

AI-Based Intrusion Detection Systems for Real-Time Identification of Network Attacks: A Literature Review

MD Shakhawat Hossain, CSE,
World University of Bangladesh

Naseef Talukder
Business department, Putra Business
School (UPM Malaysia)

Abstract:

Artificial intelligence (AI) can support network intrusion detection by learning traffic representations, classifying known attacks, and identifying deviations from expected behavior. Deep belief networks, convolutional and recurrent architectures, incremental mining, and open-set recognition have each been applied to reduce dependence on fixed signatures. Yet reported accuracy does not establish operational reliability. Dataset imbalance, simulated traffic, distribution drift, adversarial evasion, and incomplete latency measurements limit conclusions about deployment. This review examines AI-based intrusion detection as a systems problem in which classification quality must be assessed alongside processing rate, detection delay, adaptation to unseen attacks, and resistance to manipulation. The evidence supports hybrid and incremental designs, but it also shows that real-world validation and comparable evaluation protocols remain incomplete.

Keywords:

AI-based Intrusion Detection, Real-Time Network Security, Deep Learning Algorithms, Anomaly Detection Techniques, Zero-Day Attack Identification, Machine Learning Approaches

Introduction: From High-Speed Network Traffic to Adaptive Attack Identification

Network intrusion detection systems (NIDSs) inspect traffic to identify activity that threatens confidentiality, integrity, or availability. Conventional methods struggle when traffic volumes increase and when attacks differ from previously encoded rules, motivating machine learning and deep learning approaches [1]. Real-time operation adds constraints that offline classification can conceal: features must be extracted quickly, decisions must arrive before an attack causes unacceptable harm, and false alarms must remain manageable. The central research question is therefore not whether an AI model can classify a benchmark dataset, but whether an integrated detector can maintain useful accuracy while traffic, attack behavior, and adversarial pressure change.

AI Architectures for High-Speed and Real-Time Network Intrusion Detection Supervised, Hybrid, and Deep Learning Models for Known Attack Classification

AI-based NIDS architectures differ in how they represent traffic and divide detection tasks. Earlier supervised work selected 12 network features using information gain and implemented a decision-tree RT-IDS that classified normal activity, probes, and denial-of-service traffic with a detection rate above 98% within two seconds [2]. This design illustrates the value of compact features and interpretable classification when response time matters. Deep models extend this approach by learning representations from less structured inputs. AI-IDS combines normalized UTF-8 payload encoding for

spatial learning with CNN-LSTM processing of HTTP traffic, while entropy and compression measurements supplement payload analysis [3].

Hybrid architectures assign different functions to different learners. Zhang et al. combine flow calculation and frequent-pattern detection for stream processing with a DBN-SVM classifier for attack categorization [4]. Their CICIDS2017 experiments reported classification accuracy 0.7 percentage points above a DBN and two percentage points above integrated boosting and bagging methods [4]. Such separation can reduce the computational burden on the deep classifier because preliminary flow analysis narrows the data requiring detailed classification. It also creates a possible failure boundary: attacks that evade flow-level filtering may never reach the second stage.

Scalable and Incremental Detection for Evolving Network Traffic

Scalability requires more than a fast prediction function. BigFlow processes evolving traffic at large packet rates and includes verification of classifier outcomes, allowing an expert to revise the model when suspicious packets expose an error [5]. In experiments on traffic spanning a full year, the system used as little as 4% of the storage and between 0.05% and 4% of the training time required by comparison approaches, while processing 10-Gbps traffic on a 40-core commodity cluster [5]. These results connect computational efficiency with temporal reliability rather than treating them as separate objectives.

Incremental learning also appears in fuzzy association-rule detection and broad learning systems. Su et al. compare online and attack-free rule sets and produce decisions every two seconds, although their header-based features focus on large-scale attacks, particularly denial-of-service activity [6]. Li et al. evaluate recurrent networks and broad learning systems, including an incremental variant, on routing and connection records [7]. The literature therefore supports a layered architecture: rapid stream filtering, learned classification, verification, and controlled updating. The next issue is whether the data and metrics used to evaluate these layers represent operational conditions.

Data, Datasets, and Detection Latency in Evaluating AI-Based IDS Performance

Dataset Imbalance, Simulated Environments, and the Limits of Benchmark Evidence

Data validity constrains every reported AI-IDS result. Attack traffic commonly forms minority classes relative to benign traffic, so aggregate accuracy can conceal poor recognition of low-frequency attacks [8]. Generative augmentation addresses one part of this problem. Park et al. use Wasserstein-distance and reconstruction-error GAN formulations with autoencoder-based models to generate synthetic samples for minor attack traffic, reporting improved performance across several datasets [9]. Synthetic balance may improve representation of rare classes, but it does not establish that generated traffic preserves the causal and temporal properties of attacks in operational networks. Benchmark provenance creates a second limitation. Surveys report that many NID datasets originate in simulated or sandbox environments rather than real networks, which can produce unrepresentative traffic distributions [10]. The systematic review by Yang et al. identifies dataset selection, preprocessing, metrics, and attack-detection techniques as recurring dimensions across 119 anomaly-based studies, while also identifying unresolved research topics [11]. Public datasets such as CICIDS2017, CSE-CIC-IDS2018, CSIC-2010, NSL-KDD, and CIC-MalMem-2022 support reproducibility, but performance on them cannot automatically be transferred across organizations, protocols, encryption regimes, or traffic loads [3], [12], [13], [14].

Real-Time Processing, Reliability, and Latency-Sensitive Evaluation

Real-time evaluation must measure when a detector identifies an attack, not only whether it eventually assigns the correct label. Puccetti and Ceccarelli argue that temporal latency should be assessed together with false-positive rate, precision, and recall, because a correct decision delivered too late may have limited defensive value [15]. Their industrial cases show why latency-sensitive assessment requires datasets that preserve attack timing and support temporal metrics [15]. Earlier systems demonstrate different operational targets: the decision-tree RT-IDS reports detection within two seconds [2], the incremental rule system makes a decision every two seconds [6], and BigFlow addresses 10-Gbps throughput on clustered hardware [5].

These measures are not interchangeable. Throughput describes how much traffic a system can process, inference latency describes the time required to classify an observation, and attack-detection

latency describes the delay from attack onset to a usable alert. A model can achieve high offline recall while buffering traffic, or maintain low per-packet latency while producing excessive false alarms. Evaluation should therefore report the traffic representation, hardware context, processing rate, alert delay, class-specific recall, false-positive rate, and updating cost. This measurement discipline leads directly to the harder question of whether unseen attacks can be detected without destabilizing known-attack classification.

Zero-Day Generalization and Adaptation to Unseen Network Attacks **Anomaly-Based Detection of Novel and Zero-Day Attack Behaviors**

Zero-day generalization requires a detector to distinguish unfamiliar behavior from benign variation while retaining useful labels for known attacks. Autoencoders can provide an anomaly-detection layer before conventional classifiers. Dai et al. integrate an autoencoder with XGBoost and Random Forest models and report very high performance on unseen data from CIC-MalMem-2022, including 99.9892% accuracy for Random Forest-AE in the reported experiment [13]. That result indicates the value of combining representation-based anomaly detection with supervised classification, but its interpretation remains tied to one dataset and its definition of unseen traffic.

Open-set recognition offers a more explicit formulation. Soltani et al. use DOC++ to identify unknown samples while classifying multiple known attack categories, then cluster novel samples to reduce the burden of expert labeling before model updating [12]. Their evaluation uses CIC-IDS2017 and CSE-CIC-IDS2018, and the reported clustering completeness and homogeneity support the use of grouped novel samples for subsequent supervision [12]. Anomaly detection and open-set recognition address different decisions: the former signals deviation, whereas the latter attempts to separate unknown instances from recognized classes. Neither approach proves that a detector will identify every operational zero-day attack, because novelty depends on the training distribution, feature representation, and threshold policy.[16], [17]

Self-Sustaining and Adaptable Models for Changing Attack Distributions

Adaptation introduces a feedback process in which detection, review, labeling, and retraining affect future decisions. BigFlow uses outcome verification and expert assistance to incrementally modify its classifier as traffic evolves [5]. AIS-NIDS extends this concept through packet-level analysis, autonomous distinction between benign and unknown attacks, and incremental learning of new attack classes in a simulated real-world scenario [18]. These systems reduce reliance on a permanently fixed model, but autonomy must be bounded by safeguards against incorrect labels and contaminated updates.

The literature identifies a tension between adaptation speed and label quality. Expert labeling is costly, which motivates clustering and grouped review in DOC++ [12]. Automatic updating can reduce delay, yet an attacker who induces systematic misclassification may influence the data used for retraining. Deep-learning reviews identify concept distribution drift, high-dimensional traffic, class imbalance, and real-time interpretability as continuing problems [19]. A deployable design should therefore separate novelty detection from model promotion, retain evidence for audit, and test updates against established traffic before deployment. Adaptation without robustness creates a new attack surface, examined in the next section.[20], [21], [22]

Adversarial Robustness, Evasion, and Trustworthiness of AI-Based IDSs

Adversarial Examples, Black-Box Evasion, and Model Vulnerability

Adversarial robustness concerns traffic deliberately modified to change a detector's decision while preserving the attacker's objective. The survey by He et al. shows that deep NIDS models are exposed to adversarial examples, while methods borrowed from computer vision may not reflect the discrete, constrained feature spaces and processing pipelines of network traffic [23]. Network-specific perturbations can include changes to packet timing or inserted forged packets. Zhu et al. construct black-box evasive traffic by generating forged packets and inserting them into malicious flows; the reported evasion rate exceeded 90% for half of the tested anomaly detectors [24].

These findings qualify accuracy claims based only on clean test sets. MANDA detects adversarial examples through inconsistency between manifold proximity and model inference, together with uncertainty under small perturbations; on NSL-KDD and CICIDS evaluations, it reported a 98.41%

true-positive rate at a 5% false-positive rate [25]. The result is promising, but robustness remains threat-model dependent. Defenses must be assessed against adaptive, black-box, and feature-valid perturbations rather than a single attack generation method.

Generative Augmentation and Defensive Strategies for Robust Detection

Generative models can support robustness by expanding minority attack classes, although augmentation for imbalance is not equivalent to adversarial training [9]. Defensive strategies also include improved feature extraction, uncertainty estimation, manifold analysis, verification, and controlled updating [5], [24], [25], [26]. Trustworthiness consequently depends on the whole pipeline: what traffic is observed, how features are formed, how uncertainty is handled, and who authorizes adaptation. These requirements must be balanced against the processing and operational costs considered below.[27]

Synthesis: Balancing Detection Accuracy, Scalability, Novelty Detection, and Robustness Trade-Offs Between Classification Performance and Operational Responsiveness

The evidence does not support a single superior AI architecture. Compact supervised systems can deliver rapid binary or small-scale multiclass decisions, as shown by the decision-tree RT-IDS [2]. Deep and hybrid models can improve representation and attack classification, but they require more processing and depend heavily on training data [3], [4]. High-throughput systems address this constraint through stream processing, incremental computation, and selective detailed analysis [5], [6]. Accuracy, throughput, alert latency, false positives, and model-update cost should therefore be treated as a vector of operational endpoints rather than compressed into one score.

Novelty detection introduces another trade-off. Thresholds that identify more deviations may increase false alarms, while conservative thresholds may miss low-volume or slowly evolving attacks. Reported high results on unseen benchmark data demonstrate feasibility, not universal zero-day detection [13]. Robustness further changes the ranking of models because a detector with high clean-data accuracy may be vulnerable to black-box packet insertion or timing perturbation [24]. Comparative studies should evaluate clean and manipulated traffic under the same latency and workload conditions.[28], [29]

Converging Design Principles Across Hybrid, Deep, and Adaptive Systems

Across the studies, four design principles recur. First, preprocessing and feature selection must match the traffic source, whether headers, flows, payloads, or temporal sequences [2], [3]. Second, layered systems can allocate fast screening, detailed classification, novelty analysis, and expert review to different components [4], [12]. Third, incremental learning should be paired with verification or controlled labeling, as in BigFlow, DOC++, and AIS-NIDS [5], [12], [18]. Fourth, security evaluation must include adversarial behavior because model vulnerability can invalidate clean-test performance [23], [25].

A coherent architecture would thus maintain separate paths for known attacks and unknown behavior, record temporal evidence, and expose uncertainty to analysts. It would measure latency at attack level, preserve rollback capability for updates, and test feature extraction under evasion. The research gaps concern whether these principles have been demonstrated together outside controlled benchmarks.

Research Gaps in Deployable Real-Time AI-Based Intrusion Detection

Insufficient Real-World Validation and Cross-Environment Generalization

The principal gap is external validity. Many studies rely on public or simulated datasets, limiting evidence about changing user behavior, encrypted traffic, heterogeneous infrastructure, and organizational differences [10], [11]. Even studies reporting unseen-attack performance generally evaluate withheld data from established datasets rather than independently collected attacks in unfamiliar environments [12], [13]. Real-time claims also vary in meaning, from a two-second decision interval to backbone processing and clustered throughput [5], [6], [30].

Cross-environment studies should test transfer without assuming identical feature distributions. They should report performance decay, recalibration requirements, analyst workload, and update frequency. Packet-level systems may expose payload information unavailable to flow-based detectors, so

comparisons must align observability as well as algorithms [18]. Without this evidence, deployment claims remain conditional.[31], [32]

Unresolved Comparability of Latency, Robustness, and Zero-Day Performance

No common protocol currently aligns attack-detection latency, throughput, false alarms, open-set recognition, adaptation cost, and adversarial evasion. Reviews identify inconsistent datasets, metrics, preprocessing, and model categories as barriers to comparison [1], [11]. Future evaluations should define attack onset, alert usability, novelty labels, threat models, and hardware conditions before reporting results. Such protocols would make accuracy claims more informative and expose whether adaptation and robustness survive real-time constraints.

Conclusion: Toward Reliable and Adaptive Real-Time Network Attack Identification Implications for the Design of Future AI-Driven Intrusion Detection Systems

AI-based NIDS research has progressed from compact supervised classifiers toward hybrid, deep, open-set, and incremental systems. The strongest designs combine fast traffic processing with detailed classification, novelty analysis, verification, and controlled updating. Clean-data accuracy alone is insufficient because latency, data representativeness, drift, and adversarial evasion determine operational value.

Priorities for Validation, Adaptation, and Adversarial Resilience

Future work should prioritize longitudinal real-network validation, explicit attack-level latency, cross-environment testing, and adversarial evaluation. Model updates should retain expert oversight, uncertainty records, and rollback mechanisms rather than treating adaptation as unrestricted automation. Reliable real-time detection will emerge from jointly engineered data, computation, evaluation, and security controls, not from selecting a deeper classifier in isolation.

References

- [1] Z. Ahmad, A. Shahid Khan, C. Wai Shiang, J. Abdullah, and F. Ahmad, "Network intrusion detection system: A systematic study of machine learning and deep learning approaches," *Transactions on Emerging Telecommunications Technologies*, vol. 32, no. 1, Oct. 2020, doi: 10.1002/ett.4150.
- [10] D. Chou and M. Jiang, "A survey on data-driven network intrusion detection," *ACM Computing Surveys*, vol. 54, no. 9, pp. 1–36, Oct. 2021, doi: 10.1145/3472753.
- [11] Z. Yang *et al.*, "A systematic literature review of methods and datasets for anomaly-based network intrusion detection," *Computers & Security*, vol. 116, p. 102675, May 2022, doi: 10.1016/j.cose.2022.102675.
- [12] M. Soltani, B. Ousat, M. Jafari Siavoshani, and A. H. Jahangir, "An adaptable deep learning-based intrusion detection system to zero-day attacks," *Journal of Information Security and Applications*, vol. 76, p. 103516, Aug. 2023, doi: 10.1016/j.jisa.2023.103516.
- [13] Z. Dai *et al.*, "An intrusion detection model to detect zero-day attacks in unseen data using machine learning," *PLOS ONE*, vol. 19, no. 9, p. e0308469, Sep. 2024, doi: 10.1371/journal.pone.0308469.
- [14] R. Haider, Md. R. Amin, Md. S. Arafat, I. Hossain, and R. Ahmad, "Smart Automation: Revolutionizing Business Operations, Advertising Costs and Customer Satisfaction Through Technology," *European Economic Letters*, vol. 16, no. 1, p. null, 2026.
- [15] T. Puccetti and A. Ceccarelli, "On detection latencies of network intrusion detectors – discussion and application," *Empirical Software Engineering*, vol. 31, no. 1, Nov. 2025, doi: 10.1007/s10664-025-10766-3.
- [16] Raiyan Haider, Md. Farhan Ishrak Shaif, Raiyan Ahmed, Nahid Hasan Nafi, Mahmudul Reza Sumon, and Mushfiqur Rahman, "The influence of social media branding on consumer purchase behavior: A comprehensive empirical and thematic analysis," *International Journal of Science and Research Archive*, vol. 16, no. 2, pp. 460–470, Aug. 2025, doi: 10.30574/ijrsra.2025.16.2.2354.
- [17] Raiyan Haider, Jafia Tasnim Juba, and Satu Chowdhury, "Exploring the link between suicidal ideation and digital environments: The hidden impact of marketing content," *International Journal of*

- Science and Research Archive*, vol. 16, no. 2, pp. 607–614, Aug. 2025, doi: 10.30574/ijrsra.2025.16.2.2353.
- 10.[18] Y. A. Farrukh, S. Wali, I. Khan, and N. D. Bastian, “AIS-NIDS: An intelligent and self-sustaining network intrusion detection system,” *Computers & Security*, vol. 144, p. 103982, Sep. 2024, doi: 10.1016/j.cose.2024.103982.
- 11.[19] T. Yi, X. Chen, Y. Zhu, W. Ge, and Z. Han, “Review on the application of deep learning in network attack detection,” *Journal of Network and Computer Applications*, vol. 212, p. 103580, Mar. 2023, doi: 10.1016/j.jnca.2022.103580.
- 12.[2] P. Sangkatsanee, N. Wattanapongsakorn, and C. Charnsripinyo, “Practical real-time intrusion detection using machine learning approaches,” *Computer Communications*, vol. 34, no. 18, pp. 2227–2235, Dec. 2011, doi: 10.1016/j.comcom.2011.07.001.
- 13.[20] Raiyan Haider, Md Farhan Abrar Ibne Bari, Md. Farhan Israk Shaif, and Mushfiquir Rahman, “Engineering hyper-personalization: Software challenges and brand performance in AI-driven digital marketing management: An empirical study,” *International Journal of Science and Research Archive*, vol. 15, no. 2, pp. 1122–1141, May 2025, doi: 10.30574/ijrsra.2025.15.2.1525.
- 14.[21] Raiyan Haider, Md Farhan Abrar Ibne Bari, Osru, Nishat Afia, and Tanjim Karim, “Illuminating the black box: Explainable AI for enhanced customer behavior prediction and trust,” *International Journal of Science and Research Archive*, vol. 15, no. 3, pp. 247–268, 2025, doi: 10.30574/ijrsra.2025.15.3.1674.
- 15.[22] Raiyan Haider, “Navigating the Digital Political Landscape: How Social Media Marketing Shapes Voter Perceptions and Political Brand Equity in the 21st Century,” *International Journal of Science and Research Archive*, vol. 15, no. 1, pp. 1736–1744, 2025, doi: 10.30574/ijrsra.2025.15.1.1217.
- 16.[23] K. He, D. S. Kim, and M. R. Asghar, “Adversarial Machine Learning for Network Intrusion Detection Systems: A Comprehensive Survey,” *IEEE Communications Surveys & Tutorials*, vol. 25, pp. 538–566, 2023, doi: 10.1109/comst.2022.3233793.
- 17.[24] Y. Zhu, L. Cui, Z. Ding, L. Li, Y. Liu, and Z. Hao, “Black box attack and network intrusion detection using machine learning for malicious traffic,” *Computers & Security*, vol. 123, p. 102922, Dec. 2022, doi: 10.1016/j.cose.2022.102922.
- 18.[25] N. Wang, Y. Chen, Y. Xiao, Y. Hu, W. Lou, and Y. T. Hou, “MANDA: On adversarial example detection for network intrusion detection system,” *IEEE Transactions on Dependable and Secure Computing*, vol. 20, no. 2, pp. 1139–1153, Mar. 2023, doi: 10.1109/tdsc.2022.3148990.
- 19.[26] “Generative AI and quantum-inspired optimization: Redefining portfolio risk management and real-time capital allocation in volatile markets,” *Journal of Informatics Education and Research*, Mar. 2026, doi: 10.52783/jier.v6i1.4540.
- 20.[27] Raiyan Haider, Wahida Ahmed Megha, Jafia Tasnim Juba, Aroa Alamgir, and Labib Ahmad, “The conversational revolution in health promotion: Investigating chatbot impact on healthcare marketing, patient engagement, and service reach,” *International Journal of Science and Research Archive*, vol. 15, no. 3, pp. 1585–1592, Jun. 2025, doi: 10.30574/ijrsra.2025.15.3.1937.
- 21.[28] R. Haider, T. Dwivedi, A. G. Girish, N. Verma, B. Kashyap, and V. S. Dubey, “Neuromarketing Approaches to Shaping Healthy Consumer Choices: An Integrative Analysis of Methods, Efficacy, and Ethical Considerations,” *International Journal of Drug Delivery Technology*, vol. 16, no. 4s, pp. 512–524, 2026, doi: 10.25258/ijddt.16.4s.62.
- 22.[29] Raiyan Haider, Md. Farhan Israk Shaif, Raiyan Ahmed, Nahid Hasan Nafi, Mahmudul Reza Sumon, and Mushfiquir Rahman, “Assessing the impact of influencer marketing on brand value and business revenue: An empirical and thematic analysis,” *International Journal of Science and Research Archive*, vol. 16, no. 2, pp. 471–482, Aug. 2025, doi: 10.30574/ijrsra.2025.16.2.2355.
- 23.[3] A. Kim, M. Park, and D. H. Lee, “AI-IDS: Application of deep learning to real-time web intrusion detection,” *IEEE Access*, vol. 8, pp. 70245–70261, 2020, doi: 10.1109/access.2020.2986882.
- 24.[30] C. Callegari, S. Giordano, and M. Pagano, “A real time deep learning based approach for detecting network attacks,” *Big Data Research*, vol. 36, p. 100446, May 2024, doi: 10.1016/j.bdr.2024.100446.
- 25.[31] Raiyan Haider, Farhan Abrar Ibne Bari, Osru, Nishat Afia, and Mohammad Abiduzzaman Khan Mugdho, “Leveraging internet of things data for real-time marketing: Opportunities, challenges, and

- strategic implications,” *International Journal of Science and Research Archive*, vol. 15, no. 3, pp. 1657–1663, 2025, doi: 10.30574/ijrsra.2025.15.3.1936.
- 26.[32] Raiyan Haider, Md Farhan Abrar Ibne Bari, Md. Farhan Israk Shaif, Mushfiqur Rahman, Md. Nahid Hossain Ohi, and Kazi Md Mashrur Rahman, “Quantifying the Impact: Leveraging AI-Powered Sentiment Analysis for Strategic Digital Marketing and Enhanced Brand Reputation Management,” *International Journal of Science and Research Archive*, vol. 15, no. 2, pp. 1103–1121, 2025, doi: 10.30574/ijrsra.2025.15.2.1524.
- 27.[4] H. Zhang, Y. Li, Z. Lv, A. K. Sangaiah, and T. Huang, “A real-time and ubiquitous network attack detection based on deep belief network and support vector machine,” *IEEE/CAA Journal of Automatica Sinica*, vol. 7, pp. 790–799, 2020, doi: 10.1109/jas.2020.1003099.
- 28.[5] E. K. Viegas, A. O. Santin, A. Bessani, and N. Neves, “BigFlow: Real-time and reliable anomaly-based intrusion detection for high-speed networks,” *Future Generation Computer Systems*, vol. 93, pp. 473–485, 2018, doi: 10.1016/j.future.2018.09.051.
- 29.[6] M. Su, G. Yu, and C.-Y. Lin, “A real-time network intrusion detection system for large-scale attacks based on an incremental mining approach,” *Computers & Security*, vol. 28, pp. 301–309, 2008, doi: 10.1016/j.cose.2008.12.001.
- 30.[7] Z. Li, A. L. G. Rios, and L. Trajkovic, “Machine learning for detecting anomalies and intrusions in communication networks,” *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 7, pp. 2254–2264, Jul. 2021, doi: 10.1109/jsac.2021.3078497.
- 31.[8] P. Mishra, V. Varadharajan, U. Tupakula, and E. S. Pilli, “A detailed investigation and analysis of using machine learning techniques for intrusion detection,” *IEEE Communications Surveys & Tutorials*, vol. 21, no. 1, pp. 686–728, 2019, doi: 10.1109/comst.2018.2847722.
- 32.[9] C. Park, J. Lee, Y. Kim, J.-G. Park, H. Kim, and D. Hong, “An enhanced AI-Based network intrusion detection system using generative adversarial networks,” *IEEE Internet of Things Journal*, vol. 10, no. 3, pp. 2330–2345, Feb. 2023, doi: 10.1109/jiot.2022.3211346.