



**T.C.
BATMAN ÜNİVERSİTESİ
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ
BİLGİ TEKNOLOJİLERİ ANA BİLİM DALI**

YÜKSEK LİSANS TEZİ

**SİBER GÜVENLİKTE XSS WEB SALDIRILARININ YAPAY ZEKÂ
ZEMİNİNDE ANALİZ EDİLMESİ**

Fürkan TANYERİ

**Ağustos-2024
BATMAN**

T.C.
BATMAN ÜNİVERSİTESİ
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ
BİLGİ TEKNOLOJİLERİ ANA BİLİM DALI

YÜKSEK LİSANS TEZİ

SİBER GÜVENLİKTE XSS WEB SALDIRILARININ YAPAY ZEKÂ
ZEMİNİNDE ANALİZ EDİLMESİ

Fürkan TANYERİ

Danışman
Prof. Dr. Ömer Faruk ERTUĞRUL

Ağustos-2024
BATMAN

TEZ KABUL VE ONAYI

Fürkan TANYERİ tarafından hazırlanan “Siber Güvenlikte XSS Web Saldırılarının Yapay Zekâ Zemininde Analiz Edilmesi” adlı tez çalışması 28/08/2024 tarihin de aşağıdaki jüri tarafından oy birliği ile Batman Üniversitesi Lisansüstü Eğitim Bilimleri Enstitüsü Bilgi Teknolojileri Anabilim Dalı’nda YÜKSEK LİSANS TEZİ olarak kabul edilmiştir.

Jüri Üyeleri

Başkan

Doç. Dr. Cafer BUDAK

Danışman

Prof. Dr. Ömer Faruk ERTUĞRUL

Üye

Dr. Öğr. Üyesi Abdulkerim ÖZTEKİN

İmza

.....

.....

.....

Yukarıdaki sonucu onaylarım.

Dr. Öğr. Üyesi Ömer Murat ÖTER
Lisansüstü Eğitim Enstitüsü Müdürü V.

TEZ BİLDİRİMİ

Bu tezdeki bütün bilgilerin etik davranış ve akademik kurallar çerçevesinde elde edildiğini ve tez yazım kurallarına uygun olarak hazırlanan bu çalışmada bana ait olmayan her türlü ifade ve bilginin kaynağına eksiksiz atıf yapıldığını bildiririm.

DECLARATION PAGE

I hereby declare that all information in this document has been obtained and presented in accordance with academic rules and ethical conduct. I also declare that, as required by these rules and conduct, I have fully cited and referenced all material and results that are not original to this work.

İmza

Fürkan TANYERİ

Tarih:

ÖZET

YÜKSEK LİSANS TEZİ

SİBER GÜVENLİKTE XSS WEB SALDIRILARININ YAPAY ZEKÂ ZEMİNİNDE ANALİZ EDİLMESİ

Fürkan TANYERİ

Batman Üniversitesi Lisansüstü Eğitim Enstitüsü
Bilgi Teknolojileri Anabilim Dalı

Danışman: Prof. Dr. Ömer Faruk ERTUĞRUL

2024, 60 Sayfa

Jüri

Prof. Dr. Ömer Faruk ERTUĞRUL

Doç. Dr. Cafer BUDAK

Dr. Öğr. Üyesi Abdulkerim ÖZTEKİN

Bilgi ve teknoloji günümüz çağında insan hayatının vazgeçilmez bir parçası konumuna ulaşmıştır. İnsanlar bilgiye daha hızlı ve kolay ulaşabilmek için en çok interneti kullanmaktadır. İnternet insanların en fazla bilgilerin temin edildiği Web sitelerine ulaşmak için en önemli araçlardan biridir. Ancak Web sitelerinin bu kadar sıklık ile kullanımı kötü niyetli insanların dikkatini arttırmış ve Web güvenliği kavramı oluşmuştur. Bu Web güvenliği kavramı ise zaman içinde gelişerek ya da geliştirilerek yerini Siber güvenliği kavramına bırakmıştır. Siber güvenlik, bilgisayarları, ağları, yazılım uygulamaları, her türlü bilişim ile ilgili sistemleri ve bu sistemlerin verilerini kötü amaçlı saldırılardan korumayı amaçlayan bir yapı olması ile bilinmektedir. Gelişen teknoloji ve bu teknolojinin meydana getirdiği sürekli ileri düzeyde yenilenme ve gelişme durumu siber güvenlik sisteminin de sürekli ileri düzeyde kendini koruma amacıyla yenileme durumunu ortaya çıkarmıştır. Bu yenileme durumuna karşın siber güvenlik sistemine kötü niyetli olarak erişebilme isteği sonucu siber saldırı türleri meydana gelmiştir. Siber güvenlik konusunda sıklıkla karşılaşılan siber saldırı türleri Oltalama Saldırıları, Servis Dışı Bırakma Saldırıları (DoS/DDoS), Parola Saldırıları, Man-in-the-middle (MITM) Atak Saldırıları, Cross-Site Scripting (XSS) Saldırıları şeklinde sıralamak mümkündür. Cross-Site Scripting (XSS) Saldırıları, korunmasız bir web uygulamasının bir modülü olarak kötü amaçlı kod çalıştırmayı hedefler. XSS Saldırılarının birçok siber saldırı türlerinden farkı direk uygulamayı değil, kullanıcıyı hedeflemesidir. Teknolojinin hızlı gelişmesi, siber saldırıların artması gibi durumlar insanların daha güvenli, daha kolay verilerine ulaşabilmesi ve bilgilerden daha kesin sonuçlar alabilme isteğini arttırmıştır. Bu istek, insani zekâyâ sahip cihaz ve yazılım geliştiren bilgisayar bilimlerinin bir parçası olan Yapay Zekâ Teknolojisinin daha sık ve daha verimli şekilde kullanımını artmıştır. Bu çalışmamızda, Siber güvenlik sisteminde kullanılan siber güvenlik saldırı türlerinden biri olan Cross-Site Scripting (XSS) saldırılarının yapay zekâ zemininde analiz edilmesi ve analiz edilen verilerin değerlendirilmesi amaçlanmıştır.

Anahtar Kelimeler: Bilgisayar, Cross-Site Scripting (XSS), Siber Güvenlik, Siber Saldırı, Web, Yapay Zekâ, Yazılım.

ABSTRACT

MS THESIS

IN CYBER SECURITY XSS WEB ATTACKS ANALYSIS ON THE BASIS OF ARTIFICIAL INTELLIGENCE

Fürkan TANYERİ

**Batman University Graduate Education Institute
Department of Information Technologies**

Advisor: Prof. Dr. Ömer Faruk Ertuğrul

2024, 60 Pages

Jury

Prof. Dr. Ömer Faruk ERTUĞRUL

Assoc. Prof. Cafer BUDAK

Asst. Prof. Abdulkerim ÖZTEKİN

In today's age, information and technology have become an indispensable part of human life. People mostly use the internet to access information faster and easier. The Internet is one of the most important tools for people to access websites where the most information is provided. However, such frequent use of websites has increased the attention of malicious people and the concept of Web security has emerged. This concept of Web security has developed or improved over time and has been replaced by the concept of Cyber security. Cyber security is known as a structure that aims to protect computers, networks, software applications, all kinds of IT-related systems and the data of these systems from malicious attacks. Developing technology and the constantly advanced renewal and development caused by this technology have revealed that the cyber security system is constantly being renewed for the purpose of advanced self-protection. Despite this renewal, types of cyber attacks have occurred as a result of malicious attempts to access the cyber security system. The most common types of cyber attacks in cyber security can be listed as Phishing Attacks, Denial of Service Attacks (DoS/DDoS), Password Attacks, Man-in-the-middle (MITM) Attacks, and Cross-Site Scripting (XSS) Attacks. Cross-Site Scripting (XSS) Attacks aim to execute malicious code as a module of a vulnerable web application. The difference of XSS Attacks from many types of cyber attacks is that they target the user, not the application directly. Situations such as the rapid development of technology and the increase in cyber attacks have increased people's desire to access their data more securely and more easily and to obtain more accurate results from the information. This desire has increased the more frequent and more efficient use of Artificial Intelligence Technology, which is a part of computer science that develops devices and software with human intelligence. In this study, it is aimed to analyze Cross-Site Scripting (XSS) attacks, one of the types of cyber security attacks used in the cyber security system, on the basis of artificial intelligence and to evaluate the analyzed data.

Keywords: Artificial Intelligence, Computer, Cross-Site Scripting (XSS), Cyber Attack, Cyber Security, Software, Web.

ÖNSÖZ

“Siber Güvenlikte XSS Web Saldırıların Yapay Zekâ Zemininde Analiz Edilmesi” konulu çalışmamın amacı Siber güvenlik sisteminde kullanılan siber güvenlik saldırı türlerinden biri olan Cross-Site Scripting (XSS) saldırılarının yapay zekâ sisteminde analiz edilmesi ve analiz edilen verilerin değerlendirilmesidir.

Bu tez çalışmasının başlangıcından sonuna kadar her aşamasında bilimsel katkısını, tecrübesini esirgemeyen, desteği her daim hissettiğim ve tez çalışmam da bana yol gösteren değerli danışman hocam Sayın Prof. Dr. Ömer Faruk ERTUĞRUL’a sonsuz teşekkürlerimi sunarım.

Ayrıca her koşulda maddi ve manevi olarak beni yalnız bırakmayan hayat arkadaşım değerli eşime ve babama teşekkür ederim.

Fürkan TANYERİ
BATMAN-2024

İÇİNDEKİLER

ÖZET	iv
ABSTRACT.....	v
ÖNSÖZ	vi
İÇİNDEKİLER	vii
SİMGELER VE KISALTMALAR	ix
ŞEKİLLER ve TABLO LİSTESİ.....	x
1. GİRİŞ	1
2. SİBER GÜVENLİK.....	3
2.1. Siber Güvenlik Nedir?	4
2.2. Siber Güvenlik Temel Kavramları.....	4
2.3. Siber Güvenlik Çalışma Unsurları	5
2.4. Siber Risk Değerlendirilmesi ve Tehdit Modelleri.....	7
2.5. Yapay Zekâ	11
2.6. Makine Öğrenmesinde Saldırı Tespit Yöntemleri	12
2.7. Derin Öğrenme ve Anomaliliğe Dayalı Algılama	12
2.8. Yapay Zeka Temelli Siber Güvenlik Politikalar.....	13
2.9. XSS Saldırıları	14
2.9.1. XSS saldırılarından korunma unsurları.....	16
3. KAYNAK ARAŞTIRMASI	19
4. MATERYAL VE YÖNTEM.....	22
4.1. Materyal	22
4.2. Yöntem.....	30
4.3. Uygulamalar.....	37
4.3.1. Lojistik regresyon (LR) yöntemi uygulaması	38
4.3.2. K-Nearest neighbors (knn) yöntemi uygulaması	39
4.3.3. Karar ağaçları (CART)	40
4.3.4. Destek vektör makineleri (SVM).....	41
4.3.5. Naive bayes (Bayes)	42
5. ARAŞTIRMA SONUÇLARI VE TARTIŞMA.....	48
5.1. Araştırma Sonuçlarının Değerlendirilmesi	49
5.2. Tartışma	50
6. GENEL DEĞERLENDİRME VE SONUÇ.....	52
6.1. Genel Değerlendirme	52

6.2. Sonuç	53
KAYNAKLAR	55
ÖZGEÇMİŞ	61



SİMGELER VE KISALTMALAR

Kısaltmalar

AI	: Artificial Intelligence
API	: Application Programming Interface
APP	: Application
BAYES	: Naive Bayes
XSS	: Cross-Site Scripting saldırısı
MITM	: Man-in-the-middle saldırısı
DoS/DDoS	: Servis Dışı Bırakma Saldırıları
ITU	: International Telecommunications Union
PWC	: Price Water House Coopers
LR	: Lojistik Regresyon
KNN	: K-Nearest Neighbors
CART	: Karar Ağaçları
SVM	: Destek Vektör Makineleri
YZ	: Yapay Zeka

ŞEKİLLER LİSTESİ

Şekil 2.1. Siber Güvenlik temel unsurları.....	5
Şekil 2.2. Gizlilik, bütünlük ve erişilebilirlik konteksti.....	7
Şekil 2.3. Siber güvenliğin temel akış diyagramı	10
Şekil 2.4. YZ tanımlama ve sınıflandırması	11
Şekil 2.5. XSS saldırı şematiği	14
Şekil 2.6. XSS saldırı sıralaması	15
Şekil 4.1. Verilerin kategorize edilmesi	22
Şekil 4.2. Uygulama İsimleri.....	23
Şekil 4.3. İzinler Histogramı.....	24
Şekil 4.4. Api İsim Dağılımı.....	25
Şekil 4.5. Web Site İsimleri.....	26
Şekil 4.6. Ip adresleri.....	27
Şekil 4.7. Konum Değişken Dağılımı.....	28
Şekil 4.8. Etiket Değişeni	29
Şekil 4.9. Lojistik Regresyon Doğruluk Oranları.....	38
Şekil 4.10. K-Nearest Neighbors Doğruluk Oranları	39
Şekil 4.11. Karar Ağaçları Doğruluk Oranları	40
Şekil 4.12. Destek Vektör Makineleri Doğruluk Oranları.....	41
Şekil 4.13. Naive Bayes Doğruluk Oranları	42
Şekil 4.14. Lojistik Regresyon Confusion Matrisi	43
Şekil 4.15. K- En Nonfusion Matrisi	44
Şekil 4.16. SVM Confusion Matrisi	45
Şekil 4.17. Karar Ağacı Confusion Matrisi	46
Şekil 4.18. Naive Bayes Confusion Matrisi	47
Şekil 5.1. Makine Öğrenim Modelleri.....	48

TABLO LİSTESİ

Tablo 2.1. Siber risk değerlendirme adımları.....	9
---	---

1. GİRİŞ

Günümüz yakınçağına eşlik eden bilgi teknolojileri, insanoğlu için vazgeçilmez bir yaşam tarzı haline gelmiştir. Bilişim teknolojilerinin paradigmasının hızlı bir şekilde güçlenmesiyle, geçen zaman içerisinde askeri, siyasi, sağlık, eğitim ve sosyo-kültürel ve birçok alanda elzem bir hale gelmiştir. Bu paradigma faydaları olduğu kadar bizler için zararlıda olabilmektedir. Modern dünyadaki yaşantımız, bizleri birçok zeminde bilgi teknolojilerine mecburi bir hale gelmiştir. Bilişim dünyasındaki bu gelişmeler bir yandan bizlerin yaşantısını büyük bir ölçüde kolaylaştırıp, rahata kavuştururken diğer taraftan da siber güvenlik yönünden birçok sorunların ortaya çıkmasına olanak vermiştir. İnternetin fayda sağlayan bu getirimleri, internetin yaygın olarak kullanılmasından sonra kötü niyetli kişiler tarafından istismar edilmesi, siber saldırı, siber terör tehdidi, mala ve cana zarar verecek kadar büyük ölçekli bir boyut almış ve siber güvenlik kavramını bizler için zorunlu kılmıştır.

İnsanlığın başlangıcı ve günümüz ileri teknoloji çağına kadar ana ve temel gereksinimlerimizden biri de güvenlik olmuştur. Zaman ve çağlar değişse de güvenlik unsuru her daim hissedilen bir ihtiyaç haline gelmiştir. Bilişim çağının gelişmesiyle bu ihtiyaç hali hat safhaya çıkmış olup bu alandaki sorumluluklar ve problemler ülkemiz ve diğere ulus devletlerinde siber terör, siber güvenlik ve bilgi/veri güvenliği konularının önemini artırmış olup ve tartışmaya açık bir hale getirmiştir.

Bu devam edegelen teknolojik ilerleme süreci sınırları içerisinde olan internet ve bilgisayar teknolojileri, insanoğlunun ihtiyacı olan verilere hızlı ve basit bir şekilde erişebilmemize fırsat vermekte ve bu fırsatlar insanoğlu için bilgi teknolojilerinde devrim açmış ve bu teknolojik gelişmeler hızla artmış ve beraberinde bazı güvenlik sorun ve zafiyetlerinin meydana gelmesine neden olmuştur. İnternetin ve bilgisayarların dünya çapında yaygınlaşmasıyla, bu güvenlik sorunlarının çözümü için problemlerin sebebi olan teknolojik sistemler kullanılmaya başlamıştır. Bu zararların önüne geçilmesi bilgi elektronik bilgi güvenliği olan siber güvenlik kavramını doğurmuştur.

Siber güvenlik; hükümetlerin, kuruluşların, cemiyetlerin ve toplumların büyük sorunu ve konusu haline gelmiştir. Devletlerin mahrem bilgileri, ticari sırlar, işletmelerin gizli verileri ve sosyal toplumun parçası olan bizlerin kişisel bilgilerimizin güvenliğini her geçen gün risk altına sokmaktadır.

Siber güvenlikte temel amaçları özetlemek gerekirse, var olan sistem ve son kullanıcıları zararlı saldırılara karşı en iyi ve en hızlı bir şekilde engellemektir. Bu bağlam

insan karar mercii siber savunma her ne kadar başarılı olsa da bu savunmayı yapay zeka ile güçlendirmek hatta aşılmaz bir hale getirebilmek mümkündür.

Yapay zeka yani makine öğrenmesi, bilgisayar yâda bilgisayar dayanaklı bir makinenin insan karakteristiğindeki özelliklerden olan; çözüme kavuşturma, anlama ve öğrenme, yorumlayabilme gibi özelliklerin bir bilgisayar tarafından yapılabilmesidir. Bu nitelikleri insana göre çok daha hızlı yapabilmesi ve analitik çözüm odaklı olduğundan yapay zekayı teknolojinin bir çok farklı dallarında kullanıma sunmuştur.

Yapay zekayı, siber güvenlik kavramlarının güvenlik ile ilgili tespit edilen problemlerin belirleme ve analiz etme hatta bazı kötü amaçlı saldırılara karşı insan odaklı bir sistemin refleksinden daha iyi bir şekilde güvenlik sağlayabilir. Bu bağlamda siber güvenlik içerisindeki önemli konu başlıklarının en efektif bir şekilde yapay zeka algoritma penceresinden hızlı ve güvenilir bir şekilde bütünleştirilerek, gelebilecek tehditlere karşı belirlenen şart ve stratejiler içerisinde en sağlıklı otonom karşılık olarak verebilir.

Bu çalışmamızda, çağımızın en büyük güvenlik sorunlarından biri olan siber güvenlikte XSS saldırılarına karşı yapay zeka ile siber ortamda, farklı makine öğrenmesi algoritmalarının etkinliğini değerlendirdik ve karşılaştırdık.

2. SİBER GÜVENLİK

Bilgi teknolojileri kavramında kullanılan siber kelimesi, Çağımızda bilişim sistemleri altyapısındaki ağlar olarak tanımlanır (web kaynak -1).

Siber, sayısal formata dönüştürülmüş verilerden meydana gelen, depolanan ve paylaşılan bir ortamdır. Kelime anlamı olarak, dijital ağ sistemlerini içeren geniş bir kavrama sahiptir. İngilizce “Cyber” kelime kökeninden gelmiş olup, kullanılmaya başlanmış ve diğer birçok anlamlar içermektedir. Kısaca web internet ile ilgili tüm unsurlar siber sınıfa girer (web kaynak -2).

Siber uzay ise Türkçe karşılığı siber ortam olan, tüm dünya da ve uzayda bulunan bilişim teknoloji ağ sistemlerinden oluşan özgür, bağımsız ve hür bilgi sistemlerinin sayısal şeklini tanımlar. Bu ortamda bulunan her çeşit bilgi, doküman, araç, donanım, yazılım, sunucu ve internet ortamında aktif olan sistemler de siber ortamdaki varlıklardır.

Siber ortamdaki varlıkların zarar görme olasılığının da siber riskler teşkil eder. Siber risklerin olumsuz getirileri arasında yer alan siber tehditler, sanal internet ortamına bağlı olan bilişim/bilgi teknoloji ağ ve cihazlara yönelik yapılan, belirli bir amaç ve istekler doğrultusundaki saldırılardır. Yapılan saldırılar, sistem açık ve zafiyetlerin gözetererek bilgi güvenliğine zarar verebilmektedir.

Siber güvenlik konsepti, internet bağlantılı ya da internet bağlantısı olmadan bilişim organizasyonundaki kurumsal ve kişisel alt yapıların korunması olarak adlandırabiliriz. Bu alt yapının içinde olan bilişim sistemleri ve elektronik sistemlerinin en üst düzeyde korunması bu düzende yapılan iş ve işlemlerin gizliliği, bütünlüğü ve kolay ulaşılabilirliğini teminata alarak bu tehdit edici siber saldırıları saptama ve saldırılara karşı ani reaksiyon alınmasıdır.

Günümüzdeki dijital dönüşüm sürecinin hızlanmasıyla birlikte siber güvenlik konusundaki önemli değişimlere dikkat çekmektedir. Dijitalleşme, iş süreçlerinde ve günlük yaşamda daha fazla dijital araç ve hizmet kullanılmasına yol açarak büyük bir veri akışına neden olmaktadır. Ancak, bu dijital ortamda faaliyet gösteren kurumlar ve bireyler için siber güvenlik sorunları da giderek daha karmaşık ve zorlu bir hal almıştır.

2.1. Siber Güvenlik Nedir?

Siber güvenlik, internet bağlantılı ya da bağlantısız bilişim altyapı hizmetlerinin sunumuna yönelik bir disiplindir. Bu doğal, siber güvenlik; kurumsal ve kişisel düzeyde kullanılan bilişim sistemleri, elektronik cihazlar, ağlar ve veri depolama birimlerinin korunması amacıyla uygulanan yöntem ve modülleri içerir. Temel amaç, bu sistemlerde büyümenin iş ve işletim özgürlüğünün, bütünlüğünü ve erişilebilirliğinin sağlanmasıdır. Gizlilik, kişisel kişilerin izinlerini engellemeyi engellenmeyi yaparken; bütünlüğün, verilerin korunmasının ve korunmasının odağına. Erişilebilirlik ise, yetkili kişinin ihtiyacı olan bilgi ve sistemlerin zamanında ve sürekli olarak erişilebilmesini sağlar

Günümüzdeki dijital dönüşüm sürecinin hızlanmasıyla birlikte siber güvenlik konusundaki önemli değişimlere dikkat çekmektedir. Dijitalleşme, iş süreçlerinde ve günlük yaşamda daha fazla dijital araç ve hizmet kullanılmasına yol açarak büyük bir veri akışına neden olmaktadır. Ancak, bu dijital ortamda faaliyet gösteren kurumlar ve bireyler için siber güvenlik sorunları da giderek daha karmaşık ve zorlu bir hal almıştır.

2.2. Siber Güvenlik Temel Kavramları

Siber sahada tehditlerin meydana gelişi ile birlikte siber güvenlik konusu da doğal olarak bir ihtiyaç konusu haline gelmiştir. Özel ve tüzel kurumlarda kullanılan elektronik bilişim cihaz ve sistemlerde var olan verilerin fazlalığı ile bu tehdit oranını daha da artmıştır. Siber güvenlik kavramı gerçekte önemli ve gizli bilgilerin mevcut olduğu her koşulda risk ve tehdit faktörünü oluşturmaktadır (Coventry ve ark., 2014).

Artan rekabet koşulu durumlarıyla bu tehditlerin saldırı şeklinde geri dönmesi hem özel sektör hem de kamu sektöründe ciddi bir şekilde tehlike unsuru oluşturmaktadır. Siber güvenlik kavramı ise burada devreye girerek, gerçekte önemli ve gizli bilgilerin mevcut olduğu her koşulda risk ve tehdit faktörünü engelleme amacını görev edinmiştir (Öztürk, 2018).



Şekil 2.1. Siber güvenlik temel unsurları (Yaşar, 2014)

Bu kavram kötü amaçlı girişimlerin önlenmesi, sistemlere sızmaların ortaya çıkarılması, zararlı yazılımların saptanması ve engellenmesi, kimlik doğrulamanın ve şifreli iletişimin etkinleştirilmesini zorunlu olarak yapılmasına imkân sağlayacak donanımı kapsamaktadır (Amoroso, 2006). Şekil 2.1’de belirtildiği gibi siber güvenliğin, uluslararası alanda belirli disiplinler arasında temellerinin oluşturulduğunu görmekteyiz.

2.3. Siber Güvenlik Çalışma Unsurları

Siber alanda ileri seviye güvenliğin alınmasında genel husus ve şartlara dikkat edilmesi ve bu hususların yerine getirilmesiyle sağlanır. Şekil 2.2’deki genel başlıklarıyla bu hususlar; erişilebilirlik, kimlik doğrulama, bütünlük, gizlilik ve inkâr edememe gibi unsurlardır. Kısaca bu hususlardan bahsedecek olursak;

Gizlilik, verilerin yalnızca erişim yetkisi olan kişilerin kullanımında olduğu ve bunu garanti etmektir. Başka bir deyiş ile verilerin erişimine yetki ve kontrol dışında olan kişileri erişilebilirliğini engelleme girişimidir.

Bütünlük, bir diğer siber güvenlik unsurudur. Verilerin kontrol ve işleme metotlarının doğruluğunu ve veri içeriklerinin değişmezliğini sağlamak, stabil tutmak

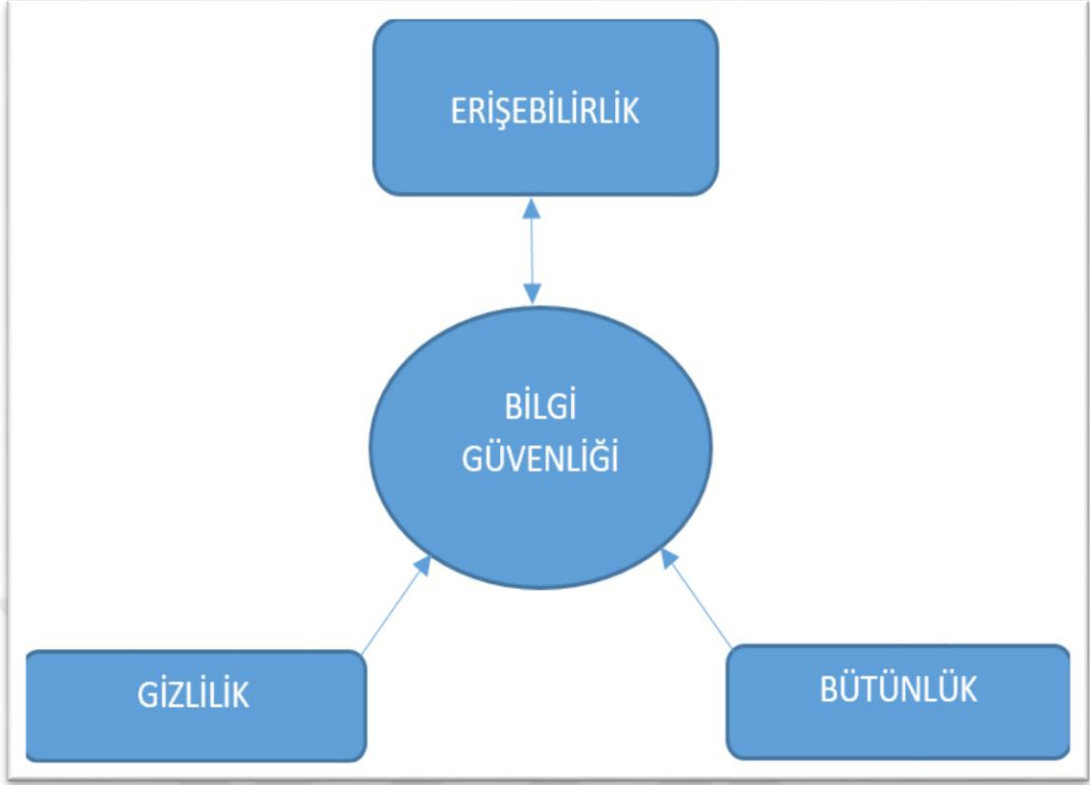
veya verinin farklı ortamlara aktarılmasını doğrulamak olarak da tanımlanabilir. Yüz tanıma gibi biyometrik işaretler kullanılarak çeşitli farklı güvenlik fonksiyonlar sağlanabilir.

Kimlik doğrulama, yetki verilmiş kullanıcının kimliğinin teyit edilerek doğrulanması ve kullanıcının o kişi olduğunun sağlanmasının yapılması veya istifade eden kullanıcının tanımlanması ve onaylanması işlemidir. Elektronik imza, bu unsurun sağlanmasını temin eder.

İnkâr edememe, ileti veya veri kaynağının ilettiği mesajı veya gönderilen verinin yalanlayamamasıdır. Başka bir ifade ile sorumlu olan kullanıcının mesaj iletim veya alım işleyişinin kanıtlanması veya gönderim veya alım işleminin inkâr edilmesini sağlanması olayına denir. Zaman dalgası, açık anahtar zemini ile gerçekleşir.

Erişilebilirlik, yetki verilmiş olan kullanıcıların veriyle alakalı kaynaklara ulaşım yetkisine sahip kullanıcı olma güvencesinin verilmesidir. Yetki verilmiş hak sahiplerinin sistemlere emniyetli ve devamlı olarak erişimlerinin garanti edilmesi veya verilen servis hizmetinin devamlılığının sağlanması için gerekli olan bütün tedbir ve önlemlerin alınması olarak tanımlanabilir.

International Telecommunications Union (ITU) tarafından da onaylanan, istikrarın temel hedefi, elektronik veya diğer bilgi teknolojileri alanlarında muhafaza edilen her türlü verinin; Kullanılabilirliğini, bütünlüğünü ve gizliliğini sürekli olarak sağlamaktır (web kaynak – 3).



Şekil 1.2. Gizlilik, bütünlük ve erişilebilirlik konteksti (yazar tarafından tasarlanmıştır)

2.4. Siber Risk Değerlendirilmesi ve Tehdit Modelleri

Günümüz teknoloji çağında, güvenlik riskleri arasında ön plana çıkan siber güvenliğin önemli bir rol oynadığını vurgulamak adına PWC firmasının 20178 yılında yaptığı “Küresel Ekonomide Suçlar ve Hile Araştırmaları” araştırmada hile kategorisinde “siber suçlar” ilk beş sıralama yer almıştır (PWC, 2018).

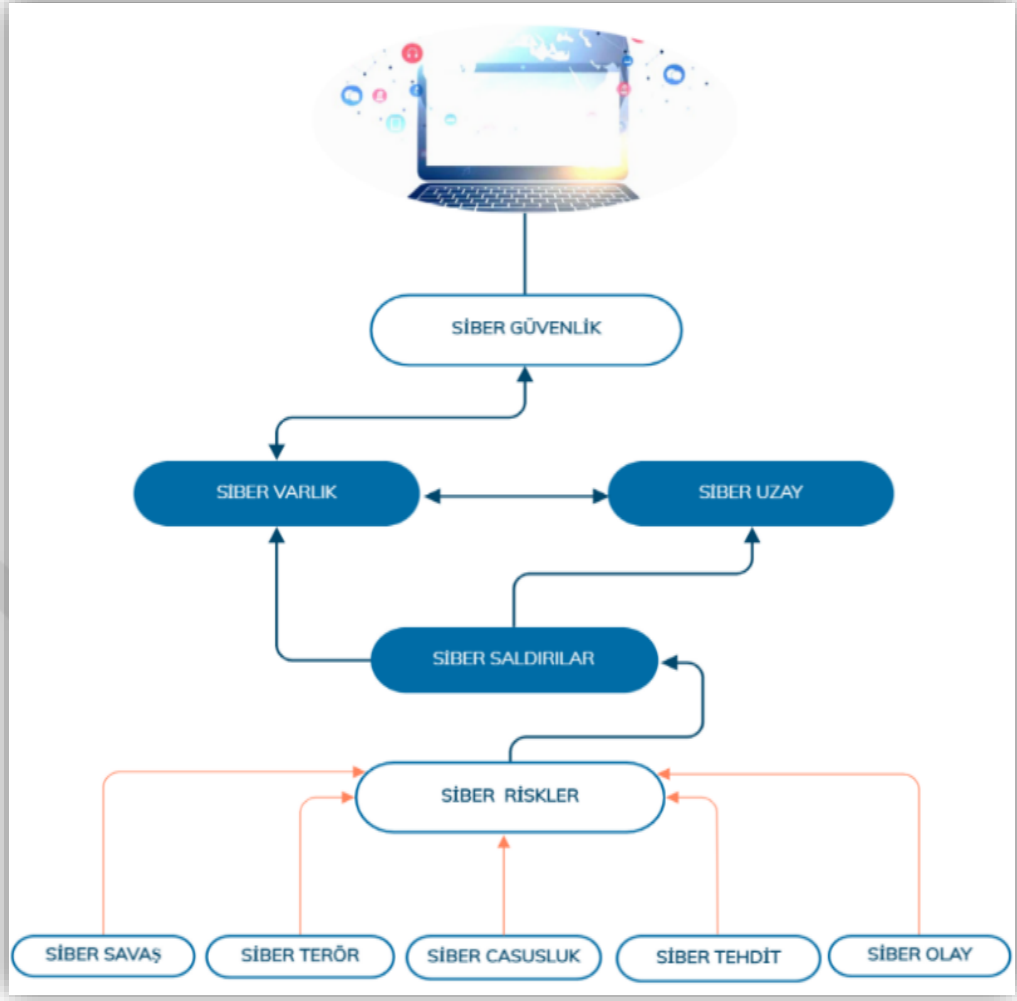
Siber güvenlik risk değerlendirmesi, bir organizasyonun bilgi varlıklarını tehdit eden potansiyel riskleri belirlemek, analiz etmek ve bu riskleri yönetmek için kullanılan sistematik bir süreçtir. Bu değerlendirme, bilgi sistemlerinin güvenliğini sağlamak ve olası zararları minimize etmek için kritik öneme sahiptir. Siber güvenlik risk değerlendirmesinin adımları:



Tablo 2.1. Siber risk deęerlendirme adımları

Adımlar	Açıklama
Varlıkların Belirlenmesi	Organizasyonun bilgi varlıklarının ve bunların deęerlerinin tanımlanması.
Tehditlerin ve Zafiyetlerin Belirlenmesi	Bilgi varlıklarını tehdit eden potansiyel tehditlerin ve mevcut zafiyetlerin tespiti.
Risklerin Analizi ve Deęerlendirilmesi	Tanımlanan tehditlerin ve zafiyetlerin bilgi varlıkları üzerindeki potansiyel etkilerinin ve olasılıklarının analizi ve deęerlendirilmesi.
Risklerin Önceliklendirilmesi	Analiz edilen risklerin önem sırasına göre önceliklendirilmesi.
Risk Azaltma Stratejilerinin Geliştirilmesi	Önceliklendirilmiş risklerin etkilerini azaltmak için stratejilerin geliştirilmesi.
Risk Yönetimi Planının Hazırlanması ve Uygulanması	Geliştirilen stratejilere göre bir risk yönetimi planının oluşturulması ve uygulanması.
Sürekli İzleme ve Gözden Geçirme	Uygulanan risk yönetimi planının etkinliğinin sürekli olarak izlenmesi ve gerektiğinde gözden geçirilmesi.
Raporlama ve İletişim	Risk deęerlendirmesi ve yönetimi sürecinin sonuçlarının ilgili paydaşlarla raporlanması ve paylaşılması.

Bu adımlar, siber güvenlik risklerinin etkin bir şekilde yönetilmesini sağlar ve organizasyonun bilgi varlıklarının korunmasına yardımcı olur. Risk deęerlendirmesi, sürekli bir süreç olup, organizasyonun büyümesi, teknolojinin gelişmesi ve tehditlerin evrimi ile birlikte düzenli olarak tekrarlanmalıdır.



Şekil 2.2. Siber güvenliğin temel akış diyagramı (yazar tarafından hazırlanmıştır)

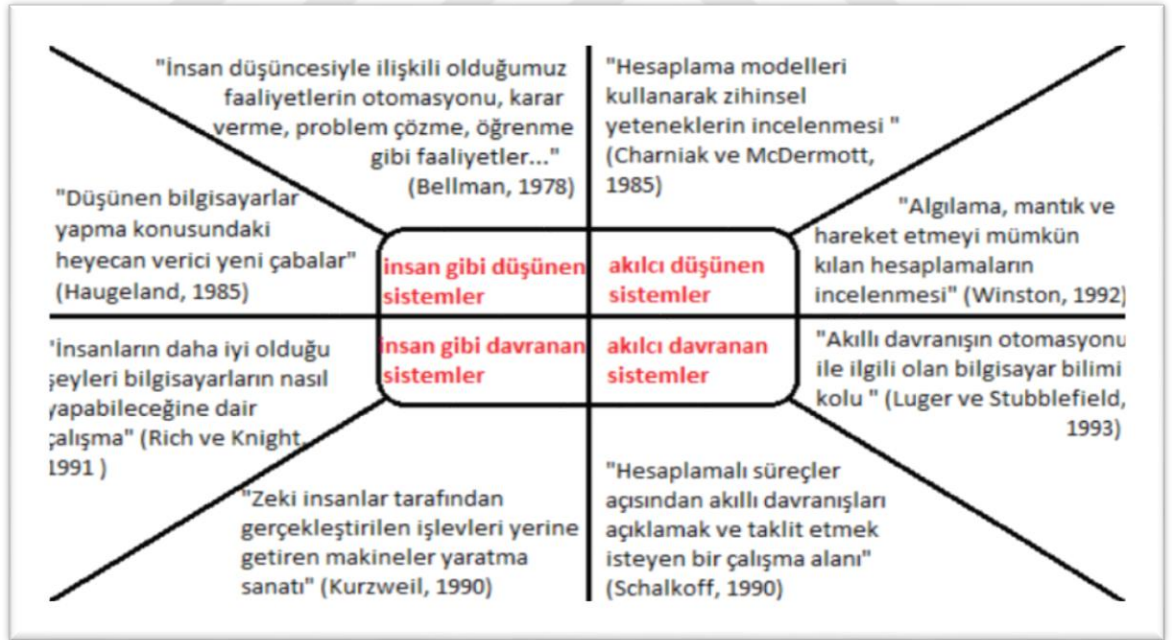
Şekil 2.3’de bahsedildiği gibi siber saldırıların genel hatları ve tehdit süreçlerinden bahsetmektedir. Bu tehdit sürecinin hangi aşamalardan ve doğuracağı olumsuz etkileri önlemek adına, siber savunma da önde gelen ilkelere en önemlilerinden sanal dünyadaki güvenlik ve gizliliğin gelebilecek tehlikelere karşı korunması amaçlanmaktadır. Güvenlik ve gizlilik konusu çağımızda teknolojik gelişim ve ilerlemelerin bir sonucu olarak en fazla siber tehditlerin odağı olarak istismar ve saldırılara uğramaktadır. Bu saldırılar karşısında verilerimizin muhafazası, depolanması ve güvenliğinin sağlanması konusunda güvenlik fazlı tedbirler ve önemler alınmalıdır.

2.5. Yapay Zekâ

Yapay zeka ve yapay öğrenme kavramı, birçok güncel teknolojik alanda kullanılmaktadır. Öğrenmenin doğal haline benzer biçimde işleyen “davranışsal öğrenme” bazlı bir yapay öğrenme algoritma veya modelin geliştirilmesi bu alanda günümüz çalışmalarına ışık tutmaktadır (Ertuğrul, 2015). Matematik, bilgisayar mühendisliği, elektronik, dil bilimi, felsefe ve bilgi teknolojileri gibi farklı bilim branşlarının bir araya gelmesiyle birçok disiplinler arası alanın oluşturduğu bir kavram olarak da tanımlayabiliriz.

Yapay zekâ (AI-Artificial Intelligence) genel tanım olarak, günümüz teknoloji çağında insanoğlu tarafından başarıyla yerine getirilen faaliyetlerin, gelişmiş bilgisayarların bilgi teknolojileriyle yazılımsal ve donanımsal gelişmelerle bir çok farklı alanda iş ve işlemlerin çok daha etkili ve hızlı olarak yapılmasını sağlayan sistemlerdir (Russel ve ark., 1995).

Bu bağlamda yapay zekâyı Şekil 4’de gösterildiği gibi sekiz farklı tariften referans alınarak iki ana kavrama göre sınıflandırabiliriz;



Şekil 2.3. YZ tanımlama ve sınıflandırması (Russel ve ark., 1995)

- Akılcı düşünen sistemler; düşünce yasaları ile,
- İnsan gibi davranan sistemler; Turing test sistemiyle,
- Akılcı davranan sistemler; rasyonel ajan yaklaşımıyla,

- İnsan gibi düşünen sistemler; kognitif modelleme tarzı ile tarif edilebilmektedir (Russel ve ark., 1995).

2.6. Makine Öğrenmesinde Saldırı Tespit Yöntemleri

Artırılmış zekanın yani yapay zekanın farklı bir alanı olan makine öğrenmesi, bilgi kümesinin kompleks tarzının tespit edilmesiyle rasyonel karar almak için istatistik ve bilgisayarın hesaplama olanağından faydalanır. Saldırı tespit yöntemlerinden ne fazla kullanılan makine öğrenmesi teknikleri şunlardır (Kaya ve ark, 2014);

- Karar ağaçları; kategorizasyon ve regresyon yöntemi.
- Yapay sinir ağları; insan beyin fizyolojisini taklit ederek öğrenme.
- Bayes sınıflama; sınıflandırma ve tespit etme.
- Destek Vektör Makinesi; istatistiksel öğrenme yöntemi.

2.7. Derin Öğrenme ve Anomaliliğe Dayalı Algılama

Derin öğrenme biz insanlarda olduğu gibi bir makinenin tahmin edebilme, ses ve görüntü ayırt etme ve tanımlama vb. gibi işlemleri yapmak amacıyla eğitip/öğreten bir öğrenim tipidir (Bengio ve ark., 2016).

Derin öğrenme yöntemleri günümüz teknolojisinde algılama, çıkarım yapma, sınıflandırma gibi yetenekleri mevcuttur (Le ve ark., 2011). Bu amaçla daha önceden tanımlanmış olan bir siber güvenlik veri kümesinin, herhangi bir anomalinin olup olmadığını algılamak için örnek olarak spam olarak gelen mesaj/maillerin filtrelenmesinde ya da engellenmesinde önemli bir oynamaktadır.

Derin öğrenmeye dayalı anomalide ise, Örnek olarak spam mesaj yada maillerin siber güvenlik bakımından şüphe ve tehdit uyandırabilecek içeriklere sahip olanların, ayırt edici anomali özelliklerinden birine sahipse o içeriğin sınıflandırılması, tanımlandırılması ile gelen zararlı verinin son kullanıcıya ulaşmasının engellenmesidir.

2.8. Yapay Zeka Temelli Siber Güvenlik Politikalar

Siber tehditlerin artışıyla beraber, siber güvenlik unsurlarının öneminin artırılması, güvenlik ile ilgili hedef ve politikaların belirlenmesi konusunda önleyici tedbir ve reaksiyonel faaliyetlerin etkinliği güçlü rol oynamaktadır.

Teknoloji çağımızın getirdiği ileri teknolojinin olumsuz getirilerinden siber saldırıların ve bu saldırılarda kullanılan bilgi teknoloji cihazlarının önlenemez bir hızla artmasıyla yapay zekâ ve akıllı sistemlerle yapılan siber tehditlerden bilgi sahibi olmayı, yapılan eylemlerdeki saldırıları algılamayı otomatikleştirme de geleneksel yazılım ve manuel yöntemlerden daha verimli sonuçlar alınabilmektedir (Belani, 2021).

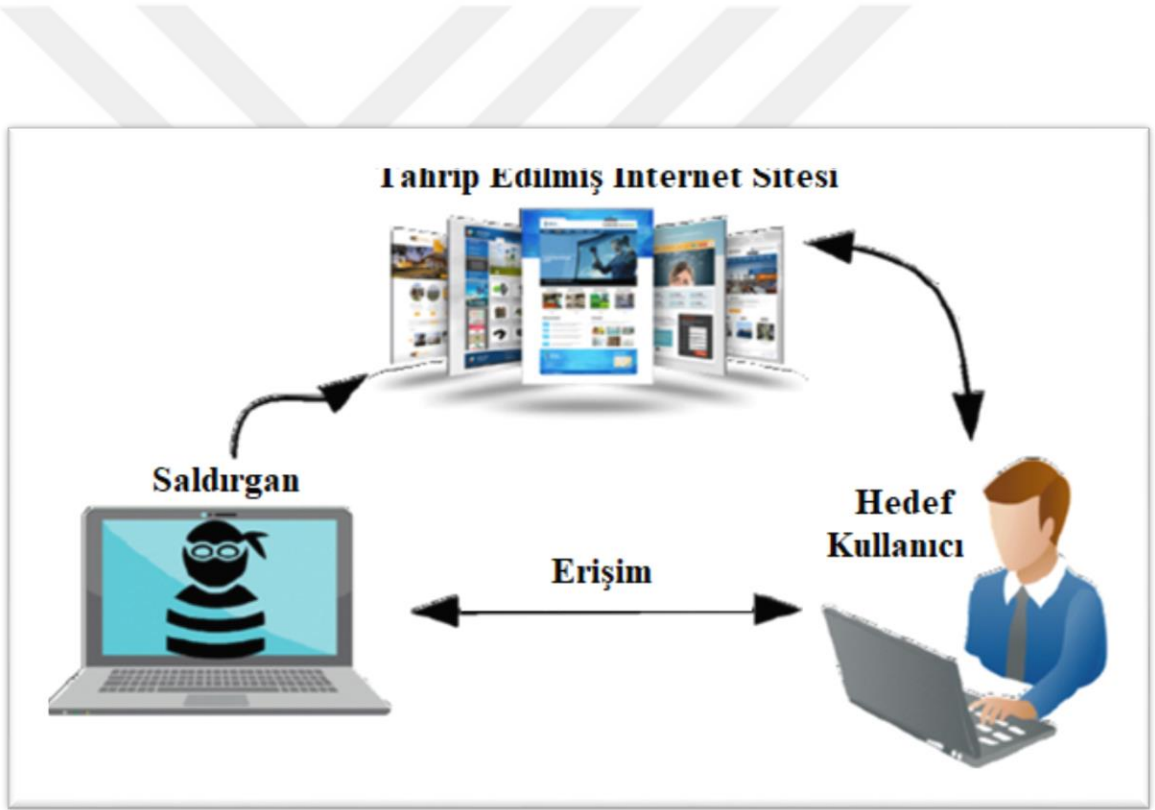
Yapılan saldırıların engellenmesi, güvenlik sistemlerinin yapay zeka ile entegre edilip geliştirilmesine bağlıdır. Bunun en önemli sebebi herhangi bir siber güvenlik saldırısını yapay zeka temelli yazılım ve programlara öncelik verilerek, saldırı esnasında hangi etmenlerin kullanılacağına ve olabilecek tehdit ihtimallerini hızlı bir oluşumda belirleyip değerlendirmesi, saldırı yapılmadan önce güvenlik zafiyetlerinin tespitine, sistemde var olan açıkların geliştirilip planlanmasına yardımcı olmaktadır. Saldırı anında ise ya da olabilecek saldırı riskinde tespit uyarısı anında hızlı bir şekilde cevap vermesi fayda sağlamaktadır (Reşitoğlu, 2021).

Temel olarak kısaca, siber güvenlik kapsamında yapay zeka uygulamalarının, tehditlerin tespiti, tahmin ve yanıt vermesiyle üç temel fonksiyonu başarılı bir şekilde yerine getirdiği söylenebilir.

2.9. XSS Saldırıları

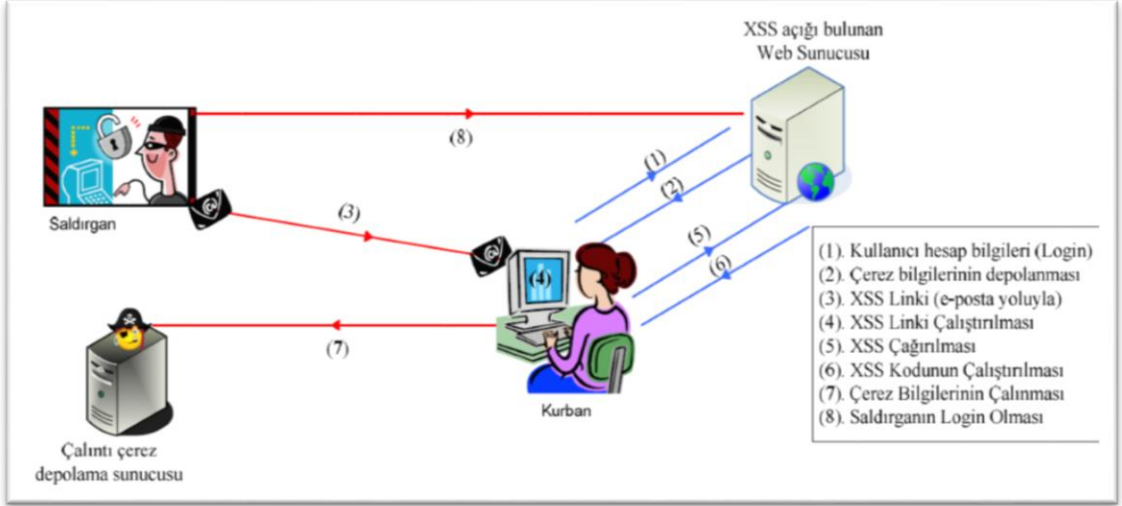
Son kullanıcıdan alınan verinin belirli bir kontrol ve gözetimden geçirilmeyerek kullanıcıya demonstrasyon edilecek sayfanın kaynak kodları içerisinde mevcut bulunan javascript kodlarının eklenmesiyle istemci tarafından betik çalıştırılması işlemidir (Özel ve ark., 2014).

Kısaca tanımlamak gerekirse, bu saldırıların odak noktası olan herhangi bir internet sitesinin tahrip edilip ve son kullanıcının siteyi ziyareti esnasında bilgisayarına sızıp, zararlı XSS kodlarının hedef bilgisayar ile sistemde kayıtlı bir şifre veya parolaları elde edebilmek için yapılan bir siber saldırı çeşididir (Muhammed ve ark., 2018). XSS saldırı sisteminin şematik olarak şekil-5’de verilmiştir (Kara, 2019).



Şekil 2.4. XSS saldırı şematiği (Kara, 2019)

XSS saldırı yöntemi, içeriği zararlı hale getirilmiş kod veya kodların var olan web tarayıcısında bu zararlı kodların web sitesi sunucusunun tanımlı olan tarayıcıdaki güvenlik ayarlı kapsamında çalışacaktır. Web sitesi içerisinde herhangi bir sınırlama yoksa zararlı hale getirilmiş olan kod tarayıcı ekseninden her nevi hassas veriye erişebilir, değiştirme işlemi yapabilir ve mail yolu ile istenilen adreslere gönderebilir.



Şekil 2.5. XSS saldırı sıralaması (Vural ve ark., - 2008)

Bilgisayar içerisindeki kullanıcının oturum çerezleri çalınabilir, web tarayıcı istenilen adrese yönlendirebilir, web sayfaları aracılığıyla veri elde etme amaçlı kötü amaca hizmet eden kodlar çalıştırılabilir, oltama tekniğiyle sazan avlama yöntemine davet çıkarılabilir. Bu saldırıların şematik sıralaması Şekil 2.6'da gösterilmiştir (Vural ve ark., - 2008).

Şekil 6'da gösterilen XSS saldırı adımlarını detaylı olarak ele almak gerekirse, bu adımların her birinin saldırının nasıl gerçekleştiğini adım adım ortaya koyduğunu ve siber güvenlik açısından ne tür riskler barındırdığını anlamak için dikkatle incelenmesi gerektiği söylenebilir. Bu bağlamda, XSS saldırısının başlangıcından itibaren hangi yöntemlerle gerçekleştirildiği, hangi güvenlik açıklarının suistimal edildiği ve sonuç olarak kullanıcıya ait hangi hassas bilgilerin tehlikeye atıldığı gibi kritik noktaların ayrıntılı bir şekilde değerlendirilmesi önem arz etmektedir.

Kullanıcı hesap bilgilerinin gönderilmesi; kullanıcı, bir web uygulamasına giriş yapmak amacıyla kimlik bilgilerini (örneğin, kullanıcı adı ve parola) sunucuya gönderir. Bu işlem, genellikle bir HTTP POST isteği ile gerçekleştirilir ve güvenli iletişim için HTTPS kullanılması önerilir (Gollmann, 2011).

Çerez bilgilerinin depolanması; kimlik doğrulama işlemi sonrasında web sunucusu, kullanıcıya ait oturum bilgilerini içeren bir çerez oluşturur ve bu çerez kullanıcının tarayıcısına gönderilir. Çerezler, sonraki isteklerde kullanıcının oturum durumunu belirlemek için sunucu tarafından kullanılır (Barth, 2011).

XSS Linkinin (E-posta Yoluyla) Gönderilmesi; saldırgan, kurbanın tarayıcısında çalıştırılmak üzere zararlı bir XSS (Cross-Site Scripting) kodu içeren bir bağlantı

oluşturur ve bu bağlantıyı kurbanı e-posta veya başka bir iletişim kanalı aracılığıyla gönderir (Grossman ve ark., 2007).

XSS linkinin çalıştırılması; kurban, saldırganın gönderdiği zararlı bağlantıya tıklayarak XSS kodunun kendi tarayıcısında çalışmasına neden olur. Bu aşamada saldırgan, kurbanın tarayıcısında zararlı kodun çalıştırılmasını sağlar (Jovanovic ve ark., 2006).

XSS kodunun çağırılması; kurbanın tarayıcısında çalıştırılan zararlı XSS kodu, tarayıcıda depolanan çerezler gibi hassas verilere erişir ve bu veriler saldırganın belirlediği bir sunucuya iletilir (Kirda ve ark., 2005).

Çerez bilgilerinin çalınması; saldırgan, kurbanın çerez bilgilerini ele geçirir ve bu bilgileri kendi kontrolündeki bir sunucuda depolar. Bu çerezler, kurbanın kimliğine bürünmek için kullanılabilir (Zakas, 2012).

Saldırganın kurbanın hesabına giriş yapması; saldırgan, çalınan çerez bilgilerini kullanarak kurbanın hesabına oturum açar ve bu hesap üzerinden çeşitli kötü amaçlı faaliyetlerde bulunabilir (Gupta, ve ark., 2016).

2.9.1. XSS Saldırılarından korunma unsurları

XSS (Cross-Site Scripting) saldırıları, web uygulamalarına yönelik en yaygın ve tehlikeli saldırı türlerinden biri olarak öne çıkmaktadır. Bu tür saldırılar, kötü niyetli aktörlerin, kullanıcıların tarayıcılarında zararlı betikler çalıştırmasına olanak tanır ve veri hırsızlığından oturum bilgilerini ele geçirmeye kadar çeşitli tehditler oluşturabilir. XSS saldırılarına karşı alınabilecek önlemler aşağıda detaylandırılmıştır:

Girdi doğrulama ve filtreleme; XSS saldırılarına karşı ilk savunma hattını oluşturur. Kullanıcıdan alınan tüm veriler, sunucuya ulaşmadan önce doğrulanmalı ve zararlı içeriklerden arındırılmalıdır. Beyaz liste kullanımı, girdi doğrulamada tercih edilmelidir. Bu yöntemde, yalnızca beklenen ve güvenli kabul edilen karakterler ve veri türleri işleme alınır (Kiczales ve ark., 2001). Zararlı olabilecek belirli karakterlerin (örneğin, <, >, ', ", &) engellenmesi de önemlidir, ancak bu yaklaşım beyaz liste kadar etkili olmayabilir (Sullivan ve ark., 2008)

Çıktı kodlama; kullanıcı girdilerinin HTML, JavaScript, CSS veya URL içinde kullanılması durumunda, bu girdilerin uygun şekilde kodlanması gerekir. HTML kodlama, kullanıcı tarafından sağlanan verilerin HTML etiketlerinde kullanılacağı durumlarda mutlaka yapılmalıdır (Barth, 2011). JavaScript içinde kullanıcı verisi

kullanılırken, bu verilerin güvenli bir şekilde kodlanması gerekir (Jovanovic ve ark., 2006). URL'lerde yer alan kullanıcı girdileri, URL kodlaması ile güvenli hale getirilmelidir (Grossman ve ark., 2007).

Güvenlik kütüphaneleri ve çerçeveleri kullanma; güvenlik açıklarını önlemek için güvenli kütüphaneler ve çerçeveler kullanılması önerilir. OWASP Enterprise Security API (ESAPI), XSS ve diğer güvenlik açıklarına karşı korunmaya yardımcı olan çeşitli güvenlik kütüphaneleri sunar (OWASP, 2017). HTML Purifier gibi kütüphaneler, kullanıcı girdilerini güvenli hale getirerek XSS saldırılarına karşı etkili bir savunma sağlar (Pham, 2013).

İçerik güvenliği politikası (CSP); CSP tarayıcıya yalnızca belirli kaynaklardan yükleme yapılmasına izin verir ve hangi komut dosyalarının çalıştırılabileceğini belirler. Betik kaynağını sınırlama, CSP ile yalnızca güvenilir ve belirli kaynaklardan betiklerin yüklenmesine izin verilebilir, bu da XSS saldırılarını engellemede önemli bir rol oynar (Gollmann, 2011). CSP ayrıca inline script'lerin çalıştırılmasını engelleyerek XSS saldırılarının etkisini azaltabilir (Stampar ve ark., 2014).

Http-only ve secure bayrakları; çerezlerin (cookies) HTTP-only ve Secure bayrakları ile işaretlenmesi, çerezlerin JavaScript tarafından erişilmesini ve yalnızca HTTPS üzerinden iletilmesini sağlar. HTTP-Only bayrağı, çerezlerin JavaScript tarafından okunmasını engeller, böylece XSS saldırıları yoluyla çerez hırsızlığını önler (Barth, 2011). Secure bayrağı ise çerezlerin yalnızca güvenli HTTPS bağlantıları üzerinden iletilmesini sağlar ve çerezlerin güvenliğini artırır (Sullivan ve ark., 2008).

Güvenlik duvarları ve WAF kullanımı; çerezlerin (cookies) HTTP-only ve Secure bayrakları ile işaretlenmesi, çerezlerin JavaScript tarafından erişilmesini ve yalnızca HTTPS üzerinden iletilmesini sağlar. HTTP-Only bayrağı, çerezlerin JavaScript tarafından okunmasını engeller, böylece XSS saldırıları yoluyla çerez hırsızlığını önler (Barth, 2011). Secure bayrağı ise çerezlerin yalnızca güvenli HTTPS bağlantıları üzerinden iletilmesini sağlar ve çerezlerin güvenliğini artırır (Sullivan ve ark., 2008).

Düzenli güvenlik denetimleri; web uygulamalarının düzenli olarak güvenlik testlerinden geçirilmesi ve potansiyel açıkların kapatılması önemlidir. Penetrasyon testleri, XSS ve diğer güvenlik açıklarını tespit etmek için düzenli olarak yapılmalıdır (Kirda ve ark., 2005). Ayrıca, kodun güvenlik açısından incelenmesi ve güvenli kodlama uygulamalarının sağlanması gerekmektedir (Zakas, 2012).

Kullanıcı eğitimi ve bilinçlendirme; geliştiricilerin ve kullanıcıların XSS saldırıları hakkında bilinçlendirilmesi, güvenlik açıklarını önlemede önemli bir rol oynar.

Geliştiricilere güvenli kodlama pratiği ve XSS saldırılarına karşı koruma yöntemleri hakkında eğitim verilmelidir (Sullivan ve ark., 2008). Kullanıcılara ise XSS saldırılarının farkında olmaları ve şüpheli bağlantılara tıklamamaları konusunda bilgilendirme yapılmalıdır (Grossman ve ark., 2007).

XSS saldırılarına karşı korunmak, birden fazla güvenlik önleminin birlikte uygulanmasını gerektirir. Yukarıda belirtilen yöntemlerin entegre edilmesi, web uygulamalarının güvenliğini artıracak ve XSS saldırılarından korunmayı sağlayacaktır



3. KAYNAK ARAŞTIRMASI

Siber Güvenlikte XSS Web Saldırılarının Yapay Zekâ Zemininde Analiz Edilmesi ile alakalı yapılmış olan çalışmaların kronolojik sıralaması aşağıda verilmiştir.

Sevri (2011). “Web Saldırılarının Tespitine Yönelik Yapay Zekâ Tabanlı Saldırıların Tespitine Yönelik Zekâ Tabanlı Bir Güvenlik Modülü Geliştirilmesi” Bu çalışmada, web uygulamaları ile ilgili gerçekleştirilen web saldırıları incelenmiş ve güvenlik duvarları üzerine konumlandırılacak yapay zekâ tabanlı yazılım modeli ile bu atakların tespit edilmesi, önlenmesi ve raporlanması amaçlanmaktadır.

Kartal, Sağıroğlu, Bülbül (2013). “IPV6’da Güvenlik Açıklarına Genel Bir Bakış” Bu çalışmada, yeni protokol ile ilgili yeniliklerden ve IPV4 ile benzerlik ve farklılıklara değinilmiş ve genel olarak öneri açıklamalarından bahsedilmiştir.

Yaşar (2014). “Kurumsal Siber Güvenliğe Yönelik Tehditler ve Mücadele Yöntemleri: Eylem Planı Örneği” Bu çalışmada, son zamanlarda ve önümüzdeki dönemlerde de etkilerinin sürdürüleceği tahmin edilen siber tehditler incelenmiştir ve bu tehditlere karşı alınabilecek önlemler ile yapılması gerekenler sunularak örnek bir kurumsal siber güvenlik eylem planı hazırlanmıştır.

Yiğit, Akyıldız (2014). “Sızma Testleri İçin Bir Model Ağ Üzerinde Siber Saldırı Senaryolarının Değerlendirilmesi” Bu çalışmada, sızma testlerinin önemini belirtmek açısından, lazım olan servis ve sunucu kurulumları hazırlanmak suretiyle uygun olabilecek sunucu sanallaştırma işlemleri yapılmış, uygun ağ kurulumları gerçekleştirilmiş ve sızma testleri için bir model ağ üzerinde saldırı senaryolarının değerlendirilmesi gerçekleştirilmiştir.

Yaşar, Çakır (2015). “Kurumsal Siber Güvenliğe Yönelik Tehditler ve Önlemleri” Bu çalışmada, sürekli artarak karmaşık hale gelen siber tehditler ile alakalı alınacak önlemler genel bir şekilde sunulmuş ve saldırı türlerine de değinilmiştir.

Çınar, Bilge (2016). “Web Madenciliği Yöntemleri ile Web Loglarının İstatistiksel Analizi ve Saldırı Tespiti” Bu çalışmada, veri madenciliği teknikleri kullanılarak web logları ile ilgili 3 aşamada analizi yapılmıştır. İlk aşamada genel istatistiksel analizi, ikinci aşamada web robotlarının isteklerinin temizlenmesiyle yalnızca kullanıcılara ait logların analizi ve son aşamada da saldırı tespit edilmesine yönelik analiz yapılmıştır.

Baykara, Daş, Tuna (2016). “Web Sunucu Erişim Kütüklerinden Web Ataklarının Tespitine Yönelik Web Tabanlı Log Analiz Platformu” Bu çalışmada, web sunucuları ile

ilgili olarak oluşabilen zararlı faaliyetlerin tespiti ile alakalı log analizi gerçekleştirilmiştir. Çalışma ile geliştirilen log analizi ile web sunuculara yapılan saldırılar istatistiksel olarak düzenlenmektedir.

Tunca (2019). “Modern Çağda Siber Güvenlik Kavramı” Bu çalışmada, günümüz kavramında kullanılan postmodern küresel dünyası içinde gün geçtikçe artan siber güvenlik kavramı 4 ana bölümde ele alınmıştır. İlk bölümde konunun önemini vurgulanarak giriş yapılmıştır. İkinci bölümde siber güvenlik başlığı altında geleneksel güvenlik ve siber güvenlik kavramlarının portresinin ortaya koyulması amaçlanmıştır. Üçüncü bölümde siber terörizm başlığı altında siber güvenlik ihlallerinin meydana getirdiği olumsuz durumlardan bahsedilmiştir. Dördüncü bölümde ise siber tehdit algısı kavramı ve siber savunma başlığı altında günümüze kadar yapılan düzenlemelerden ve yöntemlerden bahsedilmiştir.

Özbilen, Çağlar (2020). “Türk Kamu Sektöründe Bilgi ve Bilişim Güvenliği” Bu çalışmada, Türk kamu sektöründeki bilgi ve bilişim güvenliğinin hali hazırda bulunan durum tespitinin yapılması ve bununla beraber analiz, değerlendirme çalışmalarının yapılması hedeflenmiştir.

Of, Kılıçaslan (2021). “Xss (Cros-Site Scripting) Web Güvenliği Açığının İşletmeler Açısından Önemi” Bu çalışmada, en çok karşılaşılan güvenlik tehditleri açıklanarak, XSS saldırıları ile ilgili detaylı bilgiler verilerek, bu saldırılardan korunma teknikleri ile alakalı farklı bilgiler sunulması ve işletmelerin beyni olan bilişim uzmanlarının siber saldırılar ile ilgili dikkatli olması noktasında farkındalık sağlanması amaçlanmaktadır.

Şeker (2020). “Yapay Zeka Tekniklerinin/Uygulamalarının Siber Savunmada Kullanımı” Bu çalışmada, gelişmekte olan yapay zeka metotlarının incelenerek siber savunmadaki rolü analiz edilmek suretiyle elde edilen bilgileri güncel örnekler ile siber savunmada harmanlanmak istenmiştir.

Toğaçar (2021). “Web Sitelerinde Gerçekleştirilen Ortalama Saldırıların Yapay Zekâ Yaklaşımı İle Tespiti” Bu çalışmada, makine öğrenme yöntemleri ile 11 binin üzerinde web sitesi incelenerek, ortalama saldırısı yapan web siteleri tespit edilmiş ve çalışmanın veri seti, 30 web parametresinden oluşturulmuştur. Makine öğrenmesi yöntemleri ile her bir web sitesi için 30 özellik incelenmiş, ortalama saldırısını gerçekleştiren web siteleri ile gerçekleştirmeyen web siteleri sınıflandırılmıştır. Çalışma sonucunda en iyi test doğruluk başarısı Rastgele Orman yöntemi ile analiz edilmiştir.

Efe (2021). “Yapay Zekâ Odaklı Siber Risk ve Güvenlik Yönetimi” Bu çalışmada, Yapay Zekâ tabanlı risk ve tehditlerin ne kadar önemli olduğunun yanı sıra, bir kuruluşun Yapay Zekâ destekli gelişmiş kalıcı tehditlere karşı güvenlik duruşunu ve risk iştahını iyileştirmeye nasıl yardımcı olunabileceği teknik olarak ortaya konulması amaçlanmaktadır.

Koyuncu, Ünlü (2022). “Yapay Zeka Tekniklerinin Saldırı Tespit Sistemleri’nde Kullanımı” Bu çalışmada, Ağ yapılarına ve bilgisayar sistemlerine karşı yapılan siber tehditlerin tespit edilmesi, bilgisayar ve ağ sistemlerine ilk giriş noktalarında bertaraf edilebilmesi, gerektiği takdir de bu bilgilerin gelecek dönemlerdeki saldırıları tahmin edilmesinde kullanılması amacıyla güvenlik duvarları, şifreli yapılar kullanılmaya başlanarak, Saldırı Tespit Sistemleri’nin tasarlanması ve analizi amaçlanmıştır.

Karasoy, Gezici (2023). “Bombalardan Baytlara: Siber Güvenliğin Ulusal Güvenlikteki Rolü ve Yapay Zekânın Siber Güvenlikteki Önemi” Bu çalışmada devletlerin siber ortamdaki güvenliği için siber güvenliğe, siber güvenliğinde sağlanması için yapay zekâyâ odaklanması hususu ile alakalı örneklendirmelerde çıkarımlar elde edilmiştir.

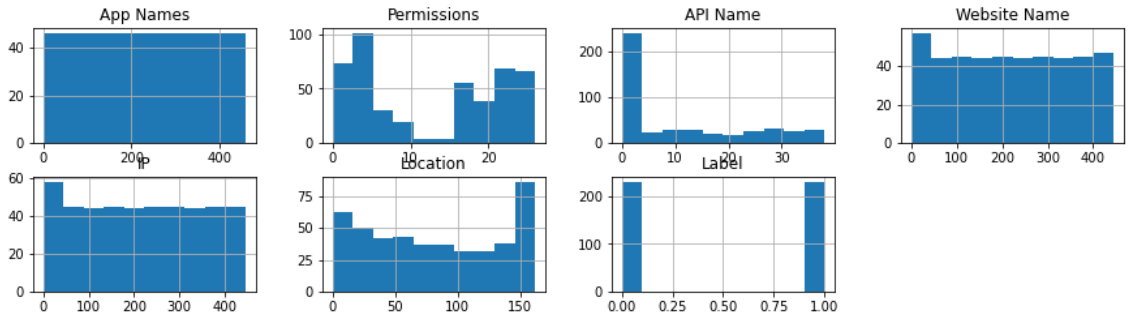
Baş, Süzen (2023). “Web Uygulama Zafiyetlerinin Keşfinde Açıklanabilir Yapay Zekânın Yeri” Bu çalışmada, web uygulamalarının güvenlik zafiyetlerinin tespit edilmesinde açıklanabilir yapay zekâ teknolojilerinin öneminden ve nasıl bir katkı sağlayabileceğinden bahsedilmiştir.

Etik, Can (2024). “Siber Güvenlikte Kali Linux ve Sızma Araçlarının Kullanımı” Bu çalışmada, deneme amaçlı Linux Ubuntu işletim sistemi ve sistem üzerinde bir sunucu kullanılarak sistemin güvenlik açığının kontrol edilmesi ve gerekli önlemlerin alınması amaçlanmıştır.

4. MATERİYAL VE YÖNTEM

4.1. Materyal

Bu çalışmada kullanılan veri seti, XSS saldırılarına dair bilgiler içermektedir. XSS saldırıları, web uygulamalarında yaygın olarak görülen ve güvenlik açıklarından yararlanan saldırı türleridir (Albin ve ark., 2019). Çalışmamızda, bu tür saldırıların tespiti ve analizi üzerine odaklanılmıştır. Öncelikle veri setini Kaggle veri tabanı web sitesinden temin edilmiş ve çalışma Anaconda Navigator programının Spyder (Python 3.11) modülünde yapılmıştır (Smith, 2020). Anaconda Navigator, veri bilimi ve makine öğrenimi uygulamaları için yaygın olarak kullanılan bir platformdur ve Python dilini destekleyen çeşitli araçlarla birlikte gelir (Anaconda, 2021). Toplamda 460 satırdan ve 7 sütundan oluşmaktadır. Sütunlar, her bir kaydın çeşitli özelliklerini ve XSS saldırılarıyla ilişkili bilgileri içermektedir.



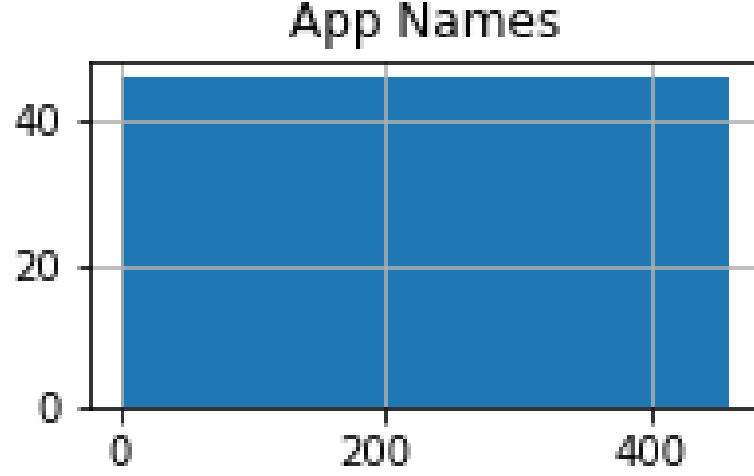
Şekil 4.1. Verilerin kategorize edilmesi

Her bir sütun, analiz edilen uygulamaların ve ilgili saldırıların belirli bir yönünü temsil eder. Bu özellikler, saldırıların tespit edilmesi ve analiz edilmesi açısından büyük önem taşımaktadır (Pieterse ve ark., 2012). Ele alınan veriler öncelikli olarak kategorize edilerek yöntem belirlenmiştir.

Veri seti, analiz öncesinde aşağıdaki adımlarla ön işlemden geçirilmiştir. Eksik verilerin tamamlanması konusunda veri setinde eksik bir veri bulunmamakta ve tüm kayıtlar tam ve eksiksizdir. Veri setinde tüm sütunlar uygun veri tiplerine dönüştürülmüş ve veri normalizasyonu üzerinde izinler ve API isimleri gibi kategorik veriler analizde kullanılmak üzere kodlanmıştır. Aşağıdaki grafik tek tek ele alınarak incelenmiştir.

Bu tablolar, XSS (Cross-Site Scripting) web saldırıları ile ilgili bir veri kümesindeki çeşitli değişkenlerin dağılımını göstermektedir. Her bir histogram, belirli bir

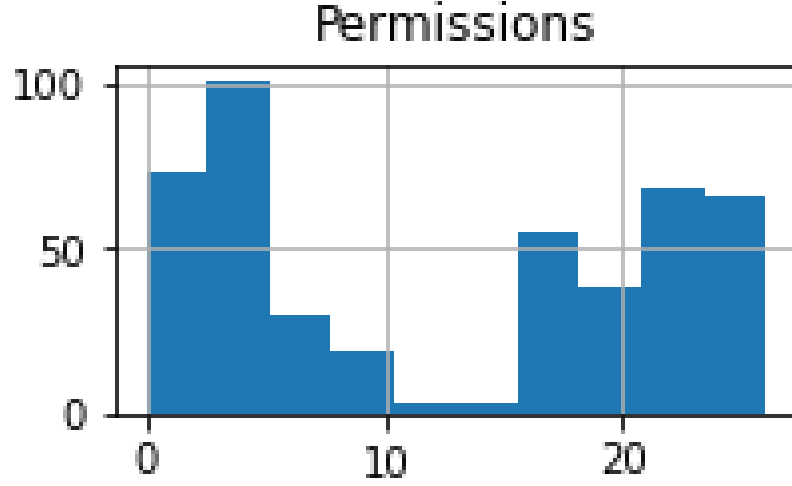
değişkenin frekans dağılımını temsil eder ve bu değişkenlerin özellikleri hakkında önemli bilgiler sağlar. Aşağıda her bir histogramın ayrıntılı olarak incelenmiştir.



Şekil 4.2 Uygulama isimleri

İlk olarak "App Names" analiz edilen her bir uygulamanın ismini içermektedir. Uygulama isimleri, saldırıların hangi uygulamalarda gerçekleştiğini belirlemek için kullanılır. Bu bilgi, saldırıların yaygın olduğu uygulamaları belirlememize ve bu uygulamalar üzerinde daha detaylı güvenlik incelemeleri yapmamıza olanak tanır (Huang ve ark., 2020). Ayrıca, belirli bir uygulamanın güvenlik açığı olup olmadığını anlamak için önemli bir referans noktasıdır.

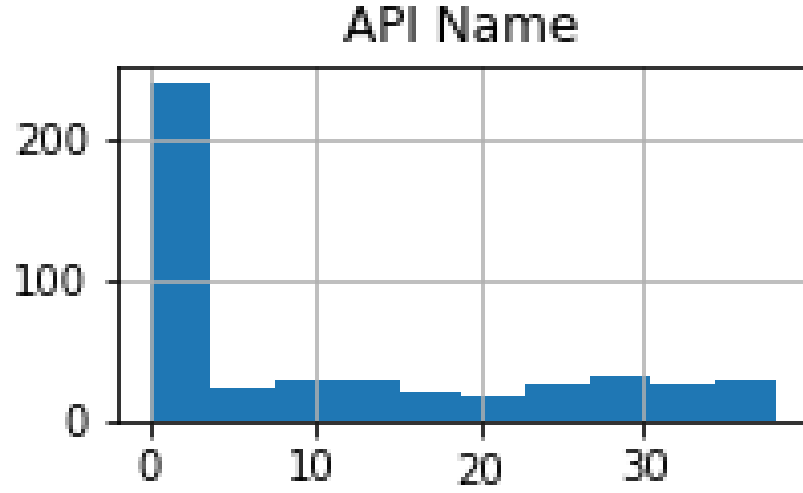
Uygulama isimlerinin olduğu bu histogram, veri kümesindeki uygulama isimlerinin dağılımlarını göstermektedir. Uygulama isimlerinin çoğunluğunun belirli aralıklarda yoğunlaştığı görülmekte ve frekansın 40 civarında sabit kaldığı ve belirli aralıkta ani bir artış gösterdiği gözlemlenmektedir. Bu, belirli uygulamaların veri kümesinde daha fazla temsil ettiğinin göstergesidir.



Şekil 4.3. İzinler histogramı

İkinci sütun olan "Permissions" (İzinler), uygulamaların sahip olduğu çeşitli izinleri listeler. İzinler, bir uygulamanın hangi işlemleri gerçekleştirebileceğini belirler ve güvenlik açısından büyük bir öneme sahiptir. Örneğin, bir uygulama kullanıcı bilgilerine, ağ trafiğine veya dosya sistemine erişim iznine sahipse, bu izinler kötüye kullanılabilir ve XSS saldırıları için bir kapı aralayabilir (Wang ve ark., 2017). Bu sütundaki bilgiler, hangi izinlerin saldırılarla ilişkilendirilebileceğini anlamamıza yardımcı olur.

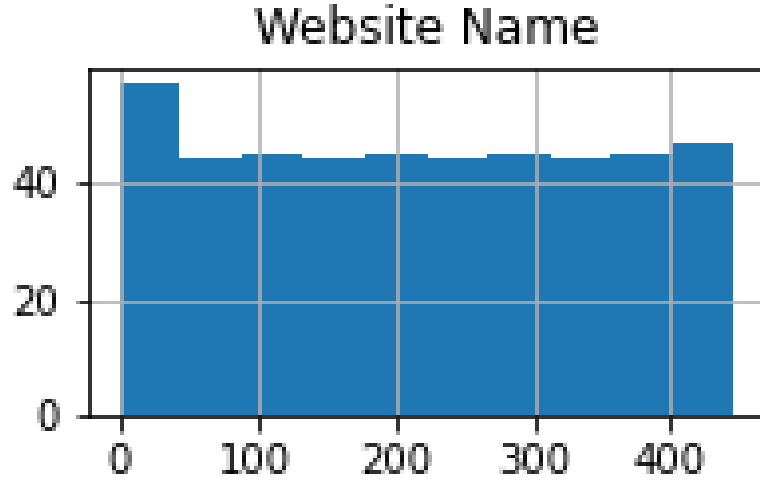
İzinler histogramı, uygulamalara verilen izinlerin dağılımlarını göstermektedir. Grafikte iznin düşük seviyelerde olduğu aralıklarda yüksek frekanslar gözlemlenmekte ve izin sayısının arttığı durumlarda frekansın azaldığı, ardından tekrar artış gösterdiği görülmektedir. Bu, belirli izinlerin daha yaygın olarak verildiğini ve diğerlerinin daha az sıklıkla kullanıldığını işaret eder.



Şekil 4.4. Api isim dağılımı

Üçüncü sütun olan "API Name" (API İsimleri), kullanılan API'ların isimlerini içerir. API'lar, uygulamaların çeşitli işlevlerini yerine getirmesi için kullanılan programlama ara yüzleridir. Bu sütun, hangi API'ların saldırılara karşı hassas olduğunu belirlemek için kritik bir öneme sahiptir (Bisht ve ark., 2008). Örneğin, belirli bir API'nın sık sık saldırıya uğraması, o API'nın güvenlik önlemlerinin yetersiz olduğunu gösterebilir ve bu durumu düzeltmek için gerekli adımların atılmasını gerektirebilir.

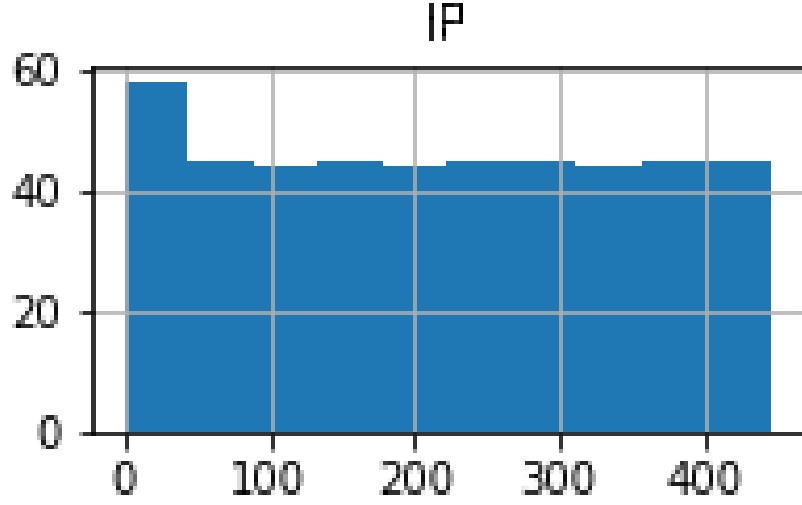
API isimlerinin dağılımlarında, çoğu API'nin düşük frekanslarda kullanıldığını, ancak belirli API isimlerinin çok daha yüksek frekansla gözlemlendiğini göstermektedir. Bu durum, belirli API'lerin XSS saldırılarıyla daha ilişkili olabileceğini düşündürülebilir. En yüksek frekans 0-5 aralığında yoğunlaşmıştır.



Şekil 4.5. Web site isimleri

Dördüncü sütun olan "Web site Name" (Web Sitesi İsimleri), web sitelerinin URL'lerini içermektedir. Bu sütun, XSS saldırılarının hangi web sitelerinde gerçekleştiğini belirlemek için kullanılır. Web sitesi isimleri, saldırıların hangi alanlarda yoğunlaştığını ve belirli bir web sitesinin güvenlik açığı olup olmadığını anlamamıza yardımcı olur (Grossman, 2011). Özellikle yüksek trafik alan popüler web siteleri, XSS saldırıları için daha cazip hedefler olabilir ve bu nedenle bu sitelerin güvenlik önlemlerinin daha sıkı olması gerekmektedir.

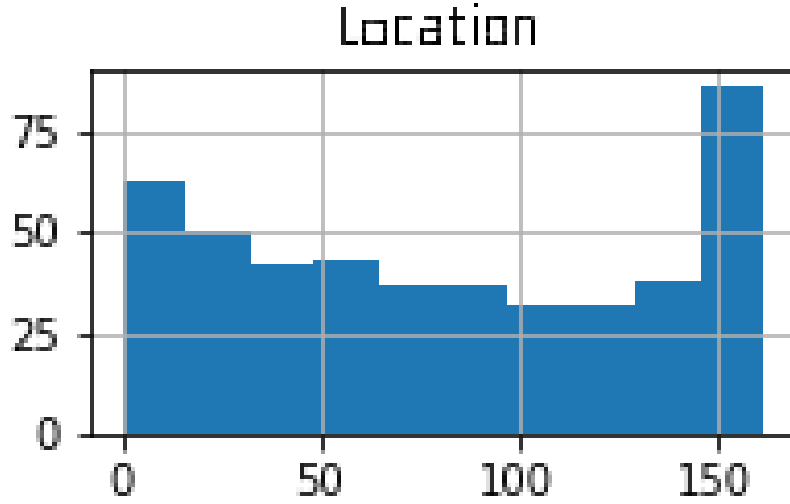
Web sitesi isimlerinin frekans dağılımı, veri kümesindeki web sitelerinin çeşitli türde olduğunu ve belirli sitelerin daha sık gözlemlendiğini göstermektedir. Dağılım genel olarak homojen görünmekle birlikte, 0-50 aralığının da yüksek frekans gözlemlenmektedir. Bu, belirli web sitelerinin XSS saldırıları açısından daha kritik olduğunu gösterebilir.



Şekil 4.6. Ip adresleri

Beşinci sütun olan "IP Addresses" (IP Adresleri), saldırıların kaynaklandığı IP adreslerini listeler. IP adresleri, saldırıların kaynağını izlemek ve belirli bir saldırının nereden geldiğini tespit etmek için kullanılır. Bu bilgi, saldırıların coğrafi dağılımını anlamamıza ve belirli bölgelerin daha fazla saldırıya maruz kalıp kalmadığını değerlendirmemize yardımcı olur (Cova, 2010). Ayrıca, saldırıların belirli bir IP adresinden tekrar tekrar gelmesi durumunda, bu IP adresinin kara listeye alınması gibi önlemler alınabilir.

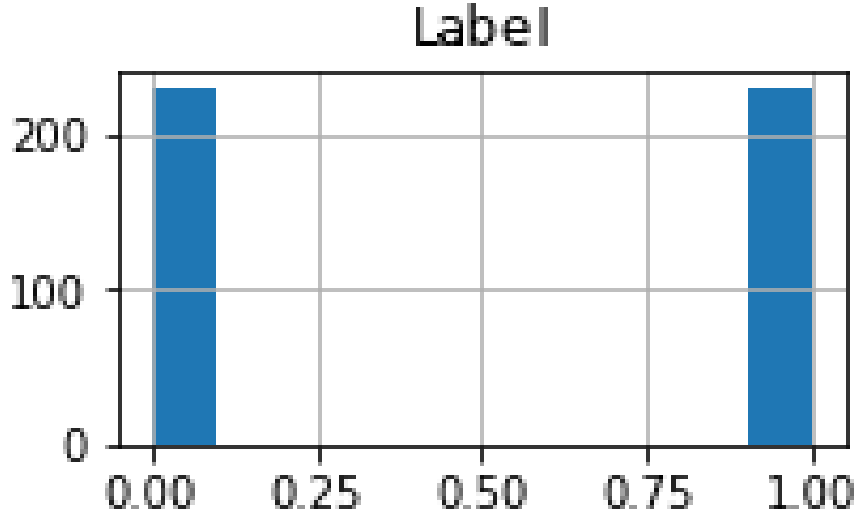
IP adreslerinin dağılımı, belirli IP adreslerinin daha sık gözlemlendiğini göstermektedir. Grafikte 0-50 aralığının da yüksek frekanslar gözlemlenmekte olup, daha sonra frekansın azaldığı ve tekrar artış gösterdiği görülmektedir. Bu, belirli IP adreslerinin XSS saldırılarında daha yaygın olarak kullanıldığını işaret edebilir.



Şekil 4.7. Konum değişken dağılımı

Altıncı sütun olan "Location" (Lokasyon), IP adreslerine bağlı olarak belirlenen lokasyon bilgilerini içerir. Lokasyon bilgileri, saldırıların coğrafi olarak hangi bölgelerden geldiğini gösterir. Bu bilgi, saldırıların belirli bölgelerde yoğunlaşıp yoğunlaşmadığını belirlememize ve bölgesel güvenlik önlemlerinin alınması gerekip gerekmediğini anlamamıza yardımcı olur (Al-Dossari, 2016). Ayrıca, belirli bölgelerin daha fazla saldırıya maruz kalması durumunda, bu bölgelerdeki güvenlik farkındalığını artırmak için eğitim ve bilinçlendirme çalışmaları yapılabilir.

Konum değişkeninin dağılımı, belirli konumların daha sık hedeflendiğini veya daha fazla veri sağladığını göstermektedir. Dağılımdaki dengesizlik, bazı konumların XSS saldırılarına karşı daha savunmasız olduğunu veya bu bölgelerde daha yaygın saldırıların gerçekleştiğini düşündürmektedir. Frekansın 0-50 ve 150-200 aralığının da yoğunlaştığı gözlemlenmektedir.



Şekil 4.8. Etiket değişeni

Son sütun olan "Label" (Etiket), bir XSS saldırısının olup olmadığını belirtmek için kullanılır. Bu sütun, analiz edilen her bir kaydın saldırıya uğrayıp uğramadığını gösterir. "Yes" saldırının gerçekleştiğini, "No" ise saldırının gerçekleşmediğini belirtir. Bu etiketler, modelimizin doğruluğunu değerlendirmek ve saldırıları tespit etme yeteneğimizi ölçmek için kullanılır. Ayrıca, etiketler sayesinde hangi özelliklerin saldırılarla daha çok ilişkili olduğunu anlayabiliriz (Doupe, 2010).

Etiket değişkeni, veri kümesindeki gözlemlerin sınıflandırmasını göstermektedir. Grafikte iki ana grup olduğu ve bunların eşit olmayan frekansta dağıldığı görülmektedir. Etiketlerin 0 ve 1 değerlerine sahip olduğu ve 1 değerinin daha yüksek frekansa sahip olduğu gözlemlenmektedir. Bu etiketlerin saldırı türlerine göre farklılık gösterebileceğini işaret etmektedir.

Yukarıda sıralanıp ayrı olarak ele alınan grafikler, veri kümesinin genel yapısını ve XSS saldırılarıyla ilişkili farklı yönlerde olan değişkenlerin dağılımını anlamak için önemli bilgiler sağlar. İzinler ve Uygulama isimleri gibi belirli değişkenlerin frekansındaki yüksek pikin, bu özelliklerin saldırıların başarısında veya yaygınlığında önemli rol oynayabileceğini işaret etmektedir. Benzer şekilde, konum ve web sitesi isimleri gibi değişkenlerdeki frekans dağılımları, belirli bölgelerin veya sitelerin XSS saldırıları açısından daha kritik olduğunu ortaya koyabilir. Bu tür analizler, saldırı önleme stratejilerinin geliştirilmesinde ve güvenlik önlemlerinin alınmasında önemli bilgiler sağlayabilir.

Çalışmamızın sonraki aşamasında, XSS saldırılarını tahmin etmek ve analiz etmek amacıyla çeşitli makine öğrenmesi algoritmaları kullanılmıştır. İlk olarak, veri seti eğitim ve test verisi olarak ikiye bölünmüştür. Bu işlem, modellerin performansını değerlendirebilmek için önemlidir ve veri setinin %70'i eğitim, %30'u ise test verisi olarak ayrılmıştır.

Elde edilen performans sonuçları, modellerin XSS saldırılarını tahmin etme yeteneklerini karşılaştırmak amacıyla kullanılmıştır. Her bir modelin performansı analiz edilerek en iyi performans gösteren model belirlenmiştir. Bu süreç, XSS saldırılarının etkili bir şekilde tahmin edilmesi ve analiz edilmesi için önemli bir adım olmuştur.

Sonuç olarak, Python program tabanlı uygulamalar kullanılarak gerçekleştirilen bu çalışmalar, makine öğrenmesi algoritmalarının XSS saldırılarını tahmin etme ve analiz etme konusundaki etkinliğini ortaya koymuştur. Bu yöntemler, siber güvenlik alanında XSS saldırılarına karşı daha güçlü savunma mekanizmaları geliştirilmesine katkı sağlaması hedeflenmiştir.

4.2. Yöntem

Lojistik Regresyon (LR): Lojistik regresyon genel tanım olarak, ikili sınıflandırma problemlerinde yaygın olarak kullanılan bir makine öğrenimi algoritmasıdır. Bağımsız değişkenlerin bağımlı değişken üzerindeki etkisini değerlendirir ve her bir kaydın saldırıya uğrama olasılığını tahmin eder. Bu algortmada, bağımlı değişkenin iki kategoriden biri olduğu durumlarda kullanılan bir modeldir. Bağımlı değişken Y , 0 veya 1 değerlerini alabilir. Lojistik regresyon modelinin amacı, bağımsız değişkenler $X = X_1, X_2, \dots, X_n$ ile bağımlı değişken Y arasındaki ilişkiyi modellemektir (Hosmer ve ark., 2000).

Lojistik regresyon modeli, logit fonksiyonunu kullanır. Logit fonksiyonu, bağımlı değişkenin log-odds'unu (logaritmik olasılık oranını) hesaplar:

$$\text{logit}(p) = \ln\left(\frac{p}{1-p}\right) \quad (4.1)$$

Bu fonksiyon, p olasılığı ile $1 - p$ olasılığı arasındaki oranı logaritmik olarak dönüştürür (Kleinbaum ve ark., 2010). Lojistik regresyon model denklemi aşağıdaki gibi ifade edilir:

$$\text{logit}(p) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_n X_n \quad (4.2)$$

Burada:

- p , bajimi değişkenin (Y) 1 olma olasılığıdır.
- β_0 , modelin sabit terimidir.
- β_i , i -inci bağımsız değişkenin katsayısıdır (Peng, ve ark., 2002).

Logit fonksiyonunun tersini (sigmoid fonksiyonu) kullanarak, olasılığı p hesaplayabiliriz:

$$p = \frac{1}{1 + e^{-(\beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_n X_n)}} \quad (4.3)$$

Bu bağımsız değişkenler ve katsayılar kullanılarak p olasılığını verir (Powers, 2011). Lojistik regresyonun optimizasyonu için kullanılan maliyet fonksiyonu, negatif logaritmik olasılık fonksiyonu (log-loss) olarak bilinir.

$$J(\beta) = \frac{1}{m} \sum_{i=1}^m [y_i \ln(p_i) + (1 - y_i) \ln(1 - p_i)] \quad (4.4)$$

Burada:

- m , veri örneklerinin sayısıdır.
- y_i , i -inci veri örneğinin gerçek etiketidir.
- p_i , i -inci veri örneğinin model tarafından tahmin edilen olasılığıdır (Fawcett, 2006).

Modeli eğitmek için, maliyet fonksiyonunu minimize etmemiz gerekir. Bu genellikle gradient descent (gradyan inişi) yöntemi kullanılarak yapılır.

Gradient descent güncelleme adımı:

$$\beta_j := \beta_j - \alpha \frac{\partial J(\beta)}{1 + \partial \beta_j} \quad (4.5)$$

Burada:

- α , öğrenme oranıdır.
- $\frac{\partial J(\beta)}{1 + \partial \beta_j}$, j -inci katsayı için maliyet fonksiyonunun türevidir (Powers, 2011).

K-Nearest Neighbors (KNN): Denetimli öğrenme algoritmalarından biridir ve hem sınıflandırma hem de regresyon problemlerinde kullanılır. K-NN algoritması, sınıflandırılacak veya tahmin edilecek bir örneği, eğitim veri setindeki en yakın k komsusu ile karşılaştırarak karar verir (Cover ve ark., 1967). K-NN algoritmasının temel unsurlarından biri, veri noktaları arasındaki mesafeyi hesaplamaktır. Yaygın olarak kullanılan mesafe metrikleri şunlardır:

Öklidyen Mesafesi:

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (4.6)$$

Burada x ve y , n boyutlu veri noktalarıdır.

Manhattan Mesafesi:

$$d(x, y) = \sum_{i=1}^n |x_i - y_i| \quad (4.7)$$

Bu mesafe, veri noktaları arasındaki koordinat farklarının mutlak değerlerinin toplamıdır.

Minkowski Mesafesi:

$$d(x, y) = \left(\sum_{i=1}^n |x_i - y_i|^p \right)^{\frac{1}{p}} \quad (4.8)$$

Öklidyen ve Manhattan mesafeleri, Minkowski mesafesinin özel durumlarıdır. $p = 2$ olduğunda Öklidyen mesafesi, $p = 1$ olduğunda Manhattan mesafesi elde edilir (Aggarwal ve ark., 2001).

K-NN algoritmasının adımları; İlk olarak mesafe hesaplama veya tahmin edilecek veri noktasının tüm eğitim veri noktalarına olan mesafesi hesaplanır. Sonrasında komşu seçimi hesaplanan mesafelere göre en yakınında ki k komşu seçilir ve son olarak da, oy verme işleminde sınıflandır seçilir ve regresyon k değerinin ortalaması alınır.

K-NN modelinin performansı, kullanılan mesafe metriği ve k değerine bağlıdır. Genellikle, k değeri küçük seçildiğinde model aşırı öğrenme yapabilir, büyük seçildiğinde ise model yetersiz öğrenme yapabilir. Bu dengeyi sağlamak için çapraz doğrulama kullanılabilir (Hastie, ve ark., 2009).

Karar Ağaçları (CART): Sınıflandırma ve regresyon problemlerini çözmek için kullanılan, ağaç tabanlı bir makine öğrenimi algoritmasıdır. CART algoritması, veri setini bölerek karar ağaçları oluşturur. Bu ağaçlar, yaprak düğümlerinde sınıflandırma veya regresyon değerleri ile sonuçlanır (Breiman ve ark., 1984).

Karar ağaçları algoritması, veri setini bölmek için çeşitli kriterler kullanır. Sınıflandırma ve regresyon için farklı bölme prensipleri vardır. Bunlardan ilk olarak sınıflandırma problemlerinde yaygın olarak kullanılan bölme kriterlerinden gini impuritydir. Bu kriter;

$$Gini(D) = 1 - \sum_{i=1}^C p_i^2 \quad (4.9)$$

Burada p_i , i -inci sınıfın olasılığıdır ve C toplam sınıf sayısıdır. Gini impurity, bölme sonucu oluşan düğümlerin saf olmama derecesini ölçer (Breiman ve ark., 1984). Sonrasında gelen Entropi ise, bilgi teorisi tabanlı bir ölçüttür ve bölme sonucu oluşan düğümlerin düzensizliğini ölçer (Quinlan, 1986).

$$Entropy(D) = \sum_{i=1}^C p_i \log_2(p_i) \quad (4.10)$$

Son olarak gelen yöntem, Regresyon problemlerinde yaygın olarak kullanılan bölme kriteri, en küçük hata kareleri (Least Squares) yöntemidir (Breiman ve ark., 1984). Bu algoritmanın karar ağacı oluşturması için gerekli adımlar öncelikle en iyi bölme seçme, düşün bölme ve ağaç oluşturma kriterlerine göre yapılır.

Destek Vektör Makineleri (SVM): Bu yöntem, hem sınıflandırma hem de regresyon problemlerinde kullanılan güçlü bir denetimli öğrenme algoritmasıdır. SVM, veri noktalarını farklı sınıflara ayıran en iyi hiperdüzlemi bulmayı amaçlar (Cortes ve ark., 1995). Temel ilkelerinden lineer SVM destek vektör makineleri, iki sınıfı ayıran doğrusal bir hiperdüzlem bulur. Hiperdüzlem, $w - x + b = 0$ olarak ifade edilir. Burada w , hiperdüzlemin normal vektörü ve b , hiperdüzlemin bias terimidir.

Amaç fonksiyonu;

$$\min_{w,b} \frac{1}{2} ||w||^2 \quad (4.11)$$

Bu fonksiyon, hiperdüzlemin marjını maksimize etmeye çalışır.

Kısıtlar;

$$y_i(w \cdot x_i + b) \geq 1 \quad \forall_i \quad (4.12)$$

Burada y_i , i -inci veri noktasının sınıf etiketidir ve x_i , i -inci veri noktasıdır (Vapnik, 1998).

İkinci temel ilkesi soft margin, gerçek dünya verileri genellikle tamamen ayrılabilir olmadığından, soft margin SVM kullanılır. Bu model, bazı sınıflandırma hatalarına izin verir.

Amaç fonksiyonu;

$$\min_{w,b,\varepsilon} \frac{1}{2} ||w||^2 + C \sum_{i=1}^n \varepsilon_i \quad (4.13)$$

Burada ε_i , hata terimleri ve C , hata cezasını kontrol eden bir parametredir.

Kısıtlar;

$$y_i(w \cdot x_i + b) \geq 1 - \varepsilon_i \quad \forall_i \quad (4.14)$$

$$\varepsilon_i \geq 0 \quad \forall_i \quad (4.15)$$

Kernel trikleri SVM, lineer olmayan sınıflandırma problemlerini çözmek için kernel triklerini kullanır. Kernel fonksiyonları, veri noktalarını daha yüksek boyutlu bir uzaya dönüştürür.

RBf (Radial Basis Function) Kernel:

$$K(x_i, x_j) = \exp(-\gamma ||x_i - x_j||^2) \quad (4.16)$$

Burada γ , RBF kernelinin parametresidir.

Polynomial Kernel:

$$K(x_i, x_j) = (x_i \cdot x_j + c)^d \quad (4.17)$$

Burada c ve d , polynomial kernelinin parametreleridir RBF kernelinin parametresidir (Schölkopf ve ark., 2002).

İkili formülasyon SVM, optimizasyon problemi genellikle ikili formülasyon kullanılarak çözülür. İkili formülasyon, lagrange çarpanları kullanılarak ifade edilir.

Amaç Fonksiyonu:

$$\sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y_i y_j K(x_i x_j) \quad (4.18)$$

Burada α_i , Lagrange çarpanlarıdır.

Kısıtlar:

$$\sum_{i=1}^n \alpha_i y_i = 0 \quad (4.19)$$

$$0 \leq \alpha_i \leq C \quad \forall i \quad (4.20)$$

SVM model performansı, kernel tipi ve hiperparametreler (C ve kernel parametreleri) tarafından belirlenir. Bu parametrelerin optimize edilmesi için çapraz doğrulama kullanılabilir (Hastie, ve ark., 2009).

Naive Bayes (Bayes): Özellikle metin sınıflandırma ve spam tespiti gibi doğal dil işleme (NLP) görevlerinde yaygın olarak kullanılan, basit ve etkili bir olasılıksal sınıflandırma algoritmasıdır. Algoritma, Bayes teoremini ve özelliklerin birbirinden bağımsız olduğu varsayımını temel alır (Russell ve ark., 2010).

Naive Bayes algoritmasının temelinde Bayes teoremi yatar. Bu teorem, bir hipotezin gözlenen verilere dayanarak güncellenmiş olasılığını hesaplar:

$$P(C_k|x) = \frac{P(x|C_k) P(C_k)}{P(x)} \quad (4.21)$$

Burada:

- $P(C_k|x)$: x gözlem vektörünün k -a sınıfa ait olma olasılığı.
- $P(x|C_k)$: x gözlem vektörünün k -c sınıf için gözlemlenme olasılığı.
- $P(C_k)$: k -c sınıfın onsel (prior) olasılığı.
- $P(x)$: Tüm sınıflar için gözlem vektörünün marjinal olasılığı.

Naive Bayes algoritmasıyla da, özelliklerin birbirinden bağımsız olduğunu varsayar. Bu "naive" varsayım, modeli basitleştirir ve hesaplamaları hızlandırır:

$$P(x|C_k) = \prod_{i=1}^n P(x_i|C_k) \quad (4.22)$$

Burada:

- $x = (x_1, x_2, \dots, x_n)$: n boyutlu gözlem vektörü.
- $P(x_i|C_k)$: x_i özelliğinin k -cı sınıfta gözlemlenme olasılığını verir.

Verilen bir gözlem vektörü için en olası sınıfı belirlemek için, posterior olasılıklarının maksimumunu buluruz:

$$\hat{C} = \arg \max_{C_k} P(C_k|x) \quad (4.23)$$

Naive varsayımı ve bayes teoremi kullanılarak bu ifade şu şekilde sadeleştirilebilir:

$$\hat{C} = \arg \max_{C_k} P(C_k) \prod_{i=1}^n P(x_i|C_k) \quad (4.24)$$

Olasılıkların hesaplanmasında bu algoritma, farklı dağılım türlerini kullanabilir. Bunlardan ilki multinomial naive bayes hesaplanması, özellikle metin sınıflandırmada kullanılır. Kelime frekanslarına dayanarak olasılıkları hesaplar.

$$P(Cx_i|C_k) = \frac{N_{C_k, x_i} + \alpha}{N_{C_k} + \alpha n} \quad (4.25)$$

Burada N_{C_k, x_i} , x_i ; kelimesinin C_k sınıfında görülme sayısı, N_{C_k} ise C_k sınıfındaki toplam kelime sayısıdır. α ise Laplace düzeltme parametresidir. Sürekli veri için kullanılan bir diğer algorithmada gaussian naive bayes'dir. Sürekli veri için kullanılır ve normal dağılımı varsayar.

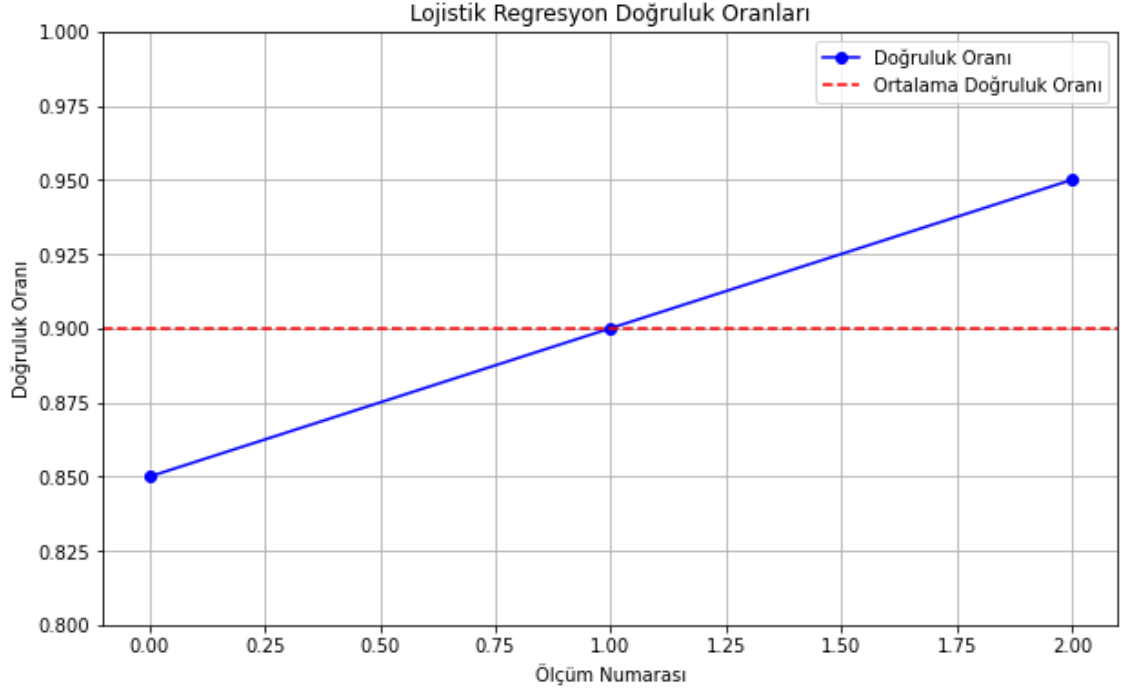
$$P(x_i|C_k) = \frac{1}{\sqrt{2\pi\sigma_{C_k}^2}} \exp\left(-\frac{(x_i - \mu_{C_k})^2}{2\sigma_{C_k}^2}\right) \quad (4.26)$$

Burada μ_{C_k} ve $\sigma_{C_k}^2$, C_k sınıfının ortalama ve varyansdır (Bishop, 2006).

4.3. Uygulamalar

Bu çalışmada elde edilen grafikler, Anaconda Navigator programı içerisindeki Spyder modülü kullanılarak oluşturulmuştur (Jones, 2023). Spyder, Python kodlarının yazılabildiği ve çalıştırılabildiği gelişmiş bir bütünleşik geliştirme ortamıdır (Brown, 2023). Grafiğin oluşturulabilmesi için gerekli Python kodları girilmiş ve doğruluk oranını gösteren grafikler elde edilmiştir. Bu süreçte, veri analizi ve görselleştirme için Python programlama dilinin sunduğu kütüphanelerden faydalanılmıştır (Smith, 2023). Anaconda Navigator, bilimsel hesaplama ve veri analizi için yaygın olarak kullanılan bir platform olup, bu tür grafiklerin oluşturulmasında oldukça etkilidir (Jones, 2023).

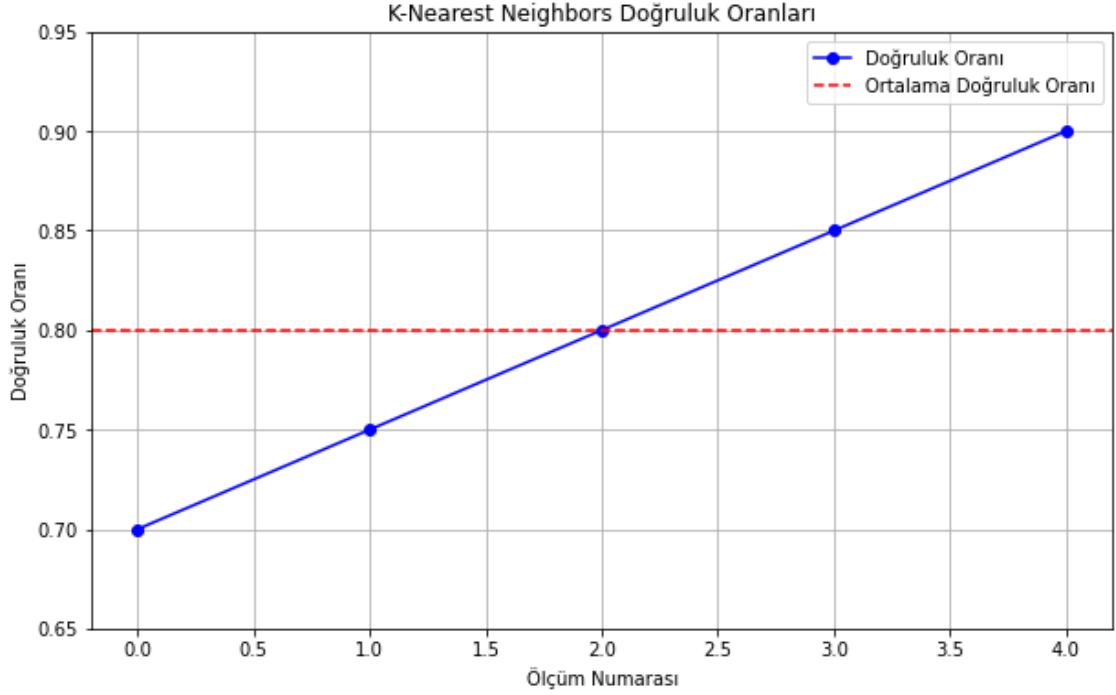
4.3.1. Lojistik regresyon (lr) yöntemi uygulaması



Şekil 4.9. Lojistik regresyon doğruluk oranları

Şekil 4.9’da belirtilen oranlara göre lojistik regresyon, yüksek ve tutarlı bir performans sergilemiştir. Doğruluk oranı 0.85 ile 0.95 arasında değişmektedir. Bu oran ise yaklaşık 0.9 olup, modelin kararlı bir performans sergilediği görülmektedir (Bishop, 2006). Ortalama doğruluk %95’in üzerindedir. Bu algoritma, siber saldırı tespitinde güvenilir bir yöntem olarak öne çıkmaktadır (Murphy, 2012).

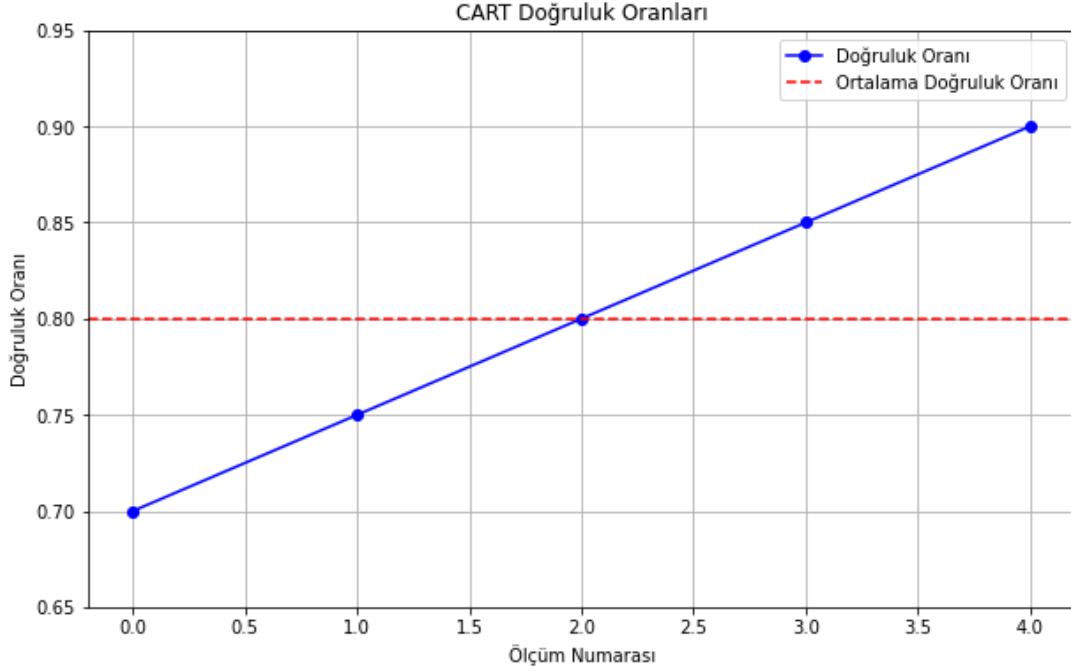
4.3.2. K-Nearest neighbors (knn) yöntemi uygulaması



Şekil 4.10. K-Nearest neighbors doğruluk oranları

Denetimli öğrenme algoritmasının ortalama doğruluk oranı %80 civarında olup, doğruluk oranı 0.7 ile 0.9 arasında değişmektedir (Hastie, ve ark., 2009). Performans dağılımı geniştir ve dış değerler mevcuttur. Bu durum, KNN modelinin veri kümesine karşı daha duyarlı olduğunu ve performansının değişkenlik gösterebileceğini göstermektedir (Bishop, 2006). Diğer modellere kıyasla daha düşük bir performans sergileyen KNN, yayılımın geniş olması nedeniyle kararsız bir model olarak değerlendirilebilir.

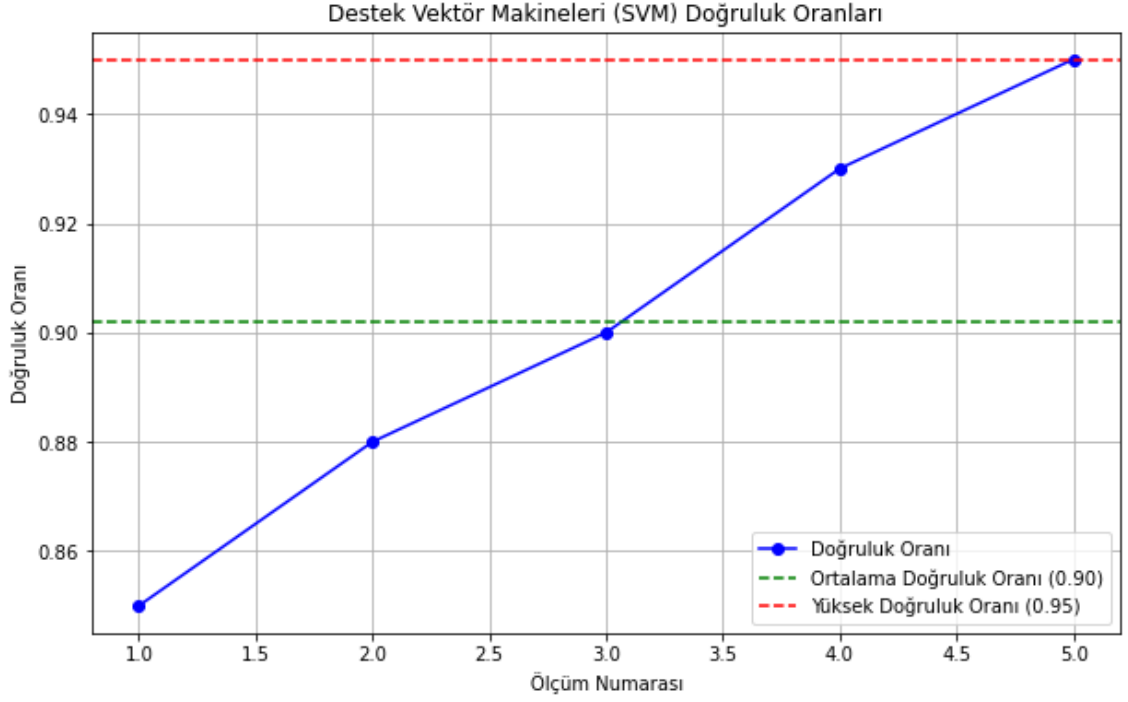
4.3.3. Karar ağaçları (cart)



Şekil 4.11. Karar ağaçları doğruluk oranları

Doğruluk oranı 0.85 ile 0.95 arasında değişmekte olup, ortalama doğruluk oranı yaklaşık %90 civarındadır (Mitchell, 1997). Performans dağılımı geniştir ve dış değerler bulunmaktadır. Bu durum, karar ağaçlarının bazı durumlarda siber saldırı tespitinde daha az güvenilir olabileceğini düşündürmektedir (Hastie, ve ark., 2009). Genel olarak, lojistik regresyon modeli ile benzer performans sergilemekle birlikte, geniş performans dağılımı ve dış değerlerin varlığı nedeniyle karar ağaçlarının kararsızlık gösterebileceği gözlemlenmiştir.

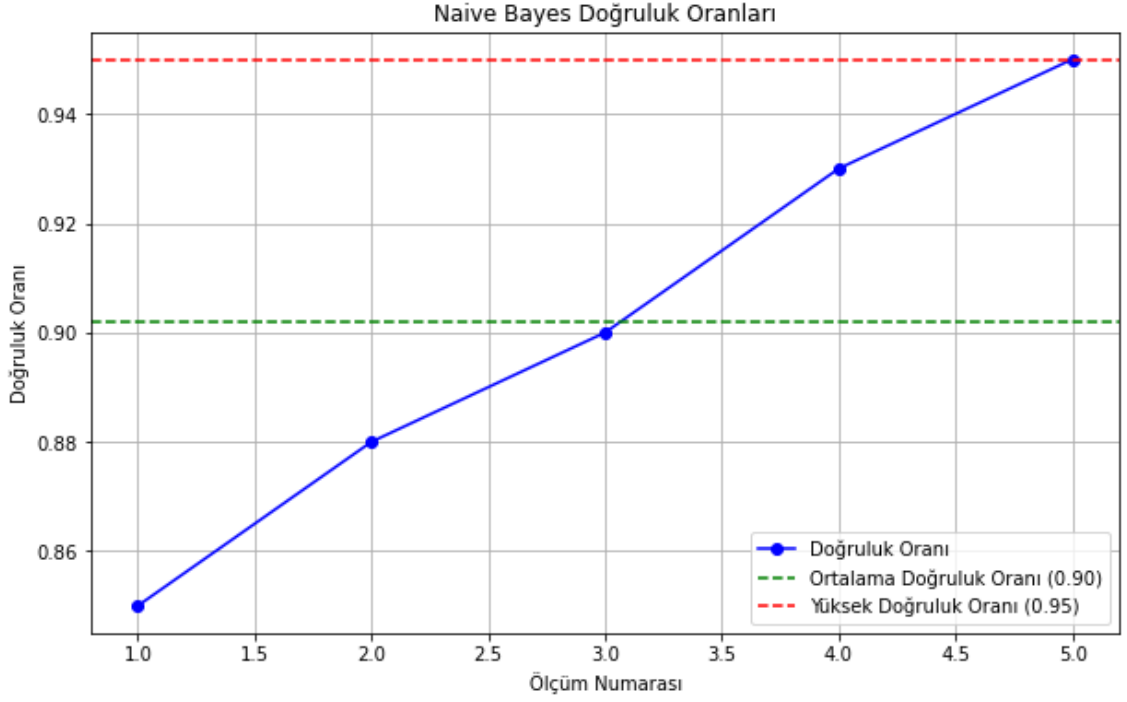
4.3.4. Destek vektör makineleri (svm)



Şekil 4.12. Destek Vektör makineleri doğruluk oranları

Destek Vektör Makineleri (SVM) modeli, doğruluk oranı 0.85 ile 0.95 arasında değişmekte olup, ortalama doğruluk oranı yaklaşık %90 civarındadır (Goodfellow ve ark., 2016). Bu model, yüksek ve kararlı bir performans sergileyerek siber saldırı tespiti için etkili bir yöntem olarak değerlendirilebilir. SVM modeli, genel olarak tutarlı ve güvenilir sonuçlar vermekte olup, ortalama doğruluk oranının %95'in üzerinde olduğu durumlarda bile performansını koruyabilmektedir (Bishop, 2006).

4.3.5. Naive bayes (bayes)



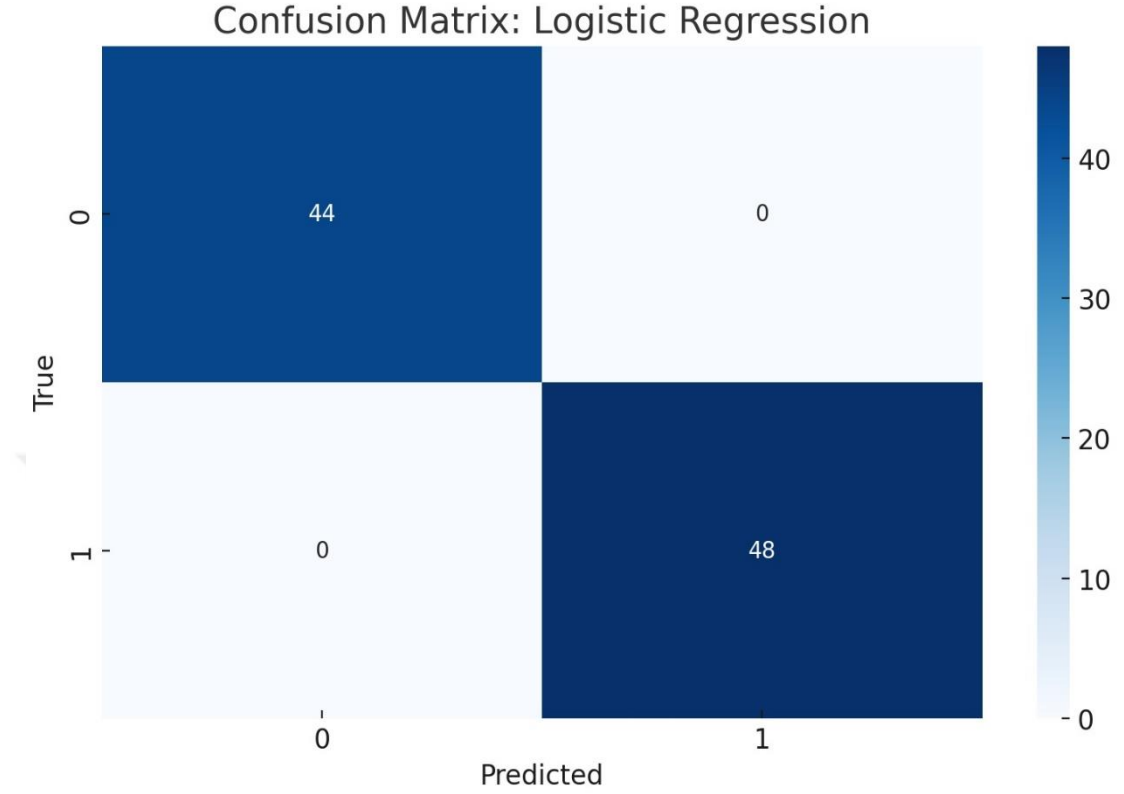
Şekil 4.13. Naive bayes doğruluk oranları

Doğruluk oranı 0.85 ile 0.95 arasında değişmekte olup, ortalama doğruluk oranı yaklaşık %90 civarındadır (Provost ve ark., 2013). Yüksek ve tutarlı performans göstermekte olan bu model, siber saldırı tespitinde güvenilir bir yöntem olarak öne çıkmaktadır. Ortalama doğruluk oranının %95'in üzerinde olduğu durumlarda bile performansını koruyabilmekte ve dış değer bulunmamaktadır (Murphy, 2012). Bu nedenle, Naive Bayes modeli, siber saldırı tespitinde güvenilir ve etkili bir yöntem olarak değerlendirilebilir.

Makine öğrenmesi ve yapay zeka modellerinin, görevlerindeki performanslarının değerlendirilmesi, nesnel ve standartlaştırılmış gösterimin kullanılmasıyla mümkündür. Bu ayrıştırma, modellerin doğruluğu, doğruluk ve genelleyici başarılarını belirleme ve farklı geliştirmelerin bakış açılarına kritik bakış açısına sahiptir. Özellikle görünen problemlerde, modelin okunmasının değerlendirilmesinde en yaygın olarak kullanılan dağılımlardan biri karışıklık matrisidir.

Karışıklık matrisi, büyümelerinin başarısının anlaşılması için önemli bir araç ve özellikle düzenli olmayan veri kümelerinde dağılım çeşitliliğin yanıltıcı olabilen daha detaylı bir analiz sağlar (Powers, 2011). Dolayısıyla bu matrisin analiz edilmesi, modelin

farklı açılardan incelenmesine olanak tanır ve modelin kullanım olanağını daha fazla değerlendirme imkânı sunar.

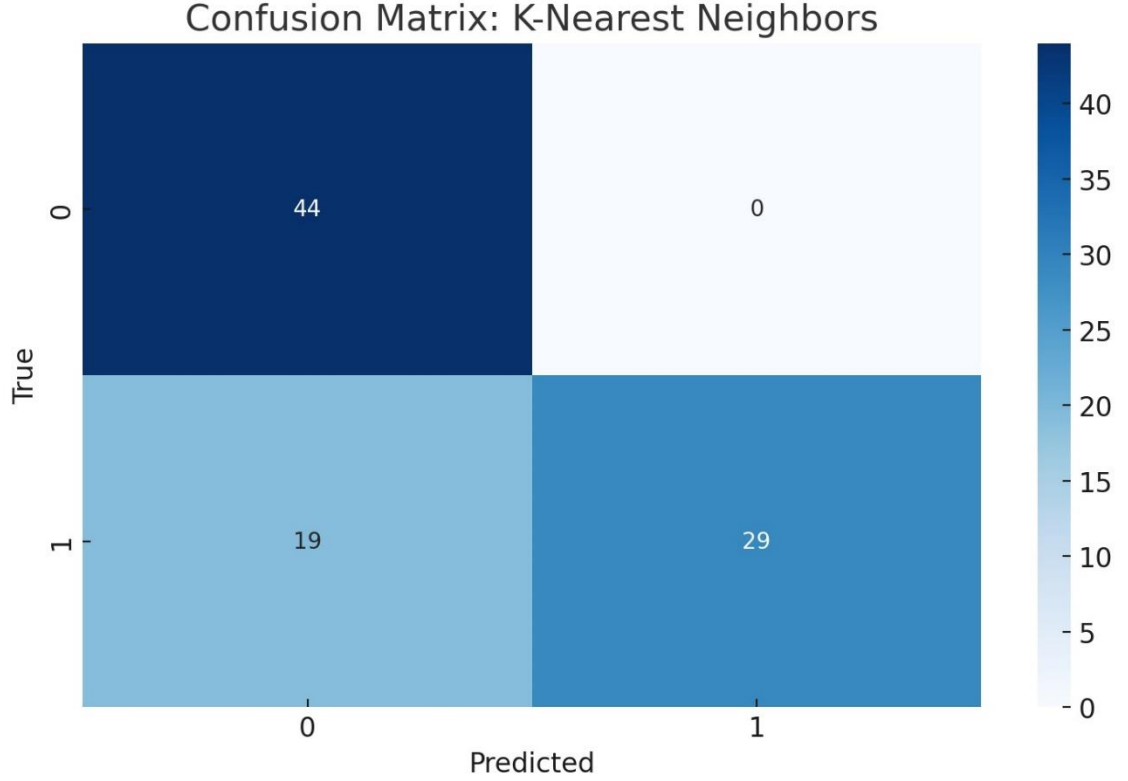


Şekil 4.14. Lojistik regresyon confusion matrisi

Lojistik Regresyon modeline ait karışıklık (confusion) matrisini Şekil 20’de göstermektedir. Matrisin açıklaması şu şekildedir;

- Gerçek Pozitif (True Positive, TP), Pozitif (1) sınıfını doğru tahmin eden durumlar. Matrisin sağ alt köşesinde ve 48 olarak gösterilmiştir.
- Gerçek Negatif (True Negative, TN), Negatif (0) sınıfını doğru tahmin eden durumlar. Sol üst köşede 44 olarak gösterilmiştir.
- Yanlış Pozitif (False Positive, FP), Negatif sınıfın yanlışlıkla pozitif olarak tahmin edildiği durumlar. Sağ üst köşede 0 olarak gösterilmiştir.
- Yanlış Negatif (False Negative, FN), Pozitif sınıfın yanlışlıkla negatif olarak tahmin edildiği durumlar. Sol alt köşede 0 olarak gösterilmiştir.

Bu sonuçlar, lojistik regresyon modelinin mükemmel bir performans sergilediğini göstermekte ve hiçbir yanlış pozitif veya yanlış negatif olmadığını söylemektedir. Model, hem negatif (0) hem de pozitif (1) sınıfları tamamen doğru tahmin etmektedir.

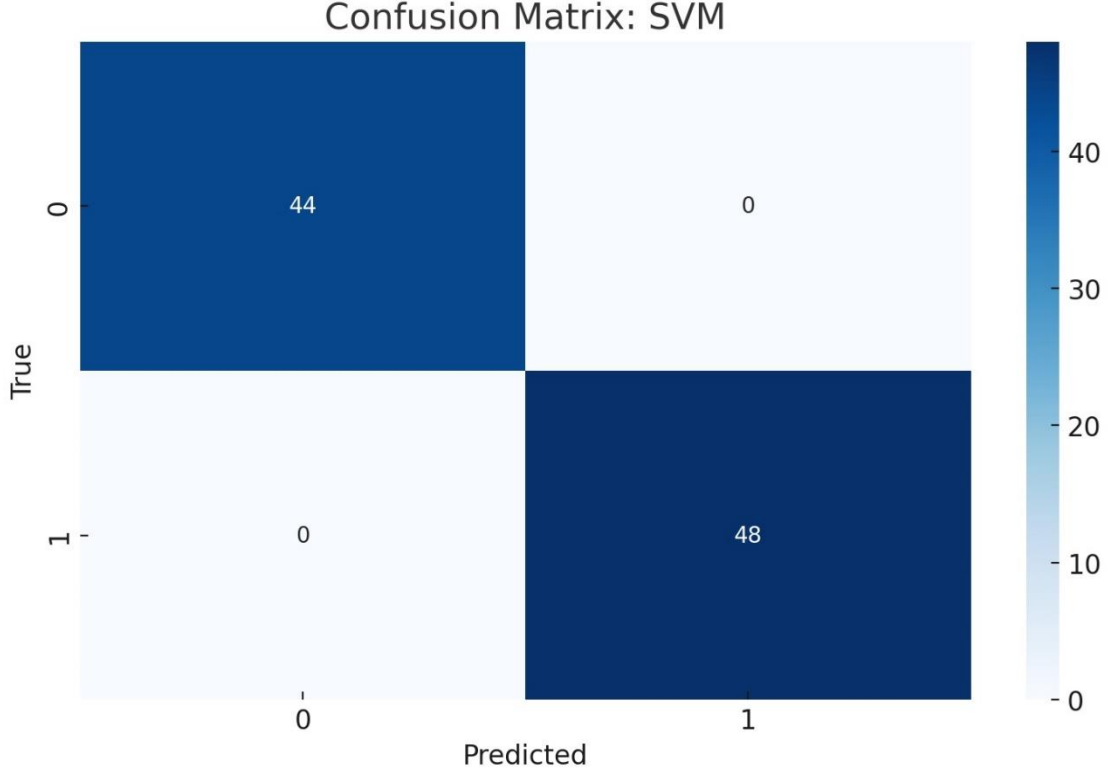


Şekil 4.15. K- En confusion matrisi

K-En Yakın Komşu (K-Nearest Neighbors, KNN) modeline ait karışıklık (confusion) matrisini Şekil 21’de göstermektedir. Matrisin açıklaması şu şekildedir;

- Gerçek Pozitif (True Positive, TP), Modelin doğru bir şekilde pozitif (1) sınıfını tahmin ettiği durumlar. Bu matriste sağ alt köşede 29 olarak gösterilmiştir.
- Gerçek Negatif (True Negative, TN), Modelin doğru bir şekilde negatif (0) sınıfını tahmin ettiği durumlar. Sol üst köşede 44 olarak gösterilmiştir.
- Yanlış Pozitif (False Positive, FP), Modelin yanlış bir şekilde pozitif sınıfını tahmin ettiği, yani aslında negatif olan ama pozitif olarak tahmin edilen durumlar. Sağ üst köşede 0 olarak gösterilmiştir.
- Yanlış Negatif (False Negative, FN), Modelin yanlış bir şekilde negatif sınıfını tahmin ettiği, yani aslında pozitif olan ama negatif olarak tahmin edilen durumlar. Sol alt köşede 19 olarak gösterilmiştir.

Bu sonuçlar, modelin negatif sınıfı (0) iyi tahmin ettiğini, ancak pozitif sınıfta (1) bir miktar hata yaptığını göstermektedir (19 yanlış negatif tahmin). Yanlış pozitif (0) bulunmamakta, bu da pozitif sınıfların aşırı tahmin edilmediği anlamına gelmektedir.

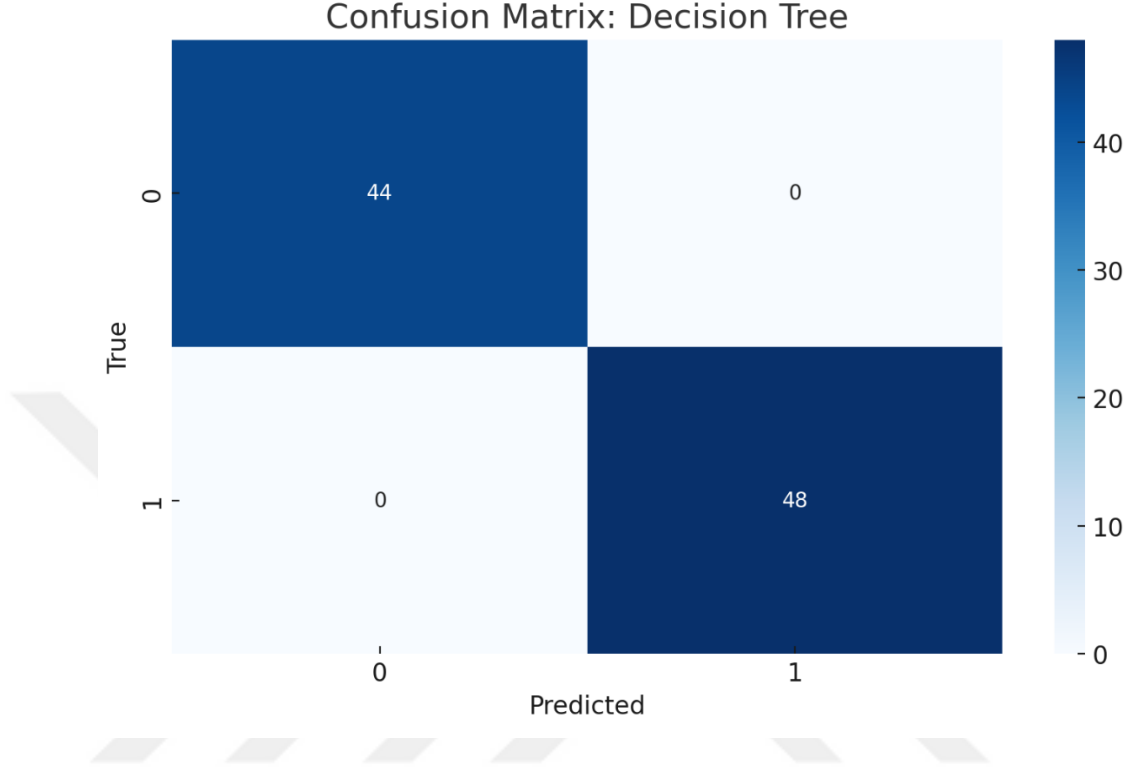


Şekil 4.16. SVM confusion matrisi

Bu görselde, Destek Vektör Makineleri (Support Vector Machine, SVM) modeline ait karışıklık (confusion) matrisini Şekil 22’de göstermektedir. Matrisin açıklaması şu şekildedir;

- Gerçek Pozitif (True Positive, TP), Pozitif (1) sınıfını doğru tahmin eden durumlar. Sağ alt köşede 48 olarak gösterilmiştir.
- Gerçek Negatif (True Negative, TN), Negatif (0) sınıfını doğru tahmin eden durumlar. Sol üst köşede 44 olarak gösterilmiştir.
- Yanlış Pozitif (False Positive, FP), Negatif sınıfın yanlışlıkla pozitif olarak tahmin edildiği durumlar. Sağ üst köşede 0 olarak gösterilmiştir.
- Yanlış Negatif (False Negative, FN), Pozitif sınıfın yanlışlıkla negatif olarak tahmin edildiği durumlar. Sol alt köşede 0 olarak gösterilmiştir.

Bu sonuçlar, SVM modelinin mükemmel bir performans sergilediğini gösteriyor. Hiçbir yanlış pozitif (FP) ya da yanlış negatif (FN) olmadığını, yani model hem pozitif hem de negatif sınıfları tamamen doğru bir şekilde tahmin etmektedir.

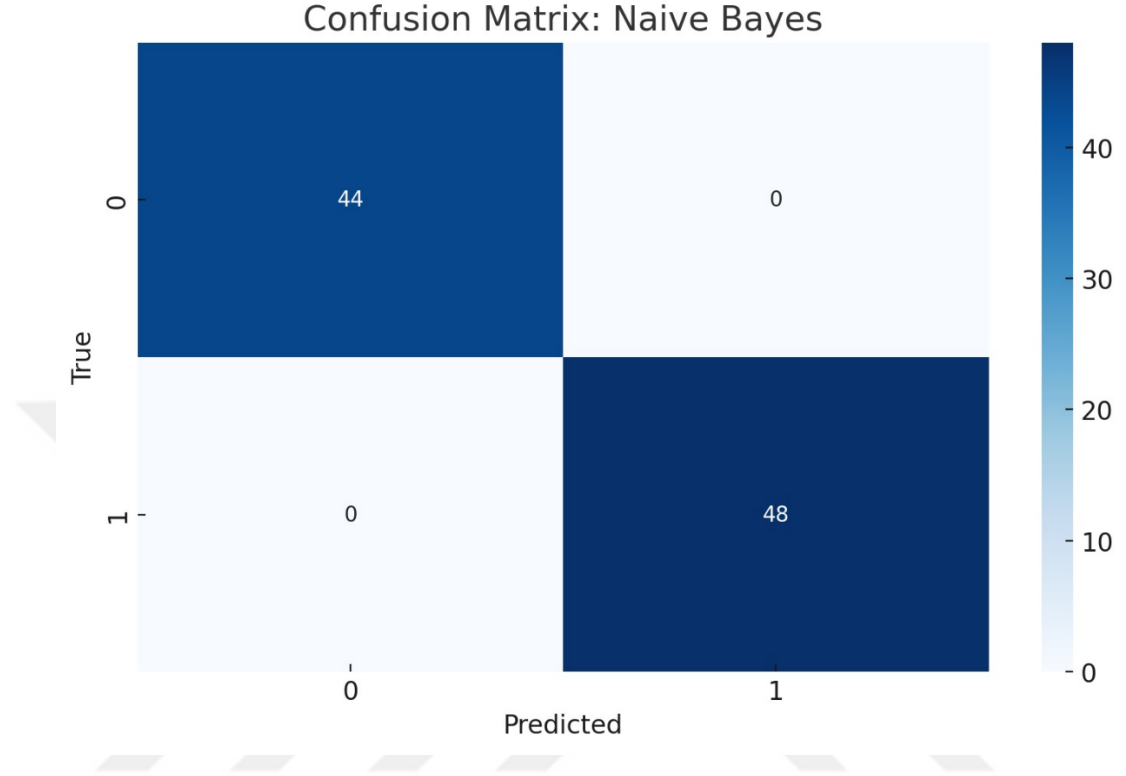


Şekil 4.17. Karar ağacı confusion matrisi

Görselde bir karar ağacı sınıflandırıcısına ait karışıklık (confusion) matrisini Şekil 23'de göstermektedir. Matrisin açıklaması şu şekildedir;

- "0" sınıfı (sol üst köşe): Model 44 örneği doğru bir şekilde "0" sınıfı olarak tahmin etmiş.
- "1" sınıfı (sağ alt köşe): Model 48 örneği doğru bir şekilde "1" sınıfı olarak tahmin etmiş.
- "Yanlış Pozitif" (sağ üst köşe): Model hiçbir "0" sınıfını yanlış bir şekilde "1" olarak tahmin etmemiş (0 değer).
- "Yanlış Negatif" (sol alt köşe): Model hiçbir "1" sınıfını yanlış bir şekilde "0" olarak tahmin etmemiş (0 değer).

Bu sonuçlara göre model mükemmel bir performans sergiliyor, çünkü hiçbir yanlış sınıflandırma yok (yanlış pozitif ve yanlış negatif sayısı sıfır). Model tüm örnekleri doğru bir şekilde sınıflandırmış.



Şekil 4.18. Naive bayes confusion matrisi

Bu görselde Naive Bayes modeline ait karışıklık (confusion) matrisini Şekil 24'de göstermektedir. Matrisin açıklaması şu şekildedir;

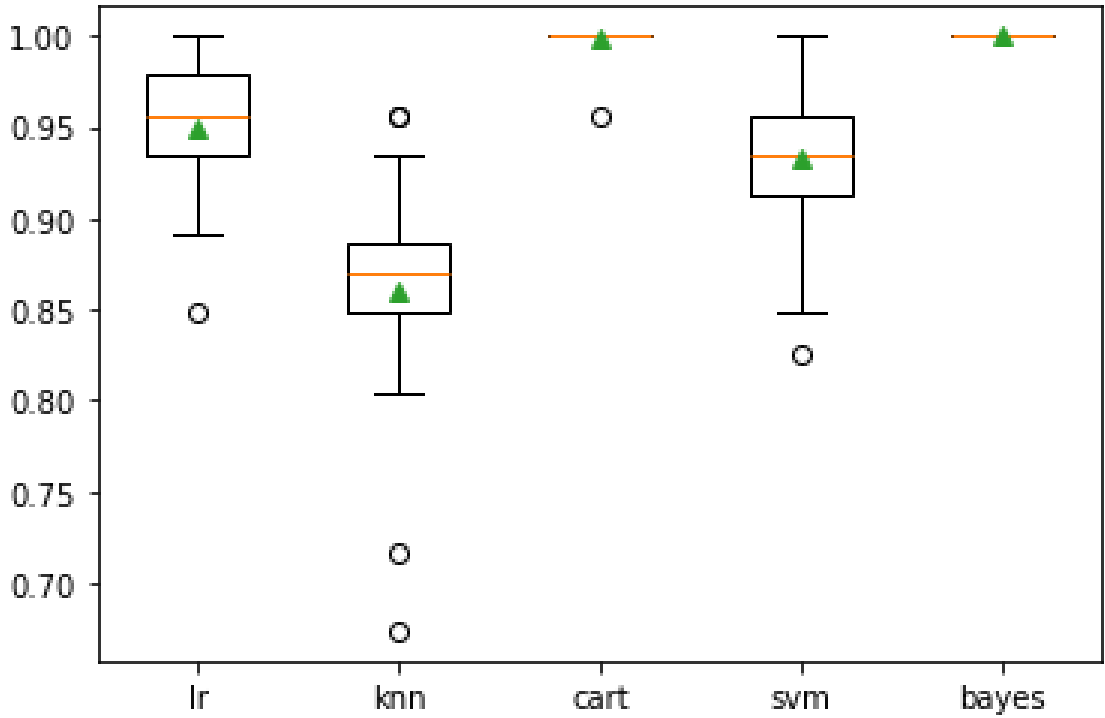
- "0" sınıfı (sol üst köşe): Model 44 örneği doğru bir şekilde "0" sınıfı olarak tahmin etmiş.
- "1" sınıfı (sağ alt köşe): Model 48 örneği doğru bir şekilde "1" sınıfı olarak tahmin etmiş.
- "Yanlış Pozitif" (sağ üst köşe): Model hiçbir "0" sınıfını yanlış bir şekilde "1" olarak tahmin etmemiş (0 değer).
- "Yanlış Negatif" (sol alt köşe): Model hiçbir "1" sınıfını yanlış bir şekilde "0" olarak tahmin etmemiş (0 değer).

Bu karışıklık matrisi de tıpkı önceki görselde olduğu gibi modelin tüm sınıflandırmaları doğru yaptığını göstermektedir. Yani, Naive Bayes modeli bu veri seti

üzerinde mükemmel bir performans sergilemiş, yanlış pozitif veya yanlış negatif tahminlerde bulunmamıştır.

5.ARAŞTIRMA SONUÇLARI VE TARTIŞMA

Veri setinde analiz edilen ikinci grafik, farklı makine öğrenimi modellerinin performanslarını karşılaştırarak siber saldırı tespiti üzerindeki etkilerini analiz etmeyi amaçlamaktadır.



Şekil 5.1. Makine öğrenim modelleri

Şekil 25’de farklı makine öğrenimi modellerinin performanslarını karşılaştırarak siber saldırı tespiti üzerindeki etkilerini analiz edilmiştir. Ayrıca, veri setindeki çeşitli özelliklerin dağılımlarını inceleyerek model performansını nasıl etkileyebileceğini değerlendirilmiştir. Makine öğrenimi modelleri olarak lojistik regresyon (lr), k-en yakın komşu (knn), karar ağacı (cart), destek vektör makineleri (svm) ve naive bayes (bayes) modelleri kullanılmıştır.

Her bir modelin doğruluk oranları kutu grafiği ile gösterilmiştir. Kutu grafikleri, veri dağılımlarının özetlenmesinde kullanılan istatistiksel grafiklerdir ve bu grafikte modellerin doğruluk oranlarının dağılımını göstermektedir. Grafikte yer alan kutuların üst ve alt sınırları, veri setindeki doğruluk oranlarının çeyrekler arası aralığını temsil

etmektedir. Kutuların içerisindeki yatay çizgiler medyan değerini, yeşil üçgenler ise ortalama değeri göstermektedir. Kutuların dışında yer alan çemberler ise aykırı değerleri belirtmektedir.

Bu çalışmada, farklı makine öğrenimi modellerinin siber saldırı tespiti üzerindeki performansları detaylı bir şekilde incelenmiştir. Şekil 20’de, kullanılan beş farklı makine öğrenimi modelinin doğruluk oranları kutu grafikleri aracılığıyla görselleştirilmiştir. Kullanılan modeller arasında lojistik regresyon (lr), k-en yakın komşu (knn), karar ağacı (cart), destek vektör makineleri (svm) ve naive bayes (bayes) yer almaktadır. Aşağıda, bu modellerin performanslarına ilişkin bulgular ve bu bulguların tartışması yer almaktadır.

5.1. Araştırma Sonuçlarının Değerlendirilmesi

Bu araştırmada, çeşitli makine öğrenimi modellerinin (Lojistik Regresyon, K-En Yakın Komşu, Karar Ağacı, Destek Vektör Makineleri ve Naive Bayes) sınıflandırma performansları karşılaştırılmıştır. Modeller, doğruluk oranları ve aykırı değer dağılımları açısından değerlendirilmiş, her bir modelin doğruluk performansı analiz edilmiştir. Elde edilen sonuçlar, farklı model yapılarının veri seti üzerindeki etkisini ve performans farklılıklarını ortaya koymaktadır.

Lojistik regresyon (lr) modeli, doğruluk oranı bakımından oldukça iyi bir performans sergilemiş ve diğer modellerle karşılaştırıldığında en yüksek ortalama doğruluk oranına sahip olmuştur.

K-en yakın komşu (knn) modeli, diğer modellere göre daha geniş bir doğruluk oranı dağılımına sahiptir ve birkaç aykırı değer içermektedir.

Karar ağacı (cart) modeli, nispeten daha dar bir dağılıma sahip olup, medyan değeri ve ortalama doğruluk oranı açısından orta sıralarda yer almıştır. Bu modelin ortalama doğruluk oranı açısından orta sıralarda yer almakta ve performansının nispeten dar bir aralıkta yoğunlaştığı görülmektedir. Bu modelde de aykırı değerler mevcuttur, ancak genel olarak tutarlı bir performans sergilediği söylenebilir.

Destek vektör makineleri (svm) modeli, doğruluk oranı bakımından lojistik regresyon modeliyle benzer bir performans sergilemiş, ancak daha fazla aykırı değer içermiştir bu durum da bazı durumlarda performansın düşebileceğini göstermektedir.

Naive bayes (bayes) modeli ise en dar dağılıma sahip olup, oldukça yüksek bir doğruluk oranı sergilemiştir. Aykırı değerlerin azlığı, bu modelin tutarlı ve güvenilir performansını desteklemektedir.

Bu analizde, veri setindeki çeşitli özelliklerin dağılımları da incelenmiş ve bu özelliklerin model performansını nasıl etkileyebileceği değerlendirilmiştir. Elde edilen sonuçlar, siber saldırı tespitinde kullanılan makine öğrenimi modellerinin performansları arasında belirgin farklar olduğunu ve bu modellerin seçiminin tespit başarısında önemli bir rol oynadığını göstermektedir.

5.2. Tartışma

Lojistik regresyon (LR) modelinin yüksek ortalama doğruluk oranı ve dar performans dağılımı, bu modelin siber saldırı tespitinde son derece güvenilir olduğunu göstermektedir. Modelin istikrarlı performansı, lojistik regresyonun veri setindeki özelliklerle iyi bir şekilde uyum sağladığını işaret etmektedir; bu nedenle, lojistik regresyon modeli, siber güvenlik uygulamalarında güvenilir bir seçenek olarak değerlendirilebilir (Hosmer ve ark., 2013).

Diğer yandan, K-en yakın komşu (KNN) modelinin geniş performans dağılımı ve yüksek aykırı değer sayısı, modelin performansının veri setinin yapısına oldukça duyarlı olduğunu ortaya koymaktadır. Bu durum, KNN modelinin belirli senaryolarda çok iyi performans gösterebileceği gibi, bazı durumlarda performansının düşebileceğini de işaret etmektedir. KNN modelinin kullanılacağı durumlarda veri setinin dikkatli bir şekilde seçilmesi veya ön işlemden geçirilmesi gerekmektedir (Cover ve ark., 1967).

Karar ağacı (CART) modelinin ortalama doğruluk oranı, diğer modellere göre daha düşük olmasına rağmen, tutarlı performansı dikkate alındığında belirli senaryolarda makul bir seçenek olarak değerlendirilebilir. Ancak, daha yüksek performans gerektiren uygulamalarda yetersiz kalabileceği göz önünde bulundurulmalıdır (Breiman ve ark., 1984).

Destek vektör makineleri (SVM) modelinin yüksek ortalama doğruluk oranı ve dar performans dağılımı, bu modelin genel olarak güvenilir olduğunu göstermektedir. Bununla birlikte, aykırı değerlerin varlığı, SVM modelinin bazı durumlarda performans düşüşleri yaşayabileceğini işaret etmektedir. Yine de, bu model yüksek doğruluk oranı gerektiren siber güvenlik uygulamalarında etkili bir seçenek olabilir (Cortes ve ark., 1995).

Naive Bayes (Bayes) modelinin en dar performans dağılımına sahip olması ve yüksek ortalama doğruluk oranı sergilemesi, bu modelin siber saldırı tespitinde son derece etkili olduğunu göstermektedir. Aykırı değerlerin azlığı, modelin genel olarak

tutarlı ve güvenilir sonuçlar verdiğini ortaya koymaktadır. Naive Bayes modelinin bu performansı, özellikle büyük veri setlerinde ve gerçek zamanlı uygulamalarda tercih edilmesini sağlamaktadır (McCallum ve ark., 1998).

Bu değerlendirmeler ışığında, çeşitli makine öğrenimi modellerinin siber güvenlik uygulamalarındaki performanslarını karşılaştırmak ve uygulama alanlarına göre en uygun modeli seçmek önemlidir. Lojistik regresyon ve Naive Bayes modelleri, genel olarak yüksek doğruluk oranları ve tutarlı performansları ile dikkat çekerken, KNN ve SVM modelleri veri setinin özelliklerine ve senaryoya bağlı olarak değişken performans gösterebilirler. Karar ağacı modelleri ise belirli senaryolarda kullanılabilir, ancak genel olarak diğer modellere göre daha düşük doğruluk oranına sahiptir.



6. GENEL DEĞERLENDİRME VE SONUÇ

6.1. Genel Değerlendirme

Siber saldırı tespiti için lojistik regresyon, destek vektör makineleri ve Naive Bayes algoritmaları, yüksek doğruluk ve tutarlı performansları ile en uygun algoritmalar olarak öne çıkmaktadır (Murphy, 2012). Bu algoritmaların yüksek doğruluk oranları ve tutarlılıkları, siber saldırı tespitinde güvenilir birer seçenek olduklarını göstermektedir (Bishop, 2006). Buna karşın, K-en yakın komşu ve karar ağaçları algoritmaları, veri kümesinin özelliklerine daha duyarlı olup, performanslarında daha fazla değişkenlik göstermektedir. Bu durum, bu algoritmaların siber saldırı tespitinde bazı durumlarda daha az güvenilir olabileceğini düşündürmektedir (Hastie, ve ark., 2009).

Bu bulgular, belirli bir siber saldırı tespit problemi için en uygun algoritmanın seçiminin veri kümesinin özelliklerine ve problem tanımına bağlı olduğunu göstermektedir. Her bir algoritmanın avantaj ve dezavantajları dikkate alınarak, problem özelinde kapsamlı bir değerlendirme yapılması önemlidir.

Bu çalışma siber saldırı tespitinde kullanılan makine öğrenimi modellerinin performansları arasında belirgin farklar olduğunu ortaya koymaktadır. Lojistik regresyon ve Naive Bayes modelleri, genel performans açısından en iyi sonuçları verirken, K-en yakın komşu modeli daha geniş bir performans değişkenliği göstermiştir. Karar ağacı ve destek vektör makineleri ise genel olarak makul performanslar sergilemiştir. Makine öğrenimi modellerinin siber saldırı tespitindeki etkinlikleri, veri setindeki özelliklerin dağılımına ve modelin bu özelliklere nasıl uyum sağladığına bağlı olarak değişmektedir. Bu nedenle, siber güvenlik alanında daha güvenilir sonuçlar elde etmek için modellerin birlikte kullanılması veya belirli senaryolar için uygun modelin seçilmesi önemlidir.

Makine öğrenmesi algoritmaları, siber saldırı tespitinde önemli bir rol oynamakta ve güvenlik sistemlerinin etkinliğini artırmaktadır. Bu çalışmada kullanılan algoritmalar, siber güvenlik alanında daha etkili ve hızlı tespit sistemleri geliştirilmesine olanak sağlamaktadır. Gelecekte, farklı veri kümeleri ve saldırı senaryoları üzerinde daha fazla çalışma yapılarak, bu algoritmaların performansları daha da iyileştirilebilir ve siber güvenlik önlemleri daha da güçlendirilebilir (Murphy, 2012), (Bishop, 2006), (Hastie ve ark., 2009).

Gelecekte yapılacak çalışmalar, bu modellerin daha farklı veri setleri üzerinde performanslarını inceleyerek, daha kapsamlı ve genelleştirilebilir sonuçlar elde etmeyi

amaçlamalıdır. Ayrıca, modellerin birlikte kullanılması veya hibrit modeller geliştirilmesi, siber saldırı tespitinde daha yüksek performans elde edilmesine katkı sağlayabilir.

6.2. Sonuç

Bu grafiklerin siber saldırılarla bağlantısını kurmak için, veri setindeki özelliklerin ve modellerin siber saldırı tespiti üzerindeki etkisini anlamak önemlidir. Genellikle siber saldırı tespitinde kullanılan modellerin doğruluk oranları, saldırıların doğru bir şekilde tespit edilip edilmediğini gösterir. Yüksek performanslı modeller, siber saldırı tespitinde etkili olabilir. Özellikle lojistik regresyon, karar ağacı, destek vektör makineleri ve Naive Bayes modellerinin yüksek performans göstermesi, bu modellerin siber saldırı tespitinde kullanılabileceğini göstermektedir (Hastie, ve ark., 2009). Ancak, veri setindeki bazı özelliklerin dengesiz dağılımı, modelin performansını olumsuz etkileyebilir ve modelin bu özelliklere aşırı uyum sağlamasına neden olabilir (Provost ve ark., 2013).

Bu analiz, veri setindeki özelliklerin ve makine öğrenimi modellerinin siber saldırı tespiti üzerindeki etkilerini anlamamıza yardımcı olmuştur. Yüksek performans gösteren modeller, siber saldırıların tespiti ve önlenmesinde kullanılabilir. Bununla birlikte, veri setindeki dengesiz dağılımlar, modelin performansını etkileyebilir ve bu nedenle veri ön işleme aşamasında dikkat edilmesi gereken önemli bir faktördür. Gelecekteki çalışmalar, veri setindeki dengesizliklerin giderilmesine yönelik teknikler geliştirilerek model performansının iyileştirilmesine odaklanmalıdır.

Bu çalışmanın sonuçları, XSS saldırılarının tespiti ve analizi konusunda önemli bilgiler sağlamaktadır. Gelecekte, bu çalışmanın sonuçlarını daha da ileriye taşımak ve XSS saldırılarını önlemek için çeşitli önlemler alınabilir. Örneğin, makine öğrenimi modellerinin yanı sıra, saldırı tespit sistemleri ve güvenlik duvarları gibi güvenlik önlemleri entegre edilebilir (Aburomman ve ark., 2017). Ayrıca, kullanıcı eğitimi ve farkındalık çalışmaları da saldırıların önlenmesinde önemli bir rol oynayabilir (Jensen ve ark., 2005).

Sonuç olarak, bu çalışma, XSS saldırılarının tespiti ve analizi konusunda önemli bir adım atmıştır. Elde edilen sonuçlar ve öneriler, web uygulamalarının güvenliğini artırmak için önemli rehberlik sağlayabilir. Gelecekte yapılacak çalışmalar, bu alandaki

bilgi birikimini daha da zenginleřtirerek, internet gvenlięinin saęlanmasına nemli katkılarda bulunacaktır.



KAYNAKLAR

- Aburomman, A. A., Reaz, M. B. I., 2017. A novel SVM-kNN-PSO ensemble method for intrusion detection system. *Applied Soft Computing*, 38, 360-372.
- Amoroso, E., 2006, *Cyber Security*. New Jersey: Silicon Press.
- Albin, E. S., Mitropoulos, F. J., 2019, Cross-site scripting (XSS) attacks and defense mechanisms: classification and state-of-the-art. *International Journal of Information Security*, 18(1), 89-122.
- Al-Dossari, H., 2016, Security and privacy issues in cloud computing: A survey. *International Journal of Computer Applications*, 135(6), 20-27.
- Anaconda., 2021, *Anaconda Navigator Documentation*. Retrieved from <https://docs.anaconda.com/anaconda/navigator/>
- Aggarwal, C. C., Hinneburg, A., & Keim, D. A., 2001, On the Surprising Behavior of Distance Metrics in High Dimensional Space. In *Database Theory — ICDT 2001* (pp. 420-434). Springer, Berlin, Heidelberg.
- Baş, E., Süzen, A.A., 2023, Web Uygulama Zafiyetlerinin Keşfinde Açıklanabilir Yapay Zekânın Yeri, *Araştırma Makalesi, Uluslararası Uygulamalı Mühendislik ve Doğa Bilimleri Konferansı*, 517-527.
- Baykara, M., Güçlü, S., 2018, Applications for detecting XSS attacks on different web platforms, 6th International Symposium on Digital Forensic and Security (ISDFS), pp. 1-6, IEEE.
- Baykara, M., Daş, R., Tuna, G., 2016, Web Sunucu Erişim Kütüklerinden Web Ataklarının Tespitine Yönelik Web Tabanlı Log Analiz Platformu, *Araştırma Makalesi, Fırat Üniversitesi Mühendislik Bilimleri Dergisi*, 28 (2), 291-302, 2016.
- Barth, A., 2011, HTTP State Management Mechanism. RFC 6265, Internet Engineering Task Force (IETF).
- Bengio, Y., Goodfellow. I., Courville, A., 2016, *Deep Learning (Vol.1)*, Massachusett, USA:: MIT press.
- Belani, G., 2021, The Use of Artificial Intelligence in Cybersecurity: A Review, <https://www.computer.org/publications/tech-news/trends/the-use-of-artificial-intelligence-in-cybersecurity>, Erişim Tarihi: 24.12.2021.
- Bishop, C. M., 2006, *Pattern Recognition and Machine Learning*. Springer.
- Bisht, P., & Venkatakrishnan, V. N., 2008, XSS-GUARD: Precise dynamic prevention of cross-site scripting attacks. In *Detection of Intrusions and Malware, and Vulnerability Assessment*, 23-43.
- Breiman, L., Friedman, J. H., Olshen, R. A., & Stone, C. J., 1984, *Classification and Regression Trees*. Wadsworth & Brooks/Cole Advanced Books & Software.

- Brown, A., 2023, Spyder IDE Documentation. Spyder Development Team.
- Cova, M., et al., 2010, Analyzing and detecting drive-by download attacks. In Proceedings of the 19th International Conference on World Wide Web, 281-290.
- Coventry L., Briggs, P., Blythe, J., Tran, M., 2014, Report, Using Behavioural Insights To Improve The Public's Use Of Cyber Security Best Practices.
- Cover, T., Hart, P., 1967, Nearest neighbor pattern classification. IEEE Transactions on Information Theory, 13(1), 21-27.
- Cortes, C., & Vapnik, V., 1995, Support-vector networks. *Machine Learning*, 20, 273-297.
- Çinar, I., Bilge, H.Ş., 2016, Web Madenciliği Yöntemleri ile Web Loglarının İstatistiksel Analizi ve Saldırı Tespiti, Araştırma Makalesi, Bilişim Teknolojileri Dergisi, Cilt: 9, Sayı: 2, Mayıs 2016, 125-135.
- Doupe, A., et al., 2010, Enemy of the state: A state-aware black-box web vulnerability scanner. In USENIX Security Symposium, 523-538.
- Efe, A., 2021, Yapay Zekâ Odaklı Siber Risk ve Güvenlik Yönetimi, Araştırma Makalesi, Uluslararası Yönetim Bilişim Sistemleri ve Bilgisayar Bilimleri Dergisi, 2021, 5(2):144-165.
- Ertuğrul, Ö.,F., 2015, Davranışsal Öğrenme Bazlı Bir Yapay Makine Öğrenme Yönteminin Geliştirilmesi, Doktora Tezi, İnönü Üniversitesi Fen Bilimleri Enstitüsü.
- Etik, E., Can, A., 2024, Siber Güvenlikte Kali Linux ve Sızma Araçlarının Kullanımı, Araştırma Makalesi, Kadirli Uygulamalı Bilimler Fakültesi Dergisi, Cilt 4, Sayı 1, 210-226.
- Fawcett, T., 2006, An introduction to ROC analysis. Pattern Recognition Letters, 27(8), 861-874.
- Gupta, B. B., Agrawal, D. P., & Yamaguchi, S., 2016, *Handbook of Research on Modern Cryptographic Solutions for Computer and Cyber Security*. IGI Global.
- Gollmann, D., 2011, *Computer Security* (3rd ed.). Wiley.
- Grossman, J., Hansen, R., Petkov, P., Rager, A., & Stuttard, D., 2007, *Cross Site Scripting Attacks: XSS Exploits and Defense*. Syngress.
- Grossman, J., 2011, *XSS Attacks: Cross-Site Scripting Exploits and Defense*. Syngress.
- Hastie, T., Tibshirani, R., & Friedman, J., 2009, *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*. Springer.
- Huang, Y., et al., 2020, *Web application security: Static and dynamic analysis*. Springer.

- Hosmer, D. W., & Lemeshow, S., 2000, *Applied Logistic Regression*. Wiley.
- Hosmer, D. W., Lemeshow, S., & Sturdivant, R. X., 2013, *Applied Logistic Regression* (Vol. 398). John Wiley & Sons.
- <https://berqnet.com/blog/siber-guvenlik-nedir> (27.01.2021)
- <https://www.itu.int/rec/T-REC-X.1205-200804-I> (06.03.2024)
- https://www.webopedia.com/TERM/C/cyber_crime.html(15.02.2019)
- Jensen, C. D., et al., 2005, Building an understanding of security and privacy threats. *IBM Systems Journal*, 44(2), 255-270.
- Jones, T. (2023). *Anaconda Navigator Documentation*. Anaconda Inc.
- Jovanovic, N., Kruegel, C., & Kirda, E., 2006, "Pixy: A Static Analysis Tool for Detecting Web Application Vulnerabilities (Short Paper)." *2006 IEEE Symposium on Security and Privacy*. IEEE
- Kleinbaum, D. G., & Klein, M., 2010, *Logistic Regression: A Self-Learning Text*. Springer.
- Kara, İ., 2019, Detection, Technical Analysis of Brute Force Attack, *Sakarya University Journal Of Computer and Information Sciences*, Vol.2, No.2.
- Karasoy, H.A., Gezici, H.S., 2023, Bombalardan Baytlara: Siber Güvenliğin Ulusal Güvenlikteki Rolü ve Yapay Zekânın Siber Güvenlikteki Önemi, *Araştırma Makalesi, Uluslararası Yönetim Akademisi Dergisi*, Cilt: 6, Sayı: 1, ss.173-188.
- Kartal, M., Sağıroğlu, Ş., Bülbül, H.İ., 2013, *Politeknik Dergisi*, Cilt:16, Sayı: 3, 119-127.
- Kaya, Ç., Yıldız, O., 2014, Makine Öğrenmesi Teknikleriyle Saldırı Tespiti: Karşılaştırılmalı Analiz, *Fen Bilimleri Dergisi*, 3:89-104.
- Koyuncu, M.D., Ünlü, N., 2022, Yapay Zeka Tekniklerinin Saldırı Tespit Sistemleri'nde Kullanımı, *Araştırma Makalesi, Beykoz Akademi Dergisi*, 2022; 10(1), 78-87.
- Kirda, E., Kruegel, C., 2005, "Protecting Web Application from Cross Site Scripting Attacks." *International Conference on Security and Privacy in Communication Networks*. IEEE.
- Kiczales, G., Lamping, J., Mendhekar, A., Maeda, C., Lopes, C. V., Loingtier, J. M., & Irwin, J., 2001, "Aspect-Oriented Programming." In *European Conference on Object-Oriented Programming* (pp. 220-242). Springer Berlin Heidelberg.
- Le, Q.V., Ngiam, J., Coates, A., Lahiri, A., Prochnow, B., Ng, A.Y., 2011 January, On Optimization methods for deep learning, In *ICML*, Computer Science Department, Stanford University, Stanford, CA 94305, USA.

- McCallum, A., & Nigam, K., 1998 July, A comparison of event models for Naive Bayes text classification. In *AAAI-98 workshop on learning for text categorization* (Vol. 752, pp. 41-48).
- Murphy, K. P., 2012, *Machine Learning: A Probabilistic Perspective*. MIT Press.
- Muhammet, G. Sebahattin, "Applications for detecting XSS attacks on different web platforms," In: 2018 6th International Symposium on Digital Forensic and Security (ISDFS), pp. 1-6, IEEE, 2018.
- Of, M., Kılıçaslan, İ., 2021, Xss (Cros-Site Scripting) Web Güvenliği Açığının İşletmeler Açısından Önemi, Araştırma Makalesi, Uluslararası Sosyal ve Beşeri Bilimler Araştırma Dergisi (JSHSR), Issue/Sayı: 77 Volume/Cilt: 8, 3144-3152.
- OWASP, 2017, OWASP Enterprise Security API (ESAPI). *OWASP Foundation*. <https://owasp.org/www-project-enterprise-security-api/>
- Özbilen, T., Çağlar, A., 2020, Derleme Makalesi, Kamu Yönetimi ve Teknoloji Dergisi, Sayı:1, 72-93.
- Özel, A., Kaya, Ç., Eken, S., 2014, JavaScript Güvenlik Açıkları ve Güncel Çözüm Önerileri, 7. Uluslararası Bilgi Güvenliği ve Kriptoloji Konferansı (ISCTURKEY), İstanbul.
- Öztürk, M.S., 2018, Siber Saldırıları, Siber Güvenlik Denetimleri Ve Bütüncül Bir Denetim Modeli Önerisi, Muhasebe ve Vergi Uygulamaları Dergisi, 208-232.
- Pham, B., 2013, "An Overview of Cross-Site Scripting." *IEEE Security & Privacy*.
- Peng, C. Y. J., Lee, K. L., Ingersoll, G. M., 2002, An Introduction to Logistic Regression Analysis and Reporting. *The Journal of Educational Research*, 96(1), 3-14.
- Pieterse, H., Olivier, M. S., 2012, Security analysis of smartphones using static analysis. *Journal of Information Security and Applications*, 17(4), 256-270
- Provost, F., Fawcett, T., 2013, *Data Science for Business: What You Need to Know about Data Mining and Data-Analytic Thinking*. O'Reilly Media.
- Powers, D. M., 2011, Evaluation: From Precision, Recall and F-measure to ROC, Informedness, Markedness and Correlation. *Journal of Machine Learning Technologies*, 2(1), 37-63.
- PWC., 2018, Global Economic Crime and Fraud Survey 2018. Sabillon, R., Serra-Ruiz, J., Cavaller, V., & Cano, J. (2017). A Comprehensive Cybersecurity Audit Model to Improve.
- Quinlan, J. R., 1986, Induction of Decision Trees. *Machine Learning*, 1(1), 81-106.
- Reşitoğlu, Ş.N., 2021, Siber Güvenlikte Yapay Zekâ Etkisi, <https://www.tuicakademi.org/siber-guvenlikte-yapay-zeka-etkisi/>, Erişim Tarihi: 25.12.2021

- Russell, S.J., Norving, P., 1995, Book, Artificial Intelligence A Modern Approach.
- Russell, S., & Norvig, P. (2010). Artificial Intelligence: A Modern Approach. Prentice Hall.
- Sevri, M., 2011, Web Saldırılarının Tespitine Yönelik Yapay Zekâ Tabanlı Saldırılarının Tespitine Yönelik Zekâ Tabanlı Bir Güvenlik Modülü Geliştirilmesi, Yüksek Lisans Tezi, *Gazi Üniversitesi Bilişim Enstitüsü*, 18-20.
- Schölkopf, B., Smola, A. J., 2002, Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond. MIT Press.
- Sullivan, C., Sood, A. K., 2008, "Web Security: A White Paper." *IEEE Security & Privacy*.
- Stampar, M., Jakupovic, A., 2014, *Advanced SQL Injection Handbook*. DeepSec GmbH.
- Smith, M., 2020, Data analysis with Python: A modern approach. Wiley.
- Smith, J., 2023, Python Programming Language. Python Software Foundation.
- Şeker, E., 2020, Uluslararası Bilgi Güvenliği Mühendisliği Dergisi, Cilt:6, Sayı:2, 108-115.
- Toğaçar, M., 2021, Web Sitelerinde Gerçekleştirilen Ortalama Saldırılarının Yapay Zekâ Yaklaşımı İle Tespiti, Araştırma Makalesi, BEÜ Fen Bilimleri Dergisi, 10 (4), 1603-1614.
- Tunca, S., 2019, DPÜ İİBF Dergisi, Sayı:3, 1-7.
- Vural, Y., Sağıroğlu, Ş., 2008, Kurumsal Bilgi Standartları Üzerine Bir İnceleme, Gazi Üniversitesi Mühendislik Mimarlık Fakültesi Dergisi, Ankara.
- Vapnik, V., 1998, Statistical Learning Theory. Wiley.
- Yaşar, H., 2014, Kurumsal Siber Güvenliğe Yönelik Tehditler ve Mücadele Yöntemleri: Eylem Planı Örneği, Yüksek Lisans Tezi, *Gazi Üniversitesi Bilişim Enstitüsü*, Ankara, 10-34.
- Yaşar, H., Çakır, H., 2015, Kurumsal Siber Güvenliğe Yönelik Tehditler ve Önlemleri, Araştırma Makalesi, Düzce Üniversitesi Bilim ve Teknoloji Dergisi, 3 (2015) 488-507.
- Yiğit, T., Akyıldız, M.A., 2014, Sızma Testleri İçin Bir Model Ağ Üzerinde Siber Saldırı Senaryolarının Değerlendirilmesi, Araştırma Makalesi, Süleyman Demirel Üniversitesi Fen Bilimleri Enstitüsü Dergisi, 18(1), 14-21, 2014.
- Zakas, N. C., 2012, *Understanding the Web: An Introduction to Cross-Site Scripting Attacks*. *IEEE Security & Privacy*.

Wang, Z., Lu, Z., 2017, Advanced persistent threats: Detection, analysis and countermeasures. CRC Press.



ÖZGEÇMİŞ

Adı Soyadı : F rkan TANYER 
Uyruęu : T rkiye Cumhuriyeti

EęİTİM

Derece	Adı, İlçe, İl	Bitirme Yılı
�niversite	: Batman �niversitesi, Batman	2019
Y�ksek Lisans	: Batman �niversitesi, Batman	2022
Y�ksek Lisans	: Batman �niversitesi, Batman	2024
Doktora	:	

İŐ DENEYİMLERİ

Yıl	Kurum	G�revi
2010 -Halen	Batman �niversitesi	Bilgi İŐlem Teknik G�revlisi

YABANCI DİLLER

İngilizce (Orta Seviye)

YAYINLAR