



**T.C.
İSTANBUL ÜNİVERSİTESİ-CERRAHPAŞA
LİSANSÜSTÜ EĞİTİM ENSTİTÜSÜ**



DOKTORA TEZİ

**AĞ GÜVENLİĞİ YÖNETİMİ İÇİN AKILLI AJANLAR TEKNOLOJİSİ
KULLANILARAK SALDIRI TESPİTİNE YÖNELİK YENİ BİR YAKLAŞIM**

Hakan AYDIN

DANIŞMAN

Doç. Dr. Muhammed Ali AYDIN

Bilgisayar Mühendisliği Anabilim Dalı

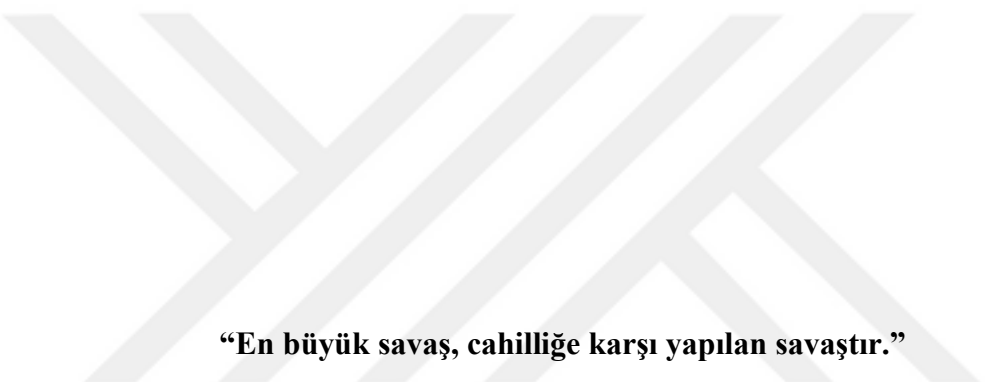
Bilgisayar Mühendisliği, Doktora Programı

Eylül, 2024

Bir öge seçin.







“En büyük savař, cahillięe karşı yapılan savařtır.”

Mustafa Kemal ATATÜRK

BÜTÇE DESTEKLERİ

AĞ GÜVENLİĞİ YÖNETİMİ İÇİN AKILLI AJANLAR TEKNOLOJİSİ KULLANILARAK SALDIRI TESPİTİNE YÖNELİK YENİ BİR YAKLAŞIM

Bu tez çalışması için herhangi bir kurumdan bütçe desteği alınmamıştır.

TEŞEKKÜR

Tez sürecinde bana yol gösteren çok kıymetli Hocam sayın Doç.Dr. Muhammed Ali AYDIN'a en içten duygularıyla teşekkür ederim. Ayrıca Tez düzenleme ve kontrol sürecinde destek olan Prof.Dr. Ahmet SERTBAŞ ve Prof.Dr.Selim AKYOKUŞ'a teşekkür ederim.

Hayatın her alanında olduğu gibi Doktora çalışmalarım sırasında da hep yanımda olan canım annem Meliha AYDIN, sevgili eşim Meral AYDIN, canım kızım Bengisu Eda AYDIN ve canım oğullarım Batuhan AYDIN ve Oğuzhan AYDIN'a destekleri için teşekkür ederim

Eylül 2024

Hakan AYDIN

İÇİNDEKİLER

Sayfa No

BEYAN	Hata! Yer işareti tanımlanmamış.
BÜTÇE DESTEKLERİ	v
TEŞEKKÜR.....	vi
İÇİNDEKİLER.....	vii
ŞEKİL LİSTESİ	ix
TABLO LİSTESİ.....	xi
SİMGE VE KISALTMA LİSTESİ.....	xii
ÖZET	xiv
ABSTRACT	xv
1. GİRİŞ.....	1
1.1. Araştırmanın Amacı.....	6
1.2. Katkılar	6
1.3. Araştırmanın Düzeni	7
2. KAVRAMSAL ÇERÇEVE	9
2.1. Ağ Güvenliği.....	9
2.1.1. Ağ Güvenliği Tanımı, Temelleri ve İlkeleri.....	9
2.1.2. Ağlara Yönelik Siber Tehdit ve Saldırıları.....	10
2.1.3. Ağ Güvenliği için Kullanılan Saldırı Tespit/Önleme Sistemleri (IDS/IDPS).....	20
2.2. Ajanlar ve Ağ Güvenliği.....	20
2.2.1. Ajan Kavramı ve Temel Özellikleri	20
2.2.2. Akıllı Ajanların Ağ Güvenliğinde Rolü ve Kullanımı	21
2.2.3. İlgili Çalışmalar Konusunda Literatür Taraması.....	22
3. YÖNTEM	28
3.1. Derin Öğrenme (DÖ) Tabanlı Model Önerisi.....	30
3.1.1. Modelin (LSTM-CLOUD) Tasarımı ve Yapılandırılması	32
3.2. Ajan Tabanlı Derin Öğrenme (DÖ) Tabanlı Model Önerisi.....	38
3.2.1. Nesnelerin İnterneti (IoT) Ağları için Model Önerisi	38
3.2.2. Bulut Altyapısı için Model Önerisi	52

3.3. Derin Pekiştirmeli Öğrenme (DRL) Tabanlı Model Önerisi	58
3.3.1. Modelin (NS-IoT) Tasarımı	58
3.3.2. Modelin (DQN-IoT) Yapılandırılması	62
4. BULGULAR	70
4.1. Derin Öğrenme (DÖ) Tabanlı Modelin Analizi	70
4.2. Ajan Tabanlı Derin Öğrenme (DÖ) Tabanlı Modelin Analizi	73
4.2.1. Bulut Altyapısı için Önerilen Modelin Analizi	73
4.2.2. IoT Ağları için Önerilen Modelin Analizi.....	81
4.3. Derin Pekiştirmeli Öğrenme (DRL) Tabanlı Modelin Analizi	86
4.3.1. CIC-IoT-2022 Veri Seti Üzerinde Önerilen Modelin Analizi	86
4.3.2. CIC-IoT-2023 Veri Seti Üzerinde Önerilen Modelin Analizi	88
4.4. Önerilen Modellerin Karşılaştırılmalı Analizi	93
4.4.1. Yöntem Türü ve Model Özellikleri Analizi	93
4.4.2. Performans Ölçütleri Analizi	97
4.4.3. Karmaşıklık Analizi	97
5. TARTIŞMA.....	99
5.1. Ajan Tabanlı olarak Siber Saldırı Tespitinde Otomatik Akıl Yürütme	106
5.2. Basitleştirilmiş Ödül Mekanizması.....	107
5.3. Ajan Tabanlı olarak Siber Saldırı Tespiti ve İnsan Ağ Operatörleri	108
5.4. Ajanların Bireysel Karar Verme Süreci	108
5.5. Kolektif Zekâ ve Karar Verme	108
5.6. Ajanların Diğer Varlıklarla İş birliği	109
5.7. Ajanların Gizliliği ve Dayanıklılığı	109
6. SONUÇ VE ÖNERİLER	111
KAYNAKLAR.....	117
İNTİHAL RAPORU İLK SAYFASI	123
ETİK KURUL İZİN YAZISI	Hata! Yer işareti tanımlanmamış.
KURUM İZNİ YAZILARI.....	Hata! Yer işareti tanımlanmamış.

ŞEKİL LİSTESİ

Sayfa No

Şekil 1.1: DDoS Saldırı Tahmini	2
Şekil 2.1: Genel Bulut Ağı Ortamı.....	12
Şekil 2.2: ABD Savunma Bakanlığı Kurumsal Bulutu.....	13
Şekil 2.3: DDoS Saldırı Vektörleri	15
Şekil 2.4: Tüketicilerin Toplam Cihaz Sayısı İçindeki Payı.....	17
Şekil 2.5: Endüstriyel Nesnelerin İnterneti (IIoT) Mimarisi Örneği.....	18
Şekil 3.1: Unutma Kapısına Sahip LSTM Mimarisi.....	30
Şekil 3.2: LSTM-CLOUD Mimarisi.....	34
Şekil 3.3: LSTM-CLOUD Sistemi Kullanım Senaryosu.....	35
Şekil 3.4: LSTM Tahmin Modeli.....	36
Şekil 3.5: CICDoS2019 Veri Seti	37
Şekil 3.6: Saldırı Türleri.....	37
Şekil 3.7: CICDoS2019 Veri Seti Ağ Trafiği	38
Şekil 3.8: MAS-IoT Sistem Mimarisi	39
Şekil 3.9: MAS-IoT Sistemi Ajanları	40
Şekil 3.10: Algılayıcı Ajan.....	41
Şekil 3.11: Tespit Ajanı	42
Şekil 3.12: Savunma Ajanı.....	43
Şekil 3.13: MAS-IoT Ajanların Kolektif Zekâsı	44
Şekil 3.14: Ajanlar Arası Şifreli Mesajlaşma Mimarisi	49
Şekil 3.15: MAS-IoT Sistemi Süreç Diyagramı	50
Şekil 3.16: Saldırı ve Normal Trafik Akışı	51

Şekil 3.17: Veri Seti Ağ Akış Dağılımı	52
Şekil 3.18: CDA-CLOUD Mimarisi	55
Şekil 3.19: CDA-CLOUD Sisteminin Kullanım Senaryosu	56
Şekil 3.20: CICDDoS2019 Veri Seti	57
Şekil 3.21: Hacimsel DDoS ve Normal Ağ Trafığı	57
Şekil 3.22: Ağ Trafik Akışı Dağılımı.....	58
Şekil 3.23: NS-IoT Sistem Mimarisi.....	59
Şekil 3.24: NS-IoT Sistemi Ajan Yapılanması	60
Şekil 3.25: NS-IoT Kullanımı Senaryosu	62
Şekil 3.26: DQN-IoT Ajanı Yapılanması	67
Şekil 4.1: Deney-2'ye ait Başarı ve Kayıp Dağılımları	71
Şekil 4.2: Deney-4 ve Deney-8 Başarı ve Kayıp Dağılımları.....	72
Şekil 4.3: Doğruluk Oranlarının Karşılaştırılması	72
Şekil 4.4: Doğruluk ve Kayıp Değerleri	76
Şekil 4.5: Deney-6 Başarı ve Kayıp Oranları	77
Şekil 4.6: Başarı Oranı En Yüksek Deneyler.....	78
Şekil 4.7: Deney 8 Doğruluk ve Kayıp Değerleri	82
Şekil 4.8: Deney-1 Doğruluk ve Kayıp Değerleri.....	82
Şekil 4.9: Deneylere İlişkin Doğruluk Değerleri	83
Şekil 4.10: Deneylere İlişkin Eğitim ve Test Süresi Değerleri (msn).....	83
Şekil 4.11: CIC-IoT-2023 Veri Seti	86
Şekil 4.12: CIC-IoT-2022 Veri Seti Akış Dağılımı	87
Şekil 4.13: 1'nci Grup Deneyler	88
Şekil 4.14: CIC-IoT-2023 Veri Seti Akış Dağılımı	89
Şekil 4.15: 2'nci Grup Deneyler	90
Şekil 4.16: Veri Seti Bazında En Yüksek Doğruluk Oranları.....	92
Şekil 4.17: Ajanların Performans Karşılaştırmaları	92

TABLO LİSTESİ

	Sayfa No
Tablo 3.1: Önerilen Modellerin Aşamalı Gelişim Yöntemi	29
Tablo 3.2: DQN Tabanlı Öğrenme Algoritması	65
Tablo 3.3: DQN-IoT Durum-Eylem-Ödül Matrisi	68
Tablo 4.1: Model Üzerinde Yapılan Deneylerin Analizi	70
Tablo 4.2: LSTM-IoT Deneysel Sonuçları	74
Tablo 4.3: Deneylere İlişkin Sonuçlar	81
Tablo 4.4: Model Üzerinde Yapılan Deneylerin Değerlendirilmesi	87
Tablo 4.5: Model Üzerinde Yapılan Deneylerin Değerlendirilmesi	89
Tablo 4.6: Deneylerde Kullanılan Veri Setleri Karşılaştırması	91
Tablo 4.7: Önerilen Yöntemlerin Karşılaştırmalı Analizi	96
Tablo 5.1: 1'inci Aşamanın İlgili Bazı Çalışmalar ile Karşılaştırılması	100
Tablo 5.2: 2'inci Aşamanın İlgili Bazı Çalışmalar ile Karşılaştırılması	101
Tablo 5.3: 3'üncü Aşamanın İlgili Bazı Çalışmalar ile Karşılaştırılması	103
Tablo 5.4: 4'üncü Aşamanın İlgili Bazı Çalışmalar ile Karşılaştırılması	104

SİMGE VE KISALTMA LİSTESİ

Kısaltmalar Açıklama

AICA: Otonom Akıllı Siber Savunma Ajanları

AES: Gelişmiş Şifreleme Standardı

CI: Kritik Altyapıların

CNN: Konvolüsyonel Sinir Ağları

CPS: Siber-Fiziksel Sistemler

DNN: Derin Sinir Ağları

DDoS: Dağıtılmış Hizmet Reddi

DDQN: Double Deep Q-Network

DRL: Derin Pekiştirmeli Öğrenme

DÖ: Derin Öğrenme

DQN: Derin Q-Ağları

IaaS: Altyapının Hizmet Olarak Sunulması

IIoT: Endüstriyel Nesnelerin İnterneti

ICS: Endüstriyel Kontrol Sistemlerinin

IDPS: İzinsiz Giriş Tespit ve Önleme Sistemleri

IDS: Saldırı Tespit Sistemleri

IoT: Nesnelerin İnterneti

İHA: İnsansız Hava Araçları

IT: Bilgi Teknolojileri

LSTM: Uzun Kısa Süreli Bellek Hafızası Sinir Ağları

MÖ: Makine Öğrenimi

MDP: Markov Karar Süreci

MES: Üretim Yürütme Sistemleri

NLP: Doğal Dil İşleme

NIST: Ulusal Standartlar ve Teknoloji Enstitüsü

OT: Operasyonel Teknoloji

PaaS: Platformun Hizmet Olarak Sunulması

PLC: Programlanabilir Lojik Denetleyici

RL: Pekiştirmeli Öğrenme

RNN: Yinelenebilir Sinir Ağları

SaaS: Yazılımların Hizmet Olarak Sunulması

SCADA: Denetleme ve Veri Toplama Sistemleri

SDN: Yazılım Tanımlı Ağ

SSL: Güvenli Yuva Katmanı

TLS: Aktarım Katmanı Güvenliği

WSN: Kablosuz Sensör Ağı

YZ: Yapay Zekâ

YSA: Yapay Sinir Ağları

ÖZET

[DOKTORA TEZİ]

[AĞ GÜVENLİĞİ YÖNETİMİ İÇİN AKILLI AJANLAR TEKNOLOJİSİ KULLANILARAK SALDIRI TESPİTİNE YÖNELİK YENİ BİR YAKLAŞIM]

[Hakan AYDIN]

İstanbul Üniversitesi-Cerrahpaşa

Lisansüstü Eğitim Enstitüsü

Bilgisayar Mühendisliği Anabilim Dalı

Bilgisayar Mühendisliği, Doktora Programı

[Danışman: Doç. Dr. Muhammed Ali AYDIN]

Bilgisayar ağ güvenliği için siber saldırıların tespit edilmesi ve önlenmesi hayati öneme haizdir. Bununla birlikte günümüzde siber saldırılar, geleneksel saldırı tespit sistemleri (IDS) ile önlenemeyecek kadar hızlı, karmaşık, Yapay Zekâ (YZ) destekli ve hatta otonom olarak gerçekleşebilmektedir. Bu nedenle otomatik akıl yürütebilen otonom IDS'lere ihtiyaç vardır. Bu bağlamda, akıllı ajan teknolojilerinden ağ güvenliği yönetiminde yararlanılması en umut verici yaklaşımlardan birisidir. Bu tez çalışmasında, ağ güvenliği yönetimi için akıllı ajan teknolojilerini kullanarak siber saldırı tespitine yönelik yenilikçi bir yaklaşım önerilmiştir. Çalışmada, dört aşamalı geliştirme yöntemiyle çeşitli özgün ajan tabanlı saldırı tespit modelleri geliştirilmiş, bu modellerin performans değerleri hesaplanarak yöntem türü ve model özellikleri, performans ölçütleri ve de karmaşıklık değerleri açısından karşılaştırılmış ve böylelikle her bir model analiz edilmiştir. Bu tez çalışmasında elde edilen sonuçlar, araştırmanın dördüncü ve son aşamasında önerilen Derin Q-Ağları (DQN) tabanlı özgün modelin (DQN-IoT) diğer önerilen özgün ajan tabanlı siber saldırı tespit modellerine kıyasla daha ön plana çıktığını ortaya koymaktadır. |

Eylül 2024 , [143] sayfa.

Anahtar kelimeler: [Ağ Güvenliği, Akıllı Ajanlar, Yapay Zekâ, Derin Q-Öğrenme]

ABSTRACT

Ph.D. THESIS

A NOVEL APPROACH FOR INTRUSION DETECTION USING INTELLIGENT AGENTS TECHNOLOGY FOR NETWORK SECURITY MANAGEMENT

Hakan AYDIN

İstanbul University-Cerrahpaşa

Institute of Graduate Studies

Department of Computer Engineering

Computer Engineering

Supervisor : Assoc. Prof. Dr. Muhammed Ali AYDIN

Detecting and preventing cyber attacks is of vital importance for computer network security. However, today's cyber attacks can occur so fast, complex, AI-supported, and even capable of being autonomous that they cannot be prevented by traditional Intrusion Detection Systems (IDS). Therefore, there is a need for autonomous IDSs that can perform automated reasoning. In this context, leveraging intelligent agent technologies for network security management is one of the most promising approaches. This thesis proposes an innovative approach to cyber attack detection using intelligent agent technologies for network security management. The study developed various unique agent-based attack detection models through a four-phase development method. The performance metrics of these models were calculated, and the methods, model features, performance criteria, and complexity values were compared, thereby analyzing each model. The results of this thesis indicate that the novel Deep Q-Networks (DQN)-based model (DQN-IoT) proposed in the fourth and final phase of the research stands out compared to other proposed unique agent-based cyber attack detection models.

September 2024, 143 pages.

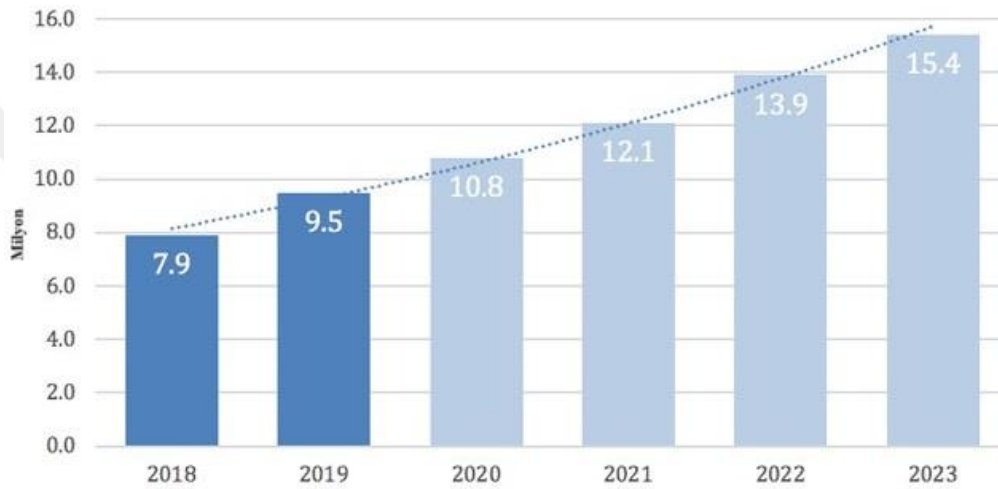
Keywords: Network Security, Intelligent Agents, Artificial Intelligence, Deep Q-Learning

1. GİRİŞ

Bilgisayar ağı, farklı bilgisayarların veya cihazların birbirine bağlanarak veri ve kaynak paylaşımı yapabildiği iletişim ağıdır. Bilgisayar ağ güvenliği, Bilgi ve İletişim Teknolojilerinden (BİT) faydalanarak ağ sistem ve kaynaklarının siber saldırılardan korunmasını amaçlayan hayati bir gerekliliktir. Ancak her geçen gün sayıları ve türleri artan bilgisayar ağlarını hedef alan siber saldırılar, ağ kaynakları ve sistemleri üzerindeki bilgilere yetkisiz erişim sağlamak, bu bilgilerin bütünlüğünü ihlal etmek veya erişimi engellemek amacıyla gerçekleştirilmektedir. Siber güvenlik, siber ortamı oluşturan bilgi sistemlerinin siber saldırılara karşı korunmasını ifade etmektedir (Von Solms ve Niekerk, 2013). İnternet ve bilgisayar sistemleri aracılığıyla hedeflenen sistemleri bozmak, siber terör saldırılarının amaçları arasında yer alır (Zanini ve Edwards, 2001). Siber savunma, siber saldırılara ve siber terörizme karşı önlem almak için uygulanan güvenlik yöntemleri olarak tanımlanabilir. Gizlilik, bütünlük ve erişilebilirlik (CIA), ağ güvenlik sistemlerinde üç temel güvenlik hizmetidir (Simmonds ve diğ., 2004). CIA üçlüsünde yer alan her bir kavram, siber güvenliğin farklı bir yönünü ele alan temel ilkeleri ifade eder. Gizlilik ilkesi, bilgiye ve sisteme yalnızca yetkili kişilerin erişebilmesini ifade eder. Bütünlük ilkesi, bilgilerin herhangi bir değişikliğe uğramamış, bozulmamış veya yok edilmemiş olmasını ifade eder. Erişilebilirlik ilkesi ise bilgiye ihtiyaç duyan kullanıcıların bu bilgilere erişebilmesi anlamına gelir. Bilgisayar ağı, geniş erişim imkânı ve karmaşık yapıları nedeniyle siber saldırılara karşı hassastır. Siber saldırılar, ağ üzerindeki veri bütünlüğünü, gizliliğini ve erişilebilirliğini ciddi şekilde tehdit edebilir. Bilgisayar ağlarına yönelik siber saldırılar büyük kesintilere, veri ihlallerine, dolandırıcılığa, sistem çökmelerine, ticari sırların çalınmasına, kontrollerin tehlikeye atılmasına ve daha fazlasına neden olabilir. Bu nedenle, ağ ortamlarındaki saldırıların tespit edilmesi ve durdurulması hayati önem taşımaktadır.

Günümüzde bilgisayar ağlarını hedef alan çok sayıda ve çeşitli siber saldırı türü bulunmaktadır. Siber saldırı türleri arasında Dağıtılmış Hizmet Reddi (DDoS) saldırıları, ağlara olası zararlı etkileri nedeniyle en yıkıcı ve şiddetli saldırılardan biri olarak öne çıkmaktadır. Özellikle DDoS saldırıları, çevrimiçi uygulama veya hizmetlerin çalışmasını engellemek amacıyla sistemlerin bant genişliğini kullanarak yanıt vermelerini engellemeyi hedefleyen saldırılardır (Boyd, 2009;

Johnson, 2015). DoS ve DDoS saldırıları, uygulamaları aşırı trafik veya isteklerle doldurmayı, kapasitelerini aşmayı ve bunları meşru kullanıcılar için erişilemez hale getirmeyi içerir (Mishra ve diğ., 2023). NetScout'un yayınladığı tehdit raporuna göre, 2019 yılında 8,4 milyon DDoS saldırısı gerçekleşmiş olup, bu da ortalama olarak ayda 670.000, günde 23.000 ve dakikada 16 saldırı anlamına gelmektedir (Global DDoS Tehdit Manzarası Raporu, 2019). Cisco'nun raporunda ise toplam DDoS saldırı sayısının 2018 yılında 7,9 milyondan 2023 yılına kadar 15 milyonun üzerine çıkacağı tahmin edilmektedir (Cisco Yıllık İnternet Raporu, 2018–2023) (Şekil 1.1).



Şekil 1.1: DDoS Saldırı Tahmini

Günümüzde bilgisayar ağ güvenliğinin sağlanabilmesi için bilgisayar ağlarını hedef alan siber saldırıların tespit edilmesi ve önlenmesi hayati önem taşımaktadır. Ancak, ağ savunma sistemleri, özellikle CIA ilkelerine bağlı kalarak ağ güvenliğini sağlamaya çalışırken çeşitli zorluklarla karşılaşabilmektedir. Bu zorluklardan en önemlisi, sürekli gelişen siber saldırı ve tehdit ortamının ağ güvenlik sistemleri tarafından dikkate alınması gerekliliğidir. Güvenlik duvarları, Saldırı Tespit Sistemleri (IDS) ve Saldırı Tespit ve Önleme Sistemleri (IDPS) gibi geleneksel ağ savunma mekanizmaları genellikle bilinen imzalara veya güncel olmayan kalıplara dayanmaktadır ve bu nedenle yeni ve bilinmeyen saldırılara karşı etkili bir savunma sunamayabilirler. Bu sebeple, ağ güvenlik sistemlerinin sürekli olarak gelişmesi ve yeni kötü amaçlı yazılımlara uyum sağlaması gerekmektedir. Bir ağdaki IDS'in temel rolü, bilgisayar sistemlerine yönelik ağ saldırılarını tespit etmek ve bu saldırılara müdahale etmek konusunda yardımcı olmaktır (Ashoor ve Gore, 2011). İmza tabanlı izinsiz giriş tespiti, izinsiz girişleri tanımlamak için faaliyetlerden toplanan verileri mevcut imzalarla eşleştirmeyi içerir (Sabahi ve

Movaghar, 2008). Anormallik tespiti ise daha önce bilinmeyen saldırıları keşfetmeyi amaçlar (Mirkovic ve Reiher, 2004). İmza tabanlı algılama, saldırıları tanımlamak için bilinen modellerden oluşan bir veri tabanı kullanırken, anormallik tabanlı algılama, trafik akışının özellikleri ve etiketlerine dayalı olarak trafiği iyi ya da kötü olarak sınıflandırmak için makine öğrenimini (MÖ) kullanır. Bu nedenle, ağ güvenlik yöneticilerinin hem bilinen hem de bilinmeyen saldırıları tespit edebilmesi için anormallik tespiti yaklaşımının kullanılması gerekmektedir (Biermann ve diğ., 2001). Mevcut geleneksel IDS'ler, ağ güvenliği gereksinimlerini yeterince karşılamamakta ve etkili izinsiz giriş tespiti için stratejik bir dağıtım yaklaşımı gerektirmektedir (Nuaimi ve diğ., 2023). Günümüzde bilgisayar ağ güvenliği için hem imza tabanlı hem de anomali tabanlı sistemlerin yaygın olarak kullanıldığı görülmektedir. Anormallik tespiti ise şu anda aktif bir araştırma alanıdır (Rodríguez ve diğ., 2023).

Ağ savunma mekanizmalarının bilgisayar ağ güvenliğini sağlamaya çalışırken karşılaşılabileceği diğer önemli zorluklardan birisi de gelişmiş akıllı siber saldırıların tespit edilmesi ve durdurulmasıdır. Günümüzde her geçen gün ortaya çıkan yeni siber saldırılar, makine hızlarında gerçekleşen, son derece karmaşık, Yapay Zekâ (YZ) destekli ve hatta otonom yeteneklere sahip olabilmektedirler. Siber saldırıların karmaşıklığı ve hızının artmasıyla birlikte, ağ savunma sistemlerinin bu saldırıları tespit etme ve durdurma yetenekleri de zorlaşmaktadır. Gelişen bilişim teknolojileri sayesinde siber saldırganlar artık daha geniş bir araç ve yetenek yelpazesine erişim sağlayabilmekte ve bu durum bilgisayar ağlarına yönelik YZ tabanlı, otonom ve ileri düzey tekniklerle gerçekleştirilen karmaşık saldırılara yol açabilmektedir. Gelişmiş düşmanca ortamlarda, saldırılar makine öğrenimi modellerindeki güvenlik açıklarından yararlanarak CIA üçlüsünü (Confidentiality, Integrity, Availability) tehlikeye atabilir (Biggio ve Roli, 2018). Siber saldırganların operasyonlarını YZ kullanarak geliştirmesi, savunucuların da benzer teknolojileri geliştirmelerini zorunlu kılar ve bu, insan karar verme sürecini azaltmaktadır (Musman ve diğ., 2019). Siber saldırganlar daha karmaşık hale geldikçe, savunucuların avantaj elde etmek için geleneksel yaklaşımlarını proaktif önlemlerle güçlendirmesi gerekmektedir (Eghtesad ve diğ., 2020). Bahse konu bu durum DDoS saldırıları için de geçerlidir. YZ, hedef sistemlerin zayıf noktalarını hızlı bir şekilde tespit ederek ve saldırıları optimize ederek DDoS saldırılarını daha otonom ve uyarlanabilir hale getirebilir. Bu, DDoS saldırılarının hızını ve ölçeğini artırarak savunma sistemleri için önemli zorluklar oluşturabilir. YZ tabanlı DDoS saldırıları genellikle ağ trafiğini sürekli olarak analiz ederek en etkili saldırı zamanını ve yöntemini belirler. Ayrıca, daha karmaşık ve tespit edilmesi

zor saldırılar oluşturmak için hacimsel, protokol tabanlı ve uygulama katmanı saldırıları gibi farklı DDoS saldırı türlerini birleştirebilir. Tüm bunların yanında gerçek zamanlı siber saldırı senaryolarında, DDoS saldırılarının tespit hızı ve bu saldırılara müdahale etme hızı kritik öneme sahiptir. Mevcut siber müdahale uygulamaları genellikle önceden tanımlanmış tepkilere, kararlara ve refleksif yanıtlara dayanır ki bu durum da potansiyel olarak saldırının etkisini artırabilir (Zonouz ve diğ., 2013). Geleneksel savunma sistemleri, yetenekli insan ağ operatörleri olsa bile, bu tür gelişmiş saldırıları tespit etmek ve önlemek için yeterli olamayabilir. Dolayısıyla, insan ağ operatörlerinin hayati önemi devam etse bile, savunma amaçlı YZ tabanlı ağ savunma sistemlerinin, özellikle bu koşullar altında gelişmiş siber saldırılara karşı önemli bir özerkliğe sahip olması gerekir. Yani, ağ savunma sistemleri, otonom adaptasyon ve tepki verebilme yeteneğine sahip YZ tabanlı teknolojilerle entegre edilmelidir. Otonom çıkarım ve karar verme yeteneğine dayalı bu ağ savunma yaklaşımı, insan müdahalesine olan bağımlılığı azaltırken yeni ortaya çıkan siber tehditlere karşı dayanıklılığı artıran bir yaklaşımı içerir. Günümüzde, gelişen siber tehditlere otonom olarak uyum sağlayabilen, insan müdahalesini en aza indirirken yeni saldırı vektörlerine karşı dayanıklılığı artıran otomatik ve akıllı izinsiz giriş tespit çözümlerine ihtiyaç duyulmaktadır (Alavizadeh ve diğ., 2022). Siber savunma alanında tespit yöntemlerini geliştirmek, muhakeme ve otomasyon yoluyla güvenlik ekibinin etkinliğini artırmak ve YZ destekli saldırılara etkili bir şekilde karşı koymak için YZ yeteneklerinden yararlanılmalıdır (Stoecklin ve diğ., 2018). Bu bağlamda, ağ güvenliği yönetimi için akıllı ajan teknolojilerinden yararlanılması umut verici bir çözümdür. Akıllı ajan teknolojisinin yer aldığı bu siber güvenlik yaklaşımı siber saldırıların karmaşıklığını yönetirken otomatik tepki ve adaptasyon yetenekleriyle savunma kapasitesini artırma, insan müdahalesine olan bağımlılığı azaltma ve siber güvenliği güçlendirme imkânı sağlayabilme imkân ve kabiliyetine sahiptir.

Russell ve Norvig (2009), ajanları çevresiyle etkileşim halinde olan ve belirli hedeflere ulaşmak için kendi başına karar alabilen varlıklar olarak tanımlamışlardır. Ajanlar, bir ortama entegre olmuş ve hedeflerine ulaşmak ile gelecekteki sonuçları etkilemek için çevresini algılayıp ona göre hareket eden özerk sistemlerdir (Franklin ve Graesser, 1996). Bir ajanın gerçekten zeki olabilmesi için, içsel bilgi ve fayda işlevini edinilen deneyime göre uyarlaması ve ondan öğrenmesi gerekir (Guarino, 2013). Literatürde, savunma amaçlı geliştirilen akıllı ajan teknolojilerinin ağ güvenliğinde de kullanıldığı görülmektedir. Bir siber ajan, ağ trafiğini izleme, siber tehditleri belirleme ve dijital ortamda meydana gelen siber olaylara hızla yanıt

verme konularında uzmanlaşmış özerk bir yazılım varlığıdır (Kott, 2018). Bu tür otonom yazılım varlıkları, minimum insan müdahalesi ile gelişmiş saldırılara karşı koymak için umut verici çözümler sunmaktadır (Dutta ve diğ., 2021). Gelecekte, özellikle insan müdahalesinin pratik olmadığı durumlarda, Otonom Akıllı Siber Savunma Ajanları (AICA) gibi YZ ajanları, kötü amaçlı yazılımların ortadan kaldırılmasında önemli bir rol oynayacaktır (Kott ve Theron, 2020). Bu bağlamda, NATO'nun Bilim ve Teknoloji Örgütü'nün Araştırma Grubu IST-152 tarafından yürütülen güncel araştırmalar, ajan teknolojilerinin siber savunmada nasıl kullanılabileceğine dair devam eden çabaları ortaya koymaktadır. Uluslararası bu ekip ayrıca, AICA'nın vizyonunu özetleyen ve üst düzey bir referans mimarisi sunan bir rapor da yayımlamıştır (Kott, 2018). Akıllı ajanlardan siber güvenlik bağlamında yararlanmayı amaçlayan bu gelişmeler, yenilikçi teknolojiler aracılığıyla ağ güvenliğini sağlama çabalarını vurgulamaktadır.

Bu tez çalışmasında, ağ güvenliği yönetimi için akıllı ajan teknolojilerinden yararlanarak saldırı tespitine yönelik yeni bir yaklaşım olarak ajan tabanlı bir siber saldırı tespit modeli önerilmiştir. Bu amaç doğrultusunda çalışmada aşamalı geliştirme yöntemiyle ajan tabanlı farklı saldırı tespit modelleri önerilmiş, önerilen modellerin performans değerleri hesaplanmış ve bu değerler birbirleriyle ve literatürdeki diğer çalışmalarla karşılaştırılmıştır. Karşılaştırma neticesinde, çalışmanın son aşamasında Derin Takviyeli Öğrenme (DRL) kullanılarak Derin Q-Ağları (DQN) tabanlı olarak geliştirilen modelin (DQN-IoT) sahip olduğu yenilikçi yaklaşımıyla diğer modellere kıyasla ön plana çıktığı tespit edilmiştir. Önerilen bu yenilikçi yaklaşım, DDoS saldırılarını sınıflandırma problemini, doğru tahminler için olumlu ödüller ve yanlış tahminler için olumsuz ödüller verilen bir tahmin oyunu şeklinde ele almaktadır. DQN-IoT ajanı, eylemlerinden ve çevresel geri bildirimlerden öğrenerek zamanla karar verme sürecini geliştirebilir ve amacını birikimli ödülleri maksimize etmek ve sınıflandırma doğruluğunu artırmak olarak belirleyebilir. Önerilen bu yenilikçi yaklaşımda Markov Karar Süreci (MDP), DQN-IoT ajanının eylem seçimlerini yönlendirmek için her adımda çevresel durumu değerlendirmektedir. Esasen bu süreç, mevcut çevresel durumu göz önünde bulundurarak ajanın bir sonraki adımda atması gereken en uygun eylemi belirlemeyi içermektedir. MDP kullanarak, DQN-IoT ajanı, DDoS saldırılarını tespit etme ve hafifletme bağlamında ödülleri en yüksek seviyeye çıkarmayı veya önceden tanımlanmış hedeflere ulaşmayı amaçlamakta ve böylece bilinçli kararlar verebilmektedir. Siber savunma amaçlı akıllı ajan tabanlı bir yapıda önerilen bu yenilikçi ağ güvenliği yaklaşımı, DQN-IoT ajanının siber tehditlere karşı savunma

etkinliğini artırarak öğrenme ve uyum sürecini zamanla optimize etmeye yardımcı olmaktadır. Bu tez çalışmasında elde edilen sonuçlar, bu çalışmada önerilen DQN-IoT modelinin siber saldırıların tespit edilmesi ve durdurulmasındaki etkinliğini göstermektedir.

1.1. Araştırmanın Amacı

Bu çalışmanın amacı, "Geleneksel bir ağ ortamında ağ güvenliği yönetimi için siber atakların tespiti maksadıyla otonom akıllı siber savunma ajanları kullanılarak ajan tabanlı özgün bir saldırı tespit sisteminin geliştirilmesi" olarak belirlenmiştir.

Araştırmanın yanıtlamaya çalıştığı iki temel sorusu bulunmaktadır:

- Bunlardan birincisi, akıllı ajanların ağ güvenliği için saldırı tespit sistemlerinde kullanımına ilişkin literatürde güncel hangi yaklaşım, araştırma ve çalışmaların olduğuna ilişkindir.
- Araştırmanın ikinci temel sorusu ise ajan tabanlı özgün ve yeni bir siber saldırı tespit sisteminin tasarımı ve geliştirilmesine ilişkindir.

1.2. Katkılar

Çalışmada elde edilen katkılar aşağıda özetlenmiştir:

- Derin Öğrenme (DÖ) tabanlı LSTM ağı kullanılarak gerçekleştirilen çalışmalar kapsamında, CICDDoS2019 veri seti ile %99,83 başarı oranına sahip imza tabanlı bir saldırı tespit sistemi (LSTM-CLOUD) önerilmiştir. Önerilen LSTM-CLOUD modelinin performansı, literatürde mevcut benzer yöntemlerle karşılaştırılmıştır.
- Ajan tabanlı olarak DÖ'den yararlanılarak LSTM ağı ile yapılan çalışmalar kapsamında, IoT ağlarını hedef alan DDoS saldırılarına karşı çoklu ajan tabanlı bir siber savunma sistemi (MAS-IoT) önerilmiştir. Bu savunma sisteminde kullanılmak üzere CIC-IoT-2022 veri seti üzerinde %99,48 doğruluk oranına sahip bir LSTM modeli (LSTM-IoT) tasarlanmıştır. Önerilen LSTM-IoT modelinin performansı, literatürde mevcut benzer yöntemlerle karşılaştırılmıştır.
- Ajan tabanlı LSTM ağı ile gerçekleştirilen çalışmalar kapsamında, bulut ağlarını hedef alan hacimsel DDoS saldırılarının tespit etmek için çoklu ajan tabanlı bir siber savunma

sistemi (CDA-CLOUD) önerilmiştir. Bu savunma sisteminde kullanılmak üzere CICDDoS2019 veri seti üzerinde %99,89 doğruluk oranına sahip bir LSTM modeli (CDA-CLOUD) tasarlanmıştır. Önerilen CDA-CLOUD modelinin performansı, literatürde mevcut benzer yöntemlerle karşılaştırılmıştır.

- Ajan tabanlı Derin Pekiştirmeli Öğrenme (DRL) DQN ağı ile yapılan çalışmalar kapsamında, Endüstriyel Nesnelerin İnterneti (IIoT) güvenliği için çoklu ajanlı DQN tabanlı bir siber savunma sistemi (NS-IIoT) önerilmiş ve DRL teknikleri ile bir DQN tabanlı ajan (DQN-IIoT) geliştirilmiştir. Ayrıca, CR-IIoT ve CM-IIoT ajanları tasarlanmıştır. DQN-IIoT ajanının etkinliği, CIC-IIoT-2022 ve CIC-IIoT-2023 veri setleri kullanılarak test edilmiştir. Önerilen CDA-CLOUD modelinin performansı, literatürde mevcut benzer yöntemlerle karşılaştırılmıştır. Karşılaştırma neticesinde, çalışmanın son aşaması olan dördüncü aşamasında DRL kullanılarak DQN tabanlı olarak önerilen özgün ajan tabanlı DQN-IIoT modelinin “yöntem türü ve model özellikleri”, “performans ölçütleri” ile “karmaşıklık değerleri” bağlamında çalışmada önerilen diğer özgün saldırı tespiti modellerine kıyasla daha ön plana çıktığını göstermektedir.

Bu katkılar ve sonuçlar ışığında, ağ güvenliği yönetimi için akıllı ajan teknolojileri kullanılarak saldırı tespitine yönelik yeni bir yaklaşım (DQN-IIoT) önerilmiştir.

1.3. Araştırmanın Düzeni

Tez çalışmasının 2. bölümünde, ağ güvenliği ile ilgili temel kavramlar, ilkeler ve tanımlar ele alınmıştır. Bu bölümde, ağ güvenliği kapsamında ağların karşılaştığı siber tehditler ve saldırı türleri detaylı bir şekilde açıklanmış ve ağ güvenliği için kullanılan saldırı tespit ve önleme sistemleri incelenmiştir. Ayrıca, “Ajanlar ve Ağ Güvenliği” alt başlığında, ajan kavramı ve temel özellikleri tanıtılmış, akıllı ajanların ağ güvenliğindeki rolü ve kullanım şekilleri açıklanmıştır. Bölümün sonunda, literatür taraması yapılarak mevcut araştırmalar ve teoriler özetlenmiş ve araştırma konusunun mevcut bilgi birikimi içindeki yeri belirlenmiştir.

Tez çalışmasının 3. bölümünde, siber saldırıların tespit edilmesine yönelik ajan tabanlı çeşitli model önerileri ayrıntılı bir şekilde sunulmuştur. Bu bölümde, ilk olarak “Derin Öğrenme (DÖ) Tabanlı Model Önerisi” başlığı altında önerilen modelin tasarımı ve yapılandırması detaylandırılmıştır. Ayrıca, önerilen modelin LSTM-CLOUD olarak adlandırılan spesifik bir versiyonu üzerinde durulmuştur. Sonraki kısımda, “Ajan Tabanlı Derin Öğrenme (DÖ) Tabanlı

Model Önerisi” başlığı altında, Nesnelerin İnterneti (IoT) ağları için önerilen MAS-IoT ve LSTM-IoT modellerinin tasarımı ve yapılandırılması ele alınmıştır. Ayrıca, bulut altyapısı için CDA-CLOUD ve LSTM-CLOUD modellerinin tasarım ve yapılandırmaları açıklanmıştır. Bölümün sonunda, “Derin Pekiştirmeli Öğrenme (DRL) Tabanlı Model Önerisi” başlığı altında NS-IoT ve DQN-IoT modellerinin tasarımı ve yapılandırması detaylandırılmıştır.

Tez çalışmasının 4. bölümünde, DÖ tabanlı modelin analizi yapılmış, ajan tabanlı DÖ modelleri için elde edilen bulgular detaylandırılmış ve özellikle bulut altyapısı ile IoT ağları için önerilen modellerin analizleri gerçekleştirilmiştir. Ayrıca, DRL tabanlı modelin analizi de yer almıştır. Bu bölümdeki veriler, önerilen modellerin etkinliğini ve performansını ortaya koymak amacıyla karşılaştırmalı bir şekilde analiz edilmiştir.

Tez çalışmasının 5. bölümünde “Tartışma” kısmı yer almaktadır. Bu bölümde, elde edilen bulgular yorumlanmış ve araştırma kapsamında belirtilen tüm model önerilerinin güçlü ve zayıf yönleri ile elde edilen sonuçlar genel bağlamda detaylandırılmıştır. “Tartışma” bölümü, araştırma bulgularının literatür ve mevcut teorilerle nasıl ilişkilendiğini tartışarak derinlemesine bir analiz sunmaktadır.

Tez çalışmasının 6. bölümünde “Sonuç ve Öneriler” bölümü bulunmaktadır. Bu bölümde, araştırmanın genel sonuçları özetlenmiş ve gelecekteki çalışmalar için çeşitli önerilerde bulunulmuştur. Bölümde yer verilen öneriler, hem mevcut araştırmanın sınırlamalarını göz önünde bulundurarak hem de gelecekteki geliştirmeler için bir yol haritası sunma amacını taşımaktadır.

2. KAVRAMSAL ÇERÇEVE

2.1. Ağ Güvenliği

2.1.1. Ağ Güvenliği Tanımı, Temelleri ve İlkeleri

Ağ güvenliği, bir bilgisayar ağı içindeki verilerin, sistemlerin ve kullanıcıların güvenliğini sağlamak amacıyla uygulanan yöntemler, politikalar ve teknolojiler bütünüdür. Temel hedefi, ağa bağlı kaynakları, verileri ve bilgileri yetkisiz erişim, saldırılar, veri kaybı veya diğer zarar verici etkilerden korumaktır. Ağ güvenliğinin amacı, bireylerin bilgisayar ağlarını hak ve çıkarlarını tehlikeye atmadan özgürce kullanabilmesini sağlamak için ağ sistemlerini ve elektronik verileri korumaktır (Wang, 2009). Bu bağlamda, İnternet ağı milyonlarca bilgisayarı birbirine bağlayan ve herkesin kolayca erişim sağlayabildiği geniş bir ağ sistemidir. Ağ güvenliği, çeşitli tehditleri ve riskleri önlemek için belirli ilkeler ve teknikler kullanır. Bu ilkeler arasında temel olarak gizlilik, bütünlük ve erişilebilirlik sayılabilir:

- **Gizlilik (Confidentiality):** Verilerin yalnızca yetkili kişiler tarafından erişilmesini sağlamayı amaçlar. Gizlilik, verilerin yetkisiz kişiler tarafından görüntülenmesini, ele geçirilmesini veya paylaşılmasını önler. Şifreleme, veri maskeleyme ve erişim kontrol mekanizmaları bu ilkeye hizmet eden yöntemlerdir.
- **Bütünlük (Integrity):** Verilerin doğruluğunu ve tutarlılığını korur. Bu ilke, verilerin yetkisiz değişikliklerden, bozulmalardan veya hatalardan korunmasını sağlar. Veri bütünlüğünü sağlamak için kriptografik hash fonksiyonları, dijital imzalar ve veri bütünlüğü denetimleri gibi teknikler kullanılır.
- **Erişilebilirlik (Availability):** Yetkili kullanıcıların gerekli bilgilere ve sistemlere herhangi bir kesinti olmaksızın zamanında erişimini garanti eder. Bu ilke, hizmet kesintilerine, siber saldırılara ve diğer tehditlere karşı ağı ve kaynakların sürekli olarak erişilebilir olmasını sağlamaya odaklanır. Yedekleme, felaket kurtarma planları ve sistem performans izleme bu ilkeye yardımcı olur.
- **Kimlik Doğrulama (Authentication):** Bir kullanıcının veya sistemin iddia ettiği kimliğini doğrulamak için kullanılan süreçtir. Bu ilke, yalnızca yetkili kişilerin veya sistemlerin ağa erişimini ve işlemleri gerçekleştirmesini sağlar. Parolalar, biyometrik

veriler ve iki faktörlü kimlik doğrulama gibi yöntemler kimlik doğrulamanın araçlarındandır.

- **Yetkilendirme (Authorization):** Bir kullanıcının veya sistemin belirli kaynaklara ve verilere erişim yetkisini belirler. Bu ilke, kimlik doğrulama sürecinden sonra gelir ve kullanıcılara veya sistemlere belirli izinler atar.

2.1.2. Ağlara Yönelik Siber Tehdit ve Saldırıları

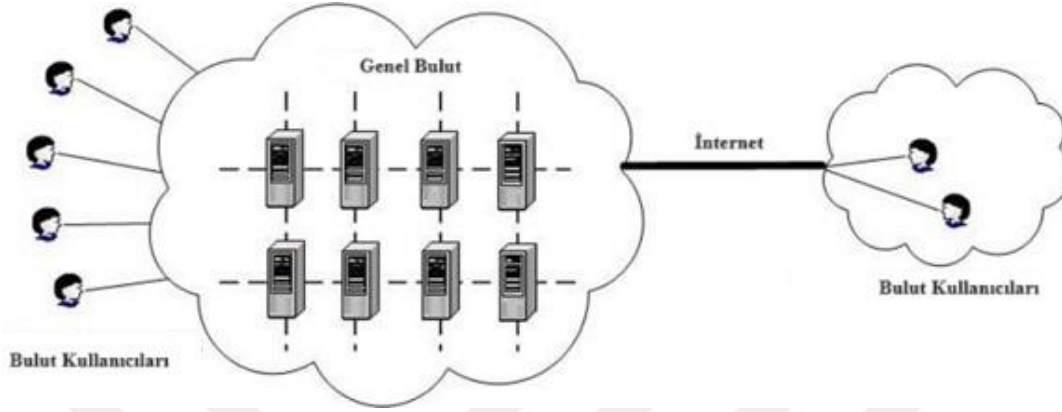
Ağlara yönelik siber tehditler ve saldırılar, günümüz dijital dünyasında sürekli olarak evrim geçiren ve giderek daha karmaşık hale gelen bir risk alanını temsil etmektedir. Bu tehditler, bireysel kullanıcıları, işletmeleri ve devlet kurumlarını hedef alarak ciddi güvenlik açıklarına yol açabilir. Özellikle DDoS saldırıları, bilgisayar ağlarını istenmeyen trafikle boğarak ağ kaynaklarını kullanıcılarına kapatabilir. Ağ sistemlerine yönelik siber saldırıların artmasıyla birlikte, bu sistemlerin güvenliği oldukça kritik bir hale gelmiştir. Siber saldırılar genel olarak sel saldırıları ve mantıksal saldırılar olarak sınıflandırılabilir (Srivastava ve diğ., 2011). DDoS saldırıları ise taşkın saldırıları, protokol saldırıları ve uygulama katmanı saldırıları olmak üzere üç temel kategoride incelenebilir (Singh ve diğ., 2017). Bu tür saldırılar genellikle kullanılabilirlik bazlı saldırılar arasında sayılmaktadır (Fakeeh, 2016). DDoS saldırılarının temel amacı, ağ iletişimini engelleyerek veya hizmetlere erişimi keserek kaynakları tüketmektir (Darwish ve diğ., 2013). DDoS saldırılarında, saldırı hedefinin tüm kaynaklarının tüketilmesi ve devre dışı bırakılması amacıyla İletim Kontrol Protokolü/İnternet Protokolü (TCP/IP) protokolünün açıklarından yararlanılarak düzenli olarak alt saldırılar yapılır. Bu saldırıların amacı, bilgi sistemlerine aşırı yük bindirerek temel hizmetleri aksatmak ve zayıflamış sistemlere sızmadır. Disk belleğini doldurmak, işlemci gibi kaynakları tüketmek, yüksek ağ trafiği oluşturmak ve yetkisiz erişim sağlamak bu tür saldırılara örnek teşkil eder. Sel saldırıları, bulut sistemlerinin kullanılabilirliğini, gizliliğini ve bütünlüğünü etkileyen yaygın saldırılar arasındadır (Modi ve diğ., 2013). DDoS saldırıları önemli miktarda para, zaman, prestij, değerli bilgi ve hatta insan yaşamına zarar verebilir. Bu nedenle, bilgisayar ağlarını hedef alan DDoS saldırılarının tespit edilmesi ve önlenmesi, ağ güvenliğinin sağlanması açısından hayati öneme sahiptir. Ağ sistemlerinde bilgi güvenliğinin korunabilmesi için siber saldırıların doğru ve zamanında tespit edilmesi, etkili bir şekilde önlenmesi, azaltılması veya durdurulması gerekmektedir. Bu bağlamda, siber saldırılara karşı önlem alabilecek ve bu saldırıları öngörüp önleyebilecek IDS'lere büyük bir ihtiyaç duyulmaktadır.

2.1.2.1. Bulut Ağları ve DDoS Saldırıları

Bulut bilişim, yüksek derecede yapılandırılabilir bilgi işlem kaynaklarından oluşan ortak bir havuza uygun ve isteğe bağlı ağ erişimi sunan bir paradigmadır (NIST, 2011). Bu model, işletmelere esneklik ve verimlilik sağlarken aynı zamanda maliyet tasarrufu, ölçeklenebilirlik, kolay erişim ve hızlı devreye alma gibi birçok avantaj sunar (Buyya ve diğ., 2010). Özellikle donanım ve yazılım giderlerini azaltma konusundaki ekonomik faydası, bulut bilişimin tercih edilmesinin başlıca nedenlerinden biridir (Erl ve diğ., 2013). Bulut kullanıcıları, yazılım, platform veya altyapı bulut hizmetlerini kullanan bireyler veya kuruluşlardır. Bulut hizmet sağlayıcıları, bu hizmetleri planlayıp üreten ve sunan birimlerdir. Bulut denetçisi ise, bulut ortamında sağlanan bilgi hizmetlerini güvenlik ve performans açısından yöneten birimdir. Kullanıcılar, doğrudan hizmet sağlayıcılarla değil, aracı bulut komisyoncuları aracılığıyla bulut hizmetlerinden yararlanabilirler. Bulut taşıyıcı, bulut bilişim hizmetlerinin sağlandığı teknolojik bilgisayar ağ altyapısını ifade eder. İnternet, bulut kullanıcılarına bilgi işlem kaynaklarına ve bulut hizmetlerine erişim sağlayan küresel bir merkez oluşturur. Bulut taşıyıcısı, bulut hizmetlerinin bağlantısından ve taşınmasından sorumlu bir katmandır (Liu ve diğ., 2011). Artan hesaplama gücü ihtiyacı, ölçeklenebilir hesaplama üzerindeki ilerlemeleri dağıtılmış hesaplama yöneltmiştir (Chetty ve Buyya, 2002). Bulut bilişim, sanallaştırma ve grid bilişim gibi mevcut teknolojiler üzerine inşa edilen bir modeldir. Hesaplamalı ızgaralar, coğrafi olarak dağıtılmış kaynakları birleştirerek bilgi işlem gücü sunmayı amaçlar (Abramson ve diğ., 2002). Bulut sistemleri genellikle bilgisayar ağlarına bağımlıdır. İsteğe bağlı self-servis, geniş ağ erişimi, kaynak havuzu, hızlı esneklik ve ölçülebilir hizmetler, bulut bilişimin temel özelliklerindendir.

Bulut hizmetleri üç ana modelde sunulur: Uygulama Hizmeti Olarak (SaaS), Altyapı Hizmeti Olarak (IaaS) ve Platform Hizmeti Olarak (PaaS). SaaS modelinde, sunucu, bellek, veri tabanı, yedekleme ve içerik dağıtım ağı gibi bilgi hizmetleri sağlanır. PaaS modelinde, bilgisayar altyapısı üzerinde bilgi hizmetleri verilirken, IaaS modelinde bulut altyapısı üzerinde çalışan yazılımlar sunulur. IaaS modeli, kullanıcıların yazılım kurulumu, bakımı ve lisanslanması gibi işlemlerle uğraşmalarına gerek kalmadan maliyet avantajı sağlar. Bulut bilişim katmanlarında sunulan SaaS, PaaS ve IaaS hizmetleri incelendiğinde, PaaS'ın genellikle sistem yöneticilerine, SaaS'ın uygulama geliştiricilere ve IaaS'ın ise son kullanıcılara hizmet verdiği görülür. Bulut bilişimin dört ana dağıtım modeli vardır: Genel, Özel, Topluluk ve Hibrit Bulut Dağıtım Modelleri (Goyal, 2014). Genel bulut yapısında, bilgi hizmetleri geniş bir ağ üzerinden halka

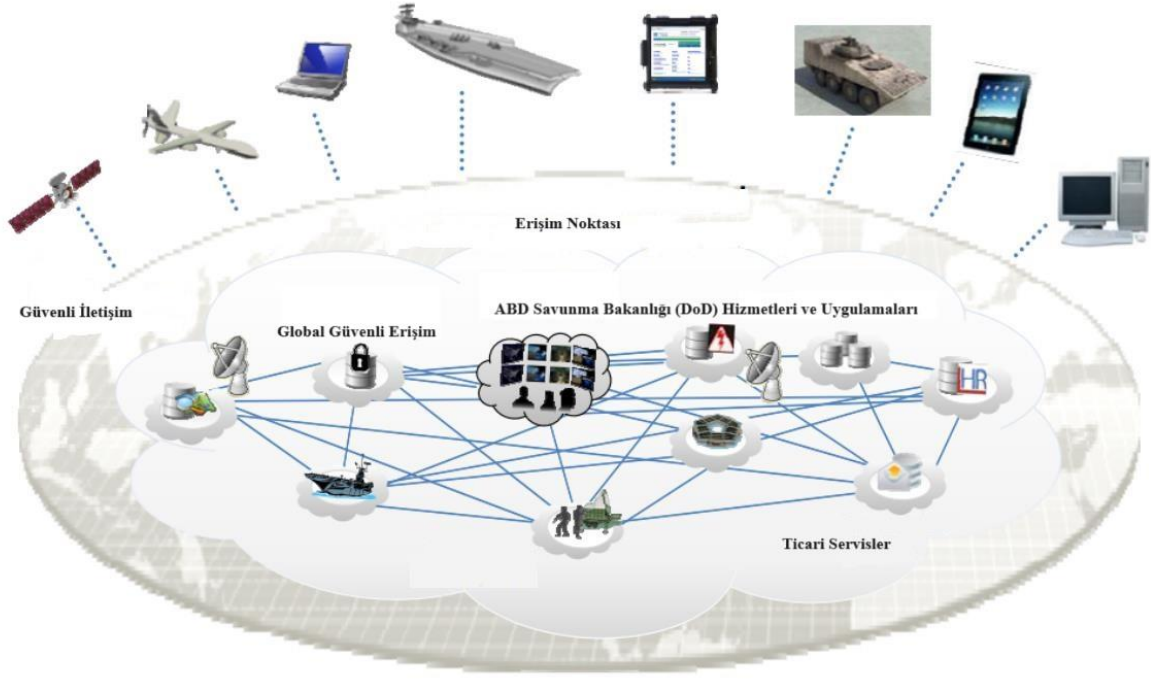
açık olarak sunulur ve kullanıcılar bu ağ aracılığıyla çeşitli hizmetlere erişebilirler. **Şekil 2.1**, genel bulut yapısının ve müşterilerinin basitleştirilmiş bir görünümünü sunmaktadır (Liu ve diğ., 2011).



Şekil 2.1: Genel Bulut Ağı Ortamı

Ulusal Standartlar ve Teknoloji Enstitüsü (NIST), bulut bilişimi, minimum yönetim çabasıyla veya hiç yönetim gerektirmeden hızlı bir şekilde erişilebilen ve kullanılabilen yapılandırılabilir bilgi işlem kaynakları olarak tanımlamaktadır (Mell ve Grance, 2011). Bulut bilişim hem bir platform hem de bir uygulama biçimi olarak kabul edilebilir (Boss ve diğ., 2007). Bu teknoloji, sanallaştırma ve dağıtılmış bilgi işlem gibi yaklaşımlardan yararlanarak hizmetlerin İnternet üzerinden barındırılmasını ve sunulmasını sağlayan yeni bir paradigma olarak ortaya çıkmıştır (Liao ve diğ., 2013). Bulut bilişim, zamanla uygun maliyetli ve ölçeklenebilir kurumsal bilgi hizmetlerinin sunulmasında önemli bir trend haline gelmiştir (Wang ve diğ., 2015). Küme, ızgara ve bulut bilişim teknolojileri, büyük miktarda bilgi işlem gücüne sanallaştırılmış bir şekilde erişim sağlamayı amaçlamaktadır (Buyya ve diğ., 2010). Bulut bilişimin temel hedeflerinden biri, daha uygun maliyetli bilişim altyapıları sunmaktır (Dudash, 2016). Bu teknolojinin önemli özelliklerinden biri, kullanıcıların pahalı bilişim altyapılarına yatırım yapma ihtiyacını ortadan kaldırması ve hizmetlerin maliyet etkinliğidir (Singh ve Chatterjee, 2017). Bulut teknolojileri, bilgi hizmetlerinin edinilmesinden sunulmasına kadar pek çok alanda uygulama bulmaktadır. Bulut bilişim, e-iş ve sosyal ağlar gibi çeşitli alanlarda önemli bir rol oynamaktadır (Boss ve diğ., 2007). Son yıllarda bulut bilişime geçiş hızlanmış; birçok kuruluş, kritik sistemler dahil olmak üzere altyapılarını doğrudan kontrollerinden çıkarıp bulut hizmet sağlayıcıları tarafından sunulan ve desteklenen platformlara taşımıştır (Haber ve diğ., 2022). İşletmelerin bulut altyapısına geçişinin temel motivasyonu maliyet etkinliğidir (Somani ve diğ., 2017). **Şekil 2.2**'de, DoD kurumsal bulut ortamının mantıksal bir tasviri sunulmakta ve

bu ortamın karmaşıklığı göz önüne alındığında bulut tabanlı güvenlik gereksinimlerinin önemi vurgulanmaktadır (Takai, 2012).



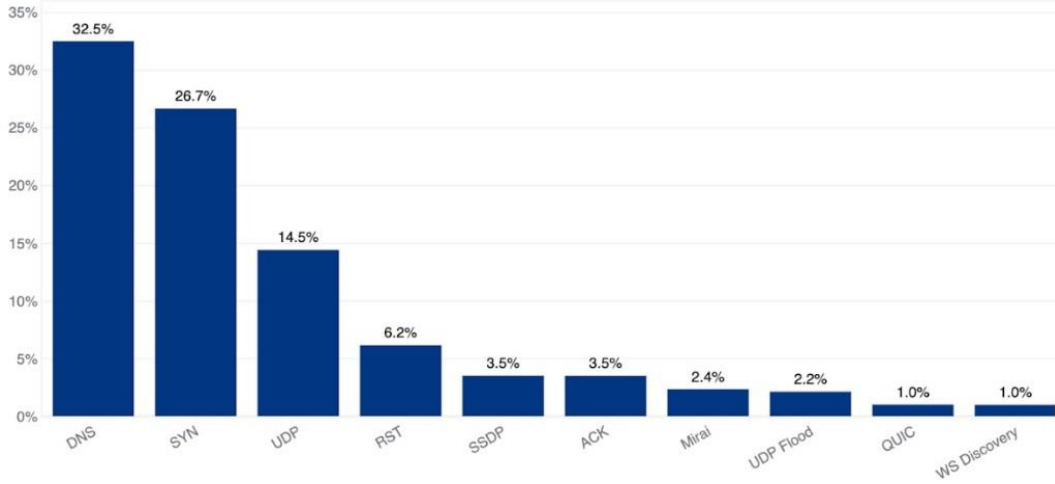
Şekil 2.2: ABD Savunma Bakanlığı Kurumsal Bulutu

Bulut tabanlı çözümlere geçiş, kuruluşları yeni siber güvenlik sorunlarıyla karşı karşıya bırakmış ve siber saldırılar büyük bir endişe kaynağı haline gelmiştir. Genel bulut bilişimdeki güvenlik ve gizlilik sorunları, sistem karmaşıklığı, paylaşılan çok kiracılı ortamlar, internet üzerindeki hizmetlere maruz kalma ve veriler üzerindeki kontrol kaybını içermektedir. Bu faktörler, güvenlik açıkları, yetkisiz erişim ve azalan kontrol ve hesap verebilirlik ile ilgili endişelere yol açmaktadır (Jansen ve Grance, 2011). Ayrıca, bulut bilişimin dezavantajlarından biri, internet bağlantısının kesilmesi durumunda veri gizliliği ve güvenliği ile ilgili endişeler yanı sıra erişim sorunlarına da yol açma potansiyelidir (Buyya ve diğ., 2010). Diğer dezavantajlar arasında ise güvenlik kaygıları ve bağımlılık riskleri öne çıkmaktadır (Hurwitz ve Kirsch, 2020). Bulut hizmet sağlayıcıları, kullanıcı verilerini korumak için önemli güvenlik önlemleri almakta, ancak bulut verilerine internet üzerinden erişilebilmesi ve büyük miktarda verinin depolanabilmesi, bu verileri potansiyel olarak savunmasız hale getirmektedir. Buluta yönelik önde gelen güvenlik sorunları ve tehditler arasında veri güvenliği, kimlik doğrulama, yetkilendirme, veri şifreleme, hizmet kesintileri ve denetim sorunları bulunmaktadır (Zhou ve

diğ., 2010). Bu nedenle, hassas verilerin bulut ortamında saklanması ve yönetilmesi dikkatli bir şekilde yapılmalıdır.

Geleneksel bilgisayar ağlarında olduğu gibi, bulut ağları da siber saldırılara karşı savunmasızdır. Günümüzde bulut bilişim ortamları, siber saldırganlar için cazip hedefler haline gelmiştir. Kullanılabilirliğe yönelik tehditler arasında DoS saldırıları, ekipman kesintileri ve doğal afetler bulunmaktadır (Jansen ve Grance, 2011). DDoS saldırıları, çevrimiçi oyunlar, hizmet sağlayıcılar, devlet kurumları, finansal hizmetler, çevrimiçi perakendeciler ve bulut ağları gibi çeşitli sektörleri hedef alarak bulut bilişim altyapılarına ciddi tehditler oluşturmaktadır. Bu saldırılar genellikle birden fazla aşamayı içerir: güvenlik açıklarını taramak için makineler kullanılır, bu güvenlik açıklarından yararlanarak seçilen makinelere saldırı kodu bulaştırılır ve saldırı yazılımı dağıtılır. Bazen, meşru uygulamalar olarak gizlenir ve güvenliği ihlal edilmiş makineleri kullanırlar (Mirkovic ve Reiher, 2004). DDoS saldırıları, bilgisayar ağlarını hedef alan tehditlerin başında gelmektedir. Bu saldırılar, bulut ağlarında hizmetlerin kesintiye uğramasına, ağ performansının düşmesine ve ağ altyapısının aşırı kötü amaçlı trafik akışıyla bunaltılmasına neden olabilir. Bu bağlamda mali kayıplar, güvenlik ihlalleri ve itibar kaybı gibi ciddi hasarlara yol açabilir. Ayrıca, bulut ortamı hacimsel DDoS saldırıları riski altındadır; bu tür saldırılar, hedef sistemi veya ağı aşırı trafikle doldurmayı amaçlar (Li ve diğ., 2020). DDoS saldırıları genellikle ağ kaynaklarını hedef alarak hizmetleri aksatmayı amaçlarken, hacimsel DDoS saldırıları ağ kaynaklarını aşırı trafik göndererek bunaltmayı hedefler. Bu tür saldırılar özellikle bulut ağlarını hedef alır çünkü bulut ağlarında eş zamanlı olarak birçok hizmet ve web sitesinde büyük ölçekli kesintilere yol açabilirler.

Şekil 2.3, hacimsel DDoS saldırı oranlarına genel bir bakış sunmaktadır. İstatistiklere göre, hacimsel DDoS tehditlerinin yaklaşık üçte biri DNS saldırıları, üçte dördü SYN saldırıları ve yaklaşık yedi saldırıdan biri UDP saldırılarıyla ilişkilidir (Cloudflare, 2023). Bu veriler, bulut ağlarının güvenliğini ve sürdürülebilirliğini sağlamak için hacimsel DDoS saldırılarına karşı güçlü önlemler geliştirmenin önemini vurgulamaktadır.



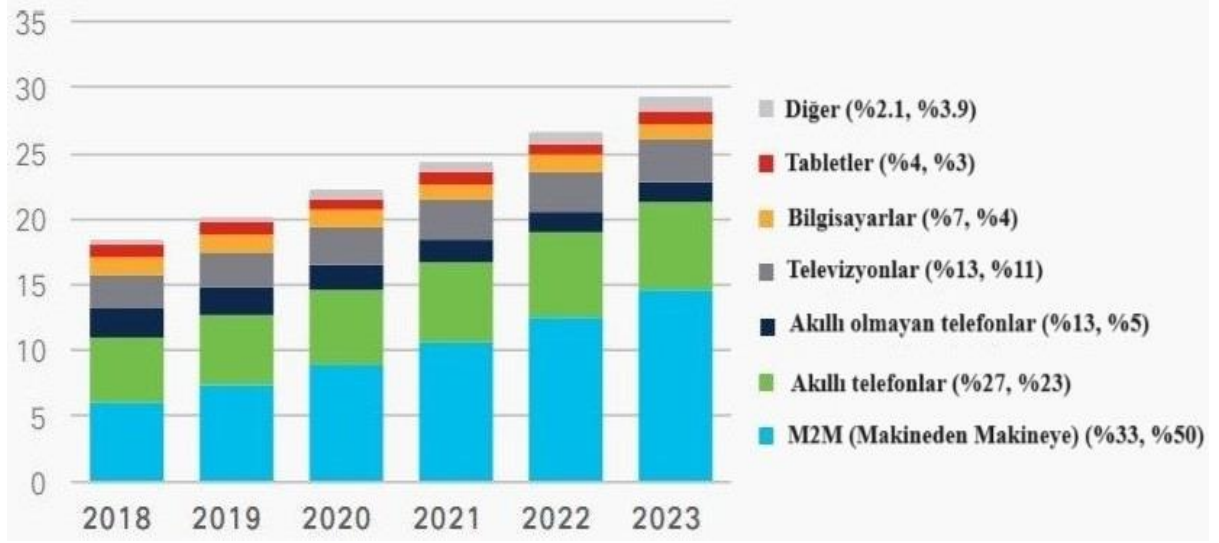
Şekil 2.3: DDoS Saldırı Vektörleri

Ajan tabanlı ağ güvenliği savunma mekanizmaları, DDoS saldırılarının etkilerini en aza indirmede kritik bir rol oynayabilir. Yeni bir yaklaşım olmamakla birlikte, siber güvenlik amacıyla tasarlanan ve geliştirilen ajanlar, bulut altyapılarının güvenliğini sağlama konusunda umut verici bir çözüm sunmaktadır. Ajanlar, bağımsız olarak çalışan varlıklar olarak tanımlanır ve karar verme, çevresine uyum sağlama ve diğer ajanlarla iş birliği yapma gibi çeşitli görevleri bağımsız bir şekilde yerine getirebilir. Ajan tabanlı siber savunma sistemleri, ister tek bir otonom ajan olarak ister çok ajanlı sistemler şeklinde uygulansın, çevrelerini algılayabilir, buna göre hareket edebilir, otonom olarak kararlar alabilir, uyum sağlayabilir, iletişim kurabilir ve iş birliği yapabilir. Bu sistemlerdeki mobil ajanlar ve diğer ajanlar, siber tehditlere karşı koymada ve bulut bilişim ortamlarının güvenliğini artırmada önemli bir rol oynayabilir. Özerklik, uyarlanabilirlik ve esneklik gibi özelliklerle tanımlanan akıllı ajanlar, ağ gelişiminin kontrollü bir şekilde yönetilmesini mümkün kılar (Boudaoud ve diğ., 2000). Geleneksel siber silahlarla karşılaştırıldığında, ajanlar daha az bilgiye ihtiyaç duyarak güvenlik açıklarını ve istismarları analiz edebilir ve zekalarını geliştirebilir (Guarino, 2013). Gelecekte, YZ destekli siber savunma ajanlarının (AICA) kullanımı, özellikle insan müdahalesinin mümkün olmadığı durumlarda zorunlu hale gelebilir (Kott ve diğ., 2018). Bulut bilişim güvenliği alanında, siber saldırıları hızlı ve etkili bir şekilde tespit etmek, bu saldırılara karşı koymak ve yanıt vermek için ajanlardan yararlanılabilir. Otonom olarak çalışan, bağımsız kararlar alabilen, çevresine uyum sağlayabilen ve diğerleriyle iş birliği yapabilen ajanlar, bulut ağlarının güvenliğini artırmada önemli bir rol oynayabilir. İster tek otonom varlıklar olarak ister çok ajanlı bir sistemin parçası olarak konuşlandırılınsınlar, siber savunma ajanları bulut çevresini algılamak, otonom kararlar alma, uyum sağlama, iletişim kurma ve etkili iş birliği yapma becerisine sahip

olabilir. Bulut güvenliğinin sürekli değişen ortamında, özellikle insan müdahalesinin pratik olmadığı senaryolarda, ajan tabanlı savunma yapılarının konuşlandırılması vazgeçilmez olabilir.

2.1.2.2. Nesnelerin İnterneti (IoT) Ağları ve DDoS Saldırıları

Nesnelerin İnterneti (IoT), elektronik cihazlar, devreler, yazılımlar, sensörler ve ağ bağlantılarıyla donatılmış fiziksel nesnelerin veri toplamasını ve bu verileri mevcut ağ altyapısı üzerinden uzaktan kontrol edilmesini sağlayarak fiziksel dünyayı bilgisayar tabanlı sistemlerle doğrudan entegre eder. Bu entegrasyon, verimliliği ve doğruluğu artırmayı hedefler (Gokhale ve diğ., 2018). Bilişim teknolojileri, günlük yaşamımızı etkileyerek çeşitli cihazları birbirine bağlar ve internetin sürekli gelişimi, IoT gibi yeni teknolojilerin ortaya çıkmasına neden olur (Abdul-Qawy ve diğ., 2015). Kablosuz Sensör Ağı (WSN) teknolojileriyle sağlanan her yerde bulunan sensörler, çevresel göstergeleri ölçme ve anlama yeteneği sunarak, çevremizdeki sensörlerin ve aktüatörlerin sorunsuz bir şekilde entegre olduğu IoT ağını oluşturur (Gubbi ve diğ., 2013). Günümüzde IoT ağlarını hedef alan DDoS saldırıları giderek daha önemli hale gelmiştir. Bu saldırı türü, IoT ağ güvenliği açısından kritik bir sorun olarak kabul edilmektedir. DDoS saldırıları önemli kayıplara neden olabileceğinden, bu saldırıların tespit edilmesi ve engellenmesi bilgisayar ağı güvenliği açısından hayati önem taşır. DDoS saldırıları, otomatik taktiklerle son derece hızlı ve büyük ölçekli bir şekilde gerçekleştirilebilir. Siber saldırganların, özellikle YZ desteğiyle daha akıllı hale geldiği gözlemlenmektedir. Her geçen gün artan DDoS saldırıları, makine hızında ve karmaşık yöntemlerle gerçekleştirilmektedir. IoT ağ operatörlerinin, yalnızca ağ yönetimiyle değil, aynı zamanda ağ güvenliği ile ilgili olarak da bu saldırılara karşı mücadele etmeleri gerekmektedir. Bu durum, giderek artan sayıda birbirine bağlı heterojen cihaz, ağ, yazılım ve insan unsuru ile dolu karmaşık ve dinamik IoT ağlarını yönettikleri anlamına gelmektedir. IoT ağlarının günümüzde hem miktar hem de büyüklük açısından hızla büyümesi, teknolojinin çeşitliliğini ve pazarın genişlemesini yansıtmaktadır. Küresel IoT pazarı 2022 yılında 478,36 milyar ABD doları seviyesindeyken, bu rakamın 2029 yılında 2.465,26 milyar ABD dolarına ulaşacağı öngörülmektedir (Fortune Business Insights, 2022). Ayrıca, 2023 yılı sonunda tüketicilerin toplam cihazlar içindeki payının %74, kalan %26'nın ise işletmelere ait olması beklenmektedir (Cisco Yıllık İnternet Raporu, 2018–2023) (Şekil 2.4).



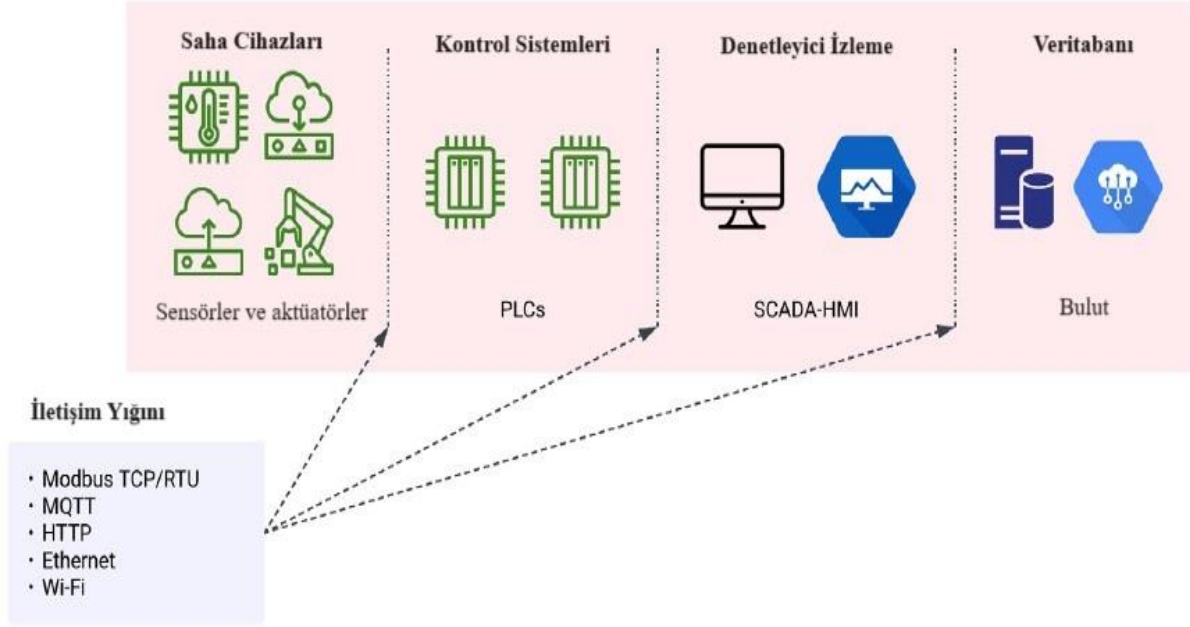
Şekil 2.4: Tüketicilerin Toplam Cihaz Sayısı İçindeki Payı

2.1.2.3. Endüstriyel Nesnelerin İnterneti (IIoT) ve DDoS Saldırıları

Günümüzde Endüstriyel Nesnelerin İnterneti (IIoT) ağları, çeşitli endüstriyel sektörlerde giderek daha yaygın hale gelmekte ve bu ağlar karmaşık ve dinamik yapılar halini almaktadır. IIoT ağları, üretim tesislerinden enerji sektörüne, ulaşımdan sağlık hizmetlerine kadar geniş bir uygulama yelpazesi sunmaktadır. Bu ağlar, endüstriyel sistemlerdeki cihazlar arasında iletişimi ve veri alışverişini kolaylaştıran bir çerçeve sağlamayı amaçlamaktadır. IIoT, endüstriyel cihazları internete bağlayarak gerçek zamanlı veri alışverişini ve gelişmiş otomasyonu mümkün kılan bir teknoloji olarak öne çıkmaktadır (Rojko, 2017). IIoT, IoT'nin bir alt kümesini oluşturur (Sisinni ve diğ., 2018; Tange ve diğ., 2020) ve endüstriyel otomasyon sistemlerinin entegre uygulamasını ifade eder (Lin ve diğ., 2015). IIoT ağları, endüstriyel ortamlarda Siber-Fiziksel Sistemler (CPS) içindeki belirli IoT teknolojilerini ve akıllı nesne türlerini kullanarak endüstriye özgü hedefleri ilerletir (Boyce ve diğ., 2018).

IIoT uygulamaları, sağlık, üretim, ulaşım ve eğitim gibi sektörlerde çeşitli avantajlar sunar. Ancak, IIoT ağlarının hızlı bir şekilde yaygınlaşması, bu ağlardaki cihazların kötü amaçlı yazılımlara karşı savunmasızlığı ve etkin bir şekilde yönetilmesiyle ilgili zorluklar, önemli güvenlik endişelerini beraberinde getirmektedir (Sengupta ve diğ., 2020). IIoT ağlarının başarılı bir şekilde uygulanması, Bilgi Teknolojileri (IT) ve Operasyonel Teknoloji (OT) alanları arasında köprü kurmayı gerektirir (Richards ve diğ., 2021). Purdue modeli, IIoT ve Siber-Fiziksel Sistemler (CPS) mimarisi için yaygın olarak kullanılan bir çerçevedir (Williams, 1994). IIoT, sensörler ve aktüatörler gibi cihazları, Programlanabilir Lojik Denetleyici (PLC),

Denetleyici Kontrol ve Veri Toplama (SCADA), Üretim Yürütme Sistemleri (MES) gibi çeşitli bileşenleri birleştiren ve entegre eden hiyerarşik bir çerçeve sunar. Purdue modelini temel alan bir IIoT mimarisi örneği **Şekil 2.5**'de gösterilmektedir (Mekala ve diğ., 2023).



Şekil 2.5: Endüstriyel Nesnelerin İnterneti (IIoT) Mimarisi Örneği

IIoT'nin yükselmesi, IT ve OT alanlarının birbirine daha yakın hale gelmesiyle birlikte üretimde yeni güvenlik zorluklarını da beraberinde getirmektedir (Mantravadi ve diğ., 2020). IIoT sistemleri, geniş ağlar üzerinden sürekli erişilebilirlik sağladığından, bu durum potansiyel saldırganlar için genişletilmiş bir saldırı yüzeyi oluşturmakta ve güvenlik risklerini artırmaktadır. IIoT teknolojileri, teknolojik ilerlemeler ve gelişen sektör ihtiyaçları doğrultusunda evrim geçirirken, bu sistemleri hedef alan siber tehditler de gelişmekte ve yeni saldırı türlerinin ortaya çıkmasına neden olmaktadır. Ayrıca, Kritik Altyapıların (CI) halka açık ağlara entegrasyonu, Endüstriyel Kontrol Sistemlerinin (ICS) çeşitli saldırı vektörlerine maruz kalmasını artırmıştır (Zimba ve diğ., 2018). Denetleme ve Veri Toplama Sistemleri (SCADA), kentsel altyapı için IIoT'nin güvenliğini sağlamak açısından kritik bir rol oynamaktadır çünkü genellikle saldırıların hedefi olup önemli sosyal aksaklıklara yol açabilmektedir (Falco ve diğ., 2018). SCADA sistemlerinin modernizasyonu, iletişim protokollerinin standartlaştırılması ve artan ara bağlantı, bu sistemlere yönelik saldırı risklerini önemli ölçüde artırmıştır (Yadav ve Paul, 2021).

İnternete bağlı cihazların çeşitli alanlarda yaygınlaşmasıyla birlikte, IIoT ağlarında DDoS saldırılarının giderek daha büyük bir endişe kaynağı haline gelmesi beklenmektedir. Siber saldırılar, IIoT ağlarının birbirine bağlı doğasını kullanarak hedeflenen sistemleri veya ağları büyük miktarda trafikle doldurup, meşru kullanıcıların erişimini engelleyebilir. Geleneksel DDoS saldırılarının aksine, IoT tabanlı DDoS saldırıları, yönlendiriciler, kameralar ve akıllı cihazlar gibi güvenliği ihlal edilmiş IIoT cihazlarından yararlanarak etkilerini artırabilir. UDP, ICMP, TCP SYN, Ping of Death, HTTP seli, NTP amplifikasyonu, Smurf saldırısı ve TCP sıfırlama gibi DDoS saldırı türleri hem IoT hem de IIoT ekosistemlerinde yaygındır (Chaudhary ve Mishra, 2023). IIoT ortamlarındaki DDoS saldırıları, veri ihlalleri veya kötü amaçlı yazılım sızması gibi daha sinsi faaliyetler için bir örtü görevi görebilir ve bu durum hassas bilgilerin gizliliğini ve bütünlüğünü tehlikeye atabilir. Mirai ve benzeri botnetler, sektörü IoT cihazlarını daha iyi koruma veya internet altyapısına yönelik artan DDoS tehditleriyle yüzleşme konusunda harekete geçmeye zorlamaktadır (Kolias ve diğ., 2017). Affinito ve diğ. (2023), Mirai imzasının aktif kaldığını ve yeni botnet varyantlarının ek bağlantı noktalarını hedef aldığını belirtmektedir. IIoT'deki sınırlı güvenlik önlemleri, cihazları DDoS saldırıları başlatmak için savunmasız hale getirir ve bu durum, felaketle sonuçlanabilecek saldırılara yol açabilir (Zhou ve diğ., 2021). DDoS saldırılarının karmaşıklığı, tüm tehdit ortamını anlamada zorluk yaratmaktadır (Mirkovic ve Reiher, 2004). Bu saldırılar, uygulamaları aşırı trafik veya isteklerle doldurarak kapasitelerini aşar ve bu nedenle yasal kullanıcılar için erişilemez hale getirir. Bir DDoS saldırısı, bir saldırganın çok sayıda IoT cihazını kontrol altına alarak bir sunucuyu büyük miktarda istekle doldurmasını ve böylece onu etkisiz hale getirmesini gerektirir (Serror ve diğ., 2020). IoT tabanlı DDoS saldırılarının kurbanı olan kuruluşlar, müşterilerinin ürün ve hizmetlerine olan güvenini zedeleyerek ciddi itibar kaybı riskiyle karşı karşıya kalmaktadır. DDoS saldırıları, IIoT sistemlerini hedef alarak kritik endüstriyel operasyonları aksatabilir, ağ kesintilerine ve işlevsellik kayıplarına neden olarak önemli bir tehdit oluşturabilir.

DDoS saldırılarını tespit etmek, günümüzdeki en önemli zorluklardan biri olmaya devam etmektedir. Antonakakis ve diğ. (2017), Mirai botnet'ini analiz ederek, bu botnet'in DDoS faaliyetleriyle uğraşan IIoT ortamlarındaki cihazlar da dahil olmak üzere IoT cihazlarına yönelik tehdit oluşturduğunu vurgulamışlardır. Zhou ve diğ. (2019), minimum yanlış alarmla hızlı yanıt sağlamak amacıyla IIoT'de DDoS azaltımı için trafik izleme ve analiz görevlerini uç cihazlar ile merkezi bulut sunucuları arasında dağıtan Sis bilişim tabanlı bir yaklaşım

önermişlerdir. Li ve diğ. (2021), DDoS tehditlerini azaltmayı hedefleyen IIoT için Birleşik Öğrenmeyle Güçlendirilmiş bir Mimari sunmuşlardır. Zhou ve diğ. (2021), IIoT sistemlerinde DDoS'u azaltmak için proaktif ve etkili bir yaklaşım geliştirmişlerdir. Mishra ve diğ. (2023), IIoT'ye özgü güvenlik ve kriptografik perspektifleri incelemişlerdir. Chaudhary ve Mishra (2023), IIoT ortamlarındaki DDoS saldırıları üzerine ayrıntılı bir çalışma gerçekleştirerek bu tür saldırıların etkilerine dikkat çekmişlerdir. Huang ve diğ. (2023), DDoS saldırılarının etkili bir şekilde tespiti, sınıflandırılması ve hafifletilmesi için "EdgeDefense" adlı çok noktalı iş birliğine dayalı bir savunma mekanizması önererek, IIoT cihazlarının ortaya çıkardığı artan siber güvenlik zorluklarını ele almışlardır.

2.1.3. Ağ Güvenliği için Kullanılan Saldırı Tespit/Önleme Sistemleri (IDS/IDPS)

Ağ güvenliği alanında kullanılan IDS'ler, ağları ve sistemleri çeşitli siber tehditlere karşı korumak amacıyla tasarlanmış kritik güvenlik araçlarıdır. Bu sistemler, ağ trafiğini ve sistem etkinliklerini sürekli olarak izleyerek şüpheli aktiviteleri tespit eder ve bu tehditlere karşı önlem alır. IDS, genellikle anormal ve kötü niyetli aktiviteleri algılayarak ağ yöneticilerini uyarır. Bu sistemler, imza tabanlı, anomali tabanlı ve davranışsal analiz yöntemleri gibi çeşitli tespit teknikleri kullanarak bilinen saldırı imzalarını, alışılmış trafik desenlerinden sapmaları veya şüpheli kullanıcı davranışlarını tanımlar. Öte yandan, IDPS, IDS'lerin ötesinde hareket eder ve yalnızca tehditleri tespit etmekle kalmaz, aynı zamanda bu tehditleri gerçek zamanlı olarak engeller veya sınırlandırır. IDPS, otomatik olarak ağ trafiğini engelleme, kötü amaçlı aktiviteleri durdurma ve güvenlik politikalarını uygulama gibi proaktif güvenlik önlemleri alır. Bu sistemler hem ağ tabanlı hem de host tabanlı yapılandırmalarla uygulanabilir ve organizasyonların güvenlik stratejilerini güçlendirerek veri güvenliğini artırır. IDS ve IDPS sistemleri, ağ güvenliğinde kritik bir rol oynayarak veri ihlalleri, hizmet kesintileri ve diğer güvenlik tehditlerine karşı savunma sağlar.

2.2. Ajanlar ve Ağ Güvenliği

2.2.1. Ajan Kavramı ve Temel Özellikleri

Ajan kavramı, YZ ve siber güvenlik gibi çeşitli alanlarda önemli bir rol oynayan bir teknolojidir ve genellikle bağımsız olarak hareket edebilen, çevresine tepki verebilen ve belirli görevleri yerine getirebilen yazılım veya donanım birimi olarak tanımlanır. Russell ve Norvig (2009), ajan kavramını hedeflere ulaşmak için esnek ve özerk eylemde bulunabilen bilgisayar sistemleri

olarak tanımlanmaktadır. Ajanlar, basit refleks ajanları, model bazlı refleks ajanları, hedef bazlı ajanlar ve fayda bazlı ajanlar olmak üzere çeşitli türlerde sınıflandırılabilir. Basit refleks ajanları, kararlarını yalnızca anlık verilere dayanarak alır ve genellikle minimum kurallarla tanımlanmış basit ortamlar için uygundur. Buna karşın, model bazlı refleks ajanları, karar alma süreçlerinde bir dünya modeline başvurur ve gözlemlenemeyen ortamlarda kullanılır. Hedef bazlı ajanlar, önceden belirlenmiş hedeflere sahiptir ve karmaşık ortamlarda çeşitli hedeflere ulaşmak için uygun eylemleri seçebilirler. Fayda bazlı ajanlar ise, tercihlerine göre beklenen faydaları en üst düzeye çıkarmak amacıyla eylemlerini seçerler ve belirsiz sonuçlar ile farklı hedefler arasındaki ödünleşimlerin olduğu ortamlarda idealdirler. Ajanlar, temel olarak üç ana özelliğe sahiptir: algılama, karar verme ve eyleme geçme. İlk olarak, ajanlar çevrelerinden veri toplar ve bu verileri işleyerek çevresel durumları algılar. Algılama süreci, ajanların çevrelerinden gelen sinyalleri anlamlı bilgiye dönüştürmesini sağlar. İkinci olarak, ajanlar topladıkları verileri analiz ederek karar verme mekanizmalarını çalıştırır ve elde edilen bilgilere dayanarak stratejik kararlar alır. Karar verme süreci, ajanın hedeflerine ulaşmasını sağlamak için en uygun hareket planını belirlemeye yöneliktir. Son olarak, ajanlar karar verdikten sonra eyleme geçer; yani belirlenen strateji doğrultusunda hareket eder ve gerekli görevleri yerine getirir. Bu süreç, ajanların belirli hedeflere ulaşmak veya sorunları çözmek için etkili bir şekilde çalışmasını sağlar. Ajanlar ayrıca genellikle özerklik, adaptasyon ve iş birliği yetenekleriyle de karakterize edilir. Özerklik, ajanların kendi başlarına çalışabilme yeteneğini ifade ederken, adaptasyon yeteneği ajanların değişen koşullara uyum sağlayabilmesini ve iş birliği yeteneği, birden fazla ajan arasında koordineli çalışma kapasitesini ifade eder. Bu özellikler, ajanların karmaşık sistemlerde ve dinamik çevrelerde etkili bir şekilde çalışabilmesini mümkün kılar..

2.2.2. Akıllı Ajanların Ağ Güvenliğinde Rolü ve Kullanımı

Siber ajanlar, ajan türleri arasında son derece önemli bir yere sahiptir (Kott, 2018). Otonom bir ajan çevresini algılayabilir ve bu bilgilere göre tepki verebilir (Guarino, 2013). Akıllı siber ajanlar, ortaya çıkan yan etkileri öngörerek ve minimize ederek karmaşık kötü amaçlı yazılımlarla mücadele edebilir ve karmaşık, çok adımlı faaliyetleri planlayıp gerçekleştirebilir (Kott, 2018). Bir siber ajanın temel görevleri arasında ağ trafiğini tanımlamak ve analiz etmek, siber saldırıları tespit ve azaltmak, sistem performansını izlemek ve insan operatörlere durumsal farkındalık sağlamak yer alır. Ayrıca, tahmine dayalı analiz, tehdit avcılığı ve olay müdahalesi gibi diğer görevlerde de kullanılabilirler. Siber ajanlar, askeri ağların ve diğer kritik siber altyapıların güvenliğini ve dayanıklılığını sağlamak açısından kritik öneme sahiptir. Akıllı

otonom ajanların, geleceğin savaş alanlarında önemli bir akıllı varlık olarak yaygın şekilde kullanılacağı söylenebilir (Kott, 2018). Ajanlar, güvenlik açıklarını analiz ederek ve bu açıklara yönelik istismarlar geliştirerek istihbaratlarını geliştirebilirler. Bu nedenle, geleneksel siber silahlara kıyasla bir siber saldırı başlatmak için çok daha az bilgiye ihtiyaç duyarlar (Guarino, 2013). YZ ajanları, özellikle insan müdahalesinin mümkün olmadığı durumlarda gelecekteki iletişim ortamlarında kötü amaçlı yazılımları ortadan kaldırmaları gerekecektir (Kott ve Theron, 2020). ACDA'lar siber saldırı stratejilerine karşı dinamik olarak karar alma süreçlerini uyarlayabilir ve geliştirebilir (Dutta ve diğ., 2021). AICA'lar tek bir otonom ajan ya da çoklu ajan sistemleri olabilir (Fruchart ve Blanc, 2020). NATO Bilim ve Teknoloji Örgütü Araştırma Grubu tarafından yapılan son araştırmalar, ajan teknolojilerinin kullanımında devam eden çabalara örnek teşkil etmektedir. Bu uluslararası ekip, ayrıca AICA vizyonunu özetleyen ve üst düzey bir referans mimarisi sunan bir rapor yayımlamıştır (Kott, 2018). Özellikle siber saldırılar yüksek makine hızlarında, YZ destekli ve hatta otonom olarak gerçekleştiğinde, ajanlara dayalı ağ güvenlik sistemleri bulut güvenliğini sağlamak için karşı koyma tedbirleri oluşturma potansiyeline sahip olabilirler. Özellikle gelişmiş gerçek zamanlı siber saldırılara karşı, minimum insan müdahalesiyle siber savunma sağlamak amacıyla otonom siber savunma ajanlarının kullanımı büyük umut vaat etmektedir (Dutta ve diğ., 2021). AICA'lar, siber saldırı işaretlerini tespit ederek karşı önlemler geliştirmeli, bu önlemleri gerçek zamanlı olarak taktiksel bir şekilde uygulamalı ve insan operatörlere geri bildirim sağlamalıdır (Fruchart ve Blanc, 2020). Çoklu ajan tabanlı siber savunma sistemleri de siber saldırıların tespit edilmesi ve önlenmesinde önemli bir potansiyele sahiptir ve güvenli ağların geliştirilmesi ve sürdürülmesi için değerli bir yaklaşım sunma imkân ve kabiliyetlerine sahiptirler (Boudaoud ve diğ., 2000).

2.2.3. İlgili Çalışmalar Konusunda Literatür Taraması

2.2.3.1. LSTM Ağı ile Derin Öğrenme (DÖ) Tabanlı Siber Saldırı Tespiti Yaklaşımı

Novaes ve diğ. (2020), DDoS ve Portscan saldırılarını tespit etmek ve azaltmak için bir LSTM-Fuzzy sistemi önermişlerdir. Alguliyev ve diğ. (2019), CNN ve LSTM tabanlı DÖ algoritmalarını kullanarak sosyal ağ metinlerinde duygu analizi gerçekleştirmiş ve bu analizler aracılığıyla DDoS saldırılarının olasılığını tahmin etmişlerdir. Fang ve diğ. (2018), LSTM tabanlı bir "kötü amaçlı JavaScript kodu algılama" modeli önermiş, Haider ve diğ. (2020) ise DDoS tespiti için bir CNN çerçevesi geliştirmişlerdir. Sahi ve diğ. (2017), genel bulut ağları mimarisinde DDoS saldırılarını tespit etmek için CS_DDoS adlı bir sınıflandırma modeli

geliştirmiş ve çalışmalarında, farklı derin öğrenme algoritmaları arasında LS-SVM sınıflandırıcısının en iyi performansı sağladığını ortaya koymuşlardır. Tan ve diğ. (2020), K-Means ve KNN tabanlı birleşik bir makine öğrenme algoritması kullanarak DDoS saldırılarını tespit etmeye yönelik bir çerçeve önermişlerdir. Hwang ve diğ. (2019), kelime yerleştirme ve LSTM modelini birleştirerek kötü amaçlı yazılımların sınıflandırılması için yeni bir yaklaşım geliştirmişlerdir. Elsayed ve diğ. (2020), CICDDoS2019 veri kümesi üzerinde RNN kullanarak SDN ağına yönelik DDoS saldırılarını tespit etmek için DDoSNet adlı bir model önermişlerdir. Çil ve diğ. (2021), ağ trafiğine yönelik DDoS saldırılarının tespiti ve sınıflandırılmasında daha etkili olması beklenen bir derin öğrenme modeli önermişlerdir. Bu çalışmada önerilen DNN modeli, DDoS saldırı tespitinde %100'e yakın doğruluk sağlamakta ve CICDoS2019 veri seti üzerinde DDoS saldırılarını %95'e yakın doğrulukla sınıflandırmada başarılı bulunmuştur. Rehman ve diğ. (2021), GRU ve RNN gibi derin öğrenme algoritmalarının yanı sıra geleneksel makine öğrenme algoritmaları olan NB ve SMO'yu kullanarak CICDoS2019 veri kümesi üzerinde ağ üzerinden DDoS saldırılarını tespit etmek ve tanımlamak için DIDDOS yaklaşımını önermişlerdir. Yazarlar, GRU algoritması kullanılarak SSDP saldırıları durumunda elde edilen en yüksek doğruluğun %99,91, diğer tüm saldırılarda ise ortalama %99,7 olduğunu belirtmişlerdir. Patil ve diğ. (2019), sanal makine ağ trafiğini izleyen bir güvenlik çerçevesi tasarlamışlardır. Zhou ve diğ. (2021), uzun vadeli bağımlılıkların tespit zorluğunu aşmak için LSTM kullanan ve çok aşamalı saldırıları tespit etmeyi amaçlayan bir saldırı tespit modeli geliştirmişlerdir. Literatürdeki bu çalışmalar, DDoS saldırılarının tespiti ve engellenmesi amacıyla çeşitli derin öğrenme algoritmalarının ve yöntemlerinin kullanıldığını göstermektedir.

2.2.3.2. LSTM Ağı ile Akıllı Ajan Tabanlı Siber Saldırı Tespiti

Mezghani ve Ktata (2020) tarafından yapılan çalışmada, üç ana ajan içeren bir sistem önerilmiştir: Toplayıcı ajanı, konfigürasyon ajanı ve eğitim ajanı. Yazarlar, bu sistemde NSL-KDD veri kümesindeki CNN ve RNN ağlarını kullanmışlardır. Thamilarasu ve diğ. (2020), birbirine bağlı tıbbi cihaz ağlarının güvenliğini artırmak amacıyla mobil ajan tabanlı bir IDS önermişlerdir. Bu çalışma, sensör ajanı, küme ajanı ve dedektif ajan olmak üzere üç ajandan oluşan bir mimariyi kapsamaktadır. Suwannalai ve Polprasert (2020), ağ IDS zorluklarını ele almak için DRL etkinliğini araştırmış ve anormallik tabanlı bir ağ IDS'si için DQN kullanmışlardır. Sethi ve diğ. (2020) tarafından geliştirilen çalışmada, birden fazla dağıtılmış ajan ortamında DQN mantığını kullanan ve gelişmiş ağ saldırılarını tanımlamak ve sınıflandırmak için dikkat mekanizmalarını birleştiren bir DRL tabanlı siber saldırı tespit

sistemi önerilmiştir. Bu çok ajanlı IDS yapısının, ajanların ölçeklenebilir, hataya dayanıklı ve çok görüntülü mimari destekli bir güvenlik sistemi oluşturmak için iş birliği yaptığı dağıtılmış bir izinsiz giriş tespit platformuna dayandığı görülmektedir. Benlloch-Caballero ve diğ. (2023), DDoS saldırılarıyla mücadele etmek için dağıtılmış ve çift katmanlı kendini koruma yetenekleri sağlayan bilişsel bir kapalı döngü sistemi geliştirmiştir. Aydın ve diğ. (2022), bulut ortamında DDoS saldırılarıyla mücadele için derin öğrenme tabanlı bir siber savunma mekanizması önermişlerdir. Alasmay ve diğ. (2022), IoT cihazlarını korumak için ajan tabanlı bir yaklaşım tanıtmışlardır. Bu sistem, iki farklı makine öğrenme sınıflandırıcısının kullanımını içermektedir. Alayizadeh ve diğ. (2022) tarafından gerçekleştirilen çalışmada, Q-öğrenme tabanlı model ile derin ileri beslemeli sinir ağı yaklaşımının birleştirilmesi hedeflenmiştir. Butt ve diğ. (2022), akıllı ev ağlarında anormallik tabanlı izinsiz giriş tespitine yönelik bir yöntem önermişlerdir. Haider ve diğ. (2020) tarafından gerçekleştirilen çalışmada, Yazılım Tanımlı Ağlar (SDN) için DDoS saldırılarının etkili bir şekilde tespit edilmesini sağlayan bir CNN topluluğu çerçevesi geliştirilmiştir. Yazarlar, önerdikleri sistemin SDN'lerde gerçek zamanlı DDoS tespiti için etkili bir çözüm sunduğunu vurgulamaktadırlar. Protogerou ve diğ. (2021), IoT'de dağıtılmış anormallik tespiti için Grafik Sinir Ağı (GNN) yöntemini önermişlerdir. Bu yöntem, IoT ağının grafiksel bir temsili için her cihazı bir düğüm ve aralarındaki bağlantıları kenarlar olarak gösteren bir grafik kullanmaktadır. Alzubi ve diğ. (2023) tarafından geliştirilen çalışmada, kötü amaçlı yazılım tespiti ve sınıflandırması için Optimal Kodlayıcı-Kod Çözücüye Dayalı bir LSTM Ağı (OELSTM-MDC) geliştirilmiştir. Bu yeni teknik, eğitim ve test veri kümelerinde kötü amaçlı yazılım sınıfı tespiti için %97,14, iyi huylu sınıf tespiti için ise %98,33 doğruluk sağlamış ve geleneksel yöntemleri geride bırakmıştır. Alweshah ve diğ. (2023) tarafından yapılan çalışmada, problem alanını keşfetmek için İmparator Penguen Kolonisi (EPC) yöntemini ve IoT uygulamalarındaki özellik seçimi zorluklarını ele almak amacıyla KNN kullanan benzersiz bir özellik seçim modeli tanıtılmıştır.

Spafford ve Zamboni (2000) tarafından geliştirilen Saldırı Tespiti için Otonom Ajan (AAFID) sistemi, saldırı tespiti amacıyla otonom ajanların kullanılmasını önermektedir. Yazarlar, bu sistemin dört temel bileşenden oluştuğunu vurgulamaktadır: ajanlar, filtreler, alıcı-vericiler ve monitörler. Fenet ve Hassas (2002), mobil ajanları kullanan ve dağıtılmış bir yapıya sahip gizli bir IDS mimarisi geliştirmiştir. Dasgupta ve diğ. (2005), "Cougar" adlı bilişsel ajan mimarisini temel alan Cougaar tabanlı bir IDS tasarlamışlardır. Katata ve diğ. (2009) tarafından yapılan çalışmada, ağlardaki kötü amaçlı etkinliklerin belirlenmesindeki sınırlamaları aşmak amacıyla

bir model geliştirilmiş ve bu model SYN Flooding saldırıları için test edilmiştir. Guarino (2013), literatürdeki otonom ajanlarla ilgili mevcut çalışmaları incelemiştir. Balogh ve diğ. (2014), sertifikalı güvenilir ajanlara dayalı bulut altyapıları için yeni bir yaklaşım önermişlerdir. Boudaoud ve diğ. (2000), akıllı ajanlar kullanarak izinsiz girişleri tespit etmeyi amaçlayan bir IDS geliştirmiştir. Mezghani ve Ktata (2020) tarafından yapılan çalışmada, beş ana ajandan oluşan bir IDS sistemi geliştirilmiştir. Lyu ve diğ. (2020), geleneksel otonom bilgi işlem çerçevesini modellemek için birden fazla araç geliştiren çok ajan tabanlı bir bilgi işlem modeli oluşturmuşlardır. Fruchart ve Blanc (2020) tarafından geliştirilen çalışmada, yeni bir IDS mimarisi önerilmiştir. Sethi ve diğ. (2020), ağ boyunca dağıtılan birden fazla bağımsız DRL ajanı kullanarak yeni ve karmaşık saldırıların doğru tespiti ve sınıflandırılması için bir IDS geliştirmiştir. Thamilarasu ve diğ. (2020), bağlı tıbbi cihazlardan oluşan ağların güvenliğini sağlamak amacıyla mobil ajan tabanlı bir IDS tasarlamışlardır. Önerilen sistem, sensör verilerindeki hiyerarşik, otonom ve ağ düzeyindeki izinsiz girişleri ve anormallikleri tespit etmek için makine öğrenimi ve regresyon algoritmalarını kullanmaktadır. Shurman ve diğ. (2020) tarafından yapılan çalışmada, IoT ortamında siber saldırıları tespit etmek için iki farklı yaklaşım önerilmiştir. Novaes ve diğ. (2020), SDN kurulumlarına yönelik DDoS saldırılarını tespit etmek ve azaltmak için LSTM-Fuzzy adlı bir sistem önermiştir. Kumar ve Behal (2020), DDoS saldırılarını tespit etmek amacıyla CICDDoS2019 veri seti üzerinde LSTM tabanlı bir model geliştirmiştir. Zhang ve diğ. (2020), programlanabilir anahtarlar kullanarak hacimsel DDoS saldırılarına karşı savunma sağlamak için POSEIDON adlı bir tasarım geliştirmiştir. Dutta ve diğ. (2021), RL tabanlı savunma politikalarının optimizasyonu ve doğrulaması amacıyla yeni bir hibrit otonom ajan mimarisi geliştirmiştir. Metge ve diğ. (2021), operatör-acente ortak karar verme görevini simüle ederek dağıtım sırasında AICA ile insanlar arasındaki iş birliğini araştırmıştır. Thomson ve diğ. (2021), insan benzeri öğrenme ve uyarlanabilirliği mümkün kılan bilişsel mekanizmaları tanıtmış ve özellikle sembolik bilgi temsili altındaki alt sembolik aktivasyon mekanizmalarına odaklanmıştır. Sethi ve diğ. (2021), gelişmiş ağ saldırılarının etkili bir şekilde algılanması ve sınıflandırılması için birden fazla dağıtılmış ajanları ve dikkat mekanizmalarını içeren, DQN mantığını kullanan bir IDS önermişlerdir.

Aydın ve diğ. (2021) tarafından yapılan çalışmada, genel bulut ağı ortamında DDoS saldırılarını tespit etmek ve önlemek amacıyla önerilen LSTM-CLOUD adlı sistem tanıtılmıştır. Prasad ve diğ. (2022), hacimsel DDoS saldırılarıyla mücadele etmek için geliştirdikleri oylamaya dayalı çok modlu bir IDS'i sunmuşlardır. Gaur ve diğ. (2022) tarafından yapılan çalışmada,

CICDoS2019 değerlendirme veri setini kullanarak DDoS saldırılarını tespit etmek için bir LSTM modeli önerilmiştir. Subramanian ve diğ. (2021), özellikle uygulama, ağ ve aktarım katmanlarındaki DDoS saldırılarını tespit etmek ve kategorize etmek amacıyla bir DÖ modeli kullanmışlardır. Sinthuja ve Suthendran (2022) tarafından yapılan çalışmada, BoT-IoT ve CICDDoS2019 veri kümelerini kullanarak IoT bağlamında DDoS saldırılarını tespit etmek için bir LSTM tabanlı derin öğrenme modeli sunulmuştur. Ahamed ve diğ. (2022), DDoS saldırıları ile normal trafik arasındaki ayrımı yapmak için YSA kullanarak akıllı ev ağlarında DDoS saldırılarını tespit etmeye yönelik etkili ve hafif bir yöntem önermişlerdir.

2.2.3.3. Derin Pekiştirmeli Öğrenme (DRL) ile Akıllı Ajan Tabanlı Siber Saldırı Tespiti

Literatürdeki mevcut çalışmalar, özellikle ağ saldırılarının tespitinde ve yeni saldırılara karşı savunmada DRL tekniklerinin kullanımına olan ilginin arttığını göstermektedir. Wiering ve diğ. (2011), sınıflandırma görevleri için RL ve çok katmanlı algılayıcıları birleştiren ve geleneksel yöntemlerle karşılaştırılabilir performans sergileyen yeni bir RL çerçevesi sunmuşlardır. Sethi ve diğ. (2020), izinsiz giriş girişimlerini belirleme ve bunlara yanıt verme etkinliğini artırmak için RL'ye dayalı bir yaklaşım kullanan bir IDS önermişlerdir. Suwannalai ve Polprasert (2020), ağ güvenliğini güçlendirmede umut verici sonuçlar gösteren RL'yi kullanan IDS'ler için yeni bir yaklaşım geliştirmişlerdir. Feng ve diğ. (2020), uygulama katmanında DDoS saldırılarıyla mücadele etmek için RL tabanlı yeni bir yöntem sunmuş ve bu yöntemin ağ ortamlarında hizmet kalitesini sağlama potansiyelini ortaya koymuşlardır. Hsu ve Matsuoka (2020), anormallik tabanlı IDS'ler için DRL öğrenmeyi kullanan yeni bir yaklaşım önererek, özellikle bulut ağ ortamlarında ağ saldırılarını etkili bir şekilde tanımlama ve azaltma potansiyelini göstermişlerdir. Sethi ve diğ. (2021), dikkat tabanlı çok ajanlı IDS'leri RL ile entegre ederek bilgisayar ağlarında izinsiz giriş tespiti için yenilikçi bir yaklaşım geliştirmişlerdir. Gülmez ve Angın (2021), izinsiz giriş tespiti için DRL etkinliğini değerlendiren bir çalışma yürütmüşlerdir. Bouhamed ve diğ. (2021), periyodik DRL öğrenmeye dayalı bir yaklaşım kullanarak özellikle İnsansız Hava Araçları (İHA) ağları için özel olarak tasarlanmış bir IDS önermişlerdir. Alavizadeh ve diğ. (2022), derin Q-öğrenmeye dayalı bir RL yaklaşımı geliştirmişlerdir. Yungaicela-Naula ve diğ. (2022), DDoS saldırılarını tespit etmek ve durdurmak için esnek bir SDN çerçevesi geliştirmişlerdir. Li ve diğ. (2023), Araçların İnterneti (IoV) için Double Deep Q-Network (DDQN) tabanlı bir DDoS algılama yöntemi önermişlerdir. Al-Fawa'reh ve diğ. (2023), DRL kullanarak IoT ağlarındaki kötü amaçlı yazılım botnet'lerini tespit etmek için bir yöntem geliştirmişlerdir. Mishra ve diğ. (2023), bilişsel yetenekleri ve dağıtılmış DRL

tekniklerini kullanarak tıbbi siber-fiziksel sistemlerde güvenliğini artırmayı amaçlayan bir model sunmuşlardır. Rookard ve Khojandi (2023), IoT cihazlarını hedef alan saldırıları tespit etmek için DRL tekniklerinin uygulanmasını incelemiş ve bu yaklaşımın IoT ağlarının çeşitli siber tehditlere karşı güvenliğini artırmadaki potansiyelini ortaya koymuşlardır. Malik ve Saini (2023) ise DRL tekniklerinden yararlanan bir ağ IDS geliştirmişlerdir.



3. YÖNTEM

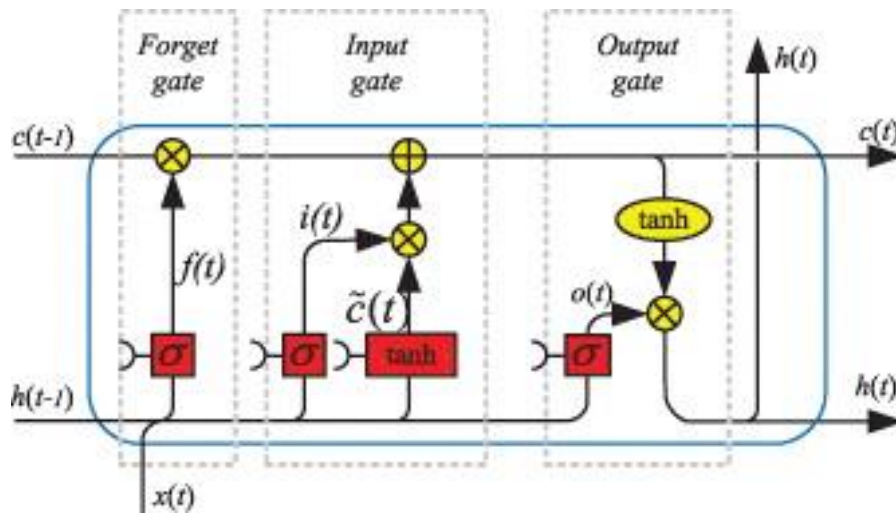
Araştırmanın metodolojisinde dört aşamalı bir gelişim yöntemi tanımlanmış ve uygulanmıştır. **Tablo 3.1**'de, ajan tabanlı siber saldırı tespiti modelinin aşamalı gelişim süreci özetlenmiştir. Birinci aşamada, LSTM ağı kullanılarak DÖ tabanlı bir siber saldırı tespiti yaklaşımı önerilmiştir. İmza tabanlı bu model, CICDDoS2019 veri seti ile test edilmiştir. Bu aşamanın amacı, imza tabanlı derin öğrenme modelinin etkinliğini ve doğruluğunu araştırmak ve daha karmaşık sistemlere geçiş için temel oluşturmaktır. İkinci aşamada, IoT ortamında, LSTM ağı ile ajan tabanlı bir siber saldırı tespiti yaklaşımı benimsenmiştir. Ayrıca, çoklu ajanlı sistemlerin imza tabanlı siber saldırı tespiti yönetimi incelenmiştir. CICDDoS2019 veri seti kullanılarak özgün bir model geliştirilmiş ve çoklu ajan sistemlerinin etkinliği detaylı olarak analiz edilmiştir. Üçüncü aşamada, bulut ortamında LSTM ağı ile ajan tabanlı bir siber saldırı tespiti modeli önerilmiştir. Ayrıca, büyük ölçekli ağlarda imza tabanlı siber saldırı tespiti ele alınmıştır. Hacimsel DDoS saldırılarına karşı etkili savunma stratejileri geliştirmeye yönelik araştırmalar yapılmış ve CICDDoS2019 veri seti kullanılmıştır. Dördüncü aşamada, endüstriyel IoT ortamında, ajan tabanlı DRL yaklaşımı uygulanmıştır. Anomali tabanlı DRL algoritmaları, CIC-IoT-2022 ve CIC-IoT-2023 veri setleri ile test edilmiştir. Bu aşamanın amacı, otonom ve etkili endüstriyel IoT güvenliği sağlamak ve hızlı tepkiler geliştirmektir. Her aşama, özgün bir saldırı tespit modeli geliştirmeyi amaçlarken, aynı zamanda ilgili veri setleri ve teknolojik yöntemler kullanılarak siber saldırılara karşı daha etkin savunma stratejileri oluşturulmuştur.

Tablo 3.1: Önerilen Modellerin Aşamalı Gelişim Yöntemi

Aşama	Siber Saldırı Tespit Yaklaşımı	Kullanılan Teknolojik Yöntem	Yaklaşım Türü	Kullanılan Veri Seti	Hedeflenen Kazanımlar	Etki Faktörü / Yorum
1'inci Aşama	LSTM Ağı ile Derin Öğrenme (DÖ) Tabanlı Siber Saldırı Tespiti Yaklaşımı	Derin Öğrenme (LSTM Ağı)	İmza Tabanlı	CICDDoS2019	Siber saldırı tespitinde imza tabanlı Derin Öğrenme (LSTM Ağı) modelinin etkinliği ve doğruluğu üzerine detaylı araştırma ile daha karmaşık ve özelleştirilmiş sistemlere geçiş için temel oluşturma. Bu bağlamda özgün bir saldırı tespit modeli geliştirme.	Modelin doğruluğu ve karmaşıklığı üzerine yapılan detaylı değerlendirme
2'nci Aşama	LSTM Ağı ile Ajan Tabanlı Siber Saldırı Tespiti (IoT Ortamı) Yaklaşımı	Ajan Tabanlı Derin Öğrenme (DDoS Saldırıları)	İmza Tabanlı	CICDDoS2019	Çoklu ajanlı sistemlerin imza tabanlı siber saldırı tespiti yönetimi ve etkinliği üzerine detaylı araştırma. Bu bağlamda özgün bir saldırı tespit modeli geliştirme.	IoT ortamında çoklu ajanlı sistemlerin yönetimi ve etkinliği üzerine derinlemesine analiz
3'üncü Aşama	LSTM Ağı ile Ajan Tabanlı Siber Saldırı Tespiti (Bulut Ortamı) Yaklaşımı	Ajan Tabanlı Derin Öğrenme (Hacimsel DDoS Saldırıları)	İmza Tabanlı	CICDDoS2019	Büyük ölçekli ağlarda çoklu ajanlı sistemlerin imza tabanlı siber saldırı tespiti üzerine detaylı araştırma. Bu bağlamda özgün bir saldırı tespit modeli geliştirme.	Bulut ortamında DDoS saldırılarına karşı etkili savunma stratejileri geliştirme üzerine odaklanma
4'üncü Aşama	Ajan Tabanlı Derin Pekiştirmeli Öğrenme (DRL) ile Siber Saldırı Tespiti (Endüstriyel IoT Ortamı) Yaklaşımı	Ajan Tabanlı Derin Pekiştirmeli Öğrenme (DRL)	Anomali Tabanlı	CIC-IoT-2022 ve CIC-IoT-2023	Çoklu Ajan ile anomali Tabanlı Derin Pekiştirmeli Öğrenme (DRL) algoritmalarının uygulanabilirliği ve etkinliği üzerine detaylı araştırma. Bu bağlamda özgün bir saldırı tespit modeli geliştirme.	Otonom ve etkili IIoT güvenliği sağlama, hızlı ve çevik tepkiler üzerine odaklanma

3.1. Derin Öğrenme (DÖ) Tabanlı Model Önerisi

DÖ algoritmaları, geçmiş ağ trafiği verilerine dayanarak gelecekteki siber saldırıları tahmin etme imkânı sunma yeteneğine sahip algoritmalarlardır. Bu algoritmalar, geçmiş verilerden hareketle ağ profilini karakterize eden bir model oluşturmaya olanak tanır. Saldırıları tanımlamak için, saldırı tespit veri kümeleri üzerinde DÖ teknikleri ile eğitim ve test süreçleri gerçekleştirilebilir. Örneğin, LSTM ağı, uzun vadeli bağımlılıkları öğrenebilen bir DÖ algoritmasıdır. Klasik DÖ modelleri arasında Derin Sinir Ağları (DNN), Konvolüsyonel Sinir Ağları (CNN) ve Tekrarlayan Sinir Ağları (RNN) yer alır. DNN'ler, birden fazla gizli katman kullanarak karmaşık ilişkileri öğrenebilir. CNN'ler, özellikle görsel ve zaman serisi verilerinde etkili olup, ağ trafiğindeki örüntüleri tanımlamada kullanılabilir. RNN'ler, zaman serisi verilerindeki geçmiş bilgileri hatırlama yeteneği ile zaman bağımlı örüntüleri yakalayabilir. CNN, konumdan bağımsız özelliklerin çıkarılmasında yüksek performans sergileyen bir DÖ mimarisi olarak kabul edilir (Yin ve diğ., 2017). Bu nedenle, Doğal Dil İşleme (NLP), görüntü ve ses işleme gibi birçok alanda kullanılmaktadır. RNN'ler ise genellikle bir sonraki adımı tahmin etmek için kullanılan derin öğrenme yapılarıdır. Eğitilirken, girdiler arasındaki ilişkileri kurarak bir sonraki adımı tahmin etme ve tüm ilişkileri hatırlama yeteneğine sahiptirler. Özellikle, LSTM ağı bir tür RNN olup, uzun vadeli bağımlılıkları öğrenme yeteneğine sahiptir ve metin tanıma, konuşma tanıma, makine çevirisi ve siber güvenlik gibi birçok alanda kullanılmaktadır. **Şekil 3.1**'de, örnek bir LSTM mimarisi gösterilmektedir.



Şekil 3.1: Unutma Kapısına Sahip LSTM Mimarisi

Her LSTM hücresinin içinde önceki adımlardan hangi bilgilerin unutulacağına, hangilerinin hatırlanacağına karar vermek gibi farklı işlemleri gerçekleştiren kapılar bulunmaktadır. Giriş kapısı (i_t , çıkış kapısı (O_t) ve unutma kapısı (f_t) ile LSTM ağı, girişler arasındaki ilişkileri uzun mesafelerden yakalayabilir ve ilgisiz ve anlamsız noktaları öğrenmez. Hücredeki ilk adım, hücrenin durumundan hangi bilgilerin çıkarılacağına karar vermektir. Unutma kapısı adı verilen kapıda bu işlem Denklem 3.1 kullanılarak gerçekleştirilir:

$$f(t) = \sigma(W_f[h_t - 1, x_t] + b_f) \quad (\text{Denklem 3.1})$$

Bir önceki adımın gizli durumu ($h_t - 1$) ve o adımdaki hücre girişi (x_t) bu sürecin girdilerini oluşturur. (W_f) ve (b_f) bu adımda kullanılacak ağırlık ve ek giriş (bias) değerleridir. İşlem sonucunu hesaplamak için 0 ile 1 arasında değer üreten $f(t)$ sigmoid aktivasyon fonksiyonu kullanılır σ . Bu fonksiyonun sonucunun 1'e yaklaşması gelen bilginin saklanması, 0'a yaklaşması ise gelen bilginin unutulması anlamına gelir. Hücrede gerçekleştirilecek bir sonraki adım, yeni gelen bilgilerden hangisinin hücrenin durumunda tutulacağına karar vermektir. Bu süreçte giriş kapısında hangi değerlerin güncelleneceğine karar verilirken (**Denklem 3.2**), C_t hücre durumuna eklenecek yeni aday bilgileri ayrı bir tanh katmanında (**Denklem 3.3**) belirlenir.

$$i_t = \sigma(W_i \cdot [h_t - 1, x_1] + b_i) \quad (\text{Denklem 3.2})$$

$$C_t = \tanh(W_c \cdot [h_t - 1, x_1] + b_c) \quad (\text{Denklem 3.3})$$

Giriş olarak unutma kapısına benzer şekilde, her iki işlem de bir önceki adımın gizli durumunu ve $h_t - 1$ adımın hücre girişini alır. W_i ve W_c bu adımların ağırlık parametreleridir b_i ve b_c ek giriş değerleridir. Bu iki adım sonucunda elde edilen değerler (i_t ve C_t), daha önce elde edilen unutma kapısı sonucu kullanılarak Denklem 3.4'deki gibi birleştirilecektir. Böylece C_t hücrenin eski durumu (-1) yeni durumuna (C_t) aktarılmış olacaktır.

$$C_t = f_t \cdot C_{t-1} + i_t * C_t \quad (\text{Denklem 3.4})$$

Çıkış kapısı, çıktıya yalnızca kararlaştırılan bilginin çıkmasını sağlar. Bu kapı da diğer kapılar gibi bir önceki adımın gizli durumunu ve $h_t - 1$ o adımın hücre girişini girdi olarak alır. Denklem 3.5'de x_t , çıkış kapısından gelen bilgiye σ_t ve Denklem 3.6 ile bulunan hücrenin durumuna bağlı olan LSTM hücresinin çıkışını gösterir.

$$\sigma_t = \sigma(W_i \cdot [h_t - 1, x_1] + b_\sigma) \quad (\text{Denklem 3.5})$$

$$h_t = \sigma_t \cdot \tan h(C_t) \quad (\text{Denklem 3.6})$$

DÖ sürecinde eğitim süreleri uzun olsa da test aşamalarında oldukça olumlu sonuçlar alınabilmektedir. LSTM yapısında ağdaki verilerin unutulup unutulmayacağına karar verilir. Unutma katmanı bir ortamdan diğerine aktarıldığında, modelin gidişatı için eski ortam gerekli değilse o ortam kaldırılır. Zamana bağlı olarak tahmin süreçlerinde ve her türlü numunede kullanılırlar. Kayan pencere yöntemi, zaman serilerini denetimli öğrenme yaklaşımına dönüştürmek için kullanılabilir. Bu yöntemde önceki zaman adımları kullanılarak bir sonraki zaman adımındaki değer tahmin edilir. Kayan pencere yöntemini kullanmanın temel amacı, çıktının zaman serisine göre kesin olduğu durumlarda derin öğrenme algoritmalarını kullanmaktır. Bu yaklaşımın temel mantığı, o andaki (t) değerini inceleyerek bir sonraki t anındaki (t+1) değerini tahmin etmektir. Zaman serisinin derin öğrenme modelinde kullanılabilmesi için denetimli öğrenme yöntemine dönüştürülmesi gerekmektedir (Brownlee, 2017).

3.1.1. Modelin (LSTM-CLOUD) Tasarımı ve Yapılandırılması

LSTM-CLOUD sisteminin temel amacı, genel bulut ağ ortamındaki DDoS saldırılarını tespit etmek ve önlemektir. DDoS saldırılarını tespit etmek için ağ trafiği, imza tabanlı bir tespit yaklaşımıyla karakterize edilmekte ve bu iş LSTM-CLOUD sistemi tarafından yapılmaktadır. Bu bölümde ele alınan çalışmanın amacı, halka açık bir bulut ağ ortamında DDoS saldırılarını tespit etmek ve önlemek için LSTM tabanlı bir sistem (LSTM-CLOUD) tasarlamak ve bu sistemde kullanılacak bir LSTM tahmin modeli geliştirmektir. Bu amaca ulaşmak için araştırmada iki temel soruya yanıt aranmıştır:

- DDoS saldırılarını tespit etmek ve önlemek için LSTM algoritması kullanarak herkese açık bulut ortamında tasarlanacak sistemle ilgili olarak nasıl bir yaklaşım benimsenmelidir?
- Bu sistemde kullanılacak bir LSTM modelini geliştirmek için hangi deneysel çalışmalar yapılmalıdır?

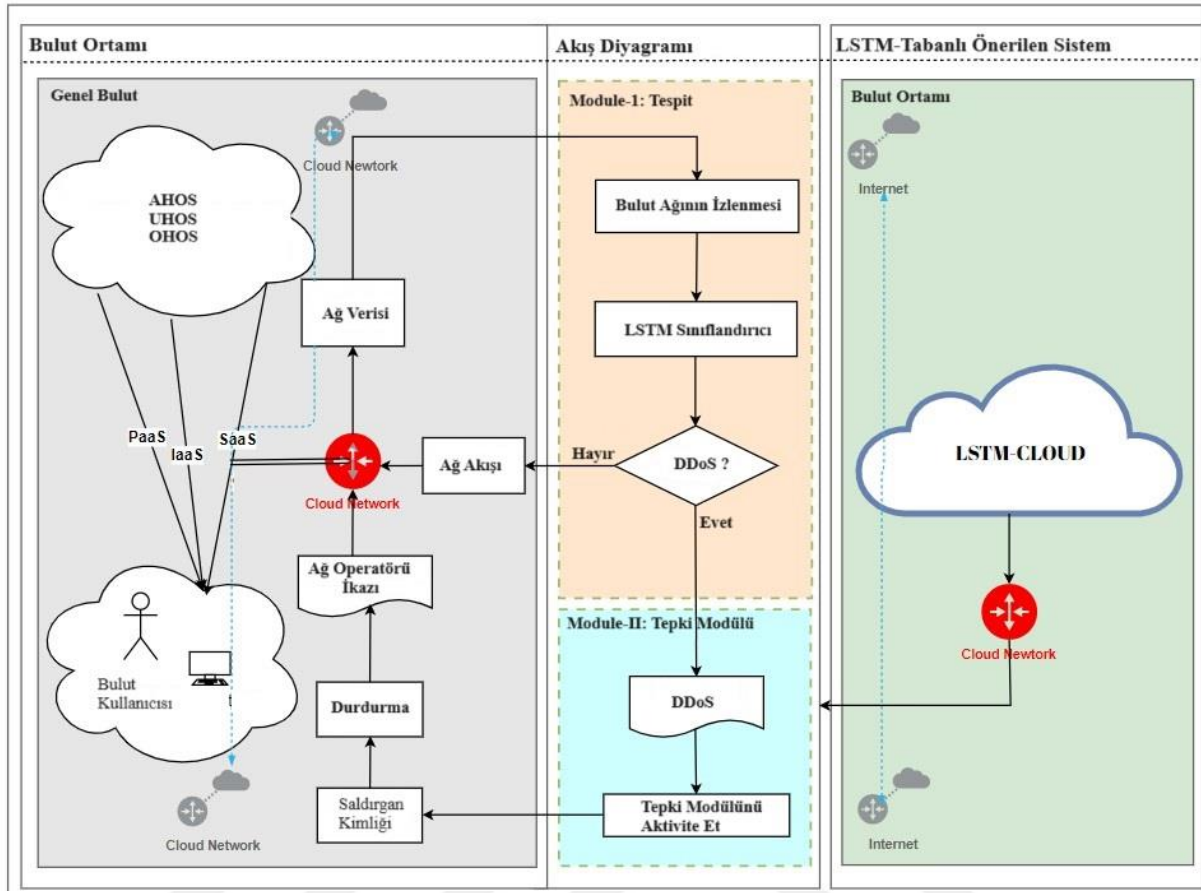
Araştırmada nicel yöntemler kullanılmıştır. Konuyla ilgili daha önce yapılan çalışmalardan yararlanmak amacıyla ilk olarak literatür taraması gerçekleştirilmiştir. Ardından, deneysel

araştırma yöntemi kullanılarak yüksek başarı oranına sahip bir LSTM modeli geliştirmek için çeşitli deneyler yapılmıştır. Katman ve çağ sayılarını değiştirerek yapılan deneylerde başarı oranlarında farklılıklar gözlemlenmiştir.

LSTM-CLOUD sistemi, DDoS saldırılarının tespiti ve önlenmesi için iki ana aşamaya ayrılmıştır: Siber Saldırı Tespiti Modülü ve Savunma Modülü.

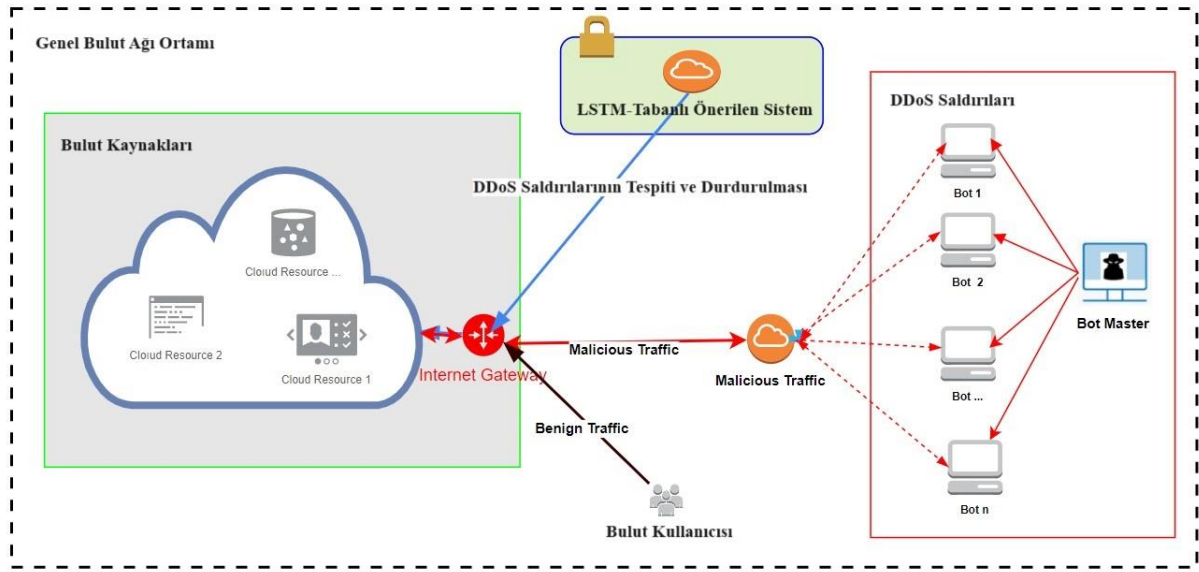
- **Siber Saldırı Tespiti Modülü:** Sistemin ilk modülü algılama modülüdür ve kötü niyetli ağ trafiğini hızlı, doğru ve gerçek zamanlı olarak tespit etmek bilgi güvenliği açısından kritik öneme sahiptir. Bu modül, anormalliklerin tespit edilmesi ve tanımlanması işlevini üstlenir. Çalışmada, bu modülün LSTM derin öğrenme modeli kullanarak DDoS saldırılarının oluşumunu tespit etmesi hedeflenmiştir. Diğer bir deyişle, kötü amaçlı ağ akışını tespit etmek için LSTM tabanlı bir sınıflandırma modeli geliştirilmiştir.
- **Savunma Modülü:** Sistemin ikinci modülü olan savunma modülü, tespit edilen anormalliklerin neden olduğu etkileri otonom politikalar uygulayarak en aza indirmeyi amaçlar. Bu modül, DDoS saldırılarının etkilerini insan müdahalesine gerek kalmadan azaltacak hafifletici politikaların uygulanmasından sorumludur. Bulut ağının güvenli ve emniyetli çalışmasını sağlamak için önceden belirlenen savunma mekanizmaları ve güvenlik politikaları bu modül tarafından devreye alınacaktır. Kötü amaçlı paketler tespit edildiğinde, analiz aralığındaki IP adresleri ve portlar şüpheli olarak değerlendirilecektir. Ardından, kötü amaçlı ağ paketlerinin düşürülmesi ve saldırganın bilgilerini içeren kara liste tablosunun oluşturulması sağlanacaktır. Yanlış pozitiflerden kaynaklanan hatalar ve arızalar bu modül tarafından ele alınacak, ayrıca ağ yöneticisine alarm verilmesi de bu modülün sorumluluğundadır.

Çalışma kapsamında tasarlanan DDoS tespit ve önleme sisteminin (LSTM-CLOUD) mimarisi **Şekil 3.2**'de sunulmuştur.



Şekil 3.2: LSTM-CLOUD Mimarisi

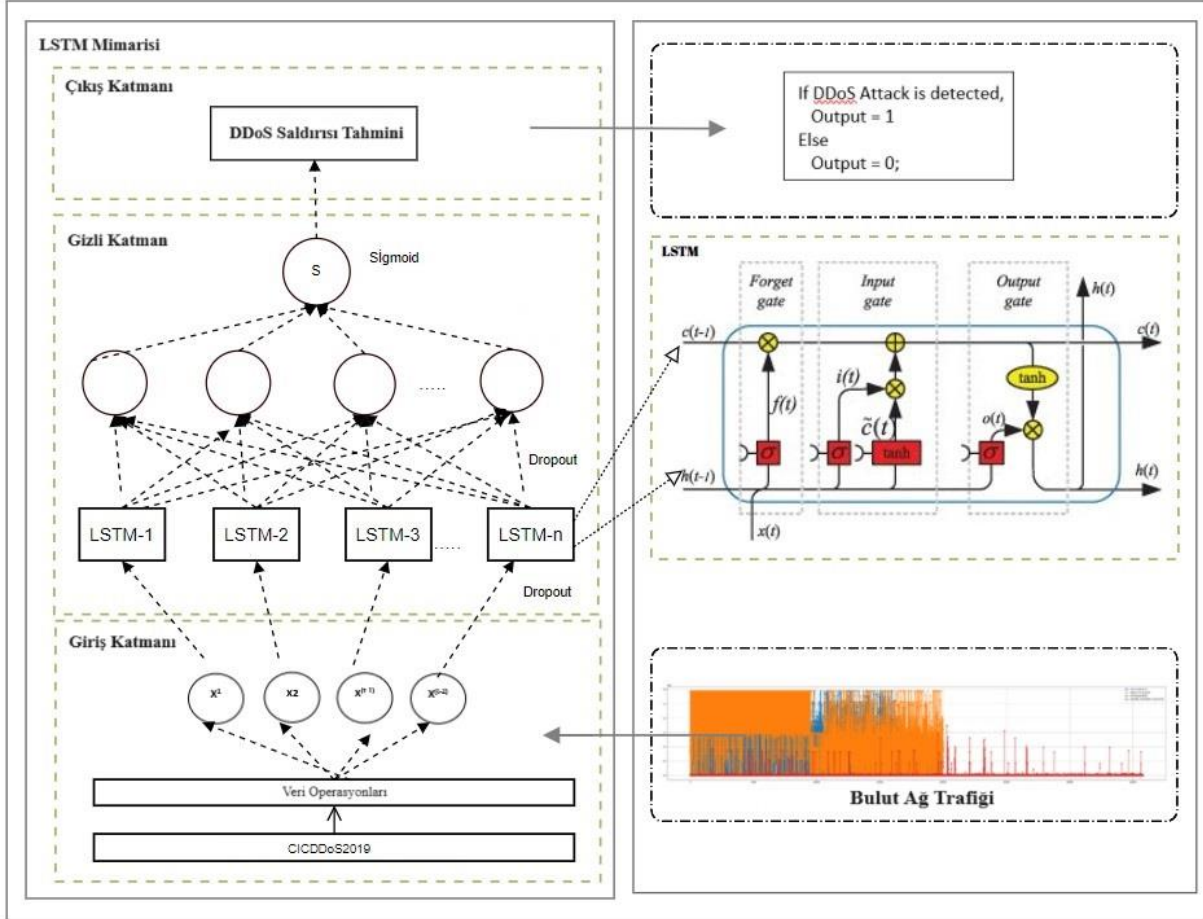
LSTM-CLOUD sisteminin kullanımına ilişkin bir senaryo Şekil 3.3’de sunulmuştur. Bu senaryo, çalışma kapsamında hayata geçirilip test edilmemiş olsa bile önerilen sistemin çalışma mekanizmasını açıklamak amacıyla örnek olarak verilmiştir. Senaryoya göre, siber saldırgan genel bulut ağ ortamında botnet adı verilen botların (bot 1, bot 2, ... bot n) uzaktan kontrolünü ele geçirmiştir. Bu şekilde siber saldırgan, olağanüstü sayıda istekle DDoS saldırıları gerçekleştirme yeteneği kazanmıştır. Böyle bir durumda, sistem tespit modülü aracılığıyla DDoS saldırılarını tespit edebilir. Siber saldırılar tespit edildiğinde, sistemin savunma modülü tespit edilen şüpheli akışları analiz eder ve bu saldırıların etkilerini en aza indirecek uygun önlemleri uygular.



Şekil 3.3: LSTM-CLOUD Sistemi Kullanım Senaryosu

LSTM-CLOUD sisteminin ilk modülü olan algılama modülünde kullanılmak üzere bu çalışmada bir LSTM modeli geliştirilmiştir. Bu model, LSTM hücresinin çalışma yapısına uygun olarak giriş, unutma ve çıkış katmanlarından oluşmaktadır. Model, geçmiş zaman serisi verilerini baz alarak saldırıları sınıflandırmakta ve gelecekteki saldırılara ilişkin tahminlerde bulunabilmektedir. Anormallik tespiti amacıyla geliştirilen bu model, LSTM tarafından tahmin edilen geçmiş trafik verilerini kullanır. Bir LSTM algoritması, yinelemeli olarak birbirine bağlı bir veya daha fazla bellek hücresinden oluşur. Bu hücreler, sırasıyla yazma, okuma ve silme işlemlerini gerçekleştirerek hafıza hücresinin davranışını kontrol eder. Bu yapıdaki bellek blokları, klasik sinir ağlarındaki gizli katmanların yerini alır. Giriş ve çıkış kapıları, türevlerin azalmasını ve kaybolmasını engeller. Unutma kapısı, bellek bloğundan bellek hücresine bilgi akışını kontrol ederek sonsuz döngüyü önler ve böylece ağın sınırlı bir belleğe sahip olmasını sağlar. Bellek bloğunun çıkışı, bir sonraki adımda hem bloğun girişine hem de kapılara yönlendirilir. Bu LSTM modeli, imza tabanlı algılama yaklaşımıyla ağ trafiğini karakterize ederek anormallik tespiti yapabilme yeteneği sağlar. Başlangıçta, LSTM modeli iki gizli katman, üç yoğun katman ve üç çıkış katmanından oluşan bir yapı ile tasarlanmıştır. Modelin tasarımı, eğitimi ve testi sırasında Keras Python kütüphanesi kullanılmıştır. Öğrenme oranını uyarlamak için Adam algoritması tercih edilmiştir. Ağın hafızada tutulmasını önlemek amacıyla uygun aktivasyon fonksiyonu ve dropout değeri belirlenmiştir. Modelin eğitiminde kullanılacak numune sayısını belirlemek için parti (epoch) büyüklüğü ayarlanmıştır. Doğruluk değeri ve ikili çapraz entropi kaybı fonksiyonu, temel performans ölçütleri olarak kullanılmıştır.

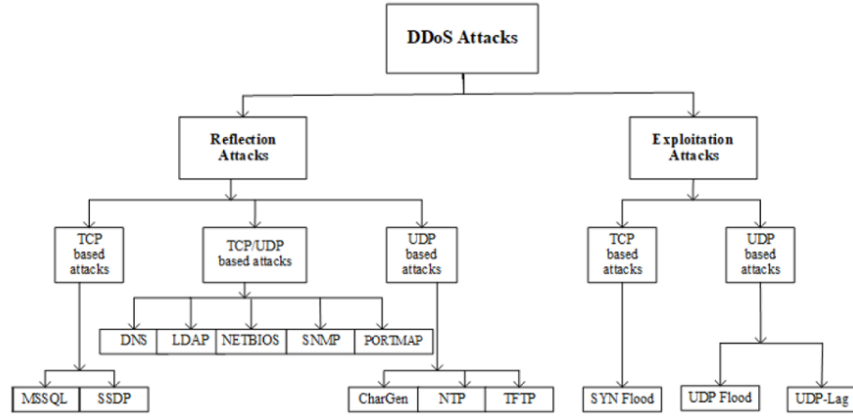
Veri setleriyle işlemler için Pandas kütüphanesi kullanılmıştır; ayrıca, Pandas kütüphanesindeki Shift fonksiyonu zaman serisini ileri veya geri taşımak için kullanılmıştır. Çalışmada kullanılan LSTM modelinin çerçevesi **Şekil 3.4**'de sunulmuştur.



Şekil 3.4: LSTM Tahmin Modeli

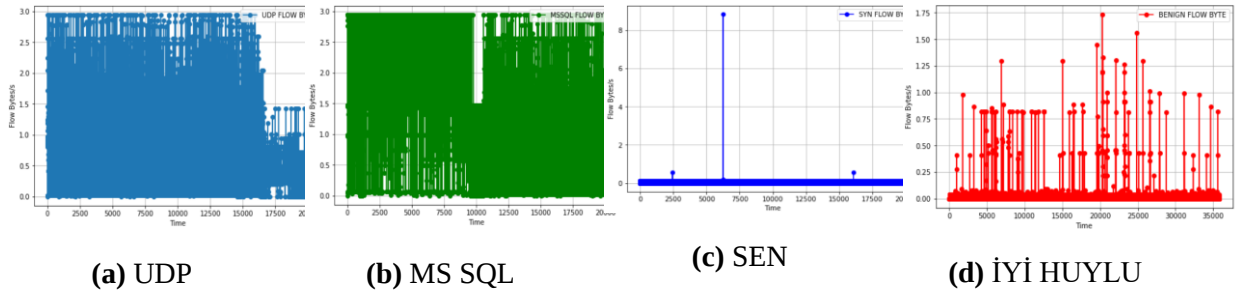
Geliştirilen LSTM modelinin etkinliğini ve verimliliğini değerlendirmek amacıyla çeşitli deneyler gerçekleştirilmiştir. Model, oluşturulduktan sonra sırasıyla eğitim ve test setleri kullanılarak eğitilmiş ve test edilmiştir. Deneylerde, modele 2, 4 ve 6 katmanlı yapılar uygulanmış ve her yapı için sırasıyla 5, 10 ve 60 eğitim dönemi (epoch) kullanılmıştır. Deneyler sırasında modelin eğitim ve test işlemlerinin hesaplama maliyeti yüksek olmuştur. Ancak, deneylerde kullanılan tüm kodlar, Google Colaboratory'nin sağladığı yüksek performanslı GPU ve TPU hizmetleri sayesinde başarıyla tamamlanabilmiştir. Modelin geliştirilmesi, eğitim ve test süreçleri deneyler sırasında birçok kez tekrarlanmıştır. Yapılan deneylerde CICDDoS2019 veri seti kullanılmıştır (Kanada Siber Güvenlik Enstitüsü, 2019). Bu veri seti, farklı türde DDoS saldırıları ve gerçekçi trafik profilleri içerir. Veri setindeki iyi huylu ve yaygın DDoS saldırıları,

CSV dosya formatında mevcuttur. DDoS saldırılarının CICDDoS2019 veri setindeki dağılımı Şekil 3.5'te gösterilmektedir.



Şekil 3.5: CICDoS2019 Veri Seti

CICDDoS2019 veri seti, 50.006.249 DDoS saldırısı ve 56.863 iyi huylu ağ trafiği örneği ile birlikte toplamda 50.063.112 gerçek dünya DDoS saldırı örneği içermektedir (Hussain, 2020). Veri setinde, UDP, SYN, DNS, LDAP, NetBIOS, NTP, MS SQL, UDP-Gecikme, PortMap ve SNMP gibi çeşitli saldırı türleri bulunmaktadır. Deneylerde MSSQL, UDP, SYN ve normal ağ trafiği verileri, 12 farklı DDoS saldırısı türünden seçilmiştir. Çalışmada kullanılan DDoS saldırıları ve normal trafik akışları Şekil 3.6'da sunulmuştur.



Şekil 3.6: Saldırı Türleri

Deneylerde, veri seti eğitim, doğrulama ve test setleri olarak üçe ayrılmıştır. Eğitim seti, sınır ağının ağırlıklarını ayarlamak için kullanılmıştır. Doğrulama seti, modelin parametrelerine ince ayar yapmak amacıyla kullanılmıştır. Test seti ise modelin performansını değerlendirmek için doğruluk tahmini amacıyla kullanılmıştır. Verilerin eğitimi öncesinde bazı ön işleme işlemleri uygulanmıştır. Bu süreçte, veriler temizlenmiş ve eksik (NaN) ile sonsuzluk değerleri giderilmiştir. Veriler, minimum değeri sıfır ve maksimum değeri bir olacak şekilde normalize

edilmiştir. İyi huylu trafik sıfır olarak, saldırılar ise bir olarak kodlanmıştır. Saldırı türleri, veri setinde sayısal değerlerle kodlanmıştır. Önerilen modelde, optimal özellikleri bulmak için derin öğrenme modeli çeşitli parametrelerle test edilmiştir. Veri setinde mevcut olan nitelikler arasından aşağıdaki özellikler seçilmiştir: "bit/byte", "paket", "kaynak IP adresi", "hedef IP adresi", "kaynak ve hedef portlar", "akış süresi" ve "fwd - paket uzunluğu - ortalaması". Bu bağlamda, çalışmada kullanılan ağ akışlarının dağılımı **Şekil 3.7**'de sunulmaktadır.



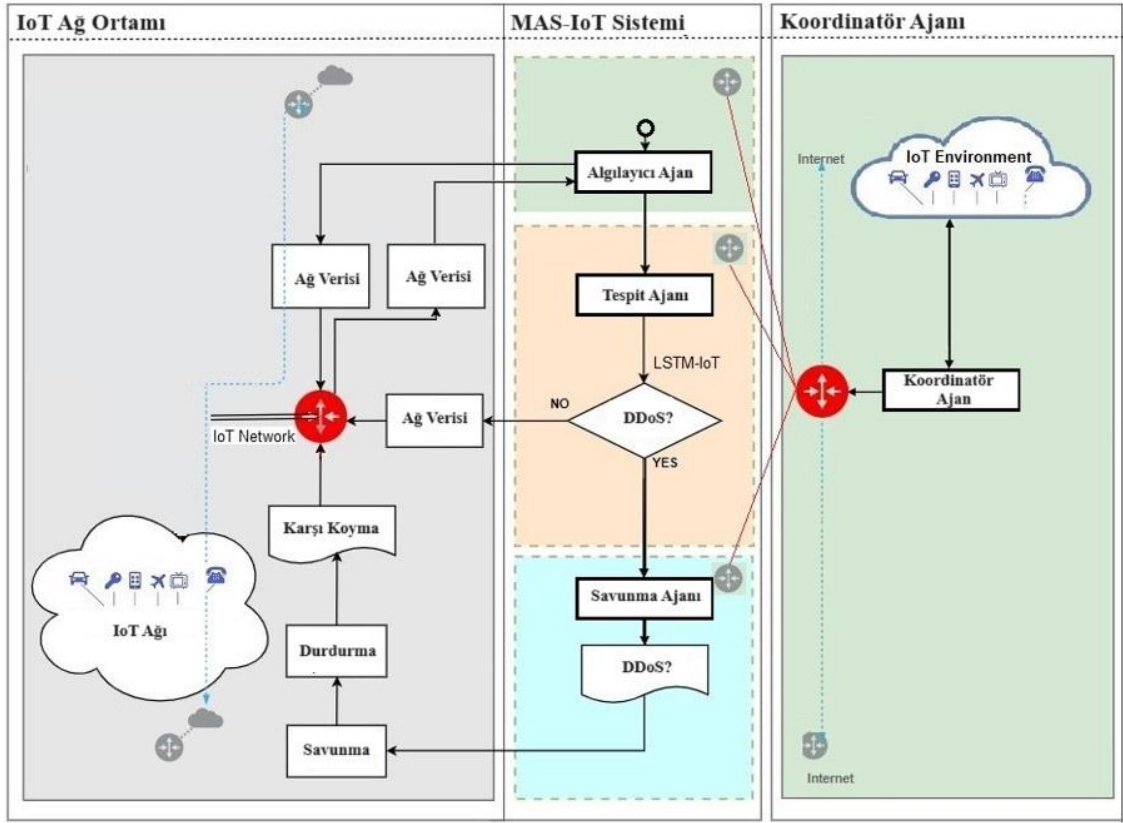
Şekil 3.7: CICDoS2019 Veri Seti Ağ Trafik

3.2. Ajan Tabanlı Derin Öğrenme (DÖ) Tabanlı Model Önerisi

3.2.1. Nesnelerin İnterneti (IoT) Ağları için Model Önerisi

3.2.1.1. Modelin (MAS-IoT) Tasarımı

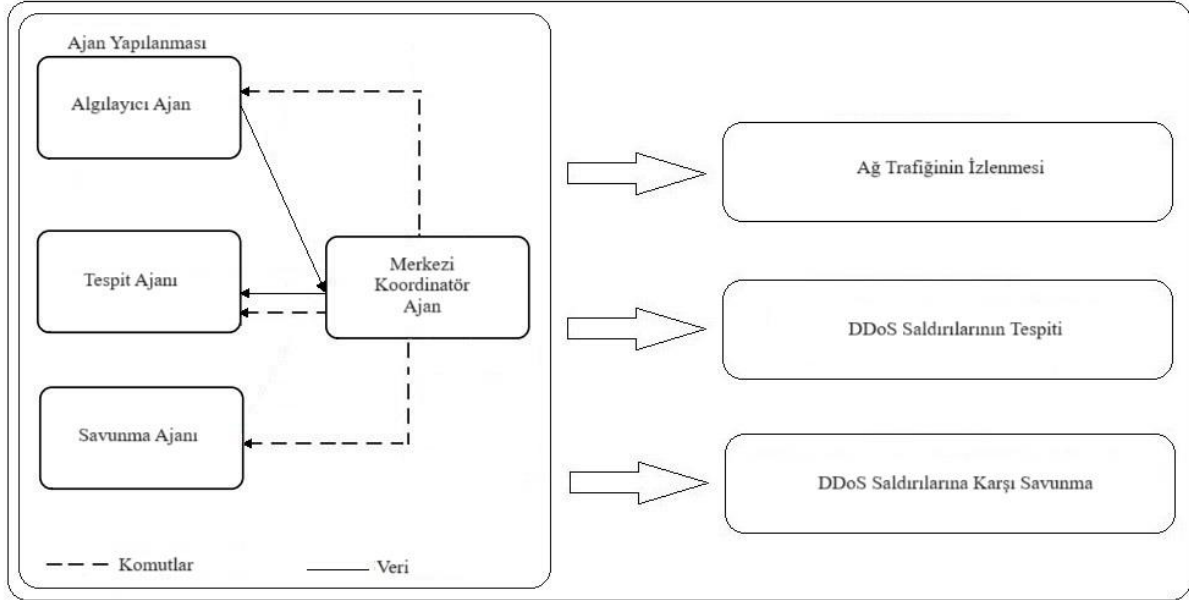
MAS-IoT sisteminin amacı, tanımlanmış ajanları kullanarak bir IoT ağında DDoS saldırılarını tespit etmek ve önlemektir. Bu sistemin mimarisi **Şekil 3.8**'de gösterilmektedir. Sistem, çeşitli türdeki ajanların iş birliği yoluyla çalışır: *Algılayıcı Ajan*, *Tespit Ajanı*, *Savunma Ajanı* ve *Koordinatör Ajan*.



Şekil 3.8: MAS-IoT Sistem Mimarisi

Sistemde tanımlanan ajan türlerinden koordinatör ajan, algılayıcı ajana ilk komutunu göndererek sistemin çalışmasını başlatır. Komutu aldıktan sonra, algılama ajanı belirli bir süre içinde IoT ortamından ağ trafiğine ilişkin verileri toplamaya başlar. Veri toplama işlemi tamamlandıktan sonra, algılayıcı ajanı koordinatör ajana bildirimde bulunur ve eş zamanlı olarak toplanan veri paketlerini tespit ajanına iletir. Bir sonraki aşamada, koordinatör ajan ikinci komutunu tespit ajanına gönderir. LSTM-IoT modelini kullanan algılama ajanı, potansiyel anormallikleri etkili bir şekilde tanımlayarak ve DDoS saldırılarını tespit ederek ağ paketlerini sınıflandırır. Bir DDoS saldırısı tespit edilirse, tespit ajanı derhal koordinatör ajana bilgi verir. Koordinatör ajan, savunma sistemini aktif hale getirmek için savunma ajanına gerekli komutu gönderir. Savunma ajanı komutu aldıktan sonra, MAS-IoT'nin siber savunma sistemini etkinleştirir. Bu aşamada, şüpheli IP'leri ve bağlantı noktalarını belirleyerek potansiyel saldırganları engeller ve ağ yöneticisini ile güvenlik personelini bilgilendirir. Tüm süreç boyunca, MAS-IoT sistemi IoT ağ trafiğini sürekli olarak izleyerek kolektif bir yaklaşımı kolaylaştırır. Algılama ajanı ağ paketlerini yakalar ve bunları diğer ajanlara iletir. LSTM tabanlı derin öğrenme modeliyle donatılmış algılama ajanı, ağ anormalliklerini etkili bir şekilde sınıflandırır ve bu sayede DDoS saldırılarını tanımlamak için önemli bir bileşen sağlar. Son

olarak, savunma ajanının hızlı eylemleri, IoT ağını saldırılardan korur ve ağın güvenliğini ile kesintisiz çalışmasını sağlar. Bu bağlamda, sistemdeki ajan türleri **Şekil 3.9**'da gösterilmektedir.

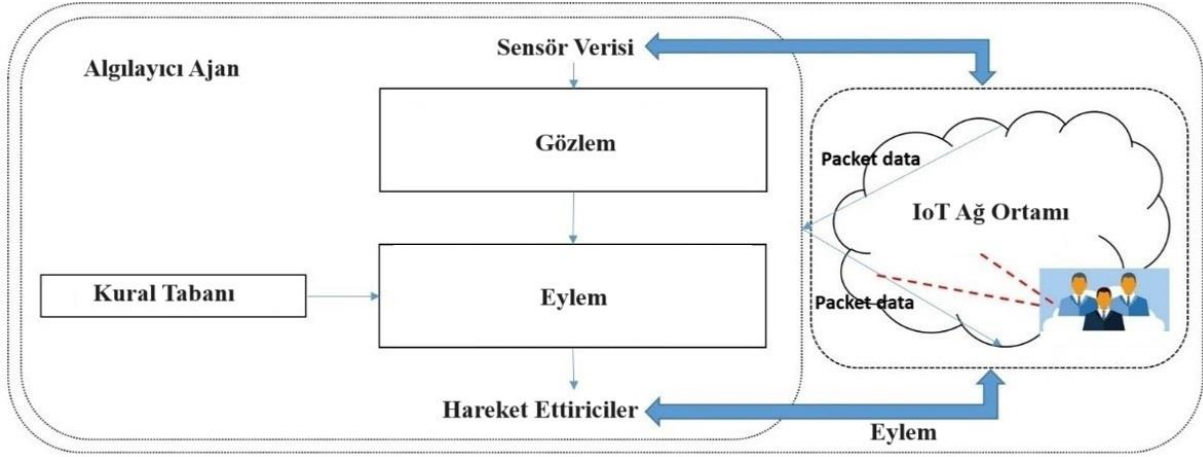


Şekil 3.9: MAS-IoT Sistemi Ajanları

Önerilen MAS-IoT sisteminde kullanılan ajanların açıklamaları aşağıda detaylı olarak sunulmuştur:

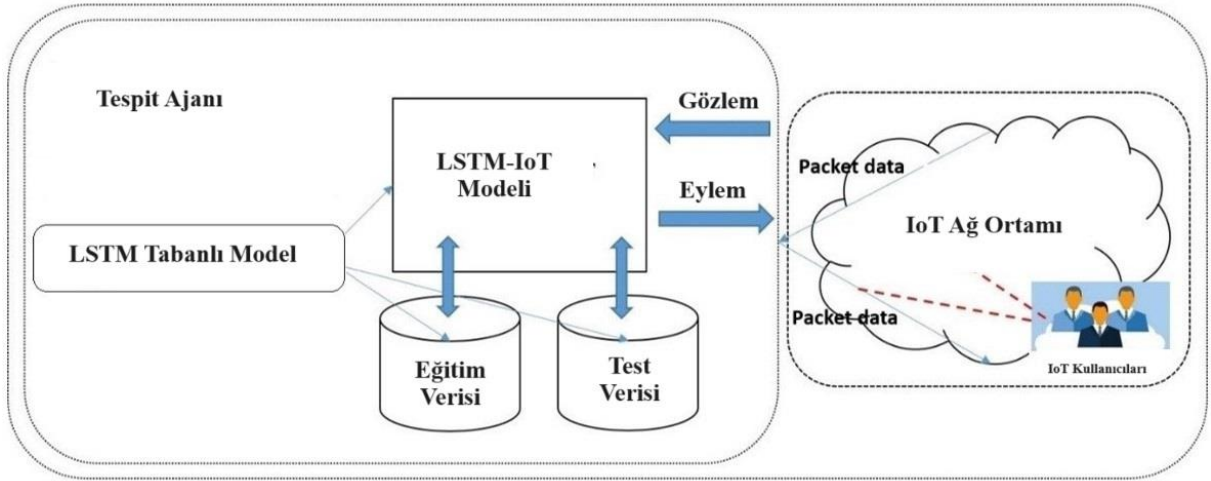
- **Algılayıcı Ajanı:** Bu ajanın birincil rolü, IoT ortamına giren paketleri izlemek ve bu paketlerden veri toplamaktır. Ajan, ağa bağlanarak IoT ağ paketlerini toplar, izler, okur ve yeterli alan varsa kendi deposuna ekler. Daha sonra bu paketleri sistemin ikinci tür ajanı olan algılama ajanına gönderir. Toplanan paketler, bir veri tabanı dosyasında saklanır ve tespit ajanının analiz ve tespit sürecini başlatması için girdi olarak kullanılır. Algılama ajanının ana işlevi, koordinasyon ajanıyla iletişim kurmak, ondan komutlar almak ve algılama ajanından topladığı ağ veri paketlerini göndermektir. Koordinasyon ajanı, algılayıcı ve tespit ajanları arasındaki iletişimi şifreli biçimde yönetir. Algılayıcı ajan, koordinasyon ajanının müdahalesiyle çalışmaya başlar ve ajan onu durdurana kadar görevine devam eder. Algılama ajanı, gerçek zamanlı algılamaya dayanarak doğru kararlar verdiğinde ve ortam tamamen gözlemlenebilir olduğunda etkili bir şekilde çalışır. Russell ve Norvig (2009) tarafından tanıtılan yaygın ajan modellerine uygun olarak, algılayıcı ajan basit bir refleks ajan olarak sınıflandırılmıştır. Refleks ajanı, mevcut algıya yanıt olarak bir eylemi seçen ve yalnızca algı geçmişini kullanan en basit ajan türüdür. Eylemler, algı geçmişine dayalı olarak gerçekleştirilen bir koşul-

eylem kuralına dayanır. Refleks ajanı, eylemlerinin gelecekteki sonuçlarını veya çevrenin mevcut durumunu mevcut algının ötesinde dikkate almaz. Algılayıcı ajan, ağ trafiğini izleyerek, paketleri toplayarak ve bunları tespit ajanına ileterek ağ güvenliğinde anormallik tespiti için gerekli verileri sağlar. Bu görevini yerine getirirken mevcut durumuna göre kararlar verir ve verilen eylemi gerçekleştirir. Bu nedenle, bu ajan türüne refleks ajan adı verilmiştir. Şekil 3.10’da bu ajan yapısı gösterilmektedir.



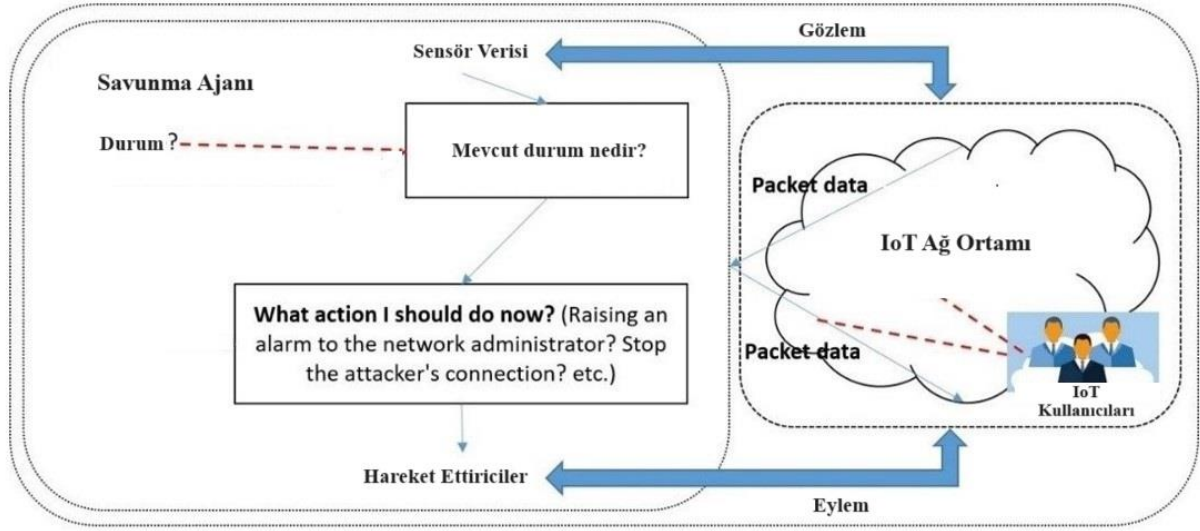
Şekil 3.10: Algılayıcı Ajan

- **Tespit Ajanı:** Bu ajanın birincil rolü DDoS saldırılarını tespit etmek olarak belirlenmiştir. Bu ajan türü bu fonksiyonu gerçekleştirmek için öncelikle bir siber saldırının gerçekleşip gerçekleşmediğini tespit eder. Eğer öyleyse, olayı koordinasyon görevlisine bildirir. Bu ajan, algılayıcı ajandan elde edilen paketleri LSTM-IoT modelini kullanarak analiz eden koordinasyon ajanı aracılığıyla algılayıcı ajanından ağ paketlerini talep eder. Bu paketler daha sonra LSTM-IoT modeliyle analiz edilir. LSTM algoritması, siber saldırılarla ilişkili geçmiş ağ trafiği verilerinin incelenmesi yoluyla gelecekteki saldırı modellerini tahmin etme konusunda uzmandır. Bazı özelliklerinden dolayı, bir tespit ajanı Russell ve Norvig (2009) bahsettiği ajan sınıflandırmasına tam olarak uymayabilir. Ancak bu tür ajan genel olarak “Model Tabanlı Refleks Ajanı” olarak kabul edilebilir. Reflex ajanlarından farklı olarak bu ajanlar, önceden belirlenmiş bir model veya plana göre eylemleri belirlemek için mevcut duruma ek olarak bir model kullanır. Şekil 3.11’de bu ajan yapısı yer almaktadır.



Şekil 3.11: Tespit Ajanı

- Savunma Ajanı:** Bu ajanın görevi tespit edilen DDoS saldırılarına karşı koymak olarak belirlenmiştir. Bu ajan, koordinasyon ajanından aldığı komutla savunma mekanizmalarını ve güvenlik politikalarını devreye alır. Bu modüldeki ajanlar, siber saldırganlara ait IP ve port gibi bazı bilgileri şüpheli olarak değerlendirip bloke etmektedir. Ayrıca koordinasyon ajanı aracılığıyla ağ yöneticisini uyarı alarmıyla bilgilendirmek, saldırganın bağlantısını kesmek, sistem güvenlik yöneticisine bildirimde bulunmak gibi bazı görevlerden de sorumludurlar. Russell ve Norvig (2009) ajan sınıflandırmasına göre, bir savunma ajanı model tabanlı refleks ajanı olarak sınıflandırılabilir. Bu ajan türü, koordinasyon ajanından aldığı komutlarla savunma mekanizmalarını ve güvenlik süreçlerini harekete geçirerek hareket eder. Ayrıca siber saldırganların IP adreslerini ve portlarını tespit ederek tespit eder ve engeller. Bu ajan türü ayrıca yanlış pozitiflerin neden olduğu hataları da ele alır. Özetle savunma ajanı, ağ yöneticisini uyarmak, saldırganın bağlantısını kesmek, koordinasyon ajanı aracılığıyla sistem güvenlik yöneticisine bilgi vermek gibi görevlerden sorumludur. **Şekil 3.12**'de bu ajan türünün yapısı yer almaktadır.



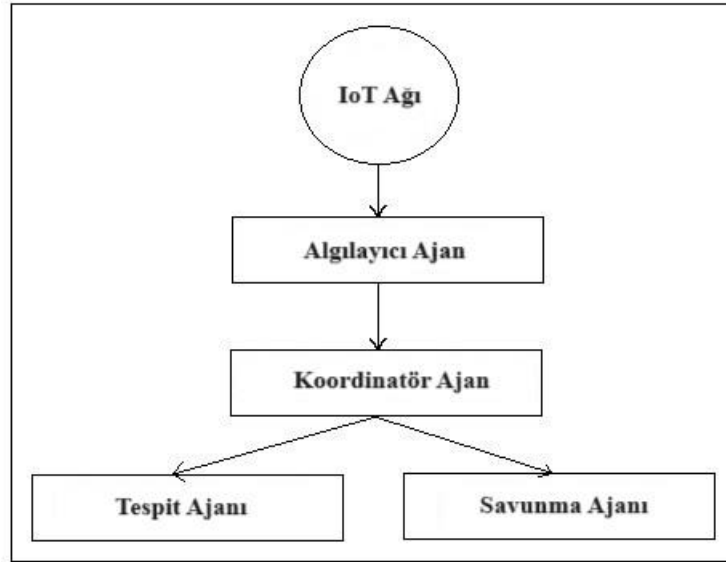
Şekil 3.12: Savunma Ajanı

- **Koordinasyon Ajanı:** Bu ajan türü, diğer ajanlar arasındaki faaliyetlerin karşılıklı bağımlılıklarını yönetmekle görevlendirilmiştir. IoT ortamında IoT cihazları üzerinde çalışan ajanların ortak çalışmasını düzenler. MAS-IoT'ta tanımlanan tüm ajanlar, bu ajan onları durdurana veya farklı bir görev atayana kadar çalışmaya devam eder. MAS-IoT mimarisinde ajanlar görevlerini gerçekleştirmek için iş birliği yapar. Sistemde ajanlar doğrudan iletişim kurmak yerine koordinasyon ajanına şifreli mesajlar gönderir. Bu nedenle bu temsilci, diğer temsilcilerden gelen bilgilere göre ne yapacağına karar verir. Bu bağlamda MAS-IoT sisteminde ana görevler ajanlar arasında paylaşılmaktadır. Ayrıca bu temsilci sistemin tüm operasyonlarından sorumludur. Bu ajan sistemdeki her ajanı kontrol eder ve bu bağlamda sistemde tanımlanan her ajan hem bu ajanla hem de birbirleriyle etkileşim halindedir. Bu ajan dışındaki tüm ajanlar, alana özel ajanlar olarak tanımlanırlar ve belirli görevleri yürütebilirler.

MAS-IoT sisteminde tanımlı ajanlar arasındaki işbirlikçi faaliyetler, birincil başlangıç noktası olarak algılayıcı ajan tarafından toplanan ağ paketlerine dayanır. Algılama ajanı, anormal trafik düzenlerini belirlemek ve bir siber saldırı olasılığını gösterebilecek işaretleri tespit etmek için ağdaki veri paketlerini dikkatlice analiz eder ve toplar. Önerilen sistemimizde toplanan bu veri paketleri, DDoS saldırılarının belirlenmesinde ve bunlara etkili bir şekilde karşı koymak için uygun yanıtların geliştirilmesinde önemli bir rol oynamaktadır. Üstelik bu paketlerin analizi, IoT ağının normal trafik davranışının anlaşılması ve dolayısıyla olası saldırılara karşı dayanıklılığının artırılması açısından oldukça önemlidir. Algılayıcı ajanı, bu veri paketlerini toplayıp ileterek, ajanlar arasındaki koordineli çalışmaları kolaylaştırır ve onların verilere

dayanarak daha etkili kararlar vermelerini sağlar. Algılama ajanı tarafından toplanan bu ağ paketleri çok çeşitli özellikler içermesine rağmen, bu çalışma imza tabanlı analiz için şu özelliklere odaklanmaktadır: “Bwd Paket Uzunluğu Ortalaması”, “Toplam Fwd Paketleri”, “Kaynak Portu”, “Hedef Portu”, “Akış Paketleri”, “Toplam Geriye Dönük Paketler”, “Akış Süresi”, “İleri Paket Uzunluğu Ortalaması”, “Akış Baytları/sn”, “Kaynak IP’si”, “Hedef IP’si”.

Şekil 3.13, bu ajanların her bir IoT cihazındaki düzenlemelerini göstermektedir. Burada amaç öncelikle bir IoT cihazındaki trafiği izlemek, ağ trafiğini sınıflandırarak bir saldırı olup olmadığını ortaya çıkarmak ve gerekli güvenlik görevlerini ortaklaşa uygulamaktır. Burada koordinatör temsilci esas olarak tüm ajanlar arasındaki iş birliğinden sorumludur. Bu yapı, MAS-IoT ajanlarının kolektif zekasını ve karar verme sürecini açıklamaktadır. Bu yapı, MAS-IoT temsilcilerinin kolektif zekâ ve karar alma süreci için veri paylaşımına veya alışverişine olanak tanıdığından, sistemde tanımlanan ajanlar, tek başına hareket eden ajanlara kıyasla toplu olarak saldırılara karşı daha etkili savunma yapabilir. Ayrıca planladıkları aksiyonlar dahil daha birçok konuyu birbirleriyle paylaşıp raporlayabiliyorlar.



Şekil 3.13: MAS-IoT Ajanların Kolektif Zekâsı

Bu çalışmada sunulan siber ajanlar, spesifik siber savunma yeteneklerine sahip küçük yazılımlar olarak tanımlanabilir. Onları diğer ajan ve yazılım türlerinden ayıran temel farklar, IoT siber uzayında otonom olarak hareket etmek, IoT siber uzayına uyum sağlamak, uyum içinde hareket etmek, diğer siber faktörlerle iletişim kurmak ve DDoS saldırılarına karşı uygun

siber savunma önlemlerini almak olarak sıralanabilir. Bu açıdan bakıldığında tanımlanan siber ajanlar, siber savunma faaliyetlerini insan müdahalesiyle karşılaştırılamayacak kadar inanılmaz bir hız ve doğrulukla yürütmeyi bilen yazılımlardır. Önerilen sistemde koordinasyon ajanı dışındaki tüm ajanlar mobil olarak tasarlanmıştır. Yani ajanlar IoT ortamına uyum sağlayabiliyor, diğer ajanlar ile iletişim kurabiliyor, iş birliği yapabiliyor, zekâ ve karar verme özelliklerine sahip olabiliyor. Bu özellikler ajanları diğerlerinden ayıran öne çıkan özelliklerdir. Yani bir ajan, bulunduğu IoT alanıyla sınırlı olmayıp, IoT ağındaki uzak noktalara hareket edebilir. Bu model aynı zamanda ajanlara heterojen ortamlarda veri işleme konusunda yeni bir anlayış sunar. Bu mobil ajanlar bulundukları ortamla sınırlı değildir ve kendi kararlarıyla belirli bir anda ağ üzerindeki başka bir düğüme geçebilirler. Taşıma sırasında çalıştırılabilir ve dinamik kodun yanı sıra, ajanların o ana kadar elde ettiği veriler de ajanlar ile birlikte taşınacaktır. Sistemdeki tüm ajan türleri, kodlarını ve verilerini uzak düğümlere taşıyabilecek yazılımlar olarak tasarlandıkları için statik değildirler. Sistemdeki mobil ajanlar sürekli hareket halinde olduklarından ve ağ üzerinde heterojen ortamlarda çalışabildiklerinden ortam koşulları değişebilmektedir. Sisteme yeni IoT ağı düğümlerinin eklenmesi veya kaldırılması, mevcut IoT düğümlerindeki kaynakların değiştirilmesi veya güvenlik gereksinimlerinin değiştirilmesi çevresel değişikliklere örnektir. İş paylaşan ajanlar için süreklilik esas olduğundan, merkezi sistemdeki bir mobil ajan sistemi, düğüm çökmelerini önlemek için siber saldırı tespitini farklı IoT düğümlerine kaydırabilmelidir. Önerilen sistemin sağladığı bir diğer önemli özellik ise sistemin bir tarayıcı yardımıyla uzaktan yönetilebilmesi ve izlenebilmesidir. Yani ağ üzerinde çalışan mobil ajan sistemi gerektiğinde uzak nodlardan yönetilebilmekte ve sistemdeki tüm ajanlar tek bir pencereden ve herhangi bir noddan izlenebilmekte ve kontrol edilebilmektedir.

MAS-IoT'de tanımlanan ajan türlerinin, IoT ağ trafiğini izlemek, DDoS saldırılarını tespit etmek ve bunlarla mücadele etmek için hızlı ve doğru kararlar alabilmesi gerekir. Ayrıca bu karar verme süreci çeşitli işlevlerden oluşur: Ağ trafiği verilerinin toplanması, ağ trafiğinin kötü amaçlı veya iyi huylu olarak sınıflandırılması, kötü amaçlı ağ paketlerine karşı gerekli eylemin planlanması, DDoS saldırılarına karşı doğru eylemin seçilmesi ve seçilen eylemin doğrulanmasının sağlanması. Bu amaçla her MAS-IoT ajanlarının, insanın karar verme sürecini taklit eden bireysel karar verme yollarına ihtiyacı vardır. Bu nedenle insan bilişine benzer şekilde ML veya DL algoritmalarını ve esnek karar verme süreçlerini kullanmaları gerekiyor. Örneğin IoT ağ trafiği DL algoritmaları ile sınıflandırılabilir. Bu aynı zamanda bu çalışmada IoT ağ trafiğini sınıflandırmak için neden LSTM tabanlı bir modelin geliştirildiğini de

açıklamaktadır. LSTM, önceki yinelemeli ağ algoritmaları tarafından çözülemeyen karmaşık ve uzun gecikmeli görevleri çözme yeteneğine sahiptir. LSTM mimarisi içerisinde verilerin ağ içerisinde unutulması konusunda bir tespit yapılır. Önceki ortamdaki veriler, modelin yeni ortamdaki işlevselliği için artık gerekli değilse kaldırılır. Bu kararlar zamana bağlıdır ve tahmin görevlerinde ve çeşitli veri örnekleri türlerinde kullanılır. Bu nedenle bu çalışmada öncelikle LSTM ile model oluşturulmuş, daha sonra eğitilip test edilmiştir. LSTM modelinin zaman serilerindeki başarısı dikkate alınarak çok değişkenli zaman serileri kullanılarak sınıflandırma yapılmıştır. MAS-IoT sistem tasarımı, IoT ağ güvenliğini artırmayı ve DDoS saldırılarına etkili bir şekilde karşı koymayı amaçlayan kapsamlı bir dizi temel unsuru kapsar. Bu unsurlar, kapsamlı insan müdahalesine olan ihtiyacı azaltan ve tehditlere gerçek zamanlı yanıtlar sağlayan gelişmiş DDoS saldırı tespit mekanizmalarını içerir. İşbirliğine dayalı temsilci yapısı, ajanlar arasında kesintisiz iletişim ve iş birliğini kolaylaştırarak olası saldırıları tespit etmek ve önlemek için kolektif ve koordineli bir yaklaşım sağlar. Ayrıca şifreli iletişim kanalları, sistemin genel güvenliğini artırarak ajanlar arasındaki hassas bilgi alışverişini korur. IoT ağlarının gerçek zamanlı izlenmesi, sistemin proaktif savunma yeteneklerini önemli ölçüde artırır. Bu, sistemin DDoS saldırılarını anında tespit etmesine ve karşı koymasına olanak tanır. Özellikle LSTM tabanlı modelin gücü, MAS-IoT sistemini DDoS saldırılarını hızlı ve doğru bir şekilde tespit etme yeteneğiyle donatarak etkili savunma stratejilerinin hızlı bir şekilde devreye alınmasını kolaylaştırıyor. MAS-IoT sistemi, bu unsurları ve özellikle LSTM-IoT modelini entegre ederek, DDoS saldırılarına karşı IoT ağ güvenliğini artırmak için umut verici bir çözüm sunduğu görülmektedir.

3.2.1.2. Modelin (LSTM-IoT) Yapılandırılması

LSTM-IoT modeli, bu çalışma kapsamında tasarlanan MAS-IoT sisteminin önemli bir bileşenidir. Bu model, sistem içindeki algılama ajanları için özel olarak geliştirilmiş derin bir yapay sinir ağıdır. LSTM-IoT, geleneksel RNN'lerde karşılaşılan kaybolan gradyan sorununu çözmek için LSTM mimarisini kullanır. Bu, modelin sıralı verilerdeki uzun vadeli bağımlılıkları ve zamansal kalıpları etkili bir şekilde yakalamasını sağlar ve bu nedenle IoT ortamlarındaki ağ trafiği gibi zaman serisi verilerini analiz etmek için oldukça uygundur. LSTM hücresinin çalışma yapısı gereği, LSTM-IoT modeli giriş, unutma ve çıkış katmanlarından oluşur. Bu mimarideki bellek hücreleri, bilgi akışını düzenleyen birbirine bağlı kapılar içerir: yeni bilgi için giriş kapısı, ilgisiz verileri atmak için unutma kapısı ve ilgili çıkış bilgisini belirlemek için çıkış kapısı. Bu bellek hücreleri, modelin önceki zaman adımlarından gelen

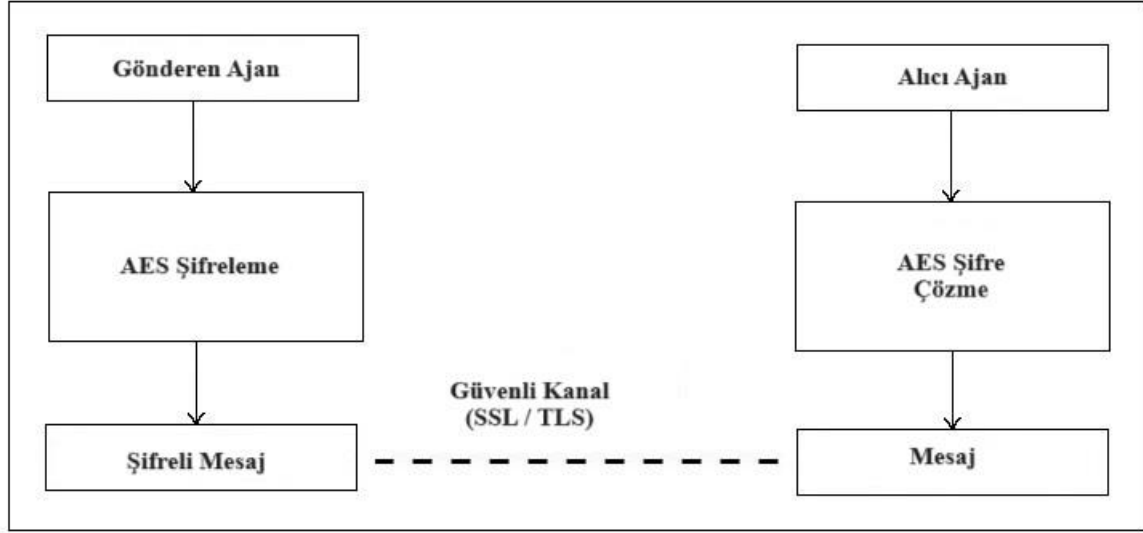
bilgileri depolamasına ve kullanmasına olanak tanır. Bu sayede model, DDoS saldırılarının doğru ve etkili bir şekilde tespit edilmesi için kritik kalıpları ve uzun vadeli bağımlılıkları kolayca tanıyabilir. Çalışmanın ana katkılarından biri, geçmiş zaman serisi saldırı verilerine dayalı sınıflandırma ve gelecekteki saldırıların tahmin edilmesini sağlayan LSTM-IoT modelinin geliştirilmesidir. Model, geçmiş trafik verilerini ve LSTM tarafından yapılan tahminleri kullanarak anormallikleri tespit etme ve potansiyel DDoS saldırılarını yüksek doğrulukla belirleme konusunda öne çıkmaktadır. Modelin ağ trafiği verilerindeki zamansal kalıpları ve değişiklikleri etkili bir şekilde yakalama yeteneği, önerilen MAS-IoT sisteminin genel performansını artırmaktadır. LSTM-IoT modelinin performansını optimize etmek için iki gizli katman, üç yoğun katman ve üç dropout (bırakma) katmanı eklenmiştir. Bu yapı, modelin derinliğini artırarak karmaşık desenleri ve özellikleri tanımlama kapasitesini geliştirmeyi hedefler. Bu bağlamda, Keras Python kütüphanesi kullanılarak LSTM sınıflandırma modeli oluşturulmuş, eğitilmiş ve test edilmiştir. Adam optimizasyon algoritması, uyarlanabilir öğrenme oranıyla optimizasyon sürecini hızlandırarak hızlı ve verimli model eğitimi sağlamıştır. Aşırı uyumu önlemek ve genelleme yeteneklerini artırmak için uygun dropout değerleri dikkatlice seçilmiştir. Ayrıca, model eğitimi için toplu iş boyutunun tanımlanması, eğitim verimliliğini ve genel model performansını artırmıştır. LSTM-IoT modelini değerlendirmek için doğruluk ve ikili çapraz entropi kaybı fonksiyonu gibi temel ölçümler kullanılmış ve bu bağlamda modelin güvenilirliği ve etkinliği doğrulanmıştır.

Genel olarak, LSTM-IoT modeli, IoT ortamlarındaki DDoS saldırılarını tespit etmek ve azaltmak için sağlam ve kapsamlı bir yaklaşım sunarak, önerilen MAS-IoT sisteminin temel yapı taşı olarak hizmet etmektedir. Esasen IoT güvenliği ile şifreleme arasında da yakın bir ilişki vardır. IoT güvenliği, IoT cihazlarının ve ağlarının korunmasını içerir ve şifreleme yöntemleri, IoT verilerinin gizliliğini ve bütünlüğünü korumada ve yetkisiz erişime veya kurcalamaya karşı korumada hayati bir rol oynar. IoT cihazları birbirleriyle iletişim kurduğunda ve veri alışverişinde bulunduğunda güvenliğe ihtiyaç duyar. Bu cihazlar hassas kişisel verileri, kullanıcı kimliklerini ve sağlık verilerini işleyip iletebilmektedir. Şifreleme, IoT verilerinin şifreli biçimde iletilmesine olanak tanıyarak, yalnızca yetkili alıcıların anlayabilmesini sağlayarak ve potansiyel kötü niyetli kişilerin veya saldırganların bu verilere erişmesini önleyerek bu verilerin korunmasına yardımcı olur. Şifreleme algoritmaları, verileri şifreleyen ve şifresini çözen matematiksel işlemleri içerir. Verinin bütünlüğünü sağlayan şifreleme aynı zamanda verinin iletim sırasında değiştirilmemesini de garanti eder. Ayrıca kullanıcıların,

verilerin bozulduğu veya değiştirildiği durumları tespit etmesine olanak tanır. Ayrıca literatürde IoT için geçerli, uygulanabilir bir şifreleme sisteminin tasarlanması konusunda tartışmalar bulunmaktadır. Alzubi ve diğ. (2023), IoT için Hermitian eğrilerine dayalı bir kriptosistem önermişlerdir. Önerilen MAS-IoT sisteminde ajanlar ayrıca her bir IoT düğümünden ajanlara, ajanlardan da IoT düğümlerine veya ajanlar arası olabilecek siber saldırılara da maruz kalabilirler. Bu nedenle ajanların güvenli bir şekilde çalışabilecekleri güvenli bir ortama ihtiyaç vardır. İletişim, kolektif zekâ ve karar almayı mümkün kıldığından MAS-IoT temsilcileri için bir zorunluluktur. Bu maksatla çalışmada en güvenli simetrik şifreleme algoritmalarından birisi olarak kabul edilen Gelişmiş Şifreleme Standardı (AES) kullanılmıştır.

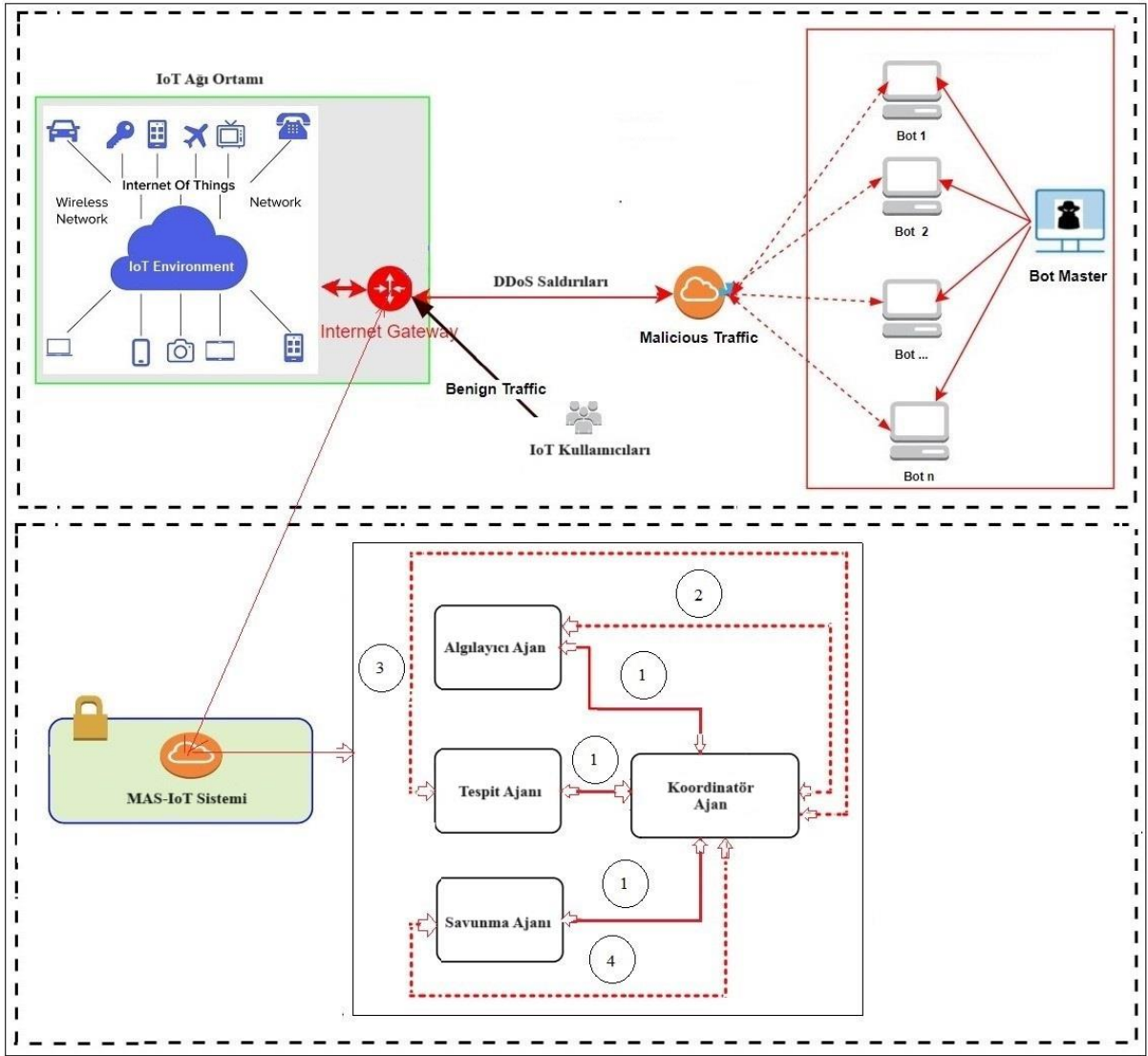
Ajan tabanlı siber savunma sistemi bağlamında, ajanlar arasında gönderilen mesajların gizli veya hassas bilgiler içerebileceği göz önüne alındığında, güvenli iletişimin sağlanması büyük önem taşımaktadır. Bu endişeyi gidermek için MAS-IoT ajanları iki temel şifreleme yöntemiyle donatılmıştır: Simetrik şifreleme için AES ve güvenli iletişim için genel/özel anahtar çiftlerine sahip SSL/TLS (Güvenli Yuva Katmanı/Aktarım Katmanı Güvenliği). AES şifrelemesinin uygulanması, ajanlar arasında paylaşılan mesajların gizliliğini ve bütünlüğünü garanti ederken, SSL/TLS, güvenli bir iletişim kanalı kurarak sistemin genel güvenliğini artırır. **Şekil 3.14**, ajanlar arasında AES şifrelemenin ve SSL/TLS tabanlı güvenli iletişimin entegrasyonunu göstermektedir. Bu şemada "Gönderen Ajan" mesajı şifreleyerek ağ üzerinden "Alıcı Ajana" göndererek AES şifreleme ve SSL/TLS ile güvenli bir iletişim kanalı oluşturur. "Alıcı Ajana" şifrelenmiş mesajı alır, SSL/TLS ile güvenli iletişim sağlar ve ardından orijinal mesajı elde etmek için AES şifrelemesinin şifresini çözer. Bu sayede iletişim güvenliği sağlanmakta ve mesajlar güvenli bir şekilde iletilmektedir. AES ve SSL/TLS şifreleme yöntemleri, mesajların IoT cihazlarına düz metin olarak iletilmemesini sağlayarak olası riskleri azaltır. Çalışmada, iletişim güvenliğini daha da artırmak için, elektronik veri şifrelemede yaygın olarak tanınan bir standart olan ve AES-128, AES-192 ve AES-256 gibi çeşitli anahtar uzunlukları sunan AES'i kullanılmıştır. Bu önlem, ağ içindeki ajanları tespit etmeye ve tehlikeye atmaya çalışan kötü amaçlı yazılımlardan kaynaklanan potansiyel tehditlere karşı koruma sağlar ve ajanların tespit edilme şansını azaltır. Ayrıca ajanların güvenli iletişimini sağlamak için sertifika bazlı sunucu modeli benimsenmiştir. Bu, SSL protokolünü kullanarak iletişimi kolaylaştırır ve ajanların güvenli bir şekilde mesaj alışverişinde bulunabilmesini sağlar. AES şifreleme ve SSL/TLS tabanlı güvenli iletişimden oluşan bu birleşik yaklaşımı benimseyerek, ajan-ajan etkileşimleri

için sağlam bir çerçeve oluşturmaktadır. Bu sayede değiştirilen mesajların gizliliğinin ve bütünlüğünün etkili bir şekilde korunmasının sağlanması mümkün olabilmektedir.



Şekil 3.14: Ajanlar Arası Şifreli Mesajlaşma Mimarisi

MAS-IoT sisteminin kullanımına ilişkin olası bir senaryo **Şekil 3.15**'de gösterilmektedir. Bu senaryoda siber saldırganların IoT ağındaki botları uzaktan kontrol ettiği varsayılmaktadır. Ayrıca siber saldırganlar DDoS saldırıları gerçekleştirebilmektedir. Böyle bir senaryoda, MAS-IoT sistemindeki her bir ajan, ilk olarak sıradan bir operasyon sırasında koordinatör ajan ve diğer komşu ajanlarla iletişim sağlığını kontrol eder (şekilde 1 rakamıyla gösterilen oklar). Bu operasyonun amacı, ajanlar arasındaki şifreli iletişim durumunu periyodik olarak izlemek ve ajanların iletişim arızalarını tespit etmektir.

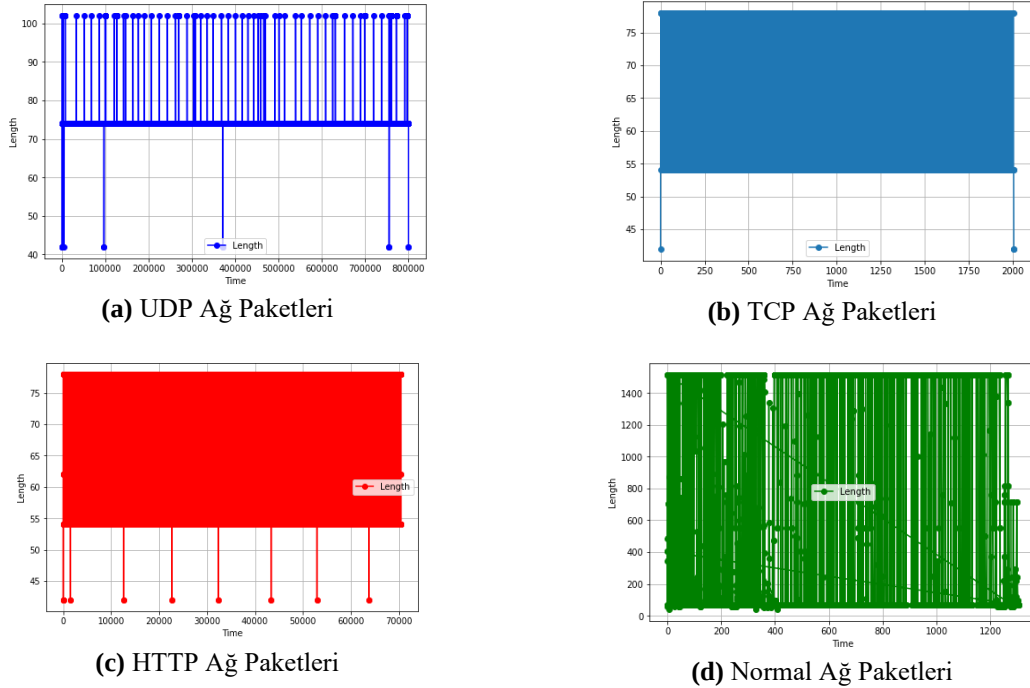


Şekil 3.15: MAS-IoT Sistemi Süreç Diyagramı

Çalışma kapsamında LSTM-IoT modelinin etkinliğini ve verimliliğini göstermek için çeşitli deneyler yapılmıştır. Deneylerde CIC-IoT-2022 veri seti güncel bir set olup IoT ile ilgili saldırı verilerini içermektedir. Bu veri kümesi, lambalar, kahve makineleri ve kameralar gibi farklı IoT fiziksel cihazlarından toplanan bilgileri içerir ve Kanada Siber Güvenlik Enstitüsü tarafından 2022'de yayımlanmıştır. CIC-IoT-2022 veri kümesi, profil oluşturma, davranış analizi ve güvenlik açığı için özel olarak oluşturulmuştur. IEEE 802.11, Zigbee tabanlı ve Z-Wave protokollerini kullanan çeşitli IoT cihazlarının test edilmesi. Veri seti; güç analizi, boşa kalma süresi analizi, etkileşim analizi, senaryo bazlı analiz, aktif kullanım analizi ve saldırı analizi dahil olmak üzere çeşitli testlerden geçmiştir. Özellikle Flood ve RTSP-Kaba Kuvvet gibi önemli saldırı türlerini yakalar. Veri seti IoT için büyük önem taşıyor çünkü güvenlik açıklarını

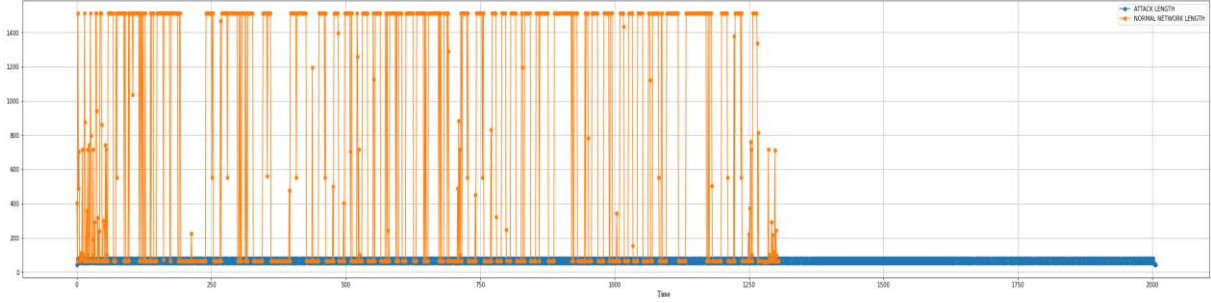
tespit etmek, cihaz davranışlarını analiz etmek ve zayıf noktaları belirlemek için kapsamlı bir kaynak görevi görmektedir.

Çalışma kapsamında bu veri setinde mevcut olan DoS saldırı verileri kullanılmıştır. Saldırı verileri ilk olarak “Wireshark” yazılımı kullanılarak “pcap” formatından “CSV” dosyasına dönüştürülmüştür. Daha sonra normal trafik verilerini saldırı trafiği verilerinden ayırt edebilmek için normal ağ verileri “0”, saldırı verileri ise “1” sayısı ile etiketlenmiştir. Bu süreç, tüm parametrelerin aynı pozitif ölçeğe sahip olması gereken bazı durumlarda faydalı olabilir. Veriler DL işleminden sorumlu olduğu için her türlü sorundan arındırılmış olması gerekir. Bu nedenle DL modeli üzerinde çalışmaya başlamadan önce veriler üzerinde bazı işlemler yapılmıştır. Bu işlemler arasında veri setinde tespit edilen saldırganların temizlenmesi, eksik verilerin tamamlanması, kontrollerin yapılması ve işlemlerin gerçekleştirilmesi yer almaktadır. Daha sonra veriler veri setinde “0” ile “1” arasında normalleştirilmiştir. Daha sonra kodlama işlemleri ve veri kümesindeki tüm saldırı türü adlarının dijitalleştirilmesi işlemleri tamamlanmıştır. DL modeli en uygun özelliklerin bulunması için farklı parametrelerle test edilmiştir. Veri setinde kullanılan özellikler; zaman, protokol, uzunluk, hedef, bilgi ve kaynaktır. Bu veri setinde yer alan normal ve saldırı trafik akışları **Şekil 3.16**'da gösterilmektedir.



Şekil 3.16: Saldırı ve Normal Trafik Akışı

Deneyler için kullanılan veri seti üç farklı set halinde dikkate alınmıştır: Eğitim, doğrulama ve test setleri. İmzaların karakterizasyonları, LSTM-IoT modeli için giriş değerleri görevi gören ve potansiyel saldırıların tanımlanmasına olanak tanıyan veri kümesi içindeki özniteliklerden türetilmiştir. **Şekil 3.17**'de ağ akışlarının dağılımını yer almaktadır.



Şekil 3.17: Veri Seti Ağ Akış Dağılımı

Deneyler sırasında LSTM-IoT modeli, her biri sırasıyla 5, 15 ve 60 farklı epoch ayarına sahip 2, 4 ve 6 katmandan oluşan çeşitli mimariler kullanılarak eğitilmiştir. Bu deneyler, kullanmak için herhangi bir kurulum gerektirmeyen ve GPU'lar ve TPU'lar dahil bilgi işlem kaynaklarına ücretsiz erişim sağlayan, barındırılan bir Jupyter Notebook hizmeti olan Google Colaboratory (Colab) tarafından gerçekleştirilmiştir.

3.2.2. Bulut Altyapısı için Model Önerisi

3.2.2.1. Modelin (CDA-CLOUD) Tasarımı

CDA-CLOUD sistemi, hacimsel DDoS saldırılarının tespiti ve azaltılması için özel olarak tasarlanmıştır. Bu bağlamda sistemde dört temel ajan türü bulunmaktadır:

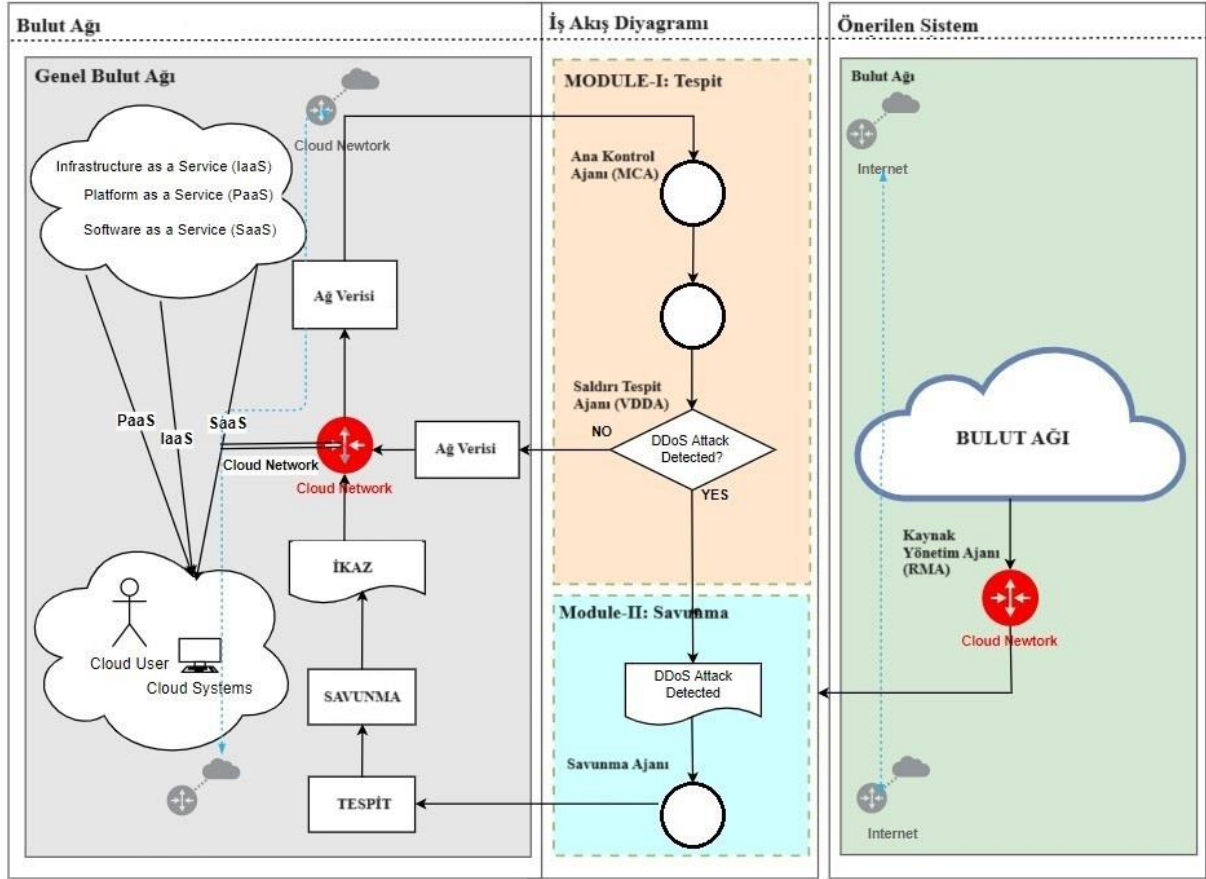
- **Hacimsel DDoS Saldırı Tespit Ajanı (VDDA):** Bu ajan, ağ trafiği analizi ve saldırı tespiti için bu çalışma kapsamında geliştirilen LSTM-CLOUD modelini kullanmaktadır. LSTM-CLOUD'un yüksek hassasiyetli saldırı sınıflandırma yeteneğinden yararlanarak, bulut ağı trafiğindeki saldırıları etkili bir şekilde tanımlar. Russell ve Norvig (2009) "Refleks Ajanları"nı çevresel koşullara anında tepki verebilen kararlar veren ajanlar olarak tanımlamıştır. Sonuç olarak VDDA'yı "Refleks Ajanı" olarak sınıflandırıyoruz. Bu sınıflandırma, ajanın değişen bulut ortamı koşullarına hızlı ve anında tepki verme kapasitesini ifade eder.
- **Savunma Ajanı (DA):** Tespit edilen saldırılara karşı ağ direncini artırmak için bu ajan, ağ operasyon düzenlemesini, trafik filtrelemeyi ve kaynak tahsisi optimizasyonunu

içerebilecek savunma önlemleri kullanır. Amacı ağ saldırıları nedeniyle oluşabilecek zararları en aza indirmek ve kesintisiz çalışmayı sağlamaktır. Russell ve Norvig (2009) "Model Tabanlı Refleks Ajanları", eylemleri yönlendirmek için önceden tanımlanmış modelleri veya planları kullanan, çevresel koşullara dayalı olarak daha bilinçli karar vermeyi kolaylaştıran ajanlar olarak tanımlamaktadır. Sonuç olarak, DA'yı "Model Tabanlı Refleks Ajanı" olarak sınıflandırıyoruz çünkü önceden belirlenmiş modellere veya ağ tehditlerine yanıt verirken iyi bilgilendirilmiş kararlar alma planlarına dayanıyor.

- **Kaynak Yönetimi Ajanı (RMA):** RMA, bulut ağı kaynaklarını verimli bir şekilde yönetme, saldırılar sırasında kaynak tükenmesini önleme ve ağ kesintilerini en aza indirmek için optimum kaynak tahsisini sağlama sorumluluğunu üstlenir. Özellikle saldırılar sırasında verimli bulut kaynak yönetimindeki rolü çok önemlidir. RMA, bir saldırı tespit edildiğinde anında harekete geçer. İlk olarak saldırı türü ve şiddeti gibi faktörleri değerlendirir. Daha sonra belirli saldırıya göre uyarlanmış en etkili savunma stratejilerini belirlemek için mevcut bulut kaynaklarını inceler. Bu süreç sırasında RMA, çeşitli bulut bölgeleri ve hizmetlerindeki kaynak kullanımını araştırır. Örneğin, saldırı ağ trafiğini artırırsa, ağ kaynaklarını bu trafiği verimli bir şekilde yönlendirecek şekilde yeniden yapılandırabilir. Ayrıca, belirli bulut hizmetleri veya bulut sunucuları saldırılara karşı daha savunmasızsa bunları izole edebilir veya gerektiğinde ek kaynaklar tahsis edebilir. Bu eylemler diğer bulut ağı hizmetlerinin kesintisiz çalışmasını sağlar. RMA'nın sıkı bulut kaynak yönetimi, bir saldırı sırasında kaynakların tükenmesini önler ve ağ genelinde optimum kaynak tahsisini sağlar. Temel olarak ana rolü, bulut kaynaklarını verimli kullanmak, kaynak tükenmesini önlemek ve saldırılar sırasında ağ dayanıklılığını arttırmaktır. Russell ve Norvig (2009), "Hedef Tabanlı Ajanlar" kategorisini, ajanların belirli hedeflere ulaşmak için tasarlandığı bir sınıf olarak tanımlamıştır. Bu çalışma kapsamında RMA ajanını, siber güvenlik önlemleri kapsamında kaynakları optimize ederken çevresel koşulları ve mevcut durumları analiz etmesi nedeniyle "Hedef Tabanlı Ajan" olarak sınıflandırılmıştır.
- **Ana Kontrol Ajanı (MCA):** Bu ajan, sistemin merkezi koordinatörü olarak görev yapar, diğer ajanların yönetimini denetler ve genel sistem işleyişini kolaylaştırır. MCA ajanının temel sorumlulukları arasında diğer ajanlar için görev tahsisi, iletişimin kolaylaştırılması, karar alma ve faaliyetlerin koordinasyonu faaliyetleri yer almaktadır. Russell ve Norvig (2009), belirli bir fayda fonksiyonunu optimize etmek için kararlar

veren ajanlar olan "Yardımcı Programa Dayalı Ajanlar" kavramını belirtmektedirler. Ek olarak, "Öğrenen Ajanlar"ı deneyimsel öğrenme yoluyla zaman içinde performanslarını artıran ajanlar olarak tanımlamaktadırlar. Sonuç olarak MCA, "Yardımcı Program Tabanlı Ajan" veya "Öğrenme Ajanı" olarak sınıflandırılabilir.

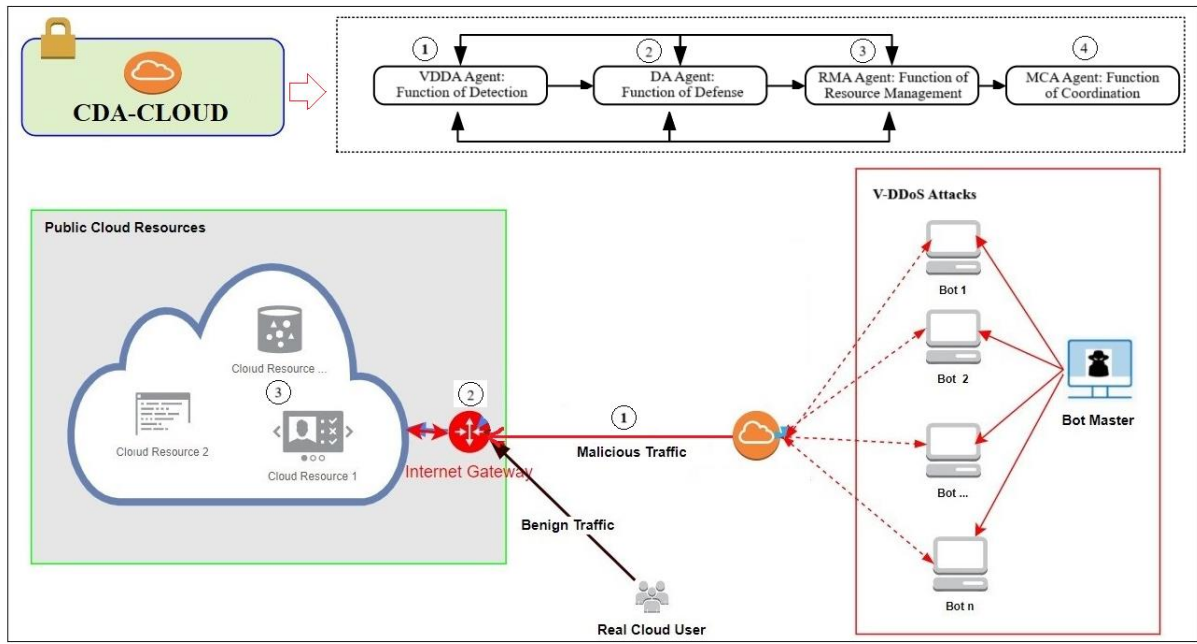
Çalışma kapsamında tasarımı yapılan CDA-CLOUD Mimarisi **Şekil 3.18**'de gösterilmektedir. Bu tasarıma uygun olarak VDDA ajanı, ağ trafiğini bulut sunucularına yönlendirir ve anormallik tespiti için YZ tabanlı analizden yararlanır. Bir saldırı tespit edildiğinde bu bilgi MCA'ya iletilir. Saldırı bildirimini aldıktan sonra MCA, etkili savunma önlemlerini uygulamak için DA'yı etkinleştirir. DA ajanı, ağ trafiğini yönetir, saldırı trafiğini filtreler ve kötü amaçlı kaynakları engeller, böylece ağın dayanıklılığını artırır. Eş zamanlı olarak RMA ajanı, ağ kaynağı tahsisini verimli bir şekilde denetlemek, dengeyi korumak ve ağ performansını optimize etmek amacıyla kaynak tükenmesini önlemek için devreye girer. Koordinatör görevi gören MCA ajanı, tüm diğer ajanları senkronize eder ve sistem yönetimini gerçekleştirir. Saldırı tespiti, savunma önlemleri, kaynak yönetimi ve genel koordinasyon arasındaki bu etkileşim, bulut tabanlı hacimsel DDoS saldırı önleme sisteminin temel iş akışını oluşturur. Etkili saldırı tespiti ve müdahalesi, bu ajanlar arasında yakın iş birliği ve koordinasyon gerektirir.



Şekil 3.18: CDA-CLOUD Mimarisi

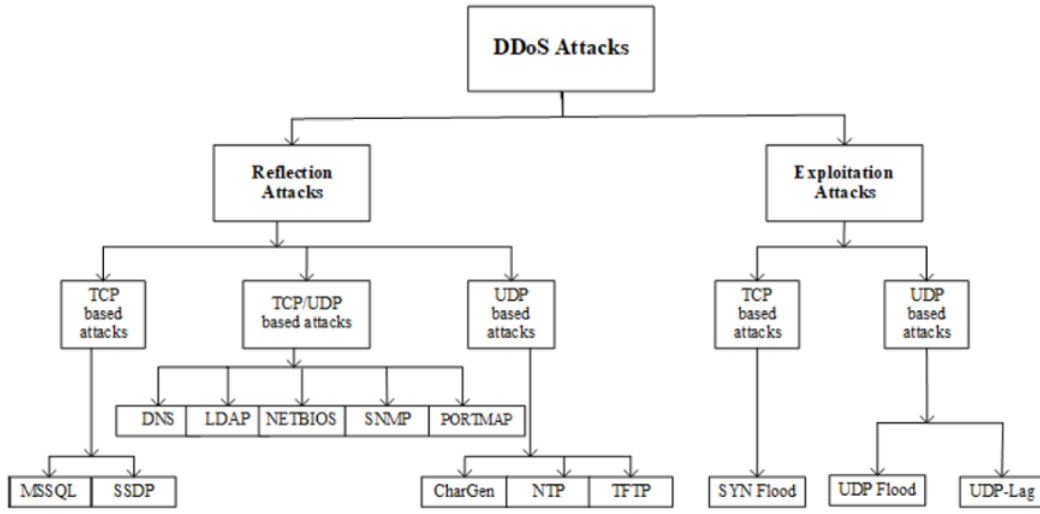
Hacimsel DDoS saldırılarına etkili bir şekilde karşı koymak ve azaltmak için CDA-CLOUD sisteminin işleyişini gösteren pratik bir senaryo Şekil 3.19’da sunulmuştur. Bu senaryoda, bir bulut ağı içerisinde birkaç uzaktan kumandalı botun (bot 1, bot 2, ...) iş birliğini içermektedir. Bu saldırganlar, bulut bilişim ağındaki hizmetleri kesintiye uğratabilecek saldırılar düzenleyebilir. Bu potansiyel tehdide yanıt olarak uzman ajanlar ile donatılmış CDA-CLOUD sistemi harekete geçmektedir. Siber saldırı tespitine yönelik olarak tasarlanan sistemde her bir ajanın farklı bir rolü vardır ve yaklaşan saldırılara karşı yapılandırılmış bir savunma mekanizması oluşturmak için birlikte çalışırlar. Ağ trafiğini analiz etmek için tasarlanan LSTM-CLOUD modelini kullanan VDDA ajanları, ağ trafiğindeki olağandışı kalıpları hızla incelerler. Saldırı işaretlerini belirledikten sonra anında alarm verirler veya tehditlere karşı önlem alırlar. Saldırı tespiti üzerine devreye giren DA ajanları, ağ trafiğini filtreleyerek, ağ operasyonlarını kontrol ederler ve ağ kaynaklarını dengeli kullanarak ağ dayanıklılığını artırır. Sistemde tanımlı tüm ajanlar esasen önceden tanımlanmış modellere dayalı savunma stratejileri uygularlar. Saldırı sırasında ağ kaynaklarını ustaca yöneten RMA ajanları, kaynakları en iyi şekilde dağıtarak kaynak tükenmesini önlerler. Bu ajanların amaçları ağ performansını optimize etmektir. MCA

ajanları diğer ajanları denetleyerek sistemin genel koordinasyonunu sağlarlar. Görev atamaktan iletişimi ve karar almayı kolaylaştırmaya kadar çeşitli görevleri yerine getirirler. Sistemin düzgün çalışmasını sağlamadaki rolleri çok önemlidir. Özet olarak VDDA ajanları, bir saldırı tespit edildiğinde MCA ajanlarının karar vermesine yardımcı olmak için anormal trafik hakkındaki bilgileri iletirler. MCA ajanı, ağ hasarını en aza indirmek ve saldırıları filtrelemek için RMA ajanı ile iş birliği yaparken, etkili savunma önlemlerini uygulamak için DA ajanını etkinleştirir. MCA ajanı, RMA ajanından elde edilen bilgileri kullanarak genel sistem performansını optimize etmek için ağ kaynaklarını düzenler.



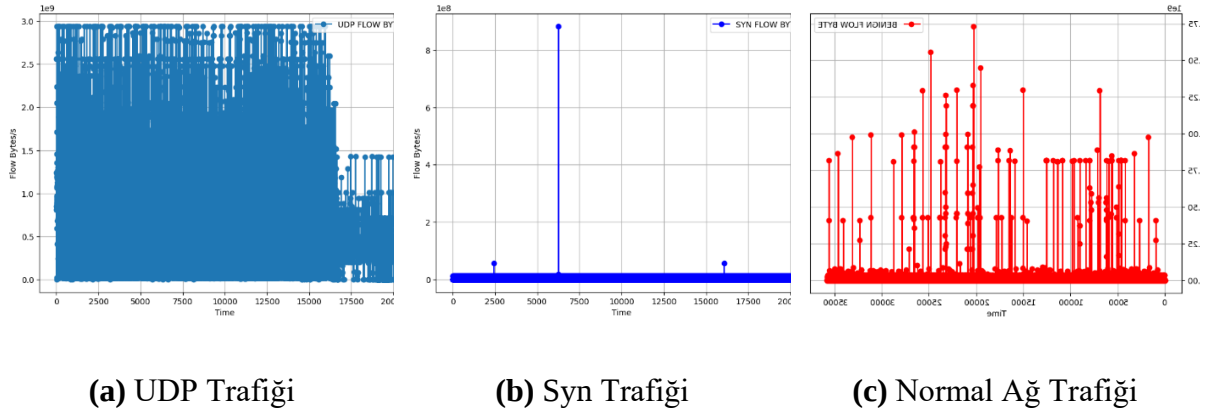
Şekil 3.19: CDA-CLOUD Sisteminin Kullanım Senaryosu

Çalışma kapsamında LSTM-CLOUD modelinin performansını değerlendirmek için bir dizi deney gerçekleştirilmiştir. Araştırmamızın deneysel aşaması için "Google Colaboratory" tarafından sunulan TPU hizmetlerinin hesaplama yeteneklerinden yararlanılmıştır. Bu deneyleri gerçekleştirmek için hem DDoS saldırılarını hem de zararsız ağ trafiği örneklerini içeren CICDDoS2019 veri kümesini kullanılmıştır. Araştırmamız özellikle hacimsel DDoS saldırılarını hedef aldığından, veri kümesindeki mevcut saldırı türlerinden yalnızca hacimsel ve normal ağ trafiği verileri kullanılmıştır. **Şekil 3.20'**de DDoS saldırılarının CICDoS2019 veri kümesi içindeki dağılımının grafiksel gösterimi gösterilmektedir.



Şekil 3.20: CICDDoS2019 Veri Seti

UDP ve SYN, CICDDoS2019 veri kümesindeki hacimsel DDoS saldırı türlerini temsil ettikleri için özel olarak seçilmiştir. Veri kümesindeki hacimsel verilerin özel olarak kullanılması, hacimsel DDoS saldırılarının gerçek zamanlı tespiti ve azaltılması için tasarlanmış çok ajanlı bir savunma sisteminin geliştirilmesini kolaylaştırdığı için bu çalışmada çok önemlidir. Çalışmada kullanılan saldırı türlerine ait trafik akış hızları Şekil 3.21'de üç ayrı grafikte gösterilmektedir.

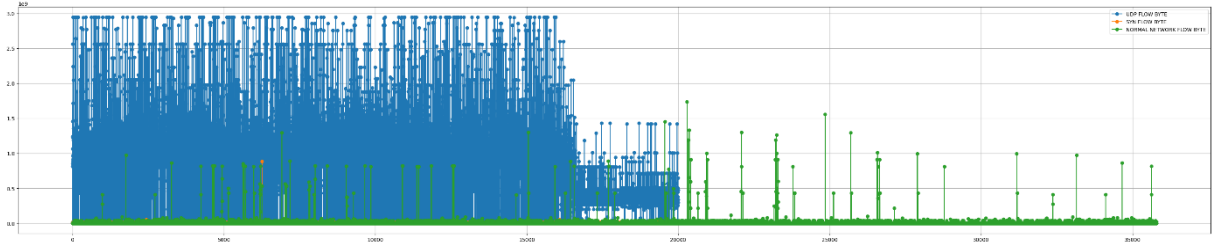


Şekil 3.21: Hacimsel DDoS ve Normal Ağ Trafığı

3.2.2.2. Modelin (LSTM-CLOUD) Yapılandırılması

Yapılan deneysel çalışmalar kapsamında saldırı verileri ve zararsız veriler üç ayrı gruba kategorize edilmiştir. Eğitim seti, eğitim sırasında LSTM-CLOUD modelinin sinir ağı ağırlıklarının yinelemeli olarak yeniden hesaplanması için kullanılmıştır. Doğrulama seti,

eğitim sırasında LSTM-CLOUD modelinin performansını değerlendirmek ve aşırı uyumu tespit etmek için kullanılmıştır. Son olarak test verileri, modelin eğitim ve doğrulama süreçlerinden sonra gerçek dünyada nasıl performans gösterdiğini değerlendirmek ve genelleme yeteneğini test etmek için kullanılmıştır. Veri setinin kalitesini artırmak, anlamlı sonuçlar elde etmek ve analizlerin doğruluğunu korumak amacıyla veri temizleme işlemleri gerçekleştirilmiştir. Spesifik olarak, veri kümesindeki "iyi huylu" terimini "0" rakamıyla değiştirilmiş ve saldırılar için "1" rakamını kullanılmıştır. Bu dönüşümde, saldırı türlerinin sayısal değerlerle temsil edilmesi amaçlanmıştır. LSTM-CLOUD modeli çeşitli parametrelerle test edilerek geliştirilmiştir. Kullanılan özellikler arasında saniye başına bit sayısı, saniye başına paket sayısı, kaynak ve hedef IP adresleri, kaynak ve hedef bağlantı noktaları, akış süreleri ve iletilen paketlerin ortalama uzunluğu yer almaktadır. LSTM-CLOUD modelinin geliştirilmesinde kullanılan karşılaştırmalı saldırı trafiğine ilişkin trafik akışlarının dağılımı **Şekil 3.22**'de gösterilmektedir.



Şekil 3.22: Ağ Trafik Akışı Dağılımı

3.3. Derin Pekiştirmeli Öğrenme (DRL) Tabanlı Model Önerisi

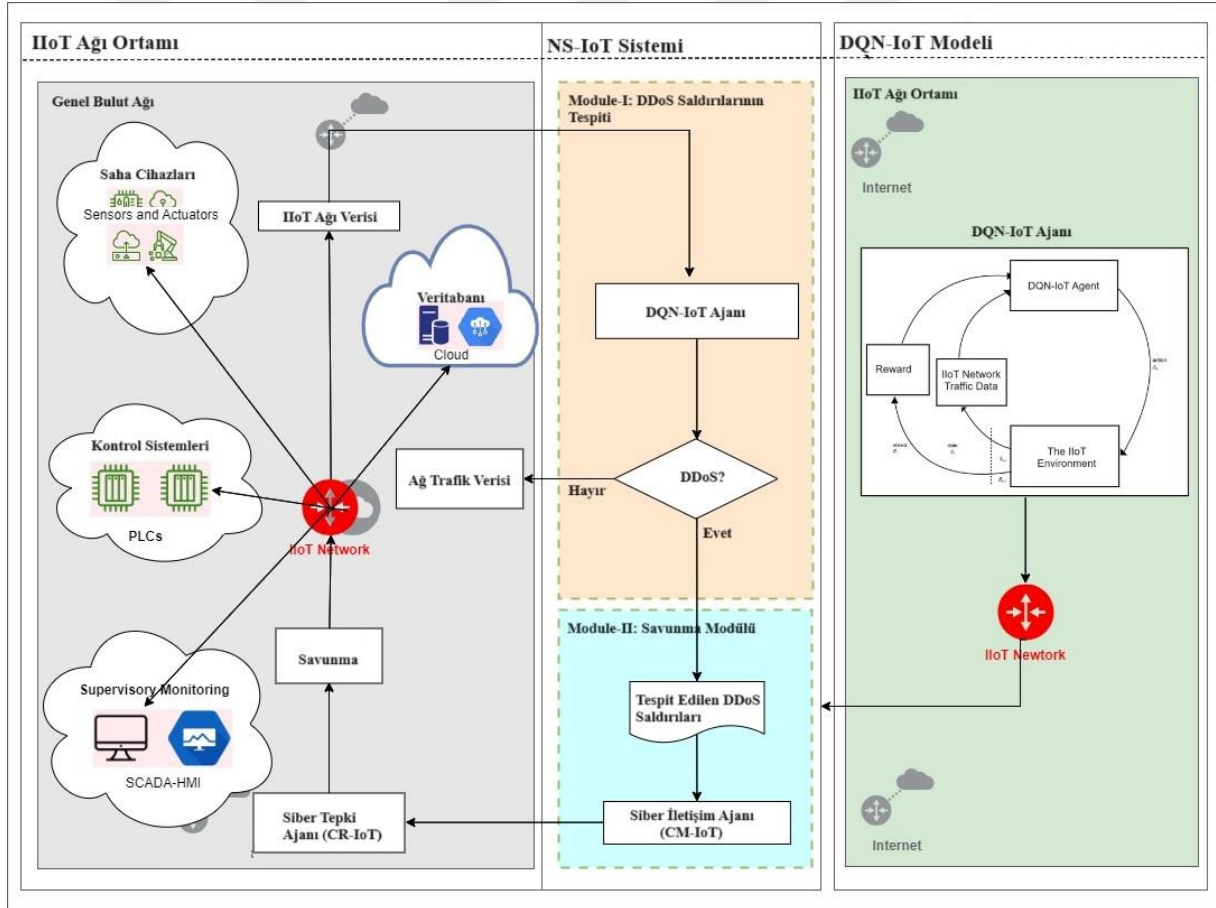
3.3.1. Modelin (NS-IoT) Tasarımı

NS-IoT sisteminin temel amacı, IIoT ortamındaki DDoS saldırılarını tespit etmek ve önlemektir. NS-IoT sisteminin mimarisi **Şekil 3.23**'te gösterilmektedir. NS-IoT sistemi iki temel modülden oluşur: Algılama Modülü ve Savunma Modülü.

- **Algılama Modülü:** Bu modülün amacı, bu çalışmada geliştirilen DQN-IoT ajanını kullanarak DDoS saldırılarını tespit etmektir. Ayrıca bu modül, DDoS saldırılarının ötesindeki anormallikleri tespit etmek için ağ trafiğinin sürekli izlenmesinden, olası güvenlik ihlalleri için sistem davranışının analiz edilmesinden, şüpheli faaliyetler için gerçek zamanlı uyarılar sağlanmasından ve algoritmalarını güncelleyerek gelişen saldırı yöntemlerine uyum sağlamaktan sorumludur. Bu amaçla RL tabanlı DQN yaklaşımını

temel alan bir ağ trafiği sınıflandırma ajanı geliştirilmiştir. DQN-IoT ajanının değerlendirmesi sonraki bölümde özetlenmiştir.

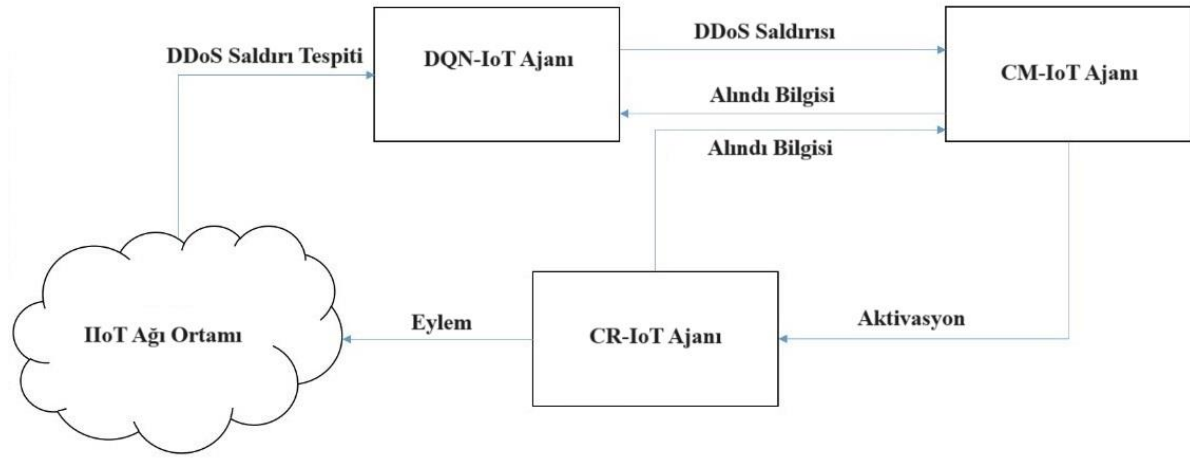
- **Savunma Modülü:** Bu modülün amacı, tespit edilen DDoS saldırılarının etkisini insan müdahalesine gerek duymadan en aza indirecek azaltma politikaları uygulamaktır. DDoS saldırısı durumunda bu modül, saldırıya karşı koymak için önceden tanımlanmış savunma mekanizmalarını ve güvenlik politikalarını otomatik olarak etkinleştirecektir. Örneğin, ağ trafiğinde kötü amaçlı paketler tespit edildiğinde sistem, ilgili IP'leri ve bağlantı noktalarını şüpheli olarak işaretleyecektir. Daha sonra kötü amaçlı paketler düşürülecek ve saldırganın bilgileri kara listeye kaydedilecektir. Bu modül için tasarlanan CR-IoT ajanı ve CM-IoT ajanının değerlendirilmesi aşağıdaki bölümde sunulmaktadır.



Şekil 3.23: NS-IoT Sistem Mimarisi

NS-IoT sisteminde önerilen ajanlar Şekil 3.24'te gösterilmektedir. Bu ajanlar IIoT ağ etkinliklerini izler, anormallikleri tespit eder ve önceden tanımlanmış yanıtları bağımsız olarak yürütür. IIoT uç noktalarına, sunuculara ve ağ altyapısına dağıtılan bu ajanlar, tehditleri

tanımlamak, kontrol altına almak ve etkisiz hale getirmek için iş birliği yaparak genel güvenliği artırır. Bu sistem, otomasyon sayesinde müdahaleleri kolaylaştırır, insan müdahalesini azaltır ve kritik altyapı ve endüstriyel veri varlıkları üzerindeki saldırı etkilerini en aza indirir. NS-IoT sisteminde bu ajanlar arasındaki iş birliği ve koordinasyon çok önemlidir. Her temsilcinin benzersiz bir rol oynaması nedeniyle aralarında etkili iletişim ve senkronizasyon çok önemlidir. Güçlü iş birliği, tespit edilen tehditlere karşı hızlı ve etkili siber yanıtlar sağlayarak sistem güvenilirliğini artırır. Ayrıca bunu yaparak kaynak kullanımını optimize eder ve genel sistem performansını artırır. Bu nedenle, NS-IoT'nin ajanlar arasındaki iş birliği ve koordinasyonu, sistemin güvenilirliği ve etkinliği açısından kritik öneme sahiptir.



Şekil 3.24: NS-IoT Sistemi Ajan Yapılanması

NS- Iot ajanlarının açıklamaları aşağıda sunulmaktadır:

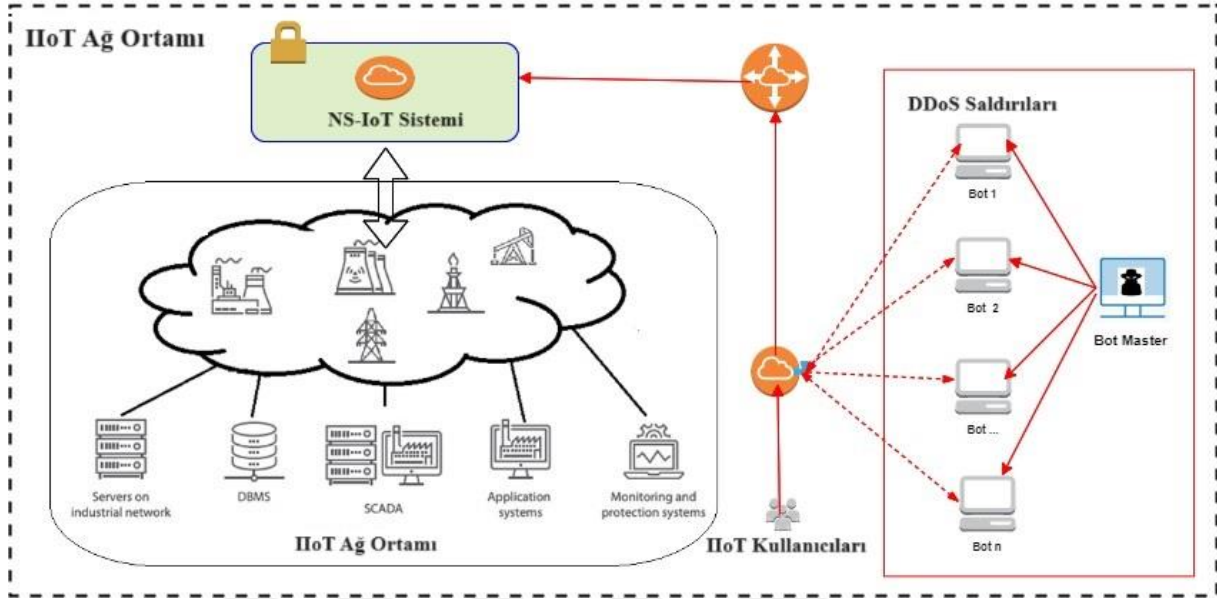
- **DQN-IoT Ajanı:** Bu ajan, öncelikle IIoT ortamındaki DDoS saldırılarını tanımlamaya odaklanarak izinsiz giriş tespitinde uzmanlaşmıştır. İlgili bilgileri paylaşmak için diğer ajanlarla iş birliği yaparak ağ trafiğini sürekli olarak izler. RL tabanlı bir anormallik algılama yaklaşımı kullanan DQN-IoT modeli, ağ trafiğini etkili bir şekilde karakterize eder. Tespit modülünün içinde yer alan bu modülün rolü, kötü amaçlı etkinliklerin hızlı ve gerçek zamanlı tanımlanmasında çok önemlidir. Temel işlevler arasında, özel olarak geliştirilen DQN modelinin kolaylaştırdığı ağ içindeki anormalliklerin tespit edilmesi ve tanımlanması yer alır. Bu ajan bir sonraki bölümde ayrıntılı olarak açıklanmaktadır.
- **CM-IoT Ajanı:** Bu ajan, NS-IoT sistemini koordine etmede, diğer ajanlar arasındaki iletişimi yönetmede ve izinsiz giriş tespiti sonrasında yanıt eylemlerini düzenlemede çok önemli bir rol oynar. DQN-IoT ajanlardan verileri alır, analiz eder ve ardından

saldırı azaltma prosedürlerini tetiklemek için ilgili bilgileri CR-IoT ajanına iletir. CR-IoT saldırıyı etkisiz hale getirdiğinde CM-IoT durum raporunu alır ve bunu DQN-IoT'ye iletir, böylece etkin saldırı tespiti ve hafifletilmesi için kapalı döngü iş birliği sağlanır. Ayrıca CM-IoT, DQN-IoT ve CR-IoT ajanları arasında kesintisiz iletişim sağlayan bir iletişim yöneticisi olarak görev yapar. Gerekli veri görevlerini işleyen, DDoS saldırı bilgilerini aktaran ve ağ trafiğini yöneten bir ajanı görevi görür. CM-IoT'nin birincil rolü, ajanlar arasında etkili iletişimi kolaylaştırarak, iş birliğini kolaylaştırarak ve iletişim karmaşıklığını azaltarak NS-IoT'de sistem performansını artırmaktır.

- **CR-IoT Ajanı:** Bu ajan, MAS-IoT sisteminin savunma modülünün koruyucusu olarak hareket ederek, tespit edilen anormalliklerin, özellikle DDoS saldırılarının, otonom olarak ve insan müdahalesine gerek kalmadan azaltılmasına öncelik verir. Temel amacı, savunma mekanizmalarını ve güvenlik politikalarını otonom bir şekilde uygulayarak IIoT ortamının güvenli çalışmasını sağlamaktır. Ajan, kötü amaçlı paketleri tespit ettikten sonra şüpheli IP'leri ve bağlantı noktalarını tanımlar, kötü amaçlı ağ paketlerini bırakır ve kara liste tablosunu saldırgan bilgileriyle günceller. Ek olarak, hataları ve arızaları ele alarak ağ yöneticisini zamanında yanıt vermesi için derhal uyarır. Bu ajanı, gerçek zamanlı bilgiye dayalı kararları etkinleştirerek, güvenlik tehditlerine karşı uyarlanabilir ve proaktif savunma sağlayarak IIoT ekosisteminin bütünlüğünü ve işlevselliğini korur.

Şekil 3.25, NS-IoT'nin IIoT ortamındaki bir DDoS saldırısına verdiği tepkiyi gösteren bir saldırı senaryosunu göstermektedir. Bu senaryoda IIoT ortamı, ortamı gözlemleyen ve DDoS saldırıları gerçekleştiren bir saldırganı içerir. Endüstriyel bir ortamda DDoS saldırıları, botnet tabanlı saldırıların bileşeni olarak kullanılabilir. Örneğin akıllı bir fabrikada, altyapıdaki birden fazla dağıtılmış düğüm, bir DDoS saldırısı başlatmak için iş birliği yapabilir. Bu tür bir saldırı, hedeflenen makineye veya üretim hattına önemli sayıda istek yağdırabilir, bu da sistem performansının düşmesine ve hatta ağ kesintilerine neden olabilir. Önerilen NS-IoT savunma sistemi, senaryo gereği IIoT altyapılarını DDoS tehdidine karşı etkili bir şekilde koruyor. Saldırı meydana geldiğinde NS-IoT'nin otonom ajanları sorunsuz bir şekilde iş birliği yapar. İlk olarak DQN-IoT ajanı, gelişmiş anormallik tabanlı algılamayı kullanarak saldırı modellerini hızlı bir şekilde tanımlar. Tespit üzerine CM-IoT ajanı yanıt sürecini yönetirken, CR-IoT ajanı DQN-IoT'den gelen öngörülerin rehberliğinde kötü amaçlı trafiği filtrelemek için otonom politikalar uygular. Bu koordineli çaba, NS-IoT'nin DDoS saldırılarını gerçek zamanlı olarak

tespit etme, analiz etme ve azaltma yeteneğini vurgulayarak IIoT ortamının bütünlüğünü ve işlevselliğini korur.



Şekil 3.25: NS-IoT Kullanımı Senaryosu

3.3.2. Modelin (DQN-IoT) Yapılandırılması

Takviyeli Öğrenme (RL), bir ajanın kümülatif ödülleri en üst düzeye çıkarmak için bir ortamla etkileşime girerek karar vermeyi öğrendiği bir makine öğrenimi paradigmasıdır. RL, yinelemeli deneme yanılma yoluyla en uygun politikayı belirlemeyi amaçlayan bir ajanın çevresi ile etkileşimi etrafında döner (Li, 2017). RL, temsilcilerin çevreyle etkileşimler yoluyla ödülleri optimize eden politikaları öğrenmesini sağlayarak uyarlanabilir özerkliği kolaylaştırır. RL, eksik bilgilerin ve belirsiz ortamların ele alınmasına ve sonuçta tam özerkliğe ulaşılmasına olanak tanır (Zhang ve diğ., 2021). RL, sürekli gelişen bir çevre ile deneme-yanılma etkileşimlerine girerek davranış kazanmaya çalışan bir ajanın karşılaştığı zorluğu özetlemektedir (Kaelbling ve diğ., 1996). RL, kümülatif ödülleri en üst düzeye çıkarmak için bir ortamda sıralı kararlar alma konusunda eğitim temsilcilerini içerir. RL algoritmaları, DDoS savunmasına uygulandığında, gelişmiş ve dinamik saldırı stratejilerine karşı uyarlanabilir bir şekilde öğrenebilir ve karşı önlemleri geliştirebilir. DDoS savunması için RL'nin benimsenmesi çeşitli avantajlar sunar. İlk olarak, RL tabanlı savunma mekanizmaları, değişen saldırı modellerini ve ağ koşullarını gerçek zamanlı olarak sürekli olarak öğrenip bunlara uyum sağlayabiliyor. Geleneksel kural tabanlı veya imza tabanlı savunma sistemlerinden farklı olarak RL ajanları, yeni ve görünmeyen DDoS saldırı türlerine karşı koymak için stratejilerini dinamik

olarak ayarlayabilir ve ağ savunmalarının dayanıklılığını artırabilir. Ek olarak, RL tabanlı yaklaşımlar, önceden tanımlanmış kurallara veya insan müdahalesine dayanmadan en uygun savunma stratejilerini özerk bir şekilde keşfetme potansiyeline sahiptir, böylece manuel konfigürasyona olan bağımlılığı ve uzman bilgisine duyulan ihtiyacı azaltır. Kot ve diğ. (2018), istenen YZ Siber Ajansı (AICA) ajanına ulaşmak için üç metodolojiyi özetlemektedir: uzmanlaşmış ajanlarından oluşan bir toplum kullanmak, çoklu ajanlı bir sistem kullanmak veya özerk bir işbirlikçi ajan geliştirmek. RL, çevre ile etkileşime giren tek bir ajanı içerir. Çoklu yinelemeler (dönemler) yoluyla, ajanı en iyi uzun vadeli getiriye elde etmek için eğitim alır (Denklem 3.7).

$$Q(s, a) = R(s, a) + \gamma \cdot \max_{a'} Q(s', a') \quad (\text{Denklem 3.7})$$

Q-öğrenme güncelleme kuralı, Bellman denklemine dayalı olarak Q değerlerini güncellemek için kullanılır. Bir durum-eylem çiftinin değeri ile bir ajanın s' durum içinde sharekete geçerek a ve ardından optimal politikayı takip ederek elde edebileceği beklenen kümülatif ödül arasındaki ilişkiyi ifade eder (Denklem 3.8).

$$Q(s, a) \leftarrow Q(s, a) + \alpha \cdot [R(s, a) + \gamma \cdot \max_{a'} Q(s', a') - Q(s, a)] \quad (\text{Denklem 3.8})$$

DQN-IoT ajanının geliştirilmesinde bir Q-öğrenme algoritması kullanılmıştır. Modelden bağımsız bir RL yaklaşımı olan O-Learning, temsilcilerin Markov ortamlarında doğrudan eylem-sonuç deneyimleri yoluyla en uygun kararları öğrenmelerine olanak tanımaktadır (Watkins ve Dayan, 1992). Başlangıçta sonsuz ufuk problemlerinde optimal karar stratejilerini tahmin etmeye yönelik bir algoritma olan Q-öğrenme, istatistik ve yapay zekada yaygın olarak uygulanan geniş bir RL yöntemleri kategorisine dönüştü (Clifton ve Laber, 2020). DQN ajanı, uçtan uca RL aracılığıyla doğrudan karmaşık duyuşal girdilerden etkili politikalar elde etmekte ve ödüller ve gözlemler gibi IIoT ortamı geri bildirimlerini aktif olarak özümsemektedir. Her DQN yinelemesinde durumları, eylemleri, ödülleri ve sonraki durumları içeren bir mini grup tekrarlanan bellekten örneklenir (Fan ve diğ., 2020).

DQN, takviyeli öğrenme alanında, özellikle de YZ ve MÖ alanında önemli bir ilerlemedir. DQN, özünde, durumları belirli bir ortamdaki eylemlerle eşleştiren, Q işlevi olarak bilinen optimum eylem-değer işlevine yaklaşmak için bir sinir ağı mimarisi kullanır. Bu, temsilcinin zaman içinde beklenen kümülatif ödülü en üst düzeye çıkaracak eylemleri seçerek bilinçli kararlar almasına olanak tanır. DQN'yi geleneksel Q-öğrenme algoritmalarından ayıran şey,

anlamli özellikleri otomatik olarak çıkarmak için sinir ağlarından yararlanarak görüntülerden ham piksel girişleri veya yüksek boyutlu sensör verileri gibi yüksek boyutlu durum alanlarını işleme yeteneğidir. Bu, DQN'nin, video oyunları ve robotik kontrol görevleri de dahil olmak üzere, geleneksel tablolu Q-öğrenme yöntemlerinin boyutluluğun laneti nedeniyle zorlandığı karmaşık ortamlarda üstün performans elde etmesine olanak tanır. DQN'de kararlı Q değerlerini hesaplamak için ikinci dereceden yöntemler kullanılır. Hedef Q değeri ile tahmin edilen Q değeri arasındaki fark olan zamansal fark (TD) hatasının hesaplanmasını ve bir dizi deneyim üzerinden ortalama kareli TD hatasının en aza indirilmesini içerir. Matematiksel olarak Denklem 3.9'daki gibi temsil edilebilir:

$$\delta = R + \gamma \cdot \max_{a'} Q(s', a'; \theta-) - Q(s, a; \theta) \quad (\text{Denklem 3.9})$$

Deneyim tekrarı, RL algoritmalarında eğitimin verimliliğini dengelemek ve geliştirmek için kullanılan bir tekniktir. Temsilcinin deneyimlerini (s, a, r, s') bir tekrar hafıza arabelleğinde depolamayı ve eğitim sırasında mini deneyim gruplarının rastgele örneklenmesini içerir. Bu, ardışık deneyimler arasındaki zamansal korelasyonların kırılmasına yardımcı olur ve ajanın daha çeşitli deneyimlerden öğrenmesine olanak tanır (Denklem 3.10).

$$D \leftarrow D \cup \{(s, a, r, s')\} \quad (\text{Denklem 3.10})$$

Öncelikli deneyim tekrarı, deneyimlerin önemlerine göre farklı olasılıklarla örneklandığı deneyim tekrar tekniğinin geliştirilmiş halidir. Sürpriz veya öğrenme potansiyeline işaret eden daha yüksek TD hatalarına sahip deneyimlere, tekrar arabelleğinde daha yüksek öncelik verilir. Bu öncelikli örnekleme şeması, ajanının bilgilendirici deneyimlerden öğrenmeye daha fazla odaklanmasına olanak tanıyarak daha hızlı ve daha verimli öğrenmeye yol açar (Denklem 3.11).

$$P(i) = |\delta i| + \epsilon \quad (\text{Denklem 3.11})$$

DQN, eğitimi stabilize etmek ve örnek verimliliğini artırmak için deneyim tekrarı ve hedef ağ tekniklerini içerir. Deneyim tekrarı, geçmiş deneyimlerin (yani durum-eylem-ödül-sonraki durum kayıtları) bir tekrar hafıza arabelleğinde saklanmasını ve eğitim sırasında deneyim gruplarının rastgele örneklenmesini içerir; bu, ardışık güncellemeler arasındaki zamansal korelasyonu keser ve veri verimliliğini artırır. Öte yandan hedef ağlar, Q değeri tahminleri için daha tutarlı bir hedef sağlayarak eğitimi stabilize etmek için sabit parametrelere sahip ayrı bir hedef Q ağına korunmasını içerir. DQN, hedef ve çevrimiçi ağları ayırarak, geleneksel Q-

öğrenme algoritmalarında bulunan fazla tahmin yanlışlığı sorununu hafifleterek daha istikrarlı ve güvenilir öğrenmeye yol açar. Ayrıca DQN, diğerlerinin yanı sıra Double DQN, Dueling DQN, Rainbow DQN ve dağıtımsal DQN dahil olmak üzere takviyeli öğrenmedeki çeşitli zorluklara ve uygulamalara hitap edecek şekilde genişletildi ve uyarlandı. Bu uzantılar, farklı etki alanları ve senaryolarda DQN'nin performansını, sağlamlığını ve örnek verimliliğini daha da artırmayı amaçlamaktadır. DQN, derin takviyeli öğrenme alanında temel bir çerçeveyi temsil eder ve geniş bir etki alanı yelpazesinde karmaşık gerçek dünya sorunlarıyla başa çıkabilen daha karmaşık ve yetenekli özerk ajanların geliştirilmesinin önünü açar.

DQN-IoT ajanının geliştirilme aşamalarına ilişkin algoritma **Tablo 3.2**'de sunulmuştur. Bu algoritma, DQN kullanarak eylem seçim ve ödül güncelleme süreçlerini optimize etmeyi amaçlayan bir yöntemdir. İlk olarak, bir dizi özellik vektörü tanımlanır ve epsilon parametresi, keşif ve sömürü dengesini sağlamak üzere ayarlanır. Eğitim süreci, belirli bir döngü sayısı boyunca tekrarlanır. Her döngüde, özellik vektörleri karıştırılarak eğitim verilerinin rastgeleliği artırılır. Her karıştırılmış özellik vektörü, modelin değerlendirmesi için durum vektörü olarak belirlenir ve bu durum vektörü DQN modeline iletilir. Model, durum vektörüne dayalı olarak çıkış değerleri üretir ve bu değerler arasından en yüksek olanı seçilerek en uygun eylem belirlenir. Seçilen eyleme ilişkin ödül alınır ve bu ödül, Bellman denklemi yardımıyla modelin ağırlıkları güncellenir. Öğrenme aşaması boyunca, Q-öğrenme güncellemeleri, çeşitli bir örnek havuzundan eşit ve rastgele alınan deneyim örneklerine uygulanır (Mnih ve diğ., 2015). Bu yöntem, modelin öğrenme sürecinde daha geniş bir deneyim çeşitliliği sunarak daha geliştirilebilir ve sağlam bir öğrenme sağlar. Hata vektörü, modelin öngörülleri ile gerçek sonuçlar arasındaki farkı ölçmek için hesaplanır ve bu hata değeri, modelin hafızasında saklanır. Elde edilen bilgiler, eğitim sürecinde tekrar kullanılmak üzere hafızaya kaydedilir ve bu sayede modelin öğrenme ve performans iyileştirmesi sağlanır. Bu süreç, algoritmanın daha etkin ve doğru eylem seçimleri yapmasını ve ödül güncellemelerini optimize etmesini sağlar.

Tablo 3.2: DQN Tabanlı Öğrenme Algoritması

Input: Bir dizi özellik vektörü ($X_1, X_2, \dots, X_N$)

Output: Temsilcinin eylemleri ve elde edilen ödüller

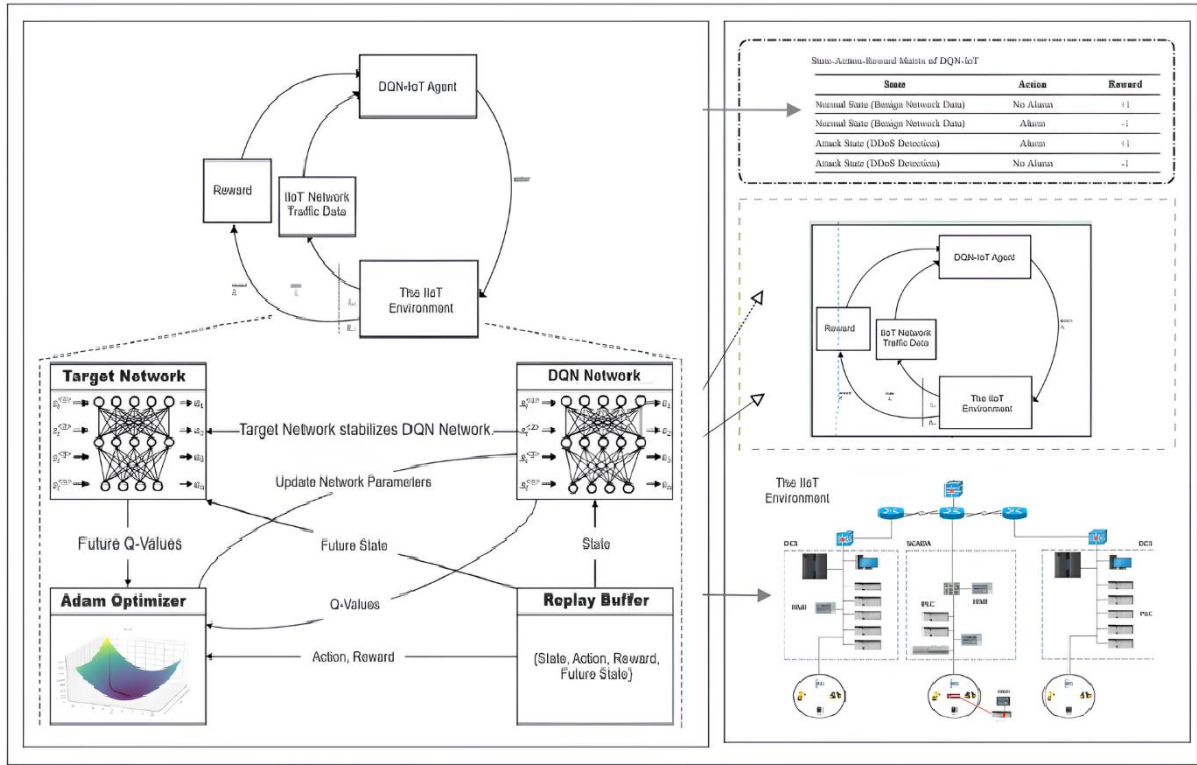
1: Özellik vektörleri kümesini başlatın: $X_1, X_2, \dots, X_N$

2: epsilon açgözlü yaklaşımı için epsilon (ϵ) değerini ayarlayın

3 : 1'den E'ye kadar olan aralıktaki i için

- 4: Durum sırasını ve özellik vektörleri kümesini karıştırın
 - 5: Karıştırılmış kümedeki her X_j özellik vektörü için şunu yapın
 - 6: durum_ vektörü $\leftarrow X_j$ 'yi ayarlayın
 - 7: Geçerli durum_ vektörünü aracının DQN'sine gönderin
 - 8 : çıkış değerlerini toplayın
 - 9 : Eylem = maksimum (olası eylemlerin Q değerleri)
 - 10: Seçilen eyleme göre ödül alın
 - 11: Bellman denklemi ile DQN ağırlıklarını güncelleyin
 - 12: $Q_{\text{new}}(\text{durum}, \text{eylem}, \text{ödül}) = \text{ödül} + \text{gamma}$
 - 13: error_vector = | değerini hesaplayın
 - 14: $\Sigma_i = \text{error_vector}[i] * \text{omega}$ değerini hesaplayın
 - 15 : current_state_loss_weight > tupleını M belleğinde saklayın
 - 16 : M hafızası örneklenen partideki eğitim sürecini tekrarla
 - 17 : son
-

Şekil 3.26, DQN-IoT ajanı mimarisinin yapısal bileşenlerini ve ağ içindeki bilgi akışını gösteren anlık görüntüsünü sunmaktadır. Bu çerçevede, veri kümeleri DQN'nin farklı durumlarını temsil eder; her sütun sistem durumunun çeşitli yönlerini veya özelliklerini temsil eder ve son sütun, model tahminlerine dayalı ödül vektörlerinin hesaplanması için bir etiket görevi görür. DQN-IoT ajanı, eğitim için DQN tabanlı bir yaklaşım kullanarak, gözlemlenen IIoT ortam durumlarına dayalı olarak gerçek zamanlı kararlar alma becerisini kazandı ve güvenlik tehditlerine karşı uyarlanabilir ve proaktif savunmayı mümkün kıldı. DQN'yi belirsizlik çerçeveleri bağlamında bir tahmin oyunu olarak ele almak, bir tahmin görevi olarak DDoS saldırı tespitini yapmak. Bu kurulumda model, belirsiz ortamlarda arzu edilen sonuçlara yol açan çeşitli eylemlerin olasılığını tahmin etmeyi öğrenir. Modeli belirsizlik altında doğru tahminler yapacak şekilde eğiten NS-IoT sistemi, gelişen tehditlere ustaca uyum sağlayabilir ve bilgiye dayalı gerçek zamanlı kararlar alabilir.



Şekil 3.26: DQN-IoT Ajanı Yapılanması

DDoS saldırı sınıflandırmasında DQN-IoT'de DQL aşağıdaki bileşenlerle uygulanır.

- **Durum:** Çevrenin durumu, s sınıflandırma için girdi görevi gören, ağ trafiği verilerinden çıkarılan özellikleri içerir.
- **Eylem:** Ajanın eylemi, a ağ trafiği verilerine bir etiket atamaya karşılık gelir; burada, a_0 'ın iyi huylu trafiği temsil ettiği ve 1 'in DDoS saldırı trafiğini temsil ettiği eylemler kümesinden seçilir.
- **Ödül:** DQN-IoT modeli, iyi huylu ağ verilerini doğru bir şekilde tespit etmek için +1 ödül ve hatalı bir şekilde alarm vermek için -1 ceza verirken, bir DDoS saldırısını başarılı bir şekilde tespit etmek için +1 ödüllendirir ve tespit edememek için -1 ceza verir. .
- **İndirim faktörü:** DQN-IoT modeli için indirim faktörü γ , her bir ağ trafiği örneğinin bağımsızlığı ve doğru tespitlerin +1 ve yanlış tespitlerin ödüllendirildiği ödül yapısı nedeniyle uzun vadeli değerlendirmeler yerine anlık ödüllere öncelik vermek üzere düşük ayarlanmıştır.

- **Keşif Oranı:** DQN-IoT'deki keşif oranı ϵ , keşif ve kullanımı dengeler; daha yüksek değerler, etkili stratejileri keşfetmeyi teşvik ederken, daha düşük değerler, bilinen stratejilerin kullanılmasını teşvik eder.
- **Bölüm:** Her bölüm, ajanın veri kümesindeki tüm ağ trafiği örneklerini sınıflandırdığında sona erer ve sınıflandırma politikasını iyileştirmek için eğitim verilerinin tamamından yinelemeli öğrenmeyi mümkün kılar.

Tablo 3.3'de DQN-IoT modelinin "Durum-Eylem-Ödül Matrisi" sunulmaktadır. Bu ödül mekanizmasına göre, ajanın bir saldırıyı doğru bir şekilde tespit etmesi veya başarılı bir şekilde hafifletmesi durumunda olumlu bir ödül alması gerekmektedir. Tersine, yanlış alarmların meydana geldiği veya saldırıların etkili bir şekilde önlenemediği durumlarda ajanın olumsuz bir ödül alması gerekir. Bu, temsilciyi doğru kararlar almaya ve zaman içinde performansını sürekli olarak geliştirmeye teşvik eder. Normal durumda, bir alarmın tespit edilmemesi bir ceza (+1), bir alarmın tespit edilmesi ise bir ödül (-1) sağlar. DDoS tespiti ile saldırı durumunda, alarmın tespit edilmesi ödül (+1) kazandırırken, alarmın tespit edilmemesi ceza (-1) gerektirir. Araştırmada kullanılan bu ödül modellemesi, IIoT güvenliğinde DDoS saldırı tespitini iyileştirmek için ödülleri uyarlıyor. Alan hedefleriyle uyumlu olarak hatalı pozitifleri/negatifleri cezalandırır. Bu yaklaşım, belirli hedeflere odaklanarak ve zaman içinde performansı artırarak ajanın öğrenmesini geliştirir.

Tablo 3.3: DQN-IoT Durum-Eylem-Ödül Matrisi

Durum	Aksiyon	Ödül
Normal Durum (İyi Niyetli Ağ Verileri)	Alarm Yok	+1
Normal Durum (İyi Niyetli Ağ Verileri)	Alarm	-1
Saldırı Durumu (DDoS Tespiti)	Alarm	+1
Saldırı Durumu (DDoS Tespiti)	Alarm Yok	-1

DQN-IoT ajanın geliştirilmesinde aşağıdaki teknikler kullanılmıştır:

- **(1) Tekrar Oynatma Mekanizması:** Tekrar oynatma mekanizması, DQN-IoT ajanın geçişleri (durum, eylem, ödül, sonraki durum) bir veri tabanında depolayarak deneyimden öğrenmesini sağlar. Sinir ağı eğitimi için yeniden oynatma arabelleği ve geçiş gruplarının örneklenmesi. Bu mekanizma, ardışık örnekler arasındaki

korelasyonları kırarak ve geçmiş deneyimlerin daha verimli kullanılmasını sağlayarak eğitimin dengelenmesine yardımcı olur.

- **(2) Hedef Ağ:** Hedef ağ, eğitim sırasında hedef Q değerlerini tahmin etmek için kullanılan ayrı bir sinir ağıdır. Hedef Q değerlerini mevcut Q değerleriyle güncellemeden önce belirli sayıda yineleme boyunca sabit tutarak eğitimin dengelenmesine yardımcı olur. Bu, hedef Q değerlerinin eğitim sırasında çok fazla dalgalanmasını önleyerek daha istikrarlı ve etkili öğrenmeye yol açar.
- **(3) Mini Toplu Eğitim:** Mini toplu eğitim, sinir ağını eğitmek için tekrar arabelleğinden küçük bir deneyim kümesinin rastgele örneklenmesini içerir. Bu yaklaşım, ağı her yinelemede öğrenebileceği çeşitli deneyimler sağlayarak öğrenme verimliliğini artırır ve belirli deneyimlerin gereğinden fazla uyma riskini azaltır.

(4) Markov Karar Süreci (MDP): Markov Karar Süreci (MDP), bir ajanın eylemler gerçekleştirerek ve ödüller alarak etkileşime girdiği belirsiz ortamlarda karar almayı modellemeye yönelik matematiksel bir çerçevedir. DQN-IoT ajanı çerçevesi içinde MDP, durumları, eylemleri ve ödülleri tanımlar ve ajanın karar verme sürecini zaman içinde kümülatif ödülleri en üst düzeye çıkaracak şekilde yönlendirir. Genel olarak bu mekanizmalar, DQN-IoT ajanının IIoT ortamını etkili bir şekilde öğrenmesini ve ona uyum sağlamasını sağlamak için birlikte çalışarak tehditleri bağımsız olarak tespit etmek ve azaltmak için bilinçli kararlar alır

4. BULGULAR

Bu bölümde, 3. Bölümde önerilen ajan tabanlı saldırı tespit modellerinin eğitim ve test süreleri, kullanılan metrikler (başarı oranı, kayıp oranı), eğitim ve test zamanı ile karmaşıklık değerleri analiz edilmiştir. Daha sonra elde edilen bu sonuçlar, birbirleri ve literatürdeki çalışmalarla karşılaştırılarak en etkili ve yenilikçi saldırı tespit modeli önerilmiştir.

4.1. Derin Öğrenme (DÖ) Tabanlı Modelin Analizi

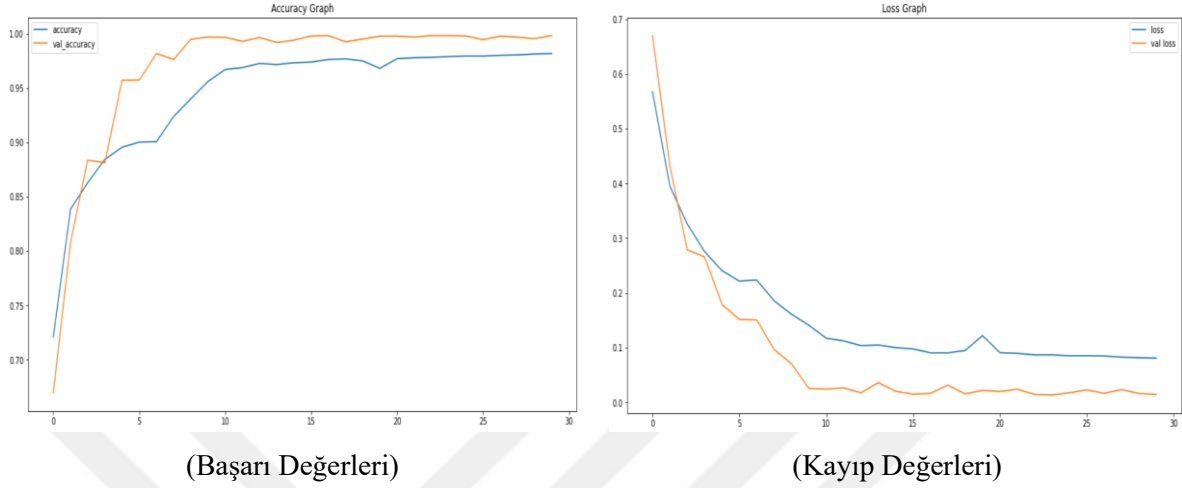
Çalışma kapsamında LSTM-CLOUD modeli kapsamında gerçekleştirilen deneyler sonucunda en doğru tahmin işleminin Deney-2'de gerçekleştirildiği belirlenmiştir. 15 dönemlik iki katmanlı bir yapıda gerçekleştirilen bu deneyde DDoS saldırıları %99,83 gibi yüksek bir başarı oranıyla tahmin edilmiş ve sınıflandırılmıştır. Bu deneyin aynı zamanda RMSE değerinin en küçük olarak elde edildiği deney olduğu görülmektedir. Bu da geliştirilen modelin DDoS saldırılarını doğru tespit ettiğini gösteriyor. Esasen bu yüksek başarı tahmin oranının arkasındaki temel faktör, dizi tahmin problemlerini çözerken doğru parametrelerle eğitilmiş LSTM derin öğrenme algoritmasının gücüdür. Bu, LSTM'nin geçmiş bilgileri saklama ve geri çağırma gücüyle, ağ trafiği akışının önceki normal veya kötü niyetli davranışının etkili bir şekilde kullanılabileceğini göstermektedir. Çalışmada yapılan deneyler ve bu deneylerin sonuçları **Tablo 4.1**'de gösterilmektedir.

Tablo 4.1: Model Üzerinde Yapılan Deneylerin Analizi

Deney Etiketi	Gizli Katmanlar	Çağlar	RMSE	Kesinlik (%)	Kayıp (%)
1.	2	5	0,194	%93.85	%6.15
2.	2	15	0,039	%99.83	%0.17
3.	2	60	0,194	%92.79	%7.21
4.	4	5	0,302	%89.61	%10.39
5.	4	15	0,102	%98.85	%1.15
6.	4	60	0,261	%90.28	%9.72
7.	6	5	0,293	%89.10	%10.9
8.	6	15	0,190	%95.56	%4.44
9.	6	60	0,203	%91.72	%8.28

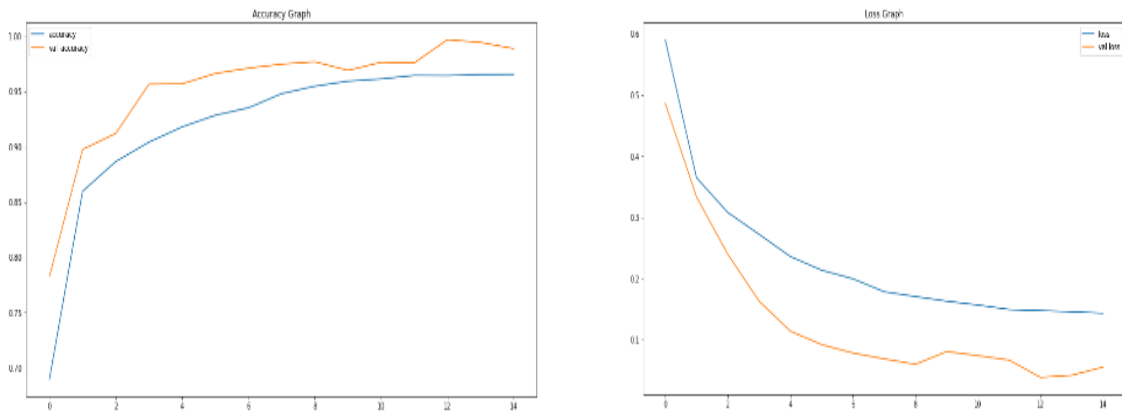
Şekil 4.1'de sunulan grafiklerde Deney-2'nin doğruluk ve kayıp oranlarını grafiksel olarak yer almaktadır. İyi tahmin modellerinde eğitim çizgisinin dinamikleri doğrulama çizgisi yönünde

ve ona mümkün olduğu kadar yakın olmalıdır. Söz konusu grafikte aşırı öğrenmenin (ağın ezberleme durumu) olmadığı görülebilmektedir.

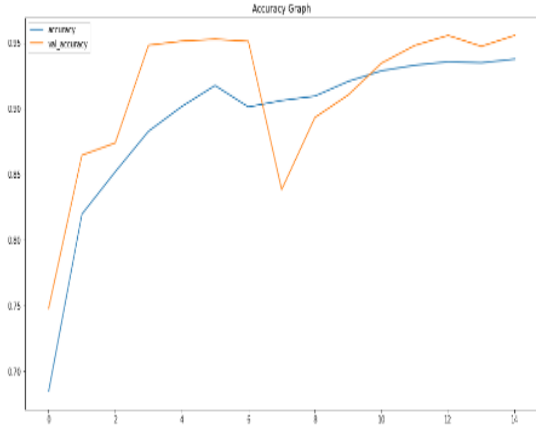


Şekil 4.1: Deney-2'ye ait Başarı ve Kayıp Dağılımları

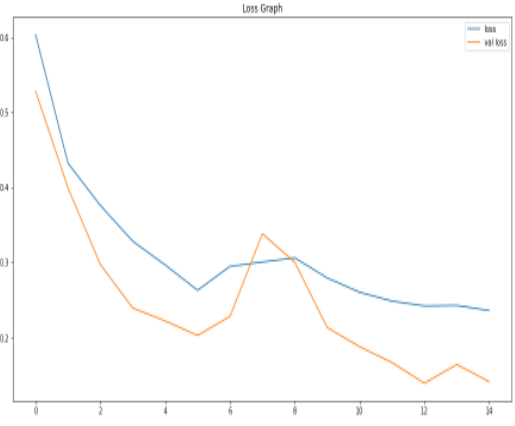
Deney-2 dışındaki deney sonuçları, RMSE oranının artık düşmediğini gösterdi. Dönem sayısı artırılarak gerçekleştirilen Deney-3'te katman sayısı artırılmamasına rağmen aynı yüksek başarı elde edilememiştir. Deney-4 ve diğer deneylerde de aynı yüksek başarı elde edilememiştir. Katman ve çağ sayısını arttırdığımızda modelin doğruluğunun da buna paralel olarak artmadığı görüldü. Bu sonuç aynı zamanda modelin katman sayısını değiştirmenin yanı sıra daha birçok faktörün de daha iyi bir başarı oranına ulaşmada etkili olduğunu göstermektedir. Tüm deneylerin sonuçları değerlendirildiğinde doğruluk oranları ile RMSE değerleri arasında ters bir ilişkinin olduğu görülmüştür. Yani deneylerde başarı oranı arttıkça RMSE değeri düşmektedir. Deney-2 dışındaki deneyler arasında en yüksek başarı oranına sahip olan Deney-4 ve Deney-5'te elde edilen başarı ve kayıp oranları **Şekil 4.2**'de gösterilmektedir.



(Deney-4 Başarı Değerleri)



(Deney-4 Kayıp Değerleri)



(Deney-8 Başarı Değerleri)

Şekil 4.2: Deney-4 ve Deney-8 Başarı ve Kayıp Dağılımları

(Deney-8 Kayıp Değerleri)

Elde edilen doğruluk oranları, gelecekteki verileri tahmin etmek için geçmiş ağ verilerini temel alan LSTM modelimizin tahmininde ne kadar doğru olduğunu göstermektedir. Yapılan deneylerde elde edilen başarı oranlarının birbirinden farklılık gösterdiği, bazen azaldığı bazen de arttığı aşikardır. Deneylerin doğruluk oranlarının karşılaştırılması Şekil 4.3'de gösterilmektedir.



Şekil 4.3: Doğruluk Oranlarının Karşılaştırılması

Bu deneylere ilişkin çalışmada, genel bir bulut ağında DDoS saldırılarının tespiti ve önlenmesine yönelik bir sistem (LSTM-CLOUD) geliştirilmiştir. Sistem tespit ve savunma olmak üzere iki modülden oluşmaktadır. Sistemin ilk modülünün işlevi, bu çalışmada

geliştirilen LSTM derin öğrenme modeli ile DDoS saldırılarının oluşumunu tespit etmek olarak belirlenirken, sistemin ikinci modülünün işlevi savunma mekanizmasını harekete geçirmek olarak belirlendi. Sistemin tespit modülünde kullanılmak üzere LSTM derin öğrenme modeli geliştirilmiştir. Deneyler, LSTM'ye dayalı tespit modelinin doğruluğunun %99,83 olduğunu göstermektedir. Geliştirilen LSTM modelini doğrulamak için çalışmada CICDDoS 2019 veri seti kullanılmıştır. Bu doğruluk oranının arkasında LSTM'nin dizi tahmin problemlerindeki gücü yatmaktadır. Karşılaştırma sonuçları, LSTM modelimizin aynı ve farklı veri setleri üzerinde farklı derin öğrenme algoritmaları ile yapılan önceki çalışmalardaki kadar iyi bir performansa sahip olduğunu gösterdi. Elde edilen sonuçlar, bu çalışmada geliştirilen LSTM modelinin saldırıların oluşumunu tespit etmedeki etkinliğini göstermektedir.

LSTM Ağı ile gerçekleştirilen LSTM-CLOUD modelinin geliştirilmesine ilişkin çalışmanın literatüre katkıları şu şekilde ifade edilebilir:

- Genel bir bulut ağ ortamının ağ trafiği karakterizasyonu için LSTM tabanlı bir sistem (LSTM-CLOUD) tasarlanmıştır. Bu sistemin modülleri ve görevleri tanımlanmıştır. LSTM-CLOUD sisteminin tasarımı imza tabanlı saldırı tespit yaklaşımına dayanmaktadır.
- LSTM-CLOUD sisteminde kullanılmak üzere çalışmada gerçekleştirilen deneysel çalışmalarda geçmiş ağ trafik değerlerine bakarak gelecekteki saldırıları tahmin edebilen %99,83 başarı oranına sahip bir LSTM modeli geliştirilmiştir. Bu da derin öğrenmeyi kullanarak bilinmeyen siber saldırıların sınıflandırılması ve tespit edilmesi konusunda avantaj sağlamaktadır.
- Bu çalışma kapsamında, LSTM-CLOUD modelini değerlendirmek için gerçek dünyadaki DDoS saldırı örneklerini içeren veri setini kullanarak farklı deneyler yapılmıştır. Deneylerde elde edilen sonuçlar, LSTM-CLOUD modelinin aynı ve farklı veri kümeleri üzerinde farklı DÖ algoritmalarıyla yürütülen literatürde mevcut önceki çalışmalardaki kadar iyi bir performans sağladığını başarıyla doğrulamıştır.

4.2. Ajan Tabanlı Derin Öğrenme (DÖ) Tabanlı Modelin Analizi

4.2.1. Bulut Altyapısı için Önerilen Modelin Analizi

Tablo 4.2'de LSTM-IoT modelinin farklı konfigürasyonları için performans ölçümlerini gösteren deneylerin sonuçlarını gösterilmiştir. Tablonun her satırı, gizli katman sayısı, dönem

sayısı ve eğitim süresi gibi faktörlere göre farklılık gösteren belirli bir konfigürasyona karşılık gelmektedir. Tablo, her konfigürasyon için RMSE (Ortalama Kare Hatanın Kökü), doğruluk oranı ve kayıp oranını sağlamaktadır. Deney sonuçları RMSE, doğruluk (%) ve kayıp değerleri dikkate alınarak analiz edilmiştir.

Tablo 4.2: LSTM-IoT Deneysel Sonuçları

Exp.No.	Num of Hidden Layers	Epoch	Training Time	Testing Time	RMSE	Accuracy (%)	Loss (%)
1.	2	5	6.09	3.86	0.182	96.60%	4.40%
2.	2	15	12.10	3.88	0.039	94.62%	5.38%
3.	2	60	28.10	3.17	0.228	94.48%	5.52%
4.	4	5	10.62	3.99	0.230	94.41%	5.09%
5.	4	15	15.4	3.24	0.228	94.55%	5.45%
6.	4	60	29.2	3.36	0.452	99.48%	0.52%
7.	6	5	14.5	3.37	0.228	94.52%	5.48%
8.	6	15	24.8	3.22	0.225	94.65%	5.35%
9.	6	60	29.8	3.25	0.227	94.58%	5.42%

Bu deneysel sonuçlara dayanarak, deneylerin ayrıntılı bir karşılaştırması aşağıda verilmiştir:

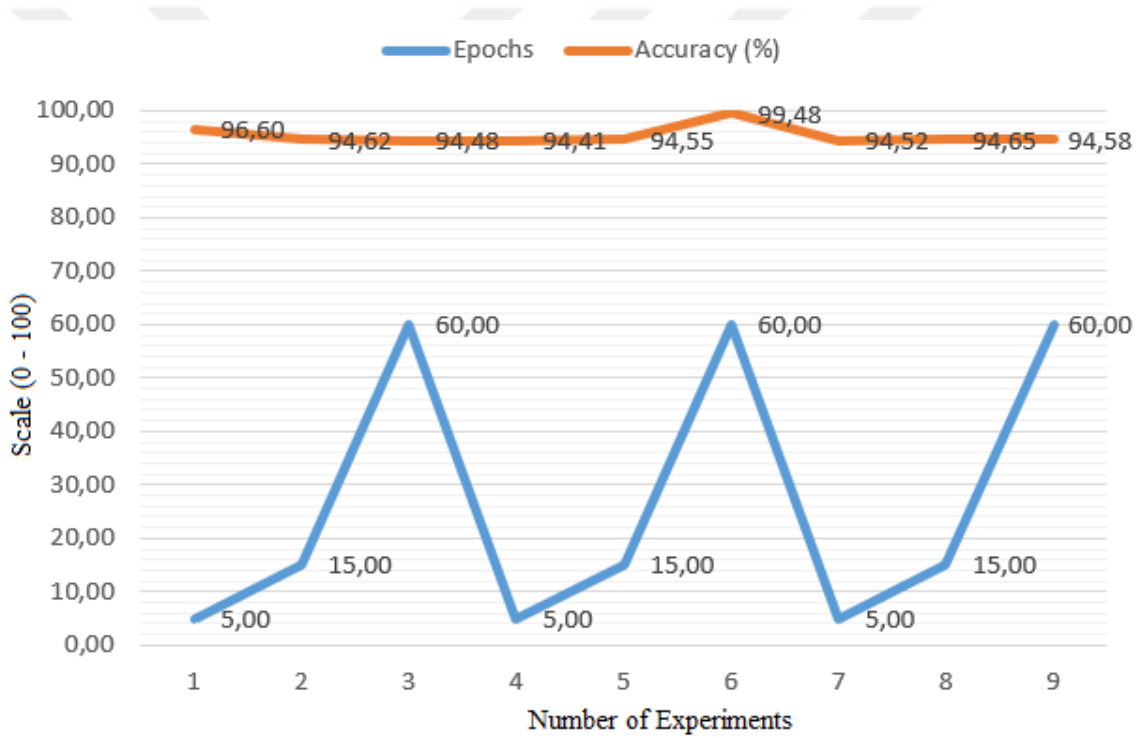
- Deney 1'de iki gizli katman, beş eğitim dönemi ve 0,182 RMSE ve %4,40 kayıp değeriyle %96,60 doğruluk oranı elde edilmiştir.
- Deney 2'de de iki gizli katman vardı ancak 15 eğitim dönemi vardı. Deney 1'e kıyasla daha düşük RMSE (%0,039) ve doğruluk (%94,62) değerlerine ve daha yüksek kayıp değeri (%5,38) elde edilmiştir.
- Deney 3, Deney 1 ve 2 ile aynı gizli katmanlara sahip olarak gerçekleştirilmiş olup 60 eğitim dönemine sahipti. Deney 3'ün, Deney 1'e göre daha yüksek RMSE değerine (0,228) ve biraz daha düşük doğruluk oranına (%94,48) sahipken, Deney 1 ve 2'ye göre daha yüksek kayıp değerine (%5,52) sahip olduğu görülmektedir.
- Deney 4'te dört gizli katman ve beş eğitim dönemi (epoch) kullanılmış ve performans ölçümleri Deney 3'e benzer (RMSE = 0,230, doğruluk = %94,41 ve kayıp = %5,09) olarak elde edilmiştir.
- Deney 5, Deney 4 ile aynı sayıda gizli katmana sahip olarak gerçekleştirilmiştir. Performans ölçümleri Deney 3 ve 4'e benzer (RMSE = 0,228, doğruluk = %94,55 ve kayıp = %5,45) olarak elde edilmiştir.

- Deney 6'da 60 eğitim dönemi içeren dört gizli katman kullanılmıştır. Diğer tüm deneylere göre en yüksek doğruluk (%99,48) ve RMSE (0,452) değerlerine ve en düşük kayıp değerine (%0,52) bu deneyde ulaşılmıştır.
- Deney 7'de altı gizli katman ve beş eğitim dönemi kullanılmıştır. Performans ölçümleri Deney 3, 4 ve 5'e benzer (RMSE = 0,228, doğruluk = %94,52 ve kayıp = %5,48) olarak elde edilmiştir.
- Deney 8, Deney 7 ile aynı sayıda gizli katmana sahip olarak gerçekleştirilmiş olup 15 eğitim dönemine sahipti. Bu deneyde performans ölçümleri Deney 1 ve 2'ye benzer (RMSE = 0,225, doğruluk = %94,65 ve kayıp = %5,35) değerlerde elde edilmiştir.
- Deney 9'da 60 eğitim dönemi içeren altı gizli katman kullanılmıştır. Performans ölçümleri Deney 3, 4 ve 5'e benzer (RMSE = 0,227, doğruluk = %94,58 ve kayıp = %5,42) olarak elde edilmiştir.

Çalışma kapsamında gerçekleştirilen deneylerde elde edilen bu sonuçlar, farklı deneyler arasında çeşitli performans farklılıklarını ortaya koymaktadır. Öncelikle daha fazla eğitim periyoduna sahip deneylerin genel olarak daha düşük RMSE değerlerine sahip olduğu söylenebilir. Ancak yüksek dönem sayıları bazen daha yüksek RMSE'ye ve daha düşük doğruluk oranlarına yol açtığı görülmektedir. Daha fazla gizli katmanlarla yapılan deneyler genellikle daha yüksek doğruluk elde ettiğinden, gizli katmanların sayısının da performansı etkilediği görülmektedir. Ancak söz konusu deneylerde gizli katman sayısının artırılması bazen daha yüksek RMSE ve kayıp değerleri ile sonuçlanmıştır. Son olarak, bazı deneylerin benzer performans ölçümleri sağladığı ancak farklı model konfigürasyonlarına sahip oldukları görülmektedir. Genel olarak deneyler, gizli katman sayısını arttırmanın veya eğitim sürelerini uzatmanın performansı sürekli olarak iyileştirmediğini ortaya çıkarmaktadır. En yüksek doğruluk ve en düşük kayıp değerlerine sahip olan Deney 6, 60 eğitim dönemi ile dört gizli katman kullanmış ancak tüm deneyler arasında en yüksek RMSE'ye sahip olmuştur. Öte yandan Deney 1 ve 2, en düşük RMSE'ye ve daha az eğitim dönemine sahip olmuştur. Optimum performansa ulaşmak için gizli katmanların sayısı ile eğitim dönemleri arasında doğru dengeyi bulmak çok önemli görünmektedir.

Deneyisel çalışmalarımızda elde edilen sonuçlar, LSTM-IoT modelinin saldırıları yüksek başarı oranıyla ve doğru bir şekilde sınıflandırdığını göstermektedir. LSTM-IoT modeli, dizi tahmin

görevleri için uygun parametrelerle eğitilmiş ve yüksek bir başarı oranı elde edilmiştir. Ayrıca deneysel sonuçlar, LSTM-IoT modelinin, zaman serisi verilerinden geçmiş bilgileri etkili bir şekilde depolayıp alarak IoT ağ trafiği akışlarını verimli bir şekilde sınıflandırma kapasitesini ortaya koymaktadır. Gerçek zamanlı siber saldırı tespitinin hassasiyetinin bilinmesi büyük önem taşımaktadır. Bu araştırma kapsamında gerçek siber saldırı verilerini içeren güncel bir veri setini kullanarak optimum parametrelere sahip bir LSTM modelin geliştirilmesi amaçlanmıştır. **Şekil 4.4**'de, her deney için değişen periyotlardaki doğruluk ve kayıp değerleri gösterilmektedir. Bu değerler incelendiğinde doğruluk oranlarının hemen hemen aynı olduğu görülmektedir. Ancak en yüksek doğruluk oranı altıncı deneyde sağlanırken, en yakın deneysel sonuçlar %96,60 başarı oranıyla birinci deneyde elde edilmiştir.

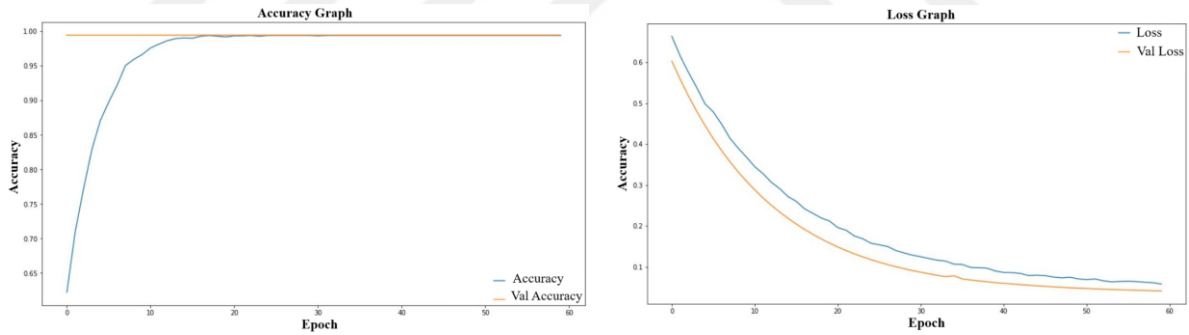


Şekil 4.4: Doğruluk ve Kayıp Değerleri

Deneysel çalışmalarda LSTM-IoT modelinden elde edilen doğruluk ve kayıp oranlarını görselleştirilmiştir. Doğruluk grafikleri modelin doğruluğunu mavi çizgiyle gösterirken kırmızı çizgi “val_accuracy” metriğini temsil etmektedir. Benzer şekilde kayıp grafiklerinde de mavi çizgi kayıp oranını, kırmızı çizgi ise “val_loss” metriğine karşılık gelmektedir. LSTM-IoT modelinin eğitim sürecinde temel değerlendirme metrikleri olarak “val_accuracy” ve “val_loss” değerleri kullanılmıştır. Bu ölçümler, modelin performansının hem eğitim hem de doğrulama veri kümeleri üzerinde değerlendirilmesi söz konusu olduğunda son derece

önemlidir. Spesifik olarak “val_accuracy”, modelin daha önce eğitim sırasında karşılaşmadığı verileri içeren doğrulama veri setinde elde edilen doğruluk oranını ölçer. Bu, modelin yeni ve görülmemiş veriler üzerinde genelleme yeteneğini değerlendirmemize olanak tanır. Tersine, “val_loss” doğrulama veri kümesinde hesaplanan kayıp değerini temsil eder ve modelin tahmin edilen değerleri ile doğrulama kümesindeki gerçek değerler arasındaki tutarsızlığı ölçer. Düşük bir “val_loss” değeri, doğrulama veri kümesinde güçlü performans anlamına gelir. Bu da modelin tahminlerinin gerçek değerlerle yakından eşleştiğini göstermektedir. Bu kritik ölçümler, LSTM-IoT modelinin genel performansını değerlendirmeye ve potansiyel aşırı uyum sorunlarını tespit etmeye hizmet etmektedir. Modelin etkili bir şekilde eğitildiğinden ve güçlü genelleme yetenekleri gösterdiğinden emin olmak için eğitim süreci boyunca “val_accuracy” ve “val_loss” değerleri kullanılmıştır.

Şekil 4.5'de, en yüksek başarı oranına ulaşan Deney 6'nın doğruluk ve kayıp eğilimlerini gösterilmektedir. Bu grafikte eğitim eğrisi, doğrulama eğrisinin yörüngesini yakından takip ederek trendlerine yakın bir şekilde hizalanmıştır. Grafik, eğitim kaybının doğrulama kaybıyla tutarlı bir şekilde azaldığını göstermektedir.

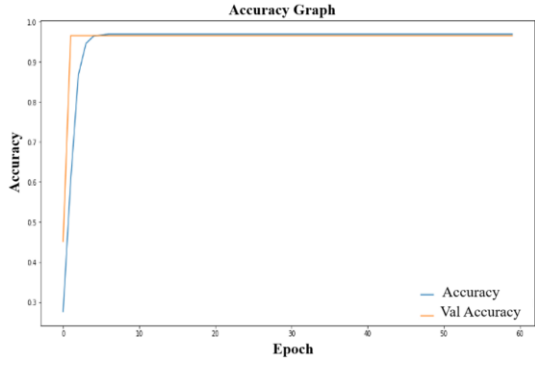


(a) Deney-6 Doğruluk Oranı

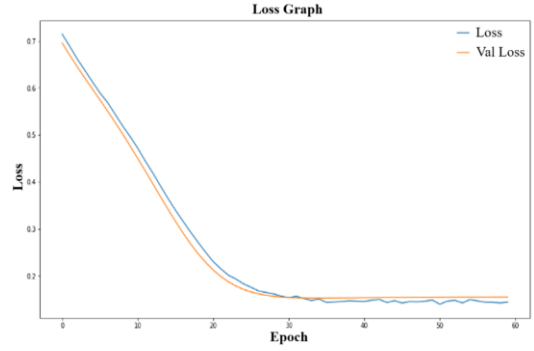
(b) Deney-6 Kayıp Oranı

Şekil 4.5: Deney-6 Başarı ve Kayıp Oranları

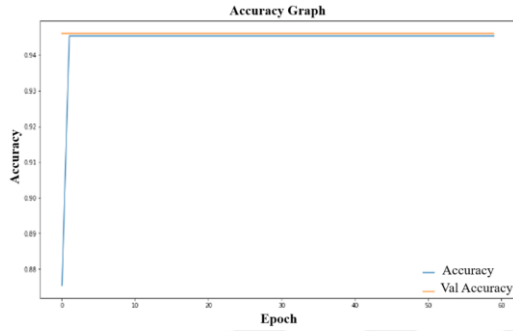
Çalışmada elde edilen bu deneysel bulgular, doğruluk oranları ile RMSE değerleri arasında ters bir ilişki olduğunu ortaya koymaktadır. Yapılan deneylerde başarı oranı arttıkça RMSE değerinin sürekli düştüğü görülmektedir. **Şekil 4.6**'da sunulan grafik, en yüksek başarı oranına sahip denemelerin başarı ve kayıp eğilimlerini göstermektedir. Bu başarı oranları genel olarak birbirine yakın görünmektedir.



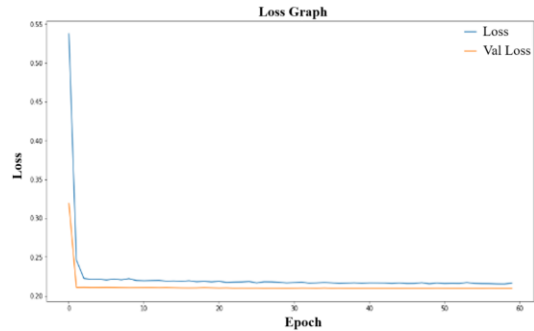
(a) Deney-1 Doğruluk Oranı



(b) Dney-1 Kayıp Oranı



(c) Deney-2 Doğruluk Oranı



(ç) Dney-2 Kayıp Oranı

Şekil 4.6: Başarı Oranı En Yüksek Deneyler

DeneySEL sonuçlar, LSTM-IoT modelinin DDoS saldırılarını tespit etmede oldukça yüksek bir doğruluk oranına ulaştığını göstermektedir. LSTM modelinin zaman serisi verilerini analiz etme yeteneği, ağ trafiğinin karmaşık dinamiklerini etkili bir şekilde öğrenmesini sağlamaktadır. LSTM-IoT modeli ağ trafiğindeki anormallikleri geçmiş bilgilere dayanarak tanımlayabilmektedir. DDoS saldırılarına karşı gerçek zamanlı savunma için yüksek doğruluk oranı çok önemlidir. LSTM-IoT modelinin DDoS saldırılarını tespit etmedeki yüksek hassasiyeti, DDoS saldırı tespiti alanında önemli bir avantajdır.

Çalışmada LSTM-IoT'nin eğitim ve test süreleri de hesaplanmıştır. Bir LSTM modeli geliştirebilmek için kısa eğitim ve test süreleri çeşitli nedenlerden dolayı önemlidir. Her şeyden önce, daha hızlı deneme ve yinleme olanağı sağlar. Daha kısa eğitim ve test süreleri sayesinde araştırmacılar modellerini daha sık eğitip test edebilir ve yaklaşımlarını daha hızlı yineleyebilir. Bu, model konfigürasyonları ve hiper parametrelerden oluşan daha geniş bir alanı keşfedebilecekleri ve daha iyi performansa yol açabilecekleri anlamına gelir. İkinci olarak, daha kısa eğitim ve test süreleri hesaplama maliyetlerinin azaltılmasına yardımcı olabilir. LSTM

modelleri, özellikle büyük veri kümeleriyle çalışırken hesaplama açısından pahalı olabilir. LSTM modellerinin gerçek zamanlı sistemler gibi pratik uygulamalarında kısa eğitim ve test sürelerine sahip olmak esastır. Eğitim ve test süreleri kısa olan bir model, gelen veri hızına ayak uydurabilmesini ve zamanında doğru tahminler yapabilmesini sağlayabilir. Deneylerimizin test ve eğitim sürelerine bakıldığında Deney 1'in en kısa eğitim süresine (6.09) sahip olduğu, Deney 6'nın ise en uzun eğitim süresine (29.2) sahip olduğu söylenebilir. Test süresi açısından ise Deney 1 en kısa test süresine (3,86) sahip olurken, Deney 2 en uzun test süresine (3,88) sahiptir. Genel olarak, Deney 1, 4 ve 7'nin diğer deneylerle karşılaştırıldığında nispeten kısa eğitim ve test sürelerine sahip olduğu sonucuna varılabilir. Buna karşılık, Deney 6 en uzun eğitim süresine sahiptir ve Deney 2 en uzun test süresine sahiptir. Bu değerler benzer çalışmalarda elde edilen değerlerden daha düşüktür. Dolayısıyla bunun iyi bir sonuç olduğu söylenebilir. Ancak, tam analiz için her deneyin RMSE, doğruluk ve kayıp gibi performans ölçümlerinin de dikkate alınması gerektiğini unutmamak önemlidir.

Son olarak çalışmada LSTM-IoT'nin zaman karmaşıklığı da hesaplanmıştır. LSTM zaman karmaşıklığı önemlidir, çünkü eğitim ve çıkarım sırasında modelin hızını ve verimliliğini etkiler. Daha fazla zaman karmaşıklığına sahip bir LSTM modeli, daha fazla hesaplama kaynağı ve verileri işlemek için daha fazla zaman gerektirir. Bu durum gerçek zamanlı uygulamalarda büyük bir dezavantaj olabilir. Ek olarak, daha fazla zaman karmaşıklığı, modelin ölçeklenebilirliğini sınırlayarak daha büyük veya daha karmaşık veri kümelerinin işlenmesini zorlaştırabilir. LSTM modelinin, genel performansının yüksek olması ve gerçek dünya uygulamalarına daha uygun olması genellikle gerekli olduğundan, düşük düzeyde zaman karmaşıklığına sahip olması gerekir. Bunun için önerilen çözümün LSTM katmanının zaman karmaşıklığının hesaplanması gerekmektedir. LSTM, zaman adımı başına hesaplama karmaşıklığı ve ağırlığın $O(1)$ olmasıyla hem uzayda hem de zamanda yerellik sergiler. Bir LSTM'nin zaman adımı başına genel karmaşıklığı, $O(w)$ olarak gösterilen ağırlıkla doğrudan orantılıdır. Sonuç olarak LSTM-IoT'nin asimptotik gösterimde O zaman karmaşıklığına sahip olduğunu söyleyebiliriz. Bu zaman karmaşıklığı, LSTM tabanlı LSTM-IoT'yi çalıştırmak için gereken süre ölçümünü kapsamaktadır.

Deneysel çalışma sonuçları, LSTM-IoT modelinin doğruluk ve özellik seçim boyutu açısından çok amaçlı parçacık sürüsü optimizasyonu (MOPSO) ve MOPSO-Lévy yöntemlerine göre açık ve önemli bir avantaj sergilediğini ve %98 başarı oranında etkileyici bir sınıflandırma doğruluğuna ulaştığını göstermektedir.

Bu çalışmada, IoT ağı ortamında DDoS saldırılarını tespit etmek ve azaltmak için çok ajan tabanlı bir sistem (MAS-IoT) önerilmiştir. MAS-IoT, dört ajandan (koordinasyon, algılayıcı, tespit ve savunma ajanı) oluşan bir güvenlik düğümüne sahip çok ajanlı bir çerçeve sağlar. Tüm sistem faaliyetleri sırasında algılayıcı, tespit ve savunma ajanlarının faaliyetleri koordinasyon ajanı tarafından koordine edilir. Tespit etme ajanı, DDoS tespiti için bu çalışmada geliştirilen LSTM tabanlı modeli kullanır. Her ajan, izlenen IoT ortamını DDoS saldırılarına karşı korumak için iş birliği içinde benzersiz işlevler gerçekleştirir. Bu çalışmada nicel araştırma yöntemleri kullanılmıştır. Bunu başarmak için siber savunmada ajan temelli yaklaşımlara odaklanan bir literatür taraması yapılmıştır. Daha sonra imza tabanlı yaklaşımla geliştirilen LSTM-IoT modelini geliştirmek için farklı deneyler yapıldı ve farklı başarı oranları gözlemlenmiştir.

Bu araştırmanın başlıca katkıları şu şekilde ifade edilebilir:

- IoT ağlarını hedef alan DDoS saldırılarına karşı çok ajanlı bir siber savunma sistemi tasarımının (MAS-IoT) önerilmesi.
- Russell ve Norvig'in (2009) yaygın olarak kabul edilen ajan modellerine dayanarak çeşitli ajan türlerinin tanımlanması.
- Sistemde tanımlanan algılama ajanları tarafından kullanılmak üzere güncel IoT veri kümesi (CIC-Dos-2022) ile imza tabanlı bir LSTM sınıflandırma modelinin (LSTM-IoT) geliştirilmesi.
- Tüm ajanlar arasında şifreli iletişim kullanılması.

MAS-IoT sistemi, IoT ağları için özel olarak tasarlanmış işbirlikçi, ajan tabanlı bir yaklaşım sunarak geleneksel siber savunma araçlarının sınırlamalarını giderir. Büyük ölçüde insan müdahalesine ve manuel yapılandırmaya dayanan geleneksel araçların aksine, sistemimiz özerkliğe ve karar verme yeteneklerine sahip siber savunma amaçlı akıllı ajanlardan yararlanır. Bu ajanlar sürekli çalışır ve IoT ağlarında DDoS saldırılarına karşı gerçek zamanlı izleme, tespit ve savunma sağlar. Saldırı tespiti için doğru bir LSTM tabanlı modeli entegre ederek ve şifreli iletişim kanallarını kullanarak sistemimiz, ajanlar arasında hassas bilgi alışverişini korurken hızlı ve doğru saldırı tespitini sağlar. MAS-IoT'nin işbirlikçi ajan yapısı, bireysel tabanlı savunma mekanizmalarının sınırlamalarının üstesinden gelerek olası saldırıların etkili toplu tespitine ve önlenmesine olanak tanır. Bu özellikleriyle MAS-IoT sistemi, IoT ağlarının hızla

gelişen ve karmaşık DDoS saldırılarına karşı korunmasında gelişmiş verimlilik, ölçeklenebilirlik ve uyarlanabilirlik sunan umut verici bir yaklaşım olarak öne çıktığı görülmektedir.

4.2.2. IoT Ağları için Önerilen Modelin Analizi

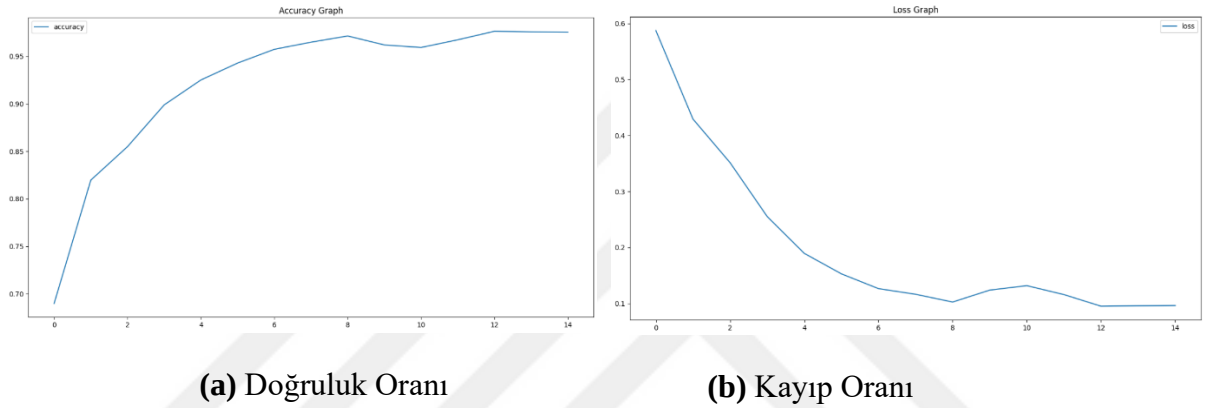
Çalışma kapsamında gerçekleştirilen deneyler ile bu deneylere ilişkin elde edilen sonuçlar **Tablo 4.3'**de sunulmuştur. Deneyler kapsamında LSTM-CLOUD modeli 5, 10 ve 60 eğitim dönemi (epoch) dahil olmak üzere farklı sayıda dönemle ve 2, 4 ve 6 katmanla eğitilmiş ve test edilmiştir. Bu deneyler üç ana başlıkta değerlendirilebilir: Doğruluk, Eğitim Süresi ve Test Süresi. 1) *Doğruluk (accuracy)*: Tablodaki deneylerin doğruluk oranları, modelin ne kadar iyi performans gösterdiğini göstermektedir. Tabloda görüldüğü gibi deneylerin doğruluk oranları %50,40 (Deney 1) ile %99,89 (Deney 8) arasında değişmektedir. Bu durum Deney 8'in diğerlerinden daha iyi performans gösterdiğini göstermektedir. 2) *Eğitim Süresi*: Eğitim süresi, bir modelin belirli bir eğitim veri seti üzerinde ne kadar süreyle eğitildiğini gösterir. Tabloda 5,30 msn (Deney 4) ile 23,24 msn (Deney 6) arasında değişen eğitim süreleri gösterilmektedir. Daha uzun eğitim süreli deneyler, daha kısa eğitim süreli deneylere göre daha fazla zaman gerektirir. 3) *Test Süresi*: Test süresi, eğitilen modelin belirli bir test veri seti üzerinde tahminlerde bulunması için gereken süreyi temsil eder. Test süreleri 0,43 msn (Deney 1) ile 0,62 msn (Deney 5) arasında değişmektedir, bu da test sürelerinin genellikle kısa olduğunu göstermektedir.

Tablo 4.3: Deneylere İlişkin Sonuçlar

Tecrübe. Hayır.	Sayı Gizli Katmanlar	Sayı Çağlar	Eğitim Zaman	Test yapmak Zaman	RMSE	Kesinlik (%)	Kayıp (%)
1.	2	5	5.37	0,43	0,494	%50.40	%49.6
2.	2	15	7.38	0,46	0,464	%72.57	%27.43
3.	2	60	18.47	0,51	0,278	%87.66	%12.34
4.	4	5	5.30	0.60	0,454	%72.80	%27.2
5.	4	15	7.35	0,62	0,296	%91.31	%8.69
6.	4	60	23.24	0,50	0,145	%97.71	%2,29
7.	6	5	8.10	0,46	0,457	%75.20	%24.8
8.	6	15	8.89	0,46	0,108	%99.89	%0,4
9.	6	60	23.04	0,51	0,080	%99.43	%0,57

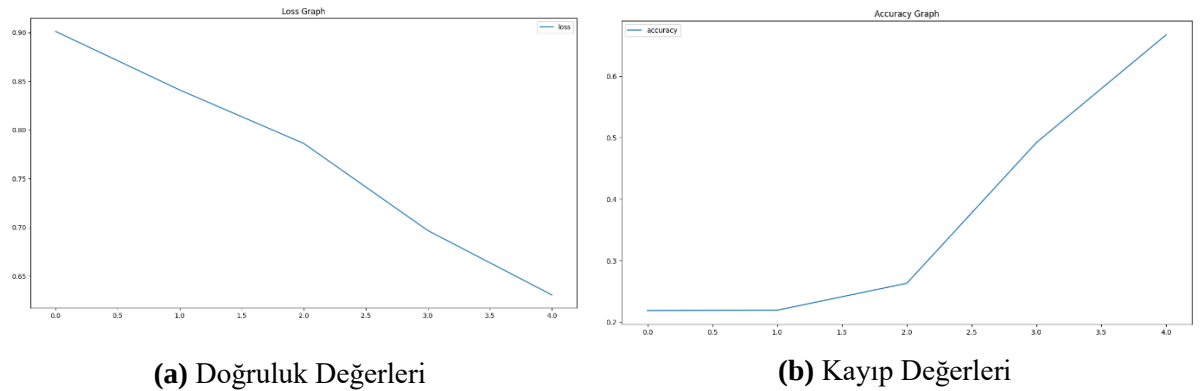
Şekil 4.7'de sunulan ve Deney 8'de elde edilen sonuçların, doğruluk açısından daha yüksek performansıyla ön plana çıktığı görülmektedir. Bu deneyde model, nispeten karmaşık bir sinir ağı mimarisine işaret eden altı gizli katmanla yapılandırılmıştır. Bu deneyin eğitim süreci 15

dönemi içermektedir. Daha karmaşık mimarisine ve yaklaşık 8,89 milisaniye süren daha uzun eğitim süresine rağmen Deney 8, görevi yalnızca 0,46 milisaniyede tamamlayarak test veri seti üzerinde tahminler yaparken etkileyici bir verimlilik sergilemiştir. Bu deneyde, modelin tahmin performansı RMSE metriği kullanılarak değerlendirilmiş ve 0,108 gibi etkileyici derecede düşük bir RMSE değeri elde edilmiştir. Bu, modelin tahminlerinin test veri kümesindeki gerçek değerlerle yakından eşleştiğini göstermektedir. Özellikle, Deney 8’de yalnızca %0,4'lük bir kayıp (loss) değeri elde edilmiştir. Bu durum modelin tahmin hatalarını en aza indirme konusundaki doğruluğunu ve yeterliliğini vurgulamaktadır.



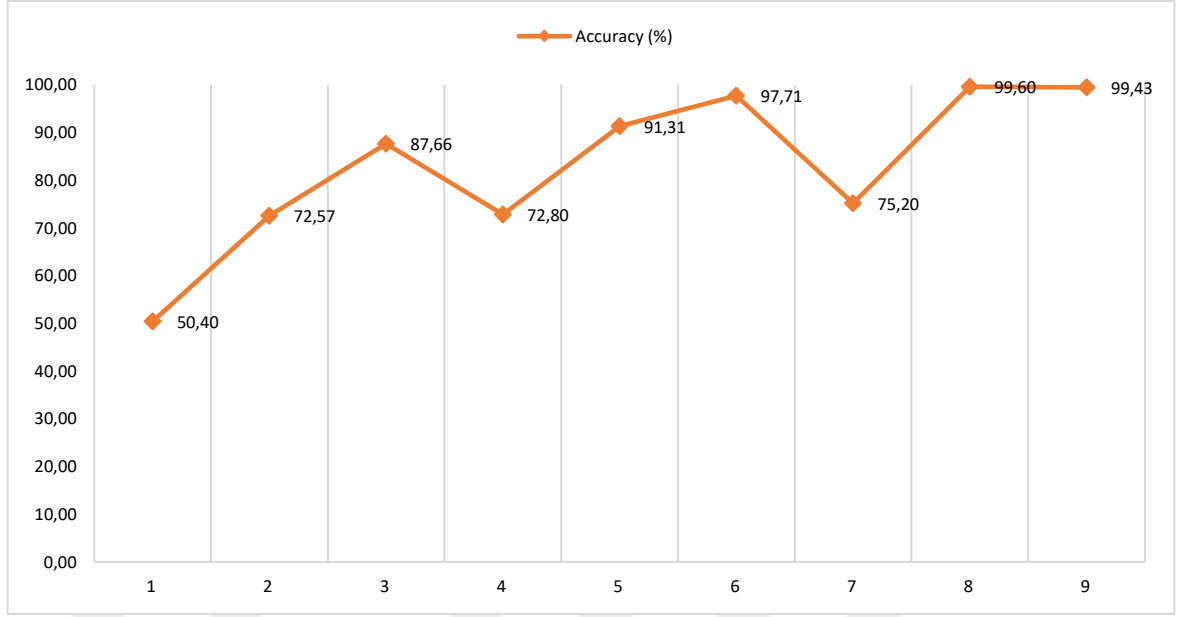
Şekil 4.7: Deney 8 Doğruluk ve Kayıp Değerleri

Şekil 4.8’de en düşük doğruluğun elde edildiği Deney 1’de elde edilen değerler grafiksel olarak gösterilmektedir.



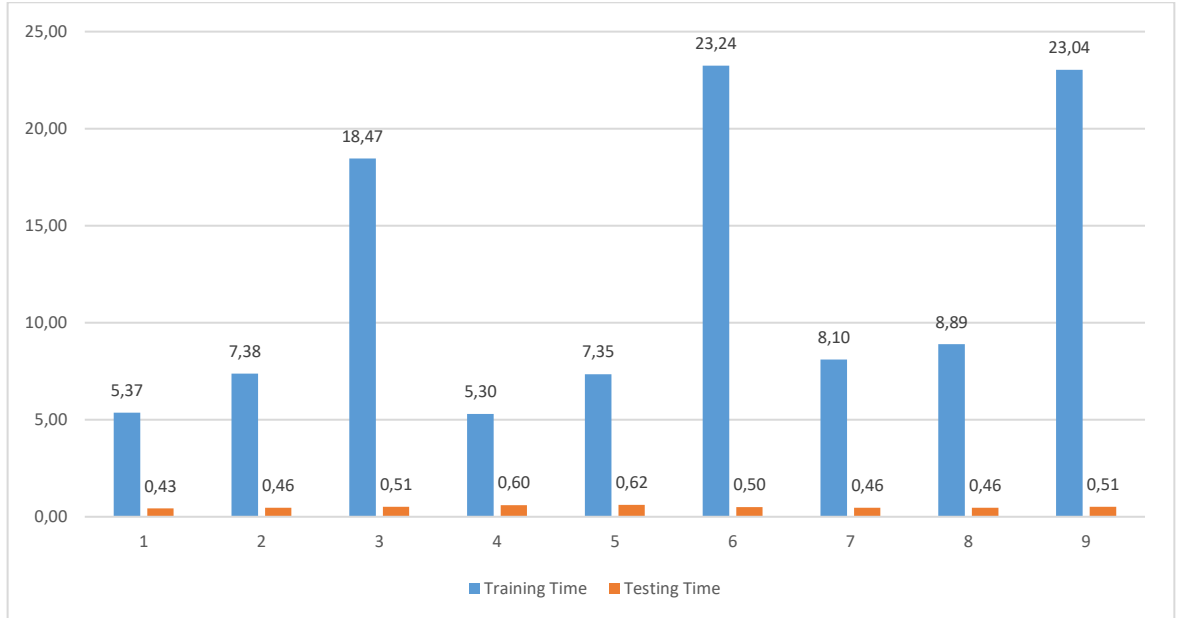
Şekil 4.8: Deney-1 Doğruluk ve Kayıp Değerleri

Deneylerin doğruluk ve kayıp değerlerine ilişkin sonuçlar Şekil 4.9’da gösterilmektedir.



Şekil 4.9: Deneylere İlişkin Doğruluk Değerleri

Gerçek zamanlı olaylar için DDoS saldırılarının tespit hızı çok önemli olduğundan modelin eğitim ve test sürelerini de çalışmalar kapsamında ayrıca hesaplanmıştır. Bu çalışmada gerçekleştirilen deneylerden elde edilen eğitim ve test sürelerinin sonuçları **Şekil 4.10**'da gösterilmektedir.



Şekil 4.10: Deneylere İlişkin Eğitim ve Test Süresi Değerleri (msn)

Son olarak modelin zaman karmaşıklığı hesaplanmıştır. LSTM mimarisi, zaman adımı ve ağırlık başına $O(1)$ hesaplama karmaşıklığına sahip olarak hem zamansal hem de uzaysal boyutlarda yerleştirilmiş verimlilik sergiler (Hochreiter & Schmidhuber, 1997). Sonuç olarak her ağırlığa $O(1)$ değerinde bir zaman karmaşıklığı atanır. LSTM'nin bireysel zaman adımı başına zaman karmaşıklığı göz önüne alındığında, $O(w)$ olarak ifade edilebilir. Burada 'w', kullanılan ağırlık sayısını temsil eder. Özetle, bu analiz, önerdiğimiz yaklaşımımızın asimptotik gösterim bağlamında bir O karmaşıklığını koruduğunu göstermektedir.

Genel bir bulut ağını hacimsel DDoS saldırılarına karşı farklı rollere sahip ajanlar kullanarak korumak, etkili bir güvenlik stratejisi ortaya koymaktadır. Siber güvenlik amacıyla geliştirilen ajanlar, çeşitli noktalarda görev yaparak saldırıyı tespit etmek, trafiği analiz etmek, kötü amaçlı trafiği filtrelemek, meşru trafiği yönlendirmek gibi farklı görevleri yerine getirebilir. Bu sayede saldırının etkisi azaltılabilir ve ağıın sürekliliği sağlanabilir. Ek olarak, ajanların minimum insan müdahalesi ile otonom olarak koordine edilmesiyle, saldırılar tespit edilip çevik müdahale sağlanarak azaltılabilir. Farklı rollere sahip ajanların kullanılması, genel bulut ağlarının hacimsel DDoS saldırılarına karşı savunma mekanizmalarını geliştirebilir. CDA-CLOUD ajanları (VDDA ajanı, DA ajanı, RMA ajanı ve MCA ajanı), özellikle ölçeklenebilir ve dinamik yapıları göz önüne alındığında, bulut ağlarının korunmasında çok önemli bir rol oynayabilir. Özellikle hızlı tespit ve yanıt talep ederek hacimsel DDoS saldırılarına karşı savaşır. VDDA ajanlarının esas görevi, acil izinsiz giriş tespitine odaklanarak anormallikleri belirlemek ve saldırıları önlemek için önemli ağ trafiğini analiz etmek olarak belirlenmiştir. Bu bağlamda DA'lar savunma önlemlerine odaklanır, RMA'lar kaynakları verimli bir şekilde yönetir ve MCA'lar sistem koordinasyonunu denetler. Bu ajanlar, etkili bir şekilde bilgi alışverişinde bulunmak ve saldırılara karşı savunmayı güçlendirmek için yakın iş birliği içinde çalışırlar. Bu dört değişik tür ajanın rolleri, genel bir bulut ağını hacimsel DDoS saldırılarına karşı korumak için belirlenmiştir. Ancak ağıın güvenliğini ve dayanıklılığını artırmak için ek görevlerin dikkate alınması gerekebilir. VDDA ajanı, saldırı tespiti için hassas bir model kullanarak saldırılara hızlı yanıt verme sorumluluğunu üstlenmektedir. Hızlı saldırı tespiti çok önemli olsa bile bu durum bu ajanın ek eğitim veya öğrenme yeteneklerinden yararlanılmayacağı anlamına gelmemektedir. DA ajanı, önceden tanımlanmış modellere veya planlara dayalı olarak saldırıları tespit ettikten sonra savunma önlemlerini uygulamak için daha bilinçli kararlar alabilecek şekilde tasarlanmıştır. Bu durum ağ dayanıklılığına katkıda bulunabilir. Ancak saldırı türlerinin ve yöntemlerinin sürekli gelişen doğası göz önüne alındığında, bu ajanın

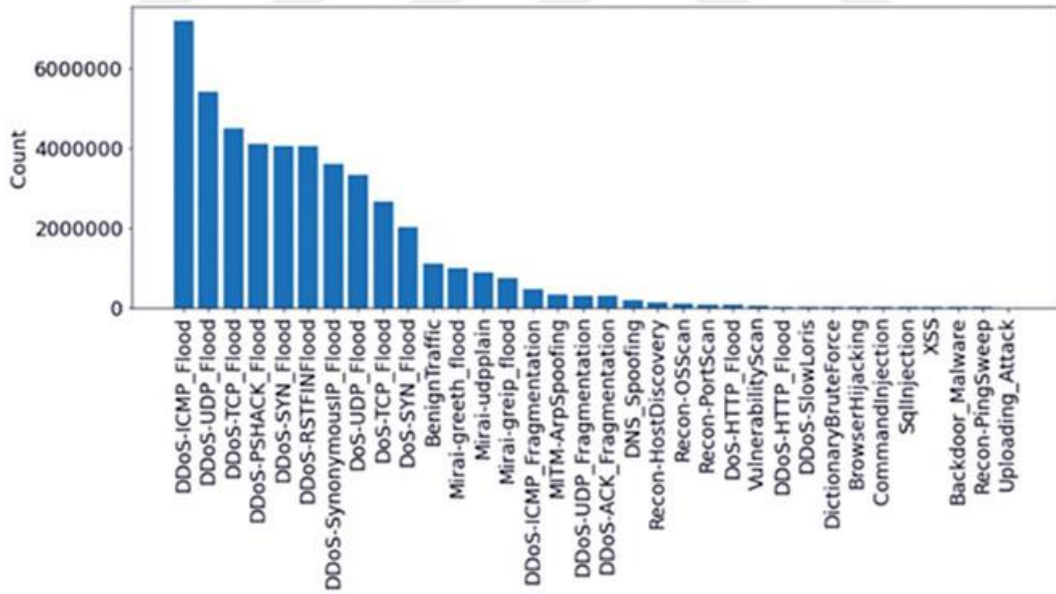
planlarını güncelleme ve öğrenme becerisine sahip olması gerekir. RMA ajanı, kaynakları verimli bir şekilde yönetebilme ve saldırılar sırasında kaynakların tükenmesini önleme yeteneğine sahip olacak şekilde tasarlanmıştır. Bu durum ağın korunması açısından çok önemlidir. Ancak, saldırı altındaki hizmetleri yeniden yönlendirmek veya yedekleme kaynaklarını otomatik olarak tahsis etmek gibi ek görevler, ağın genel sağlığını korumak için dikkate alınmalıdır. Diğer ajanları yönetmek ve koordine etmekle görevli MCA ajanının da öğrenme yetenekleri olmalıdır, çünkü ağ güvenliğinin sürekli değişen tehditlere karşı güncel kalması gerekir ve deneyimlerden ders almak bu amaç için değerlidir. VDDA ajanının hassas saldırı sınıflandırması için imza tabanlı LSTM-CLOUD modelinden yararlanma kapasitesi son derece değerlidir. Bununla birlikte, saldırıları tespit etmede imza tabanlı modellerin, önceden tanımlanmış saldırı imzalarına güvenmeleri, yanlış pozitiflere karşı duyarlılıkları, yeni tehditleri tespit etmedeki verimsizlikleri ve gelişen saldırı tekniklerine uyum sağlama konusundaki eksiklikleri de dahil olmak üzere sınırlamaları vardır ve bu da onları modern siber güvenlikte daha az etkili hale getirir. İmza tabanlı modeller bilinen saldırı modellerini tanıma konusunda üstün olduğu ancak ortaya çıkan tehditlerle mücadele ettiği için bu sınırlamaların kabul edilmesi önemlidir.

Bu çalışmada, bulut ortamlarındaki hacimsel DDoS saldırılarını tespit etmek ve azaltmak için çok ajanlı bir siber savunma sistemi (CDA-CLOUD) tasarlanmıştır. Çalışma kapsamında imza tabanlı bir yaklaşım kullanılarak LSTM-CLOUD modelinin geliştirilmesine yönelik çeşitli deneyler gerçekleştirilmiştir. Çalışma boyunca nicel araştırma yöntemleri kullanılmıştır. Bu araştırmadan öne çıkanlar şu şekilde özetlenebilir:

- Bulut ağ güvenliğinde gerçek zamanlı hacimsel DDoS tehditleriyle etkili bir şekilde mücadele etmek için çok ajanlı sistemleri imza tabanlı modellerle bütünleştiren bir siber savunma sistemi (CDA-CLOUD) tasarlanmıştır.
- Hacimsel DDoS saldırılarının tespiti ve azaltılması için özel olarak LSTM-CLOUD modeli tasarlanmıştır.
- İmza tabanlı algoritmalarda en yüksek doğruluğu elde ajanın önemi yanında ayrıca eğitim/test süresi, karmaşıklık ve hesaplama kaynakları gibi performans ölçümlerinin dikkate alınmasının önemi ortaya konmuştur.

4.3. Derin Pekiştirmeli Öğrenme (DRL) Tabanlı Modelin Analizi

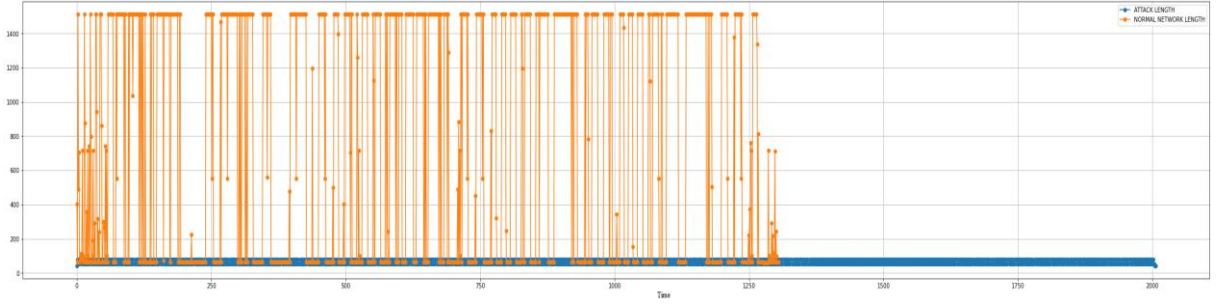
Geliştirilen DQN-IoT ajanının etkinliğini ve verimliliğini göstermek için iki grup halinde çeşitli deneyler yapıldı. Ajanlar oluşturulduktan sonra her biri kendi eğitim ve test setleri kullanılarak eğitildi ve test edildi. Deneylerdeki tüm kodlar, Google Collaboratory tarafından sağlanan GPU ve TPU hizmetlerinin yüksek performansından yararlanılarak başarıyla yürütüldü. Tüm bu eğitim ve test deneyleri birden çok kez gerçekleştirildi. Deneylerde CIC-IoT-2022 ve CIC-IoT-2023 veri setleri kullanıldı (Kanada Siber Güvenlik Enstitüsü, 2023). Bu veri kümeleri farklı türde DDoS saldırıları ve gerçekçi trafik profilleri içerir. Veri kümesi hem zararsız hem de yaygın DDoS saldırılarını içeriyor. **Şekil 4.11**'de CIC-IoT-2023 veri kümesindeki DDoS saldırılarının dağılımı gösterilmektedir. Veri kümesindeki DDoS saldırıları arasında ACK parçalanması, UDP seli, SlowLoris , ICMP seli, RSTFIN seli, PSHACK seli, HTTP seli, UDP parçalanması, TCP seli, SYN seli ve SynonymousIP seli yer alıyor. Bu veri setinin oluşturulmasında 105 cihazdan oluşan IoT topolojisinde kötü amaçlı IoT cihazları tarafından diğer IoT cihazlarını hedef alan 33 saldırı gerçekleştirildi. Bu nedenle CIC-IoT-2023, IoT ortamındaki büyük ölçekli saldırılar için gerçek zamanlı bir veri seti ve kıyaslama sağlar.



Şekil 4.11: CIC-IoT-2023 Veri Seti

4.3.1. CIC-IoT-2022 Veri Seti Üzerinde Önerilen Modelin Analizi

Birinci grup deneyler CIC-IoT-2022 veri seti üzerinde gerçekleştirildi. **Şekil 4.12**'de ağ akışlarının dağılımını göstermektedir.



Şekil 4.12: CIC-IoT-2022 Veri Seti Akış Dağılımı

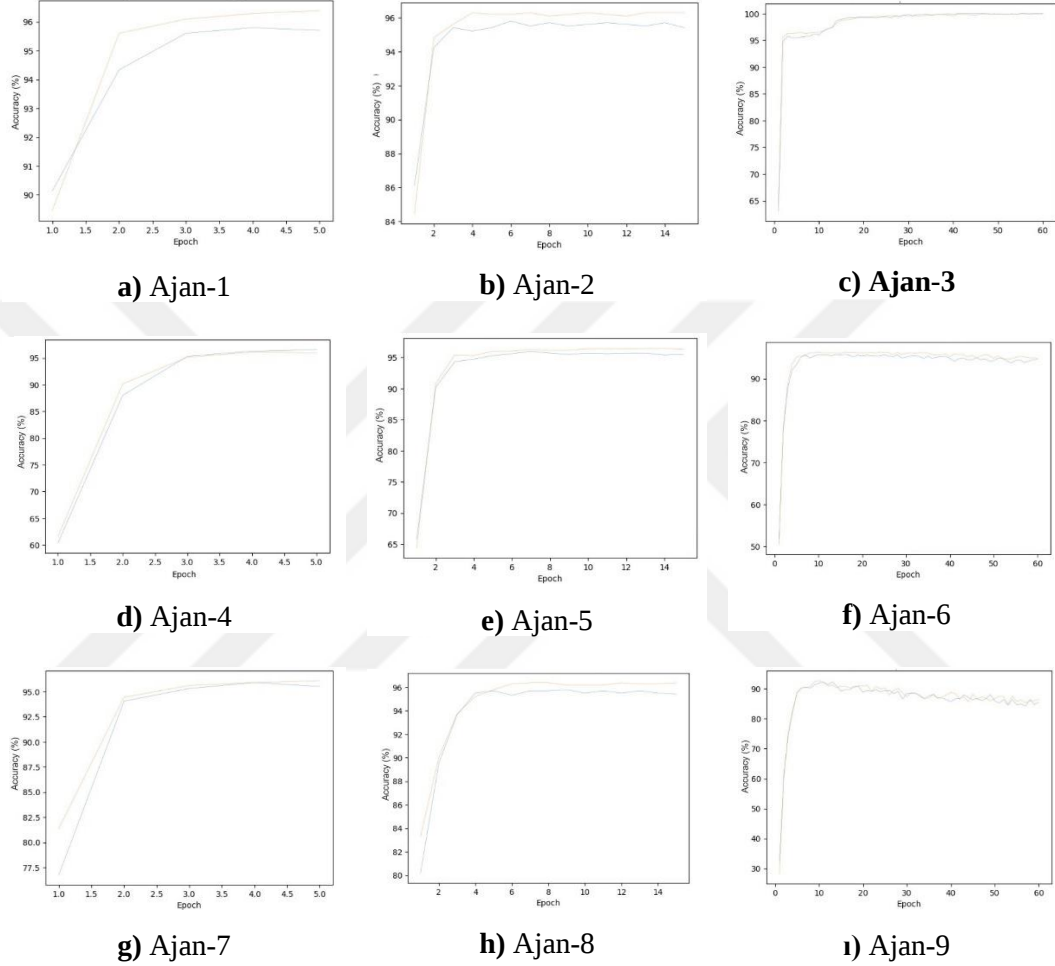
Birinci grup deneyinin bulguları **Tablo 4.4**'de gösterilmektedir. Bu deney grubunda, 2 gizli katman ve 5 dönem ile yapılandırılan Ajan-1, 0,79 saniyelik eğitim süresi ve 0,09 saniyelik test süresi ile %96,39'luk bir doğruluk elde etti. Benzer şekilde, aynı sayıda gizli katmana ancak 15 döneme sahip olan Ajan-2, sırasıyla 2,37 saniye ve 0,27 saniyelik daha uzun eğitim ve test süreleriyle %96,29'luk biraz daha düşük bir doğruluk elde etti. 2 gizli katmana sahip olan ve 60 periyot boyunca eğitilen Ajan-3, diğer ajanlara göre daha uzun eğitim ve test süreleri gerektirmesine rağmen %98,43 ile en yüksek doğruluğu sergiledi. Buna karşılık, 6 gizli katman ve 60 dönemle yapılandırılan Ajan-9, 7,61 saniyelik en uzun eğitim süresiyle birlikte %86,79'luk en düşük doğruluğu sergiledi.

Tablo 4.4: Model Üzerinde Yapılan Deneylerin Değerlendirilmesi

Deney Grupları	Tecrübe. Hayır.	Temsilciler	Gizli Katmanlar	Dönem	Antrenman vakti	Test Süresi	RMSE	Kesinlik (%)	Kayıp (%)
1. Grup Deneyleri (CIC-IoT-2022 Veri kümesi)	1.	Ajan-1	2	5	0,79s	0,09s	%0,22	96.39	3.61
	2.	Ajan-2	2	15	2.37s	0,27s	%0,21	96.29	3.71
	3.	Ajan-3	2	60	6.84s	0,144s	%0,22	98.43	1.57
	4.	Ajan-4	4	5	0.80s	0,10s	%0,21	96.09	3.91
	5.	Ajan-5	4	15	2.40s	0,28s	%0,21	95.51	4.49
	6.	Ajan-6	4	60	6.90'lar	0,145s	%0,23	94.63	5.37
	7.	Ajan-7	6	5	0.88s	0.11s	%0,22	95.90	4.10
	8.	Ajan-8	6	15	2.64s	0,33s	%0,22	95.41	4.59
	9.	Ajan-9	6	60	7.61s	0,22s	%0,38	86.79	13.21

Söz konusu deneylerde elde edilen doğruluk sonuçları **Şekil 4.13**'de sunulmaktadır. Bu bulgular, çeşitli yapılandırmaların ajanların veri kümesini yönetme performansı üzerindeki etkisini vurgulamaktadır. Bu deneyde birkaç dikkate değer nokta ortaya çıkıyor. Öncelikle Agent-3, 2 gizli katman ve 60 dönem ile %98,43'lük yüksek bir doğruluk elde etmesiyle öne çıkıyor ve bu da daha derin ağ mimarilerinin ve daha uzun eğitim sürelerinin üstün sonuçlar verebileceğini gösteriyor. Bunun tersine, Agent-9, 6 gizli katman ve 60 dönem ile %86,79'luk daha düşük bir doğruluk sergiliyor; bu da aşırı uyum veya artan model karmaşıklığı gibi

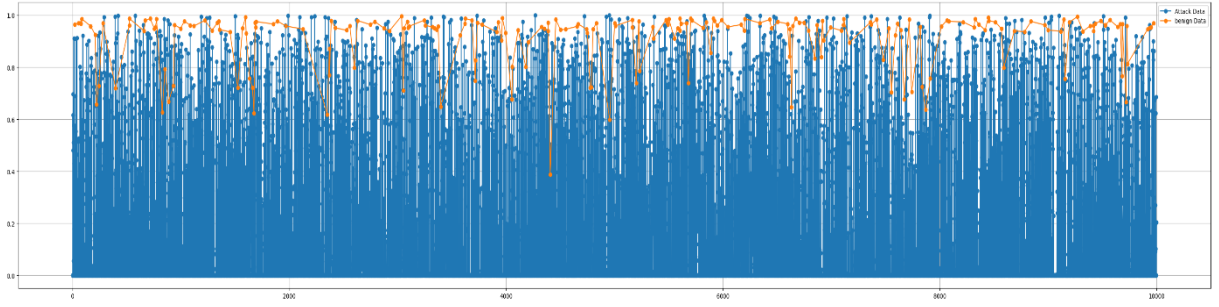
potansiyel sorunlara işaret ediyor. Ek olarak, Agent-3 gibi daha karmaşık modellerin eğitim ve test süreleri oldukça uzundur ve bunun uygulama bağlamına bağlı olarak sonuçları olabilir. Bununla birlikte, tüm ajanlar arasında RMSE değerleri nispeten düşüktür, bu da verilere genel olarak iyi bir uyumun olduğunu gösterir.



Şekil 4.13: 1'nci Grup Deneyler

4.3.2. CIC-IoT-2023 Veri Seti Üzerinde Önerilen Modelin Analizi

CIC-IoT-2023 veri seti üzerinde ikinci grup deneyler yapıldı. Şekil 4.14'de ağ akışlarının dağılımını göstermektedir.



Şekil 4.14: CIC-IoT-2023 Veri Seti Akış Dağılımı

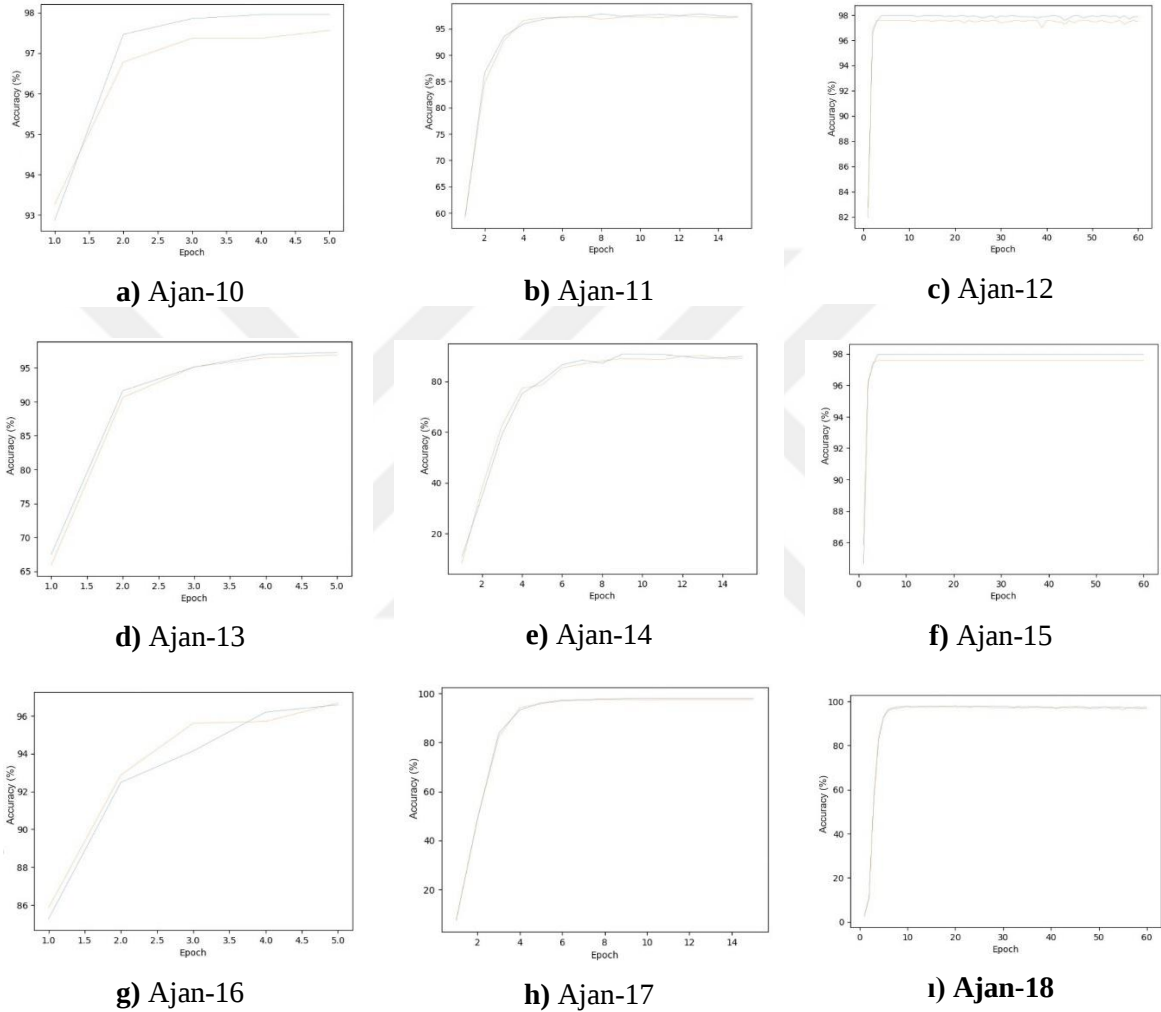
Bu deney grubunda çeşitli ajanlar farklı konfigürasyonlarla değerlendirildi (**Tablo 4.5**). 2 gizli katmanla donatılan ve 5 dönem boyunca eğitilen Agent-10, %0,14'lük minimum RMSE ile %96,46'lık etkileyici bir doğruluk elde etti. Benzer şekilde, yine 2 gizli katmana sahip olan ancak sırasıyla 15 ve 60 dönem için eğitilmiş olan Ajan-11 ve Ajan-12, dönem sayısı ile birlikte artan doğruluk sergileyerek %97,85 ve %97,56'ya ulaştı. Bunun tersine, 4 gizli katmana sahip olan ve 15 dönem için eğitilmiş olan Ajan-14, RMSE'de %0,32'ye kadar gözle görülür bir artış sergiledi ve doğrulukta da %90,23'e kadar kayda değer bir düşüş sergiledi; bu, potansiyel aşırı uyum veya karmaşıklık sorunlarına işaret ediyor. Özellikle, 6 gizli katman ve 60 dönem ile yapılandırılan Agent-18, %0,16'lık oldukça düşük bir RMSE ve %1,95'lik minimum kayıp ile %98,05'lik en yüksek doğruluğu elde etmesiyle öne çıktı.

Tablo 4.5: Model Üzerinde Yapılan Deneylerin Değerlendirilmesi

Deney Grupları	Tecrübe. Hayır.	Temsilciler	Gizli Katmanlar	Dönem	Antrenman vakti	Test Süresi	RMSE	Kesinlik (%)	Kayıp (%)
2. Grup Deneyleri (CIC-IoT-2023 Veri kümesi)	1.	Ajan-10	2	5	0.30s	0,08s	%0,14	96.46	3.54
	2.	Ajan-11	2	15	0.40s	0.30s	%0,15	97.85	2.15
	3.	Ajan-12	2	60	1.20s	0.40s	%0,14	97.56	2.44
	4.	Ajan-13	4	5	0.45s	1.20s	%0,15	96.88	3.12
	5.	Ajan-14	4	15	1.50s	0.45s	%0,32	90.23	9.77
	6.	Ajan-15	4	60	3.20s	1.50s	%0,13	97.95	2.05
	7.	Ajan-16	6	5	0.55s	5.60s	%0,17	93.34	6.66
	8.	Ajan-17	6	15	2.10s	0.55s	%0,14	97.56	2.44
	9.	Ajan-18	6	60	3.78s	0,066s	%0,16	98.05	1.95

1. Grup Deneylerinden elde edilen doğruluk sonuçları **Şekil 4.15**'de sunulmaktadır. Bu deneyden birkaç dikkate değer gözlem ortaya çıkıyor. İlk olarak, Agent-10, Agent-11 ve Agent-12 gibi daha az gizli katman ve dönemle eğitilen modeller, daha basit mimarilerine rağmen yüksek doğruluk sergiledi. Bu, daha az karmaşık modellerin bile övgüye değer sonuçlar verebileceğini göstermektedir. Bununla birlikte, 4 gizli katman ve 15 dönemle eğitilen Agent-14, potansiyel aşırı uyum veya karmaşıklık sorunlarına işaret ederek doğrulukta kayda değer bir azalmanın yanı sıra RMSE'de önemli bir artış sergiledi. Buna karşılık Agent-18, 6 gizli

katman ve 60 dönem içeren daha derin mimarisine rağmen yüksek doğruluk ve düşük RMSE ile olağanüstü bir performans sergileyerek optimize edilmiş bir model yapısının etkinliğini vurguladı. Bu bulgular, farklı konfigürasyonların model performansı üzerindeki farklı etkisinin altını çiziyor ve üstün sonuçlar elde etmek için karmaşıklık ile optimizasyonun dengelenmesinin önemini vurguluyor.



Şekil 4.15: 2'nci Grup Deneyler

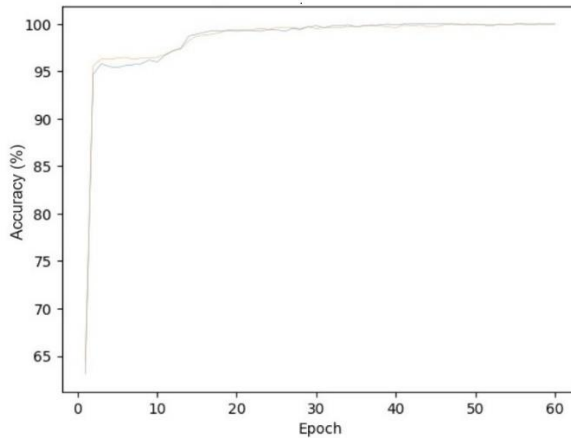
Tablo 4.6'da verilen eğitim ve test süreleri, farklı veri kümeleri üzerinde yürütülen deneylere ilişkin bilgiler sunmaktadır. Deneyleri doğruluk açısından değerlendirirken, CIC-IoT-2022 ve CIC-IoT-2023 deneyleri arasındaki hesaplama süresindeki azalmaya rağmen doğrulukta %0,22'den %0,16'ya hafif bir düşüş olması dikkat çekicidir. Bu, hesaplama verimliliği ile model performansı arasında potansiyel bir değiş tokuş olduğunu göstermektedir. Eğitim ve test sürelerinin azaltılması pratik uygulama açısından avantajlı olsa da yüksek doğruluğun

sürdürülmesi, gerçek dünya senaryolarında modellerin etkinliği açısından kritik olmaya devam etmektedir. DDoS algılama hızı çok önemlidir çünkü saldırıları azaltmak için yanıt süresini doğrudan etkileyerek potansiyel kesinti süresini ve ağ altyapısına verilen zararı en aza indirir. Hızlı algılama, hızlı eyleme olanak tanıyarak genel ağ güvenliğini ve siber tehditlere karşı dayanıklılığı artırır. Sonuçlar, CIC-IoT-2022 ve CIC-IoT-2023 etiketli deneyler arasında eğitim ve test sürelerinde fark edilebilir bir fark olduğunu gösteriyor. CIC-IoT-2022 ile karşılaştırıldığında CIC-IoT-2023 deneyinde eğitim süresinin 6,84 saniyeden 3,78 saniyeye ve test süresinin 0,144 saniyeden 0,066 saniyeye düşürülmesi, model mimarisi veya eğitim metodolojilerindeki potansiyel optimizasyonları veya iyileştirmeleri sergiliyor. Hesaplama süresindeki azalmaya rağmen deneyler minimum kayıpla yüksek doğruluk düzeylerini koruyor ve bu da modellerin veri kümelerini verimli bir şekilde işleme ve analiz etmedeki etkinliğini vurguluyor.

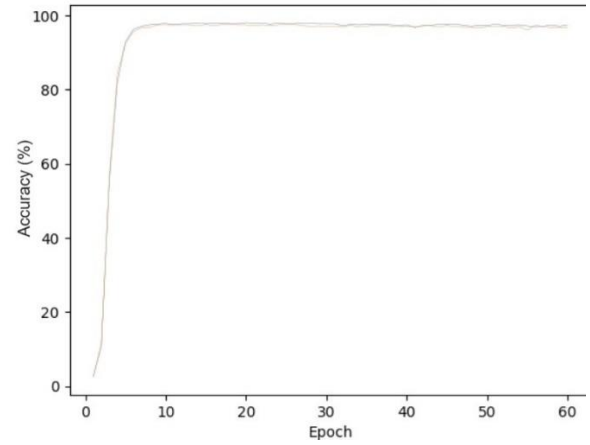
Tablo 4.6: Deneylerde Kullanılan Veri Setleri Karşılaştırması

Deney Etiketi	Gizli Katmanlar	Çağlar	Antrenman vakti	Test Süresi	Kesinlik (%)	Kayıp (%)
CIC-IoT-2022	2	60	6.84s	0,144s	%0,22	98.43
CIC-IoT-2023	6	60	3.78s	0,066s	%0,16	98.05

Şekil 4.16'da farklı veri kümeleri arasında algoritma doğruluğunun karşılaştırılmasını göstermektedir. Bu grafik, CIC-IoT-2022 ve CIC-IoT-2023 olmak üzere iki veri kümesinin doğruluk yüzdelerini gösterir. Doğruluk yüzdesi, sınıflandırma veya tespit algoritmalarının her veri kümesinde ne kadar etkili performans gösterdiğini belirtir. Siber güvenlikte, özellikle DDoS saldırılarının azaltılmasında hızlı ve doğru tespit çok önemlidir. Bir sistem bir saldırıyı ne kadar hızlı algılayabilirse, karşı önlemleri de o kadar hızlı başlatabilir, böylece ağ kullanılabilirliği ve performansı üzerindeki potansiyel etki azalır. Bu nedenle, Şekil 10'da gösterildiği gibi doğruluk yüzdeleri daha yüksek olan algoritmalar, siber tehditlere karşı güçlü savunma mekanizmalarının sağlanması açısından büyük önem taşımaktadır.



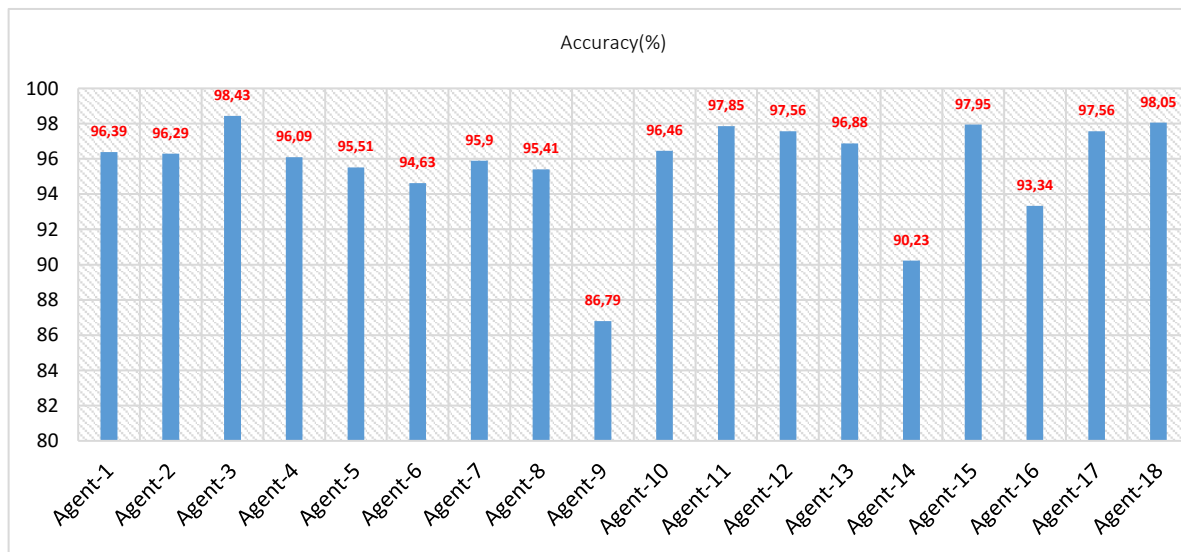
(a) CIC-IoT-2022



(b) CIC-IoT-2023

Şekil 4.16: Veri Seti Bazında En Yüksek Doğruluk Oranları

Şekil 4.17’de, çalışmada eğitilen ve test edilen tüm ajanların doğruluk oranlarının grafiksel gösterimini sunmaktadır. Doğruluk oranları ajanlar arasında değişiklik göstermektedir; bazıları %98,43’ün üzerinde yüksek doğruluk oranlarına ulaşırken diğerleri bunun altına düşmekte ve en düşük doğruluk oranı %86,79 olarak kaydedilmiştir. Bu, temsilcilerin görevlerini veya sorumluluklarını yerine getirirken performans veya etkililik açısından farklılıklar olduğunu göstermektedir. Daha yüksek doğruluk oranlarına sahip ajanlar, rollerini etkili bir şekilde yerine getirme konusunda daha iyi bir yetenek sergilerken, daha düşük oranlara sahip olan ajanlar, daha fazla optimizasyon veya ince ayar gerektirebilir.



Şekil 4.17: Ajanların Performans Karşılaştırmaları

4.4. Önerilen Modellerin Karşılaştırmalı Analizi

Çalışmanın bu bölümünde elde edilen bulguların karşılaştırmalı bir analizi yer almaktadır. Önerilen modellerin karşılaştırmalı analizi, üç ana başlık altında gerçekleştirilmiştir. İlk olarak, yöntem türü ve model özellikleri incelenmiştir. Bu kapsamda kullanılan metotlar, temel model ve algoritmalar olarak LSTM ve DQN tekniklerini içermektedir. Yöntemler, imza tabanlı veya anomali tabanlı yaklaşımları kapsar ve bu da tespit yöntemlerinin saldırı türlerini belirleme şekillerini ifade eder. Hedef ağ yapıları ise uygulamanın yapıldığı ağ türlerine göre sınıflandırılır; bunlar arasında Bulut Ağları, IoT Ağları ve Endüstriyel IoT Ağları bulunmaktadır.

İkinci başlık altında performans ölçütleri değerlendirilmiştir. Doğruluk oranı, modelin ne kadar doğru tahminlerde bulunduğunu gösterir ve yüksek doğruluk oranı modelin etkinliğini işaret eder. Kayıp oranı ise modelin hata oranını belirler; düşük kayıp oranları, modelin performansının yüksek olduğunu gösterir. Son olarak, karmaşıklık ve hesaplama maliyetleri ele alınmıştır.

Üçüncü başlık altında ise karmaşıklık ölçütleri değerlendirilmiştir. Karmaşıklık, algoritmanın hesaplama karmaşıklığını (örneğin, $O(n)$ veya $O(m^2)$) belirtir ve modelin eğitim sürecinde gereken hesaplama kaynaklarını belirler. Eğitim ve test zamanı ise modelin bu süreçlerde geçirdiği süreyi ifade eder; daha kısa süreler, modelin eğitim ve test aşamalarında daha hızlı ve verimli olduğunu gösterir.

Tablo 4.7'de siber saldırı tespiti amacıyla bu araştırmada aşamalı olarak önerilen ve geliştirilen yöntemler sunulmaktadır.

4.4.1. Yöntem Türü ve Model Özellikleri Analizi

Birinci aşama, uzun kısa süreli bellek (LSTM) tabanlı bir model kullanarak bulut ağlarında imza tabanlı bir siber saldırı tespit yöntemi önermektedir. Bu model, tanımlanmış saldırı desenlerini kullanarak belirli saldırı türlerini tespit etmeye yöneliktir ve CICDDoS2019 veri seti üzerinde test edilmiştir. İmza tabanlı yöntemler, belirlenmiş saldırı desenlerini tanıyarak bilinen saldırıları algılar. İkinci aşamada, uzun kısa süreli bellek (LSTM) modeline ek olarak çok ajanlı sistemler kullanılmıştır. Bu model, IoT ağları üzerinde imza tabanlı bir yaklaşımı benimsemekte olup, çok ajanlı sistemler çeşitli ajanın iş birliği içinde çalışmasını sağlayarak tespit sürecine dinamik bir katman ekler. Veri seti olarak CIC-IoT-2022 kullanılmıştır.

Buradaki çok ajanlı sistemler, saldırı tespitinde daha esnek ve etkili bir yöntem sunar. Üçüncü aşama, yine uzun kısa süreli bellek (LSTM) ve çok ajanlı sistemler kullanarak imza tabanlı bir yöntemi bulut ağları üzerinde uygular. Ancak, bu model önceki aşamadan farklı olarak daha yüksek bir performans göstermektedir ve CICDDoS2019 veri seti üzerinde değerlendirilmiştir. Bu aşamada, modelin performansı daha iyidir ve imza tabanlı yöntemlerin bulut ağlarındaki etkinliği artırılmıştır. Dördüncü aşama, derin Q-ağları (DQN) ve pekiştirmeli öğrenme (RL) yöntemlerini kullanarak anomali tabanlı bir siber saldırı tespit modeli önermektedir. Endüstriyel IoT ağlarında test edilen bu model, alışılmadık davranışları tespit ederek bilinmeyen saldırı türlerini tanımlamak için anomali tabanlı yöntemler kullanır. CIC-IoT-2022 ve CIC-IoT-2023 veri setleri üzerinde değerlendirilmiş olup, bu yaklaşım, gelişmiş ve bilinmeyen saldırılara karşı daha esnek bir savunma sunar.

Dördüncü ve son aşamada, derin Q-ağları (DQN) ve pekiştirmeli öğrenme (RL) yöntemleri kullanılarak anomali tabanlı bir tespit modeli geliştirilmiştir. Model, endüstriyel IoT ağları üzerinde CIC-IoT-2022 ve CIC-IoT-2023 veri setleri ile test edilmiştir. Anomali tabanlı yöntemler, alışılmadık davranışları tespit ederek bilinen saldırı desenlerinden bağımsız olarak yeni ve bilinmeyen saldırıları tanımlama kapasitesine sahiptir. Bu, DQN ve RL tabanlı modelin, geleneksel imza tabanlı yaklaşımların ötesine geçerek, daha esnek ve adaptif bir savunma stratejisi sunduğunu göstermektedir. Eğitim süresi 6.84 ms ve test süresi 3.78 ms ile oldukça hızlı sonuçlar sunmasına rağmen, doğruluk oranları %98.43 (CIC-IoT-2022) ve %98.05 (CIC-IoT-2023) gibi diğer aşamalardan biraz daha düşük olmakla birlikte, modelin anomali tespiti konusundaki yeteneği ve adaptifliği, bilinmeyen ve gelişmiş saldırılara karşı önemli bir avantaj sağlamaktadır.

Yöntem türü ve model analizi bağlamında yapılan karşılaştırma sonuçları, dördüncü aşamanın siber saldırı tespiti konusunda diğer aşamalardan belirgin bir şekilde ayrıştığını ortaya koymaktadır. Bu aşamanın ön plana çıkması, siber güvenlik alanında yenilikçi ve dinamik bir yaklaşımı temsil etmekte ve genişletilmiş saldırı senaryolarına karşı daha kapsamlı bir savunma mekanizması sunmaktadır. İlk olarak, bu modelin (DQN-IoT) doğruluk oranı, %98.43 ve %98.05 olarak belirlenmiş olup, diğer modellerin doğruluk oranlarına yakın ancak daha düşük değildir. Bu, modelin siber saldırıları tespit etme konusunda yüksek bir performans sergilediğini göstermektedir. Bununla birlikte, kayıp oranları da %1.57 ve %1.50 gibi kabul edilebilir seviyelerde kalmaktadır. Bu durum modelin yanlış alarm oranının yönetilebilir olduğunu ve güvenilir bir sonuç sağladığını işaret etmektedir. İkinci olarak, dördüncü aşamanın

veri seti ve hedef ağ yapısı geniş bir kapsamı içermektedir. IIoT ağları üzerinde yapılan analiz, siber güvenliğin kritik olduğu endüstriyel alanlarda uygulama açısından büyük bir önem taşır. Bu durum, modelin çeşitli ve kompleks ağ yapılarına karşı etkili olduğunu ve gerçek dünya uygulamaları için uygun olduğunu göstermektedir. Son olarak, eğitim ve test zamanı açısından, DQN-Iot modelinin performansı etkileyici bir şekilde optimize edilmiştir. Eğitim/test süresi, 6.84 ve 3.78 ms olarak ölçülmüş olup, bu süreler diğer aşamalara kıyasla oldukça kısa ve verimlidir. Bununla birlikte, bu modelin karmaşıklığı $O(m^2)$ olarak sınıflandırılmıştır ki bu, modelin işlem yükünün diğer modellere göre daha yüksek olduğunu gösterse de performans ve doğruluk açısından bu karmaşıklığın getirdiği faydalar öne çıkmaktadır. Görüldüğü üzere, dördüncü aşamanın diğer aşamalara göre daha ön plana çıkmasının sebepleri arasında yüksek doğruluk oranları, yönetilebilir kayıp oranları, geniş veri setleri ve hızlı eğitim/test süreleri gibi faktörler yer almaktadır. Araştırmada önerilen yenilikçi DQN-IoT modeli yaklaşımı DRL tabanlı DQN tekniklerini ve çok ajanlı sistemleri kullanarak, gerçek dünya uygulamaları için güçlü ve etkili bir çözüm sunmaktadır.

Tablo 4.7: Önerilen Yöntemlerin Karşılaştırmalı Analizi

Aşama / Önerilen Model	Metot	Yöntem	Hedef Ağ Yapısı	Veri Seti	Doğruluk Oranı (Accuracy)	Kayıp Oranı (Loss)	Karmaşıklık	Eğitim / Test Zamanı (msec)
1'inci Aşama: Derin Öğrenme (DÖ) Tabanlı Model Önerisi	Uzun Kısa Süreli Bellek (LSTM)	İmza Tabanlı Yöntem	Bulut Ağları	CICDDoS2019	%99.83	%0,17	$O(n)$	11.90, 3.68
2'inci Aşama: Ajan Tabanlı Derin Öğrenme (DÖ) Tabanlı Model Önerisi	Uzun Kısa Süreli Bellek (LSTM), Çok Ajanlı Sistemler	İmza Tabanlı Yöntem	IoT Ağları	CIC-IoT-2022	%99.48	%0.52	$O(n)$	29.2, 3.36
3'üncü Aşama: Ajan Tabanlı Derin Öğrenme (DÖ) Tabanlı Model Önerisi	Uzun Kısa Süreli Bellek (LSTM), Çok Ajanlı Sistemler	İmza Tabanlı Yöntem	Bulut Ağları	CICDDoS2019	%99.89	%0.11	$O(n)$	8.89, 0.46
4'üncü Aşama: Ajan Tabanlı Derin Pekiştirmeli Öğrenme (DRL) ile Siber Saldırı Tespiti Model Önerisi	Derin Q-Ağları (DQN), Pekiştirmeli Öğrenme (RL), Çok Ajanlı Sistemler	Anomali Tabanlı Yöntem	Endüstriyel IoT Ağları (IIoT)	CIC-IoT-2022, CIC-IoT-2023	% 98.43, % 98.05	%1.57, %1.50	$O(m^2)$	6.84, 3.78

4.4.2. Performans Ölçütleri Analizi

Performans açısından yapılan değerlendirmelerde, modellerin doğruluk oranları ve kayıp oranları dikkatle incelenmiştir. Üçüncü aşama, %99.89 doğruluk oranı ile en yüksek performansı sergileyerek diğer aşamalardan üstün bir başarı göstermektedir. İlk aşama da yüksek bir doğruluk oranına sahip olup %99.83 seviyesindedir. İkinci aşama, %99.48 doğruluk oranı ile biraz daha düşük bir performans sunmaktadır. Dördüncü aşama ise anomali tespit yöntemini kullanarak %98.43 (CIC-IoT-2022) ve %98.05 (CIC-IoT-2023) doğruluk oranları ile diğer aşamalardan daha düşük bir performans sergilemektedir. Kayıp oranları açısından, üçüncü aşama en düşük kayıp oranını %0.11 ile elde etmiştir. İlk aşama %0.17 kayıp oranı ile iyi bir performans sunarken, ikinci aşamada kayıp oranı %0.52 seviyesine çıkmıştır. Dördüncü aşama ise daha yüksek kayıp oranlarına sahiptir; CIC-IoT-2022 ve CIC-IoT-2023 veri setlerinde sırasıyla %1.57 ve %1.50 kayıp oranları gözlemlenmiştir. Bu yüksek kayıp oranları, dördüncü aşamanın doğruluk oranlarındaki düşüşü ile paralellik göstermektedir ve modelin performansında belirli bir azalmayı işaret etmektedir.

Performans analizi bağlamında yapılan karşılaştırma sonuçları, dördüncü aşamanın siber saldırı tespiti konusunda belirgin bir avantaj sunarak diğer aşamalardan ayrıldığını göstermektedir. Bu aşamada kullanılan anomali tespit yöntemleri, %98.43 (CIC-IoT-2022) ve %98.05 (CIC-IoT-2023) doğruluk oranları ile diğer aşamalardan biraz daha düşük performans sergilemesine rağmen, bilinmeyen ve gelişmiş saldırılara karşı yüksek derecede esneklik sağlama kapasitesine sahiptir. Üçüncü aşama en yüksek doğruluk oranını (%99.89) sunarken, dördüncü aşamanın doğruluk oranlarındaki küçük düşüş, modelin özellikle endüstriyel IoT ağlarında karşılaşılabilecek çeşitli anomali ve bilinmeyen tehditleri tanımlama yeteneğinin gelişmişliğini ve kapsamını ortaya koyar. Ayrıca, dördüncü aşamanın daha yüksek kayıp oranları (%1.57 ve %1.50) bulunmasına rağmen, bu aşama geniş bir veri seti üzerinde etkili bir şekilde çalışabilme ve bilinmeyen saldırılara karşı yüksek adaptasyon yeteneği sağlama avantajı sunmaktadır. Bu özellikler, dördüncü aşamanın siber güvenlik uygulamalarında önemli bir avantaj sunduğunu ve karmaşık ve bilinmeyen tehditlerle başa çıkma kapasitesinin yüksek olduğunu göstermektedir.

4.4.3. Karmaşıklık Analizi

Performans değerlendirmesinde, modellerin doğruluk oranları ve kayıp oranları detaylı bir şekilde incelenmiştir. Üçüncü aşama, %99.89 doğruluk oranı ile en yüksek performansı sunarak

en başarılı sonuçları sağlamaktadır. İlk aşama da oldukça yüksek bir doğruluk oranına sahip olup %99.83 seviyesindedir. İkinci aşama ise %99.48 doğruluk oranı ile biraz daha düşük bir performans göstermektedir. Dördüncü aşama ise anomali tespit yöntemiyle çalıştığı için doğruluk oranları %98.43 (CIC-IoT-2022) ve %98.05 (CIC-IoT-2023) olarak belirlenmiştir, bu da diğer aşamalara kıyasla daha düşük bir performans seviyesi ifade etmektedir. Kayıp oranları açısından, üçüncü aşama %0.11 ile en düşük kayıp oranını elde etmiştir. İlk aşama %0.17 kayıp oranı ile iyi bir performans sergilerken, ikinci aşama %0.52 kayıp oranına sahiptir ve daha yüksek bir kayıp göstermektedir. Dördüncü aşama ise daha yüksek kayıp oranlarına sahiptir; CIC-IoT-2022 veri setinde %1.57 ve CIC-IoT-2023 veri setinde %1.50 kayıp oranları gözlemlenmiştir. Bu yüksek kayıp oranları, dördüncü aşamanın doğruluk oranlarındaki düşüşle paralel bir performans kaybını işaret etmektedir.

Karmaşıklık ve hesaplama maliyetleri analizi bağlamında yapılan karşılaştırma sonuçları, dördüncü aşamanın siber saldırı tespiti konusunda diğer aşamalardan belirgin bir şekilde ayrılarak öne çıktığını ortaya koymaktadır. Bu aşamada kullanılan DQN ve DRL yöntemleri, anomali tabanlı bir yaklaşımı benimsemekte olup, daha karmaşık bir hesaplama yapısı ile dikkat çekmektedir. Karmaşıklık açısından $O(m^2)$ ile yüksek bir hesaplama yükü gerektirir, bu da modelin daha kapsamlı ve detaylı analizler gerçekleştirdiğini göstermektedir. Eğitim süresi 6.84 ms ile oldukça hızlı olup, test süresi 3.78 ms ile daha uzun olmasına rağmen, modelin geniş veri setleri üzerinde etkili bir şekilde çalışmasını sağlamaktadır. Dördüncü aşamanın, diğer aşamalara göre daha karmaşık bir yapıya sahip olması, bilinmeyen saldırı türlerine karşı daha esnek ve dinamik bir savunma sunma yeteneğini de ayrıca artırmaktadır. Dördüncü ve son aşamada sunulan yenilikçi DQN-IoT modeli yaklaşımı, yüksek hesaplama maliyetleri ve karmaşıklığına rağmen, anomali tespiti konusunda önemli bir avantaj sağlamakta ve IIoT ağlarında karşılaşılabilecek gelişmiş saldırılar karşısında güçlü bir koruma mekanizması sunmaktadır.

5. TARTIŞMA

Geleneksel savunma mekanizmaları birçok eksiklikten mustarıdır. İlk olarak, saldırganın ağ altyapısını güvenlik açıkları için taramak için yeterli zamanı vardır. İkincisi, güvenlik duvarı, IDS, IDPS gibi savunma mekanizmaları genel olarak imza tabanlıdır. Bu nedenle, yalnızca bilinen saldırıları tespit edip durdurabilirler. Ancak geleceğin bilgisayar ağları, cihazlar, ağlar, yazılımlar ve insanlar dahil olmak üzere büyük sayıda birbirine bağlı, heterojen ve bazen otonom ajanlarla dolu oldukça karmaşık ve dinamik yapılar olacaktır. Ve bu yapıları sadece bilinen değil ancak hiç bilinmeyen siber saldırılar da hedef alabilecektir. Bu tür karmaşık ve/veya otonom sistemleri hızlı, karmaşık, YZ destekli ve hatta otonom yapıda gelişmiş siber saldırılara karşı savunulması yalnızca geleneksel saldırı tespit sistemleri (IDS) ile yapılması çok zor ve hatta imkânsız hale gelecektir. Bu bağlamda akıllı siber savunma ajanları bilgisayar ağ sistemlerini savunmada kilit rol oynayacaktır. Bu tez çalışmasında, ağ güvenliği yönetimi için akıllı ajan teknolojilerini kullanarak siber saldırı tespitine yönelik yenilikçi bir yaklaşım önerilmiştir. Çalışmada, dört aşamalı geliştirme yöntemiyle çeşitli özgün ajan tabanlı saldırı tespit modelleri geliştirilmiş, bu modellerin performans değerleri hesaplanarak yöntem türü ve model özellikleri, performans ölçütleri ve de karmaşıklık değerleri açısından karşılaştırılmış ve böylelikle her bir model analiz edilmiştir. Çalışmada önerilen ilk üç aşamada geleneksel yöntemlerin sınırlamaları göz önüne alındığında esasen başarılı sonuçlar elde edilmiştir. İlk aşama, DÖ tabanlı imza tespit yöntemini kullanarak bulut ağları üzerinde etkileyici bir performans sergilemiştir. Ancak, eğitim ve test süreleri diğer aşamalara kıyasla daha uzundur ki bu da hız ve verimlilik açısından dezavantaj oluşturabilir. İkinci aşama, LSTM ve çok ajanlı sistemlerin birleşimiyle IoT ağlarında imza tabanlı tespit gerçekleştirmiştir. Bu aşamada doğruluk oranı biraz daha düşük olmakla birlikte, eğitim süresi oldukça uzun ve test süresi daha verimlidir. Üçüncü aşama, yine LSTM ve çok ajanlı sistemlerle Bulut Ağları üzerinde imza tabanlı tespit sağlamıştır. Bu aşama, en yüksek doğruluk ve en düşük kayıp oranını sunarak performans açısından önemli bir avantaj sağlamaktadır. Eğitim ve test süreleri ise diğer aşamalara göre en hızlı olanıdır. Ancak, imza tabanlı yaklaşımların değişken saldırı türlerine karşı esneklik sunmadığı da göz önünde bulundurulmalıdır. Bu tez çalışmasının dördüncü ve son aşamasında önerilen modelde ise diğer aşamalardan farklı olarak DRL ve DQN kullanılarak anomali tabanlı bir saldırı tespit ve önleme modeli (DQN-IoT) önerilmiştir. Bu tez çalışmasında

edilen sonuçlar, araştırmanın dördüncü ve son aşamasında önerilen özgün ajan tabanlı DQN-IoT modelinin diğer önerilen özgün ajan tabanlı siber saldırı tespit modellerine kıyasla daha ön plana çıktığını ortaya koymaktadır.

Tez çalışmasının 1'inci Aşamada önerilen “LSTM Ağı ile Derin Öğrenme (DÖ) Tabanlı Siber Saldırı Tespiti Yaklaşımının” literatürde mevcut diğer çalışmalar ile karşılaştırılması **Tablo 5.1**'de sunulmuştur.

Tablo 5.1: 1'inci Aşamanın İlgili Bazı Çalışmalar ile Karşılaştırılması

Araştırma	Yıl	Yöntem	Veri kümesi	Başarı oranı
Sahi ve diğ.	2017	LS-SVM, Saf Bayes, K-En yakın, Çok Katmanlı Algılayıcı	Ağ Paketleri	%97
Fang ve diğ.	2018	LSTM	İnternet Sayfaları	%99
Patil ve diğ.	2019	Rastgele Orman	UNSW-NB15, CICIDS-2017	%97
Haider ve diğ.	2020	CNN	CICIDS2017	%99
Tan ve diğ.	2020	K-Anlamına gelir, KNN	NSL-KDD, KDD 99	%98
Novaes ve diğ.	2020	LSTM-Bulanık	CICDoS 2019	%97.89
Elsayed ve diğ.	2020	RNN	CICDoS 2019	%99
Çil ve diğ.	2021	DNN	CICDoS 2019	%95
Rehman ve diğ.	2021	GRU, RNN, Not: SMO	CICDoS 2019	%99.9
Bu Araştırma	2021	LSTM	CICDoS 2019	%99.83

Literatür taramasında incelenen söz konusu çalışmalar, anomali tespitinin, uygulamaları nedeniyle önemli bir sorun olarak görüldüğünü ve farklı veri setleri üzerinde farklı derin öğrenme algoritmaları kullanılarak yaygın olarak çalışıldığını göstermektedir. Önceki çalışmalarla karşılaştırıldığında bu çalışmanın farklılıkları şu şekilde ifade edilebilir: (1) *Performans*: Bu çalışmanın modeli, farklı veya aynı derin öğrenme algoritmaları ve veri kümeleri ile yapılan önceki çalışmalar kadar iyi bir performansa sahipti. Novaes ve diğ. (2020) da aynı veri setini kullanmıştır. Ancak bizim çalışmamızdan farklı olarak LSTM kombinasyonu ile bulanık mantık kullanılmıştır. Elsayed ve diğ. (2020) da aynı veri seti üzerinde RNN algoritmasını kullanarak çalışmalarını gerçekleştirmişlerdir. (2) *Paket Düzeyinde Sınıflandırma*: Çalışmamız, CCICDDoS2019 veri seti üzerinde DDoS saldırı sınıflandırmasının yüksek doğruluk oranlarıyla yapılabileceğini ortaya koymaktadır. (3) *LSTM-CLOUD sistemi*: Önceki çalışmalardan farklı olarak, mevcut çalışmada genel bulut ağ

ortamının ağ trafiği karakterizasyonu için bir LSTM-CLOUD sistemi de tasarlanmış ve önerilmiştir.

Tez çalışmasının 2'inci Aşamasında IoT altyapısı için önerilen yaklaşımın literatürde mevcut diğer çalışmalar ile karşılaştırılması **Tablo 5.2'**de sunulmuştur.

Tablo 5.2: 2'inci Aşamanın İlgili Bazı Çalışmalar ile Karşılaştırılması

Araştırma	Yöntem	Veri kümesi	Eğitim Süresi (sn)	Test Süresi (sn)	Kesinlik
Mezgani ve Ktata [9]	CNN ve RNN	NSL-KDD	-	-	%99, %97
Suwannalai ve Polprasert [11]	RL	NLS-KDD	-	-	%79
Sethi ve ark. [12]	RL	NSL-KDD ve CICIDS2017	-	-	%96, %97,5
Aydın ve ark. [14]	LSTM	CICDoS-2019	11.90	3.68	%99,83
Alasmary ve ark. [15]	RNN/LSTM	CIC-IDS2017-Çar, CIC-IDS2017-Cum, CIC IoT 2022	-	-	%99,9, 1 ve %99,98
Alavizadeh ve ark. [16]	RL	NSL-KDD	-	-	%78
Butt ve ark. [17]	KNN, DT ve LSTM'nin hibridizasyonu.	CIC-IoT-2022 ve UNSW-NB15	-	-	%100, %99,9
Haider ve ark. [18]	CNN	CICIDS2017	39.52	0,061	%99
Protogerou ve diğerleri. [19]	GNN	Simüle edilmiş veriler	89.69	-	%99
Ebu Bakar ve ark. [20]	Makine öğrenimi	CICDoS-2019	-	-	%99,7
Bu çalışma	LSTM	CIC-IoT-2022	29.2	3.36	%99,48

Araştırma kapsamında literatürde mevcut olan ve incelenen çalışmaların bazı önemli sınırlılıklarının bulunduğu görülmektedir. Bu sınırlamalar şu şekilde kategorize edilebilir: 1) *Eğitim ve Test Süresi:* Bazı çalışmalarda yer alan modellerin eğitim ve test süreleri hakkında ayrıntılı bilgi yer almadığı görülmektedir. Buna karşılık bu çalışmada ve diğer birkaç çalışma da önerilen siber saldırı tespit modellerinin gerçek dünyadaki yeteneklerinin daha iyi anlaşılmasını sağlayan eğitim ve test süreleri yer almaktadır. 2) *Sınırlı Veri Kümesi Temsili:* İncelenen çalışmalar arasında yer alan belirli çalışmaların sadece NSL-KDD veri kümesine dayandığı görülmektedir. Ancak bu durum gerçek dünyadaki DDoS saldırı senaryolarını tam olarak temsil etmeyebilir. Her ne kadar NSL-KDD veri kümesi DDoS saldırılarını tespit etmek için kullanılıyor olsa bile esasen bu veri seti gerçek dünyadaki olaylar yerine simüle edilmiş veriler içermektedir. Bu nedenle bu veri seti bazı DDoS saldırı türlerini tam olarak temsil etmeyebilir ve modelin gerçek dünya koşullarındaki etkililiğinin doğru şekilde değerlendirilmesini zorlaştırabilir. Gerçek DDoS saldırılarının karmaşıklığını ve çeşitliliğini

tam olarak yakalayamadığı için daha gerçekçi ve çeşitli veri kümelerinin kullanılması, güvenilir sonuçlar elde etmek için şarttır. 3) *Metodolojik Çeşitlilik*: Araştırma kapsamında incelenen çalışmalar üç ana gruba ayrılabilir. Birinci gruptaki çalışmalar DÖ yöntemlerini içermektedir. İkinci gruptaki çalışmalar ise RL yaklaşımını kapsamaktadır. Son olarak üçüncü grupta yer alan çalışmalar ise hibrit yöntemleri içermektedir. Bu çalışma ise literatürde mevcut diğer çalışmalardan birkaç temel açıdan farklılık göstermektedir: 1) *Çoklu Ajan Tabanlı Sistem*: Bu çalışmada önerilen sistem, çoklu ajan yaklaşımı içinde LSTM tabanlı bir DL modelini içermektedir. Deneyler, bu modelin zaman karmaşıklığının, insan ağ savunucularının yeteneklerinden daha iyi performans göstererek DDoS saldırılarına karşı gerçek zamanlı siber savunmaya olanak sağladığını ortaya koymaktadır. Esasen çoklu ajan tabanlı siber saldırı tespiti yaklaşımı, ağ savunmasında insan müdahalesine olan ihtiyacı da azaltmaya yönelik bir yaklaşımdır. 2) *Paket Seviyesi Sınıflandırması*: Bu çalışmada DDoS saldırı sınıflandırmasında yüksek doğruluk elde etmek için CIC-IoT-2022 veri seti kullanılmıştır. 3) *Performans*: Bu çalışmadaki LSTM tabanlı model, farklı veya benzer DL algoritmaları ve veri kümeleri kullanan diğer çalışmalarla karşılaştırıldığında daha iyi bir performans sergilemiştir. Önceki bazı çalışmalardan farklı olarak bu çalışmada geliştirilen yaklaşımımız model tasarımında LSTM'nin gücüne odaklanmaktadır. Ek olarak çalışmamız, onu CIC-IoT-2022 veri kümesini kullanan diğerlerinden ayıran ajan tabanlı bir yaklaşım ortaya koymaktadır. 4) *Eğitim/Test Süreleri*: LSTM-IoT modelini geliştirirken her deney için eğitim ve test sürelerini hesaplanmıştır. Bu değerler diğer çalışmaların değerleriyle karşılaştırıldığında bu çalışmada geliştirdiğimiz modelimizin eşit etkinlikte performans gösterdiğini ortaya koymaktadır. Bu sınırlamalar giderildikten ve söz konusu iyileştirmeler de dahil edildikten sonra çalışmamızın DDoS tespiti ve azaltılması konusunda umut verici ve etkili bir yaklaşım sunduğu görülmektedir. Bu çalışmada geliştirilen özgün LSTM tabanlı çok ajanlı siber saldırı tespiti sisteminde CIC-IoT-2022 veri kümesinin kullanılması, modelin performansını artırmakta ve pratik uygulanabilirliğinin altını çizmektedir.

Tez çalışmasının 3'üncü Aşamasında bulut altyapısı için önerilen yaklaşımın literatürde mevcut diğer çalışmalar ile karşılaştırılması **Tablo 5.3**'de sunulmuştur.

Tablo 5.3: 3'üncü Aşamanın İlgili Bazı Çalışmalar ile Karşılaştırılması

Araştırma	Yöntem	Veri kümesi	Eğitim Süresi (sn)	Test Süresi (sn)	Kesinlik
Aydın ve diğ. (2021)	LSTM	CICDoS2019	11.90	3.68	%99,83
Novaes ve diğ. (2020)	LSTM-Bulanık	CICDoS2019	-	-	%97,89
Ebu Bakar ve diğ. (2023)	Makine öğrenimi	CICDoS2019	-	-	%99,7
Gaur ve diğ. (2022)	LSTM	CICDoS2019	330	-	%99,16
Subramanian ve diğ. (2022)	DNN, LSTM ve GRU	CICDoS2019	-	-	Sırasıyla %99,32, %99,4 ve %92,5.
Shurman ve diğ. (2020)	LSTM	CICDoS2019	-	-	%99,85
Kumar ve Behal (2020)	LSTM Çift Yönlü, LSTM, Yığılmış LSTM, GRU ve CNN	CICDoS2019	-	-	Sırasıyla 99,41, 99,43, 99,55 ve 99,52.
Sinhuja ve Suthendran (2022)	LSTM	BoT-IoT ve CICDoS2019	-	-	Sırasıyla 99.37 ve 98.47.
Zainudin ve diğ. (2022)	XGBoost ve hibrit bir CNN-LSTM modeli	CICDoS2019	0,124 ve 0,179.	-	Sırasıyla 99.57 ve 99.50.
Ahamed ve diğ. (2022)	YSA	IoT-23, DS2OS, NUSW-NB15GT ve CICDDOS2019.	Sırasıyla 521, 401, 117 ve 81.	-	Sırasıyla %99,07, %99,23, %9,71 ve 99,98(.
Bu çalışma	LSTM ile çoklu ajanlar	CICDoS2019	8.89	0,46	%99,89

Bu çalışmada incelenen çalışmalar iki ana kategoride toplanabilir. Birinci grupta imza tabanlı modellerin eğitim ve test sürelerini raporlayan çalışmalar yer almaktadır. Örneğin Aydın ve diğ. (2021) ve Gaur ve diğ. (2022) modellerinin eğitim ve test sürelerini belirtmiştir. Bu çalışma da eğitim ve test dönemlerini içermesi nedeniyle bu kategoriye girmektedir. İkinci grupta ise eğitim ve test sürelerinin belirtilmediği çalışmalar yer almaktadır. Bu duruma örnek olarak Novaes ve diğ. (2020) ile Subramanian ve diğ. (2022) çalışmaları verilebilir. Bu gruptaki zaman ölçümlerinin eksikliği, özellikle hızlı yanıtların kritik olduğu gerçek zamanlı DDoS saldırılarında, savunma modeli performansının kapsamlı ve karşılaştırmalı olarak değerlendirilmesini engellemektedir. Yüksek doğruluklu imza tabanlı bir DDoS saldırı sınıflandırma modeline sahip olmak, ağ güvenlik sistemleri için hayati öneme sahipken, bu

yeteneğin gerçek zamanlı DDoS olayları sırasında anında kullanılabilir olması da aynı derecede kritiktir. Gerçek zamanlı olaylarda DDoS saldırılarının hızlı tespiti çok önemlidir. Bu nedenle imza tabanlı izinsiz giriş tespit modelinin verimliliğini değerlendirirken eğitim/test sürelerini, karmaşıklığı ve hesaplama kaynaklarını da dikkate almak önemlidir.

Bu araştırma kendisini diğer araştırmalardan çeşitli yönlerden ayırmaktadır: 1) *Çok Ajanlı Tabanlı Sistem*: Bu araştırma, hacimsel DDoS saldırılarına karşı bulut güvenliği için çok ajanlı, çevik bir siber savunma stratejisi kullanırken LSTM'nin gücünden yararlanmaktadır. 2) *Paket Düzeyinde Bulut Trafiği Sınıflandırması*: Bu çalışmayı diğer çalışmalardan ayıran özelliklerden biri, özellikle paket düzeyinde hacimsel DDoS saldırılarının imza bazlı sınıflandırılması için veri kümesindeki hacimsel DDoS ağ trafiği verilerini kullanmasıdır. LSTM tabanlı saldırı sınıflandırma modeli geliştirilirken, veri kümesindeki tüm saldırı verilerinin kullanılması yerine yalnızca hacimsel DDoS ağ trafiği verilerinin kullanılmasındaki önemli fark, daha spesifik ve odaklı analiz yapılabilmesinde yatmaktadır. 3) *Eğitim/Test Süreleri*: Bu çalışma, imza tabanlı algoritmalarda yüksek doğruluğun ötesinde eğitim/test süreleri, karmaşıklık ve hesaplama kaynakları gibi performans metriklerinin dikkate alınmasının öneminin altını çizmektedir. Yüksek doğruluğun ötesinde performans ölçümlerinin dikkate alınmasına yapılan bu vurgu, imza tabanlı algoritmaların optimizasyonunu gerçek dünya uygulamalarında kaynak açısından daha verimli ve pratik hale getirmektedir.

Tez çalışmasının 4'üncü Aşamasında bulut altyapısı için önerilen yaklaşımın literatürde mevcut diğer çalışmalar ile karşılaştırılması **Tablo 5.4**'de sunulmuştur.

Tablo 5.4: 4'üncü Aşamanın İlgili Bazı Çalışmalar ile Karşılaştırılması

Araştırma	Yöntem	Veri kümesi	Eğitim Zaman (sn)	Test Zaman (sn)	Kesinlik
Alavizadeh ve ark. (2022)	DQN	NLS-KDD	-	-	%92,47
Li ve diğerleri. (2023)	Çift DQN Transferi	CICDoS2019	-	-	%98,5
Al- Fawa'reh ve ark. (2023)	DQN	MedBioT , N- BaIoT	12.6	8.4	%99,80 ve %99,40
Mishra ve ark. (2023)	Blockchain, Dağıtılmış RL	Kitsune Ağ Saldırısı Veri Kümesi	-	-	%94,5 ve %89,1
Sethi ve ark. (2021)	A-DQN	NL-KDD, CICIDS2017	-	-	%97,2 ve %98,7
Yungaicela -Naula ve ark. (2022)	DRL	Orta ölçekli veri merkezi ağı	-	-	%98

Suwannalai ve Polprasert (2020)	RL	NLS-KDD	-	-	%79
Sethi ve ark. (2020)	DRL	NSL-KDD, UNSW-NB15, AWID	-	-	% 81,80, % 96,12 ve %85,09
Gülmez ve Angın (2021)	DRL	NSL-KDD, UNSW-NB15	-	-	%97, %96
Feng ve diğerleri. (2020)	RL	Uygulama Katmanı DDoS Simülatörü	-	-	% 98.73
Hsu ve ark. (2020)	DRL	NSL-KDD, UNSW-NB15	-	-	% 91.4, % 91.8
Ahsan ve ark. (2022)	RL	NSL-KDD	-	-	% 86
Rookard ve ark. (2023)	DQN	TON-IoT	-	-	% 87.82
Malik ve Saini (2023)	DQN	NSL-KDD, CICIDS2017	-	-	% 97,9
Bouhamed ve ark. (2021)	DRL	CICIDS2017	-	-	% 99,7
Ren ve ark. (2022)	Q-Öğrenim	CSE-CIC-IDS2018	-	-	% 99.83
Mevcut çalışma	DQN	CICIoT2022, CICIoT2023	6.84, 3.78	0.14, 0.06	% 98.43, % 98.05

Literatürdeki diğer çalışmalarla karşılaştırıldığında bu araştırma birkaç açıdan öne çıkmaktadır:

(1) *Metodoloji*: Bu araştırma, IIoT ağlarında DDoS tespiti için bir DQN yaklaşımının kullanılmasına odaklanmaktadır. Hedef ağ, tekrarlama mekanizması, mini toplu eğitim ve model optimizasyonu gibi temel bileşenleri entegre eder. Üstelik DDoS saldırı sınıflandırmasını ödülleriyle zenginleştirerek MDP içerisinde stratejik bir tahmin oyunu olarak ele alıyor. Ayrıca bu çalışmada Nesnelerin İnterneti ile ilgili CICIoT2022 ve CICIoT2023 veri kümeleri kullanılmaktadır. (2) *Performans*: Önceki çalışmalarla karşılaştırıldığında, bu çalışma veri kümeleri üzerinde yüksek doğruluk oranlarına ulaşmaktadır. Ayrıca, eğitim ve test sürelerinin optimizasyonu, DDoS azaltma stratejilerinde gerçek zamanlı olay tespiti ve yanıt hızının kritik öneminin altını çiziyor. İncelenen çalışmaların çoğunda eğitim ve test süreleri genellikle raporlanmamaktadır. Ancak bu çalışmada diğerlerinden farklı olarak bu süreler titizlikle hesaplandı çünkü gerçek zamanlı olaylarda DDoS tespit hızı çok önemli. Eğitim ve test sürelerinin raporlanması, bir algoritmanın ve modelin gerçek dünya uygulamalarında ne kadar pratik olduğunu değerlendirmek için önemlidir. Bu süreler bir modelin ne kadar hızlı eğitildiğini ve ne kadar hızlı tahmin yapabildiğini gösterir. Bu, özellikle gerçek zamanlı uygulamalarda ve güvenlik açısından kritik sistemlerde çok önemlidir. Daha kısa eğitim süreleri, daha hızlı model yinelenmelerine ve güncellemelere olanak tanırken, daha kısa test süreleri, gerçek zamanlı uygulamalarda hızlı yanıt verilmesini sağlar. Bu nedenle, eğitim ve test

sürelerinin raporlanması, bir çalışmanın pratik uygulanabilirliği ve gerçek dünya performansı hakkında değerli bilgiler sağlar. Bu tür bilgiler, diğer araştırmacılara ve endüstri uygulayıcılarına çalışmamızı değerlendirmede ve gerektiğinde benzer yöntemleri uygulama konusunda yardımcı olabilir. (3) *Trafik Düzeyinde Sınıflandırma*: Bu araştırma, DDoS saldırısı algılama duyarlılığını ve ayrıntı düzeyini artırmak için trafik düzeyi sınıflandırmasına odaklanmaktadır. Aracı, DDoS saldırı sınıflandırmasını MDP içerisinde stratejik bir tahmin oyunu olarak ele alarak, normal ve kötü niyetli trafik kalıplarını ustalıkla ayırt eder; bu şekilde önceden tanımlanmış örüntü tanımayla ilgili sınırlamaların üstesinden gelinebilir. (4) *NS-IoT Sistemi*: Geleneksel siber savunma sistemlerinin aksine, NS-IoT sistemi otonom akıllı siber savunma araçlarına dayanır. IIoT ağları için özel olarak tasarlanan bu iş birliğine dayalı ve özerk çok aracıli mimari yaklaşımı, saldırılara hızlı ve uyarlanabilir otomatik yanıtlar sağlayarak insan müdahalesini azaltır ve sistem verimliliğini artırır.

5.1. Ajan Tabanlı olarak Siber Saldırı Tespitinde Otomatik Akıl Yürütme

Araştırmada önerilen dördüncü aşamada yer alan NS-IoT sisteminin uygulanmasında önemli bir unsur, siber savunma stratejilerinde karar verme yeteneklerini ve adaptasyon kabiliyetlerini artırmak için otomatik akıl yürütme mekanizmalarının oluşturulmasıdır. Otomatik akıl yürütme, siber güvenlikte kritik bir rol oynar; insan müdahalesini en aza indirir ve tehditlerin tespitini, analizini ve yanıtını hızlandırır. YZ ve MÖ (özellikle RL ve DQN) kullanılarak yapılan otomatik akıl yürütme, siber tehditleri hızlı bir şekilde tespit edip hafifletebilir. NS-IoT sisteminde, DQN-IoT ajanı RL kullanarak DDoS saldırılarını tespit eder ve IIoT veri setlerinden öğrenerek tehditleri daha etkili bir şekilde sınıflandırır ve analiz eder. Bu, gerçek zamanlı ve otomatik bir savunma sağlar ve IIoT güvenlik stratejilerinin genel doğruluğunu ve etkinliğini artırır. Siber tehditler gelişmeye devam ettikçe, otomatik akıl yürütme dinamik ve güçlü savunma mekanizmaları için önemini koruyacaktır.

NS-IoT sisteminde önerilen ajanların ayrıca şu gereksinimleri de karşılaması gerekliliği dikkate alınmalıdır:

- *Kapsamlı Entegrasyon*: Ajanlar, yerleştirilecekleri karmaşık veya otonom sistemlerin mühendislik sürecine entegre edilmelidir.
- *Özel Metodoloji*: Ajanların doğası, amacı, referans mimarisi ve operasyonel bağlamına uygun belirli bir metodolojiye göre tasarlanmalıdır.

- *Standartlara Uygunluk:* Ajanlar, uyumluluk, kalite ve güvenliği sağlamak için tanınmış standartlara göre inşa edilmelidir.

Ayrıca NS-IoT ajanlarının olanakları ve sınırlamaları dikkatlice ve kapsamlı bir şekilde test edilmelidir. Bu teknoloji güvenilir, güvenli ve etkili olduğunu garanti etmek için test veri setleri ve protokoller oluşturulmalıdır. Bu, şüpheli ve problemlili operasyon koşullarında bile geçerliliği sağlamayı amaçlamalıdır.

NS-IoT ajanlarının sadece karmaşık ağlara değil ancak aynı zamanda otonom sistemlere ve İnternet alanında dağılmış sensörler gibi oldukça basit cihazlara entegre edilebileceği de göz önünde bulundurulmalıdır. Bu durum, uzun süreler boyunca iletişim kurmadan devam edebilir ve ajanlar arasında, merkezi bir siber komuta merkeziyle veya insan operatörleriyle bağlantı kurulamayabilir. Algoritmalarını ve veri tabanlarını güncelleyememek, ajanların verimliliğinin hızla ve köklü bir şekilde azalmasına neden olabilir. Bu varsayım altında, ajanların yeni operasyon koşullarına uyum sağlamak için işlevlerini kendi kendine geliştirme yeteneğine sahip olmaları kritik öneme sahiptir. Ayrıca, kalite, güvenilirlik, bütünlük gibi özelliklerini garanti edebilmek için kendini güvence altına alma kapasitesine de ihtiyaç duyarlar. Bu tür kapasiteler, otonom öğrenme bağlamında da geçerli olacaktır. Bu süreç, ajanların veri tabanlarını yeni bilgilerle besleyerek, gerekirse yeni algoritmalar geliştirme ihtiyaçlarını ortaya koyacaktır.

5.2. Basitleştirilmiş Ödül Mekanizması

Ödül modelinin basitleştirilmesi, NS-IoT'de etkili öğrenme ve karar verme için kritik öneme sahiptir. NS-IoT sistemindeki DQN-IoT ajanı, ağ trafiği saldırı sınıflandırmasını bir MDP içindeki tahmin oyunu olarak çerçeveleyen bir ödül yapısı kullanır. Doğru sınıflandırmalar olumlu ödüller getirirken, yanlış sınıflandırmalar cezalara neden olur ve DDoS saldırılarının durum-uzay patlaması sorununu çözer. Ancak, bu ödül yapısının karmaşıklığı, modelin çeşitli DDoS varyasyonları ve IoT ortamlarında genelleme yapma yeteneğini kısıtlayabilir. Bu çalışmada, ödül mekanizmasını basitleştirerek genelleştirmeyi artırmayı ve eğitim sürecini optimize etmeyi hedefledik. Bu sadeleştirilmiş yaklaşım, eğitim verimliliğini artırır, hata ayıklamayı kolaylaştırır ve modelin gerçek dünya uygulamalarına uyumunu geliştirir, nihayetinde uygulama ve bakım maliyetlerini azaltır.

5.3. Ajan Tabanlı olarak Siber Saldırı Tespiti ve İnsan Ağ Operatörleri

İnsan ağ operatörlerinin entegrasyonu, NS-IoT ve IIoT sistemlerinin güvenliğini sağlamak için son derece önemlidir. İnsan operatörlerinin IIoT siber savunma döngüsüne entegrasyonu, NS-IoT sisteminin otonom yeteneklerini tamamlamak ve karmaşık siber savunma senaryolarında etkili gözetim ve yanıt sağlamak için kritik bir rol oynar. DQN-IoT ajanlarının otomasyonuna rağmen, insan müdahalesi, karmaşık siber tehditlere hızlı ve doğru bir yanıt vermek için gereklidir. İnsan müdahalesi, kritik bağlamsal anlayış, stratejik içgörü ve gelişmiş DDoS saldırılarına karşı otomatik savunmalar sağlama konusunda denetim sunar. Ağ operatörleri, sistemin izlenmesi ve sürdürülmesi, evrilen tehditlere uyum sağlanması ve etkili iletişim ile sistem bütünlüğünün korunması açısından hayati öneme sahiptir. Ayrıca, insan operatörleri, DQN-IoT ajanlarının kalibrasyonunda ve gelişiminde, hedeflerle ve etik standartlarla uyumlu hale getirilmesinde önemli rol oynar. Otonom sistemlerle insan operatörleri arasındaki bu iş birliği, IIoT siber savunma stratejilerinin genel etkinliğini ve dayanıklılığını artırmak için gereklidir.

5.4. Ajanların Bireysel Karar Verme Süreci

NS-IoT ajanlarının güvenilirliği, karar verme süreçlerine bağlıdır. Esasen bu ajanların karar verme süreci bir dizi işlevden oluşur: algılama/veri toplama, durum farkındalığı, eylem planlama, eylem seçimi ve eylem etkinleştirme. Çeşitli taktiksel ve teknik durumlarla başa çıkmak ve mümkün olduğunca kontrol edilemeyen felaket sonuçlarını engellemek için, NS-IoT ajanlarının insan karar verme süreçlerini taklit eden bireysel karar verme yollarına ihtiyaçları vardır. Bu amaçla, yalnızca kararlar önermek veya bir uyarıcıya tepki vermek için MÖ algoritmalarına ihtiyaç duymayacaklar, aynı zamanda insan bilişimini andıran, karmaşık durumlara sürekli olarak uyum sağlayacak kadar esnek bir “derin karar verme” sürecine de ihtiyaç duyacaklardır. Derin karar verme sürecinin işlevleri, AI teknikleri kullanılarak uygulanacaktır. Derin karar verme süreçleri ve işlevleri, bilişsel bilimlerin yeni ve karmaşık bir odak noktasıdır. Ayrıca, alınan kararlar ve karar verme süreci insan karar vericilere açıklanabilir olmalıdır.

5.5. Kolektif Zekâ ve Karar Verme

NS-IoT ajanlarının tekil ajanlara üstünlük sağlaması için, sürüler veya topluluklar halinde hareket etmeleri gerekmektedir. Bu, kolektif kararlar alarak tekil ajarlardan daha etkili sonuçlar

elde etmelerini sağlar. Bunun için, ajanların veri paylaşımı ve/veya veri alışverişi yapmaları gerekecektir. Birbirlerine yardımcı olmaları, hesaplama, veri ve eylem yükünü paylaşmaları da muhtemelen gerekecektir. Ayrıca, güven sorunuyla ilişkili olarak, bu konu kolektif karar verme ve zekâ sürecinin bir parçası olmalıdır.

NS-IoT ajanları, kötü amaçlı yazılımları ve kendi davranışlarını ve eylemlerini öğrenmeye devam etmelidirler. Öğrenme iki şekilde gerçekleştirilebilir: çevrimiçi (online) ve çevrimdışı (offline). Çevrimiçi öğrenme büyük hesaplama ve depolama kapasiteleri gerektirirken, çevrimdışı öğrenme büyük bilgisayarlarda yapılabilir ve sonuçlar ajanlara, kötü amaçlı yazılımlarla savaşmadan önce yüklenebilir. Her iki seçenek de mevcut olmalıdır ki bu durum ajanların ölçeklenmesini de etkileyecektir. Aksi takdirde, ölçekleme öğrenme sürecinin çevrimiçi veya çevrimdışı olma seçimini sınırlandıracaktır. Ajanlar küçük boyutlu olduğunda, çevrimdışı öğrenme tercih edilebilir. Çevrimdışı öğrenme ve savaşmadan önce yükleme, ajanların karar verme yeteneklerini basit kararlarla sınırlandırabilir. Ayrıca, ajanlar uzun süreli olarak ana sistemlere entegre edilecekse ve düzenli güncellemeler veya öğrenme olmadan verimlilikleri azalabilecekse, bu ajanların ömrünü sorun haline getirebilir.

5.6. Ajanların Diğer Varlıklarla İş birliği

NS-IoT ajanlarının diğer ajanlar, bir siber merkez veya insan operatörleri ile işbirliği yapmaları gerekebilir. Bu iş birliği, güven sorunlarının yanı sıra, iş birliği koşulları, ihtiyaçlar, işlevsellikler, protokoller, veri akışları ve bilgi doygunluğu, ergonomi, güvenlik ve aksaklık durumlarında süreklilik gibi birçok önemli konuyu içerir. Bu araştırma alanı, ajanların diğer varlıklarla etkili bir şekilde iş birliği yapabilmesi için gerekli olan tüm bu koşulları ve süreçleri ele almalıdır.

5.7. Ajanların Gizliliği ve Dayanıklılığı

NS-IoT ajanları, özellikle otonom akıllı kötü yazılım sınıfından olanlar ile düşman kötü yazılımlarının birincil hedefi haline gelecektir. Bu nedenle, NS-IoT ajanlarının düşman saldırılarından korunması ve savunulması kritik önem taşır. Siber tehdit ve risk analizleri bu süreçte önemli bir rol oynayabilir ve gerektiğinde veya düzenli olarak ajanlar tarafından gerçekleştirilmesi mümkün olabilir. Ajanların, düşman kötü yazılımını saldırıya geçmekten caydırma kapasitesine sahip olması düşünülebilir. Ajanların koruma mekanizmalarının içinde gizlilik (stealth) sağlama ve saldırılara karşı dayanıklılık (robustness) geliştirme gibi özellikler

bulunmalıdır. Ayrıca, ajanlar arası şifreleme, ajanlar arasında güvenli veri iletişimi sağlamak için hayati öneme sahiptir. Bu şifreleme, ajanların iletişimlerinin ve verilerinin yetkisiz erişimlerden korunmasına yardımcı olur, böylece düşmanların ajanlar arasında veri sızdırmasını veya manipülasyonunu önler. NS-IoT ajanları, saldırılar gerçekleştiğinde dayanıklı olmalı ve bu saldırılara karşı etkili bir şekilde savunma yapabilmelidir. Dolayısıyla NS-IoT ajanları bireysel veya grup olarak hedef alınan saldırılara karşı geri savunma yapabilme kapasitesine sahip olmalıdırlar. Ajanlar, saldırılardan öğrenerek bireysel ve kolektif kapasitelerini güçlendirmeli ve gelecekteki saldırılara karşı daha iyi bir savunma mekanizması geliştirmelidirler. Dayanıklılık, tehdit ve risk analizi, caydırıcılık, koruma, izleme ve tespit, saldırı yanıtı, öğrenme ve iyileştirme gibi altı mekanizma, ajanların genel dayanıklılığını sağlamalıdır. Bu mekanizmalar, ajanların iletişimini de kapsayacak şekilde genişletilmelidir.

Ayrıca dost ve düşman ajanlar ile karşılaştığında ya da bilinmeyen bir ajan ile karşı karşıya kalındığında NS-IoT ajanlarının iyi ajanları kötü ajanlardan ayırt edebilme yeteneğine sahip olmaları gerekmektedir. NS-IoT ajanlarının kümeleri (sürüleri) bu şekilde sızmalara karşı kendilerini savunabilirler. Yeni bir üye kabul edilirken iş birliği bağlantılarını yeniden yapılandırabilirler. Ayrıca, dostane ajanların birbirlerine güvenebilmeleri gerekir. Güvenilmez ajanlar ise yasaklanmalı veya imha edilmelidir ve bu stratejiler ajanların görevleri, öncelikli çıkarları ve operasyon kuralları ile uyumlu olmalıdır. Güvenilmez ajanlar diğer ajanlara da bildirilmelidir, böylece diğer ajanlar bu ajanlarla iş birliği yapmaktan kaçınabilirler. Bu sosyal dinamikler, ajanların aralarındaki güveni ve iş birliği süreçlerini yönetmek için kritik bir rol oynar. Ajanlar arasındaki etkileşimlerde güven mekanizmaları oluşturmak, siber güvenlik operasyonlarının etkinliğini artırır ve olası tehditlere karşı koruma sağlar.

6. SONUÇ VE ÖNERİLER

Gelecekteki bilgisayar ağları giderek daha karmaşık, dinamik, YZ destekli ve makine hızında hatta otonom siber saldırılarla karşılaşacaktır. Bu bağlamda, bu tür gelişmiş saldırılara karşı etkili bir savunma sağlamak için otonom siber savunma sistemleri, otomatik akıl yürütme ve yanıt mekanizmalarıyla donatılmalıdır. Bilgisayar ağlarını hedef alan bu tür siber saldırıların tespitinde, akıllı ajan teknolojilerinden yararlanmak umut verici bir yaklaşım olarak öne çıkmaktadır. Bu tez çalışmasında, ağ güvenliği yönetimi için akıllı ajan teknolojisinin kullanıldığı saldırı tespitine yönelik yeni bir yöntem önerilmiştir. Araştırma, dört aşamalı bir süreç uygulayarak farklı saldırı tespit yöntemlerini değerlendirmiş ve her aşamanın güçlü ve zayıf yönlerini detaylı bir şekilde analiz etmiştir.

Araştırmanın ilk aşamasında, LSTM ağı kullanılarak imza tabanlı bir saldırı tespit modeli geliştirilmiştir. Bu model, CICDDoS2019 veri seti üzerinde %99.83 doğruluk oranı elde ederek yüksek performans göstermiştir. LSTM modelinin siber saldırıları yüksek doğrulukla tespit edebilme yeteneği, etkili bir savunma mekanizması sağladığını ortaya koymaktadır. Ancak, bu aşamanın imza tabanlı tespit yöntemleri nedeniyle bazı sınırlamaları ve değişken saldırı türlerine karşı esneklik sağlayamama gibi dezavantajları bulunmaktadır.

Araştırmanın ikinci aşamasında, IoT ortamında siber saldırıları tespit etmek için LSTM tabanlı bir ajan sistemi geliştirilmiştir. Bu sistem, çok ajanlı bir yapıyla entegre edilmiş olup CIC-IoT-2022 veri seti üzerinde %99.48 doğruluk oranı elde etmiştir. Bu aşama, IoT ortamlarında minimum insan müdahalesi ile etkili bir saldırı tespiti sağlamada başarılı olmuştur, ancak sistemin çevikliği ve adaptasyon yeteneği sınırlı kalmıştır.

Araştırmanın üçüncü aşamasında, bulut ortamındaki hacimsel DDoS saldırılarına karşı LSTM tabanlı bir ajan sistemi geliştirilmiştir. Model, hacimsel veri kullanarak ve çok ajanlı sistemlerle imza tabanlı tespit yöntemlerini birleştirerek %99.89 doğruluk oranı elde etmiştir. Bu aşama, bulut ortamındaki saldırılara zamanında ve işbirlikçi bir yanıt verme kapasitesini göstermiştir, ancak veri hacmi ve işlem süresi açısından optimizasyon gerektirmiştir.

Araştırmanın dördüncü ve son aşamasında ise, IIoT ortamlarına yönelik olarak ajan tabanlı DRL yaklaşımı kullanılarak bir siber saldırı tespit sistemi geliştirilmiştir. DQN ve çok ajanlı sistemler kullanılarak oluşturulan bu model, CIC-IoT-2022 ve CIC-IoT-2023 veri setlerinde sırasıyla %98.43 ve %98.05 doğruluk oranları sağlamıştır. Bu yaklaşım, yüksek otonomi ve çeviklik sunarak endüstriyel IoT ortamındaki saldırılara karşı etkili bir savunma sağlama potansiyeline sahiptir. Bu sistem, saldırı tespit ve müdahale süreçlerinde büyük hız ve esneklik sunarak, güvenlik önlemlerini güçlendirmeye olanak tanımaktadır.

Araştırmada elde edilen bilgiler kapsamında önerilen dört aşamalı yöntemler arasında en etkili yaklaşım dördüncü aşama olarak belirlenmiştir:

- **Otonomi ve En Az İnsan Müdahalesi:** Dördüncü aşama, insan müdahalesini en aza indirerek tamamen otonom bir savunma sağlar. Otonom ajanlar, ağdaki saldırıları tespit etmek ve müdahale etmek için sürekli olarak çalışır, bu da insanların sürekli olarak sistem üzerinde kontrol sağlamasına gerek kalmadan etkili bir güvenlik yönetimi sunar. Bu hem insan hatalarını azaltır hem de sürekli ve kesintisiz bir koruma sağlar.
- **Etkili Koordinasyon ve Bilinmeyen Saldırlara Yanıt:** DQN-IoT, CR-IoT ve CM-IoT ajanları arasındaki koordinasyon, savunma süreçlerini entegre eder. DQN-IoT ajanı, DDoS gibi bilinen saldırıları tespit ederken; CR-IoT, tespit edilen anormalliklerin etkilerini hafifletir; CM-IoT ise ajanlar arasındaki iletişimi koordine eder. Ayrıca, DRL tabanlı yaklaşım, tamamen bilinmeyen siber saldırıları tanıma yeteneğine sahiptir. DRL algoritmaları, daha önce karşılaşılmamış tehditleri öğrenerek ve bu tehditlere karşı stratejiler geliştirerek dinamik ve bilinmeyen saldırılara karşı etkili savunmalar sunar.
- **Yüksek Doğruluk:** Dördüncü aşama, Derin Pekiştirmeli Öğrenme (DRL) ve çok ajanlı sistemleri birleştirerek saldırıları yüksek doğrulukla tespit eder. DRL, ajanların çevresindeki verileri işleyip öğrenmelerine olanak tanırken, çok ajanlı sistemler bu bilgileri birden fazla ajan arasında paylaşarak analiz eder. Bu birleşim, hem doğru hem de güvenilir tespit sonuçları sağlar, yanlış pozitif ve yanlış negatif oranlarını azaltır.
- **Hızlı Tepki:** Bu yöntem, saldırılara hızlı bir şekilde yanıt vererek ağın zarar görmesini önler. Çok ajanlı sistemler, aynı anda birden fazla tehdidi değerlendirebilir ve etkili bir şekilde yanıt verebilir. DRL algoritmalarının sunduğu öğrenme yeteneği sayesinde,

ajanlar saldırıları hızlı bir şekilde tanır ve uygun savunma stratejilerini uygulayarak hızla tepki verir. Bu, ağın güvenliğini anında korur ve saldırıların etkisini en aza indirir.

- **Dinamik ve Ölçeklenebilir Savunma:** DRL tabanlı yaklaşım, ajanların çevreleriyle etkileşimlerinden öğrenmelerini sağlar ve bu sayede dinamik ortamlarda etkili bir şekilde çalışabilir. Bu teknoloji, sistemin çevresindeki değişikliklere uyum sağlayabilir ve ölçeklenebilir bir savunma sunar. Yani, ağda meydana gelen yeni tehditler veya değişen koşullara göre savunma stratejilerini sürekli olarak optimize edebilir, bu da sistemin uzun vadeli güvenliğini artırır.

En etkili yaklaşım olarak belirlenen dördüncü aşama, DRL ve çok ajanlı sistemlerin birleşimini kullanarak yüksek doğruluk ve hızlı tepki sağlamakta en başarılı sonuçları vermiştir. Ağ güvenliği yönetimi için akıllı ajan teknolojisi kullanarak saldırı tespitine yönelik olarak önerilen bu yöntem, RL hassasiyetini çoklu ajan sistemlerinin çevikliğiyle birleştirerek etkili bir savunma mekanizması oluşturma potansiyeline sahiptir. Bu tez çalışmasında geliştirilen NS-IoT sistemi, bilgisayar ağları için tasarlanmış işbirlikçi, otonom çoklu ajan DQN tabanlı bir siber savunma sistemidir ve DDoS saldırılarıyla mücadele etmeyi amaçlamaktadır. DQN teknolojisini kullanan otonom ajanlar arasındaki iş birliği, NS-IoT sisteminin proaktif tehdit tespiti, hızlı karar alma ve etkili yanıt yürütme yeteneklerini artırmayı hedeflemektedir. Bu entegre yaklaşım, bilgisayar ağlarının zeki ve ileri düzey saldırılara karşı dayanıklılığını artırırken, minimal insan müdahalesi ile ağ sistemlerinin bütünsel doğasını koruyabilecek otonom siber savunma sistemleri için bir temel oluşturmaktadır. Ayrıca, NS-IoT sistemi, DQN-IoT, CR-IoT ve CM-IoT ajanlarından oluşan tespit ve savunma modüllerine sahiptir. DQN-IoT ajanı, DDoS saldırılarını DRL ve Markov karar süreçlerini kullanarak tespit ederken; CR-IoT ajanı, tespit edilen anormalliklerin etkilerini hafifletmek için savunma modülünde çalışmakta; CM-IoT ajanı ise sistem ajanları arasındaki iletişimi koordine etmektedir. IoT ile ilgili CIC-IoT-2022 ve CIC-IoT-2023 veri setleriyle yapılan deneysel değerlendirmeler, DQN-IoT ajanının DDoS saldırılarını tespit etme konusundaki etkinliğini ortaya koymuştur. Çalışmada kullanılan DRL tabanlı yaklaşım, DQN-IoT ajanının çevresiyle etkileşimlerinden öğrenerek DDoS saldırılarını tespit etme yeteneğini geliştirmekte ve dinamik ortamlarda ölçeklenebilirlik ve gerçek zamanlı savunma optimizasyonu sunmaktadır. Sonuç olarak, NS-IoT sistemi, minimal insan müdahalesi ile gerçek zamanlı DDoS saldırılarını tespit etme ve hafifletme yeteneğine sahip çevik ve otonom bir IIoT güvenlik çözümü sunmaktadır. Ayrıca, bu çalışma

ajan mühendisliği ve sertifikasyonu gibi diğer araştırma ve teknoloji akışları için de bir temel oluşturabilir.

Araştırmamızda elde edilen sonuçlarla bağlantılı olarak aşağıdaki konularda öneriler sunmak olanaklıdır:

- **Akıllı Ajanlar ve YZ Entegrasyonu:** Akıllı ajanların saldırı tespit ve ağ güvenliği yönetiminde etkinliğini artırmak için YZ tekniklerinin entegrasyonu sağlanmalıdır. Özellikle, YZ algoritmalarının akıllı ajanların davranış ve karar verme süreçlerine dahil edilmesi, veri analizi ve tehdit tanımlama yeteneklerini önemli ölçüde geliştirebilir. YZ tabanlı teknikler, ajanların anormal ağ trafiği ve sofistike saldırı desenlerini daha hızlı ve doğru bir şekilde tanımlamasına olanak tanır. Ayrıca, MÖ ve DÖ yöntemleri kullanılarak ajanların sürekli öğrenme ve adaptasyon yetenekleri artırılabilir. Böylece yeni ve evrimleşen tehditlere karşı hızlı, çevik ve otonom olarak daha etkili bir savunma sağlanabilir.
- **Akıllı Ajanları Gerçek Zamanlı Veri İşleme Kapasiteleri:** Akıllı ajanların gerçek zamanlı veri analizi yapabilme kapasiteleri, düşük gecikme süreleri ve yüksek işlem gücü sağlayacak şekilde optimize edilmelidir. Böylelikle ajanların siber saldırıları makine hızında tespit edebilmeleri ve yanıt sürelerinin kısaltılması sağlanabilir.
- **Akıllı Ajanlar ile Çok Katmanlı Güvenlik Entegrasyonu:** Akıllı ajanlar, ağ güvenliğinin bir bileşeni olarak sadece IDS veya IDPS ile değil ancak diğer güvenlik çözümleriyle de (örneğin firewall vb.) ayrıca entegre edilmelidir. Bu entegre yapı, çok katmanlı bir güvenlik sağlayarak saldırı tespit oranını artırabilir.
- **Akıllı Ajanlar ve Kullanıcı Davranış Analizi Özellikleri:** Akıllı ajanlar, ağ üzerindeki kullanıcı davranışlarını analiz edebilen özelliklerle donatılmalıdır. Bu, olağandışı kullanıcı aktivitelerini tespit ederek potansiyel tehditlerin daha iyi değerlendirilmesine yardımcı olabilir.

Gelecekteki araştırmalar çerçevesinde, bu Tez çalışmasında ağ güvenliği yönetimi için akıllı ajanlar teknolojisi kullanılarak saldırı tespitine yönelik yeni bir yaklaşım olarak önerilen NS-IoT sisteminin, gerçek dünya altyapısını temsil eden gerçek bir bilgisayar ağı ortamında uygulanması, test edilmesi ve elde edilen test sonuçlarının analiz edilerek sisteme yeni

özellikler eklenmesi planlanmaktadır. Bu yaklaşım, mevcut araştırmanın kapsamını genişleterek, sistemin performansını, uygulama pratikliğini ve dayanıklılığını derinlemesine değerlendirmeyi hedeflemektedir. Önerilen sistemin değerlendirilmesi sürecinde, canlı ve gerçek bir bilgisayar ağı ortamında gerçekleştirilecek testler, çeşitli performans kriterlerine göre sistemin yeteneklerini ölçmeye olanak tanıyacaktır. Bu test ortamında, sistemin çeşitli senaryolar altında karşılaşılabileceği potansiyel sorunlar, performans metrikleri ve kullanıcı geri bildirimleri sistematik bir şekilde toplanacak ve belgelenerek kapsamlı raporlar hazırlanacaktır. Test sonuçlarının analizi, sistemin mevcut zayıflıklarını ve performans sınırlamalarını belirleyecek ve bu doğrultuda sistem üzerinde gerekli iyileştirme stratejileri geliştirilecektir. İyileştirme sürecinde, performans artırma ve güvenlik açıklarını kapatma amacıyla yenilikçi çözümler ve teknikler kullanılacaktır. Gelecekteki araştırmalar çerçevesinde son olarak, geliştirilmiş sistemin performans standartlarını karşılayacak şekilde yeniden test edilmesi sağlanacaktır. Bu aşamada, sistemin performans verileri ve güvenlik özellikleri güncellenmiş standartlara göre değerlendirilecek ve gereken son iyileştirmeler yapılacaktır. Bu süreç, önerilen NS-IoT sisteminin ağ güvenliği ve saldırı tespiti konularındaki etkinliğini artırarak, daha sağlam ve güvenilir bir çözüm sunmasını sağlayacaktır.

Aşağıda önerilen liste akıllı ajanlar ve siber savunma konularında araştırma ve geliştirme için önerilen konuları kapsamaktadır. Her bir madde, akıllı ajanların performansını, güvenliğini, etkileşimini ve uygulama alanlarını incelemeye yönelik potansiyel çalışma alanlarını temsil eder. Liste hem mevcut konuları hem de yeni önerileri içerir ve bu alanlarda derinlemesine araştırma yapılmasını teşvik etmek amacıyla hazırlanmıştır.

- **Kötü Amaçlı Otonom Yazılımlara Karşı Savunma:** Akıllı ajanlar, otonom ve akıllı kötü amaçlı yazılımlarla, hatta bunlar sürü halinde çalışıyorsa, nasıl etkileşime geçebilirler?
- **Ajan İş birliği:** Akıllı ajanlar, siber savaş alanındaki anlayışlarını artırmak ve dayanıklılık ile yeteneklerini geliştirmek için nasıl iş birliği yapmalıdır?
- **Siber Saldırıları:** Akıllı ajanlar ve çoklu ajan sistemleri hangi saldırılara maruz kalabilir?
- **Saldırıları Karşı Dayanıklılık:** Akıllı ajanlar, saldırıları nasıl öngörüp engelleyebilir ve bu saldırılara karşı nasıl dayanıklı olabilirler?

- **Sınırlı Kaynaklarla Savunma:** Akıllı ajanlar, sınırlı bellek ve işlem kaynakları ile nasıl etkili bir şekilde öğrenebilir ve savunma yapabilir?
- **Akıllı Ajan Tabanlı Siber Savunma Modelleri:** Akıllı ajanlar için uygun modeller neler olabilir?
- **Dinamik Savunma Amaçlı Ekip Oluşumu:** Akıllı ajanlar ile dinamik olarak heterojen ekipler nasıl oluşturulabilir ve bunlarla nasıl etkileşime geçilebilir?
- **Siber Savunma Ajanı - İnsan Etkileşimi:** Akıllı ajan sistemleri insanlar ile nasıl
- **Adaptif Öğrenme:** Siber savunma amaçlı akıllı ajanların çevresel değişikliklere ve yeni tehditlere karşı adaptif öğrenme yetenekleri nasıl geliştirilebilir?
- **Veri Mahremiyeti ve Güvenlik:** Akıllı ajanların veri güvenliğini sağlamak için hangi veri şifreleme ve anonimleştirme teknikleri kullanılabilir?
- **Güvenlik Amaçlı Karar Destek Sistemleri:** Akıllı ajanlar, insan karar vericilere nasıl etkili destek sağlayabilir ve bu süreçte nasıl entegre olabilirler?
- **Karmaşık Ağların Güvenlik Yönetimi:** Akıllı ajanlar, karmaşık ağ yapılarındaki performansı ve güvenliği nasıl yönetebilir?
- **Siber Tehdit Analizi ve Tahmin:** Akıllı ajanlar, gelecekteki siber tehditleri nasıl tahmin edebilir ve bu tehditlere karşı nasıl hazırlıklı olabilirler?

KAYNAKLAR

- Abdul-Qawy, A.S., Pramod, P.J., Magesh, E. & Srinivasulu, T., 2015. The internet of things (IoT): An overview. *International Journal of Engineering Research and Applications*, 5(12), pp.71-82.
- Abramson, D., Buyya, R. & Giddy, J., 2002. A computational economy for grid computing and its implementation in the Nimrod-G resource broker. *Future Generation Computer Systems*, 18(8), pp.1061-1074.
- Affinito, A., Zinno, S., Stanco, G., Botta, A. & Ventre, G., 2023. The evolution of Mirai botnet scans over a six-year period. *Journal of Information Security and Applications*, 79, 103629.
- Ahamed, I., Ahamad, F., Palade, V. & Ahamed, A., 2022. Neural network-based distributed denial of service (DDoS) attack detection in smart home networks. In: *6th Smart Cities Symposium (SCS 2022)*, December 2022, pp.174-179. IET.
- Alasmary, F., Alraddadi, S., Al-Ahmadi, S. & Al-Muhtadi, J., 2022. ShieldRNN: A Distributed Flow-Based DDoS Detection Solution for IoT Using Sequence Majority Voting. *IEEE Access*, 10, pp.88263-88275. doi: 10.1109/ACCESS.2022.3200477.
- Alavizadeh, H., Alavizadeh, H. & Jang-Jaccard, J., 2022. Deep Q-learning based reinforcement learning approach for network intrusion detection. *Computers*, 11(3), p.41. doi: 10.3390/computers11030041.
- Al-Fawa'reh, M., Abu-Khalaf, J., Szewczyk, P. & Kang, J.J., 2023. MalBoT-DRL: Malware Botnet Detection Using Deep Reinforcement Learning in IoT Networks. *IEEE Internet of Things Journal*.
- Alguliyev, R.M., Aliguliyev, R.M. & Abdullayeva, F.J., 2019. The improved LSTM and CNN Models for DDoS attacks prediction in social media. *International Journal of Cyber Warfare and Terrorism (IJCWT)*, 9(1), pp.1-18.
- Alweshah, M., Hammouri, A., Alkhalaileh, S. et al., 2023. Intrusion detection for the internet of things (IoT) based on the emperor penguin colony optimization algorithm. *Journal of Ambient Intelligence and Humanized Computing*, 14, pp.6349-6366. doi: 10.1007/s12652-022-04407-6.
- Alzubi, O.A., Alzubi, J.A., Alzubi, T.M. et al., 2023. Quantum Mayfly Optimization with Encoder-Decoder Driven LSTM Networks for Malware Detection and Classification Model. *Mobile Networks and Applications*. doi: 10.1007/s11036-023-02105-x.

- Antonakakis, M., April, T., Bailey, M., Bernhard, M., Bursztein, E., Cochran, J. & Zhou, Y., 2017. Understanding the Mirai botnet. In: *26th USENIX Security Symposium (USENIX Security 17)*, pp.1093-1110.
- Ashoor, A.S. & Gore, S., 2011. Importance of intrusion detection system (IDS). *International Journal of Scientific and Engineering Research*, 2(1), pp.1-4.
- Aydin, H., Orman, Z. & Aydin, M.A., 2022. A long short-term memory (LSTM)-based distributed denial of service (DDoS) detection and defense system design in public cloud network environment. *Computers & Security*, 118, 102725. doi: 10.1016/j.cose.2022.102725.
- Balogh, Z., Gatia, E., Hluchý, L., Toegl, R., Pirker, M. & Hein, D., 2014. Agent-based cloud resource management for secure cloud infrastructures. *Computing and Informatics*, 33(6), pp.1333-1355.
- Benlloch-Caballero, P., Wang, Q. & Calero, J.M.A., 2023. Distributed dual-layer autonomous closed loops for self-protection of 5G/6G IoT networks from distributed denial of service attacks. *Computer Networks*, 222, 109526. doi: 10.1016/j.comnet.2022.109526.
- Biermann, E., Cloete, E. & Venter, L.M., 2001. A comparison of intrusion detection systems. *Computers & Security*, 20(8), pp.676-683.
- Biggio, B., Corona, I., Fumera, G., Giacinto, G. & Roli, F., 2011. Bagging classifiers for fighting poisoning attacks in adversarial classification tasks. In: *Multiple Classifier Systems: 10th International Workshop, MCS 2011, Naples, Italy, June 15-17, 2011. Proceedings 10*, pp.350-359. Springer Berlin Heidelberg.
- Booker, L.B. & Musman, S.A., 2020. A model-based, decision-theoretic perspective on automated cyber response. *arXiv preprint arXiv:2002.08957*.
- Boss, G., Malladi, P., Quan, D., Legregni, L. & Hall, H., 2007. Cloud computing. *IBM white paper*, 321, pp.224-231.
- Boudaoud, K., Labiod, H., Boutaba, R. & Guessoum, Z., 2000. Network security management with intelligent agents. In: *NOMS 2000. 2000 IEEE/IFIP Network Operations and Management Symposium 'The Networked Planet: Management Beyond 2000'*, pp.579-592. IEEE.
- Bouhamed, O., Bouachir, O., Aloqaily, M. & Al Ridhawi, I., 2021. Lightweight IDS for UAV networks: A periodic deep reinforcement learning-based approach. In: *2021 IFIP/IEEE International Symposium on Integrated Network Management (IM)*, pp.1032-1037. IEEE.
- Boyd, B.L., 2009. Cyber warfare: Armageddon in a Teacup? *Army Command and General Staff College Fort Leavenworth KS*.
- Boyes, H., Hallaq, B., Cunningham, J. & Watson, T., 2018. The industrial internet of things (IIoT): An analysis framework. *Computers in Industry*, 101, pp.1-12.

Brownlee, J., 2017. *Introduction to Time Series Forecasting with Python: How to Prepare Data and Develop Models to Predict the Future*. Machine Learning Mastery.

Butt, N., Shahid, A., Qureshi, K.N., Haider, S., Ibrahim, A.O., Binzagr, F. & Arshad, N., 2022. Intelligent deep learning for anomaly-based intrusion detection in IoT smart home networks. *Mathematics*, 10(23), 4598. doi: 10.3390/math10234598.

Buyya, R., Broberg, J. & Goscinski, A.M. (Eds.), 2010. *Cloud computing: Principles and paradigms*. John Wiley & Sons.

Chaudhary, S. & Mishra, P.K., 2023. DDoS attacks in Industrial IoT: A survey. *Computer Networks*, 110015.

Chetty, M. & Buyya, R., 2002. Weaving computational grids: how analogous are they with electrical grids?. *Computing in Science & Engineering*, 4(4), pp.61-71.

Çil, A.E., Yildiz, K. & Buldu, A., 2021. Detection of DDoS attacks with feed forward based deep neural network model. *Expert Systems with Applications*, 169, 114520.

Cloudflare, 2023. DDoS threat report for 2023 Q2. Available at: <https://blog.cloudflare.com/ddos-threat-report-2023-q2/> (Accessed: 27 August 2024).

Cisco, 2018. Cisco Annual Internet Report, 2018–2023. Available at: <https://www.cisco.com/c/en/us/solutions/collateral/executive-perspectives/annual-internet-report/white-paper-c11-741490.html> (27 Ağustos 2024).

Darwish, M., Ouda, A. & Capretz, L.F., 2013. Cloud-based DDoS attacks and defenses. In: *International Conference on Information Society (i-Society 2013)*, pp.159-163. IEEE.

Dufresne, C., Soler, P., Checinski, G. & Valois, T., 2023. A survey of deep learning-based intrusion detection systems in industrial environments. *Computers*, 12(2), p.17. doi: 10.3390/computers12020017.

Elia, V., Rak, J., Giacomini, M. & Weiss, M., 2022. A survey of machine learning techniques for cybersecurity in smart cities. *ACM Computing Surveys (CSUR)*, 55(11), pp.1-37.

Etzioni, O., 2018. The role of AI in combating cyberattacks. *Communications of the ACM*, 61(10), pp.6-8.

Ghazali, S. & Yap, K.H., 2019. Cloud computing and cyber security: A review. *Proceedings of the 2019 IEEE 8th Global Conference on Consumer Electronics (GCCE)*, pp.198-200. IEEE.

Giacinto, G., Fumera, G. & Roli, F., 2002. A comparison of techniques for enhancing the robustness of neural network-based classifiers. *Pattern Recognition Letters*, 23(1), pp.1-12.

Goh, G., 2020. Cyber security threats and how to counter them. *Journal of Information Security*, 11(2), pp.89-105.

- Gündüz, E., Topcu, A., İsmail, S. & Ekinici, G., 2022. Intrusion detection in IoT networks: A comprehensive review and research agenda. *Journal of Ambient Intelligence and Humanized Computing*, 13(8), pp.3297-3317. doi: 10.1007/s12652-021-03640-w.
- Hadi, F., Shah, H.A., Gillani, H.A. & Younis, I., 2023. A comprehensive review of DDoS attack mitigation strategies in cloud computing. *Journal of Cloud Computing: Advances, Systems and Applications*, 12(1), 40. doi: 10.1186/s13677-023-00368-4.
- Han, D., Han, Z. & Zhang, Y., 2019. Distributed denial-of-service (DDoS) attack detection using machine learning: A review. *Journal of Computational Science*, 31, p.102701. doi: 10.1016/j.jocs.2018.10.018.
- Hsiao, W.H., Zhang, X., Hsu, Y. & Chiueh, T.C., 2014. A lightweight and efficient DDoS detection system. *Computer Networks*, 66, pp.124-137.
- Huang, Q., Wu, X. & Zhang, X., 2022. Adaptive ensemble learning for cyber attack detection in industrial control systems. *Computers*, 11(12), 174. doi: 10.3390/computers11120174.
- Jayaraman, P.P., Zhang, H., Tseng, P.H. & Tzeng, Y., 2023. An integrated framework for cloud-based distributed denial of service (DDoS) attack mitigation. *IEEE Transactions on Network and Service Management*, 20(1), pp.341-354. doi: 10.1109/TNSM.2022.3218354.
- Joshi, S., 2020. Review of malware detection techniques using machine learning algorithms. *Journal of Information Security*, 11(4), pp.183-195.
- Kamara, S. & Lauter, K., 2011. Cryptographic cloud storage. In: *International Conference on Financial Cryptography and Data Security*, pp.136-149. Springer Berlin Heidelberg.
- Kim, S., Lee, K., Kim, J., Shin, H. & Kim, H., 2022. A hybrid deep learning model for detecting and classifying DDoS attacks. *Computers*, 11(5), p.69. doi: 10.3390/computers11050069.
- Kumar, K., Prasad, R. & Singh, K., 2020. A review of cybersecurity in cloud computing and future research directions. *Journal of Cloud Computing: Advances, Systems and Applications*, 9(1), 39. doi: 10.1186/s13677-020-00232-x.
- Lee, C., Liu, F. & Li, J., 2019. Survey and evaluation of machine learning algorithms for intrusion detection. *Journal of Computer Security*, 86, p.102291. doi: 10.1016/j.jocs.2019.102291.
- Li, X., Shen, W., Luo, H. & Guo, L., 2023. A deep learning-based anomaly detection framework for distributed denial of service attacks in cloud environments. *Future Generation Computer Systems*, 138, pp.94-105. doi: 10.1016/j.future.2022.10.003.
- Liu, L., Li, X. & Yu, Y., 2023. Network anomaly detection using long short-term memory network and feature attention. *Journal of Computer Networks and Communications*, 2023, 9684931. doi: 10.1155/2023/9684931.

- Masi, T. & Masek, P., 2022. A review of machine learning and deep learning techniques for intrusion detection systems. *Computers*, 11(12), 178. doi: 10.3390/computers11120178.
- Mollah, M.A., Rahman, M.M., Rahman, M.M., Hasan, M.M. & Abedin, M.J., 2023. Machine learning-based approach for detecting and mitigating DDoS attacks. *Journal of Information Security and Applications*, 72, 103497. doi: 10.1016/j.jisa.2023.103497.
- Nasser, M.A., Mahmoud, K. & Saeed, S., 2016. A new approach for detecting and mitigating DDoS attacks using a hybrid model. *Computers & Security*, 59, pp.39-50. doi: 10.1016/j.cose.2015.12.004.
- Nguyen, H.T., Le, D.T., Nguyen, N.A. & Nguyen, T.V., 2021. Deep learning for intrusion detection in industrial IoT environments. *Journal of Computer Networks and Communications*, 2021, 8899441. doi: 10.1155/2021/8899441.
- Obaid, H., Agrawal, N. & Gupta, D., 2022. An ensemble approach for DDoS attack detection using machine learning and feature engineering. *Journal of Cloud Computing: Advances, Systems and Applications*, 11(1), 46. doi: 10.1186/s13677-022-00345-2.
- Omer, M., Khan, M., Jan, S., Jan, I. & Rashid, M., 2023. A novel hybrid DDoS detection mechanism using ensemble learning for cloud computing. *Security and Privacy*, 1(4), e226. doi: 10.1002/spy2.226.
- Park, J., Kwon, D. & Kim, Y., 2022. An effective approach for DDoS attack detection using deep learning. *Journal of Computational Science*, 55, 102612. doi: 10.1016/j.jocs.2022.102612.
- Raza, H., Younis, M., Zafar, M.A., Shahid, M.F., Javaid, N. & Ali Shah, S.Z., 2022. Comprehensive survey on detection and mitigation techniques of distributed denial of service attacks. *Journal of Ambient Intelligence and Humanized Computing*, 13, pp.1049-1080. doi: 10.1007/s12652-021-03577-9.
- Sadiq, M., Rasheed, A., Khan, M.A. & Afaq, M., 2021. A review of machine learning techniques for DDoS attack detection. *Journal of Computer Security*, 103, p.102128. doi: 10.1016/j.jocs.2021.102128.
- Shah, A., Khan, M.A., Ullah, M.F. & Ahmed, E., 2022. Detecting distributed denial-of-service (DDoS) attacks in cloud computing: A survey and research agenda. *Journal of Network and Computer Applications*, 197, 103232. doi: 10.1016/j.jnca.2022.103232.
- Singh, R., Singh, D., Chaudhary, S. & Gaur, S., 2022. DDoS attack detection using machine learning in cloud computing. *Computers & Security*, 115, 102680. doi: 10.1016/j.cose.2022.102680.
- Soualhi, K., Mnaouer, M., Ben-Ncer, A., Sarih, R. & Ghozali, S., 2021. A survey on intrusion detection and prevention systems in the IoT environment. *Journal of Network and Computer Applications*, 187, 103137. doi: 10.1016/j.jnca.2021.103137.

Sultan, S., Ahmed, A., Hussain, A., Bibi, S. & Manzoor, A., 2022. A survey of machine learning techniques for malware detection. *Journal of Cyber Security and Privacy*, 2(4), pp.549-567.

Wang, X., Yang, Z., Liu, Z. & Zhang, C., 2023. Anomaly detection in cloud computing environments: A comprehensive survey. *Future Generation Computer Systems*, 134, pp.162-176. doi: 10.1016/j.future.2022.06.015.

Wu, J., Wu, X., Zhang, Y., Liu, X. & Liu, Z., 2022. A deep learning-based framework for network anomaly detection and mitigation in cloud computing environments. *Journal of Cloud Computing: Advances, Systems and Applications*, 11(1), 33. doi: 10.1186/s13677-022-00335-4.

Xue, J., Wang, Y., Yang, J. & Xu, B., 2023. Research on the application of deep learning techniques in network security. *Security and Privacy*, 1(3), e236. doi: 10.1002/spy2.236.

Xu, X., Qian, W. & Zhang, Y., 2021. Cloud-based intrusion detection system for IoT networks. *Journal of Computer Networks and Communications*, 2021, 8846211. doi: 10.1155/2021/8846211.

Yang, Y., Li, Q., Liu, S. & Zheng, Y., 2020. A review on DDoS attack detection methods using machine learning techniques. *Computer Networks*, 171, 107089. doi: 10.1016/j.comnet.2020.107089.

Yoon, S., Park, C. & Lee, S., 2023. Advanced techniques for DDoS attack detection and mitigation: A survey. *Journal of Computer Security*, 106, p.103580. doi: 10.1016/j.jocs.2023.103580.

Zhang, X., Zhang, M., Wang, W. & Liu, S., 2022. A hybrid deep learning model for DDoS attack detection. *Future Generation Computer Systems*, 128, pp.263-274. doi: 10.1016/j.future.2021.11.025.

Zhao, J., Yang, S., Li, Z., Xu, J. & Guo, Y., 2022. Ensemble learning for intrusion detection in industrial control systems. *Journal of Computer Security*, 104, p.102224. doi: 10.1016/j.jocs.2022.102224.

Zheng, Z., Dong, C. & Zhang, X., 2022. A survey of DDoS attack detection and mitigation techniques in cloud computing environments. *Future Generation Computer Systems*, 130, pp.190-205. doi: 10.1016/j.future.2022.08.016.

Zhou, Y., Yang, Y., Li, W., Zhou, L. & Wang, Z., 2023. A deep learning approach for detecting DDoS attacks in cloud computing environments. *IEEE Transactions on Network and Service Management*, 20(3), pp.3298-3310. doi: 10.1109/TNSM.2023.3252135.

İNTİHAL RAPORU İLK SAYFASI

ORJİNALLİK RAPORU

%**4**

BENZERLİK ENDEKSİ

%**3**

İNTERNET KAYNAKLARI

%**1**

YAYINLAR

%**1**

ÖĞRENCİ ÖDEVLERİ

BİRİNCİL KAYNAKLAR

1

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

Öğrenci Ödevi

%**1**

2

www.seruvenyayinevi.com

İnternet Kaynağı

<%**1**

3

dergipark.org.tr

İnternet Kaynağı

<%**1**

4

Hakan Aydın, Gülsüm Zeynep Gürkaş Aydın,
Ahmet Sertbaş, Muhammed Ali Aydın.

"Internet of things security: A multi-agent-
based defense system design", Computers
and Electrical Engineering, 2023

Yayın

<%**1**

5

acikbilim.yok.gov.tr

İnternet Kaynağı

<%**1**

6

www.researchgate.net

İnternet Kaynağı

<%**1**

7

www.openaccess.hacettepe.edu.tr:8080

İnternet Kaynağı

<%**1**

