

# AI-Augmented Intrusion Detection: Advancements in Real-Time Cyber Threat Recognition

EDIMA DAVID ETIM<sup>1</sup>, IBORO AKPAN ESSIEN<sup>2</sup>, JOSHUA OLUWAGBENGA AJAYI<sup>3</sup>,  
ESEOGHENE DANIEL ERIGHA<sup>4</sup>, EHIMAH OBUSE<sup>5</sup>

<sup>1</sup>Core IP Engineer, Cobranet Ltd, Lekki, Lagos, Nigeria

<sup>2</sup>Mobil Producing Nigeria Unlimited, Eket, Nigeria

<sup>3</sup>Kobo360, Lagos, Nigeria

<sup>4</sup>Senior Software Engineer, Eroo Consulting, Dubai, UAE

<sup>5</sup>Lead Software Engineer, Choco, Berlin, Germany

*Abstract- The rapid evolution of cyber threats demands innovative approaches to safeguarding digital infrastructures. AI-augmented intrusion detection systems (IDS) represent a paradigm shift in real-time cyber threat recognition, integrating advanced machine learning algorithms, deep learning architectures, and intelligent data analytics to detect, classify, and mitigate threats with unprecedented speed and accuracy. This study examines recent advancements in AI-driven IDS, focusing on their capacity to process vast, heterogeneous network data streams in real time, identify complex attack patterns, and adapt to emerging threats through continuous learning mechanisms. The integration of anomaly detection, behavioral analysis, and threat intelligence feeds enables these systems to recognize subtle deviations from normal activity, even in encrypted traffic, reducing false positives and enhancing situational awareness. Additionally, the research highlights the role of reinforcement learning in optimizing detection policies and response strategies, ensuring adaptive defense against polymorphic and zero-day attacks. Implementation challenges such as data quality, computational overhead, algorithm interpretability, and adversarial evasion are critically assessed, alongside potential solutions including federated learning, explainable AI, and hybrid signature-anomaly detection models. The study further explores real-world deployments in enterprise, cloud, and IoT environments, illustrating performance metrics such as detection rate, precision, recall, and mean time to detect (MTTD). These case analyses underscore the transformative impact of AI in accelerating intrusion detection response times, minimizing operational disruption, and strengthening cyber resilience. The paper*

*concludes by identifying research gaps and recommending future directions, including energy-efficient AI models, integration with security orchestration and automated response (SOAR) platforms, and the development of standardized benchmarks for AI-based IDS evaluation. By bridging the gap between traditional security paradigms and intelligent automation, AI-augmented intrusion detection systems offer a robust pathway toward proactive, adaptive, and scalable cyber defense in an era of increasingly sophisticated threats.*

*Index Terms- AI-Augmented Intrusion Detection, Real-Time Cyber Threat Recognition, Machine Learning, Deep Learning, Anomaly Detection, Behavioral Analysis, Zero-Day Attacks, Explainable AI, Cybersecurity Resilience, Adaptive Defense Systems*

## I. INTRODUCTION

The rapid expansion of digital infrastructures and the proliferation of interconnected devices have significantly transformed the cyber threat landscape, creating an environment where malicious actors continuously develop increasingly sophisticated attack techniques. Modern cyber threats are no longer confined to simple malware or easily detectable exploits; instead, they often involve multi-stage, stealthy, and adaptive tactics capable of evading conventional security measures. Advanced Persistent Threats (APTs), zero-day exploits, polymorphic malware, and coordinated distributed denial-of-service (DDoS) attacks have become prevalent, exploiting vulnerabilities across enterprise networks,

cloud platforms, Internet of Things (IoT) ecosystems, and critical infrastructure systems. This growing complexity places immense pressure on cybersecurity defenses to detect and respond to malicious activities in real time (Dogho, 2011, Oni, et al., 2018).

Traditional intrusion detection systems (IDS), whether signature-based or anomaly-based, face inherent limitations in meeting this challenge. Signature-based IDS rely on predefined patterns of known threats, making them ineffective against novel or evolving attacks. Anomaly-based IDS, while capable of identifying unusual patterns, often suffer from high false positive rates and lack the contextual intelligence required to distinguish between benign anomalies and genuine threats. Both approaches struggle to adapt rapidly to the dynamic nature of modern cyberattacks, resulting in delayed detection, inefficient incident response, and increased risk to organizational assets (AdeniyiAjonbadi, et al., 2015).

Integrating Artificial Intelligence (AI) into intrusion detection offers a compelling solution to these shortcomings by enabling systems to learn from large volumes of heterogeneous data, adapt to evolving threat patterns, and provide more accurate, context-aware analyses. AI-powered IDS can leverage machine learning, deep learning, and advanced analytics to detect both known and unknown threats with reduced false positives, improved scalability, and faster decision-making. Furthermore, AI integration allows for the incorporation of behavioral analysis, threat intelligence feeds, and automated response mechanisms, creating a proactive defense posture capable of mitigating threats before they cause significant damage (Oni, et al., 2018).

This study aims to investigate the advancements in AI-augmented intrusion detection, with a particular focus on real-time cyber threat recognition. It explores the underlying technologies, architectural models, implementation strategies, and operational challenges associated with AI-driven IDS. The scope encompasses a critical evaluation of recent research, emerging techniques, and practical deployment scenarios, providing a comprehensive understanding of how AI can revolutionize intrusion detection to meet the demands of today's complex and fast-

evolving threat environment (Otoum, 2019, Pauwels & Denton, 2018, Yarali, et al., 2019).

## 2.1. Literature Review

Intrusion detection systems (IDS) have undergone a significant evolution since their inception, driven by the necessity to address increasingly complex cyber threats targeting modern digital infrastructures. Early IDS implementations were primarily signature-based, relying on databases of known attack patterns or "signatures" to identify malicious activities. While effective against known threats, these systems lacked the adaptability to detect novel or evolving attacks, resulting in a reactive rather than proactive security posture (Orren, 2019, Renda, 2019, Tobiyama, et al., 2016). This limitation led to the development of anomaly-based IDS, which establish baselines of normal network or system behavior and flag deviations as potential intrusions. Although anomaly detection broadened the detection scope to include unknown threats, it also introduced a high rate of false positives, as legitimate but unusual activities were frequently misclassified as malicious (Adenuga, Ayobami & Okolo, 2019). Over time, hybrid IDS models emerged, combining signature and anomaly detection to improve accuracy, yet even these approaches faced scalability issues and challenges in real-time analysis as network traffic volumes and attack sophistication increased.

The integration of Artificial Intelligence (AI) into intrusion detection represents a pivotal shift in the evolution of IDS, providing systems with the capacity to learn from large and diverse datasets, adapt to dynamic threat landscapes, and deliver context-aware, real-time insights. AI applications in cybersecurity extend beyond intrusion detection to include malware classification, phishing detection, fraud prevention, vulnerability assessment, and automated incident response. In the context of IDS, AI techniques such as machine learning, deep learning, and natural language processing are employed to identify complex and subtle attack patterns that traditional methods might overlook. Machine learning algorithms, including decision trees, random forests, support vector machines, and k-nearest neighbors, have been widely applied for feature-based classification of network

traffic. These methods excel in detecting known threats and some anomalies by learning decision boundaries from labeled datasets, but their reliance on predefined features can limit adaptability when facing evolving attack strategies.

Deep learning approaches, on the other hand, have shown remarkable potential in intrusion detection by automatically learning hierarchical representations of data without extensive manual feature engineering. Convolutional Neural Networks (CNNs) have been applied to capture spatial correlations in traffic patterns, while Recurrent Neural Networks (RNNs) and Long Short-Term Memory (LSTM) networks are effective in modeling temporal dependencies in sequential data such as network flows and system logs. Autoencoders and Generative Adversarial Networks (GANs) have been leveraged for unsupervised anomaly detection, enabling the discovery of zero-day attacks without prior labelling (Olasehinde, 2018). Comparative studies consistently indicate that deep learning models often outperform traditional machine learning in terms of detection accuracy and the ability to generalize to new threat types. However, they also come with challenges, including higher computational requirements, longer training times, and a need for large volumes of high-quality labeled data for optimal performance.

In comparing machine learning and deep learning for intrusion detection, several key distinctions emerge. Machine learning models are generally easier to interpret, which is crucial for compliance, auditability, and human analyst trust in security operations. They can be trained relatively quickly on smaller datasets, making them suitable for environments with limited computational resources or where explainability is a priority. Deep learning models, while more resource-intensive, excel in complex, high-dimensional data environments, such as large-scale enterprise or cloud networks, where patterns of malicious activity are deeply embedded in noisy datasets (Mohit, 2018, Sareddy & Hemnath, 2019). Their capacity to integrate multiple data modalities including network traffic, endpoint telemetry, and threat intelligence further enhances their value for comprehensive intrusion detection. Nevertheless, explainability remains a significant barrier to their adoption in regulated industries, prompting research into

explainable AI (XAI) techniques to make deep learning outputs more transparent and actionable.

Despite these advancements, critical research gaps persist in the field of AI-augmented intrusion detection for real-time threat recognition. One of the most pressing challenges is the issue of timeliness. Many AI models, particularly deep learning architectures, are optimized for accuracy but not necessarily for speed, leading to latency in detection that can undermine their effectiveness in stopping fast-moving attacks. Achieving both high detection accuracy and low latency remains an unresolved problem, particularly in high-bandwidth, low-latency environments such as 5G networks or industrial control systems. Another gap lies in the ability to handle concept drift the phenomenon where the statistical properties of network traffic and attack patterns change over time. Static models, even when highly accurate initially, degrade in performance as attackers adapt and infrastructure evolves (Hao, et al., 2019, Xu, et al., 2019). This necessitates ongoing model retraining, which is resource-intensive and operationally challenging, especially in mission-critical environments. Figure 1 shows main components of intrusion detection system presented by Karatas, Demir & Sahingoz, 2018.

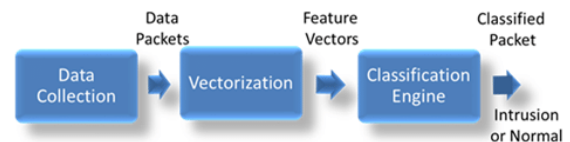


Figure 1: Main Components of Intrusion Detection System (Karatas, Demir & Sahingoz, 2018).

Data quality and availability also present persistent hurdles. Many high-performing AI models are trained on benchmark datasets such as KDD Cup 99, NSL-KDD, or UNSW-NB15, which, while useful for research, may not reflect the complexity and heterogeneity of modern real-world traffic. The scarcity of large-scale, up-to-date, and labeled datasets representing diverse attack types hampers the generalization capability of AI models in production environments. Privacy concerns further limit the sharing of real-world attack data, complicating collaborative research and cross-industry model development.

Another research gap is resilience against adversarial attacks. AI models themselves can be targeted by adversarial machine learning techniques, where small, carefully crafted perturbations to input data cause misclassification. This vulnerability raises significant concerns for the reliability of AI-augmented IDS in adversarial settings. Developing models that are robust to such attacks, while maintaining high accuracy and low false positive rates, remains an active area of investigation (Weng, et al., 2019, Zhou, et al., 2019).

Finally, the integration of AI-augmented IDS into operational cybersecurity workflows presents its own set of challenges. While research has demonstrated high-performing models in controlled environments, deployment at scale requires compatibility with existing security tools, interoperability with SIEM and SOAR platforms, and minimal disruption to established processes. Balancing automation with human oversight is also crucial to prevent overreliance on AI and ensure that analysts can interpret and act upon AI-generated alerts effectively.

In summary, the literature reflects substantial progress in the application of AI to intrusion detection, with machine learning and deep learning each offering distinct advantages and trade-offs. The shift toward AI-augmented IDS has enhanced detection capabilities, expanded the scope of recognizable threats, and opened pathways to more adaptive, real-time defenses. However, addressing latency, concept drift, data scarcity, adversarial resilience, and operational integration is essential for realizing the full potential of these systems. Continued research in these areas, coupled with advances in explainable AI and privacy-preserving techniques, will be key to developing AI-augmented intrusion detection systems capable of meeting the demands of an ever-evolving cyber threat landscape (Brynskov, Facca & Hrasko, 2018, Kumari, Hsieh & Okonkwo, 2017).

## 2.2. Methodology

This study employs an integrated approach combining deep learning, ensemble machine learning, and adaptive security analytics to advance real-time intrusion detection capabilities. Initially, raw network data is collected from heterogeneous sources,

including packet captures, system logs, and user activity streams. The collected data undergoes preprocessing involving noise reduction, missing value handling, normalization, and feature engineering to ensure compatibility with AI algorithms. Feature selection techniques such as information gain, mutual information, and dimensionality reduction are applied to retain the most discriminative attributes, thereby enhancing computational efficiency and reducing overfitting risks.

The core of the system is an AI-augmented Intrusion Detection System (IDS) that merges deep learning architectures such as Convolutional Neural Networks (CNNs) for spatial pattern recognition and Recurrent Neural Networks (RNNs) for temporal behavior modeling with traditional ensemble methods like Random Forests and Gradient Boosting. This hybrid detection engine is designed to recognize both known attack signatures and anomalous patterns indicative of zero-day exploits. Model training leverages labeled datasets from benchmark intrusion detection corpora, supplemented with synthetic attack traffic generated via adversarial machine learning techniques to improve resilience against evasion strategies.

Once deployed, the IDS performs real-time threat recognition by continuously monitoring incoming network traffic and system events. Detected threats trigger the decision and response layer, which automates incident handling through alerts, traffic blocking, and detailed forensics logging. A feedback loop is integrated to facilitate continuous model updates, incorporating newly labeled attack data and adversarial training to adapt to evolving cyber threats. This iterative refinement ensures sustained detection accuracy and robustness in dynamic network environments, aligning with current literature on AI-driven cybersecurity advancements and real-world deployment considerations.

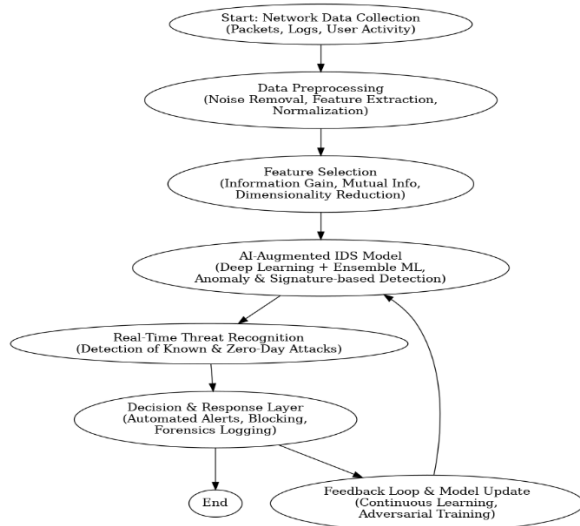


Figure 2: Flow chart of the study methodology

### 2.3. AI Techniques in Intrusion Detection

Artificial Intelligence techniques have transformed the capabilities of intrusion detection systems by enabling them to analyze vast, heterogeneous datasets, adapt to evolving threat patterns, and operate in real time with higher accuracy than traditional methods. Machine learning approaches form the foundational layer of AI-augmented intrusion detection, offering various strategies depending on the availability and nature of training data. In supervised learning, models are trained on labeled datasets containing both normal and malicious instances, allowing them to learn decision boundaries for classifying new observations. Algorithms such as support vector machines, decision trees, random forests, and gradient boosting machines have been widely applied to network traffic and log data, yielding effective detection of known attack types with relatively low computational demands. The limitation of supervised learning lies in its dependence on comprehensive and representative labeled datasets, which may be difficult to obtain in the constantly changing threat landscape. Unsupervised learning, in contrast, is designed to detect anomalies without prior labeling, making it well-suited for identifying novel or zero-day attacks. Clustering techniques such as k-means, DBSCAN, and self-organizing maps can group similar behaviors and flag deviations as suspicious. However, these methods can produce high false positive rates if normal network behavior is highly variable. Reinforcement learning introduces an

adaptive dimension to intrusion detection, where agents learn optimal detection and response strategies through trial-and-error interactions with the environment, guided by reward functions (Achar, 2018, Shah, 2017). This approach is particularly promising for dynamic network environments and automated policy optimization, although it can be computationally intensive and requires careful design to avoid undesirable behaviors.

Deep learning architectures have advanced intrusion detection further by automatically learning complex, hierarchical features from raw data, reducing the need for manual feature engineering. Convolutional Neural Networks (CNNs) have proven effective in extracting spatial correlations from transformed network traffic data, such as flow matrices or encoded packet sequences, enabling the detection of subtle attack signatures embedded in high-dimensional spaces. Recurrent Neural Networks (RNNs), particularly Long Short-Term Memory (LSTM) networks, excel in modeling temporal dependencies in sequential data, such as system logs or time-series network flows, making them valuable for identifying multi-stage attacks that unfold over time. Transformers, with their self-attention mechanisms, offer the ability to capture both local and global dependencies efficiently, enabling scalable intrusion detection across large datasets with parallelizable computation (Duddu, 2018, Ibitoye, et al., 2019). These architectures have demonstrated strong performance in real-time threat recognition scenarios, especially when combined with transfer learning to adapt pretrained models to specific network environments. Despite their power, deep learning models often require substantial computational resources, large volumes of training data, and strategies to address interpretability challenges, particularly in regulated or high-stakes domains. Figure 3 shows figure of types of intrusion detection techniques presented by Kene & Theng, 2015.

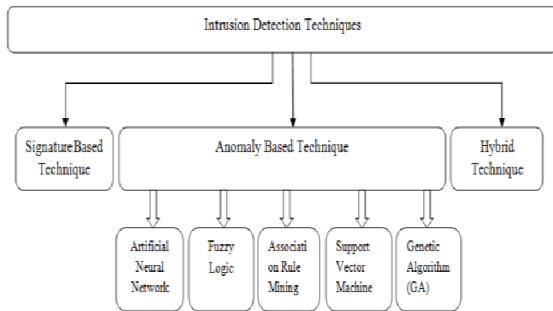


Figure 3: Types of Intrusion Detection Techniques (Kene & Theng, 2015).

Hybrid models represent an important evolution in intrusion detection, combining the precision of signature-based detection with the adaptability of anomaly-based methods. In such systems, signature-based components rapidly detect known threats by matching patterns against established databases, while anomaly-based modules leverage machine learning or deep learning to identify deviations indicative of new or modified attacks. This integration reduces the detection latency for known threats while maintaining vigilance for novel attack vectors, creating a more comprehensive defense (Biggio & Roli, 2018, Shi, et al., 2018). AI-enhanced hybrid systems can dynamically adjust detection thresholds, incorporate contextual information from threat intelligence feeds, and employ ensemble learning to aggregate outputs from multiple detection models, thereby reducing false positives and improving resilience against evasion techniques. Hybrid AI-IDS solutions are increasingly relevant in environments where both established and emerging threats are prevalent, such as cloud infrastructures, IoT deployments, and industrial control systems.

Behavioral analytics and profiling extend the capabilities of AI-augmented intrusion detection by focusing on patterns of activity associated with specific users, devices, or entities. By building behavioral baselines through continuous monitoring of network interactions, system commands, application usage, and access patterns, AI-driven systems can detect deviations that may indicate compromised accounts, insider threats, or stealthy lateral movement within a network. Machine learning algorithms can create dynamic behavioral profiles that adapt over time, accounting for normal changes in

usage while maintaining sensitivity to anomalies (Apruzzese, et al., 2019, Laskov & Lippmann, 2010). Deep learning approaches, particularly those incorporating sequence modeling and attention mechanisms, can enhance behavioral analytics by capturing the contextual relationships between events, enabling more accurate detection of subtle threats. For example, an AI system might identify that a user accessing sensitive databases outside normal working hours, in combination with an unusual volume of data transfers, constitutes a potential security incident. Figure 4 shows classification of intrusion detection techniques presented by Alhakami, et al., 2019.

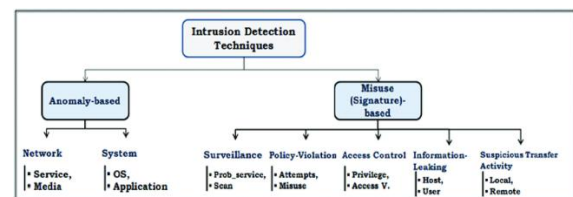


Figure 4: Classification of intrusion detection techniques (Alhakami, et al., 2019).

In practice, the integration of machine learning, deep learning, hybrid detection mechanisms, and behavioral analytics creates a layered defense model that significantly improves the speed, accuracy, and adaptability of intrusion detection. Machine learning models provide interpretable and efficient solutions for structured, labeled datasets, while deep learning architectures excel in high-dimensional, unstructured, or sequential data contexts. Hybrid approaches bridge the strengths of signature and anomaly-based methods, ensuring both known and unknown threats are addressed effectively. Behavioral analytics introduces a human and contextual dimension, allowing systems to identify threats that might evade purely technical detection strategies. The synergy among these techniques allows AI-augmented intrusion detection systems to operate effectively in complex, high-throughput environments, from enterprise networks to national critical infrastructure (Chen, et al., 2019, Dasgupta & Collins, 2019).

However, deploying these AI techniques in real-world intrusion detection scenarios requires addressing several operational challenges. Ensuring data quality and representativeness is critical for avoiding bias and maintaining high detection accuracy across different

environments. Model retraining and adaptation to concept drift are necessary to sustain performance as network behaviors and attack strategies evolve. Computational resource management is essential, particularly for deep learning models deployed in real-time, high-volume environments. Additionally, explainability remains an ongoing concern; security analysts need to understand and trust the outputs of AI-driven IDS to make informed decisions. Advances in explainable AI (XAI) and human-in-the-loop systems are helping bridge this gap, enabling transparent decision-making without sacrificing detection capability (Liu, et al., 2018, Sethi, et al., 2018).

Overall, AI techniques in intrusion detection are enabling a shift from reactive, static defenses to proactive, adaptive, and context-aware security systems. By combining the strengths of supervised, unsupervised, and reinforcement learning with the representational power of CNNs, RNNs, LSTMs, and Transformers, and by integrating hybrid detection models with behavioral analytics, organizations can build robust intrusion detection systems capable of recognizing and responding to cyber threats in real time. As cyber threats continue to grow in complexity and scale, the ongoing refinement and integration of these AI techniques will be critical to maintaining the resilience and security of digital infrastructures.

#### 2.4. Real-Time Cyber Threat Recognition Framework

A real-time cyber threat recognition framework within the context of AI-augmented intrusion detection is built on the ability to gather, process, and analyze data from diverse and heterogeneous sources with minimal latency, while maintaining high levels of accuracy and adaptability to evolving threats. The first critical component of such a framework is data acquisition, which involves collecting information from a wide range of sources, including network traffic flows, packet captures, system and application logs, authentication records, endpoint telemetry, and device-specific data from Internet of Things (IoT) ecosystems. In modern enterprise and cloud environments, this process must encompass both IT and operational technology (OT) domains, as threats often traverse traditional network boundaries to target

interconnected systems. Data acquisition in real time requires scalable architectures capable of ingesting high-throughput streams without introducing bottlenecks, often achieved through distributed data collectors, API integrations, and message queuing systems (Dalal, 2018, Mittal, Joshi & Finin, 2019). The heterogeneity of the sources means the data arrives in multiple formats and at varying levels of granularity, necessitating robust preprocessing pipelines to normalize, cleanse, and enrich the information before it enters the analytical phase.

Preprocessing is essential for ensuring that raw data, which often contains noise, redundancies, and inconsistencies, is transformed into a form suitable for machine learning and deep learning models. This stage may include tasks such as timestamp synchronization across distributed systems, removal of duplicate records, conversion of categorical variables into machine-readable encodings, and resolution of missing or corrupted values. In network traffic analysis, for example, preprocessing might involve aggregating flows over defined time windows, extracting protocol-specific metadata, and anonymizing sensitive identifiers to comply with privacy requirements. IoT device data may require additional parsing to handle proprietary formats and sensor-specific attributes, ensuring compatibility with the broader detection framework (Holzinger, et al., 2018, Mavroeidis & Bromander, 2017). The preprocessing pipeline also integrates enrichment processes, where contextual information such as geolocation, device reputation scores, or asset criticality is added to the dataset to enhance the interpretability and effectiveness of downstream models.

Feature engineering and selection play a pivotal role in improving the accuracy and efficiency of the detection models. Feature engineering involves creating informative variables that capture the behavioral and structural characteristics of network and system activity, often derived from raw telemetry. Examples include statistical measures of packet size distributions, connection frequency histograms, session duration patterns, user access frequency to sensitive resources, or entropy measures for payload inspection. The goal is to translate raw data into features that reveal patterns indicative of malicious

activity while minimizing irrelevant or redundant information. Feature selection, whether through statistical tests, embedded model techniques, or dimensionality reduction algorithms, ensures that only the most relevant attributes are retained for training and inference (Hagras, 2018, Svenmarck, et al., 2018). This step not only improves model performance by reducing overfitting and computational complexity but also enhances interpretability, allowing analysts to understand which factors are most influential in classifying threats. In environments where data is high-dimensional, such as large-scale enterprise networks or IoT deployments with thousands of devices, feature selection becomes crucial for sustaining real-time responsiveness.

Integration of anomaly detection mechanisms with threat intelligence feeds significantly amplifies the capability of the real-time threat recognition framework. Anomaly detection models, whether based on supervised, unsupervised, or hybrid learning, identify deviations from established baselines of normal activity. In an AI-augmented setting, these models continuously adapt to changing patterns, refining their baselines as legitimate behavior evolves. However, anomaly detection alone can generate false positives, particularly in dynamic environments with legitimate but unusual behavior. To address this, the framework incorporates threat intelligence feeds that provide curated, continuously updated information on known malicious IP addresses, domain names, malware hashes, vulnerabilities, and attack campaigns. This external intelligence enables the system to validate and enrich anomalies detected internally, reducing false positives and prioritizing alerts that match known threat indicators (Glomsrud, et al., 2019, Gudala, et al., 2019). Furthermore, integrating both open-source and commercial threat intelligence sources ensures broader coverage, while contextualizing anomalies within the global cyber threat landscape. For example, if an anomaly detection model identifies an unusual outbound connection to a previously unseen domain, correlating this with a threat intelligence report linking the domain to a botnet can elevate the severity of the alert and trigger an immediate automated response.

The role of streaming analytics in real-time processing is central to ensuring that detection and response occur

within timeframes that can prevent or minimize damage. Streaming analytics refers to the continuous analysis of data as it is generated, enabling immediate insight extraction without the delays associated with batch processing. This capability is vital for intrusion detection, where even seconds of delay can allow an attacker to exfiltrate data, move laterally across the network, or disable defenses. In a real-time cyber threat recognition framework, streaming analytics platforms such as Apache Kafka, Apache Flink, or Spark Streaming orchestrate the ingestion, transformation, and analysis of data streams at scale (Lawless, et al., 2019, O'Sullivan, et al., 2019). AI models deployed in this environment operate on sliding or tumbling windows, making inferences on recent activity and updating threat assessments as new events arrive. Streaming analytics pipelines can apply multi-stage processing, where initial filters remove benign events based on established whitelists, followed by feature extraction modules that feed into machine learning or deep learning classifiers. The outputs are then enriched with threat intelligence matches and risk scoring algorithms, producing actionable alerts that can be sent to security operations centers (SOCs) or automated response systems.

By combining data acquisition from heterogeneous sources, rigorous preprocessing, effective feature engineering, integrated anomaly detection, and real-time streaming analytics, the framework enables continuous monitoring and rapid recognition of both known and emerging threats. One of its defining advantages is adaptability the ability to learn from both historical incidents and evolving real-time data to refine detection strategies. This adaptability is particularly critical in defending against advanced persistent threats (APTs) and zero-day attacks, which often bypass static defenses by mimicking legitimate behavior. In such cases, the fusion of anomaly detection with live threat intelligence can uncover hidden attack vectors that would otherwise remain undetected until significant damage occurs (Otokiti, 2012).

However, implementing such a real-time recognition framework requires addressing operational and technical challenges. The need for low-latency processing must be balanced against the computational demands of complex AI models,

particularly deep learning architectures, which may require specialized hardware such as GPUs or TPUs for inference at scale. Data privacy concerns must also be addressed, particularly when integrating third-party threat intelligence or aggregating telemetry from multiple organizational domains. This may involve deploying federated learning models that allow collaborative detection without sharing raw data, or applying privacy-preserving transformations during preprocessing. Furthermore, ensuring high data quality is essential; inaccurate or incomplete data at the ingestion stage can propagate errors through the detection pipeline, leading to false alarms or missed threats (Otokiti, 2018).

Another challenge lies in maintaining interoperability with existing security infrastructure, including SIEM, SOAR, and endpoint detection and response (EDR) systems. The framework must be capable of both consuming data from and sending actionable outputs to these platforms without disrupting established workflows. This requires flexible APIs, adherence to industry data exchange standards, and modular architecture that allows integration with diverse security tools. Finally, human oversight remains a critical component of the system, ensuring that AI-generated alerts are validated, false positives are managed, and evolving threats are correctly incorporated into the model's knowledge base.

In essence, a real-time cyber threat recognition framework for AI-augmented intrusion detection is a dynamic, data-driven ecosystem that unites multiple technical disciplines: data engineering, machine learning, threat intelligence, and real-time analytics into a cohesive whole. Its strength lies in its ability to continuously learn, adapt, and act upon both known and emerging threats at speeds necessary to protect today's highly connected and fast-moving digital environments. As the complexity and velocity of cyber threats continue to rise, the refinement of such frameworks will be pivotal in enabling organizations to move from reactive defense to proactive, predictive, and automated protection strategies.

## 2.5. Advancements in AI-Augmented IDS

Advancements in AI-augmented intrusion detection systems (IDS) have been marked by the growing sophistication of learning mechanisms, optimization strategies, transparency initiatives, and privacy-preserving techniques, all aimed at enhancing their resilience against increasingly complex and adaptive cyber threats. One of the most important developments in this domain is the shift toward adaptive learning and continuous model retraining to address the evolving nature of malicious activities. Unlike static detection models that degrade over time as attackers modify their tactics, adaptive learning approaches enable IDS to refine their detection capabilities based on new data, emerging threat signatures, and observed changes in normal network behavior. This process often involves incremental or online learning techniques, where models are updated in near real time without requiring complete retraining from scratch. In practice, adaptive learning supports the rapid integration of threat intelligence from recent incidents, ensuring that the system can recognize variations of known attacks as well as brand-new exploits. For example, an IDS deployed in a large enterprise network might adjust its anomaly detection baselines following the introduction of new cloud-based applications, distinguishing between legitimate traffic changes and malicious deviations. However, adaptive retraining must be carefully managed to avoid concept drift in the wrong direction, where malicious patterns inadvertently become normalized due to insufficient labeling or inadequate validation controls (Otokiti & Akorede, 2018; Scholten, et al., 2018).

Building on adaptability, reinforcement learning (RL) introduces an advanced layer of intelligence for optimizing intrusion detection policies dynamically. In RL-based IDS, an agent interacts with the network environment, receiving feedback in the form of rewards or penalties based on the correctness and timeliness of its detection and response actions. This framework enables the IDS to explore various detection thresholds, alerting strategies, and response mechanisms, ultimately converging on policies that maximize long-term effectiveness rather than short-term gains. For instance, in a high-traffic environment, the RL agent might learn to adjust sensitivity settings to maintain high detection accuracy while minimizing

false positives that could overwhelm security analysts (Sharma, et al., 2019). RL is also highly relevant for automated incident response, where the system not only detects threats but also selects the optimal containment or mitigation strategy based on contextual factors, such as the criticality of affected assets or the potential business impact. While RL offers significant promise, its practical deployment requires careful design to prevent unintended consequences, such as overfitting to specific attack patterns or making overly aggressive responses that disrupt legitimate operations.

Transparency in decision-making has become a critical requirement for AI-augmented IDS, especially in regulated industries where explainability is tied to compliance, auditability, and trust. Explainable AI (XAI) addresses this by providing insights into how and why a model arrives at a particular detection or classification outcome. In the context of intrusion detection, XAI techniques can highlight which features, traffic patterns, or user behaviors contributed most to flagging an event as malicious, enabling security analysts to validate the detection and understand its rationale. Methods such as SHAP (SHapley Additive exPlanations), LIME (Local Interpretable Model-agnostic Explanations), and attention-based visualization in deep learning models are increasingly integrated into IDS to make their decision processes more transparent (Ajonbadi, et al., 2014). This not only improves analyst trust in automated alerts but also aids in refining detection models by revealing potential biases, irrelevant feature dependencies, or gaps in training data. Furthermore, explainable models facilitate faster incident triage, as analysts can immediately grasp the context and significance of an alert, reducing mean time to investigate (MTTI) and respond (MTTR). Despite these benefits, there is a trade-off between model complexity and interpretability, as some of the most accurate deep learning architectures are also the most opaque, requiring the development of hybrid solutions that balance predictive performance with understandable reasoning.

In parallel with advancements in adaptability, optimization, and explainability, federated learning has emerged as a transformative approach for privacy-preserving intrusion detection, enabling collaborative

model training across multiple organizations or distributed systems without the need to share raw data. In federated learning setups, each participating node be it a corporate network, IoT deployment, or cloud tenant trains a local model on its own data, then shares only model updates or gradients with a central aggregator. These updates are combined to produce a global model that benefits from the collective knowledge of all participants while ensuring that sensitive network data never leaves its original environment (Ajonbadi, Otokiti & Adebayo, 2016, Menson, et al., 2018). This approach is particularly valuable in sectors such as healthcare, finance, and critical infrastructure, where regulatory constraints and confidentiality concerns limit the sharing of security logs or operational telemetry. In the context of zero-day attack detection, federated learning allows the aggregation of insights from diverse environments, increasing the chances of recognizing novel threats that manifest differently across networks. To enhance privacy further, techniques such as secure aggregation, homomorphic encryption, and differential privacy can be applied, ensuring that even the shared model updates cannot be reverse-engineered to reveal sensitive information.

The integration of adaptive learning, reinforcement learning, XAI, and federated learning into AI-augmented IDS represents a synergistic advancement rather than a set of isolated innovations. Adaptive learning ensures that models remain current and relevant, reinforcement learning optimizes the strategic aspects of detection and response, explainable AI provides the interpretability needed for operational trust, and federated learning enables broad-based collaboration without compromising privacy. Together, these capabilities create IDS solutions that are not only technically advanced but also operationally viable in the complex realities of modern cybersecurity.

In practical deployments, these advancements often intersect in compelling ways. For example, an enterprise IDS might use federated learning to train its anomaly detection models on patterns observed across a consortium of industry peers, while reinforcement learning agents fine-tune the system's alerting and response strategies based on the organization's specific risk appetite. At the same time, explainable AI

components ensure that the SOC team can understand and justify automated decisions during security audits or incident reviews, and adaptive retraining mechanisms keep the system responsive to sudden changes in network behavior, such as those introduced by remote workforce expansions or new SaaS integrations (Mustapha, et al., 2018).

Nonetheless, implementing these advancements in real-world environments involves overcoming certain challenges. Adaptive learning and frequent model retraining can be resource-intensive, necessitating efficient scheduling and prioritization to avoid overloading computational infrastructure. Reinforcement learning agents must be constrained by well-defined safety rules to prevent harmful automated actions, and their training processes can be lengthy and data-intensive. XAI integration can increase processing overhead, especially when generating detailed explanations for high-volume alerts, requiring careful balance between depth of explanation and system performance. Federated learning deployments must contend with issues such as heterogeneous data distributions, varying computational capacities across participants, and potential poisoning attacks where malicious updates are introduced to corrupt the global model (Nsa, et al., 2018).

Despite these challenges, the trajectory of AI-augmented IDS development strongly suggests that these advancements will become increasingly standard features rather than experimental capabilities. As cyber threats continue to grow in scale, speed, and sophistication, the need for systems that can learn continuously, optimize dynamically, explain their reasoning, and collaborate securely will only intensify. Research in these areas is expanding rapidly, with promising developments in lightweight adaptive models for edge deployment, multi-agent reinforcement learning for coordinated defense, inherently interpretable neural architectures, and blockchain-integrated federated learning frameworks for secure model governance (Ajonbadi, Mojeed-Sanni & Otokiti, 2015).

The future of intrusion detection lies in seamlessly combining these advancements into unified platforms

capable of operating across heterogeneous environments while meeting the performance, transparency, and privacy demands of modern cybersecurity operations. AI-augmented IDS that embody adaptive learning, reinforcement learning, XAI, and federated learning will not only detect and mitigate threats more effectively but also foster greater trust, compliance, and collaboration in the shared fight against cyber adversaries. In doing so, they will play a critical role in moving organizations toward a proactive and resilient security posture, where detection is instantaneous, responses are optimized, decisions are explainable, and collaboration is both secure and scalable (Lawal, Ajonbadi & Otokiti, 2014).

## 2.6. Implementation Challenges

Implementing AI-augmented intrusion detection systems in real-time cyber threat recognition environments presents a number of significant challenges that extend beyond technical development to encompass operational, regulatory, and ethical dimensions. One of the most prominent obstacles is the high computational requirement associated with training and deploying advanced AI models, particularly deep learning architectures, at the scale and speed necessary for real-time security operations. Models such as convolutional neural networks, recurrent neural networks, transformers, and ensemble hybrids often demand substantial processing power, high memory bandwidth, and specialized hardware accelerators such as GPUs or TPUs to achieve low-latency inference on streaming data (Ridley, 2018, Su, et al., 2016, Zhu, Hu & Liu, 2014). In high-throughput environments such as large enterprise networks, cloud platforms, or industrial control systems network traffic volumes can reach millions of events per second, and processing this data in real time places considerable strain on available infrastructure. Scalability further complicates the issue; a model that performs well in a controlled laboratory setting may suffer significant degradation when scaled to enterprise or multi-cloud deployments due to network bottlenecks, distributed processing challenges, and the overhead of integrating multiple data sources. Efficient model compression, hardware-aware optimization, and edge-based inference can mitigate some of these concerns, but

implementing them without sacrificing detection accuracy remains a delicate balancing act.

Beyond computational demands, data quality and labeling issues represent another critical challenge in the implementation of AI-augmented intrusion detection. Machine learning and deep learning models rely heavily on large volumes of representative data to achieve robust performance, but security datasets often suffer from noise, incompleteness, and inconsistencies. Real-world network traffic can include erroneous logs, missing fields, or misaligned timestamps, all of which can degrade model accuracy. Imbalanced datasets are a particularly acute problem; in most network environments, benign activity vastly outnumbers malicious events, leading models to become biased toward the majority class and potentially missing rare but critical attack signatures (Chen, et al., 2019, Han, et al., 2018, Vinayakumar, et al., 2019). Addressing imbalance requires strategies such as oversampling minority classes, generating synthetic attack data, or applying anomaly detection methods that do not assume balanced distributions. Labeling is equally problematic, as accurately annotating network data with ground truth requires expert knowledge and is labor-intensive. Manual labeling can introduce errors, especially when complex, stealthy attacks are involved, and automated labeling tools are not yet sufficiently reliable to replace human oversight. These limitations can lead to models that perform well on benchmark datasets but fail to generalize in production settings.

The threat of adversarial attacks and model evasion tactics adds another layer of complexity to deploying AI in intrusion detection. Adversarial machine learning techniques exploit vulnerabilities in model decision boundaries, enabling attackers to subtly manipulate input data so that malicious activities are misclassified as benign. In the context of intrusion detection, this could involve crafting network packets or altering behavioral patterns in ways that fool anomaly detectors without disrupting the underlying attack (Appelt, et al., 2018, Choraś & Kozik, 2015, Ganesan, et al., 2016). Evasion tactics also include mimicry attacks, where adversaries deliberately imitate normal traffic patterns, and poisoning attacks, where attackers inject malicious data into the training set to corrupt the model's learning process. Such

threats not only undermine detection performance but also erode trust in automated systems. Defending against these tactics requires ongoing research into adversarially robust models, defensive distillation, input sanitization, and continuous monitoring for anomalies in model behavior. However, implementing these countermeasures often increases computational overhead and can introduce new complexities into model maintenance.

In addition to the technical challenges, regulatory, ethical, and privacy considerations significantly influence the design and deployment of AI-augmented intrusion detection systems. Many jurisdictions enforce stringent data protection regulations, such as the General Data Protection Regulation (GDPR) in the European Union, the Health Insurance Portability and Accountability Act (HIPAA) in the United States for healthcare data, and various sector-specific security standards such as PCI DSS for payment systems (Cybenko, et al., 2014, Huang & Zhu, 2019, Khurana & Kaul, 2019). Compliance with these frameworks often restricts the type of data that can be collected, stored, and processed, especially when it involves personally identifiable information (PII) or sensitive operational telemetry. Since AI-based IDS frequently require access to extensive and detailed network traffic data to train and operate effectively, balancing detection performance with privacy obligations is a persistent challenge. This tension is heightened in multi-tenant environments, such as cloud services, where data from different organizations must be strictly segregated while still enabling collaborative threat intelligence sharing.

Ethical considerations extend beyond compliance to encompass issues of fairness, accountability, and transparency. AI models can inadvertently inherit biases from training data, leading to disproportionate false positives or negatives for specific user groups, departments, or geographic regions. Such biases can have serious operational consequences, including the misallocation of security resources, reputational damage, and in some cases, discriminatory enforcement of security policies. Ethical deployment of AI-augmented IDS requires not only careful dataset curation and bias mitigation strategies but also mechanisms for explainability so that decisions can be understood, challenged, and audited (Feng & Xu,

2017, Kozik & Choraś, 2014, Zhang, Patras & Haddadi, 2019). Explainable AI (XAI) methods are gaining traction in this context, but they must be implemented in ways that preserve sensitive information while still providing actionable insight into model reasoning.

The interplay between privacy concerns and operational requirements also impacts data sharing and collaborative detection efforts. While sharing threat intelligence across organizations can significantly enhance the ability to detect emerging threats, regulatory and competitive considerations often prevent the exchange of raw data. Techniques such as federated learning offer a potential solution by enabling decentralized model training without transferring raw data, but these approaches are still maturing and introduce their own challenges, including communication overhead, heterogeneity of local data distributions, and the risk of model update poisoning (Mohammad, Thabtah & McCluskey, 2014, Sahingoz, Baykal & Bulut, 2018).

From an operational standpoint, integrating AI-augmented intrusion detection into existing security workflows requires alignment with established incident response protocols, SIEM (Security Information and Event Management) systems, and SOC (Security Operations Center) procedures. This integration can be complex, as it involves not only technical compatibility but also changes to organizational processes and human analyst roles. Resistance to change, lack of AI literacy among security teams, and concerns over over-reliance on automation can hinder adoption. Moreover, the “black box” nature of many high-performance models makes it difficult for analysts to validate or trust the alerts generated, which can slow incident resolution and lead to alert fatigue (Jaroszewski, Morris & Nock, 2019, Pham, et al., 2018, Smadi, Aslam & Zhang, 2018).

There is also the issue of lifecycle management for AI models in intrusion detection. Continuous retraining is necessary to adapt to new threats, but this process requires access to high-quality labeled data, significant computational resources, and rigorous testing to avoid introducing regressions or new vulnerabilities. In fast-paced environments, the operational cost of frequent

retraining may be prohibitive, and delays in updating models can leave systems exposed to undetected threats. Additionally, ensuring that retraining processes themselves are secure against manipulation or corruption is vital, as compromised training pipelines could undermine the entire defense system (Nauman, et al., 2018, Sahingoz, et al., 2019, Sowah, et al., 2019).

Finally, the global and cross-border nature of cyber threats introduces jurisdictional complexity. An AI-augmented IDS deployed by a multinational organization may need to comply with multiple, sometimes conflicting, legal regimes governing data sovereignty, surveillance, and security reporting. These requirements can impact where data is stored and processed, which in turn affects system architecture and performance. Navigating this landscape demands close collaboration between technical teams, legal experts, and compliance officers to design systems that meet both security and regulatory objectives without sacrificing operational effectiveness (Chen, et al., 2018, Gan, et al., 2017, Liao, et al., 2019).

In essence, the challenges of implementing AI-augmented intrusion detection for real-time threat recognition are multi-faceted, involving computational scalability, data integrity, model robustness, and socio-legal constraints. Addressing these issues requires a holistic approach that combines technical innovation with strong governance, ethical safeguards, and regulatory compliance strategies (Masoud, Jaradat & Ahmad, 2016, Ramaraj & Chellappan, 2019). Advances in model optimization, adversarial defense, bias mitigation, and privacy-preserving analytics hold promise for overcoming many of these hurdles, but successful deployment will depend equally on organizational readiness, cross-disciplinary collaboration, and sustained investment in both infrastructure and expertise. Without careful attention to these challenges, even the most advanced AI-augmented IDS risks falling short of its potential in the face of an ever-evolving cyber threat landscape.

## 2.7. Case Studies and Real-World Applications

Real-world adoption of AI-augmented intrusion detection systems has moved far beyond experimental laboratory environments, with organizations across multiple sectors leveraging these technologies to strengthen their real-time cyber threat recognition capabilities. At the enterprise level, deployments often focus on integrating AI-driven models into existing security operations centers (SOCs) to enhance monitoring efficiency, reduce analyst workload, and improve detection accuracy. Large corporations in sectors such as finance, telecommunications, and healthcare are implementing hybrid AI-IDS solutions that combine signature-based recognition for known threats with anomaly detection powered by machine learning and deep learning models for zero-day and emerging attacks (Bolanle & Bamigboye, 2019, Calloway, 2010, Tian, et al., 2019). For instance, a multinational financial services provider may deploy a deep learning-based IDS integrated into its SIEM platform, enabling automated analysis of billions of log entries per day. The system continuously refines its models through adaptive learning, identifying suspicious patterns such as unauthorized database queries or anomalous fund transfer behaviors in near real time. The outcome is a reduction in mean time to detect (MTTD) and mean time to respond (MTTR), enabling faster containment of potential breaches before they escalate into large-scale incidents.

Cloud-based AI-IDS systems are another area where practical deployments have gained traction. With the proliferation of multi-cloud and hybrid cloud environments, traditional perimeter-based security models are no longer sufficient to protect data and applications. Cloud-native AI-IDS solutions are built to operate within these distributed architectures, processing telemetry from virtual machines, containerized workloads, API calls, and cloud-native network flows. One example involves a software-as-a-service (SaaS) provider using a transformer-based IDS that ingests and analyzes real-time traffic metadata from multiple geographic regions. By integrating with cloud provider threat intelligence feeds, the AI system can detect abnormal behaviors, such as unexpected spikes in outbound traffic from specific microservices or repeated failed authentication attempts from unusual geographic locations. The advantage of

deploying AI-IDS in the cloud lies in its elasticity: detection models can scale processing power dynamically in response to traffic volume, ensuring low-latency analysis during peak demand without sacrificing accuracy (Dalal, 2019, Laura & James, 2019, Vinayakumar, Soman & Poornachandran, 2018). These systems also leverage federated learning across distributed cloud nodes to continuously improve threat recognition without transferring sensitive customer data between regions, thus maintaining compliance with data sovereignty regulations.

IoT and critical infrastructure security represent some of the most compelling and high-stakes applications of AI-augmented intrusion detection. Industrial control systems (ICS), smart grids, connected healthcare devices, and transportation networks all face unique vulnerabilities due to the convergence of operational technology (OT) and IT. AI-driven IDS deployments in these contexts must handle heterogeneous data sources, including sensor readings, control commands, network logs, and device-to-device communications, often under stringent latency requirements. In one critical infrastructure deployment, an AI-IDS was integrated into a national power grid control system, using recurrent neural networks (RNNs) to model normal operational sequences and detect deviations that could indicate cyber-physical attacks (He & Kim, 2019, Kolluri, et al., 2016, Mansoor, 2019). By correlating anomalies in control system commands with external threat intelligence, the system was able to flag coordinated intrusion attempts targeting both the IT and OT layers. In healthcare, hospital networks have deployed AI-IDS to monitor IoT-enabled medical devices, detecting suspicious firmware changes or unauthorized access attempts that could compromise patient safety. Similarly, in transportation, AI-IDS has been applied to connected vehicle systems to detect abnormal vehicle-to-infrastructure communications that could signal malicious interference with traffic control systems.

Evaluating the effectiveness of these deployments depends heavily on performance metrics, which are critical for justifying investment, refining models, and ensuring operational reliability. Accuracy, detection rate, false positive rate, and MTTD are among the most widely used indicators. Accuracy measures the

overall proportion of correct classifications both benign and malicious produced by the IDS, serving as a baseline indicator of model reliability. In enterprise deployments, AI-IDS models leveraging ensemble learning have achieved accuracy rates exceeding 95% on real-time traffic, providing confidence that most legitimate activity is correctly classified while malicious actions are flagged for further analysis (Mohammed, 2015, Petrov & Znati, 2018). Detection rate, or true positive rate, specifically measures the proportion of actual threats correctly identified by the system. High detection rates are crucial in environments where even a single missed intrusion can have catastrophic consequences, such as in financial trading platforms or critical infrastructure control networks.

False positive rate remains a critical metric because excessive false alarms can overwhelm security teams, leading to alert fatigue and slower response times. AI-IDS deployments have demonstrated significant improvements in this area compared to traditional systems, with reductions in false positives of up to 40% in some enterprise case studies. Techniques such as behavioral analytics, context-aware detection, and integration with threat intelligence help AI-IDS distinguish between unusual but benign behavior and genuine malicious activity. For example, in a cloud deployment scenario, a sudden surge in outbound API calls from a microservice might be benign if tied to a planned software update; AI models enriched with contextual metadata can recognize this and suppress unnecessary alerts (Gudala, et al., 2019, Konn, 2018, Zhong & Gu, 2019).

MTTD is another critical measure, representing the average time taken to detect a security incident from the moment it occurs. In real-world applications, AI-IDS systems have shown their ability to reduce MTTD from hours or even days to mere seconds or minutes, drastically improving the potential to mitigate damage. In one documented case, a global telecommunications provider reduced its average MTTD from 12 hours to under two minutes after deploying an AI-augmented intrusion detection system integrated with automated triage and response workflows. This reduction was achieved by combining deep learning-based anomaly detection with reinforcement learning-driven policy optimization, enabling the system to rapidly escalate

and respond to high-severity incidents without waiting for manual review (Elish, 2018, Hameed & Suleman, 2019, Hughes, 2015).

Across these case studies, certain common benefits and patterns emerge. First, AI-IDS consistently delivers improvements in early detection, enabling faster containment and reducing the dwell time of attackers in compromised networks. Second, the adaptability of AI models allows for better handling of evolving threats, including zero-day attacks and novel attack vectors, which are often missed by purely signature-based systems. Third, integration with operational workflows whether through SOC dashboards, automated SOAR (Security Orchestration, Automation, and Response) systems, or incident response platforms maximizes the practical impact of these detection improvements. Fourth, AI-IDS deployments often provide valuable secondary benefits, such as enhanced network visibility, improved asset inventory accuracy, and better prioritization of remediation efforts based on contextual risk scoring.

However, real-world applications also reveal ongoing challenges. In enterprise contexts, aligning AI-IDS with compliance requirements and auditability standards can be complex, especially in highly regulated sectors. In cloud deployments, latency, data sovereignty, and multi-tenant security concerns must be balanced with detection performance. In IoT and critical infrastructure settings, resource constraints, proprietary protocols, and the need for deterministic performance can limit the applicability of certain AI techniques (Aisyah, et al., 2019, Gopireddy, 2019, Thangan, Gulhane & Karale, 2019). Performance metrics, while encouraging, must be interpreted carefully, as high accuracy on historical data does not guarantee resilience against adaptive adversaries who may attempt to exploit weaknesses in model design or training data.

Despite these challenges, the trajectory of real-world AI-augmented intrusion detection deployments points to increasing maturity and integration into mainstream cybersecurity strategies. As models become more efficient, explainable, and capable of privacy-preserving learning, their adoption is likely to

accelerate across industries and national security domains. In each of the domains examined enterprise networks, cloud environments, IoT ecosystems, and critical infrastructure the combination of improved detection metrics, reduced false positives, and lower MTTD demonstrates that AI-IDS is not merely an experimental enhancement but a foundational capability for modern threat recognition (De Spiegeleire, Maas & Sweijts, 2017, Hurley, 2018). By continuing to refine their architectures, training methodologies, and operational integrations, organizations can ensure that these systems deliver not only technical excellence but also tangible, measurable improvements in their overall cybersecurity posture.

## 2.8. Conclusion and Future Directions

The growing complexity, scale, and velocity of cyber threats has made AI-augmented intrusion detection systems an indispensable component of modern security architectures. The research and practical applications reviewed demonstrate that these systems can dramatically enhance real-time threat recognition through advanced machine learning and deep learning techniques, hybrid detection models, adaptive learning, and integration with behavioral analytics. By reducing false positives, increasing detection rates, and lowering mean time to detect, AI-IDS solutions significantly strengthen an organization's ability to respond to both known and emerging threats, including zero-day exploits and multi-stage attack campaigns. Their capacity to operate across heterogeneous environments spanning enterprise networks, cloud infrastructures, IoT deployments, and critical infrastructure confirms their versatility and long-term strategic value.

Looking forward, one priority will be the development of energy-efficient AI models for cybersecurity. Current deep learning architectures can be computationally and power-intensive, limiting their deployment in edge environments, IoT ecosystems, and resource-constrained operational technology networks. Research into lightweight architectures, pruning techniques, quantization, and neuromorphic computing offers promising pathways toward

sustainable, high-performance intrusion detection that does not compromise on real-time responsiveness.

Equally important is the deeper integration of AI-IDS with Security Orchestration, Automation, and Response (SOAR) platforms. Such integration enables automated, policy-driven responses that can isolate compromised endpoints, adjust firewall rules, or trigger containment protocols within seconds of detection. By bridging detection with immediate remediation, AI-IDS can close the gap between identifying a threat and neutralizing it, thus minimizing attacker dwell time and reducing potential damage.

To ensure trust, comparability, and accountability, the development of standardized benchmarks for AI-IDS evaluation is critical. Benchmarking datasets, performance metrics, and testing protocols that reflect the complexity of real-world environments will allow organizations to assess solutions on a level playing field and drive innovation through transparent performance comparison.

Finally, cross-industry collaborative threat intelligence sharing will amplify the effectiveness of AI-IDS solutions. By pooling anonymized attack patterns, indicators of compromise, and behavioral signatures, organizations across sectors can train more robust models that detect threats earlier and with greater accuracy. Privacy-preserving mechanisms such as federated learning and secure multi-party computation can facilitate this collaboration without exposing sensitive operational data.

The potential of AI-augmented intrusion detection lies not only in its technical sophistication but also in its adaptability, scalability, and collaborative capabilities. Continued innovation, supported by interdisciplinary research and shared expertise, will be essential to maintaining a decisive edge against rapidly evolving cyber adversaries. As the threat landscape grows more complex, the fusion of advanced AI techniques, operational integration, and cooperative intelligence sharing will define the future of proactive, resilient, and globally coordinated cybersecurity defense.

## REFERENCES

- [1] Achar, S. (2018). Data Privacy-Preservation: A Method of Machine Learning. *ABC Journal of Advanced Research*, 7(2), 123-129.
- [2] AdeniyiAjonbadi, H., AboabaMojeed-Sanni, B., & Otokiti, B. O. (2015). Sustaining competitive advantage in medium-sized enterprises (MEs) through employee social interaction and helping behaviours. *Journal of Small Business and Entrepreneurship*, 3(2), 1-16.
- [3] Adenuga, T., Ayobami, A.T. & Okolo, F.C., 2019. Laying the Groundwork for Predictive Workforce Planning Through Strategic Data Analytics and Talent Modeling. *IRE Journals*, 3(3), pp.159–161. ISSN: 2456-8880.
- [4] Aisyah, N., Hidayat, R., Zulaikha, S., Rizki, A., Yusof, Z. B., Pertiwi, D., & Ismail, F. (2019). Artificial intelligence in cryptographic protocols: Securing e-commerce transactions and ensuring data integrity.
- [5] Ajonbadi, H. A., & Mojeed-Sanni, B. A & Otokiti, BO (2015). ‘Sustaining Competitive Advantage in Medium-sized Enterprises (MEs) through Employee Social Interaction and Helping Behaviours.’. *Journal of Small Business and Entrepreneurship Development*, 3(2), 89-112.
- [6] Ajonbadi, H. A., Lawal, A. A., Badmus, D. A., & Otokiti, B. O. (2014). Financial control and organisational performance of the Nigerian small and medium enterprises (SMEs): A catalyst for economic growth. *American Journal of Business, Economics and Management*, 2(2), 135-143.
- [7] Ajonbadi, H. A., Otokiti, B. O., & Adebayo, P. (2016). The efficacy of planning on organisational performance in the Nigeria SMEs. *European Journal of Business and Management*, 24(3), 25-47.
- [8] Akinbola, O. A., & Otokiti, B. O. (2012). Effects of lease options as a source of finance on profitability performance of small and medium enterprises (SMEs) in Lagos State, Nigeria. *International Journal of Economic Development Research and Investment*, 3(3), 70-76.
- [9] Alhakami, W., ALharbi, A., Bourouis, S., Alroobaea, R., & Bouguila, N. (2019). Network anomaly intrusion detection using a nonparametric Bayesian approach and feature selection. *IEEE access*, 7, 52181-52190.
- [10] Appelt, D., Nguyen, C. D., Panichella, A., & Briand, L. C. (2018). A machine-learning-driven evolutionary approach for testing web application firewalls. *IEEE Transactions on Reliability*, 67(3), 733-757.
- [11] Apruzzese, G., Colajanni, M., Ferretti, L., & Marchetti, M. (2019, May). Addressing adversarial attacks against security systems based on machine learning. In 2019 11th international conference on cyber conflict (CyCon) (Vol. 900, pp. 1-18). IEEE.
- [12] Biggio, B., & Roli, F. (2018, October). Wild patterns: Ten years after the rise of adversarial machine learning. In *Proceedings of the 2018 ACM SIGSAC Conference on Computer and Communications Security* (pp. 2154-2156).
- [13] Bolanle, O., & Bamigboye, K. (2019). AI-Powered Cloud Security: Leveraging Advanced Threat Detection for Maximum Protection. *International Journal of Trend in Scientific Research and Development*, 3(2), 1407-1412.
- [14] Brynskov, M., Facca, F. M., & Hrasko, G. (2018). Next Generation Internet of Things. *H2020 Coordination and Support Action (CSA), NGIoT Consortium, 2021*, 2019.
- [15] Calloway, M. (2010). AI-Powered Threat Detection, Intrusion Prevention, and Network Security. *International Journal of Artificial Intelligence and Machine Learning*, 10(10).
- [16] Chen, T., Liu, J., Xiang, Y., Niu, W., Tong, E., & Han, Z. (2019). Adversarial attack and defense in reinforcement learning-from AI security view. *Cybersecurity*, 2(1), 11.
- [17] Chen, T., Liu, J., Xiang, Y., Niu, W., Tong, E., & Han, Z. (2019). Adversarial attack and defense in reinforcement learning-from AI security view. *Cybersecurity*, 2(1), 11.

- [18] Chen, Z., Yan, Q., Han, H., Wang, S., Peng, L., Wang, L., & Yang, B. (2018). Machine learning based mobile malware detection using highly imbalanced network traffic. *Information Sciences*, 433, 346-364.
- [19] Choraś, M., & Kozik, R. (2015). Machine learning techniques applied to detect cyber attacks on web applications. *Logic Journal of IGPL*, 23(1), 45-56.
- [20] Cybenko, G., Jajodia, S., Wellman, M. P., & Liu, P. (2014, December). Adversarial and uncertain reasoning for adaptive cyber defense: Building the scientific foundation. In *International conference on information systems security* (pp. 1-8). Cham: Springer International Publishing.
- [21] Dalal, A. (2018). Cybersecurity And Artificial Intelligence: How AI Is Being Used in Cybersecurity To Improve Detection And Response To Cyber Threats. *Turkish Journal of Computer and Mathematics Education* Vol, 9(3), 1704-1709.
- [22] Dalal, A. (2019). AI Powered Threat Hunting in SAP and ERP Environments: Proactive Approaches to Cyber Defense. *Available at SSRN 5198746*.
- [23] Dasgupta, P., & Collins, J. (2019). A survey of game theoretic approaches for adversarial machine learning in cybersecurity tasks. *AI Magazine*, 40(2), 31-43.
- [24] De Spiegeleire, S., Maas, M., & Sweijs, T. (2017). *Artificial intelligence and the future of defense: strategic implications for small-and medium-sized force providers*. The Hague Centre for Strategic Studies.
- [25] Dogho, M. (2011). The design, fabrication and uses of bioreactors. Obafemi Awolowo University.
- [26] Duddu, V. (2018). A survey of adversarial machine learning in cyber warfare. *Defence Science Journal*, 68(4), 356.
- [27] Elish, M. C. (2018, October). The stakes of uncertainty: developing and integrating machine learning in clinical care. In *Ethnographic Praxis in Industry Conference Proceedings* (Vol. 2018, No. 1, pp. 364-380).
- [28] Feng, M., & Xu, H. (2017, November). Deep reinforcement learning based optimal defense for cyber-physical system in presence of unknown cyber-attack. In *2017 IEEE Symposium Series on Computational Intelligence (SSCI)* (pp. 1-8). IEEE.
- [29] Gan, J., Li, S., Zhai, Y., & Liu, C. (2017, March). 3d convolutional neural network based on face anti-spoofing. In *2017 2nd international conference on multimedia and image processing (ICMIP)* (pp. 1-5). IEEE.
- [30] Ganesan, R., Jajodia, S., Shah, A., & Cam, H. (2016). Dynamic scheduling of cybersecurity analysts for minimizing risk using reinforcement learning. *ACM Transactions on Intelligent Systems and Technology (TIST)*, 8(1), 1-21.
- [31] Glomsrud, J. A., Ødegårdstuen, A., Clair, A. L. S., & Smogeli, Ø. (2019, September). Trustworthy versus explainable AI in autonomous vessels. In *Proceedings of the International Seminar on Safety and Security of Autonomous Vessels (ISSAV) and European STAMP Workshop and Conference (ESWC)* (Vol. 37).
- [32] Gopireddy, S. R. (2019). AI-Augmented Honeypots for Cloud Environments: Proactive Threat Deception. *European Journal of Advances in Engineering and Technology*, 6(12), 85-89.
- [33] Gudala, L., Shaik, M., Venkataramanan, S., & Sadhu, A. K. R. (2019). Leveraging artificial intelligence for enhanced threat detection, response, and anomaly identification in resource-constrained iot networks. *Distributed Learning and Broad Applications in Scientific Research*, 5, 23-54.
- [34] Hagrass, H. (2018). Toward human-understandable, explainable AI. *Computer*, 51(9), 28-36.
- [35] Hameed, A., & Suleman, M. (2019). AI-Powered Anomaly Detection for Cloud Security: Leveraging Machine Learning and DSPM.
- [36] Han, Y., Rubinstein, B. I., Abraham, T., Alpcan, T., De Vel, O., Erfani, S., ... & Montague, P. (2018, September).

- Reinforcement learning for autonomous defence in software-defined networking. In *International conference on decision and game theory for security* (pp. 145-165). Cham: Springer International Publishing.
- [37] Hao, M., Li, H., Luo, X., Xu, G., Yang, H., & Liu, S. (2019). Efficient and privacy-enhanced federated learning for industrial artificial intelligence. *IEEE Transactions on Industrial Informatics*, 16(10), 6532-6542.
- [38] He, K., & Kim, D. S. (2019, August). Malware detection with malware images using deep learning techniques. In *2019 18th IEEE international conference on trust, security and privacy in computing and communications/13th IEEE international conference on big data science and engineering (TrustCom/BigDataSE)* (pp. 95-102). IEEE.
- [39] Holzinger, A., Kieseberg, P., Weippl, E., & Tjoa, A. M. (2018, August). Current advances, trends and challenges of machine learning and knowledge extraction: from machine learning to explainable AI. In *International cross-domain conference for machine learning and knowledge extraction* (pp. 1-8). Cham: Springer International Publishing.
- [40] Huang, L., & Zhu, Q. (2019, October). Adaptive honeypot engagement through reinforcement learning of semi-markov decision processes. In *International conference on decision and game theory for security* (pp. 196-216). Cham: Springer International Publishing.
- [41] Hughes, E. (2015). AI-Driven Cybersecurity System: Benefits and Vulnerabilities. *International Journal of Artificial Intelligence and Machine Learning*, 6(1).
- [42] Hurley, J. S. (2018). Enabling successful artificial intelligence implementation in the department of defense. *Journal of Information Warfare*, 17(2), 65-82.
- [43] Ibitoye, O., Abou-Khamis, R., Shehaby, M. E., Matrawy, A., & Shafiq, M. O. (2019). The Threat of Adversarial Attacks on Machine Learning in Network Security--A Survey. arXiv preprint arXiv:1911.02621.
- [44] Jaroszewski, A. C., Morris, R. R., & Nock, M. K. (2019). Randomized controlled trial of an online machine learning-driven risk assessment and intervention platform for increasing the use of crisis services. *Journal of consulting and clinical psychology*, 87(4), 370.
- [45] Karatas, G., Demir, O., & Sahingoz, O. K. (2018, December). Deep learning in intrusion detection systems. In *2018 international congress on big data, deep learning and fighting cyber terrorism (IBIGDELFT)* (pp. 113-116). IEEE.
- [46] Kene, S. G., & Theng, D. P. (2015, February). A review on intrusion detection techniques for cloud computing and security challenges. In *2015 2nd International Conference on Electronics and Communication Systems (ICECS)* (pp. 227-232). IEEE.
- [47] Khurana, R., & Kaul, D. (2019). Dynamic cybersecurity strategies for ai-enhanced ecommerce: A federated learning approach to data privacy. *Applied Research in Artificial Intelligence and Cloud Computing*, 2(1), 32-43.
- [48] Kolluri, V. E. N. K. A. T. E. S. W. A. R. A. N. A. I. D. U. (2016). A Pioneering Approach To Forensic Insights: Utilization AI for Cybersecurity Incident Investigations. *IJRAR-International Journal of Research and Analytical Reviews (IJRAR)*, E-ISSN, 2348-1269.
- [49] Konn, A. (2018). Next-Generation Cybersecurity: Harnessing AI for Detecting and Preventing Cyber-Attacks in Cloud Environments.
- [50] Kozik, R., & Choraś, M. (2014). Machine learning techniques for cyber attacks detection. In *Image Processing and Communications Challenges 5* (pp. 391-398). Heidelberg: Springer International Publishing.
- [51] Kumari, M., Hsieh, G., & Okonkwo, C. A. (2017, December). Deep learning approach to malware multi-class classification using image processing techniques. In *2017 International Conference on Computational Science and*

- Computational Intelligence (CSCI)* (pp. 13-18). IEEE.
- [52] Laskov, P., & Lippmann, R. (2010). Machine learning in adversarial environments. *Machine learning*, 81(2), 115-119.
- [53] Laura, M., & James, A. (2019). Cloud Security Mastery: Integrating Firewalls and AI-Powered Defenses for Enterprise Protection. *International Journal of Trend in Scientific Research and Development*, 3(3), 2000-2007.
- [54] Lawal, A. A., Ajonbadi, H. A., & Otokiti, B. O. (2014). Leadership and organisational performance in the Nigeria small and medium enterprises (SMEs). *American Journal of Business, Economics and Management*, 2(5), 121.
- [55] Lawal, A. A., Ajonbadi, H. A., & Otokiti, B. O. (2014). Strategic importance of the Nigerian small and medium enterprises (SMES): Myth or reality. *American Journal of Business, Economics and Management*, 2(4), 94-104.
- [56] Lawless, W. F., Mittu, R., Sofge, D., & Hiatt, L. (2019). Artificial intelligence, autonomy, and human-machine teams interdependence, context, and explainable AI. *Ai Magazine*, 40(3), 5-13.
- [57] Liao, R., Wen, H., Pan, F., Song, H., Xu, A., & Jiang, Y. (2019, March). A novel physical layer authentication method with convolutional neural network. In *2019 IEEE International Conference on Artificial Intelligence and Computer Applications (ICAICA)* (pp. 231-235). IEEE.
- [58] Liu, Q., Li, P., Zhao, W., Cai, W., Yu, S., & Leung, V. C. (2018). A survey on security threats and defensive techniques of machine learning: A data driven view. *IEEE access*, 6, 12103-12117.
- [59] Mansoor, A. (2019). Mitigating Cyber-Attacks with AI-Driven Cybersecurity Solutions in Cloud and Device Technologies.
- [60] Masoud, M., Jaradat, Y., & Ahmad, A. Q. (2016, December). On tackling social engineering web phishing attacks utilizing software defined networks (SDN) approach. In *2016 2nd International Conference on Open Source Software Computing (OSSCOM)* (pp. 1-6). IEEE.
- Manickam, M., Ramaraj, N., & Chellappan, C. (2019). A combined PFCM and recurrent neural network-based intrusion detection system for cloud environment. *International Journal of Business Intelligence and Data Mining*, 14(4), 504-527.
- [61] Mavroeidis, V., & Bromander, S. (2017, September). Cyber threat intelligence model: an evaluation of taxonomies, sharing standards, and ontologies within cyber threat intelligence. In *2017 European Intelligence and Security Informatics Conference (EISIC)* (pp. 91-98). IEEE.
- [62] Menson, W. N. A., Olawepo, J. O., Bruno, T., Gbadamosi, S. O., Nalda, N. F., Anyebe, V., ... & Ezeanolue, E. E. (2018). Reliability of self-reported Mobile phone ownership in rural north-Central Nigeria: cross-sectional study. *JMIR mHealth and uHealth*, 6(3), e8760.
- [63] Mittal, S., Joshi, A., & Finin, T. (2019). Cyber-all-intel: An ai for security related threat intelligence. *arXiv preprint arXiv:1905.02895*.
- [64] Mohammad, R. M., Thabtah, F., & McCluskey, L. (2014). Predicting phishing websites based on self-structuring neural network. *Neural Computing and Applications*, 25(2), 443-458.
- [65] Mohammed, I. A. (2015). A technical and state-of-the-art assessment of machine learning algorithms for cybersecurity applications. *International Journal of Current Science (IJCS PUB)* [www.ijcspub.org](http://www.ijcspub.org), ISSN, 2250-1770.
- [66] Mohit, M. (2018). Federated Learning: An Intrusion Detection Privacy Preserving Approach to Decentralized AI Model Training for IOT Security.
- [67] Mustapha, A. Y., Chianumba, E. C., Forkuo, A. Y., Osamika, D., & Komi, L. S. (2018). Systematic Review of Mobile Health (mHealth) Applications for Infectious Disease Surveillance in Developing Countries. *Methodology*, 66.
- [68] Nauman, M., Tanveer, T. A., Khan, S., & Syed, T. A. (2018). Deep neural architectures for

- large scale android malware analysis. *Cluster Computing*, 21(1), 569-588.
- [69] Nsa, B., Anyebe, V., Dimkpa, C., Aboki, D., Egbule, D., Useni, S., & Eneogu, R. (2018, November). Impact of active case finding of tuberculosis among prisoners using the WOW truck in North Central Nigeria. *The International Journal of Tuberculosis and Lung Disease*, 22(11), S444. The International Union Against Tuberculosis and Lung Disease.
- [70] Ogungbenle, H. N., & Omowole, B. M. (2012). Chemical, functional and amino acid composition of periwinkle (*Tympanotonus fuscatus* var *radula*) meat. *Int J Pharm Sci Rev Res*, 13(2), 128-132.
- [71] Olasehinde, O. (2018, December). Stock price prediction system using long short-term memory. *BlackInAI Workshop @ NeurIPS 2018*.
- [72] Oni, O., Adeshina, Y. T., Iloeje, K. F., & Olatunji, O. O. (2018). Artificial Intelligence Model Fairness Auditor For Loan Systems. *Journal ID*, 8993, 1162.
- [73] Orren, D. (2019). *Safe Employment of Augmented Reality in a Production Environment Final Report* (No. ONROLCVA).
- [74] O'Sullivan, S., Nevejans, N., Allen, C., Blyth, A., Leonard, S., Pagallo, U., ... & Ashrafian, H. (2019). Legal, regulatory, and ethical frameworks for development of standards in artificial intelligence (AI) and autonomous robotic surgery. *The international journal of medical robotics and computer assisted surgery*, 15(1), e1968.
- [75] Otokiti, B. O. (2012). Mode of entry of multinational corporation and their performance in the Nigeria market (Doctoral dissertation, Covenant University).
- [76] Otokiti, B. O. (2018). Business regulation and control in Nigeria. *Book of readings in honour of Professor SO Otokiti*, 1(2), 201-215.
- [77] Otokiti, B. O., & Akorede, A. F. (2018). Advancing sustainability through change and innovation: A co-evolutionary perspective. *Innovation: Taking creativity to the market. Book of Readings in Honour of Professor SO Otokiti*, 1(1), 161-167.
- [78] Otoum, S. (2019). *Machine learning-driven intrusion detection techniques in critical infrastructures monitored by sensor networks* (Doctoral dissertation, Université d'Ottawa/University of Ottawa).
- [79] Pauwels, E., & Denton, S. W. (2018). Searching for privacy in the Internet of Bodies. *The Wilson Quarterly*, 42(2).
- [80] Perumallapalli, R. (2017). *Federated Learning Applications in Enterprise Network Management*. Available at SSRN 5228699.
- [81] Petrov, D., & Znati, T. (2018, October). Context-aware deep learning-driven framework for mitigation of security risks in BYOD-enabled environments. In *2018 IEEE 4th International Conference on Collaboration and Internet Computing (CIC)* (pp. 166-175). IEEE.
- [82] Pham, C., Nguyen, L. A., Tran, N. H., Huh, E. N., & Hong, C. S. (2018). Phishing-aware: A neuro-fuzzy approach for anti-phishing on fog networks. *IEEE Transactions on Network and Service Management*, 15(3), 1076-1089.
- [83] Preuveneers, D., Rimmer, V., Tsingenopoulos, I., Spooren, J., Joosen, W., & Ilie-Zudor, E. (2018). Chained anomaly detection models for federated learning: An intrusion detection case study. *Applied Sciences*, 8(12), 2663.
- [84] Renda, A. (2019). The age of foodtech: Optimizing the agri-food chain with digital technologies. In *Achieving the sustainable development goals through sustainable food systems* (pp. 171-187). Cham: Springer International Publishing.
- [85] Ridley, A. (2018). Machine learning for autonomous cyber defense. *The Next Wave*, 22(1), 7-14.
- [86] Sahingoz, O. K., Baykal, S. I., & Bulut, D. (2018). Phishing detection from urls by using neural networks. *Computer Science & Information Technology (CS & IT)*, 41-54.
- [87] Sahingoz, O. K., Buber, E., Demir, O., & Diri, B. (2019). Machine learning based phishing

- detection from URLs. *Expert Systems with Applications*, 117, 345-357.
- [88] Sareddy, M. R., & Hemnath, R. (2019). Optimized federated learning for cybersecurity: Integrating split learning, graph neural networks, and hashgraph technology. *International Journal of HRM and Organizational Behavior*, 7(3), 43-54.
- [89] Scholten, J., Eneogu, R., Ogbudebe, C., Nsa, B., Anozie, I., Anyebe, V., Lawanson, A., & Mitchell, E. (2018, November). Ending the TB epidemic: Role of active TB case finding using mobile units for early diagnosis of tuberculosis in Nigeria. *The International Journal of Tuberculosis and Lung Disease*, 22(11), S392. The International Union Against Tuberculosis and Lung Disease.
- [90] Sethi, T. S., Kantardzic, M., Lyu, L., & Chen, J. (2018). A dynamic-adversarial mining approach to the security of machine learning. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 8(3), e1245.
- [91] Shah, H. (2017). Deep Learning in Cloud Environments: Innovations in AI and Cybersecurity Challenges. *Revista Espanola de Documentacion Cientifica*, 11(1), 146-160.
- [92] Sharma, A., Adekunle, B. I., Ogeawuchi, J. C., Abayomi, A. A., & Onifade, O. (2019). IoT-enabled Predictive Maintenance for Mechanical Systems: Innovations in Real-time Monitoring and Operational Excellence.
- [93] Shi, Y., Sagduyu, Y. E., Davaslioglu, K., & Levy, R. (2018). Vulnerability detection and analysis in adversarial deep learning. In *Guide to vulnerability analysis for computer networks and systems: An artificial intelligence approach* (pp. 211-234). Cham: Springer International Publishing.
- [94] Smadi, S., Aslam, N., & Zhang, L. (2018). Detection of online phishing email using dynamic evolving neural network based on reinforcement learning. *Decision Support Systems*, 107, 88-102.
- [95] Sowah, R. A., Ofori-Amanfo, K. B., Mills, G. A., & Koumadi, K. M. (2019). Detection and prevention of man-in-the-middle spoofing attacks in MANETs using predictive techniques in artificial neural networks (ANN). *Journal of Computer Networks and Communications*, 2019(1), 4683982.
- [96] Su, X., Zhang, D., Li, W., & Zhao, K. (2016, August). A deep learning approach to android malware feature learning and detection. In *2016 IEEE Trustcom/BigDataSE/ISPA* (pp. 244-251). IEEE.
- [97] Svenmarck, P., Luotsinen, L., Nilsson, M., & Schubert, J. (2018, May). Possibilities and challenges for artificial intelligence in military applications. In *Proceedings of the NATO big data and artificial intelligence for military decision making specialists' meeting* (Vol. 1).
- [98] Thangan, M. S. S., Gulhane, V. S., & Karale, N. E. (2019). Review on "Using Big Data to Defend Machines against Network Attacks".
- [99] Tian, Z., Luo, C., Qiu, J., Du, X., & Guizani, M. (2019). A distributed deep learning system for web attack detection on edge devices. *IEEE Transactions on Industrial Informatics*, 16(3), 1963-1971.
- [100] Tobiyama, S., Yamaguchi, Y., Shimada, H., Ikuse, T., & Yagi, T. (2016, June). Malware detection with deep neural network using process behavior. In *2016 IEEE 40th annual computer software and applications conference (COMPSAC)* (Vol. 2, pp. 577-582). IEEE.
- [101] Vinayakumar, R., Alazab, M., Soman, K. P., Poornachandran, P., & Venkatraman, S. (2019). Robust intelligent malware detection using deep learning. *IEEE access*, 7, 46717-46738.
- [102] Vinayakumar, R., Soman, K. P., & Poornachandran, P. (2018). Detecting malicious domain names using deep learning approaches at scale. *Journal of Intelligent & Fuzzy Systems*, 34(3), 1355-1367.
- [103] Weng, J., Weng, J., Zhang, J., Li, M., Zhang, Y., & Luo, W. (2019). Deepchain: Auditable and privacy-preserving deep learning with blockchain-based incentive. *IEEE Transactions on Dependable and Secure Computing*, 18(5), 2438-2455.

- [104] Xu, G., Li, H., Liu, S., Yang, K., & Lin, X. (2019). VerifyNet: Secure and verifiable federated learning. *IEEE Transactions on Information Forensics and Security*, 15, 911-926.
- [105] Yarali, A., Ramage, M. L., May, N., & Srinath, M. (2019, April). Uncovering the true potentials of the internet of things (IoT). In *2019 Wireless Telecommunications Symposium (WTS)* (pp. 1-6). IEEE.
- [106] Zhang, C., Patras, P., & Haddadi, H. (2019). Deep learning in mobile and wireless networking: A survey. *IEEE Communications surveys & tutorials*, 21(3), 2224-2287.
- [107] Zhong, W., & Gu, F. (2019). A multi-level deep learning system for malware detection. *Expert Systems with Applications*, 133, 151-162.
- [108] Zhou, P., Wang, K., Guo, L., Gong, S., & Zheng, B. (2019). A privacy-preserving distributed contextual federated online learning framework with big data support in social recommender systems. *IEEE Transactions on Knowledge and Data Engineering*, 33(3), 824-838.
- [109] Zhu, M., Hu, Z., & Liu, P. (2014, November). Reinforcement learning algorithms for adaptive cyber defense against heartbleed. In *Proceedings of the first ACM workshop on moving target defense* (pp. 51-58).