



American Journal of Innovation in Science and Engineering (AJISE)

ISSN: 2158-7205 (ONLINE)

VOLUME 4 ISSUE 3 (2025)



PUBLISHED BY
E-PALLI PUBLISHERS, DELAWARE, USA

Ai-Driven Threat Detection and Prevention in Cloud Computing Environments

Ifeoma Eleweke^{1*}, Michael Friday Umakor², Chimdi Walter Ndubuisi³, Clifford Godwin Amomo⁴, Samuel Adeniji⁴,
Modupe Temidayo⁶

Article Information

Received: May 07, 2025

Accepted: June 11, 2025

Published: October 25, 2025

Keywords

AI-Driven Cybersecurity, Cloud Security, Hybrid AI Models, Machine Learning, Reinforcement Learning, Supervised Learning, Threat Detection, Unsupervised Learning

ABSTRACT

Cloud computing has become a cornerstone of modern IT infrastructure, offering scalability and efficiency but also exposing organizations to evolving cyber threats such as data breaches, insider threats, and advanced persistent threats (APTs). Traditional security mechanisms struggle to address these dynamic challenges, necessitating the integration of AI-driven threat detection and prevention strategies. This conceptual paper explores the comparative effectiveness of supervised learning, unsupervised learning, reinforcement learning, and hybrid AI models in cloud security. Supervised learning excels in identifying known attack patterns, while unsupervised learning is crucial for detecting zero-day threats and anomalies. Reinforcement learning enables self-adaptive security measures, and hybrid models offer a comprehensive, multi-layered approach to cloud security. However, AI-driven cybersecurity faces significant challenges, including data privacy risks, bias in threat detection, adversarial AI attacks, and lack of model interpretability. Emerging AI trends such as federated learning, quantum security, and explainable AI (XAI) are shaping the future of cloud security, while regulatory frameworks like GDPR, NIST AI Risk Management, and the EU AI Act play a crucial role in standardizing ethical AI use. This study provides insights into the strengths, weaknesses, and future directions of AI-driven cloud security, offering recommendations for researchers, policymakers, and cybersecurity practitioners to enhance AI resilience against emerging threats.

INTRODUCTION

Cloud computing has reshaped modern IT platforms by providing both individuals and businesses with modular, affordable, and flexible solutions (Soni *et al.*, 2025). Advanced persistent threats (APTs), Distributed Denial-of-Service (DDoS) attacks, insider threats, and data breaches are merely a few of the security problems that have been brought on by its widespread use (Sharma, 2024). Cloud environments are dynamic, decentralized, and multi-tenant, making it difficult for traditional security measures like firewalls and signature-based intrusion detection systems to adjust. Due to the constant growth of cybercriminals' attack methods, cloud security solutions must go beyond static, rule-based defenses (Ganguli, 2024).

In recent times, AI is accepted as a significant component required for the understanding and management of cybersecurity. The study conducted by Gupta & Srivastava (2025) indicated that AI has the potential to provide avenues for advanced threat detection, predictive analytics, and automated response capabilities. Most AI-linked models employ machine learning (ML), deep learning (DL), and anomaly detection algorithms (ADA) to identify existing and emerging threats in real time, which are better technologies when compared to

the traditional security systems that are only functional for the identification of signatures for known attacks (Olowu *et al.*, 2024). AI operates to enhance cloud security and risk minimization through the automation of threat identification, acceleration of time and reduction of human intervention (David & Edoise, 2025). The operation of cloud security can be migrated by AI technologies from reactive defensive strategies to proactive and flexible protective measures (Oloyede, 2024).

Based on the literature, the demonstration of the economic significant and statistical significance of AI in cybersecurity have been observed with some positive outcomes. According to Gartner (2023), the complexity of cyber threats is currently growing at exponential rate and some advantages can be trapped through the boasting of the roles of AI and the expansion of worldwide AI-driven cybersecurity industry's growth from the previously estimation of \$17.4 billion in 2022 to expected value of \$46.3 billion by 2027. According to a Capgemini poll, 69% of businesses believe AI would be crucial in the near future for dealing with cyberattacks (Ajala *et al.*, 2024). According to a survey by IBM Security (2023), companies that use AI-powered security solutions reduce the average cost of data breaches by \$1.76 million when

¹ Department of Computer Science, Westcliff University, USA

² School of Computer Science, Western Illinois University, USA

³ Department of Electrical and Computer Science Engineering, University of Missouri-Columbia, USA

⁴ Department of Computer Science, Stephen F. Austin State University, USA

⁵ College of Business, Information Assurance, Bowie State University, USA

* Corresponding author's e-mail: cleweke.ifeoma.25@gmail.com

compared to those who use traditional security methods. AI-powered cloud security systems use a variety of ML methods including supervised, unsupervised learning and reinforcement learning for threat classification, anomaly detection and adaptive security responses respectively (Nwachukwu *et al.*, 2024). The listed techniques are fundamental for the enhancement of AI models to capacity to identify complex patterns that are associated with attack, and also identification of internal threats, and automate extensive security tasks (Alzaabi & Mehmood 2024). Also supports uninterrupted authentication and risk assessment based on user behavior analytics, AI enhances Zero Trust security frameworks. AI also has the capacity to thwart complex cyber threats because it can instantly assess large volume of cloud data (Olabanji *et al.*, 2024). Despite the above listed and other potential, AI-powered cloud security is still experiencing some challenges, such as explainability conundrums, adversarial AI attacks, and data privacy concerns. In addition, some complains upheld that security experts are currently facing challenges in justifying reasonable conclusions concerning AI models because most of them are black boxes (Ariyibi *et al.*, 2024). Additionally, scammers have started to take advantage of AI flaws using adversarial attacks, in which small changes to input data can fool AI-powered threat detection systems (Ajayi *et al.*, 2024). It therefore follows that persistent advancements in explainable AI (XAI), federated learning (FL) engaged for privacy-preserving security, as well as the resilience of AI models against harsh threats are fundamental for resolving the associated challenges (Saeed & Alsharidah 2024).

This study is a review design to and identify compare different threat detections that are powered by AI in cloud environments and to evaluate their effectiveness, limitations, and potential improvements. Through an analysis of supervised learning, unsupervised learning, reinforcement learning, and hybrid AI models, this study seeks to shed light on the developing field of AI-driven cybersecurity. It also examines emerging concerns, legislative frameworks, and trends that will affect AI-driven cloud security in the future. Businesses looking to strengthen their cloud security posture in the face of growing cyber threats must comprehend these components.

AI-Driven Threat Detection: Existing Approaches Supervised Learning Models in Cloud Security

Most AI methodology in cyber security is supervised learning, which trains models to recognize patterns of known attacks using labeled datasets (Hussain *et al.*, 2025). For classification tasks such as discriminating between malicious and normal network traffic, decision trees and random forests are commonly used (Sah & Venkatesh 2024). These models categorize threats based on trained patterns after security data decomposition to hierarchical structure, where every decision node represents a security attribute (Nnenna *et al.*, 2025). Decision trees are widely utilized in malware detection and intrusion detection

systems (IDS) due to their interpretability and efficiency (Mohale & Obagbuwa).

The SVM is a high ranked learning method that is applicable for binary classification issues in cybersecurity (Zada *et al.*, 2024). SVM apply mathematical models to transform security data to a high-dimensional space, with additional benefits of establishing an ideal boundary between permitted and malicious behavior. According to Salman *et al.* (2024), SVMs has the sophistication in the identification of spam filtering, phishing attacks, and behavioral threat detection. These and other capabilities makes SVMs to be hihly reliant on high-quality training sets with labels, rendering them weak in stopping emerging threats like zero-day attacks (Mohamed *et al.*, 2025).

Abdallah *et al.* (2024) states that supervised threat detection has been greatly improved with deep learning models and neural networks, which have the ability to detect intricate patterns of attacks in large cloud environments. The Recurrent Neural Networks (RNNs) and Convolutional Neural Networks (CNNs) are capable of handling sequential and spatial data and thus are appropriate to be used for anomaly detection in real-time network traffic and cloud logs (Abdallah *et al.*, 2024). In spite of these, DL models are demanding because of heavy computational weight and large-labeled datasets requirement. The scalability and deployment challenges of the method can also limit their performance. However, it is widely acknowledged as a significant AI powered cloud security because of its accuracy in the detection of known cyber threats (Abdallah *et al.*, 2024). 2.2 Unsupervised Learning Techniques for Anomaly Detection, In contrast to supervised learning, unsupervised learning operates with no labeled data sets, hence its suitability in detecting unknown threats and zero-day attacks in cloud environments (Sharma & Singh 2025).

Clustering methods like K-Means and DBSCAN process cloud security data by grouping similar activities and detecting anomalies (Artioli *et al.*, 2024). The highlighted methods are mostly applied in the identification of network anomaly, and can also assist businesses in discovering patterns that are strange and significant in signaling data breach or insider threat (Thapaliya & Gurung, 2025).

Zideh *et al.* (2024) argue that autoencoders, which are a form of neural network, are perhaps the best unsupervised technique currently found in cyber defense against anomalies.

Autoencoders compress input information to a reduced dimension and subsequently attempt to rebuild it (Chen *et al.*, 2025). When a difference is detected that is significant between the original and rebuilt data, the input is identified as a potential threat. The method is very effective at identifying advanced persistent threats (APTs) and intricate cyberattacks that bypass normal security measures (Ibrahim *et al.*, 2024). Yet, the challenge to enhance autoencoder models to minimize false positives persists.

Isolation forests is yet another leading unsupervised technique that operates through the random selection of features and splitting data points to pinpoint anomalies (Yepmo *et al.*, 2024).

Unlike conventional statistical methods, isolation forests are able to detect outliers in large cloud datasets without having prior knowledge of attack patterns. This flexibility renders unsupervised learning inevitable for autonomous threat detection systems to enable AI models to learn as new threats arise (Simanjuntak *et al.*, 2024).

Unsupervised models have high false-positive rates and need continuous tuning to become more reliable (Olateju *et al.*, 2024).

Reinforcement Learning & Hybrid Approaches In Cloud Security

Based on the documented article of Kheddar *et al.* (2024), reinforcement learning (RL) is a relatively recent and novel paradigm that represents adaptive security solutions because

it permits AI models to learn, correct and develop solutions through trial and error. However, predefined labels or patterns are not applicable to SL, USL and RL but operates on consistent interaction taffter several rewards and punishments (Mvula *et al.*, 2024). On the other hand, the Deep Q-Networks (DQN) and Policy Gradient approaches are key functions for the automation of intrusion prevention systems (IPS). The IPS can further enhance dynamic responses of AI-based models to real time cyber threats (Louati *et al.*, 2024).

The RL can combat attacks (especially, zero day attacks and polymorphic malware) spontaneously because it can adjust to fresh cyber threats, attacks. or example, reinforcement learning-driven firewall management systems can independently modify security settings according to detected attack patterns (Pham *et al.*, 2024). The application of reinforcement learning in cloud security necessitates substantial training data and computational resources, potentially hindering deployment (Byatarayanapura *et al.*, 2024). Moreover, the design of the incentive function is crucial, as inadequately described reward structures might result in unanticipated security vulnerabilities (Miao *et al.*, 2025).

Hybrid AI models, integrating supervised, unsupervised, and reinforcement learning methodologies, signify the forthcoming advancement in AI-enhanced cloud security (Olowu *et al.*, 2024). These models utilize the advantages of each methodology: Supervised learning for the classification of known threats, unsupervised learning for anomaly detection, and reinforcement learning for real-time response and decision-making. Hybrid security systems, through the integration of various AI paradigms, provide enhanced and adaptable protection, minimizing false positives and detection time (Hernández-Rivas *et al.*, 2024). As cloud risks persist in their evolution, hybrid AI architectures are anticipated to establish the benchmark for intelligent threat prevention (Adeusi *et al.*, 2024).

Comparative Analysis of AI Methods

Strengths and Weaknesses of Each AI Approach

Based on the publication by Marengo *et al.* (2024), SL are effective, reliable and accurate for the identification of established threats including malware, phishing attempts, and intrusion attempts detection. The expected level of confidence that SL can generate implies that cybersecurity experts can take advantage of both automation and detection levels to design outcome that is characterized by confidence having minimal error margin. (Sarker *et al.*, 2024). However, the limitation can be concerned with reliance on prior history, such that the recognition of new or emerging threats, for instance, zero-day attacks, may be distorted. Moreover, supervised learning models require large labeled datasets, which dynamic cloud systems might not always provide (Ahmed, 2024).

Kaliyaperumal *et al.* (2024) contends that unsupervised learning techniques are especially effective in identifying new threats as they do not depend on pre-labeled attack signatures. They evaluate security data to discover anomalies, rendering them particularly successful in detecting zero-day vulnerabilities, insider threats, and atypical network behavior (Olawale *et al.*, 2024). Nonetheless, their principal constraint is an elevated false-positive rate. These algorithms categorize deviations as potential threats, occasionally misinterpreting valid yet atypical activity as security risks, resulting in superfluous alarms and an augmented workload for security professionals. Moreover, refining unsupervised models to enhance their precision might be arduous (Ying *et al.*, 2024).

Reinforcement learning (RL) and hybrid methodologies integrate the advantages of both supervised and unsupervised learning, incorporating an adaptive element. Reinforcement learning-based systems perpetually enhance and refine their danger detection capacities through real-time feedback (Shehzadi, 2024). This renders them exceptionally efficient for automated security responses in cloud environments. Nonetheless, reinforcement learning is computationally demanding and necessitates considerable training duration (Stranieri *et al.*, 2024). Moreover, establishing an optimal incentive structure for security risks is intricate, and inadequately built reinforcement learning models may misidentify threats or lack generalizability across many attack situations. Hybrid methodologies seek to address these deficiencies by amalgamating various AI strategies, hence enhancing accuracy and flexibility (Azevedo *et al.*, 2024).

Performance Metrics: Accuracy, False Positives, and Real-Time Efficiency

The effectiveness of AI-driven cloud security models is typically measured using three key performance metrics:

1. Accuracy – The percentage of correctly identified threats.
2. False Positives – The rate at which normal activities are incorrectly flagged as threats.
3. Real-Time Efficiency – The speed at which the model can detect and respond to security incidents.

Table 1: Physical, chemical and biological properties of experimental soil (0-20 cm)

AI Approach	Accuracy	False Positives	Real-Time Efficiency
Supervised Learning	High (85–95%)	Low (5–10%)	Moderate (relies on dataset size)
Unsupervised Learning	Moderate (65–85%)	High (20–30%)	High (fast detection but needs tuning)
Reinforcement Learning	Very High (90–98%)	Moderate (10–15%)	High (self-optimizing, real-time response)
Hybrid Models	Very High (95–99%)	Low (5–10%)	Very High (leverages multiple techniques)

Supervised models offer the best accuracy for known threats, while reinforcement learning and hybrid models provide the best real-time efficiency and adaptability. However, unsupervised learning methods struggle with false positives, which can impact security response effectiveness.

Best Use Cases for Each AI Method

Each AI approach is best suited for specific cybersecurity challenges in cloud environments:

Supervised Learning Best Use Cases

- Malware & Phishing Detection – Recognizing known malicious URLs, emails, and attack signatures.
- Intrusion Detection Systems (IDS) – Identifying predefined attack patterns in cloud traffic.
- Cloud Compliance Monitoring – Automating security audits for regulatory compliance.

Unsupervised Learning Best Use Cases

- Zero-Day Threat Detection – Identifying unknown attack behaviors in cloud environments.
- Insider Threat Detection – Monitoring user behavior anomalies in cloud systems.
- Anomaly-Based Network Security – Flagging unusual data flows or access patterns.

Reinforcement Learning & Hybrid Models Best Use Cases

- Automated Threat Response – Adaptive security measures that improve over time.
- Self-Healing Security Systems – AI models that learn from past attacks to enhance future defenses.
- Dynamic Cloud Security Policy Enforcement – Adjusting firewall and authentication policies based on evolving threats.

Challenges In AI-Driven Threat Detection

Data Privacy and Bias Concerns

According to Arif and others (2024), data privacy is among the stumbling blocks against the optimum performance of AI-based threat detection in the cloud. The requirements for the flexibility of AI algorithms is humongous volumes of data for training and actual threat identification. According to Adako *et al.* (2024), the success of AI-powered model also requires the review of some private information such as network traffic, user behavior logs, and login credentials (Adako *et al.*, 2024). These factors can lead to data exposure, legal obligations,

and potential abuse. However, the requirements for stringent data protection regulations (such as the California Consumer Privacy Act (CCPA) and the General Data Protection Regulation (GDPR)) could be challenged but are mandatory for an organizations to achieve AI-driven security and legal compliance. One of such is the tendency for an AI-based security to ignorantly process personally identifiable information (PII) without adequate safeguards, leading to a compromise of the user privacy rights.

AI model bias can also be challenged if biased training data sets are involved. A skew view of cyber threat can be displayed by AI model if the training data was specific to a given industries or attacks (Nnenna *et al.*, 2025). For example, an AI model trained on Western enterprise data may not comprehends Asian cloud network attack behaviors. An AI system running with historical attack data may also amplify biases to certain types of threats because it lacks the ingenuity of emerging cyber threats. Challenges concerning privacy and bias can be resolved by an organization through committed employment of privacy-preserving AI methods such as FL, differential privacy, and secure multi-party computation. FL can permit AI models to learn on decentralized datasets without raw data leakage. In addition, cybersecurity teams must also prioritize dataset diversity and bias auditing to ensure AI models are trained on a variety of threat scenarios across industries and geographies.

Interpretability and Explainability Problems are Associated with the Black Box Identity of AI Models

For example, models based in DL and neural network could be unintelligible in decision making pattern. Therefore, why detection is commendable, understanding the conditions surrounding the detection may be difficult. The consequences are lack of transparency, associated issues with trust and unreliability of AI model to secure the cloud.

The EU's AI Act , the NIST AI Risk Management Framework and other regulatory frameworks is necessary in the enforcement of cyber security. Explainability becomes important in forensic analysis, compliance audit, and judicial proceedings, where organizations need transparent reasons behind their decisions on security. If an AI model mistakenly denies a legitimate user's cloud access or marks normal network traffic as an attack, security teams must understand and fix the problem in a timely manner. But deep learning models for threat detection lack transparency inherently, and therefore

debugging problems and improving decision-making is difficult.

To enhance explainability of AI for cloud security, researchers and practitioners are developing XAI techniques, such as SHAP (SHapley Additive Explanations), LIME (Local Interpretable Model-agnostic Explanations), and attention-based models. They provide human-understandable explanations of AI decisions, and security professionals can view why an attack was identified or why a user had access denied. Improving interpretability is essential to building trust in AI-driven cybersecurity and ensuring AI models are commensurate with human knowledge and decision-making.

Adversarial AI Threats

Cyber attackers are crafting new attack strategies to take advantage of weaknesses in AI as AI-based cybersecurity measures enhance. In adversarial AI attacks, well-crafted inputs are submitted into AI models to manipulate them and mislead security mechanisms into creating faulty conclusions. These types of attacks are a threat to cloud security because they can evade malware classifiers, anomaly detection techniques, and AI-powered intrusion detection systems (IDS). Adversarial examples make AI models misclassify threats without raising an alarm by introducing subtle data input changes.

Evasion attacks, in which attackers tamper with malware signatures or traffic patterns in networks to deceive AI models that they are safe, is a typical example of adversarial AI in cybersecurity. A malware detection that is based on AI, for example, might not be able to detect a malware sample if it contains minor variations in its binary code. Likewise, attackers may provide tampered login requests and user behavior information to disrupt AI-based authentication systems, thereby being a significant threat to the cloud access security. Poisoning is another adversarial AI manipulation. It is a type of feeding malicious content into the training data of an AI system to taint future decisions.

Responding to the progressing adversarial AI threats, cybersecurity researchers are developing robust AI models capable of detecting and resisting adversarial manipulation. Techniques like adversarial training, input sanitization, and model uncertainty estimation enhance the resilience of AI-driven security systems. In addition, based on the real-time identification of potential adversarial attacks by continuous monitoring and anomaly detection techniques, security teams can respond to damage before it takes place. Since AI weaknesses are increasingly being exploited by cybercriminals, developing defensive AI solutions will be crucial for safe and resilient cloud environments.

The Future of AI in Cloud Security

Emerging AI Trends in Cloud Security

AI-driven security systems must adapt to improve threat detection, guard against privacy, and implement real-time reaction mechanisms as threats change and become

more complex. Most likely one of the most thrilling new developments in AI security is federated learning, a decentralized model for training AI models without transmitting raw data to a central server. This method solves data sovereignty and regulatory compliance problems by allowing cloud-based security systems to collectively work on threat detection while maintaining user privacy. In multi-cloud environments where several enterprises have to exchange security intelligence without revealing sensitive data, federated learning is especially useful.

Convergence of quantum security techniques with AI-powered cloud security solutions is another crucial trend. As quantum computing continues to mature, classical encryption technologies may be susceptible to being broken by quantum-powered solution to these challenges. The search for these challenges is being pioneered in several researches that are considering quantum-resistant encryption algorithms and AI-powered quantum security solutions. A significant attribute is the employment of the post-quantum cryptography that is controlled by AI for actual selection of quantum-resistant encryption algorithms after learning from existing threat analysis. The quantum AI models can also lead to adaptive and resilience security solutions by improving anomaly detection in cloud networks.

XAIs are also effective in the improvement of cloud security through the provision of solutions that upturn the challenges associated with the traditional AI interpretability.

XAI techniques, such as SHAP (Shapley Additive Explanations) and LIME (Local Interpretable Model-Agnostic Explanations), support the cybersecurity analysts in understanding, validations and refinement of threat response. The contributions of XAI is fundamentally significant for business security alignment, regulatory bodies, and forensic analysis. This will provide assurance that the services of AI models projects beyond the black boxes to the provision of an open explanations for security decisions.

Finally, AI-driven autonomous security operations (SecOps) are useful when self-learning AI models are applied to automate security operations.

The Role of Regulations and Standardization in AI Security

As cloud security solutions powered by AI proliferate, governments and regulatory agencies are striving to create guidelines and regulations to guarantee the ethical, open, and efficient use of AI. Laws like the EU Artificial Intelligence Act, the NIST AI Risk Management Framework, and industry-specific security standards (such PCI DSS for finance and HIPAA for healthcare) are influencing the development, application, and auditing of AI-powered security solutions. By highlighting the significance of equity, responsibility, and openness in AI-powered cybersecurity, these frameworks make sure that AI models don't add prejudice, invasions of privacy, or security flaws.

Data privacy and compliance are among the most important aspects of AI legislation. Data protection regulations like the CCPA, GDPR, and ISO/IEC 27001 must be followed by businesses implementing AI-based cloud security solutions to make sure AI models don't abuse user privacy or misuse personal information. AI-driven security frameworks are using standardized privacy-preserving AI approaches, such as homomorphic encryption, federated learning, and differential privacy, to meet stringent regulatory standards while preserving strong threat detection capabilities.

Another major focus of AI security regulations is adversarial robustness and model security. Cybercriminals are increasingly targeting AI-driven security systems with adversarial attacks, manipulating AI models to bypass security controls. Regulatory bodies are pushing for standardized testing, certification, and adversarial defense strategies to ensure that AI models deployed in cloud security systems are resilient against adversarial manipulation. The MITRE ATLAS framework, for example, provides a benchmark for assessing AI security vulnerabilities, helping organizations enhance AI robustness.

Looking ahead, global cooperation on AI security standardization will be essential for ensuring interoperability and trust in AI-powered cloud security. Initiatives like the OECD AI Principles and the United Nations' AI governance discussions are working to establish global best practices for AI security, ensuring that AI-driven threat detection systems operate ethically, securely, and in compliance with international laws. As AI continues to transform cloud security, clear regulatory frameworks and industry-wide security standards will play a pivotal role in shaping the future of AI-driven cybersecurity.

CONCLUSION

The recent global digital revolution is also advancing cloud security through AI-enhanced threat detection, improvement of predictive analytics, reaction automation, as well as threat detection. Current global position is the expanding rate and advancement of cyber attacks over the traditional security measures. This makes AI to be a fundamental tool in facing the reality regarding the protection and defence of the cloud and cyber spaces. Unsupervised learning are useful for the identification of emerging abnormalities, SL are fundamental for labeling existing attacks, and RL becomes inevitable because of its self-adaptive security features. Currently, limitations to AI cyber security are also challenging some success for example, model bias, adversarial attacks, privacy issues, and explainability issues. However, with regulation and compliance, normalize ethics-based AI deployment, innovative concepts such as explainable AI, FL, and quantum security, most of these limitations can be brought under controlled.

Scientists must develop privacy-reassuring AI techniques, such as differential privacy, homomorphic encryption, and

federated learning, to allow safe AI training at the cost of not revealing sensitive data. Development of explainable AI (XAI) techniques will also strengthen transparency, trust, and the adoption of AI-based threat detection systems. Defense mechanisms against adversarial AI must also be strengthened using AI training with robustness, adversarial detection, and continuous monitoring. Research must also explore hybrid AI models that combine supervised, unsupervised, and reinforcement learning to possess more adaptive, efficient, and scalable cloud security solutions.

Cloud and cyber security experts should supervise AI-driven security technologies in order to bridge the alignment between automated techniques and human decision-making. The avoidance of model drift and adversarial attacks can be achieved by AI-based models are continuously monitored, restored, and augmented. Privacy-first AI adoption entails making organizations GDPR, NIST, and new AI security standards compliant from the outset. Lastly, to provide proactive, cloud-scale, and adaptive security for the cloud and reduce response times to cyber threats, practitioners need to place greatest priority on AI-based real-time security orchestration.

REFERENCES

- Abdallah, A., Alkaabi, A., Alameri, G., Rafique, S. H., Musa, N. S., & Murugan, T. (2024). Cloud network anomaly detection using machine and deep learning techniques: Recent research advancements. *IEEE Access*, 12, 56749–56773. <https://doi.org/10.1109/ACCESS.2024.3390844>
- Adako, O., Adeusi, O., & Alaba, P. (2024). Integrating AI tools for enhanced autism education: A comprehensive review. *International Journal of Developmental Disabilities*, 1–13. <https://doi.org/10.1080/20473869.2024.2392983>
- Adeusi, O. C., Adebayo, Y. O., Ayodele, P. A., Onikoyi, T. T., Adebayo, K. B., & Adenekan, I. O. (2024). IT standardization in cloud computing: Security challenges, benefits, and future directions. *World Journal of Advanced Research and Reviews*, 22(3), 2050–2057.
- Ahmed, Q. O. (2024). Machine learning for intrusion detection in cloud environments: A comparative study. *Journal of Artificial Intelligence General Science (JAIGS)*, 6(1), 550–563. <https://doi.org/10.60087/jaigs.v6i1.287>
- Ajala, O. A., Okoye, C. C., Ofodile, O. C., Arinze, C. A., & Daraojimba, O. D. (2024). Review of AI and machine learning applications to predict and thwart cyberattacks in real time. *Magna Scientia Advanced Research and Reviews*, 10(1), 312–320.
- Ajayi, A. M., Omokanye, A. O., Olowu, O., Adeleye, A. O., Omole, O. M., & Wada, I. U. (2024). *Detecting insider threats in banking using AI-driven anomaly detection: A data science approach to cybersecurity*. Unpublished manuscript.
- Alzaabi, F. R., & Mehmood, A. (2024). A review of recent advances, challenges, and opportunities in malicious insider-threat detection using machine learning methods. *IEEE Access*, 12, 30907–30927. <https://doi.org/10.1109/ACCESS.2024.3390844>

- org/10.1109/ACCESS.2024.3372277
- Arif, H., Kumar, A., Fahad, M., & Hussain, H. K. (2024). Future horizons: AI-enhanced threat detection in cloud environments—Unveiling opportunities for research. *International Journal of Multidisciplinary Sciences and Arts*, 3(1), 242–251.
- Ariyibi, K. O., Bello, O. F., Ekundayo, T. F., & Ishola, O. (2024). Leveraging artificial intelligence for enhanced tax fraud detection in modern fiscal systems. *GSC Advanced Research and Reviews*, 21(2), 129–137. <https://doi.org/10.30574/gscarr.2024.21.2.0415>
- Artioli, P., Maci, A., & Magri, A. (2024). A comprehensive investigation of clustering algorithms for user and entity behavior analytics. *Frontiers in Big Data*, 7, 1375818. <https://doi.org/10.3389/fdata.2024.1375818>
- Azevedo, B. F., Rocha, A. M. A., & Pereira, A. I. (2024). Hybrid approaches to optimization and machine learning methods: A systematic literature review. *Machine Learning*, 113(7), 4055–4097. <https://doi.org/10.1007/s10994-023-06467-x>
- Byatarayanapura Venkataswamy, S., Patil, K. S., Narayanaswamy, H. K., & Veerabadrappa, K. (2024). Access management based on deep reinforcement learning for effective cloud storage security. *International Journal of System Assurance Engineering and Management*, 15(8), 1–20. <https://doi.org/10.1007/s13198-024-02596-1>
- Chen, J., Höhle, K., & Lerch, S. (2025). Learning low-dimensional representations of ensemble forecast fields using autoencoder-based methods. *arXiv*. <https://arxiv.org/abs/2502.04409>
- David, A. A., & Edoise, A. (2025). Cloud computing and machine learning for scalable predictive analytics and automation: A framework for solving real-world problems. *Communications in Physical Sciences*, 12(2), 406–416. <https://dx.doi.org/10.4314/cps.v12i2.16>
- Ganguli, P. (2024). *The rise of cybercrime-as-a-service: Implications and countermeasures* (SSRN Working Paper No. 4959188). SSRN. <https://ssrn.com/abstract=4959188>
- Gupta, R., & Srivastava, P. (2025). Artificial intelligence and machine learning in cybersecurity applications. In *Cyber security solutions for protecting and building the future smart grid* (pp. 271–296). Elsevier. <https://doi.org/10.1016/B978-0-443-14066-2.00004-9>
- Hernández Rivas, A., Morales Rocha, V., & Sánchez Solís, J. P. (2024). Towards autonomous cybersecurity: A comparative analysis of agnostic and hybrid AI approaches for advanced persistent threat detection. In *Innovative applications of artificial neural networks to data analytics and signal processing* (pp. 181–219). Springer, Cham. https://doi.org/10.1007/978-3-031-69769-2_8
- Hussain, H., Kainat, M., & Ali, T. (2025). Leveraging AI and machine learning to detect and prevent cybersecurity threats. *Dialogue Social Science Review*, 3(1), 881–895.
- Ibrahim, N., Rajalakshmi, N. R., & Hammadeh, K. (2024). Exploration of defensive strategies, detection mechanisms, and response tactics against advanced persistent threats (APTs). *Nanotechnology Perceptions*, 20(S4), 439–455. <https://doi.org/10.62441/nano-ntp.v20iS4.33>
- Kaliyaperumal, P., Periyasamy, S., Thirumalaisamy, M., Balusamy, B., & Benedetto, F. (2024). A novel hybrid unsupervised learning approach for enhanced cybersecurity in the IoT. *Future Internet*, 16(7), 253. <https://doi.org/10.3390/fi16070253>
- Kheddar, H., Dawoud, D. W., Awad, A. I., Himeur, Y., & Khan, M. K. (2024). Reinforcement learning based intrusion detection in communication networks: A review. *IEEE Communications Surveys & Tutorials*, 27(4), 2420–2469. <https://doi.org/10.1109/COMST.2024.3484491>
- Louati, F., Ktata, F. B., & Amous, I. (2024). Enhancing intrusion detection systems with reinforcement learning: A comprehensive survey of RL based approaches and techniques. *SN Computer Science*, 5(6), 665. <https://doi.org/10.1007/s42979-024-03001-1>
- Marengo, A., & Pagano, A. (2024). Machine learning for cybersecurity for detecting and preventing cyber attacks. *Machine Intelligence Research*, 18(1), 672–689. <https://doi.org/10.1016/j.mir.2023.11.015>
- Miao, Y., Zhang, S., Ding, L., Bao, R., Zhang, L., & Tao, D. (2025). INFORM: Mitigating reward hacking in RLHF via information theoretic reward modeling. In *Advances in Neural Information Processing Systems*, 37 (pp. 134387–134429).
- Mohale, V. Z., & Obagbuwa, I. C. (2025). A systematic review on the integration of explainable artificial intelligence in intrusion detection systems: Enhancing transparency and interpretability in cybersecurity. *Frontiers in Artificial Intelligence*, 8, 1526221. <https://doi.org/10.3389/frai.2025.1526221>
- Mohamed, A. A., Al Saleh, A., Sharma, S. K., & Tejani, G. G. (2025). Zero day exploits detection with adaptive WavePCA Autoencoder (AWPA) adaptive hybrid exploit detection network (AHEDNet). *Scientific Reports*, 15(1), 4036. <https://doi.org/10.1038/s41598-025-87615-2>
- Mvula, P. K., Branco, P., Jourdan, G. V., & Viktor, H. L. (2024). A survey on the applications of semi supervised learning to cyber security. *ACM Computing Surveys*, 56(10), 1–41. <https://doi.org/10.1145/3657647>
- Nnenna, J. O., Olaoye, S. A., & Samuel, A. A. (2025). Enhancing cybersecurity in communication networks using machine learning and AI: A case of 5G infrastructure security. *World Journal of Advanced Research and Reviews*, 26(1), 1210–1219. <https://doi.org/10.30574/wjarr.2025.26.1.1098>
- Nnenna, J. O., Adesola, A. A., Samuel, A. A., & Rhoda, K. T. (2025). Federated learning for privacy preserving data analytics in mobile applications. *World Journal of Advanced Research and Reviews*, 26(1), 1220–1232. <https://doi.org/10.30574/wjarr.2025.26.1.1099>
- Nwachukwu, C., Durodola Tunde, K., & Akwiwu Uzoma, C. (2024). *AI driven anomaly detection in cloud computing environments*. Unpublished manuscript.

- Olabanji, S. O., Marquis, Y., Adigwe, C. S., Ajayi, S. A., Oladoyinbo, T. O., & Olaniyi, O. O. (2024). AI driven cloud security: Examining the impact of user behavior analysis on threat detection. *Asian Journal of Research in Computer Science*, 17(3), 57–74. <https://doi.org/10.9734/ajrcos/2024/v17i3301>
- Olateju, O., Okon, S. U., Igwenagu, U., Salami, A. A., Oladoyinbo, T. O., & Olaniyi, O. O. (2024). *Combating the challenges of false positives in AI driven anomaly detection systems and enhancing data security in the cloud*. SSRN. <https://ssrn.com/abstract=4859958>
- Olawale, A., Ajoke, O., & Adeusi, C. (2020). Quality assessment and monitoring of networks using a passive technique. *Review of Computer Engineering Research*, 7(2), 54–61. <https://doi.org/10.18488/journal.76.2020.72.54.61>
- Olowu, O., Adeleye, A. O., Omokanye, A. O., Ajayi, A. M., Adepoju, A. O., Omole, O. M., & Chianumba, E. C. (2024). AI driven fraud detection in banking: A systematic review of data science approaches to enhancing cybersecurity. *GSC Advanced Research and Reviews*, 21(2), 227–237. <https://doi.org/10.30574/gscarr.2024.21.2.0418>
- Oloyede, J. (2024). *AI driven cybersecurity solutions: Enhancing defense mechanisms in the digital era*. SSRN. <https://ssrn.com/abstract=4976103>
- Pham, V. H., Do Hoang, H., Trung, P. T., To, T. N., & Duy, P. T. (2024). Raiju: Reinforcement learning guided post exploitation for automating security assessment of network systems. *Computer Networks*, 253, 110706. <https://doi.org/10.1016/j.comnet.2024.110706>
- Saeed, M. M., & Alsharidah, M. (2024). Security, privacy, and robustness for trustworthy AI systems: A review. *Computers & Electrical Engineering*, 119, 109643. <https://doi.org/10.1016/j.compeleceng.2024.109643>
- Sah, A. K., & Venkatesh, K. (2024, April). Anomaly based intrusion detection in network traffic using machine learning: A comparative study of decision trees and random forests. In *2024 2nd International Conference on Networking and Communications (ICNW/C)* (pp. 1–7). IEEE.
- Salman, M., Ikram, M., & Kaafar, M. A. (2024). Investigating evasive techniques in SMS spam filtering: A comparative analysis of machine learning models. *IEEE Access*, 12, 24306–24324. <https://doi.org/10.1109/ACCESS.2024.3364671>
- Sarker, I. H., Janicke, H., Mohsin, A., Gill, A., & Maglaras, L. (2024). Explainable AI for cybersecurity automation, intelligence, and trustworthiness in digital twin: Methods, taxonomy, challenges, and prospects. *ICT Express*. <https://doi.org/10.1016/j.icte.2024.10.012>
- Sharma, A., & Singh, U. K. (2025). Cloud computing security through detection and mitigation of zero day attacks using machine learning techniques. In *Natural Language Processing for Software Engineering* (pp. 357–388). <https://doi.org/10.1016/B978-0-443-25665-1.00014-7>
- Sharma, H. (2024). The evolution of cybersecurity challenges and mitigation strategies in cloud computing systems. *International Journal of Computer Engineering and Technology*, 15(4), 118–127. <https://doi.org/10.5281/zenodo.13140593>
- Shehzadi, T. (2024). Reinforcement learning based autonomous systems for cyber threat detection and response. *Eastern European Journal for Multidisciplinary Research*, 1(1), 123–137.
- Simanjuntak, T. (2024). Emerging cybersecurity threats in the era of AI and IoT: A risk assessment framework using machine learning for proactive threat mitigation. *International Journal of Information System and Innovative Technology*, 3(1), 15–22.
- Soni, R., Bhatia, K., & Rajput, N. (2025). A thorough analysis of cloud computing technology: Present, past, and future. In *Recent advances in sciences, engineering, information technology & management* (pp. 137–145). CRC Press. <https://doi.org/10.1201/9781003598152-19>
- Stranieri, F., Fadda, E., & Stella, F. (2024). Combining deep reinforcement learning and multi stage stochastic programming to address the supply chain inventory management problem. *International Journal of Production Economics*, 268, 109099. <https://doi.org/10.1016/j.ijpe.2023.109099>
- Thapaliya, S., & Gurung, M. R. C. (2025). Mitigating insider threats and data breaches: Enhancing data loss prevention systems with behavioral analytics and NLP. *International Journal of Multidisciplinary and Interdisciplinary Research*, 2(1).
- Yepmo, V., Smits, G., Lesot, M. J., & Pivert, O. (2024). Leveraging an isolation forest to anomaly detection and data clustering. *Data & Knowledge Engineering*, 151, 102302. <https://doi.org/10.1016/j.datak.2024.102302>
- Ying, W., Wang, D., Hu, X., Zhou, Y., Aggarwal, C. C., & Fu, Y. (2024, August). Unsupervised generative feature transformation via graph contrastive pre training and multi objective fine tuning. In *Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining* (pp. 3966–3976). <https://doi.org/10.1145/3637528.3672015>
- Zada, I., Alatawi, M. N., Saqlain, S. M., Alshahrani, A., Alshamran, A., Imran, K., & Alfraihi, H. (2024). Fine tuning cyber security defenses: Evaluating supervised machine learning classifiers for Windows malware detection. *Computers, Materials & Continua*, 80(2), 2917–2939. <https://doi.org/10.32604/cmc.2024.052835>
- Zideh, M. J., Khalghani, M. R., & Solanki, S. K. (2024). An unsupervised adversarial autoencoder for cyber attack detection in power distribution grids. *Electric Power Systems Research*, 232, 110407. <https://doi.org/10.1016/j.epsr.2024.110407>