSTRACT

With the advancement of technology, the concept of cybersecurity has also entered our lives. Over the years, cyber attacks have increased due to various reasons such as financial losses, personal data breaches, and blackmail, leading to an acceleration in efforts to combat these attacks.

This study arose from the need to detect cyber attacks. In the study, various models such as Decision Trees, ADABOOST, Artificial Neural Networks (ANN), and Support Vector Machines (SVM) were developed and evaluated using the widely used KDD99 dataset for cyber attack detection. As a result of the evaluations, it was found that the model with the highest success rate was Decision Trees. However, it was also observed that all models achieved success rates of 99% and above. While these high success rates indicate that machine learning models are effective in detecting cyber attacks, they also raise concerns about the risk of overfitting.

To reduce this risk, six different experimental methods were applied, offering different perspectives on the dataset and models. However, it was found that the models still achieved success rates of 99% and above with these methods. This finding strengthens the possibility that the risk of overfitting could not be fully prevented.

Keywords: Cyber security, cyber attack, machine learning, KDD99 dataset

İÇİNDEKİLER

	Sayfa
YÜKSEK LİSANS TEZİ SINAV SONUÇ FORMU	ii
TEŞEKKÜR.....	iii
ÖZ.....	iv
ABSTRACT.....	v
İÇİNDEKİLER	vi
ŞEKİLLER LİSTESİ.....	viii
TABLolar LİSTESİ	ix
KISALTMALAR	x
BÖLÜM BİR - GİRİŞ	X
BÖLÜM İKİ - MAKİNE ÖĞRENMESİ.....	4
2.1 Makine Öğrenmesi Süreci	6
2.1.1 Veri toplama.....	6
2.1.2 Veri Ön İşleme	6
2.1.2.1 Veri Temizleme.....	6
2.1.2.2 Veri Dönüştürme	7
2.1.2.3 Veri İndirgeme	7
2.2 Makine Öğrenmesi Yöntemleri	8
2.2.1 Karar Ağaçları	8
2.2.2 ADABOOST.....	9
2.2.3 Yapay Sinir Ağları (YSA)	10
2.2.4 Destek Vektör Makineleri (DVM)	11
2.3 Değerlendirme Ölçütleri.....	12
2.3.1 Doğruluk (Accuracy)	13
2.3.2 Kesinlik (Precision)	14
2.3.3 Duyarlılık (Sensitivity)	14
2.3.4 Özgüllük (Specificity)	14

2.3.5 F1- skoru	14
2.3.6 K-Katlı Çapraz Doğrulama	15
BÖLÜM ÜÇ - SİBER GÜVENLİK.....	16
3.1 Siber Güvenlik Nedir?.....	16
3.1.1 Siber Saldırı Tipleri	17
3.1.1.1 Hizmet Engelleme Saldırıları (DoS)	17
3.1.1.2 Kullanıcıdan Kök Kullanıcıya Saldırıları (U2R)	18
3.1.1.3 Uzaktan Yerel Kullanıcıya Saldırıları (R2L)	19
3.1.1.4 Bilgi Tarama Saldırıları	19
BÖLÜM DÖRT - VERİ VE UYGULAMA.....	21
4.1 Veri Seti.....	21
4.2 Veri Setinde Uygulanan Ön İşleme Adımları.....	23
4.3 Modelleme	26
4.3.1 DeneySEL Çalışmalar	26
4.3.1.1 Özellik Seçimi Yapılmadan	27
4.3.1.2 Farklı Örnekler Oluşturma	29
4.3.1.3 Yinelenen Değerlerin Silinmesi	31
4.3.1.4 Korelasyon Matrisi	32
4.3.1.5 Makale Referansı ile Özellik Seçimi	35
4.3.1.6 Lojistik Regresyon	37
BÖLÜM BEŞ - SONUÇ	38
KAYNAKLAR	39

ŞEKİLLER LİSTESİ

	Sayfa
Şekil 2.1 Makine öğrenmesi	5
Şekil 2.2 Karar ağaçları (Nilsson, 1998)	8
Şekil 2.3 K-katlı doğrulama	15
Şekil 4.1 Sınıf dağılımı	22
Şekil 4.2 Karar ağaçları hata matrisi	27
Şekil 4.3 ADABOOST hata matrisi	27
Şekil 4.4 DVM hata matrisi	28
Şekil 4.5 YSA hata matrisi	28
Şekil 4.6 Korelasyon Matrisi	33
Şekil 4.7 Karar ağaçları hata matrisi	34
Şekil 4.8 ADABOOST hata matrisi	34
Şekil 4.9 Karar ağaçları hata matrisi	36
Şekil 4.10 ADABOOST hata matrisi	36

TABLÖLAR LİSTESİ

Sayfa

Tablo 2.1 Hata matrisi	13
Tablo 2.2 Hata matrisine ait değerler	13
Tablo 4.1 Veri setine ait etiketler	22
Tablo 4.2 Özellik tanımları	23
Tablo 4.3 Değerlendirme ölçütleri	28
Tablo 4.4 Karar ağaçları modeli sonuçları	29
Tablo 4.5 ADABOOST modeli sonuçları	30
Tablo 4.6 DVM modeli sonuçları	30
Tablo 4.7 YSA modeli sonuçları.....	30
Tablo 4.8 Karar ağaçları modeli sonuçları	31
Tablo 4.9 ADABOOST modeli sonuçları	31
Tablo 4.10 DVM modeli sonuçları	32
Tablo 4.11 YSA modeli sonuçları.....	32
Tablo 4.12 Seçilen özellikler	35
Tablo 4.13 Çalışma süresi.....	37

KISALTMALAR

DDoS : Dağıtılmış Hizmet Reddi

MFA : Çok Faktörlü Kimlik Doğrulama

IDS/IPS : Saldırı Tespit ve Önleme Sistemleri

FCFB : Hızlı Korelasyon Tabanlı Filtre

DoS : Hizmet Engelleme Saldırıları

U2R : Kullanıcıdan Kök Kullanıcıya Saldırıları

RL2 : Uzaktan Yerel Kullanıcıya Saldırıları

DVM : Destek Vektör Makineleri

YSA : Yapay Sinir Ağları

BÖLÜM BİR

GİRİŞ

Siber güvenlik kavramı, gelişen teknoloji ve hayatlarımızın neredeyse her alanını kaplayan yazılımların gelişimi ile birlikte önem kazanmaktadır. CrowdStrike'ın 2024 Küresel Tehdit Raporu'na göre, bulut tabanlı ihlallerde %75'lik bir artış ve kimlik tabanlı tehditlerde önemli bir yükseliş gözlemlenmiştir. Bu gelişmeler, siber saldırganların daha hızlı ve gizli operasyonlar yürüterek güvenlik sistemlerini etkili bir şekilde aşabildiklerini ortaya koymaktadır (CrowdStrike, 2024).

Siber saldırıların bu denli artmasının arkasında çeşitli amaçlar yatmaktadır. Bu amaçlar arasında finansal kazanç elde etmek, casusluk ve bilgi toplamak, itibar zedelemek, hizmeti kesintiye uğratmak, ideolojik ve politik amaçlar ve rekabet avantajı sağlamak bulunmaktadır. Örneğin, Dağıtılmış Hizmet Reddi (DDoS) saldırıları, hedef sistemleri veya ağları yoğun trafikle doldurarak erişilemez hale getiren yaygın bir siber saldırı türüdür. Bu saldırılar, özellikle büyük ölçekli kuruluşlarda ciddi hizmet aksaklıklarına neden olabilir. 2023 yılında Google ve Amazon, tarihteki en büyük DDoS saldırılarından biriyle karşı karşıya kalmıştır. Bu saldırı, Ağustos ayında başlamış ve Ekim ayına kadar tam olarak kontrol altına alınamamıştır. Saldırı sırasında saniyede 398 milyon talep gönderildiği rapor edilmiş ve bu durum internet trafiğinde ciddi kesintilere neden olmuştur (World Economic Forum, 2023).

Bilişim sistemlerindeki bu gibi saldırıları önlemek ve bilgi güvenliğini sağlamak için çeşitli yöntemler geliştirilmiştir. Bunlar arasında güvenlik politikaları, şifre politikaları, veri şifreleme, çok faktörlü kimlik doğrulama (MFA), saldırı tespit ve önleme sistemleri (IDS/IPS), yedekleme, siber tehdit istihbaratı ve yasal uyumluluk yer almaktadır. Saldırı tespit sistemleri açısından KDD99 ve NSL-KDD veri setleri önemli bir yere sahiptir. KDD99 veri seti, siber güvenlik ve ağ saldırı tespit sistemleri araştırmalarında en yaygın olarak kullanılan veri setlerinden biridir. Bu veri seti, 1998 DARPA IDS değerlendirme programından türetilmiş olup, akademik ve endüstriyel araştırmalarda standart bir kıyaslama ölçütü olarak kabul edilmiştir (Tavallae vd., 2009).

Bu tezde çalışmasında KDD99 ve NSL-KDD veri setleri kullanılarak saldırı tespit sistemleri kurulmuş ve karşılaştırmaları yapılmıştır. Tavallae vd. (2009) çalışmasında, KDD99 veri setinin aşırı derecede tekrar eden kayıtlar ve dengesiz veri dağılımı gibi önemli eksikliklerinden bahsedilmektedir. NSL-KDD veri seti ise, KDD99 veri setinin içerdiği bazı temel sorunları çözmek amacıyla önerilmiştir. Aynı çalışmada, KDD99 veri setinin çok sayıda tekrar eden kayda sahip olması ve bu nedenle değerlendirilen sistemlerin performansını olumsuz yönde etkilemesi gibi sorunlara dikkat çekilmiş ve bu sorunları çözmek amacıyla NSL-KDD veri seti geliştirilmiştir (Tavallae vd., 2009). Bu çalışma, bu iddiaların doğruluğunu test etmek amacıyla, karar ağaçları, yapay sinir ağları, destek vektör makineleri ve ADABOOST makine öğrenmesi yöntemlerini her iki veri seti için de kullanarak karşılaştırmalar yapmayı hedeflemektedir.

Özgür ve Erdem (2012), siber saldırı tespitinde makine öğrenmesi algoritmalarının etkinliğini araştırmak için yaptıkları çalışmada KDD99 veri setini kullanmış ve bunu yaparken karar ağaçları, yapay sinir ağları, destek vektör makineleri ve AdaBoost yöntemlerini kullanarak bu yöntemlerin etkinliğini karşılaştırmıştır. Ayrıca bu çalışmada, KDD99 veri seti üzerinde kullanılan algoritmaların performanslarının gerçek sistemlerde yapılacak olan denemelerde aynı başarıyı göstermeyeceğine dair iddiasını da taşır.

Sapre, Ahmadi ve İslam (2019) çalışmasında, KDD99 ve NSL-KDD veri setleri kullanılmıştır. Bu veri setleri ile siber saldırı tespit performanslarını değerlendirmek için naif bayes, destek vektör makineleri, rastgele orman ve yapay sinir ağları makine öğrenmesi yöntemleri kullanılmıştır. Bu çalışmada KDD99 veri setinin yüksek tekrar içermesi ve kolay sınıflandırılabilir olması nedenleri ile yüksek doğruluk sonuçları verdiği; NSL-KDD veri setinin ise daha az tekrar eden veri içermesi ve KDD99 veri seti kadar kolay sınıflandırılabilir olmadığı için daha düşük doğruluk sonuçları verdiğini göstermiştir.

Siber saldırı tespitinde kullanılan makine öğrenmesi yöntemleri ile sınıflandırma algoritmalarının sonuçları ve parametreleri ile karşılaştırılan bu çalışmada Microsoft tarafından sağlanan veri seti kullanılmıştır.. Bu çalışmada naive bayes, karar ağacı, rastgele orman, AdaBoost, LightGBM algoritmaları kullanılmıştır. Sonuçların

performansı değerlendirildiğinde LightGBM algoritmasının en yüksek doğruluk değerini verdiği ortaya konmuştur (Güleş, 2020).

NSL-KDD veri setinin kullanıldığı başka bir çalışmada özel olarak sadece yapay sinir ağları makine öğrenmesi yöntemi kullanılmıştır. Yapay sinir ağları yöntemi ile performans analizi, ikili sınıflandırmada ve beşli sınıflandırmalarda kullanılarak karşılaştırmalar yapılmıştır. Bu sınıflandırmalar verinin içerdiği saldırı tipleri sınıflandırmalarıdır. Her iki sınıflandırmanın da yüksek tespit oranlarına sahip olduğu bulunmuştur (Ingre ve Yadav, 2015).

NSL-KDD veri setinde makine öğrenmesi yöntemlerinin kullanılarak saldırı tespit oranlarının karşılaştırıldığı başka bir çalışmada Naive Bayes, Gizli Bayes ve NBTtree algoritmaları kullanılmıştır. Bu çalışmada sınıflandırma performansını yükseltmek için veri ön işleme basamaklarında öznitelik seçimlerinde ve ayrıklaştırma yöntemleri kullanılarak doğruluk oranlarında artış olduğu gösterilmiştir. NSL-KDD veri setinin boyutu Hızlı Korelasyon Tabanlı Filtre (FCBF) algoritması kullanılarak azaltılmıştır (Deshmukh, Ghorpade ve Padiya, 2015).

UNSW-NB 15 and CI-CIDS2017 veri setleri ile kullanılan; Tait vd. (2021) tarafından yapılan çalışma da Ingre ve Yadav (2015) tarafından yapılan çalışmaya benzer olarak ikili ve çoklu sınıflandırmalar yapılmıştır. Görselleştirilen ikili ve çoklu sınıflandırmalarda daha başarılı sonuçlar veren ikili sınıflandırma yöntemleri olmuştur. Bu çalışmada destek vektör makineleri, k-en yakın komşu ve rastgele orman makine öğrenmesi yöntemleri uygulanmıştır. Ek olarak; çalışmanın sonuçlarına göre, en iyi sonuç veren makine öğrenmesi yöntemini de sınıflandırma farklılıklarına göre değişiklik göstermiştir.

Almseidn vd. (2017) çalışmasında Naive Bayes, rastgele orman, karar ağaçları, Çok katmanlı algılayıcı gibi makine öğrenmesi yöntemleri kullanarak makine öğrenmesi sınıflandırıcılarının doğruluk oranlarını ortaya koymuştur. KDD99 veri setinde yapılan bu çalışma, karşılaştırmaları yaparken doğruluk oranları, yanlış negatif ve yanlış pozitif oranları üzerinden ilerlediğini belirtmiştir.

BÖLÜM İKİ

MAKİNE ÖĞRENMESİ

Öğrenme kavramı, “bilgi, anlama veya beceri kazanma” ve “deneyimle bir davranış eğilimini değiştirme” olarak betimlenen sözlük tanımına sahiptir. Makine öğrenmesi tanımlaması ancak bilgisayarların verilerden öğrenerek ve deneyimlerden yola çıkarak kendini geliştirme yeteneğidir. Bunun ortaya çıkabilmesi için; belirli görevleri yerine getirmek amacıyla programlanan sistemlerin, verilerden örüntüler ve korelasyonlar öğrenerek performanslarını artırmaları gerekir (Mitchell, 1997). Yapay zeka alanının alt dallarından biri olan makine öğrenmesi, bilgisayarların insan müdahalesi olmadan kendilerini geliştirmeleri hedefiyle kullanılır.

Makine öğrenmesinin önemli bir alan olmasının başlıca nedenleri mevcuttur. İnsanlar, ellerindeki çoklu, büyük boyutlu veriler ya da bilgiler arasındaki tüm değişkenleri, bilgiler arası koşulları ya da korelasyonları tam olarak gözlemleyemeyebilir. Makine öğrenmesi yöntemleri kullanılarak bu koşullar ya da korelasyonlar veri madenciliği ile ortaya koyulabilir. Ek olarak yüksek boyutlardaki bilgilerin kodlanması pratik ya da verimli değildir. Makine öğrenmesi yöntemleri ile yüksek boyuttaki verilerin işlevlerini, ilişkileri sınırlayarak kullanılabilir. Bu yapay zeka alt dalı, makinelerin iş başında kendileri performanslarını arttırmalarına ve değişen ortamlara başarılı adaptasyon sağlamalarına olanak tanır (Russell ve Norvig, 2021).

Teknolojinin hızla geliştiği, yeni bilgi keşiflerinin sürekli arttığı ve yeni olayların yaşandığı dünyada; yapay zeka yöntemlerinin devamlı, tekrarlı tasarımı verimli değildir. Makine öğrenmesi yöntemleri kullanılarak, yeni bilgiler ve yeni olaylar kolaylıkla izlenebilir ayrıca sistemlerin güncel kalması da sağlanabilir. Bu nedenlerle de makinelerin daha etkili bir biçimde çalışması, var olan bilgilerin kullanımı ve yeni bilgi entegrasyonu ile sağlanır (Goodfellow, Bengio ve Courville, 2016). Şekil 2.1’de makine öğrenmesi ile ilgili bir şema gösterilmiştir.



Şekil 2.1 Makine öğrenmesi

Makine öğrenmesi yöntemleri üç ana başlık grubuna sahiptir. Bunlar; denetimli öğrenme, denetimsiz öğrenme ve pekiştirmeli öğrenme. Denetimli öğrenme yönteminde veriler etiketli bir şekilde bulunur. Bu yöntemde etiketli veri giriş ve çıktılar arası ilişkileri saptama ve ortaya çıkartmak asıl amaçtır. Bu sayede de her yeni veri girdisinin doğru çıktısı tahmin edilir. Modelin başarı oranını yükseltmek amacıyla giriş verileri ve beklenen çıktılar arası fark en aza indirilmeye çalışılır (Mitchell, 1997). lineer regresyon, karar ağaçları, rastgele orman, k-en yakın komşu, naif bayes, destek vektör makineleri algoritmaları denetimli öğrenme algoritmaları örneklerindendir.

Denetimsiz öğrenme, etiketlenmemiş veriler ya da kategoriler içerir. Bilgisayara girilen bu etiketsiz verilerin arasındaki örüntü, kategori ya da korelasyonları saptamak ve buna göre çalışması beklenir. Denetimsiz algoritmaların asıl araçları kümeleme, veri boyutu azaltmak ve veriler arası ilişkileri ortaya koymaktır. Benzer ya da ilişkili özelliklere sahip olan verileri gruplandırarak veri setini anlaşılır hale getirmek kümeleme algoritmaları, verinin ana yapısını değiştirmeden veri azaltmak da veri boyutunu azaltmaktır (Goodfellow, Bengio ve Courville, 2016).

Pekiştirmeli öğrenme, bir ajanın belirlenen ortamdaki davranışlarının bir anlamıyla ödül ceza mantığı olarak öğrenilmesidir. Yani, ortama uygun, doğru eylem ödüllendirilir (pekiştirilir) ve hatalar cezalandırılarak uzun dönemde maksimum verimin alınması sağlanır. Bu yöntemin başarılı olduğu alanlar, robotik ve oyun teorisi alanlarıdır ve bu alanlarda etkili uygulamalar ortaya koyar (Sutton ve Barto, 2018).

2.1 Makine Öğrenmesi Süreci

Bu bölümde makine öğrenmesi yöntemleri kurulmadan önceki kritik öneme sahip basamaklar olan veri toplama ve veri ön işleme adımları detaylandırılacaktır.

2.1.1 Veri toplama

Veri toplama işlemi, model kurma sürecinin başlangıç adımı olarak bilinir. Bu adımda; araştırılmak , analiz edilmek ya da tespit edilmek istenen durum ya da probleme bağlı olarak ilişkili kaynaklar araştırılır. Verilerin kaynaklarından alınıp uygun ortam ve biçimlerde kaydedilme işlemi olan temel basamaktır.

2.1.2 Veri Ön İşleme

Veri ön işleme, veri setleri kullanılarak yapılan özellikle yapay zeka alt dalları yöntemlerinin uygulanmasında oldukça önemli bir adımdır. Bu adım verinin modellenmesinden önce yapılan ve modelin çalışma başarısını yüksek ölçüde arttıran süreçleri içerir. Titizlikle uygulanan veri ön işleme yöntemleri kurulan modeller ve çalışma sonuçlarında kritik öneme sahiptir.

2.1.2.1 Veri Temizleme

Veri setinin doğruluğu ve tutarlılığı için yapılan bu adım eksik verilerin tamamlanması ya da silinmesi, hata, eksik ve yanlışların tespiti ve gereksiz verilerin kaldırılması aşamaları; veri temizleme basamaklarıdır.

Kullanılan veri setlerinin çok boyutlu ya da karmaşık olması makine öğrenmesi alanındaki çalışmalarda da modeller kurmayı zorlaştırır. Veriler kayıp, yinelenen, sonsuz ya da aykırı değerler içerebilir. Bu değerler yoğunluklarına da bağlı olarak, çalışmanın doğruluk ve güvenilirlik gibi sonuçlarını olumsuz yönde etkiler. Tüm bu nedenlerle de bahsedilen değerler farklı yöntemlerle veriyi etkileyemeyecek hale getirilir. Bu yöntemler; eksik verilerin silinmesi, eksik değerler yerine değişkenin ortalamasının kullanılması, eşdeğer kategorilerde bulunan tüm veriler için ortalama değer kullanılması, mevcut verilerin detaylı incelenerek en ideal değerle doldurmadır (Roiger ve Geatz, 2003).

2.1.2.2 Veri Dönüştürme

Veri dönüştürme işlem dizisi ile veriler, model kurulmaya ve veri madenciliği yapılmasına daha uygun hale gelir. Normalizasyon, ölçekleme, kategorik verilerin dönüştürülmesi gibi adımlar en bilinen veri dönüştürme yöntemleridir.

Bu yöntemlerden normalizasyon, verinin ortalama değer ve standart sapma değerlerinin bulunması, ardından da bu sonuçlar ile normal dağılıma yaklaşır. Bazı makine öğrenmesi yöntemleri normal dağılıma yakın veri kümeleri ile daha başarılı sonuçlar ortaya koyar.

Başka bir yöntem olan standardizasyon yöntemi, veri setinin her bir özelliğinin z-puanının hesaplanarak özelliklerin ortalamalarının 0'a, standart sapma değerlerinin ise 1'e çevrilmesidir. Bu adım sayesinde özellikle uzaklık temelli modellerde başarı oranı artırılır ve modelin kurulma sürecinde zaman tasarrufunun sağlanması, karşılaştırma başarısının artması sağlanır.

Normalizasyon ve standardizasyon işlemleri detaylandırılan veri dönüştürme adımlarının uygulanması, modellerin başarı oranlarını yükseltir, daha anlamlı ilişkiler kurulmasını ve bu ilişkilerin incelenmesini sağlar ve sağlıklı görselleştirmeler sağlar.

2.1.2.3 Veri İndirgeme

Veri setlerinin boyutunun azaltılması ve hacminin küçültülmesi hedeflenerek yapılan işlemler toplamı veri indirgeme yöntemleri olarak tanımlanır. Veri indirgeme yöntemleri arasında özellik seçimi, özellik çıkarımı, örnekleme ve kümeleme bulunur.

Özellik seçimi yöntemi veri setinin özelliklerinin ve veri setindeki etkisinin iyi analiz edilmesine bağlıdır. Ortaya konulmak istenen analize ya da tespit edilmek istenen probleme bağlı olarak veri setindeki önemli özelliklerin belirlenerek kalan özelliklerin veri setinden çıkartılması işlemidir. Önemli özelliklerin belirlenmesi istatistiksel yöntemler kullanılarak ya da algoritma yapıları kurularak belirlenir.

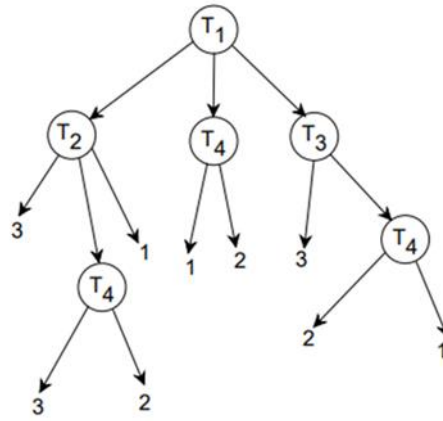
Veri setinin hacminin küçültülmesine bağlı olarak artan hız ve azalan maliyetler, başarılı genelleme oranları, depolama alanının azalması gibi nedenlerle veri ön işleme yöntemi olan veri indirgeme metotları önemli bir adımdır.

2.2 Makine Öğrenmesi Yöntemleri

Literatürde sınıflandırma ve karar alma süreçlerinde kullanılan pek çok makine öğrenmesi yöntemi vardır. Bu çalışma kapsamında kullanılacak dört yöntem ana hatları ile aşağıda açıklanmıştır.

2.2.1 Karar Ağaçları

Makine öğrenmesi ve veri madenciliği alanlarında sıkça kullanılan karar ağaçları algoritmaları güçlü bir regresyon ve sınıflandırma yöntemidir. Çalışma biçimi açısından karar ağaçları, veri集中的 verileri belirli özelliklere göre sahip olma durumuna göre analiz ederek başarılı sonuçlar vermeye çalışır. Biçim, verinin bir kök düğüme sahip olması ve her düğümün belirli bir özelliğe göre dallara ayrılması şeklindedir. Verinin belirli bir özelliğe göre kategorize edilmesi dallanan her bir yol sayesinde sağlanır ve dallanan yollardan ortaya çıkan son yaprak düğümündeki sonuç da veri örneğinin kategorize edilmesini belirlemektedir (Cios vd., 2007). Şekil 2.2’de karar ağaçları algoritmasına ait bir ağaç temsili gösterilmektedir.



Şekil 2.2 Karar ağaçları (Nilsson, 1998)

Yorumlanmasının ve yapılandırılmasının yani anlaşılabilir olma gibi temel avantajlar karar ağaçları yöntemini makine öğrenmesi yöntemleri içerisinde önemli bir konuma getirir. Bu avantaj; karar ağaçlarının bilgisayar ağlarındaki saldırıları tespit etmek için kullanıldığında hızlı ve doğru bir şekilde normal ve anormal ağ trafiği örüntülerinin ayrıştırılabildiği örneğinde gösterilmiştir (Quinlan, 1986). Verinin nasıl sınıflandırıldığı, belirli bir özelliğe bağlı olarak veri kümesini bölen dallar sayesinde ortaya koyulur. Ek olarak karar ağaçları yöntemi, kendisini esnek ve fonksiyonel hale getiren hem kategorik hem de sayısal verilerle çalışabilme özelliğine sahiptir (Mitchell, 1997). Karar ağaçları yöntemini gerçek zamanlı deneyler ve çalışmalar için optimal hale getiren de hatalara karşı dayanıklılık ve eksik veri değerleri ile iyi çalışabilme özellikleridir (Nilsson, 1998).

Karar ağaçları, sağlık alanında tanı koymaktan finans alanında risk tespitlerine kadar geniş bir aralıkta kullanılabilir bir yöntemdir. Bu denli geniş bir kullanım aralığına sahip olmasının başlıca nedenleri yüksek boyutlu ve karmaşık veri setlerindeki başarısı ve verimliliğidir.

Siber saldırı tespit sistemlerinde karar ağaçlarının kullanımı, imza tabanlı ve anormallik tabanlı tespit yöntemleri ile bağlantılıdır. Bu iki tespit yöntemi, depolanan ya da modellenen verilere uygulanan yaklaşıma bağlı olarak farklılaşır. Örneğin imza tabanlı yöntemlerde, daha öncesinden saldırı imzaları daha öncesinden de bilindiği için depolanmış halde bulunur. Bu yöntemde gelen ağ trafiği, depolanan imzalarla karşılaştırılır ve sonuç döner. Fakat anormallik tabanlı yöntemlerde, modellenen normal ağ davranışları ile gelen ağ trafiği karşılaştırılır ve sapma değerler anomali olarak döner. Karar ağaçları yöntemi esnek ve etkili bir tespit mekanizması olma özelliğini bu iki yaklaşımı da destekleyerek sağlar.

2.2.2 ADABOOST

Bir topluluk öğrenme yöntemi olan ADABOOST, güçsüz öğrencileri güçlü modellere çevirmek için kullanılmaktadır. Freund ve Schapire tarafından geliştirilen bu yöntem, güçsüz öğrencilerin öncelikli modellerinin arka arkaya oluşturulmasıdır. Bu biçim sayesinde bir sonraki adım bu hatalı modellenen öğrenciye daha da önem vererek ilerler. Modelin temel mantığı, her iterasyonda, bir önceki öğrencinin hataları

tespit edilip geliştirilir. Güçsüz öğrencilerin modellenme hatalarını en aza indirerek sonuç çıktısı olarak güçlü bir modelin kurulması, ADABOOST yönteminin ana amacıdır (Mitchell, 1997). Hem sınıflandırma hem de regresyon modellemeleri için başarılı çalışmaktadır (Solomatine ve Shrestha, 2004).

ADABOOST yöntemi, esnek çalışma biçimi ve modellediği başarı sonuçları ile büyük bir kullanım alanına sahiptir. Sınıflandırma başarısı, güçsüz öğrencilerle ortak kullanım durumunda belirgin bir artış gösterebilir (Cios vd., 2007). Yöntemin çalışma prensibi, güçsüz öğrenci tahminlerini bir araya getirerek modelin son halini ortaya koymak için bu tahminlerin ağırlıklı ortalamasını almaktır. Bu sayede algoritmanın çalışma performansı artar, karışık verilerle çalışmada daha dayanıklı hale gelir ve dengesiz verilerde de iyi sonuçlar verebilir (Nilsson, 1998).

Sınıflandırma doğruluğunu arttıran bir yöntem olması ADABOOST yönteminin en önemli avantajlarından biridir. Ancak bu avantajlı özelliği aynı zamanda aşırı uyum (overfitting) sorununun ortaya çıkma ihtimalini arttırabilir ve veri setinden veri setine değişkenlik gösterse de hesaplama zamanı yavaş olabilir. Yine de ADABOOST yöntemi, çeşitli ve geniş yelpazede kullanım alanlarına sahiptir. Başarılı model sonuçları, esnekliği ve çeşitlenmiş kullanım alanları ile makine öğrenmesi alanında ve bu bağlamdaki akademik ve pratik çevrede sıklıkla başvurulanan yöntemlerden biridir (Schapire, 2001).

2.2.3 Yapay Sinir Ağları (YSA)

Yapay sinir ağları; karmaşık verileri çözümlemede kullanılan, biyolojik sinir sistemlerinden model alınarak geliştirilmiş olan güçlü hesaplama modelleridir. Ana çalışma biçimi olarak yapay sinir ağları birbirleriyle ilişkili yapay nöronlardan oluşur. Bu nöronlar, bilgiyi işleyip aktararak öğrenme ve karar verme işlevlerini yerine getirirler. Nöronları kullanan bu ağlar makine öğrenmesi yönteminde, çoklu temel birimler ve bu birimlerin birbirleri ile olan ilişkileri doğrultusunda bilgi işlenir (Öztemel, 2012). Yapay nöronlar, çıktıları kullanılması istenen verilerin alınıp belirli bir süreç sayesinde işlenmesi sayesinde üretirler. Veriler arasındaki her bir ilişkinin ağırlık değeri atanır ve ağırlık öğrenme oranı da atanan ağırlıkların başarılı bir şekilde düzenlenmesine bağlıdır (Özgür ve Erdem, 2012). Yapay sinir ağlarının, veriler ve

başarı oranları arasındaki komplike bağlantıları öğrenmesine ve algoritma kurmasına olanak sağlayan da çalışma yöntemidir. Yapay sinir ağları, çok seviyeli bir yapıda olabileceği gibi bu durum sayesinde de doğrusal olmayan ilişkilerin algoritmasını kurabilir (Kaya ve Yıldız, 2014).

Sınıflandırma, ilişkilendirme ve tahminleme gibi problemlerde etkili olan yapay sinir ağları, bu yetenekleri sayesinde de farklı alanlarda ve geniş yelpazede kullanılmaktadır. Sağlık alanında tanı koyma, görüntü işleme ve ses tanıma, yatırım tahminlemeleri, endüstriyel otomasyon gibi alanlar bu geniş yelpazenin örneklerindendir. İşlenmesi güç verilerde ve çözümlenmesi zor ilişkilerde yöntem kurma yeteneği yüksek olan yapay sinir ağları, konvansiyonel hesaplama yöntemlerinin başarısız olduğu durumlarda başvurulacak etkili bir yöntemdir. Yapay sinir ağları aynı zamanda anormallik tespiti gibi yüksek başarı oranı yakalanması zor alanlarda etkili sonuçlar yakalayabilmektedir. Bu özelliği bağlamında da farklı saldırı türlerini tespit etme hedeflerinde iyi sonuçlar yakalamaktadır. Yüksek boyutlu veri setlerinde ilişkileri anlama ve işlemede başarılı tahmin sonuçları ortaya koyan yapay sinir ağları, bu özelliklerinin de etkisi ile veri analitiği ve yapay zeka alanlarında yapılan çalışmalarda oldukça popülerdir. (Cios vd., 2007).

2.2.4 Destek Vektör Makineleri (DVM)

Denetimli bir öğrenme algoritması olan Destek Vektör Makineleri, sınıflandırma ve regresyon problemleri ile ilgili çalışmalarda sıklıkla kullanılan ve yüksek başarı sonuçları elde edebilen bir yöntemdir. Verilerin çok boyutlu bir uzaya taşınması ve bu uzayı en başarılı şekilde ayıran hiperdüzlemi keşfedip bu hiperdüzlemde çalışma biçimi, destek vektör makinelerinin temel çalışma prensibidir. Destek vektör makineleri yönteminin bu çalışma prensibindeki amacı uzayda ayırdığı iki sınıf arasında en geniş marjini sağlamaktır. Bu hedefle en geniş marjini sağlayan hiperdüzlem yaratılmış olur. Burada bahsi geçen marjin kavramı da destek vektörleri olarak tanımlanan yani uzayda ayrılmış olan iki sınıfın arasındaki en yakın veri noktalarına olan uzaklıktır (Cortes ve Vapnik, 1995). Bu çalışma prensibi biçimi ile de destek vektör makinelerinin genelleştirme becerisi yüksektir ve yeni verilerle çalışırken de iyi başarı sonuçları verir.

Doğrusal olarak ayrılabilir olan ve olmayan veri setleri için de iyi sonuçlar ortaya çıkartabilen destek vektör makineleri bu sayede de çalışma alanını genişletir. Doğrusal olarak ayrılmayan veri setlerinin sayısal oranı doğrusal olarak ayrılan veri setlerine göre oldukça fazladır. Doğrusal olarak ayrılmayan veri setlerinde destek vektör makineleri çekirdek fonksiyonları kullanarak çalışır. Çekirdek fonksiyonu kullanımı veriyi çok katmanlı bir uzaya dönüştürüp doğrusal olmayan ayrımları da ortaya çıkartmaya olanak tanır (Soman, Loganathan ve Ajay, 2011). Çekirdek fonksiyonları çeşitli türlerde olabileceği gibi en yaygın olanları, polinom, lineer, sigmoid, radyal bazlıdır. Radyal tabanlı fonksiyonların özel olarak da iyi başarı performansı gösterdiği de bilinmektedir (Ayhan ve Erdoğan, 2014). Modelin performansı ile doğrudan ilişkili olduğu için kullanılacak çekirdek fonksiyonu dikkatli seçilmelidir.

Dayanıklılığı yüksek olan destek vektör makineleri, veri setine modelin aşırı uyumlanması (overfitting) sorununu engelleyerek veri setlerinin komplike yapısını başarılı bir şekilde başa çıkmasını sağlar (Vapnik, 1998). Boyutu yüksek verilerde iyi performans gösterme ve başarılı genelleme, destek vektör makinelerinin önemli özelliklerindendir. Tüm bu gibi önemli nedenler itibarı ile de destek vektör makineleri biyoinformatik, veri madenciliği, yüz tanıma, ve metin madenciliği, siber saldırı tespiti gibi alanlarda sıklıkla kullanılır.

2.3 Değerlendirme Ölçütleri

Bu bölümde sınıflandırma modellerinin karşılaştırılmasında kullanılan ve hangi modellerin daha iyi sonuç verdiğini yorumlamamıza yarayan değerlendirme ölçütlerinden bahsedilecektir. Bölüm Dört'te kullanılan veri seti üzerinde kurulan modellerin sonuçları aşağıda açıklanan değerlendirme ölçütlerine göre karşılaştırılmıştır. Bu ölçütleri açıklamadan önce hata matrisi (confusion matrix) kavramını anlamakta fayda vardır. Hata matrisi, verileri sınıflandırma amacıyla kurulan modellerin performanslarını gösteren bir tablodur. Modelin tahmin performansını gerçekte pozitif ve negatif olarak bulunan sınıflar üzerinden gösterir. Modellerin performans ölçütlerini hesaplarken kullanılan hata matrisi (confusion matrix) ve ona ait değerlerin açıklamaları Tablo 2.1 ve Tablo 2.2'de gösterilmiştir.

Tablo 2.1 Hata matrisi

		Gerçek Sınıf	
		Pozitif	Negatif
Tahmin Sınıfı	Pozitif	Gerçek Pozitif (GP)	Yanlış Pozitif (YP)
	Negatif	Yanlış Negatif (YN)	Gerçek Negatif (GN)

Tablo 2.2 Hata matrisine ait değerler

GERÇEK POZİTİF (GP)	Doğru tespit edilen pozitif sonuçları temsil eder.
GERÇEK NEGATİF (GN)	Doğru tespit edilen negatif sonuçları temsil eder.
YANLIŞ POZİTİF (YP)	Negatif sınıfın yanlış bir şekilde pozitif olarak tespit edilmesidir.
YANLIŞ NEGATİF (YN)	Pozitif sınıfın yanlış bir şekilde negatif olarak tespit edilmesidir.

Tablo 2.1’den faydalanılarak hesaplanan değerlendirme ölçütlerinden Doğruluk, kesinlik, duyarlılık, özgüllük ve F1 skoru aşağıda açıklanmıştır.

2.3.1 Doğruluk (Accuracy)

Modelin hangi aralıklarla doğru tahminlemede bulunduğunu gösteren metriğe doğruluk denir. Doğruluk, modelin pozitif ve negatif sınıfların tahminlenmesindeki başarı oranıdır. Gerçek pozitif ve gerçek negatif değerlerin toplam tahmine oranı denklem 2.1’deki gibi hesaplanabilir.

$$\frac{GP+GN}{GP+GN+YP+YN} \quad (2.1)$$

2.3.2 Kesinlik (Precision)

Karmaşık ve dengesiz veri setlerinde kullanımı yaygın olan kesinlik, modelin pozitif tahminlerinin doğruluk oranıdır. Kurulan modelin gerçek pozitif tahmin oranı arttıkça kesinlik oranında da artış gerçekleşir. Gerçek pozitif tahminlerin toplam pozitif tahminlere oranı Denklem 2.2’de gösterilmiştir.

$$\frac{GP}{GP+YP} \quad (2.2)$$

2.3.3 Duyarlılık (Sensitivity)

Geri çağırma (recall) olarak da bilinen duyarlılık, modelin pozitif sınıfın başarılı tahmin edilme oranıdır. Kurulan modelin gerçek pozitif tahmin oranı arttıkça duyarlılık oranında da artış gerçekleşir. Gerçek pozitif örneklerin pozitif olarak sınıflandırması olan duyarlılık; doğru tahmin edilen pozitiflerin, gerçek pozitive olan oranı olarak hesaplanır.

$$\frac{GP}{GP+YN} \quad (2.3)$$

2.3.4 Özgüllük (Specificity)

Negatif sınıf tahminin başarı göstergesi olan özgüllük, gerçek negatif örneklerin başarılı bir şekilde negatif olarak sınıflandırılmasıdır. Doğru negatif tahminlerin, toplam gerçek negatiflere oranlanması ile hesaplanan özgüllük formülü Denklem 2.4’te gösterilmiştir.

$$\left(\frac{GN}{GN+YP} \right) \quad (2.4)$$

2.3.5 F1- skoru

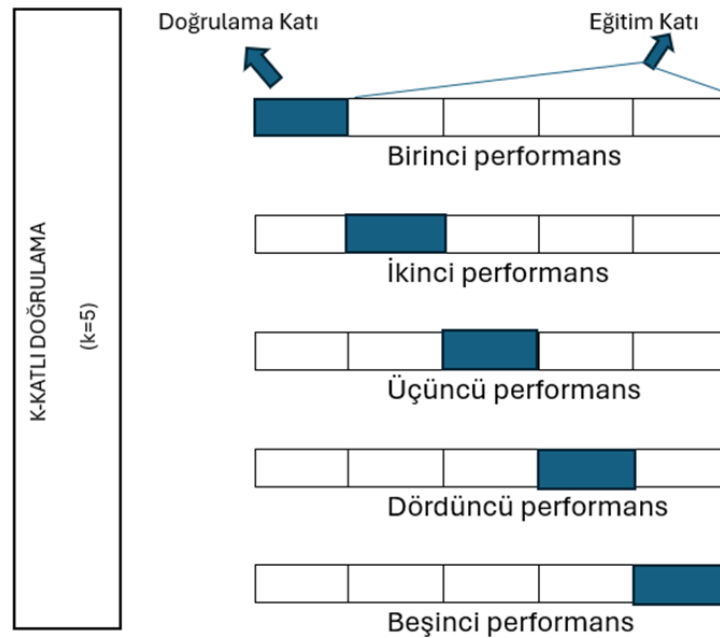
F1-skoru modelin performansının, başarı oranının değerlendirilmesinde önemli bir metrik olarak karşımıza çıkar. Dengesiz veri setlerinde ya da modelin performansının dengelenmesinde kullanılır. Kesinlik ve duyarlılık arasından bir denge kuran F1-skor, bu iki metriğin harmonik ortalamasıyla hesaplanır. Hesaplama formülü aşağıdadır.

$$2 \times \frac{\text{Kesinlik} \times \text{Duyarlılık}}{\text{Kesinlik} + \text{Duyarlılık}} \quad (2.5)$$

2.3.6 K-Katlı Çapraz Doğrulama

Model doğrulama tekniklerinden biri olan k -katlı çapraz doğrulama, veri setinin doğruluğu ve genellemeler için kullanılır. Bu yöntemde k kat sayısı belirlenir (örneğin $k=10$). Veri seti, belirlenen k sayı kadar alt kümeye bölünür ve her bir alt kümede modelleme yapılır ve performans ölçütleri uygulanır. Yapılan modellemelerin sonuçları ve ortalamaları kullanılarak genel model performansı incelenir.

Modelin performansı bakımından başarılı ve esnek bir teknik olan k -katlı çapraz doğrulamanın bazı avantaj ve dezavantajları bulunur. Modelin genelleştirme yeteneğinin artırılması, verinin tamamının kullanılması itibari ile veri kaybının önlenmesi ve dengesiz veri setlerinde başarı oranlarının artırılması anlamında bu yöntem avantajlıdır. Ancak bu yöntem, belirlenen sayı kadar alt küme oluşturma ve bu alt kümelerde modeller kurma işlemlerini uyguladığı için hesaplama maliyeti yüksek olması ve sonuca varma konusunda çalışma süresinin fazla olması gibi dezavantajları da içerir. Şekil 2.3’de k -katlı çapraz doğrulamaya ait bir illüstrasyon gösterilmiştir.



Şekil 2.3 K-katlı doğrulama

BÖLÜM ÜÇ

SİBER GÜVENLİK

3.1 Siber Güvenlik Nedir?

Bilgisayar ve bilgisayar ağları ile ilgili kavramları veya işlemleri tanımak amacıyla kullanılan siber kavramı, ilk kez 1990'lı yıllarda kullanılmıştır. Siber kavramının ilk kullanımı, bilgi teknolojileri sistemlerinin ağı bağlı güvenlik tehlikelerini tanımlamak amacıyla.

Gelişmekte olan teknolojinin önemli durum ve tartışmalarından biri siber güvenlik kavramıdır. Yetkilendirilmemiş erişim, saldırı veya hasar yaratma gibi tehditlere karşı bilgi teknoloji sistemlerinin, ağların, programların ve verilerin korunmasını hedefleyen çalışma ve analiz çalışmalarının tümüne siber güvenlik adı verilir. Bilginin gizliliğin, güvenliğin ve ulaşılabilirliğinin sağlanması siber güvenliğin temel amaçlarındandır (Hekim ve Başbüyük, 2013).

Genel anlamı itibari ile yalnızca teknoloji ile ilişkilendirilen siber güvenlik; sosyal, maddi ve politik durumları da kapsayan bir alandır. Bilginin her aşamasında karşılaşılabilen güvenlik açıkları bireylerin ve kurumların yaşamlarına ve çalışmalarına ciddi zarar verme riski taşır (Özdemirci ve Torunlar, 2018). Tüm bu nedenlerle gelişen teknolojinin kesişimi ile evrensel olarak siber güvenlik üzerine çalışmalar ve siber güvenlik politikaları artmıştır.

Siber güvenlik politikaları da yalnızca teknolojik politikalar ile sınırlandırılmamış olup hukuki düzenlemeler ve eğitim programlarını da önemli hale getirmiştir (T.C. Ulaştırma Denizcilik ve Haberleşme Bakanlığı, 2016-2019). Yukarıda da bahsedildiği üzere siber saldırılar geniş bir uygulama alanına sahiptir. Bu nedenle de bahsi geçen ve diğer pek çok alanda görev yapan uzmanların bu alanla ilgi en yeni bilgileri ve yetenekleri takip edip geliştirmeleri büyük öneme sahiptir (Özdemirci ve Torunlar, 2018).

Siber güvenlik alanının başka bir kritik noktası da uluslararası boyutudur. Siber saldırıların sınırsız yapısı itibari ile ulusal güvenlik stratejileri kurulması ihtiyacı doğar ve uluslararası anlaşmalar da gerekli hale gelir. Bu anlaşmalar, işbirliği ve

koordinasyon sayesinde siber güvenlik politikalarında kritik bir yer tutar. Devletlerin birbirleriyle olan anlaşmaları ve evrensel kuruluşlar aracılığıyla siber güvenlik alanında ortak protokoller geliştirip hayata geçirmektedir (Clarke ve Knake, 2011).

Bu bağlamda, siber güvenlik eğitimleri ve önleyici uygulama çalışmaları evrensel öneme sahiptir. Bireylerden devletlere, kurumlardan uluslararası ilişkilere kadar siber güvenlik önlemleri devamlı güncellenen ve iyileştirilen bir alan olmalıdır. Tüm boyutları (teknolojik, hukuki, sosyal, maddi vb.) ile değerlendirildiği müddetçe siber güvenliğin sağlanması mümkün olacaktır. Teknolojik gelişimin faydaların güvenli bir şekilde kullanılmasına ve devamlı bir teknolojik geleceğin yaratılması, siber güvenliğin yaratılmasında saklıdır (Özdemirci ve Torunlar, 2018).

3.1.1 Siber Saldırı Tipleri

Siber güvenlik kavramını ve önemini anlamak bakımından siber saldırıların ne amaçla ve nasıl gerçekleştirildiğini anlamak önemlidir.

Siber saldırılar, rekabette avantaj elde etmek, belirli gruplara zarar verme, dikkat çekme, kişisel tatmin, hizmeti kesintiye uğratma, şantaj ve tehditi maddi kazanç sağlama, bilgi hırsızlığı, siyasal ve ideolojik amaçlar olmak gibi geniş alanlar hedeflenerek uygulanır.

Siber saldırılar dağıtık hizmet engelleme (DDoS), içeriden tehdit, parola saldırıları, SQL enjeksiyonu, kötü amaçlı yazılım gibi birçok tip içerisinde kategorize edilebilir. Bu başlıkta çalışmanın makine öğrenmesi yöntemleri uygulanarak yaratmaya çalıştığı siber saldırı sistemi için kullanılan KDD99 veri setinde bulunan hizmet engelleme saldırıları (DoS), kullanıcıdan kök kullanıcıya saldırılar (U2R), uzaktan yerel kullanıcıya saldırılar (R2L) ve bilgi tarama saldırıları detaylandırılacaktır.

3.1.1.1 Hizmet Engelleme Saldırıları (DoS)

Hizmet engelleme saldırıları, sistem, ağ veya kuruluşu hedef alarak kullanıcıların kullanımını engellemek amacıyla yapılan saldırılardır. Bu saldırılar devlet kurumları, yüksek sermayeli kuruluşlar gibi büyük hedeflere saldırdığı gibi küçük organizasyonlara ya da bireylere de saldırı düzenleyebilir. Genel çalışma biçimi olarak

hizmet engelleme saldırıları; yüksek miktarda veri göndererek ya da çalışma sisteminin kaynaklarını yoğun kullanarak uygulanır. Tek kaynaktan saldırı şeklinde işlenen hizmet engelleme saldırıları daha sonrasında geliştirilerek çoklu kaynaklardan senkronize şekilde gerçekleştirilen ve tespiti zorlaştıran, saldırıyı derinleştiren dağıtık hizmet engelleme saldırıları (DDoS) ek biçimini de almıştır (Mirkovic ve Reiher, 2004).

Bu saldırıların temel amacı hizmet kalitesinin azaltılması ya da hizmetin tamamen durdurulmasıdır. Bu hedefin sonucu olarak saldırı gerçekleştirilen kurum ya da kurumlar veri kayıpları, kullanıcı memnuniyetsizliği, prestij kaybı ve maddi kayıplar gibi sonuçlarla karşı karşıya kalabilir. Bu saldırı tipinin tespit edilmesi ya da önlenmesi açısından güvenlik duvarlarının oluşturulması, saldırı tespit ve önleme sistemlerinin (IDS/IPS) kurulması ve devamlı olarak ağ gözlemleme yapılması gibi yöntemler mevcuttur. DoS ve DDoS saldırılarına karşı korunmak için etkili bir siber güvenlik stratejisi uygulamak büyük bir öneme sahiptir (Mirkovic vd., 2002).

3.1.1.2 Kullanıcıdan Kök Kullanıcıya Saldırıları (U2R)

Saldırganın bir kullanıcının hesabına erişerek yönetici düzeyine yükselmeyi amaçlayarak hareket etmesi, kullanıcıdan kök kullanıcıya saldırılardır. Bu saldırı türünde saldırı yönetici pozisyonuna erişmesi durumunda, sistemi kontrol edebilir pozisyona gelir ve bu durum büyük güvenlik problemleri ortaya çıkarabilir (Hoglund ve McGraw, 2004).

Sistemin yazılım eksikliklerini kullanarak saldırı gerçekleştirilebileceği gibi sosyal mühendislik becerileri yani kullanıcıları aldatarak da yetkilere erişebilir. Bunun yanı sıra güçsüz şifrelemeler ve güncel olmayan yazılım programları nedeni ile de bu saldırı türüne alan açılmış olunur. Güvenlik alanındaki bilgilerin eksik olduğu kurum ya da kişilerde, yönetici düzeyi yetkilerin ele geçirilmesi daha kolaydır (Erickson, 2008).

Siber güvenlik alanında önemli bir saldırı türü olan kullanıcıdan kök kullanıcıya saldırılar, ciddi veri kayıplarına ve mali kayıplara neden olabilir. Güvenlik duvarlarının güçlendirilmesi ve açıkların izlenmesi bu saldırıların karşısında tespit ve savunma mekanizması oluşturur (Hoglund ve McGraw, 2004).

3.1.1.3 Uzaktan Yerel Kullanıcıya Saldırıları (R2L)

Saldırganın yönetici düzeyine erişim sağlayarak kullanıcının haklarının, bilgilerinin ele geçirilmesi biçimindeki saldırılar, uzaktan yerel kullanıcıya saldırı tipleri olarak adlandırılır. Saldırgan ya da saldırganların yönetici düzeyine erişim için en çok kullandıkları biçim, parola saldırılarıdır. Sosyal mühendislik uygulamaları da gizli bilgilerin erişilmesi ve bu yolla yönetici düzeyine erişilip kullanıcı kademesine zarar vermek açısından önemlidir. Kullanıcı bilincinin yüksek olmadığı ortamlarda daha sık görülür (Gifty Jeya vd., 2012). Buna ek olarak zararlı yazılımlar da bu saldırı tipinde kullanılan yöntemlerdendir.

Uzaktan yerel kullanıcıya saldırılar, sistem güvenliğini tehlikeye atan, kişisel bilgilerin kullanılarak tehdit, şantaj gibi unsurların kullanılabilceği saldırılardır. Saldırgan ya da saldırganlar kullanıcının haklarına, bilgilerine eriştikleri zamanlarda manipölasyonlarla sistemin bir aradalığını ve gizliliğini tehlikeye düşürebilir (Gifty vd., 2012).

Bu saldırı tipinin engellenmesi de diğer saldırı tiplerine benzer önlemleri içermektedir. Çok faktörlü kimlik doğrulama (MFA) yönteminin kullanılması, kullanıcıların bu konuda ve siber güvenlik alanı anlamında eğitimli olması ve güçlü parola yaratımı bu saldırı tipi için önemlidir. Bu basamakları kullanan ve uygulayan sistem ya da kuruluşlar, yerel kullanıcıya saldırıların etkisini asgari hale getirerek sistem güvenliğinin yaratılmasını mümkün kılar.

3.1.1.4 Bilgi Tarama Saldırıları

Bilgi tarama saldırıları, diğer saldırı tiplerinden farklı olarak şantaj, tehdit ya da maddi zarar gibi tehlikeler yaratmaz. Bu saldırı tipinin asıl hedefi bilgisayar sistemleri, kuruluşlar ya da bireyler ile ilgili bilgi toplamaktır. Bilgi toplama biçimi, genel itibari ile büyük ve önlenmesi güç saldırıların ön hazırlığı olarak nitelendirilebilir.

Bu saldırı tipinin en çok kullandığı yöntemlerden biri port taramalarıdır. Bu yöntemde ağdaki açık olan portlar belirlenmeye ve açık portların içerikleri öğrenilmeye çalışılır. Ek olarak, ağ haritalama yöntemi de bilgi tarama saldırılarının yaygın kullandığı yöntemlerden biridir. Bu yöntem ile ağların birbirleri ile ilişkisini ve

ağdaki cihazların bilgisini ele geçirir. Dolayısıyla ağdaki özelliklerin, hizmetlerin öğrenilmesi ile birlikte saldırıların hedef belirlenmesi kolaylaşır.

Bilgi tarama saldırılarının önlenmesi için güvenlik duvarlarının kurulması, ağ segmentasyonunun sağlanması, ağ trafiğinin sürekli izlenip kayıt altına alınması ve yazılımların güncel tutulması gibi bir çok yöntem mevcuttur. Siber güvenlik önlemlerinin sıklaştırılması ve aktif uygulanması sadece bu saldırı tipi için değil tüm saldırı tiplerinde belirgin düşüşler sağlayacaktır.



BÖLÜM DÖRT

VERİ VE UYGULAMA

Bu bölümde, çalışmada kullanılan veri seti, veri seti üzerinde yapılan işlemler ve kurulan modeller sonuçlarıyla birlikte anlatılacaktır.

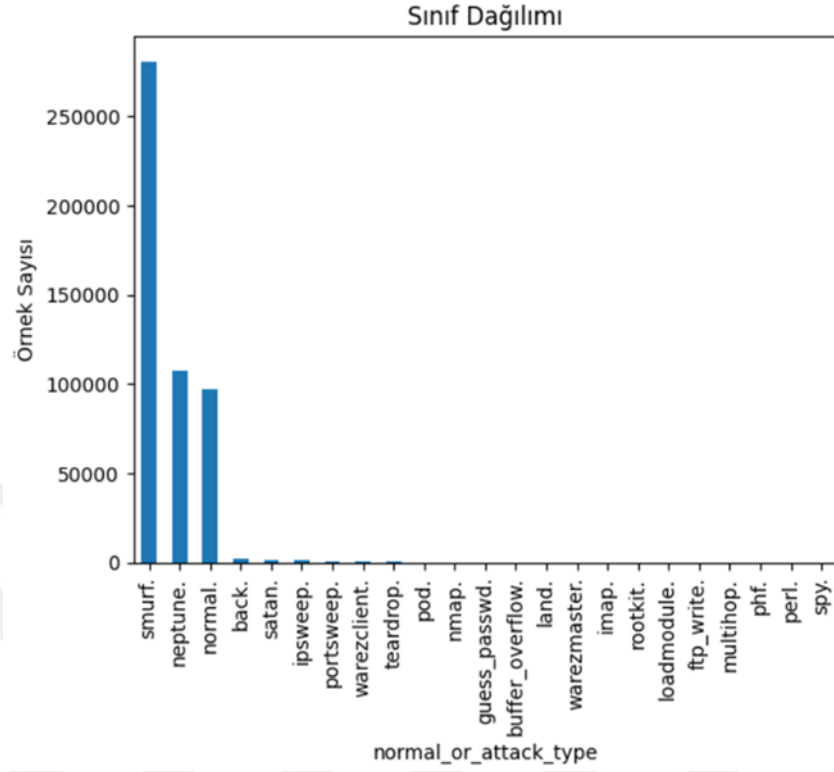
4.1 Veri Seti

KDD99 veri seti, literatür taramasının da gösterdiği üzere siber saldırı tespit sistemlerinin kurulmasında ve araştırılmasında sıklıkla kullanılan bir veri seti olma özelliğini taşır. KDD99, 1998 DARPA IDS değerlendirme programından oluşturulmuş bir veri setidir. Bu veri seti, siber saldırıları tespit sistemlerinde kullanılan modellerin başarılarının değerlendirilmesi amacıyla ortak bir referans noktası sunmaktadır (Tavallae vd., 2009).

Eğitim ve test veri seti olarak ayrılmış olan veri seti, 42 sütundan oluşmaktadır. “Label” olarak adlandırılmış olan sütun hedef sınıfıdır (target class). Hedef sınıf, normal veya saldırı adı olarak etiketlenmiştir. Bu veri seti yaklaşık 5 milyon kayıttan oluşmaktadır ve etiketlenmiş olan saldırı isimleri dört temel saldırı tipini kapsamaktadır. Bu saldırı tipleri Bölüm Üç’te açıklandığı üzere; Hizmet Engelleme Saldırıları (DoS), Kullanıcıdan Kök Kullanıcıya Saldırıları (U2R), Uzaktan Yerel Kullanıcıya Saldırıları (R2L) ve Bilgi Tarama Saldırılarıdır (Probing).

Tüm bu özelliklerin dışında veri setinde, eğitim ve test setlerinde yineleyen kayıt sayısının çok fazla olması gibi bazı eksiklikler bulunmaktadır. Veri setinin kullanılarak saldırı tespit sistemi oluştururken kurulan modellerin sonuç başarılarının literatürde de yüksek olmasına bu durumun neden olduğu tahmin edilmektedir. Bölüm 4.2’de uygulanan ön işleme yöntemleri anlatımında yinelenen kayıt sayısına dair bilgiler verilecektir.

Ek olarak; KDD99 veri setinin bu sorunları göz önüne alınarak, NSL-KDD adlı yeni bir veri seti önerilmiştir, bu veri seti KDD99’un temel sorunlarını gidermek amacıyla geliştirilmiştir (Tavallae vd., 2009). Veride bulunan sınıflara ait dağılım Şekil 4.1’de; saldırı tipleri Tablo 4.1’de gösterilmiş ve veri setindeki etiketler Tablo 4.2’de açıklanmıştır.



Şekil 4.1 Sınıf dağılımı

Tablo 4.1 Veri setine ait etiketler

SİBER SALDIRI TİPLERİ	KDD99 VERİ SETİNDE ETİKET ADI
Hizmet Engelleme Saldırıları	back, land, neptune, pod, smurf, teardrop
Kullanıcıdan Kök Kullanıcıya Saldırıları	ftp_write, guess_passwd, imap, multihop, phf, spy, warezclient, warezmaster
Uzaktan Yerel Kullanıcıya Saldırıları	buffer_overflow, loadmodule, perl, rootkit
Bilgi Tarama Saldırıları	ipsweep, nmap, portsweep, satan

Özellik tanımları aşağıda yer almaktadır.

- 1 Duration : Bağlantının süre uzunluğu
- 2 Protocol_type : Bağlantıda kullanılan protokol
- 3 Service : Varış noktası ağ servisi
- 4 Flag : Bağlantının durumu
- 5 Src_bytes : Tek bir bağlantıda kaynak ile hedef arasında aktarılan veri baytlarının sayısı
- 6 Dst_bytes : Tek bir bağlantıda hedeften kaynağa aktarılan veri baytlarının sayısı
- 7 Land : Kaynak ve hedef IP adresleri ile port numaraları eşitse, bu değişken 1 değerini alır; aksi takdirde 0 değerini alır.
- 8 Wrong_fragment : Bu bağlantıdaki toplam yanlış parça (fragment) sayısı
- 9 Urgent : Bu bağlantıdaki acil paketlerin sayısı
- 10 Hot : İçerikteki 'sıcak' göstergelerin sayısı, örneğin: bir sisteme giriş
- 11 Num_failed_logins : Başarısız giriş denemelerinin sayısı
- 12 Logged_in : Giriş Durumu (0,1 şeklinde)
- 13 Num_compromised : “Tehlikeye atılmış” durumların sayısı
- 14 Root_shell : Kök kabuğu elde edilirse 1; diğer türlü 0
- 15 Su_attempted : “su root” komutu denenirse veya kullanılırsa 1; diğer türlü 0
- 16 Num_root : Bağlantıda kök olarak gerçekleştirilen ``root" erişimlerinin veya işlem sayısının sayısı
- 17 Num_file_creations : Bağlantıdaki dosya oluşturma işlemlerinin sayısı
- 18 Num_shells : Kabuk istemlerinin sayısı
- 19 Num_access_files : Erişim kontrol dosyalarındaki işlem sayısı
- 20 Num_outbound_cmds : Bir ftp oturumunda giden komutların sayısı
- 21 Is_hot_login : Oturum açma “hot” listesine aitse, yani kök veya yönetici ise 1; diğer türlü 0
- 22 Is_guest_login : Oturum açma “guest” oturum açma ise 1; diğer türlü 0
- 23 Count : Geçtiğimiz iki saniyede geçerli bağlantıyla aynı hedef ana bilgisayara yapılan bağlantı sayısı
- 24 Srv_count : Geçtiğimiz iki saniyede geçerli bağlantıyla aynı hizmete (port numarası) yapılan bağlantı sayısı

- 25 Serror_rate : Count(23)'ta toplanan bağlantılar arasında (4) s0, s1, s2 veya s3 bayrağını etkinleştiren bağlantıların yüzdesi
- 26 Srv_serror_rate : Srv_count(24) toplanan bağlantılar arasında (4) s0, s1, s2 veya s3 bayrağını etkinleştiren bağlantıların yüzdesi
- 27 Rerror_rate : Count (23)'ta toplanan bağlantılar arasında (4) REJ bayrağını etkinleştiren bağlantıların yüzdesi
- 28 Srv_rerror_rate : Srv_count (24) toplanan bağlantılar arasında (4) REJ bayrağını etkinleştiren bağlantıların yüzdesi
- 29 Same_srv_rate : Toplanan bağlantılar arasında aynı hizmete giden bağlantıların yüzdesi
- 30 Diff_srv_rate : Toplanan bağlantılar arasında farklı hizmetlere giden bağlantıların yüzdesi
- 31 Srv_diff_host_rate : Toplanan bağlantılar arasında farklı hedef makinelere giden bağlantıların yüzdesi (srv_count)
- 32 Dst_host_count : Aynı hedef ana bilgisayar IP adresine sahip bağlantıların sayısı
- 33 Dst_host_srv_count : Aynı bağlantı noktası numarasına sahip bağlantıların sayısı
- 34 Dst_host_same_srv_rate : Toplanan bağlantılar arasında aynı hizmete giden bağlantıların yüzdesi (dst_host_count)
- 35 Dst_host_diff_srv_rate : Toplanan bağlantılar arasında farklı hizmetlere giden bağlantıların yüzdesi (dst_host_count)
- 36 Dst_host_same_src_port_rate : dst_host_srv_count'da toplanan bağlantılar arasında aynı kaynak bağlantı noktasına giden bağlantıların yüzdesi
- 37 Dst_host_srv_diff_host_rate : dst_host_srv_count'da toplanan bağlantılar arasında farklı hedef makinelere giden bağlantıların yüzdesi
- 38 Dst_host_serror_rate : dst_host_count'da toplanan bağlantılar arasında s0, s1, s2 veya s3 bayrağını etkinleştiren bağlantıların yüzdesi
- 39 Dst_host_srv_serror_rate : dst_host_srv_count'da toplanan bağlantılar arasında s0, s1, s2 veya s3 bayrağını etkinleştiren bağlantıların yüzdesi
- 40 Dst_host_rerror_rate : dst_host_srv_count'da toplanan bağlantılar arasında REJ bayrağını etkinleştiren bağlantıların yüzdesi

41 Dst_host_srv_error_rate : bağlantılar arasında REJ bayrağını etkinleştiren dst_host_count'ta toplanmış bağlantıların yüzdesi

42 Label : Hedef sınıf. Normal veya saldırı tipi olarak dönüş yapar.

4.2 Veri Setinde Uygulanan Ön İşleme Adımları

Çalışmada kullanılan veri seti KDD Cup 1999 veri tabanından elde edilmiştir (UCI KDD Archive, 2024). Veri seti zip dosyası şeklinde indirilmiş olup txt dosyası formatına çevirilmiştir. Text dosyası olarak jupyter ortamına yüklenmiş ve dosyanın okunması sağlanmıştır.

Yukarıda da bahsedildiği üzere veri seti eğitim ve test olmak üzere iki dosyaya ayrılmış şekilde bulunmaktadır. Test veri seti, ana veri setinin %10'luk kısmını içerecek şekilde oluşturulmuş biçimdedir. Test veri seti 494021 satır, 42 sütundan oluşmaktadır. Modeller, test veri seti üzerinden kurulmuştur. Veri seti üzerinde ön işleme adımları izlenmiştir.

Veri setinin ilk indirilmiş hali, özellik(feature) bilgileri eklenmiş bir düzende olmadığı için özellikler isimlendirilmiştir. Hedef sınıf olan yani normal veya saldırı adı şeklinde etiketlenmiş olan sınıf/özellik, bu çalışmada “label” olarak adlandırılmıştır.

Ardından hedef sınıfın etiket dağılımı öğrenilmek için Şekil 4.1’de gösterilen etiketlerin dağılımına bakılmıştır.

Veri setinin eksik, aykırı, “NA” gibi değerleri içerip içermediği kontrol edilmiştir. Bu kontrol sonrası veri setinin eksik, aykırı ya da “NA” gibi değerler içermediği gözlemlenmiştir.

Veri setinin başarı oranını etkileyen faktörlerden biri olan, veri setinde yinelenen değerlerin bulunması durumu da Jupyter ortamında kontrol edilmiştir. Veri setinin yineleyen satırlarının oldukça fazla olduğu ve 494021 gözlem değerine sahip olan verinin, yineleyen değer sayısının 348435 olduğu görülmüştür.

Makine öğrenmesi modelleri kurulurken hedef sınıfın yani veriye dair etiketlerin iki kategorili etiketlere sahip olması kolaylaştırıcı bir noktadır. Bu nokta modellerin

alıřma hızını da etkiler. Bu nedenle “label” stnuna ait “normal” ve “saldırđ ismi” řeklindeki etiketler, “normal” ve “saldırđ” olarak deęiřtirilmiřtir. Bu deęiřiklik sonrası ortaya ıkan tabloda, “0” yani “normal” etiketli deęer sayısı 97278, “1” yani “saldırđ” etiketli deęer sayısı da 396743’tr.

Makine ęrenmesi yntemlerinin bazıları sayılar verilerle alıřma prensibine sahiptir. Karar aęaları ve ADABOOST yntemleri bu alıřma prensibine sahip makine ęrenmesi yntemlerine rnektir. Bu nedenle kategorik verileri ikili formata dnřtrmek iin tek-semeli kodlama(one-hot encoding) kullanılır. Bu alıřmada da kategorik formatta olan "protocol_type", "service", "flag" stnları ikili formata dnřtrlmřtir. Bu stnlar bu adım sonrası “0” ve “1” deęerleri almıřtır.

4.3 Modelleme

Sınıflandırma modelleri kurulurken veriden hedef sınıf(bu veri setinde “label” stnu) ıkartılır. Bu alıřmada kullanılan veri setinde de hedef deęiřken olan label stnu ıkartılmıř ve geri kalan deęiřkenler zerinden modeller kurulmuřtur. Bu baęlamda bu makineye “label” stnunun hedef sınıf olduęu anlatılmıř ve “label” stnu ıkartılan veri yeni bir veri erevesine aktarılmıřtır.

Bu iřlemin ardından bir dięer nemli basamak olan veri setinin test ve eęitim veri setlerine ayrılması iřlemine geilmiřtir. Kullanılan veri seti, %20 test, %80 eęitim olacak řekilde ikiye ayrılmıřtır. Bu noktadan sonra veri, model kurulumuna hazır hale gelmiřtir.

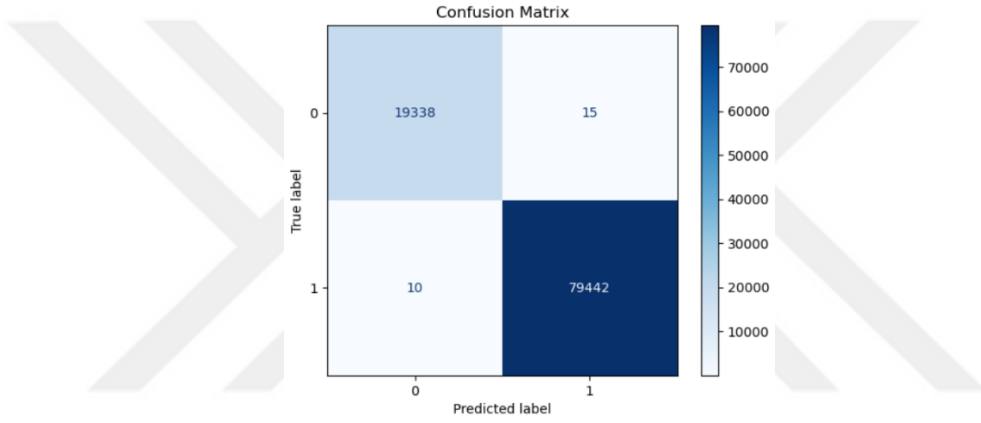
4.3.1 Deneyisel alıřmalar

Bu blmde karar aęaları, ADABOOST, yapay sinir aęları, destek vektr makineleri yntemlerinin farklı sreler ve iřlemler baęlamında modellenmeleri anlatılacak ve tabloladıtırılacaktır. Ayrıca, bu blmnde KDD99 veri setinin kullanıldıęđ farklı deneyisel alıřmaları ieren saldırđ tespit sistemleri alıřmalarının performansları deęerlendirilecektir.

4.3.1.1 Özellik Seçimi Yapılmadan

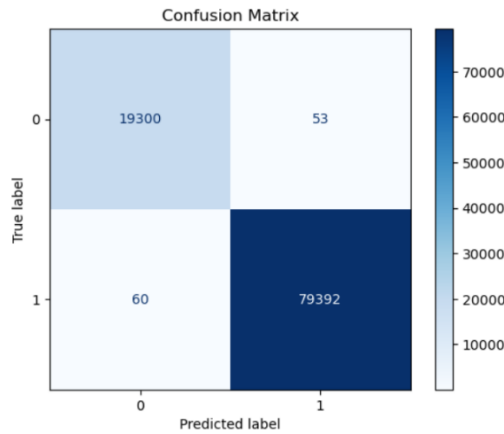
Bu deneyde, bölüm 4.2 ve 4.3'te detaylandırılan işlemler sonrası kurulan dört modelin performans sonuçları hata matrisleri ile birlikte gösterilecektir. Bu deney için en başarılı sonuç veren makine öğrenmesi yöntemi karar ağaçları olmuştur. Değerlendirme ölçütleri olan doğruluk, kesinlik, duyarlılık, F1-skoru sonuçları ile beraber aşağıda tablolastırılmıştır.

Karar ağaçları modeline ait hata matrisi Şekil 4.2'de gösterilmiştir.



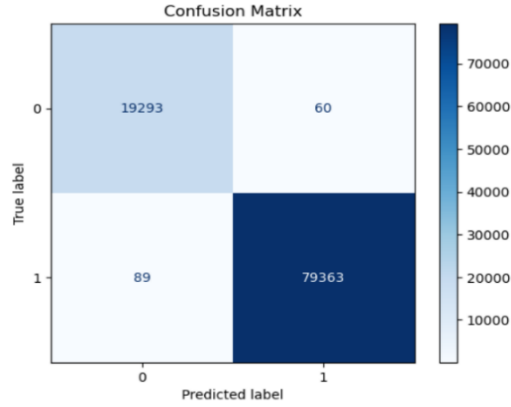
Şekil 4.2 Karar ağaçları hata matrisi

ADABOOST modeline dair hata matrisi Şekil 4.3'de gösterilmiştir.



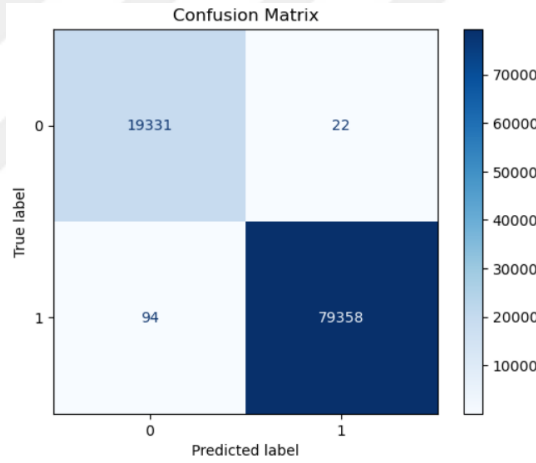
Şekil 4.3 ADABOOST hata matrisi

Destek vektör makineleri (DVM) modeline dair hata matrisi Şekil 4.4’de gösterilmiştir.



Şekil 4.4 DVM hata matrisi

Yapay sinir ağları (YSA) modeline dair hata matrisi Şekil 4.5’de gösterilmiştir.



Şekil 4.5 YSA hata matrisi

Dört modelin de performans sonuçları Tablo 4.3’te gösterilmiştir.

Tablo 4.2 Değerlendirme ölçütleri

	Karar Ağaçları	Destek Vektör Makineleri	Yapay Sinir Ağları	ADABOOST
ACCURACY	0.9993	0.9899	0.9960	0.9933
F1-SCORE	1.00	1.00	1.00	1.00
PRECISION	1.00	1.00	1.00	1.00
RECALL	1.00	1.00	1.00	1.00

Tüm modellerin performanslarına bakıldığında başarı oranlarının yaklaşık %99 ve üzeri olduğu gözlemlenmiştir. Hata oranları da hata matrisi üzerinden görselleştirilerek kolay okunabilirlik sağlanmıştır. Başarı oranlarının bu denli yüksek oluşu, modellerin performansının neredeyse %100'e yakın sonuçlar vermesi aşırı uyumlanma (overfitting) ihtimalini düşündürmüştür. Bu nedenle diğer deneysel çalışmalar bu ihtimali güçlendirme ya da zayıflatma hedefi ile yapılmıştır.

4.3.1.2 Farklı Örnekler Oluşturma

Bu çalışmada, daha önce de ifade edildiği gibi, KDD99 veri setinin tamamı yerine, yalnızca %10'luk bir bölümü seçilerek analizler gerçekleştirilmiştir. Ancak, modellerin performanslarının bu kadar yüksek olması KDD99 test veri setinin yanılabileceğini ortaya çıkartmıştır. Bu nedenle de bu bölümde bu ihtimale dair yapılan çalışmalar anlatılacaktır.

KDD99 veri setinin ana hali alınarak otuz farklı %10'luk veri seti oluşturulmuştur. Tüm modeller bu otuz örnekte tek tek denenmiştir.

Karar ağaçları modeline dair bazı örneklerin sonuçları Tablo 4.4'te gösterilmiştir.

Tablo 4.3 Karar ağaçları modeli sonuçları

Karar Ağaçları		SAMPLE1	SAMPLE2	SAMPLE3	SAMPLE4	SAMPLE5
	ACCURACY	0.999878	0.999735	0.999449	0.999816	0.999816
0	PRECISION	1.00	1.00	1.00	1.00	1.00
1	PRECISION	1.00	1.00	1.00	1.00	1.00
0	RECALL	1.00	1.00	1.00	1.00	1.00
1	RECALL	1.00	1.00	1.00	1.00	1.00
0	F1-SCORE	1.00	1.00	1.00	1.00	1.00
1	F1-SCORE	1.00	1.00	1.00	1.00	1.00

ADABOOST modeline dair bazı örneklerin sonuçları Tablo 4.5'de gösterilmiştir.

Tablo 4.4 ADABOOST modeli sonuçları

ADABOOST		SAMPLE1	SAMPLE2	SAMPLE3	SAMPLE4	SAMPLE5
	ACCURACY	0.999408	0.999255	0.999449	0.999255	0.999449
0	PRECISION	1.00	1.00	1.00	1.00	1.00
1	PRECISION	1.00	1.00	1.00	1.00	1.00
0	RECALL	1.00	1.00	1.00	1.00	1.00
1	RECALL	1.00	1.00	1.00	1.00	1.00
0	F1-SCORE	1.00	1.00	1.00	1.00	1.00
1	F1-SCORE	1.00	1.00	1.00	1.00	1.00

Destek vektör makineleri (DVM) modeline dair bazı örneklerin sonuçları Tablo 4.6'da gösterilmiştir.

Tablo 4.5 DVM modeli sonuçları

DVM		SAMPLE1	SAMPLE2	SAMPLE3	SAMPLE4	SAMPLE5
	ACCURACY	0.999092	0.999020	0.999030	0.998938	0.999224
0	PRECISION	1.00	1.00	1.00	1.00	1.00
1	PRECISION	1.00	1.00	1.00	1.00	1.00
0	RECALL	1.00	1.00	1.00	1.00	1.00
1	RECALL	1.00	1.00	1.00	1.00	1.00
0	F1-SCORE	1.00	1.00	1.00	1.00	1.00
1	F1-SCORE	1.00	1.00	1.00	1.00	1.00

Yapay sinir ağları (YSA) modeline dair bazı örneklerin sonuçları Tablo 4.7'de gösterilmiştir.

Tablo 4.6 YSA modeli sonuçları

YSA		SAMPLE1	SAMPLE2	SAMPLE3	SAMPLE4	SAMPLE5
	ACCURACY	0.9991	0.9992	0.9963	0.9989	0.9968
0	PRECISION	0.9991	0.9992	0.9963	0.9989	0.9968
0	RECALL	0.9991	0.9992	0.9963	0.9989	0.9968
0	F1-SCORE	0.9991	0.9992	0.9963	0.9989	0.9968

Bu deney sonrası sonuçların da gösterdiği üzere KDD99 test veri setinin sadece kendisinin değil ana veri setinin %10'luk örneklerinin de başarı oranlarının %99 ve üzeri sonuçlar verdiği gösterilmiştir. Başarı oranlarının farklı %10'luk örneklerde de bu denli yüksek olması 4.3.1.1 bölümündeki deneyde iddia edilen aşırı uyumlama (overfitting) iddiasını da destekler niteliktedir.

4.3.1.3 Yinelenen Değerlerin Silinmesi

Yinelenen değerlerin fazlaca olması veride modeller kurarken modelin performansında aşırı düşüşlere ya da yükselişlere neden olabilir. Bu durumun kendisi de model sonuçlarının yanlı olması ile sonuçlanabilir. Yinelenen değerlere dair 4.2 başlığında da bahsedildiği üzere 494021 değere sahip olan verinin, yineleyen değer sayısının 348435 olduğu görülmüştür. Yineleyen değerlerin bu denli yüksek oluşu da veriden yinelenen değerlerin çıkartılarak ve yine ham KDD99 verisinden %10'luk otuz örnek oluşturularak modellerin kurulması fikrini ortaya çıkartmıştır. Bu bağlamda oluşturulan yeni örnek veri setleri ile karar ağaçları, ADABOOST, yapay sinir ağları ve destek vektör makineleri makine öğrenmesi yöntemleri kurulmuştur. Her bir yöntem için bazı örneklerin sonuçları Tablo 4.8, 4.9 4.10 ve 4.11'de gösterilmiştir.

Tablo 4.7 Karar ağaçları modeli sonuçları

Karar Ağaçları		SAMPLE1	SAMPLE2	SAMPLE3	SAMPLE4	SAMPLE5
	ACCURACY	0.999209	0.999302	0.999070	0.999628	0.999442
0	PRECISION	1.00	1.00	1.00	1.00	1.00
1	PRECISION	1.00	1.00	1.00	1.00	1.00
0	RECALL	1.00	1.00	1.00	1.00	1.00
1	RECALL	1.00	1.00	1.00	1.00	1.00
0	F1-SCORE	1.00	1.00	1.00	1.00	1.00
1	F1-SCORE	1.00	1.00	1.00	1.00	1.00

Tablo 4.8 ADABOOST modeli sonuçları

ADABOOST		SAMPLE1	SAMPLE2	SAMPLE3	SAMPLE4	SAMPLE5
	ACCURACY	0.997768	0.998419	0.998186	0.997861	0.998372
0	PRECISION	1.00	1.00	1.00	1.00	1.00
1	PRECISION	1.00	1.00	1.00	1.00	1.00
0	RECALL	1.00	1.00	1.00	1.00	1.00
1	RECALL	1.00	1.00	1.00	1.00	1.00
0	F1-SCORE	1.00	1.00	1.00	1.00	1.00
1	F1-SCORE	1.00	1.00	1.00	1.00	1.00

Tablo 4.9 DVM modeli sonuçları

DVM		SAMPLE1	SAMPLE2	SAMPLE3	SAMPLE4	SAMPLE5
	ACCURACY	0.996512	0.992791	0.996419	0.996	0.996140
0	PRECISION	1.00	1.00	1.00	1.00	1.00
1	PRECISION	1.00	1.00	1.00	1.00	1.00
0	RECALL	1.00	1.00	1.00	1.00	1.00
1	RECALL	1.00	1.00	1.00	1.00	1.00
0	F1-SCORE	1.00	1.00	1.00	1.00	1.00
1	F1-SCORE	1.00	1.00	1.00	1.00	1.00
0	SUPPORT	16281	16223	16263	16243	16298
1	SUPPORT	5219	5277	5237	5257	5202

Tablo 4.10 YSA modeli sonuçları

YSA		SAMPLE1	SAMPLE2	SAMPLE3	SAMPLE4	SAMPLE5
	ACCURACY	0.9978	0.9977	0.9973	0.9971	0.9968
	PRECISION	0.9978	0.9977	0.9973	0.9971	0.9968
	RECALL	0.9978	0.9977	0.9973	0.9971	0.9968
	F1-SCORE	0.9978	0.9977	0.9973	0.9971	0.9968

Bu deney sonucunda da her bir makine öğrenmesi yönteminin performansı %99 ve üzeri sonuç vermiştir. Sonuçların yüksekliği aşırı uyumlama (overfitting) ihtimalini güçlendirmektedir.

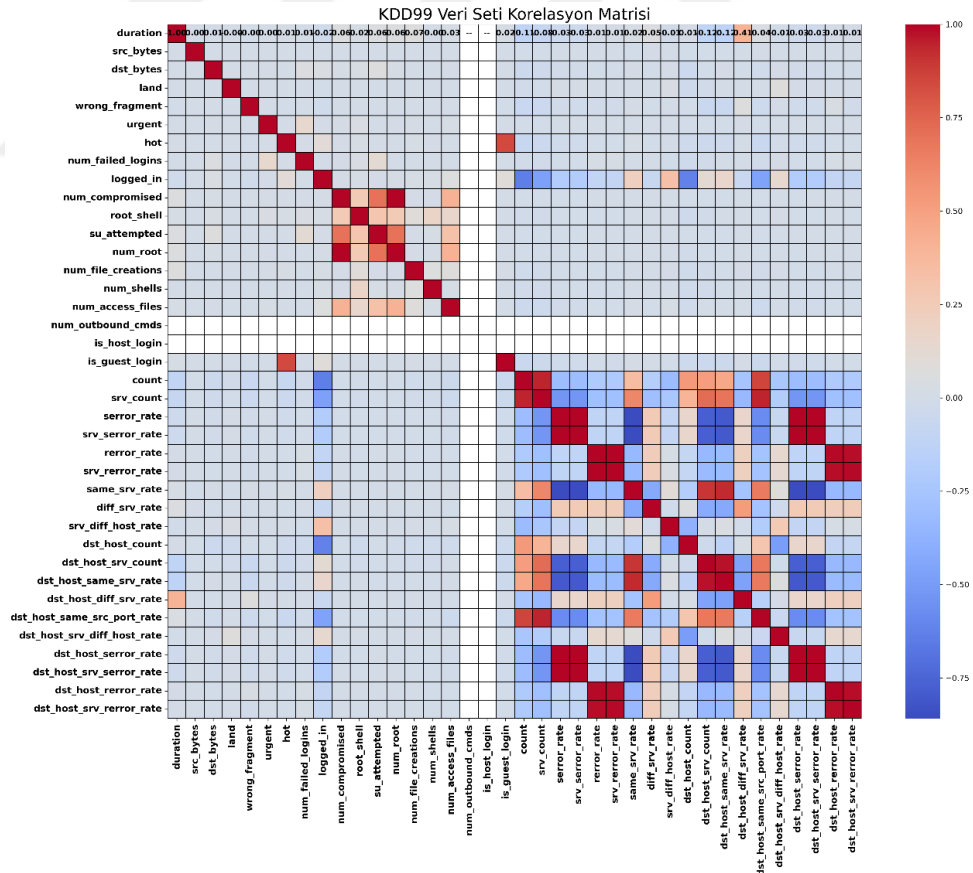
4.3.1.4 Korelasyon Matrisi

Sosyal bilimlerden pazarlama stratejilerine kadar yaygın bir kullanımı olan korelasyon matrisi, çalışmada kullanılan veri setinin içerdiği değişkenlerin birbirleri arasındaki ilişkiyi ortaya koyan bir tablodur. Tablo, her iki değişkenin ilişkisini gösteren bir hücreye sahiptir. Bu hücre korelasyon katsayısını içerir. Korelasyon katsayısı da -1 ile +1 arasında değerler içerir. Korelasyon katsayısının +1 veya +1'e yakın olması bahsi geçen iki değişkenin arasındaki ilişkinin mükemmel veya mükemmele yakın olduğunu gösterirken; -1 veya -1'e yakın olması iki değişkenin ilişkisinin negatif anlamda mükemmel veya mükemmele yakın olduğunu gösterir. Korelasyon katsayısının 0 olduğu durum da iki değişken arası herhangi bir ilişkinin bulunmadığını gösterir.

Korelasyon matrisi yönteminin bu çalışmada kullanılma amacı, farklı deneylerde kurulan makine öğrenmesi yöntemlerinin yüksek başarı sonuçlarına ulaşmasının nedenini göstermeye çalışırken başka bir yaklaşım kullanmaktır. Makine öğrenmesi

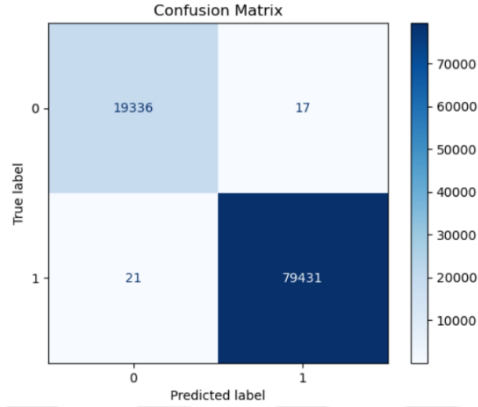
yöntemleri kurulurken de kullanılan korelasyon matrisi yöntemi, hedef değişken ile veri setinde geride kalan değişkenlerin arasındaki ilişkiyi ortaya koymaktır. Hedef değişken ile korelasyon katsayısı +1 ya da -1 olan değişkenlerin varlığı ve çalışmalarda kullanımı aşırı uyumlamaya(overfitting) ya da modellerin çalışma süresinin artmasına neden olabilmektedir. Bu gibi sorunların önüne geçebilmek için kullanılan bu yöntem, hedef değişkenin kategorik olduğu şartlarda kullanılır. Ancak bu çalışmada kullanılan KDD99 veri setinin hedef değişkeni olan “label” kategorik bir değişken değildir. Bu nedenle de korelasyon matrisi özellik seçimleri açısından bu veri setine uygulanamaz.

Öte yandan bu veri setinde korelasyon matrisi, birbirleri ile yüksek ilişkili olan değişkenleri tablolastırır. Bu yüksek ilişkili özelliklerin veri setinden çıkartılması yöntemi ile de gereksiz değişken kullanılarak oluşabilecek çalışma süresi, model performanslarının etkilenmesi gibi sorunlar ortadan kaldırılabilir. Bu amaçla kullanılacak olan KDD99 veri setinin korelasyon matrisi Şekil 4.6’da gösterilmiştir.



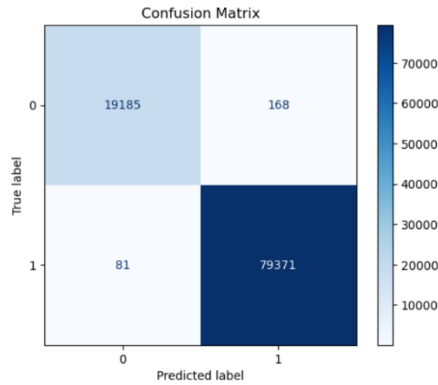
Şekil 4.6 Korelasyon matrisi

Korelasyon matrisi kullanılarak korelasyon katsayısı yüksek özelliklerin çıkartılması sonrası kurulan modellerin başarı oranlarında belirgin bir değişiklik gözlemlenmemiştir. Bu karşılaştırmayı yapabilmek adına karar ağaçları ve ADABOOST modellerinin hata matrisleri (confusion matrix) Şekil 4.7 ve 4.8’de gösterilmiştir.



Şekil 4.7 Karar ağaçları hata matrisi

Karar ağaçları modelinin başlangıçtaki başarı oranı 0.9993 iken bu deneyde kurulan modelin sonucu 0.9996’ye yükselmiştir.



Şekil 4.8 ADABOOST hata matrisi

ADABOOST modelinin başlangıçtaki başarı oranı 0.9933 iken bu deneyde kurulan modelin sonucu 0.9975’e yükselmiştir.

4.3.1.5 Makale Referansı ile Özellik Seçimi

Makine öğrenmesi yöntemleri uygulanırken veri setinin gereksiz ve fazla özellik içermesi durumu modelin performansını etkilemektedir. Model başarısını arttırma, model çalışma süresini düşürme amacıyla özellik azaltma yöntemleri denenir. Bu çalışma kapsamında da korelasyon matrisi çıktıları ile özellik azaltma ve çıkarılan özellikler sonrası model kurma yöntemi denenmiştir.

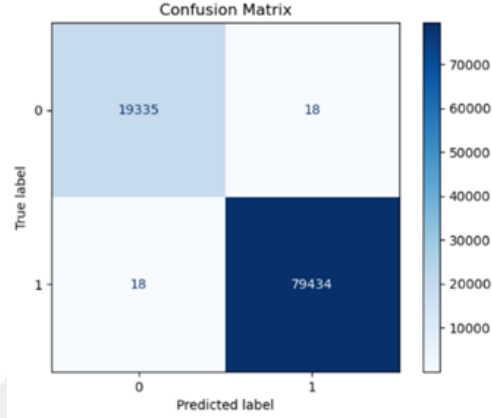
Bu deneyde de KDD99 veri setinde daha önce yapılan çalışmaların taranması sonrası ulaşılan Eldos, Siddiqui ve Kanan (2012) tarafından yapılan çalışma kullanılacaktır. Bahsi geçen çalışmada, KDD99 veri setinin içerdiği özellikler analiz edilmiş ve 41 olan özellik sayısının 13 tanesinin önemli olduğu iddia edilmektedir. Bu özellikler Tablo 4.12’de gösterilmiştir.

Tablo 4.11 Seçilen özellikler

Özelliklerin İsimleri
Protocol_type
Service
Flag
Src_bytes
Dst_bytes
Wrong_fragment
Logged_in
Num_compromised
Is_guest_login
Count
Srv_count
Dst_host_srv_count
Dst_host_same_src_port_rate

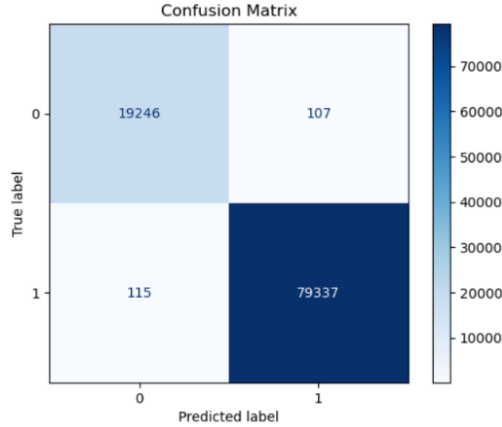
Aynı zamanda çalışma, 13 özellik harici diğer özelliklerin model kurmada etkisinin sıfıra yakın olduğunu da belirtir. Bu doğrultuda, bu deneyde karar ağaçları, ADABOOST, yapay sinir ağları ve karar destek makineleri yöntemleri 13 özellik kullanılarak kurulmuştur.

Siber saldırı tespit sistemleri açısından daha gelişkin ve gerçeğe yakın sonuçları bulabilmek ve bu alana dair çözümleri de ortaya koyabilmek için de yapılan bu deney sonuçlarında belirgin değişiklikler gözlemlenmemiştir. Sonuçları göstermek açısından karar ağaçları ve ADABOOST yöntemlerinin hata matrisleri (confusion matrix) kullanılmıştır.



Şekil 4.9 Karar ağaçları hata matrisi

Karar ağaçları modelinin başlangıçtaki başarı oranı 0.9993 iken bu deneyde kurulan modelin sonucu 0.9997'ye yükselmiştir.



Şekil 4.10 ADABOOST hata matrisi

ADABOOST modelinin başlangıçtaki başarı oranı 0.9933 iken bu deneyde kurulan modelin sonucu 0.9978'e yükselmiştir.

4.3.1.6 Lojistik Regresyon

Lojistik regresyon çoğunlukla bağımlı değişkenin iki sınıfa ayrıldığı durumlarda kullanılan bir yöntemdir. Yüksek boyutlu verilerde başarılı modelleme sağlayan bu yöntem, doğrusal regresyon modeline benzer biçimde çalışır. Bu çıktıya sigmoid fonksiyonunun uygulanması sonrası 0 ile 1 arası bir değer alır. Olasılık 0.5 değerinden büyükse model 1 olarak; diğer türlü 0 olarak sınıflandırılır. Sınıflandırma problemlerinde sıklıkla kullanılan bir model olan lojistik regresyon, bu çalışmada KDD99 veri seti kullanılarak özellik sayısı seçimi için kullanılmıştır.

Bu yöntem ile özelliklerin önem derecesi belirlenir. Özellik seçimi, her bir özellik için katsayı tahminine dayanır. Katsayılar, özelliğin hedef sınıfa olan etkisine bağlıdır. “SelectFromModel” yöntemi ile lojistik regresyon sonucu ortaya çıkan katsayılara dayandırılarak önemli özellikler seçilir. Bu sayede, belirli bir etki değerine ve üzerine sahip olan özellikler seçilmiş ve modeller bu özellikler kullanılarak kurulmuş olur.

Bu deneyin yapılma amacı özellik seçimi yöntemi ile modellerin hesaplanma sürelerini minimize etmek, modellerin yanlış sonuçlar verme ihtimalini düşürmek ve modellerin çalışma performanslarını arttırmaktır.

Bu modelin kurulması sonrası model başarısında belirgin bir değişiklik gözlemlenmemiştir. Ancak modelin çalışma süresinde düşüş gözlemlenmiştir. Bu düşüş Tablo 4.13’de gösterilmiştir.

Tablo 4.12 Çalışma süresi

Özellik Seçimi Öncesi Modelin Çalışma Süresi	Özellik Seçimi Sonrası Modelin Çalışma Süresi
16.15 Dakika	10.37 Dakika

BÖLÜM BEŞ

SONUÇ

Bu tez, makine öğrenmesi yöntemleri kullanılarak siber saldırı tespit sistemi geliştirmek üzere ilerletilmiştir. Siber saldırı alanında bilinen bir veri seti olan KDD99 veri setinin kullanıldığı çalışmada Karar Ağaçları, ADABOOST, Yapay Sinir Ağları (YSA) ve Destek Vektör Makineleri (DVM) gibi farklı modeller kurulmuş ve değerlendirilmiştir.

Kurulan modellerde en yüksek başarı oranına sahip olan modelin Karar Ağaçları yöntemi olduğu gözlemlenmiştir. Ancak her bir modelin başarı oranının %99 ve üzeri değerlere sahip olduğu ortaya çıkmıştır. Bu denli yüksek başarı oranları, kullanılan makine öğrenmesi modellerinin siber saldırı tespiti alanında başarılı olduğunu düşündürdüğü kadar aşırı öğrenme (overfitting) ihtimalini de beraberinde getirmiştir.

Aşırı uyumlanma (overfitting) ihtimali nedeniyle çalışmanın devamında altı farklı deney yöntemi uygulanmıştır. Her bir yöntemin veri setine ve modellere farklı bir bakış sağlaması hedeflenmiştir. Ancak, elde edilen modellerin başarı oranlarının yine %99 ve üzeri olduğu sonucuna ulaşılmıştır. Bu da veri setinin aşırı uyumlanma ihtimalinin tamamen engellenemediği ihtimalini arttırmıştır.

Sonuç olarak, bu çalışmada kurulan makine öğrenmesi modelleri, siber saldırı tespit sistemi açısından yüksek performans göstermiştir. Ancak sonuçlar değerlendirilirken aşırı uyumlanma riski göz önünde bulundurulmalıdır. Gelecekte yapılacak çalışmaların, sonuçların gerçek veriler üzerinde de benzer performans gösterip göstermeyeceğini doğrulamaya ve veri seti genişletilmesi, farklı özellik seçimi yöntemlerinin kullanılması gibi çalışmalarla aşırı uyumlanma riskinin azaltılmasına dair olması bu alandaki çalışmalara katkı sağlayacaktır.

KAYNAKLAR

- Ayhan, S., & Erdoğan, Ş. (2014). Destek Vektör Makineleriyle Sınıflandırma Problemlerinin Çözümü için Çekirdek Fonksiyonu Seçimi. *Eskişehir Osmangazi Üniversitesi İİBF Dergisi*, 9(1), 175-198.
- Boser, B. E., Guyon, I. M., & Vapnik, V. N. (1992). A training algorithm for optimal margin classifiers. In *Proceedings of the fifth annual workshop on Computational learning theory* (pp. 144-152).
- Cios, K., Pedrycz, W., Swiniarski, R., & Kurgan, L. A. (2007). *Data Mining: A Knowledge Discovery Approach*. Springer Science & Business Media.
- Cortes, C., & Vapnik, V. (1995). Support-vector networks. *Machine learning*, 20(3), 273-297.
- Eldos, T., Siddiqui, M. K., & Kanan, A. (2012). On the KDD'99 Dataset: Statistical Analysis for Feature Selection. *Journal of Data Mining and Knowledge Discovery*, 3(3), 88-90.
- Erickson, J. (2008). *Hacking: The Art of Exploitation*. No Starch Press.
- Gifti Jeya, P., Ravichandran, M., & Ravichandran, C. S. (2012). Efficient classifier for R2L and U2R attacks. *International Journal of Computer Applications*, 45(21), 28-32.
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press.
- Hoglund, G., & McGraw, G. (2004). *Exploiting Software: How to Break Code*. Addison-Wesley.
- Kaya, Ç., & Yıldız, O. (2014). Makine öğrenmesi teknikleriyle saldırı tespiti: Karşılaştırmalı analiz. *Marmara Fen Bilimleri Dergisi*, 3, 89-104. DOI: 10.7240/mufbed.24684.
- Mitchell, T. M. (1997). *Machine Learning*. McGraw-Hill.
- Mirkovic, J., Prier, G., & Reiher, P. (2002). Attacking DDoS at the source. *Proceedings 10th IEEE International Conference on Network Protocols*, 312-321. doi:10.1109/ICNP.2002.1181429
- Mirkovic, J., & Reiher, P. (2004). A taxonomy of DDoS attack and DDoS defense mechanisms. *ACM SIGCOMM Computer Communication Review*, 34(2), 39-53. doi:10.1145/997150.997156
- Nilsson, N. J. (1998). *Introduction to Machine Learning: An Early Draft of a Proposed Textbook*. Robotics Laboratory, Department of Computer Science, Stanford University. Retrieved from <https://ai.stanford.edu/~nilsson/MLBOOK.pdf>
- Öztemel, E. (2012). *Yapay Sinir Ağları*. Papatya Yayıncılık Eğitim.

- Roiger, R. J., & Geatz, M. W. (2003). *Data mining: A tutorial-based primer*. Pearson Education.
- Russell, S. J., & Norvig, P. (2021). *Artificial Intelligence: A Modern Approach* (4th ed.). Pearson.
- Sanayinin Dijital Dönüşümü. (2024, Temmuz 12). Yapay Zeka (Artificial Intelligence) Nedir? Sanayinin Dijital Dönüşümü. Retrieved from <https://www.sanayinindijitaldonusumu.com/yapay-zeka-artificial-intelligence-nedir/>.
- Soman, K. P., Loganathan, R., & Ajay, V. (2011). *Machine Learning with SVM and Other Kernel Methods*. PHI Learning Pvt. Ltd.
- Sutton, R. S., & Barto, A. G. (2018). *Reinforcement Learning: An Introduction* (2nd ed.). MIT Press.
- Tath, H. Ç. (2011). *Makine Öğrenmesinde Teoriden Örnek Matlab Uygulamalarına Kadar Destek Vektör Makineleri*.
- UCI KDD Arşivi. (2024). KDD Cup 1999 Verileri. <https://kdd.ics.uci.edu/databases/kddcup99/kddcup99.html>
- Vapnik, V. (1998). *Statistical learning theory*. Wiley.