

A Machine Learning Based Framework for Intelligent Cybersecurity Threat Detection and Prevention

Akash Jain

Abstract-- The rapid evolution of cyber threats has rendered traditional signature-based and rule-driven security systems increasingly inadequate. Machine learning (ML) and deep learning (DL) techniques have emerged as transformative enablers of intelligent, adaptive, and proactive cybersecurity frameworks capable of detecting novel and sophisticated attacks in real time. This survey provides a comprehensive review of ML-based approaches applied to cybersecurity threat detection and prevention, covering supervised, unsupervised, semi-supervised, and reinforcement learning paradigms. We systematically analyze major threat categories including network intrusions, malware, ransomware, distributed denial-of-service (DDoS) attacks, phishing, and advanced persistent threats (APTs). We examine widely used benchmark datasets, feature engineering strategies, model architectures, and performance evaluation metrics. Special attention is given to emerging directions such as federated learning for privacy-preserving threat intelligence, generative adversarial networks for synthetic attack data generation, and explainable AI (XAI) for transparent decision-making in security operations. We further discuss the key challenges of class imbalance, adversarial evasion attacks against ML models, concept drift, and scalability constraints. The survey concludes with an open research agenda identifying critical gaps and future opportunities in the design of robust, explainable, and operationally deployable ML-driven cybersecurity systems.

Keywords-- Machine Learning, Cybersecurity, Intrusion Detection System, Deep Learning, Anomaly Detection, Threat Intelligence, Federated Learning, Explainable AI, Malware Detection, Adversarial Robustness

I. INTRODUCTION

Cybersecurity has emerged as one of the most critical challenges of the digital age. The global cost of cybercrime exceeded \$8 trillion in 2023 and is projected to surpass \$10.5 trillion annually by 2025 (Cybersecurity Ventures, 2023). Organizations across all sectors face an unprecedented volume and diversity of cyber attacks, ranging from opportunistic malware campaigns to highly sophisticated nation-state sponsored intrusions. Traditional cybersecurity defenses, including firewalls, antivirus software, and signature-based intrusion detection systems (IDS), rely on pre-defined rules and known threat signatures.

While effective against established attack patterns, these approaches are fundamentally reactive and incapable of detecting zero-day exploits, polymorphic malware, or novel attack vectors that deviate from known signatures.

Machine learning and deep learning have fundamentally changed the paradigm of threat detection. By learning discriminative patterns from large volumes of network traffic, system logs, and behavioral data, ML-based systems can identify anomalous activity indicative of attacks without requiring explicit programming for each threat type. The ability to generalize from historical data to novel threats gives ML-based security frameworks a significant advantage over rule-based counterparts. Furthermore, advances in computational infrastructure — including graphics processing units (GPUs), cloud computing, and edge AI — have made real-time ML inference increasingly feasible in production security environments.

Despite the substantial promise of ML in cybersecurity, significant challenges remain. These include the scarcity of labeled training data, severe class imbalance between benign and malicious samples, high false alarm rates that fatigue security analysts, model opacity that hinders trust and interpretability, and susceptibility to adversarial attacks specifically designed to fool ML classifiers. Addressing these challenges requires an interdisciplinary research agenda that spans machine learning, network security, systems engineering, and human factors.

This survey makes the following contributions: (1) a structured taxonomy of ML techniques applied across the cybersecurity threat landscape; (2) a systematic comparison of benchmark datasets and their suitability for different threat detection tasks; (3) a critical analysis of model performance, limitations, and deployment considerations; (4) a discussion of cutting-edge directions including federated learning, XAI, and adversarial robustness; and (5) a forward-looking research agenda identifying open problems and promising solutions.

The remainder of this paper is organized as follows. Section 2 presents the threat landscape. Section 3 reviews ML paradigms used in cybersecurity. Section 4 discusses datasets and feature engineering.



International Journal of Recent Development in Engineering and Technology
Website: www.ijrdet.com (ISSN 2347-6435 (Online) Volume 15, Issue 08, August 2026)

Section 5 presents performance benchmarks. Section 6 covers emerging and advanced techniques. Section 7 addresses open challenges. Section 8 presents the research agenda and Section 9 concludes.

II. THE CYBER THREAT LANDSCAPE

2.1 Taxonomy of Cyber Threats

Modern cyber threats can be classified along several dimensions: attack vector (network, application, insider), attack intent (espionage, sabotage, financial gain), persistence (transient vs. advanced persistent), and automation level (manual, semi-automated, fully automated). Understanding this taxonomy is essential for selecting appropriate ML strategies.

2.2 Network Intrusion and Scanning

Network intrusions encompass unauthorized attempts to access, disrupt, or exfiltrate data from networked systems. Common forms include port scanning, brute-force authentication attacks, man-in-the-middle (MITM) attacks, and lateral movement within enterprise networks. Traditional anomaly-based IDS systems suffer from high false positive rates; ML models trained on labeled network flow data have demonstrated significant improvements in detection precision and recall.

2.3 Malware and Ransomware

Malware — including viruses, worms, trojans, spyware, and ransomware — represents one of the most persistent and economically damaging threat categories. Ransomware attacks alone caused an estimated \$20 billion in damages in 2021. Static analysis of malware binaries using feature extraction from portable executable (PE) headers combined with ML classifiers has been widely studied. Dynamic behavioral analysis using sandboxed execution combined with recurrent neural networks has shown promise for detecting obfuscated and polymorphic malware that evades static analysis.

2.4 Distributed Denial of Service (DDoS)

DDoS attacks overwhelm network infrastructure by flooding targets with traffic from compromised botnets, rendering services unavailable to legitimate users. The rise of IoT devices has dramatically expanded the attack surface for DDoS amplification. ML-based detection exploits the statistical regularity of DDoS traffic patterns — including packet rate anomalies, entropy changes in source IP distributions, and protocol feature deviations — to distinguish attack traffic from legitimate surges.

2.5 Phishing and Social Engineering

Phishing attacks exploit human psychology to harvest credentials and deploy malware through deceptive emails, web pages, and messaging platforms. URL-based phishing detection using natural language processing (NLP) and graph-based analysis of domain relationships has been extensively researched. Machine learning models trained on lexical, DNS, and WHOIS features of URLs have achieved detection accuracies exceeding 97% on benchmark datasets.

2.6 Advanced Persistent Threats (APTs)

APTs represent the most sophisticated category of cyber attacks, characterized by prolonged, stealthy intrusion campaigns sponsored by well-resourced adversaries. APT actors use a combination of zero-day exploits, spear phishing, living-off-the-land techniques, and custom malware. Detecting APTs requires behavioral analytics over extended time horizons, threat intelligence correlation, and user entity behavior analytics (UEBA), all of which are natural applications of ML.

III. MACHINE LEARNING PARADIGMS IN CYBERSECURITY

3.1 Supervised Learning

Supervised learning algorithms, including decision trees, random forests, support vector machines (SVM), k-nearest neighbors (k-NN), and neural networks, learn classification or regression functions from labeled training examples. In the cybersecurity context, supervised models are trained on datasets where network flows, file samples, or log entries are labeled as either benign or belonging to specific attack categories. Random forests have been particularly popular due to their robustness to noise, resistance to overfitting, and inherent feature importance estimation. Deep learning architectures such as convolutional neural networks (CNN) and long short-term memory (LSTM) networks have achieved state-of-the-art results on malware classification and network anomaly detection tasks respectively.

3.2 Unsupervised Learning

Unsupervised methods, including k-means clustering, DBSCAN, isolation forests, and autoencoders, do not require labeled training data. This is particularly advantageous in cybersecurity, where labeling large volumes of network traffic is both costly and time-consuming. Anomaly detection approaches model the statistical distribution of normal behavior and flag significant deviations as potential threats.

Autoencoders learn compact representations of normal activity; high reconstruction error signals anomalous input. These methods are especially effective for zero-day threat detection, where attack signatures are unknown a priori.

3.3 Semi-Supervised and Self-Supervised Learning

Semi-supervised learning bridges the gap between supervised and unsupervised approaches by leveraging a small set of labeled examples alongside large amounts of unlabeled data. Techniques such as label propagation, generative semi-supervised models, and consistency regularization have been applied to network intrusion detection to reduce dependence on expensive manual annotation. Self-supervised learning, including contrastive learning approaches, has recently shown strong results for pre-training security-relevant representations from raw log data.

3.4 Reinforcement Learning

Reinforcement learning (RL) models sequential decision-making processes through interaction with an environment, receiving reward signals for beneficial actions.

In cybersecurity, RL has been applied to adaptive firewall rule generation, automated penetration testing, intelligent honeypot management, and dynamic network reconfiguration in response to detected threats. Deep reinforcement learning combines RL with deep neural networks to handle high-dimensional state spaces typical of complex network environments.

3.5 Transfer Learning and Domain Adaptation

Transfer learning leverages knowledge acquired on one task or domain to improve performance on a related but different task. In cybersecurity, pre-trained models on large-scale network datasets can be fine-tuned for organization-specific threat detection with minimal additional labeled data. Domain adaptation techniques address distribution shifts between training and deployment environments, a common challenge when deploying models trained on public benchmarks to real-world enterprise traffic.

Table 1 below summarizes the primary ML techniques surveyed, their cybersecurity applications, key advantages, and principal limitations.

Table 1: Summary of Machine Learning Techniques for Cybersecurity

ML Technique	Primary Application	Advantages	Limitations
Decision Trees / Random Forest	Intrusion detection, malware classification	High interpretability, fast training	Moderate scalability
Support Vector Machines (SVM)	Anomaly detection, spam filtering	Effective in high-dim spaces	Slow on large datasets
Artificial Neural Networks (ANN)	Pattern recognition, threat scoring	Learns complex mappings	Needs large labeled data
Convolutional Neural Networks (CNN)	Malware image classification, log analysis	Excellent feature extraction	High computational cost
Recurrent Neural Networks / LSTM	Network traffic sequence analysis	Captures temporal dependencies	Vanishing gradient issues
Autoencoders	Anomaly/zero-day detection	Unsupervised, no labeled data	High false positives
Generative Adversarial Networks (GAN)	Synthetic attack data generation	Data augmentation for rare attacks	Training instability
Transformer / BERT	Threat intelligence NLP, log parsing	Long-range context modeling	Very high resource demand
Federated Learning	Privacy-preserving IDS across orgs	Preserves data privacy	Communication overhead
Reinforcement Learning	Adaptive firewall rule generation	Self-improving policy	Reward shaping difficulty

IV. DATASETS AND FEATURE ENGINEERING

4.1 Benchmark Datasets

The availability of high-quality, labeled datasets is foundational to the development and evaluation of ML-

based cybersecurity systems. Table 2 provides a comparative overview of widely used datasets across network intrusion, malware, and phishing detection domains.

Table 2: Benchmark Datasets Used in ML-Based Cybersecurity Research

Dataset	Domain	Size	Notes
KDD Cup 1999	Network intrusion	~5M records	Outdated, high redundancy
NSL-KDD	Network intrusion	~148K records	Imbalanced classes
CICIDS-2017	Multi-attack network traffic	~2.8M records	Benchmark standard
UNSW-NB15	Network anomaly	~2.5M records	Modern attack types
EMBER	PE malware	1M samples	Binary/multi-class
MallImg	Malware visualization	9,339 images	25 malware families
CAIDA	DDoS/traffic analysis	Packet-level	Access-restricted
PhishTank	Phishing URLs	Millions of URLs	Community-labeled

The KDD Cup 1999 dataset, while historically important, contains significant redundancy and no longer reflects modern attack patterns. The CICIDS-2017 dataset, generated by the Canadian Institute for Cybersecurity, has become the de facto benchmark for multi-class intrusion detection owing to its realistic traffic profiles and diverse attack categories. The UNSW-NB15 dataset, designed to address weaknesses of earlier benchmarks, incorporates modern attack types including backdoors, fuzzers, exploits, and reconnaissance probes.

4.2 Feature Engineering for Network Intrusion Detection

Effective feature engineering is critical to model performance. Network flow features commonly used include: packet counts and byte counts per unit time, inter-arrival time statistics, TCP flag distributions, port numbers, and protocol fields. Tools such as CICFlowMeter automate the extraction of 80+ statistical flow features from raw pcap captures. Dimensionality reduction techniques including principal component analysis (PCA), mutual information-based feature selection, and recursive feature elimination help mitigate the curse of dimensionality and reduce computational overhead.

4.3 Feature Representations for Malware Analysis

Malware analysis features span static, dynamic, and behavioral domains. Static features include API call n-grams from disassembled binary code, opcodes, entropy measurements, and PE header metadata. Dynamic features are extracted from sandbox execution traces and include system call sequences, registry modifications, network connections, and file system operations. Malware visualization converts binary files to grayscale images, enabling the application of CNN-based image classification techniques that have achieved high accuracy on the MallImg dataset.

4.4 Class Imbalance and Data Augmentation

A persistent challenge in cybersecurity datasets is severe class imbalance: benign traffic typically vastly outnumbers attack traffic, often by ratios exceeding 100:1. Standard ML classifiers trained on imbalanced data tend to be biased toward the majority class, resulting in high false negative rates for attack detection. Mitigation strategies include random oversampling, SMOTE (Synthetic Minority Over-sampling Technique), cost-sensitive learning, and GAN-based synthetic attack sample generation. Federated learning architectures also help by pooling rare attack examples across organizational boundaries without sharing raw data.

V. PERFORMANCE EVALUATION AND COMPARATIVE ANALYSIS

5.1 Evaluation Metrics

Standard evaluation metrics for cybersecurity ML models include accuracy, precision, recall (detection rate), F1-score, false alarm rate (FAR), and area under the ROC curve (AUC-ROC). The FAR, also known as the false positive rate, is of particular operational significance: in high-throughput enterprise environments, even a 1% FAR translates to thousands of daily false alerts, overwhelming

security operations centers (SOCs). Matthews Correlation Coefficient (MCC) is increasingly recommended as a balanced metric for imbalanced datasets.

5.2 Comparative Performance of ML Algorithms

Table 3 presents a comparative summary of representative ML and DL algorithms across different datasets, reflecting results reported in recent literature. All results are indicative and sourced from peer-reviewed publications. FAR = False Alarm Rate.

Table 3: Performance Comparison of ML Algorithms on Cybersecurity Datasets

Algorithm	Dataset	Accuracy	F1-Score	FAR	Complexity
Random Forest	NSL-KDD	99.1%	98.8%	0.3%	Low
SVM	UNSW-NB15	96.4%	95.9%	1.1%	Low
CNN	Mallmg	98.5%	98.2%	0.8%	Medium
LSTM	CICIDS-2017	97.8%	97.3%	1.4%	High
Autoencoder	CICIDS-2017	94.2%	93.8%	4.2%	Low
GAN + CNN	EMBER	99.3%	99.0%	0.5%	High
Transformer	NSL-KDD	99.5%	99.2%	0.4%	Very High
Federated RF	CIC+NSL	97.1%	96.6%	1.8%	Medium

The results in Table 3 demonstrate that ensemble methods such as Random Forest consistently achieve high accuracy and low FAR on established benchmarks. Deep learning architectures including CNN, LSTM, and Transformer-based models achieve the highest detection rates on complex tasks such as malware image classification and sequential traffic analysis, at the cost of substantially higher computational requirements. Autoencoder-based anomaly detection, while valuable for zero-day threat detection, tends to exhibit higher false alarm rates. GAN-augmented training pipelines show competitive performance, particularly on imbalanced datasets where minority attack class representation is sparse.

5.3 Limitations of Benchmark Evaluations

It is important to note that performance figures reported on public benchmarks often do not transfer directly to real-world deployments.

Factors contributing to the performance gap include distribution shift between benchmark and production traffic, the absence of encrypted traffic (which now constitutes over 80% of enterprise network flows), benchmark contamination through data leakage, and the lack of temporal validation splits that simulate concept drift. Researchers and practitioners should interpret benchmark results with these caveats in mind and prioritize evaluation on organization-specific traffic samples.

VI. EMERGING AND ADVANCED TECHNIQUES

6.1 Federated Learning for Privacy-Preserving Threat Intelligence

Federated learning (FL) enables multiple organizations to collaboratively train a shared ML model without exchanging raw data. This is highly relevant to cybersecurity, where organizations are often reluctant to share sensitive network telemetry due to regulatory, competitive, and privacy concerns.

In the FL paradigm, each participant trains a local model on its own data and shares only model gradients or parameters with a central aggregator. The aggregator combines updates using algorithms such as FedAvg to produce a global model that benefits from the collective threat intelligence of all participants. FL has been successfully applied to intrusion detection, malware classification, and phishing detection across healthcare, financial, and critical infrastructure sectors.

6.2 Generative Adversarial Networks for Threat Augmentation

Generative adversarial networks (GANs) consist of a generator network that learns to produce synthetic samples and a discriminator network that attempts to distinguish real from synthetic samples. In cybersecurity, GANs serve two complementary roles: (1) data augmentation — generating synthetic attack traffic to address class imbalance and improve minority class detection; and (2) adversarial attack simulation — generating adversarial network flows that evade detection to evaluate model robustness. Conditional GANs (cGANs) and Wasserstein GANs (WGANs) have been particularly effective for generating realistic attack traffic with controlled attribute distributions.

6.3 Explainable AI (XAI) for Security Operations

The opacity of complex ML models, particularly deep neural networks, is a significant barrier to their adoption in high-stakes security decision-making. Security analysts require not only accurate alerts but also interpretable explanations that enable rapid triage, root cause analysis, and incident response. Explainability techniques including LIME (Local Interpretable Model-Agnostic Explanations), SHAP (SHapley Additive exPlanations), and attention visualization for transformer models have been integrated into security analytics platforms to provide feature-level explanations for individual predictions. Explainability also plays a critical role in detecting model bias and validating that detection decisions are based on semantically meaningful security features rather than spurious correlations.

6.4 Graph Neural Networks for Threat Intelligence

Cyber threats frequently manifest as structured relational patterns — attack graphs, IP communication graphs, provenance graphs of system call chains, and knowledge graphs of threat intelligence. Graph neural networks (GNNs) are designed to learn representations of graph-structured data and have been applied to malware detection through code call graphs, lateral movement detection through enterprise network communication graphs, and phishing campaign attribution through domain registrar and hosting provider networks.

GNNs offer a natural inductive bias for security tasks involving relational reasoning that flat feature vector representations cannot capture.

6.5 Large Language Models and Transformer Architectures

The transformer architecture, originally developed for natural language processing, has found increasing application in cybersecurity. Pre-trained language models such as BERT and GPT have been fine-tuned for vulnerability report classification, threat intelligence extraction from unstructured text, log anomaly detection, and command-line shellcode analysis. Models such as SecurityBERT and CyBERT incorporate domain-specific pre-training on cybersecurity corpora to improve performance on security NLP tasks. The ability of large language models (LLMs) to reason over complex contextual relationships positions them as powerful tools for analyst augmentation in threat hunting and incident response.

6.6 Threat Hunting and Automated Response

Beyond passive detection, ML-based frameworks are increasingly being used to actively hunt for threats and automate response actions. Automated threat hunting combines unsupervised anomaly detection with threat intelligence enrichment to proactively search for indicators of compromise (IoCs) and tactics, techniques, and procedures (TTPs) within enterprise environments. Security orchestration, automation, and response (SOAR) platforms integrate ML-driven alert prioritization with automated playbook execution to reduce mean time to detect (MTTD) and mean time to respond (MTTR). Reinforcement learning agents trained in simulated cyber ranges are beginning to demonstrate the ability to autonomously remediate detected intrusions in controlled environments.

VII. OPEN CHALLENGES AND LIMITATIONS

7.1 Adversarial Robustness

ML models for cybersecurity are themselves vulnerable to adversarial attacks — carefully crafted perturbations to input data designed to cause misclassification. Adversarial examples in the network traffic domain involve manipulating packet features or timing to evade anomaly-based detectors. In the malware domain, adversarial perturbations to binary code — such as adding benign API calls or appending code segments — can fool ML classifiers while preserving malicious functionality. Defenses including adversarial training, certified robustness through randomized smoothing, and ensemble diversity have been proposed but the arms race between attackers and defenders remains an active research frontier.



7.2 Concept Drift

Cyber threat landscapes evolve continuously as attackers adopt new tools, tactics, and infrastructure. An ML model trained on historical attack data rapidly becomes stale as the statistical properties of attack traffic shift over time — a phenomenon known as concept drift. Detecting and adapting to concept drift requires online or incremental learning algorithms that can update model parameters in response to new data streams. Drift detection algorithms including Page-Hinkley, ADWIN, and drift-aware ensemble methods have been proposed for cybersecurity applications, but operational deployment of adaptive models remains challenging.

7.3 Scalability and Real-Time Processing

Enterprise networks generate enormous volumes of traffic — often hundreds of gigabits per second at major network interconnects. Real-time ML-based analysis of high-throughput traffic requires highly optimized inference pipelines, hardware acceleration through GPUs and FPGAs, and efficient stream processing architectures. Lightweight model architectures, knowledge distillation, and model quantization techniques are being actively developed to reduce inference latency and memory footprint without sacrificing detection accuracy.

7.4 Encrypted Traffic Analysis

The widespread adoption of TLS encryption — now encompassing over 80% of enterprise traffic — severely limits the availability of payload-level features traditionally used for signature-based detection. ML-based analysis of encrypted traffic must rely on flow-level metadata such as packet lengths, inter-arrival times, TLS handshake parameters, and certificate properties. While these features carry rich behavioral signatures, they are substantially less discriminative than payload content, necessitating more sophisticated modeling approaches including traffic fingerprinting and side-channel analysis.

7.5 Interpretability and Trust

Deploying ML models in security-critical contexts requires a high degree of trust from human operators. Black-box models that provide accurate alerts without explanations are difficult to validate, audit, and defend in regulatory or legal contexts. Building trusted ML security systems requires progress on multiple fronts: interpretable model architectures, post-hoc explanation methods, calibrated uncertainty estimation, and rigorous testing and red-teaming protocols.

VIII. RESEARCH AGENDA AND FUTURE DIRECTIONS

Based on the survey findings and identified challenges, we propose the following priority research directions for the ML-based cybersecurity community:

- *Adversarially Robust Models:* Development of ML architectures with provable robustness guarantees against domain-specific adversarial perturbations in both network traffic and malware domains.
- *Continual and Lifelong Learning:* Creation of security ML systems capable of learning from new threat intelligence incrementally without catastrophic forgetting of prior knowledge.
- *Encrypted Traffic Intelligence:* Novel feature engineering and deep representation learning methods for accurate threat detection within TLS-encrypted traffic without decryption.
- *Privacy-Preserving Collaborative Defense:* Advancement of federated learning, secure multi-party computation, and differential privacy techniques for cross-organizational threat intelligence sharing.
- *Causality and Root Cause Analysis:* Integration of causal inference methods with ML to support automated root cause attribution and actionable incident response recommendations.
- *Human-AI Collaborative Security:* Design of human-in-the-loop ML systems that efficiently leverage analyst expertise for model updating, anomaly triage, and alert disposition.
- *Large Language Models for Security Reasoning:* Investigation of LLM capabilities for complex threat reasoning, vulnerability assessment, penetration testing assistance, and automated security policy generation.
- *Standardized Evaluation Benchmarks:* Development of community-accepted evaluation protocols with realistic traffic diversity, temporal splits, encrypted traffic, and adversarial robustness assessments.

IX. CONCLUSION

This survey has presented a systematic and comprehensive review of machine learning based frameworks for intelligent cybersecurity threat detection and prevention. We have examined the evolving threat landscape across intrusion, malware, DDoS, phishing, and APT domains, and demonstrated how ML paradigms spanning supervised, unsupervised, semi-supervised, and reinforcement learning have been applied to address the fundamental inadequacies of traditional signature-based security systems.

Our comparative analysis of benchmark datasets, feature engineering strategies, and algorithm performance reveals that no single ML technique universally dominates across all threat categories. Ensemble methods including random forests offer strong performance with high interpretability, while deep learning architectures including CNNs, LSTMs, and transformers achieve state-of-the-art results on sequential and unstructured data at higher computational cost. Federated learning and privacy-preserving techniques are rapidly maturing as viable frameworks for collaborative threat intelligence across organizational boundaries.

The field faces substantial open challenges including adversarial robustness, concept drift, encrypted traffic analysis, model interpretability, and real-time scalability. Addressing these challenges requires close collaboration between the machine learning research community, security practitioners, and policymakers. The integration of explainable AI into security operations centers represents a particularly impactful near-term opportunity to accelerate analyst decision-making and build institutional trust in ML-driven defenses.

As cyber threats continue to evolve in sophistication and scale, the development of intelligent, adaptive, and transparent ML-based cybersecurity frameworks will remain among the most critical research imperatives in computer science. We hope this survey serves as a valuable reference and catalyst for advancing this vital research agenda.

REFERENCES

- [1] Buczak, A. L., & Guven, E. (2016). A survey of data mining and machine learning methods for cyber security intrusion detection. *IEEE Communications Surveys & Tutorials*, 18(2), 1153–1176.
- [2] Apruzzese, G., et al. (2018). On the effectiveness of machine and deep learning for cyber security. In *Proc. 10th International Conference on Cyber Conflict (CyCon)*, 371–390.
- [3] Sharafaldin, I., Lashkari, A. H., & Ghorbani, A. A. (2018). Toward generating a new intrusion detection dataset and intrusion traffic characterization. In *Proc. 4th International Conference on Information Systems Security and Privacy (ICISSP)*.
- [4] Moustafa, N., & Slay, J. (2015). UNSW-NB15: A comprehensive data set for network intrusion detection systems. In *Proc. Military Communications and Information Systems Conference (MilCIS)*, 1–6.
- [5] Vinayakumar, R., et al. (2019). Deep learning approach for intelligent intrusion detection system. *IEEE Access*, 7, 41525–41550.
- [6] Nataraj, L., et al. (2011). Malware images: Visualization and automatic classification. In *Proc. 8th International Symposium on Visualization for Cyber Security*, 4.
- [7] Anderson, H. S., & Roth, P. (2018). EMBER: An open dataset for training static PE malware machine learning models. *arXiv preprint arXiv:1804.04637*.
- [8] Li, D., & Deng, L. (2019). *Big data in complex and social networks*. CRC Press.
- [9] McMahan, B., et al. (2017). Communication-efficient learning of deep networks from decentralized data. In *Proc. 20th International Conference on Artificial Intelligence and Statistics (AISTATS)*, 1273–1282.
- [10] Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). Why should I trust you? Explaining the predictions of any classifier. In *Proc. 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 1135–1144.
- [11] Lundberg, S. M., & Lee, S. I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems (NeurIPS)*, 30.
- [12] Goodfellow, I., et al. (2014). Generative adversarial networks. *Communications of the ACM*, 63(11), 139–144.
- [13] Mirsky, Y., et al. (2018). Kitsune: An ensemble of autoencoders for online network intrusion detection. In *Proc. Network and Distributed System Security Symposium (NDSS)*.
- [14] Sommer, R., & Paxson, V. (2010). Outside the closed world: On using machine learning for network intrusion detection. In *Proc. IEEE Symposium on Security and Privacy*, 305–316.
- [15] Scarselli, F., et al. (2008). The graph neural network model. *IEEE Transactions on Neural Networks*, 20(1), 61–80.
- [16] Devlin, J., et al. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proc. NAACL-HLT*, 4171–4186.
- [17] Yin, C., et al. (2017). A deep learning approach for intrusion detection using recurrent neural networks. *IEEE Access*, 5, 21954–21961.
- [18] Shone, N., et al. (2018). A deep learning approach to network intrusion detection. *IEEE Transactions on Emerging Topics in Computational Intelligence*, 2(1), 41–50.
- [19] Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology*, 10(2), 1–19.
- [20] Cybersecurity Ventures. (2023). Cybercrime to cost the world \$8 trillion annually in 2023. *Cybercrime Magazine*.