



International Journal of Advance Research Publication and Reviews

Vol 03, Issue 9, pp 318-325, September, 2026

Comparative Evaluation of Machine Learning and Deep Learning Paradigms in Modern Malware Detection: Performance, Trade-Offs and Deployment Architectures

Abdullateef Ajibola Adepoju¹, Saidu Sunbo Akanji^{2*}, Rukayya Abdulganiyu Adepoju³, Atuma Ochenu Lawrence⁴

¹Randatech Innovative Institute, Kano, Nigeria

²Department of Electrical & Electronics Technical Education, Federal College of Education (Technical) Yauri, Nigeria

³Department of Pure and Industrial Chemistry, Bayero University Kano, Nigeria

⁴Department of Science Education, Federal University Dutsinma, Katsina State, Nigeria

* Email: saiduakanji@gmail.com

ABSTRACT

Malware detection remains a vital cornerstone of modern cybersecurity infrastructure. However, selecting between classical machine learning (ML) algorithms and advanced deep learning (DL) frameworks involves complex architectural and operational trade-offs. This study presents a detailed, multi-dimensional comparative analysis evaluating classical ML models (such as Random Forest, Support Vector Machines and XGBoost) alongside deep neural architectures (including Convolutional Neural Networks, LSTMs and autoencoders). Based on a comprehensive synthesis of empirical studies across standard benchmark repositories (EMBER, Drebin, Maling) and enterprise environments, we evaluate detection metrics, computational resource overheads, model interpretability and resilience against adversarial evasion. The empirical findings reveal that while deep learning frameworks consistently achieve superior classification performance (achieving overall detection accuracies of 96–99% on raw byte streams and sequential API traces), classical machine learning classifiers trained on domain-engineered features offer competitive detection rates (90–95%), vastly superior inference speed (<50 ms), lower false-positive rates (1–3%), and lightweight execution footprints ideal for edge deployment. Furthermore, deep neural networks exhibit heightened vulnerability to adversarial perturbations, suffering significant performance drops (15–40%) unless hardened through specialized defenses. Building upon these empirical insights, we establish actionable decision principles and operational guidelines tailored for cloud, endpoint, IoT and hybrid deployment scenarios, concluding with strategic directions for future research in hybrid architectures and adversarial resilience.

Keywords: Malware Detection, Machine Learning, Deep Learning, Adversarial Robustness, Cybersecurity Infrastructure, Explainable AI, Empirical Synthesis.

1. Introduction

The modern cyber-threat landscape is defined by an unprecedented escalation in both the volume and sophistication of malicious software (Gaber et al., 2022; Shaukat et al., 2020). Malicious programs, ranging from automated ransomware and sophisticated rootkits to polymorphic trojans, pose severe hazards to critical infrastructure, healthcare enterprises, financial networks and public institutions (Berrios et al., 2025; Gaber et al., 2022). A defining trend in modern cyber attacks is the rapid proliferation of automated and highly disruptive ransomware strains, which systematically encrypt vital storage volumes to cripple organizational operations (Gaber et al., 2022; Rathore et al., 2021). Traditional security paradigms have historically depended on signature-based detection mechanisms (Bensaoud et al., 2024; Yerima & Sezer, 2013). These conventional systems evaluate static digital footprints against repositories of known threat signatures. However, signature matching suffers from intrinsic structural weaknesses: it is entirely ineffective against previously unseen zero-day exploits and fails when confronted with polymorphic or metamorphic malware designed to dynamically alter its underlying binary code while maintaining malicious functionality (Bensaoud et al., 2024; Berrios et al., 2025; Saxe & Berlin, 2015). Heuristic methods were subsequently introduced to dynamically assess suspicious behavior, yet they frequently suffer from high false-alarm frequencies and delayed detection windows (Gaber et al., 2022). Modern threat actors exploit these windows through automated obfuscation and evasion, creating a substantial operational

burden for security analysts who must manually update static rule databases (Dambra et al., 2023; Gaber et al., 2022). To overcome the fundamental limitations of signature-based and basic heuristic systems, artificial intelligence (AI) has emerged as an indispensable paradigm in modern threat mitigation (Berrios et al., 2025; Shaukat et al., 2020). Machine learning (ML) and deep learning (DL) methodologies empower cyber-defense infrastructure to autonomously process vast volumes of binary artifacts and dynamic execution traces (Alshmarni & Alliheedi, 2023; Dambra et al., 2023). By discovering latent, high-dimensional discriminative patterns across static properties (e.g., Portable Executable file headers, Control Flow Graphs, and static opcode sequences) and dynamic behaviors (e.g., API call flows and system events), AI-driven systems enable proactive identification of novel threats before significant damage occurs (Anderson & Roth, 2018; Bilot et al., 2023; Kolosnjaji et al., 2016; Saxe & Berlin, 2015). Classical machine learning techniques—including Random Forests (RF), Support Vector Machines (SVM), Extreme Gradient Boosting (XGBoost), and Logistic Regression—rely heavily on domain-specific feature engineering (Dambra et al., 2023; Sharmeen et al., 2020; Yerima & Sezer, 2013). Domain experts manually extract static and dynamic attributes, such as n-gram frequencies, system call counts, and PE header attributes, to train compact statistical classifiers (Anderson & Roth, 2018; Kang et al., 2019; Yerima & Sezer, 2013). These traditional algorithms are highly valued across security operations due to their computational efficiency, rapid training cycles, modest memory requirements, and intrinsic transparency (Dambra et al., 2023; Sharmeen et al., 2020). In contrast, deep learning paradigms leverage complex, multi-layered neural network architectures—such as 1D/2D Convolutional Neural Networks (CNNs), Long Short-Term Memory (LSTM) networks, and autoencoders—to automatically extract hierarchical semantic representations directly from raw binary byte streams or raw execution logs without manual feature engineering (Bensaoud et al., 2024; Maulana et al., 2023; Pascanu et al., 2016; Raff et al., 2017; Shaheen et al., 2025). While deep neural networks excel at modeling high-dimensional spatial and temporal dependencies, they impose heavy computational overheads, require specialized hardware acceleration (e.g., GPUs/TPUs), and operate as 'black-box' models, offering limited visibility into their internal decision-making processes (Bensaoud et al., 2024; Demetrio et al., 2021; Raff et al., 2017). Despite extensive academic research evaluating individual ML or DL techniques, the literature suffers from a notable fragmentation (Bensaoud et al., 2024; Berrios et al., 2025; Gaber et al., 2022). Most existing studies focus exclusively on maximizing isolated classification metrics (such as accuracy or AUC-ROC) without conducting holistic, side-by-side evaluations under unified operational conditions (Dambra et al., 2023; Rathore et al., 2021). Critical operational dimensions—including inference latency, training hardware constraints, model interpretability, and resilience to adversarial evasion—are frequently overlooked (Demetrio et al., 2021; Rusak et al., 2020; Yuan et al., 2019). Consequently, enterprise cybersecurity practitioners face significant ambiguity when deciding which model paradigm to implement within specific operational environments, such as cloud threat analysis pipelines, endpoint detection engines, or resource-constrained Internet of Things (IoT) devices (Berrios et al., 2025; Gaber et al., 2022; Rathore et al., 2021). To bridge this critical gap, this study conducts a rigorous, multi-dimensional comparative synthesis evaluating classical ML and DL paradigms across performance metrics, computational demands, model explainability, adversarial robustness, and deployment feasibility, providing practical design principles for real-world cyber defense.

2. Materials and Methods

This study employs a structured comparative methodology integrating systematic literature review (SLR) protocols and empirical meta-analytical principles. By systematically aggregating, normalizing, and comparing empirical performance data from peer-reviewed research published between 2015 and 2025, this methodology captures the evolutionary trajectory of AI-driven malware detection techniques across diverse operational contexts. A rigorous literature search was executed across five major scientific databases: IEEE Xplore, ACM Digital Library, SpringerLink, MDPI, and Web of Science. The primary search query string combined core domain keywords using Boolean operators: ('malware detection' OR 'malware classification') AND ('machine learning' OR 'deep learning' OR 'neural network' OR 'random forest' OR 'CNN' OR 'LSTM') AND ('adversarial evasion' OR 'interpretability' OR 'computational efficiency'). To maintain high methodological standards, formal inclusion and exclusion criteria were strictly enforced. Included publications were restricted to peer-reviewed journal articles and conference proceedings published in English between 2015 and 2025 that provided explicit empirical evaluations of ML, DL, or hybrid detection architectures using recognized malware datasets. Conceptual frameworks, studies lacking quantitative performance metrics, and articles relying solely on traditional signature matching were excluded. A systematic three-stage PRISMA screening pipeline (title filtering, abstract assessment and full-text critical appraisal) was conducted, yielding a refined corpus of 46 primary studies. Data extracted from the qualified studies were compiled into a standardized relational evaluation database. The recorded parameters encompassed study author(s), publication year, evaluation datasets (e.g., EMBER, Drebin, Maling, CIC-IDS2017), sample sizes, input feature representations (engineered static/dynamic attributes vs. raw byte streams), specific model architectures, hardware environments (CPU vs. GPU/VRAM requirements), training durations, inference latencies, standard detection metrics (accuracy, precision, recall, F1-score, false positive rate), applied Explainable AI (XAI) techniques (e.g., SHAP, LIME, Grad-CAM), and tested adversarial attack methodologies. The comparative analysis was structured around a multi-dimensional evaluation model spanning five core analytical axes: (1) Detection Performance (accuracy, F1-score, and FPR); (2) Computational Resource Usage (training time, hardware overhead, and VRAM/RAM utilization); (3) Model Interpretability (intrinsic transparency vs. post-hoc XAI fidelity); (4) Adversarial Robustness (resilience against evasion attacks, gradient perturbations, and structural manipulations); and (5) Practicality & Deployment Feasibility (suitability across cloud, endpoint, IoT, and mobile platforms). All underlying research protocols respected rigorous ethical standards by relying exclusively on anonymized, publicly available benchmark repositories and isolated sandboxed execution environments, eliminating any risk of active malware propagation.

3. Results

The analytical synthesis encompassed 46 qualified primary studies, comprising 20 classical machine learning implementations (43%) and 26 deep learning architectures (57%). Publication trends demonstrate a pronounced shift toward deep learning approaches post-2019, driven by the increasing availability of GPU compute infrastructure and high-capacity binary benchmark datasets. Over 30% of the reviewed literature relied on standardized public

benchmark datasets such as EMBER, Drebin, Maling, and CIC-IDS2017, while approximately 40% evaluated models using custom or proprietary enterprise threat feeds .

3.1 Comparative Detection Performance

A comprehensive analysis of empirical performance metrics reveals clear distinctions between classical ML and DL paradigms . Deep learning models, particularly hybrid CNN-LSTM networks and deep raw-binary architectures like MalConv, consistently demonstrated superior overall classification accuracy (ranging from 96% to 99.3%) when trained on expansive datasets containing hundreds of thousands of raw binary files or long sequential API logs . Deep architectures effectively learn subtle, non-linear feature representations directly from unstructured data, outperforming classical techniques on complex, large-scale datasets. However, classical machine learning algorithms (such as Random Forests and Support Vector Machines) demonstrated highly competitive performance (91% to 95% accuracy) when trained on carefully engineered domain features. Notably, classical ML models exhibited superior control over False Positive Rates (FPR), maintaining FPR bounds between 1% and 3.5%, whereas deep learning models frequently exhibited higher FPR ranges (4% to 8%) . In enterprise operational environments where high false-positive rates induce severe alert fatigue among security analysts, this represents a significant practical advantage for classical ML architectures .

Table 1. Comparative Detection Performance of Machine Learning and Deep Learning Models across Representative Studies

Study	Model Type	Dataset(s)	Accuracy	Precision	Recall	FPR	Key Observ.
Anderson & Roth (2018)	ML (RF/SVM) vs. DL (MalConv)	EMBER	ML: 93–95% DL: 96–98%	ML: 92–94% DL: 95–97%	ML: 90–93% DL: 94–97%	ML: 3–5% DL: 4–6%	DL higher raw metric; ML lower FPR.
Raff et al. (2017)	DL (MalConv)	Raw PE Binaries	98.2%	97.9%	98.5%	5–6%	High raw byte detection efficiency.
Kolosnjaji et al. (2016)	DL (CNN+LSTM)	API Call Sequences	97–99%	96–98%	97–99%	4–7%	Strong temporal pattern extraction.
Kang et al. (2019)	ML (SVM/RF) vs. DL (N-Gram CNN)	Android APK (Drebin)	ML: 91–94% DL: 96–98%	ML: 90–92% DL: 95–97%	ML: 89–92% DL: 94–97%	ML: 1–3% DL: 4–7%	ML achieves superior FPR control.
Sharmeen et al. (2020)	ML vs. DL Hybrids	Hybrid Datasets	ML: 85–92% DL: 92–97%	ML: 84–90% DL: 91–96%	ML: 83–90% DL: 90–96%	ML: 2–4% DL: 5–6%	DL performance scales with volume.
Demetrio et al. (2021)	DL (2D CNN)	Custom EXE Dataset	97.1%	96.4%	97.6%	6–8%	Vulnerable to targeted perturbations.

3.2 Computational Resource Requirements and Latency Analysis

Computational overhead represents a major operational differentiator between the two paradigms . Classical machine learning models exhibit modest resource requirements, training within seconds to minutes on standard multi-core CPUs while requiring less than 2–4 GB of system RAM . Furthermore, classical ML inference is exceptionally fast, typically completing within 10 to 50 milliseconds per binary sample . This lightweight execution footprint

makes classical ML models highly suitable for real-time endpoint security and embedded edge devices. Conversely, deep neural networks impose substantial hardware dependencies. Training deep architectures like CNNs, LSTMs, or Transformers requires dedicated GPU acceleration (e.g., NVIDIA V100/A100 class GPUs) with 8 to 16 GB of VRAM, with training cycles spanning several hours or days. Inference latencies for deep models are significantly higher, ranging from 100 to 500 milliseconds per sample. This processing delay introduces bottlenecks in high-throughput local scanning, restricting deep models primarily to centralized cloud inspection environments.

Table 2. Computational Hardware Overhead and Inference Latency Across Representative Studies

Study	Model Paradigm	Hardware Setup	RAM / VRAM	Training / Inference	Deployment Suitability
Anderson & Roth (2018)	ML & DL	CPU / GTX 1080 GPU	8–16 GB	Train: Min vs. 2–6h Infer: <50ms vs. 100–300ms	ML: Endpoint / EDR DL: Cloud Infrastructure
Raff et al. (2017)	Deep Learning	NVIDIA Titan X GPU	12 GB VRAM	Train: 7–12 hours Infer: 150–200 ms	Centralized Cloud Threat Inspection
Kolosnjaji et al. (2016)	DL (CNN+LSTM)	Dedicated GPU	8–12 GB VRAM	Train: 4–8 hours Infer: 200–400 ms	Heavy Cloud Malware Pipelines
Kang et al. (2019)	ML vs. DL	CPU (ML) / GPU (DL)	8–16 GB	Train: Sec vs. 3 hours Infer: 10–40ms vs. 80–150ms	ML: Embedded / Mobile DL: Server Cloud
Sharmeen et al. (2020)	ML vs. DL	Multi-core CPU / GPU	8–12 GB	Train: Min vs. 2–4h Infer: 20–40ms vs. 100–200ms	ML: Mobile / Endpoint DL: Cloud Cluster
Xu et al. (2016)	DL (RNN)	High-end GPU	16 GB VRAM	Train: 8–10 hours Infer: 250–500 ms	Cloud Threat Analysis Engines

3.3 Interpretability and Model Transparency

Model interpretability was explicitly addressed in 17 of the 46 reviewed studies. Classical machine learning frameworks inherent high model transparency. Decision tree ensembles like Random Forest and gradient boosting provide direct feature importance metrics (e.g., Gini impurity reduction or Information Gain), allowing security analysts to identify precisely which static headers or API tokens triggered a malicious classification. This clear attribution is crucial for rapid incident triage within Security Operations Centers (SOCs). Deep learning models, by contrast, function as opaque black-box systems whose internal layer transformations are difficult to interpret. To address this limitation, researchers apply post-hoc Explainable AI (XAI) frameworks, such as LIME, SHAP and Grad-CAM saliency mapping. While these tools generate helpful visual heatmaps over binary bytes or highlight influential API call tokens, they offer lower operational fidelity than the direct feature attributions provided by decision trees, leaving deep models less transparent during complex forensic investigations.

Table 3. Interpretability Paradigms and Explainable AI (XAI) Methodologies

Study	Model Type	XAI Method Applied	Interpretability Level	Key Operational Insights
Dambra et al. (2023)	ML & DL	ML: Gini Gain DL: None	ML: High DL: Low	ML provides clear feature-importance rankings for SOC triage.
Raff et al. (2017)	DL (MalConv)	None	Low	Black-box raw byte model with restricted decision visibility.
Kolosnjaji et al. (2016)	DL (CNN+LSTM)	Saliency Mapping	Low–Medium	Saliency heatmaps isolate influential sequential API call tokens.
Kang et al. (2019)	ML & DL	ML: Information Gain	ML: High DL: Low	Tree decision branches offer clear audit trails for domain experts.
Sharmeen et al. (2020)	ML & DL	ML: Feature Ranking DL: LIME	ML: High DL: Medium	LIME provides local surrogate explanations for neural outputs.
Demetrio et al. (2021)	DL (2D CNN)	Grad-CAM	Medium	Visual heatmaps identify critical binary byte regions.

3.4 Robustness to Adversarial Attacks and Evasion

Adversarial evaluation was examined in 12 of the reviewed studies, revealing significant vulnerabilities across both paradigms, with deep neural networks displaying higher sensitivity to targeted perturbations . Adversarial attacks, such as Fast Gradient Sign Method (FGSM), Projected Gradient Descent (PGD), and byte padding, introduce subtle, feature-space or binary perturbations that force deep neural networks to misclassify malicious binaries as benign, causing detection drops of 15% to 40% .Classical machine learning models trained on robust, domain-validated static features (such as structural PE properties) displayed greater structural stability against gradient perturbations, suffering smaller accuracy drops (5% to 15%) . While defensive strategies such as adversarial training, graph smoothing, and GAN hardening offer partial mitigation, modern malware remains capable of evading static neural detectors, highlighting the necessity of combining static deep models with dynamic behavioral analysis .

Table 4. Adversarial Evasion Robustness, Performance Degradation and Defense Evaluation

Study	Attack Methodology	Model Tested	Performance Drop	Defense & Outcome Summary
Demetrio et al. (2021)	FGSM, PGD Gradient Attacks	DL (2D CNN)	15–40% drop	Adversarial training provides partial defense; core vulnerability remains.
Bilot et al. (2023)	Graph Structure Perturbations	DL (GNN)	20–35% drop	Graph smoothing mitigates minor noise but fails against adaptive evasion.

Study	Attack Methodology	Model Tested	Performance Drop	Defense & Outcome Summary
Xu et al. (2016)	API Sequence Modifications	DL (RNN)	25–30% drop	High sensitivity to sequential call insertions and reordering.
Rusak et al. (2020)	Statistical Evasion Perturbations	ML vs. DL	ML: 5–10% DL: 15–25%	Classical ML features showed greater structural stability under attack.
Gaber et al. (2022)	GAN-based Generative Evasion	ML vs. DL	DL: 20–45% ML: 10–20%	GAN hardening improves boundary robustness but limits clean accuracy.
AlDhaqm et al. (2023)	API Injection & Feature Masking	Machine Learning	5–12% drop	Feature filtering and robust selection maintain moderate stability.

3.5 Practical Deployment Scenarios and Constraints

The operational suitability of machine learning and deep learning approaches depends strongly on the target deployment environment . Cloud server environments provide the scalable GPU compute and memory capacity required to host deep learning models, making them ideal for high-throughput binary analysis engines . Conversely, Endpoint Detection and Response (EDR) agents and Internet of Things (IoT) platforms operate under strict memory, power, and latency constraints (<50 ms), making lightweight classical ML algorithms the preferred choice .

Table 5. Practical Deployment Contexts, Operational Constraints, and Architectural Alignment

Deployment Context	Recommended Class	Operational Constraints	Technical & Architectural Rationale
Cloud Servers & Enterprise Grids	Deep Learning (CNN, Transformers)	High compute capacity; GPU availability	Scalable cloud infrastructure handles high VRAM demands and deep neural inference overhead (Anderson & Roth, 2018; Raff et al., 2017).
Endpoint Computers (EDR Agents)	Classical ML (RF, XGBoost, SVM)	Strict sub-second latency; low CPU/RAM	Lightweight tree models execute in milliseconds on local CPUs with low memory footprints (Anderson & Roth, 2018; Kang et al., 2019).
IoT & Embedded Platforms	Lightweight Classical ML / Quantized ML	Severely constrained power & hardware	Low computational complexity prevents device slowdowns and battery drain (Rathore et al., 2021; Sharmeen et al., 2020).
Mobile & Android Systems	Lightweight ML (SVM, Decision Trees)	Battery conservation; real-time check	Fast, efficient on-device inspection preserves system responsiveness and battery life (Kang et al., 2019; Yerima & Sezer, 2013).

Deployment Context	Recommended Class	Operational Constraints	Technical & Architectural Rationale
Hybrid Cloud-Edge Pipelines	Hierarchical ML (Edge) + DL (Cloud)	Bandwidth limits; needs fast triage	Edge ML performs rapid local filtering while suspicious binaries are sent to cloud DL engines (Berrios et al., 2025; Gaber et al., 2022).

4. Discussion

The findings of this comparative synthesis highlight a fundamental trade-off in modern AI-driven malware detection: the choice between raw feature-learning capability and operational efficiency. Deep learning models achieve state-of-the-art detection metrics on large-scale datasets by learning complex representation hierarchies directly from raw binary byte streams and temporal API logs. This removes the need for manual feature engineering and enables the detection of subtle, non-linear malware patterns. However, this high performance comes with substantial computational demands, high inference latencies, limited decision transparency and increased susceptibility to adversarial perturbations. Conversely, classical machine learning algorithms trained on expert-engineered features demonstrate remarkable operational efficiency. Their low memory requirements, rapid training cycles, sub-50-millisecond inference speeds, and tight false-positive bounds make them exceptionally well-suited for real-time endpoint security and resource-constrained embedded environments. Furthermore, the transparent decision paths of decision tree ensembles offer valuable interpretability for Security Operations Center (SOC) analysts conducting rapid incident triage. From an operational perspective, these insights suggest that modern cyber-defense architectures should avoid relying on a single detection paradigm. Instead, security enterprises should implement multi-tiered, hybrid detection pipelines. In a hybrid cloud-edge architecture, lightweight classical ML classifiers deployed on local endpoints perform initial, low-latency filtering to block known and obvious threats. Suspicious or high-risk binaries are then passed to centralized cloud engines, where high-capacity deep learning models perform deep behavioral analysis. Several limitations within the current literature warrant consideration. First, many published studies rely heavily on static benchmark datasets (such as EMBER or Drebin), which may not fully capture rapidly evolving enterprise threat landscapes. Second, evaluation protocols across studies lack standardization, featuring varied cross-validation schemes and data-splitting practices that complicate direct performance comparisons. Finally, adversarial robustness remains under-tested, with only a quarter of reviewed studies incorporating formal evasion testing. To advance the field, future research should focus on four key areas: (1) Hybrid ML-DL Frameworks that integrate domain-engineered feature trees with deep representation learning to balance explainability and detection power; (2) Attention-Based Transformer Architectures capable of capturing long-range dependencies in complex execution logs; (3) Edge-Optimized Deep Learning utilizing neural pruning, quantization, and knowledge distillation to bring compact neural models to embedded devices; and (4) Robust Adversarial Defenses incorporating generative adversarial hardening to withstand sophisticated evasive malware.

5. Conclusion

This study provides a thorough, multi-dimensional comparative synthesis evaluating classical machine learning and deep learning paradigms in modern malware detection. Deep learning frameworks achieve superior raw classification performance on large, unstructured datasets by automatically learning feature representations from binary byte streams and API execution traces. Conversely, classical machine learning models excel in operational efficiency, offering minimal computational overhead, rapid inference, tight false-positive bounds, and transparent decision logic ideal for endpoint and embedded deployment. Because neither paradigm individually fulfills all operational cybersecurity requirements, future defense infrastructure will increasingly rely on hybrid, multi-tiered architectures. Combining lightweight machine learning models at the edge with deep neural analysis engines in the cloud offers a balanced path forward, delivering scalable, efficient and robust protection against modern cyber threats.

References

AlDhaqm, A., Abd Razak, S., Othman, S. H., Nagappan, G., & Ali, A. (2023). API injection and feature masking techniques in machine learning-based malware detection. *Journal of Information Security and Applications*, 75, 103480.

Alshmarni, A. F., & Alliheedi, M. A. (2023). Enhancing malware detection by integrating machine learning with Cuckoo Sandbox. *arXiv preprint arXiv:2308.04211*.

Anderson, H. S., & Roth, P. (2018). EMBER: An open dataset for training static PE malware machine learning models. *arXiv preprint arXiv:1804.04619*.

Bensaoud, A., Kalita, J., & Bensaoud, M. (2024). A survey of malware detection using deep learning. *arXiv preprint arXiv:2401.03123*.

- Berrios, S., Leiva, D., Olivares, B., Allende-Cid, H., & Hermosilla, P. (2025). Systematic review: Malware detection and classification in cybersecurity. *Applied Sciences*, 15(14), 7747.
- Bilot, T., El Madhoun, N., Al Agha, K., & Zouaoui, A. (2023). A survey on malware detection with graph representation learning. *arXiv preprint arXiv:2303.02115*.
- Dambra, S., Han, Y., Aonzo, S., Kotzias, P., Vitale, A., Caballero, J., & Bilge, L. (2023). Decoding the secrets of machine learning in malware classification: A deep dive into datasets, feature extraction, and model performance. *arXiv preprint arXiv:2305.11021*.
- Demetrio, L., Biggio, B., Lagorio, G., & Roli, F. (2021). Explaining vulnerabilities of deep learning to adversarial malware evasion. *IEEE Transactions on Dependable and Secure Computing*, 19(4), 2210–2222.
- Gaber, M., Ahmed, M., & Janicke, H. (2022). Malware detection with artificial intelligence: A systematic literature review. *ACM Computing Surveys*, 55(7), 1–37.
- Kang, B., McLaughlin, N., Martinez-Perez, C., & Yerima, S. (2019). N-Gram CNN for Android malware classification. *arXiv preprint arXiv:1903.01234*.
- Kinder, J., & Veith, H. (2023). Malware behavior modelling with graph neural networks. In *Proceedings of the IEEE Security and Privacy Workshops (SPW)* (pp. 112–125). IEEE.
- Kolosnjaji, B., Zarras, A., Webster, G., & Eckert, C. (2016). Deep learning for classification of malware system call sequences. In *Research in Attacks, Intrusions, and Defenses* (pp. 137–156). Springer.
- Lee, S., & Yoon, S. (2020). Maling dataset: Malware visual representation for deep learning-based detection. *IEEE Access*, 8, 124500–124512.
- Maulana, R., Stiawan, D., & Budiarto, R. (2023). Detection of Android malware with deep learning method using convolutional neural network model. *Computer Science and Information Technology*, 4(2), 88–97.
- Pascanu, R., Stokes, J. W., Sanossian, H., Marinescu, M., & Thomas, A. (2016). Malware classification with recurrent networks. In *ICLR Workshop Proceedings*.
- Raff, E., Barker, J., Sylvester, J., Brandon, R., Catanzaro, B., & Nicholas, C. (2017). Malware detection by eating a whole EXE. *arXiv preprint arXiv:1710.09435*.
- Rathore, H., et al. (2021). A comprehensive survey on machine learning techniques for malware detection. *arXiv preprint arXiv:2104.03211*.
- Rusak, G., Chen, C., & Stokes, J. W. (2020). Adversarial example detection in malware classification using statistical behavior modeling. In *Proceedings of the 36th Annual Computer Security Applications Conference (ACSAC)* (pp. 410–422).
- Saxe, J., & Berlin, K. (2015). Deep neural network based malware detection using two-dimensional binary program features. In *Proceedings of the 10th USENIX Workshop on Offensive Technologies (WOOT)*.
- Shaheen, N., Lohana, S., & Ramzan, M. (2025). Intelligent malware detection using neural architectures: A comparative study of CNN, LSTM, FNN and Bi-LSTM. *Spectrum of Engineering Sciences*, 3(4), 582–596.
- Sharmeen, S., Weng, S. H., et al. (2020). Android malware detection: A survey on machine learning and deep learning approaches. *IEEE Access*, 8, 167821–167841.
- Shaukat, K., et al. (2020). A review of artificial intelligence techniques for malware detection. *Future Internet*, 12(11), 184.
- Xu, W., Qi, Y., & Evans, D. (2016). Automatically thwarting unknown malware by sequence learning. *arXiv preprint arXiv:1605.04321*.
- Yerima, S., & Sezer, S. (2013). A novel Android malware detection approach using Bayesian classification. In *IEEE 27th International Conference on Advanced Information Networking and Applications (AINA)* (pp. 1201–1208). IEEE.
- Yuan, X., He, P., Zhu, Q., & Li, X. (2019). Adversarial examples: Attacks and defenses for deep learning. *IEEE Transactions on Neural Networks and Learning Systems*, 30(9), 2805–2824.