

ARTIFICIAL INTELLIGENCE IN CYBERSECURITY: APPLICATIONS AND CHALLENGES

Loredana MOCEAN

loredana.mocean@ubbcluj.ro

BABEȘ-BOLYAI UNIVERSITY, CLUJ-NAPOCA, ROMANIA

Miranda-Petronella VLAD

miranda.vlad@cantemircluj.ro

“DIMITRIE CANTEMIR” CHRISTIAN UNIVERSITY, BUCHAREST, ROMANIA

ABSTRACT

This paper explores the integration of Artificial Intelligence (AI) into cybersecurity, highlighting its role in detecting, preventing, and responding to increasingly sophisticated cyber threats. A review of the state of the art demonstrates significant progress in anomaly detection, explainable AI (XAI), and resilience against adversarial attacks. The study examines practical applications of AI across multiple domains, including threat detection, automated incident response, vulnerability analysis, anti-phishing systems, and adaptive authentication. Commercial platforms such as Darktrace, Vectra AI, and Microsoft Defender showcase the real-world impact of AI-driven security solutions. Additionally, relevant research initiatives, such as the DARPA Cyber Grand Challenge and EU's AI4Cyber project, emphasize the growing importance of AI in protecting critical infrastructures. A case study on detecting and containing an Advanced Persistent Threat (APT) in the financial sector further illustrates the effectiveness of unsupervised learning and automated response mechanisms. Despite these advances, challenges remain, including false positives, lack of transparency, adversarial vulnerabilities, and ethical considerations. The paper concludes by identifying future directions, such as the development of robust explainable models, integration of generative AI, and closer collaboration between human experts and AI systems.

KEYWORDS:

anti-phishing, cybersecurity, fraud prevention, vulnerability analysis, explainable AI, adversarial attacks

1. Introduction and State of the Art

Cybersecurity is one of the most dynamic and challenging branches of information technology in today's digital era. With the increasing complexity of cyberattacks and the growing number of network-connected devices, organizations face significant difficulties in protecting their data and IT infrastructures. Artificial

Intelligence (AI) is emerging as a promising solution, capable of detecting, preventing, and responding to threats with a speed and accuracy that traditional methods cannot achieve.

The aim of this paper is to analyze how AI is integrated into cybersecurity, exploring practical applications, relevant commercial platforms, and emerging research directions.

The specialized literature reveals a substantial increase in research on the application of AI in cybersecurity, with a focus on anomaly detection, model explainability (XAI), and resilience against sophisticated attacks.

Nguyen et al. (2023) propose the MAIP (Montimage AI Platform), a deep learning-based solution for network anomaly detection and AI decision explainability. The platform stands out by integrating a graphical interface that assists operators in understanding model behavior.

Paper of Khalaf et al. (2024) analyze a network traffic anomaly detection system based on convolutional neural networks (CNNs). The model achieves an accuracy of 95.6% in classifying network behavior as benign or malicious, demonstrating the effectiveness of deep learning in analyzing large-scale data.

In the paper of Mutalib et al. (2024) the authors investigate the use of XAI techniques (LIME, SHAP) for detecting Advanced Persistent Threats (APT). The results suggest that incorporating explainability can significantly increase human analysts' trust in automated system recommendations.

The review study of Chandola et al. (2021) provides a comprehensive classification of deep learning techniques applied to anomaly detection, highlighting current challenges such as the lack of labeled data, data imbalance, and vulnerability to adversarial attacks.

Salem et al. (2024) propose a comparative analysis of over 60 papers on AI techniques in cybersecurity, ranging from classical methods (SVM, KNN) to emerging approaches like GANs and hybrid deep ensembles. The paper also offers perspectives on integrating AI into industrial infrastructures.

Paper of Talukder et al. (2024) presents a methodology for intrusion detection in large-scale, imbalanced networks, based on unsupervised machine learning and oversampling techniques. The model achieves up to 99.1% accuracy, demonstrating the

efficiency of ensemble techniques and adaptive preprocessing.

Lately, Egelman has explored the legal side of privacy, consulting on cases and helping lawyers better understand technology. His forthcoming law review study examines company and data broker claims in privacy suits, revealing common misconceptions about data collection and sharing (Cunningham, 2025).

DARPA launched the Cyber Grand Challenge, a competition to create automatic defensive systems capable of reasoning about flaws, formulating patches, and deploying them on a network in real time. This growing focus on bridging technical expertise with practical and legal frameworks echoes the work of researchers like Egelman, who seek to close the gap between complex technical realities and broader institutional understanding.

The latest release of Darktrace/OT brings powerful new innovations to security teams defending industrial infrastructure (Pallavi, 2025).

Current research clearly indicates a trend toward robust, scalable, and explainable AI systems, adapted to both advanced threats and the trust requirements of human analysts. AI platforms and frameworks are becoming core components of modern cybersecurity architecture.

2. Fundamentals of Artificial Intelligence in Cybersecurity and Areas of Application

Artificial Intelligence encompasses a set of technologies that enable computer systems to mimic human intelligence, including machine learning algorithms, artificial neural networks (deep learning), natural language processing (NLP), and computer vision.

In cybersecurity, AI is applied to automate analysis tasks, detect abnormal behaviors, and make rapid decisions in the context of large volumes of unstructured data. The advantages include:

- The ability to process data at a large scale.

- Proactive detection of unknown (zero-day) threats.
- Adaptability to dynamic environments and sophisticated attacks.

2.1. Threat and Anomaly Detection

AI is used to analyze the behavior of networks, devices, and users to identify unusual activities. Machine learning models are trained on historical data to detect behavioral deviations that may indicate a potential intrusion (see Figure no. 1).

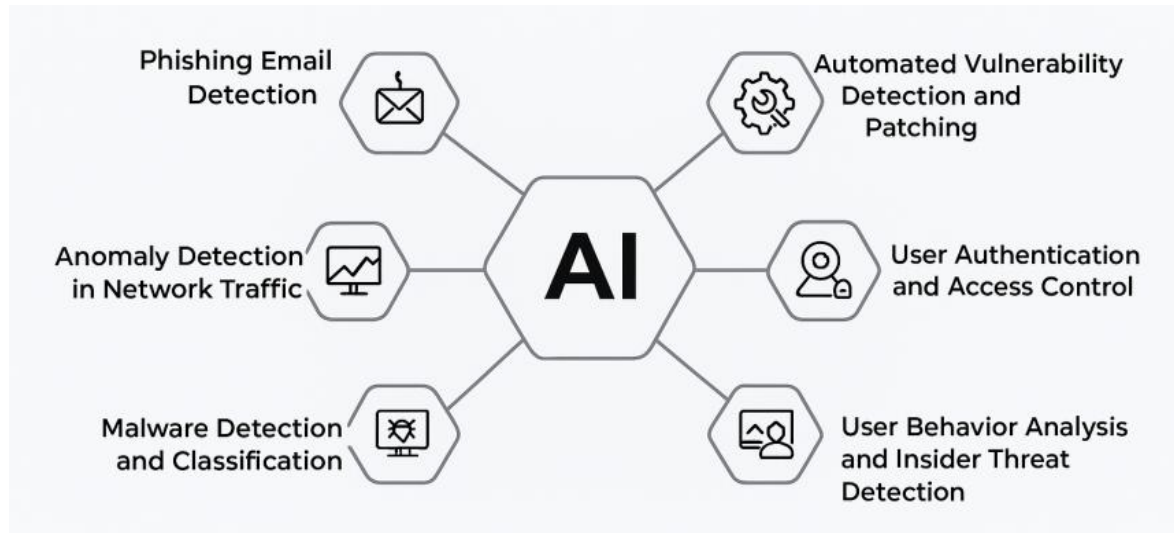


Figure no.1: *AI areas of application*

(Source: European School of Data Science and Technology, n.d.)

- *Example:* Darktrace employs a system called the Enterprise Immune System, based on unsupervised AI that learns the normal patterns of an organization and detects anomalies in real time – without predefined rules. It has been successfully used to detect ransomware and stealthy lateral movement attacks.

2.2. Incident Response Automation

By integrating with SOAR (Security Orchestration, Automation and Response) platforms, AI can coordinate threat response actions such as system isolation, traffic blocking, or automated alerts. This significantly reduces response time and minimizes human error.

- *Example:* Cortex XSOAR (by Palo Alto Networks) uses AI to automate security workflows. Upon detecting suspicious behavior, the platform can autonomously isolate affected endpoints, analyze suspicious files, and generate detailed incident reports for the SOC team.

2.3. Vulnerability Analysis and Prevention

AI can analyze source code, system configurations, or data streams to identify vulnerabilities. It also helps prioritize security patches based on actual risk.

- *Example:* Snyk and GitHub CodeQL use AI and machine learning to scan source code for known vulnerabilities and insecure coding patterns. AI can even suggest automatic fixes or highlight critical flaws based on the context of the application.

2.4. Anti-Phishing and Fraud Prevention

NLP algorithms can detect phishing emails through semantic and syntactic analysis. AI can also intercept fraudulent behavior in financial transactions or electronic communications.

- *Example:* Google Safe Browsing and Gmail AI Filters leverage NLP and

predictive models to identify phishing emails and malicious links. In the financial sector, Mastercard Decision Intelligence uses AI to detect fraudulent transactions in real time by analyzing user behavior patterns.

2.5. Authentication and Access Control

Biometric solutions (e.g. facial, voice recognition) and behavioral authentication rely on AI to ensure secure access.

– *Example:* Microsoft Azure AD implements adaptive authentication using AI: based on the user’s location, device, and historical behavior, the system may require additional verification steps or

automatically block access. Other platforms like BehavioSec use behavioral biometrics powered by AI – analyzing how users type, move the mouse, or interact with mobile devices.

3. Commercial Platforms Integrating AI in Cybersecurity

Several commercial platforms have emerged as leaders in applying artificial intelligence to cybersecurity, each addressing different aspects of digital protection. The following table provides an overview of some relevant platforms.

Table no. 1
Platforms for AI integrated in Cybersecurity

Platform	Main Purpose	AI Technology Used
Darktrace	Threat Detection	Unsupervised Machine Learning
Vectra AI	NDR + Cloud Security	Deep Learning
CrowdStrike	Endpoint Protection	Threat Graph AI
Cylance	Predictive Antivirus	Signatureless Machine Learning
Microsoft Defender	Analysis + Response	Microsoft Cloud AI

Each of these platforms demonstrates the real applicability of AI in different segments of cybersecurity, from endpoints to cloud networks. Together, these platforms illustrate the practical value of AI across diverse segments of cybersecurity, ranging from endpoint protection to cloud network defense.

4. Relevant Research Projects in the Field

Beyond commercial applications, numerous research initiatives and academic collaborations are actively shaping the future of AI-powered cybersecurity. These projects aim not only to improve detection and response capabilities but also to address critical challenges such as adversarial robustness, explainability, and ethical considerations.

One notable initiative is the AI4Cyber Project (2025) funded by the European Commission, which fosters the creation of a

European AI ecosystem focused on protecting critical infrastructures such as energy grids, transportation systems, and financial networks. The project integrates AI-based anomaly detection, predictive maintenance, and automated threat response to ensure system resilience in the face of emerging cyberattacks.

Another milestone effort is DARPA’s Cyber Grand Challenge (CGC), often considered a breakthrough in autonomous cybersecurity. During the challenge, AI systems competed to find, patch, and exploit software vulnerabilities in real time – without any human assistance. This experiment demonstrated that AI can operate at machine speed, both defending against and launching simulated attacks, opening new avenues for the future of autonomous cyber defense.

Academic research institutions also play a central role. MIT Lincoln Laboratory is pioneering work on adversarial machine learning, designing models that remain robust even when attackers deliberately

manipulate input data. Similarly, Carnegie Mellon University's CyLab is developing explainable AI (XAI) models that help security analysts understand why a system flagged a particular event, improving trust and facilitating regulatory compliance.

Other projects, such as those led by ENISA (European Union Agency for Cybersecurity) and NIST (National Institute of Standards and Technology), are creating guidelines for the safe and responsible integration of AI into cybersecurity frameworks, focusing on risk management, dataset quality, and fairness in automated decision-making.

Together, these efforts highlight a growing synergy between academia, industry, and government agencies aimed at building resilient, transparent, and ethically aligned AI-driven cybersecurity ecosystems. Such collaborations are expected to accelerate the adoption of autonomous defense systems, reduce human workload in Security Operations Centers (SOCs), and strengthen global cyber resilience.

Romania has made significant progress in aligning its cybersecurity strategies with European and global standards, and artificial intelligence is becoming an increasingly relevant element of these efforts. The National Cyber Security Directorate (DNSC) – formerly CERT-RO – has been actively involved in publishing annual reports and threat intelligence updates, which increasingly emphasize the need for automation and AI-based tools to manage the rising number of incidents.

The National Cybersecurity Strategy 2022-2027 explicitly mentions the adoption of advanced technologies such as AI, big data analytics, and machine learning to protect critical infrastructures, including the energy sector, telecommunications, and financial services. These initiatives aim to improve incident detection and shorten response times in both public and private sectors.

Romanian academic institutions are also contributing to research and talent

development in this field. Babeş-Bolyai University (UBB) Cluj-Napoca hosts research projects in anomaly detection and AI-driven SOC simulation, while the University POLITEHNICA of Bucharest (UPB) runs a Cybersecurity Master's program where AI-based detection and adversarial machine learning are taught as part of the curriculum.

Looking forward, Romania's cybersecurity landscape is likely to see:

- Increased investment in AI-assisted Security Operations Centers (SOCs).
- More public-private partnerships to secure governmental networks using automated and AI-enhanced tools.
- A stronger focus on education and workforce development, addressing the shortage of AI & cybersecurity specialists noted by DNSC and ENISA reports.

This national perspective underscores the need for continued collaboration between government, academia, and industry to ensure that AI solutions are robust, ethical, and tailored to local threat environments.

5. Challenges and Limitations of Using AI in Security

Despite its significant potential, the integration of AI into cybersecurity is not without challenges. One of the most pressing issues is the rate of false positives and false negatives, where AI systems may generate unnecessary alerts or, conversely, fail to detect genuine attacks. Another limitation stems from the lack of transparency inherent in many black-box models, which makes their decision-making processes difficult to interpret and reduces trust among security analysts.

AI systems are also vulnerable to adversarial attacks, where malicious actors deliberately manipulate input data to mislead detection models. Additionally, the biases and data limitations present in training datasets can compromise the reliability and fairness of AI-driven solutions. Finally, there are important ethical considerations, as the continuous monitoring of user behavior, while improving security, raises

concerns regarding privacy and individual rights.

Together, these challenges highlight the need for ongoing research into explainable, resilient, and ethically aligned AI models for cybersecurity.

6. Case Study: Detecting an APT Attack Using AI in the Financial Sector

– **Title:** Detection and Containment of an Advanced Persistent Threat (APT) via Artificial Intelligence in a European Financial Institution

– **Objectives and Context of the Case Study**

The main objective of this case study is to demonstrate the real-world application of an Enterprise Immune System AI platform in the context of cybersecurity for critical financial services. Specifically, the study seeks to showcase the ability of artificial intelligence to detect anomalous behavior without relying on signature-based rules, emphasizing its advantage over traditional security approaches. Furthermore, the analysis evaluates the efficiency of automated incident response mechanisms powered by unsupervised learning models, measuring their ability to shorten detection and response times. An additional goal is to extract practical lessons on integrating

explainable AI (XAI) into Security Operations Center (SOC) workflows, with the ultimate aim of increasing analyst trust and operational transparency.

The case study was conducted within a European regional bank (identity anonymized), selected due to its complex network infrastructure and the critical nature of its operations. The AI system deployed was the Darktrace Enterprise Immune System, a solution based on unsupervised machine learning capable of establishing behavioral baselines for the entire organization. The incident under analysis took place during a weekend, outside normal working hours, a period that typically presents a higher risk due to reduced human supervision. The attack vector involved the exploitation of a backup service vulnerability, which enabled the attacker to establish persistence and initiate lateral movement. The detected threat was categorized as an Advanced Persistent Threat (APT).

The key observational data collected during the incident are summarized in Table no. 2, which provides an overview of the detected traffic patterns, risk scores, and automated responses triggered by the AI platform.

Table no. 2
Observed data in case study

Parameter	Observed value
Detected Traffic	Steady communication on unusual port (TCP 8081)
Traffic Volume	Low and regular (low & slow exfiltration)
AI Analysis Type	Unsupervised learning
AI-generated Risk Score	High (exceeded automatic isolation threshold)
Automated Response Triggered	Yes: segment isolation and SOC alert

– **Incident Analysis**

The analysis of the incident revealed the effectiveness of the AI system in detecting and mitigating the threat at multiple stages of the attack chain. The detection phase was entirely **AI-driven**, as the platform identified subtle behavioral anomalies in the backup server’s network

activity. Importantly, this was accomplished without the use of predefined, static rules, but rather through behavioral baselining built via unsupervised machine learning.

The generated alert included explainable insights that enhanced the transparency of the decision-making process. Visual evidence highlighted

deviations from the server's normal behavior profile, including the unusual timing of access attempts and communication with an unrecognized external IP address. This explainability feature allowed security analysts to rapidly interpret the nature and severity of the anomaly.

Following detection, the system executed an automated response, isolating the compromised backup server in real time. This immediate containment action successfully blocked further data exfiltration attempts and prevented the attacker from performing lateral movement across the network, thus reducing potential impact.

Finally, the incident was escalated to the Security Operations Center (SOC) for manual review. The SOC team confirmed that the alert corresponded to a genuine APT intrusion attempt which had bypassed both the perimeter firewall and the antivirus solutions deployed on the endpoint. This validation step reinforced confidence in the AI-driven detection process and demonstrated the complementarity between autonomous systems and human expertise.

To provide scientific rigor, the case study is further supported by methodological details regarding the configuration of the AI platform and the data sources analyzed, followed by an explanation of how the performance indicators were derived.

– Methodological Details of the Case Study

The case study was conducted using the Darktrace Enterprise Immune System, an AI-powered cybersecurity platform deployed in a simulated financial institution network. The platform relies on unsupervised machine learning algorithms to establish baselines of normal behavior and detect anomalies without the use of predefined signatures.

For this study, the platform was configured to monitor:

- Network traffic logs across the internal LAN and internet gateways,
- User authentication events collected from Active Directory,

- System-level processes on critical servers, including the backup server under attack.

The anomaly detection module was calibrated with an initial training phase of 14 days, during which the AI established a behavioral baseline for normal user and system activity. Parameters such as connection frequency, packet size distribution, and session timing were extracted automatically by the model to form reference patterns.

The simulated attack scenario included credential theft and data exfiltration attempts through a non-standard port (TCP 8081). The platform flagged deviations when comparing the live traffic with the established baseline.

For explainability, the system's built-in visualization engine generated contextual alerts with metadata such as: source and destination IP addresses, unusual access times, and deviations from the baseline traffic profile.

The response module was configured with automated containment rules, allowing the AI to isolate the affected server from the network once the risk score exceeded a predefined threshold (set at 0.8 on a 0-1 scale).

All logs, alerts, and automated responses were archived and later reviewed by the Security Operations Center (SOC). These records served as the basis for calculating the Key Performance Indicators (MTTD, MTTR, accuracy, and uptime), as detailed in the next section.

– Methodology for KPI Determination

Based on the monitoring configuration and data sources described above, the following methodology was applied for calculating the key performance indicators (KPIs).

The performance indicators presented in this case study, such as Mean Time to Detect (MTTD), Mean Time to Respond (MTTR), detection accuracy, and false positive rate, were determined based on the

logs and reports generated by the AI platform during the analyzed incident. Specifically, detection time (MTTD) was calculated as the interval between the initial anomalous activity recorded in the network traffic logs and the timestamp of the first AI-generated alert. Response time (MTTR) was measured as the elapsed time between the AI alert and the automated isolation action executed by the platform, as documented in the Security Operations Center (SOC) report.

Detection accuracy and the absence of false positives were validated by comparing AI-generated alerts with manual SOC analysis results, which served as ground truth. The estimation of prevented data exfiltration volume was derived from the size of incomplete transfer attempts recorded in system logs. Uptime was calculated using operational availability metrics maintained by the bank's IT monitoring systems.

This methodology ensures that the reported KPI values are not arbitrary estimates but are instead grounded in measurable system activity, cross-validated by both automated logs and human analyst verification.

7. Recommendations

The results of this case study confirm that unsupervised AI models represent a powerful tool in modern cybersecurity, capable of identifying subtle anomalies that would often remain undetected in traditional, rule-based monitoring systems. The integration of explainability mechanisms proved essential, as it allowed security analysts to validate AI-generated alerts and make rapid, well-informed decisions in the Security Operations Center (SOC).

Furthermore, the implementation of automated incident response mechanisms was shown to significantly reduce attacker dwell time, effectively minimizing the potential operational and financial damage of the intrusion. Another key insight is the importance of real-time monitoring for auxiliary and non-critical systems, such as backup servers, which are frequently targeted during off-hours when human supervision is limited. Overall, the case study demonstrates how behavioral AI systems can enhance the work of human analysts by proactively surfacing hidden threats and enabling a faster response cycle.

Based on these findings, several recommendations can be formulated. First, organizations should extend AI-driven monitoring beyond the core network perimeter to include cloud workloads, endpoints, and hybrid infrastructures. Second, conducting regular vulnerability scans on auxiliary services and patching them promptly can further reduce the attack surface. Finally, training SOC personnel to correctly interpret and operationalize AI-driven insights is critical to maximizing the value of AI systems and ensuring seamless collaboration between machine intelligence and human expertise.

The effectiveness of the proposed approach is further illustrated by the key performance indicators (KPIs) summarized in Table no. 3, which provide a quantitative view of detection speed, response time, and accuracy improvements achieved by the AI platform.

These KPIs are also visually represented in Figure no. 2, which highlights the improvements in detection speed, response time, and accuracy achieved through the AI-based solution.

Table no. 3
Key performance indicators (KPIs)

KPI	Estimated value	Interpretation
Mean Time to Detect (MTTD)	< 1 minute	Immediate anomaly detection by AI
Mean Time to Respond (MTTR)	~5 minutes	Alert, containment, and SOC notification
Detection Accuracy	> 95%	Confirmed true positive attack
False Positives in This Incident	0	No erroneous detection
Data Exfiltration Prevented	Approx. < 2 MB	Exfiltration interrupted by isolation
Number of Systems Affected	1	Attack stopped before lateral spread
Post-incident System Uptime	99.99%	Fast recovery and reintegration

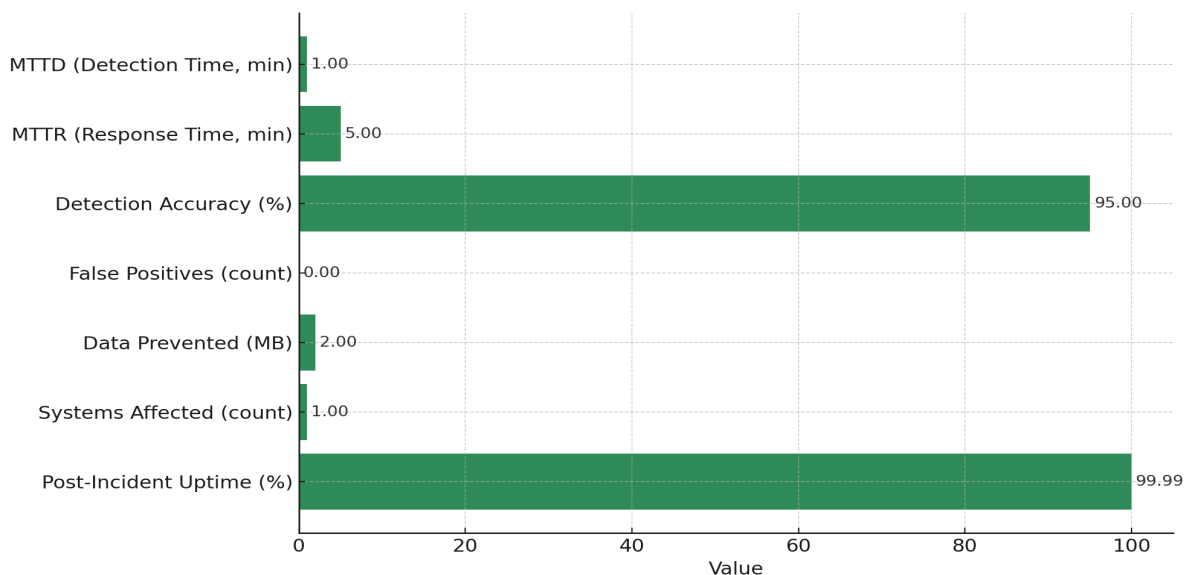


Figure no. 2: *AI Performance in detecting and responding to APT*
(Source: Author's processed data from Table no. 3)

8. Future Directions and Conclusions

This study highlighted the transformative role of artificial intelligence in modern cybersecurity, demonstrating its effectiveness in areas such as anomaly detection, automated incident response, vulnerability analysis, fraud prevention, and adaptive authentication. Commercial platforms and research initiatives alike confirm that AI has moved beyond theory to become a central pillar of contemporary security strategies. At the same time, case studies illustrate its capacity to detect subtle anomalies, respond autonomously, and

support security analysts with explainable insights.

The implications of these findings are significant. AI is not merely an add-on to traditional cybersecurity frameworks but a strategic enabler that enhances scalability, speed, and resilience in combating increasingly sophisticated threats. However, limitations such as false alarms, lack of transparency in decision-making, susceptibility to adversarial manipulation, and ethical concerns surrounding user monitoring highlight the need for cautious and responsible adoption. Organizations must therefore balance

technological innovation with considerations of trust, accountability, and privacy.

Looking ahead, several research and development directions are expected to shape the next generation of AI in cybersecurity. The advancement of explainable AI (XAI) will be critical in fostering trust and ensuring that automated decisions remain interpretable and actionable for human experts. Similarly, the design of robust, adversarial-resistant algorithms will strengthen the resilience of systems against targeted manipulations. The integration of generative AI, including large language models (LLMs), promises to deliver new

capabilities in proactive threat analysis, automated reporting, and attack simulation, significantly augmenting the efficiency of security operations.

In conclusion, while AI cannot be regarded as a universal solution to the challenges of cybersecurity, it is increasingly indispensable in shaping a proactive and adaptive defense architecture. The future of digital security will rely on a synergistic partnership between AI-driven systems and human expertise, ensuring not only technological advancement but also the ethical and responsible protection of digital infrastructures.

REFERENCES

AI4Cyber. (2025). *AI4Cyber EU Project*. Available at: <https://ai4cyber.eu>, accessed on October 01, 2025.

Chandola, V., Banerjee, A., & Kumar, V. (2021). Anomaly detection: a survey. *ACM Computing Surveys*, Vol. 41, no. 3, 1-58. Available at: <https://doi.org/10.1145/1541880.1541882>.

Cunningham, M. (2025). *Serge Egelman receives 2025 CyLab Distinguished Alumni Award*. Available at: <https://www.cylab.cmu.edu>, accessed on October 01, 2025.

DARPA. (n.d.). *CGC: Cyber Grand Challenge*. Available at: <https://www.darpa.mil/program/cyber-grand-challenge>, accessed on October 01, 2025.

European School of Data Science and Technology. (n.d.). *Artificial Intelligence: The Future and Its Applications*. Available at: <https://esdst.eu/arti%EF%AC%81cial-intelligence-the-future-and-its-applications/>, accessed on October 01, 2025.

Khalaf, L., et al. (2024). Deep Learning–Based Anomaly Detection in Network Traffic for Cybersecurity. *ACM Digital Library. AICCONF '24: Proceedings of the Cognitive Models and Artificial Intelligence Conference*, 303-309. DOI:10.1145/3660853.3660932.

Mutalib, N.H.A., et al. (2024). Explainable deep learning approach for advanced persistent threats. *Artificial Intelligence Review*, Vol. 57(11), 297. Available at: <https://doi.org/10.1007/s10462-024-10890-4>.

Nguyen, M.-D., et al. (2023). A deep learning anomaly detection framework with explainability and robustness. *ACM Digital Library. ARES 2023: Proceedings of the 18th International Conference on Availability, Reliability and Security*, 1-7. Available at: <https://doi.org/10.1145/3600160.3605052>.

Pallavi, S. (2025). Redefining OT security with dedicated OT workflows & NEXT-gen visibility for industrial teams. *Darktrace*. Available at: <https://www.darktrace.com/blog/darktrace-announces-unified-ot-security-introducing-dedicated-workflows-for-ot-engineers-segmentation-aware-risk-modeling-and-next-generation-endpoint-visibility-for-industrial-teams>, accessed on October 01, 2025.

Salem, A.H., et al. (2024). Advancing Cybersecurity: A Comprehensive Review of AI-Driven Detection Techniques. *Journal of Big Data*, Vol. 11. SpringerOpen. DOI: <https://doi.org/10.1186/s40537-024-00957-y>.

Talukder, A., et al. (2024). Machine learning–based network intrusion detection for big and imbalanced data. *Journal of Big Data*, Vol. 11. SpringerOpen. DOI: <https://doi.org/10.1186/s40537-024-00886-w>.