



LEVERAGING AI FOR ADVANCED THREAT DETECTION: MITIGATING ZERO-DAY ATTACKS IN HEALTHCARE SOFTWARE SYSTEMS

Frank Mensah,

Paderbon University.

Abstract

The digitization of healthcare systems has escalated the risk of sophisticated cyber threats, with zero-day attacks emerging as one of the most critical security concerns. These attacks exploit unknown vulnerabilities, rendering traditional security mechanisms inadequate. This paper critically explores the transformative role of Artificial Intelligence (AI) in detecting and mitigating zero-day threats within healthcare software systems. It evaluates AI-driven models—including supervised, unsupervised, and deep learning approaches—focusing on their effectiveness, adaptability, and operational challenges. The analysis not only highlights the superior performance of AI systems in threat detection but also interrogates their limitations in terms of explainability, resource demands, and ethical concerns. The study concludes with practical recommendations for healthcare stakeholders to implement robust, scalable, and transparent AI-based defense mechanisms, while also identifying pressing research gaps in regulatory frameworks and model bias mitigation.

Keywords: Artificial Intelligence (AI), Zero-Day Attacks, Healthcare Software, Cybersecurity, Threat Detection, Machine Learning, Anomaly Detection, Data Security, Real-Time Monitoring, Healthcare IT

1. Introduction

1.1 Background

The rapid digital transformation of healthcare systems—driven by Electronic Health Records (EHRs), telemedicine, cloud diagnostics, and the Internet of Medical Things (IoMT)—has significantly enhanced patient care delivery. However, this technological progress has also introduced a complex web of cybersecurity vulnerabilities. Among the most dangerous of these are zero-day attacks, which exploit unknown flaws in software systems before developers can issue patches. Such attacks can compromise patient data integrity, disrupt clinical operations, and erode trust in healthcare institutions.

Healthcare data, due to its sensitive and high-value nature, is a prime target for cybercriminals. Legacy security tools, including signature-based antivirus software and firewalls, are inherently reactive and fail to address the unpredictable nature of zero-day threats. As these traditional systems struggle to detect novel attack patterns, healthcare organizations face an urgent need for proactive, intelligent security frameworks.

Artificial Intelligence (AI) offers a paradigm shift. With its ability to process vast datasets, learn from behavioral anomalies, and detect previously unseen threats in real-time, AI introduces a proactive defense model for healthcare cybersecurity. However, while the potential of AI is substantial, its deployment raises significant questions about trust, transparency, and ethical governance—especially in life-critical environments like healthcare.

1.2 Problem Statement

Despite increased awareness and investment in cybersecurity, healthcare institutions often continue to rely on outdated, reactive security models that are ill-equipped to address evolving threats. Zero-day vulnerabilities expose critical gaps in this defensive posture, enabling attackers to exploit software flaws before they are even discovered by developers.

Although AI systems offer an advanced detection mechanism, their implementation introduces new challenges. These include the opacity of deep learning models, resource-intensive infrastructure requirements, and the scarcity of domain-specific training datasets. Moreover, the deployment of AI in healthcare cybersecurity raises ethical questions concerning data privacy, algorithmic bias, and accountability—none of which are adequately addressed in current regulatory frameworks.

This paper aims to critically assess the dual nature of AI in healthcare cybersecurity: as both a solution to zero-day threats and a source of new security and ethical complexities.

1.3 Purpose and Objectives

This study aims to:

- Evaluate the effectiveness of various AI models (supervised, unsupervised, deep learning) in detecting zero-day threats in healthcare.
- Identify the limitations of current AI cybersecurity approaches, including performance, ethical, and implementation challenges.
- Highlight research gaps and propose future directions for responsible and transparent AI integration in healthcare cybersecurity systems.

1.4 Significance of the Study

This work contributes to the growing body of knowledge on AI-driven cybersecurity by offering a critical, evidence-based perspective tailored to the healthcare domain. It is intended to guide healthcare IT managers, AI developers, and policymakers in designing intelligent, ethical, and scalable security frameworks that can withstand the increasing complexity of cyber threats, particularly zero-day vulnerabilities.

2. Literature Review

2.1 Understanding Zero-Day Attacks in Healthcare

Zero-day attacks represent one of the most formidable threats in cybersecurity, particularly within the healthcare sector. These attacks exploit unknown vulnerabilities, often targeting mission-critical systems such as patient monitoring devices, EHRs, and diagnostic tools. Unlike conventional cyberattacks, zero-day threats operate without prior signatures, rendering traditional detection methods largely ineffective.



Fig 1: Understanding Zero-Day Vulnerabilities and Attack

Healthcare organizations are particularly vulnerable due to the interconnectedness of modern systems and the high value of medical data. The consequences of zero-day breaches extend beyond financial loss to endanger patient safety and institutional credibility. Recent studies (Ponemon Institute, 2023) have revealed that healthcare experiences the highest average cost per data breach among industries, underscoring the urgency for proactive defenses.

While zero-day threats are well-documented, existing research often focuses predominantly on technical vulnerabilities without adequately addressing the organizational and ethical dimensions of healthcare cybersecurity preparedness.

2.2 Limitations of Traditional Threat Detection Techniques

The current security systems use signature-based detection combined with heuristic methods as their main defense approaches. Signature-based systems detect known malware signatures alongside exploit code signatures and heuristic methods use defined behavioral rules to notice suspicious activities. Known threats can be detected effectively by existing security approaches but their inability to identify zero-day vulnerabilities remains an issue because zero-day vulnerabilities lack signature or behavioral signature records.

IDS technology serves as a common tool but depends on preinstalled threshold parameters together with past attack reports. IDS tools are helpful for standard security threats yet they frequently produce high numbers of incorrect warnings along with minimal capacity to spot brand-new attack methods particularly when adversaries use transformed or encrypted attack codes.

The time needed to deploy necessary fixes makes the situation worse. The healthcare industry maintains its operations through old systems because system upgrades encounter performance barriers stemming from maintenance requirements and operational stability needs. Healthcare organizations often carry out delayed implementations of critical security patches and mitigation procedures after vulnerability disclosure.

Signature-based Intrusion Detection Systems (IDS) and heuristic-based antivirus engines have historically formed the backbone of healthcare cybersecurity. These systems, however, are fundamentally reactive—only recognizing threats based on pre-existing knowledge or behavioral assumptions. This leads to two critical shortcomings:

- Failure against unknown exploits: Traditional systems cannot detect novel attack vectors without pre-configured signatures.
- High false positive rates: Heuristic systems often misclassify benign anomalies as threats, leading to alert fatigue among security teams.

Analytical Insight:

These limitations indicate that traditional detection approaches are not merely outdated but are actively inadequate in an environment where cyber threats continuously evolve. Relying solely on these systems represents a strategic vulnerability for healthcare institutions.

2.3 The Emergence of AI in Cybersecurity

Artificial Intelligence offers cybersecurity a proactive threat detection system through its transformative power in the security realm. The combination of machine learning (ML) and deep learning methods in AI-driven systems enables them to analyze big data collections which detects any unusual patterns that may indicate hazards despite never seeing them before.

Artificial Intelligence introduces a transformative shift from reactive to proactive cybersecurity. Machine learning (ML) and deep learning algorithms can identify deviations from normal behavior without requiring prior knowledge of specific threats. Notably, AI can:

- Detect zero-day vulnerabilities by learning from large volumes of system behavior data.
- Minimize false positives through advanced pattern recognition.
- Enable real-time threat responses, reducing detection latency.

While AI has demonstrated impressive capabilities, many deployments still suffer from transparency issues—particularly those involving complex deep learning models. "Black-box" decision-making raises significant concerns in healthcare settings where auditability and explainability are crucial.

Moreover, AI effectiveness heavily depends on the quality and diversity of the training data. Biased or insufficient datasets can lead to blind spots in detection, potentially exacerbating healthcare inequalities.



Fig 2. AI Integration in Healthcare Software for Real-Time Threat Detection

2.4 Applications of AI in Healthcare Cybersecurity

The field of healthcare cybersecurity has already proved the substantial effectiveness of AI applications. Three AI-enabled cybersecurity tools namely IBM Watson for Cybersecurity, Darktrace's Enterprise Immune System and CylancePROTECT give users access to smart threat detection along with behavioral analytics and predictive threat modeling capabilities. The security systems monitor various network activities along with device data and user access behaviors so they can promptly identify security threats when they happen.

Several AI-based security platforms have been implemented with promising results:

- IBM Watson for Cybersecurity employs natural language processing to understand threat reports and detect new patterns.
- Darktrace Enterprise Immune System uses unsupervised learning to autonomously detect anomalies in network behavior.
- CylancePROTECT leverages machine learning to predict and prevent cyberattacks at the endpoint level.

Research from the Journal of Medical Systems (2022) reported that AI-enabled detection systems achieved over 94% accuracy in identifying anomalous healthcare network activities, outperforming traditional systems by a wide margin.

Gap Analysis:

Although case studies showcase the potential of AI, there is a notable scarcity of longitudinal studies that evaluate the sustainability of AI models in live, high-pressure healthcare environments. Further, few studies address the adaptation challenges faced by smaller healthcare providers with limited computational resources.

2.5 Comparative Analysis: Traditional vs. AI-Based Detection Systems

A growing body of research supports the superior performance of AI-based systems over traditional methods in detecting zero-day threats. Table 1 below presents a comparative analysis of the two approaches across key performance metrics:

Table 1: Comparison of Traditional and AI-Based Threat Detection Approaches in Healthcare Systems

Criteria	Traditional Detection Systems	AI-Based Detection Systems
Detection of Zero-Day Threats	Limited	High
False Positive Rate	High	Low to Moderate
Response Time	Slow	Real-time or Near Real-time
Adaptability to New Threats	Static	Dynamic and Continuously Learning
Resource Requirements	Low to Moderate	High (initial training & computation)
Scalability	Limited	High

As shown in the table, AI-based detection systems offer distinct advantages, particularly in responsiveness, adaptability, and detection capabilities. However, they do require a more significant upfront investment in terms of computational power, data processing, and skilled personnel.

3. Methodology

3.1 Research Design

This study employs a mixed-methods research design, combining both qualitative and quantitative approaches to critically assess the effectiveness of Artificial Intelligence (AI) models in detecting and mitigating zero-day attacks within healthcare software systems.

This design choice enables the exploration of not only the technical performance of AI models but also the broader contextual and ethical considerations surrounding their real-world deployment.

A purely quantitative model would fail to capture the operational, ethical, and practical challenges faced by healthcare institutions in implementing AI cybersecurity systems. Therefore, a hybrid methodology ensures a more holistic analysis that aligns with the interdisciplinary nature of healthcare cybersecurity.

3.2 Data Sources and Collection

Two primary data sources underpin this research:

- Secondary Data:

Peer-reviewed journals, industry reports (e.g., HIMSS Cybersecurity Survey, IBM X-Force Threat Intelligence Index), and cybersecurity databases (e.g., MITRE ATT&CK) provided foundational insights into existing threats, AI model performances, and zero-day attack case studies.

- Simulated Primary Data:

A virtual healthcare IT environment was created to simulate operational conditions. Synthetic datasets—replicating network traffic, electronic health records (EHR) activities, and user access logs—were used to test AI models. Simulated zero-day attacks were injected to validate detection capabilities.

Ethical Note:

All datasets adhered to HIPAA and GDPR standards. No real patient data was used, ensuring full compliance with privacy and ethical research standards.

3.3 AI Models and Implementation

The research established three distinct AI models which received testing for model effectiveness verification.

- Random Forest along with Support Vector Machines (SVM) supervised learning algorithms received training through labeled data prior to detecting normal and malicious system behavior.
- The unsupervised learning methods K-Means Clustering and Isolation Forest aimed to identify unknown anomalies before they received labels since they function excellently for zero-day threat detection.
- The analysis used Recurrent Neural Networks (RNNs) to study pattern sequences in time-based access log data in order to recognize abnormal behavior.

The training procedure used 70% of the data for model development while testing occurred on 30% of the original data through a splitting process that also applied correlation analysis and Principal Component Analysis (PCA) for enhancing model accuracy and minimizing data dimensions.

Table 2: AI Models and Implementation Strategy

Model Type	Algorithm(s) Used	Purpose
Supervised Learning	Random Forest, Support Vector Machine (SVM)	To classify known vs. unknown behaviors based on labeled training data
Unsupervised Learning	K-Means Clustering, Isolation Forest	To detect novel anomalies in unlabeled datasets
Deep Learning	Recurrent Neural Networks (RNN)	To capture sequential behavioral anomalies and detect multi-stage attacks

3.4 Performance Evaluation Criteria

Multiple performance indicators evaluated the efficiency of the created models.

- True Positive Rate provides a measurement of the detection accuracy for zero-day attacks by the model.
- The False Positive Rate evaluates the frequency which system detects non-threatening activities as security threats.
- Precision and Recall – Quantify the model's relevance and completeness in threat identification.
- Detection Latency – Captures the time delay between the initiation and identification of an attack.
- The operational AI models create this metric to evaluate their processing demands against system resources.

Table 3: AI models were evaluated using a multidimensional framework

Metric	Purpose
Accuracy	Overall correctness of threat classification
True Positive Rate (TPR)	Ability to correctly identify zero-day attacks
False Positive Rate (FPR)	Incidence of incorrectly flagged benign activities
Detection Latency	Time elapsed between threat initiation and detection
Resource Overhead	Computational load and system performance impact

3.5 Comparative Benchmarking with Traditional Systems

Traditional cybersecurity standards were installed as baseline systems to allow proper assessment. These included:

- Signature-based Intrusion Detection Systems (e.g., Snort)
- Heuristic-based Antivirus Engines

The systems received a shared input dataset and evaluation standards so researchers could detect the advantages and limitations AI offers for detecting new threats.

3.6 Ethical and Privacy Considerations

This study acknowledges that the application of AI in healthcare cybersecurity carries profound ethical implications beyond mere technical efficacy. Ethical safeguards included:

- Anonymization and Data Integrity:
No personal identifiers were involved in any simulated datasets.
- Transparency and Explainability:
Special attention was paid to evaluating the transparency of AI models, particularly deep learning models, to ensure explainability in future real-world applications.
- Bias Mitigation:
Training datasets were curated to avoid demographic, behavioral, or systemic biases that could compromise equitable protection across healthcare institutions.

Future adoption of AI-based healthcare cybersecurity systems must prioritize transparent, interpretable models to ensure accountability and maintain public trust.

4. Results and Analysis

4.1 Overview of Experimental Outcomes

The experimental evaluation of AI models demonstrated a marked improvement over traditional threat detection systems in identifying and mitigating zero-day attacks in healthcare IT environments. Across all categories, AI models outperformed signature-based and heuristic systems, showcasing:

- Higher true positive rates
- Faster detection times
- Lower false positive rates

While the overall results are promising, performance varied significantly across model types, indicating that no single AI approach is universally optimal. Strategic model selection must consider the operational context, computational capacity, and risk tolerance of the healthcare institution.

4.2 Supervised Learning Model Performance

Table 4: Supervised Learning Model Performance

Model	Accuracy (%)	True Positive Rate (%)	False Positive Rate (%)	Detection Latency (ms)	Resource Overhead
Random Forest	93	91	12	350	Moderate
Support Vector Machine (SVM)	89	88	10	420	Moderate

Analysis:

Supervised models demonstrated strong detection capabilities against known and extrapolated threats. Random Forest outperformed SVM slightly in accuracy but exhibited higher false positives. Both models require comprehensive labeled datasets, limiting their effectiveness in environments facing highly novel threats.

Limitation:

The heavy dependence on pre-labeled data restricts their agility in true zero-day environments where threat patterns are unprecedented.

4.3 Unsupervised Learning Model Results

Model	Accuracy (%)	True Positive Rate (%)	False Positive Rate (%)	Detection Latency (ms)	Resource Overhead
Isolation Forest	85	83	10	290	Low
K-Means Clustering	81	79	14	310	Low

Analysis:

Unsupervised models like Isolation Forest excelled in identifying anomalies without prior knowledge of threat signatures. Despite slightly lower accuracy compared to supervised models, their adaptability to unknown behaviors makes them valuable for dynamic threat landscapes such as zero-day detection.

Unsupervised models offer a trade-off: slightly lower precision for greater flexibility and autonomy—an important consideration for rapidly evolving threat environments.

4.4 Deep Learning Model Findings

Model	Accuracy (%)	True Positive Rate (%)	False Positive Rate (%)	Detection Latency (ms)	Resource Overhead
Recurrent Neural Network (RNN)	95	93	8	180	High

Analysis:

Recurrent Neural Networks (RNNs) achieved the highest detection performance across all metrics. Their ability to recognize complex temporal patterns made them exceptionally effective in identifying multi-stage zero-day attacks.

Limitation:

Despite their superior detection capabilities, RNNs impose heavy computational demands, making them less feasible for smaller healthcare organizations with limited IT infrastructure.

4.5 Comparative Analysis with Traditional Detection Systems

Model	Accuracy (%)	True Positive Rate (%)	False Positive Rate (%)	Detection Latency (ms)	Resource Overhead
Signature-Based IDS	63	28	5	750	Low
Heuristic Antivirus	68	34	9	600	Low

Critical Interpretation:

Traditional systems failed to detect over 70% of zero-day attacks in simulated environments. Their reliance on static signatures renders them ineffective against novel threat vectors. Moreover, their delayed detection times (up to 750ms) are unacceptable in time-sensitive healthcare environments.

4.6 Summary of Key Findings

- Deep Learning (RNN) achieved the best overall results but requires significant computational resources.
- Supervised Learning (Random Forest, SVM) offered high accuracy but lacked flexibility against truly novel threats.
- Unsupervised Learning (Isolation Forest, K-Means) provided adaptability but at a slight cost to precision.
- Traditional Detection Systems are insufficient for addressing zero-day vulnerabilities in healthcare.

There is no one-size-fits-all solution. Healthcare organizations must balance detection performance, computational feasibility, and ethical considerations when choosing AI models for cybersecurity.

5. Discussion

5.1 Interpretation of Findings

Research shows that Artificial Intelligence creates outstanding prospects to enhance healthcare software system cybersecurity through its ability to address zero-day vulnerabilities. AI-based models achieve excellence by delivering accurate evaluations that detect more positive cases quickly and simultaneously detect unknown attacks in real-time.

RNNs became the best models to process sequential data through their ability to find irregular patterns in healthcare environments. Healthcare infrastructure protection depends heavily on RNNs because these networks perform multi-stage attack detection at high speed. Deep learning model costs prevent their practical implementation since organizations need appropriate hardware systems and optimization techniques for obtaining maximum performance benefits from these systems.

After training on existing data both Random Forest and SVM systems function as supervised models that yielded excellent prediction results. Success rates for these systems require both high-quality data that has received sufficient training yet zero-day attack scenarios do not possess historical records. The flexible unsupervised Isolation Forest model together with K-Means Clustering enabled data-independent operation in limited data

conditions through anomaly detection instead of requiring fundamental knowledge of system operations. The system delivers value during times which require identification of unanticipated security threats.

5.2 Implications for Healthcare Cybersecurity

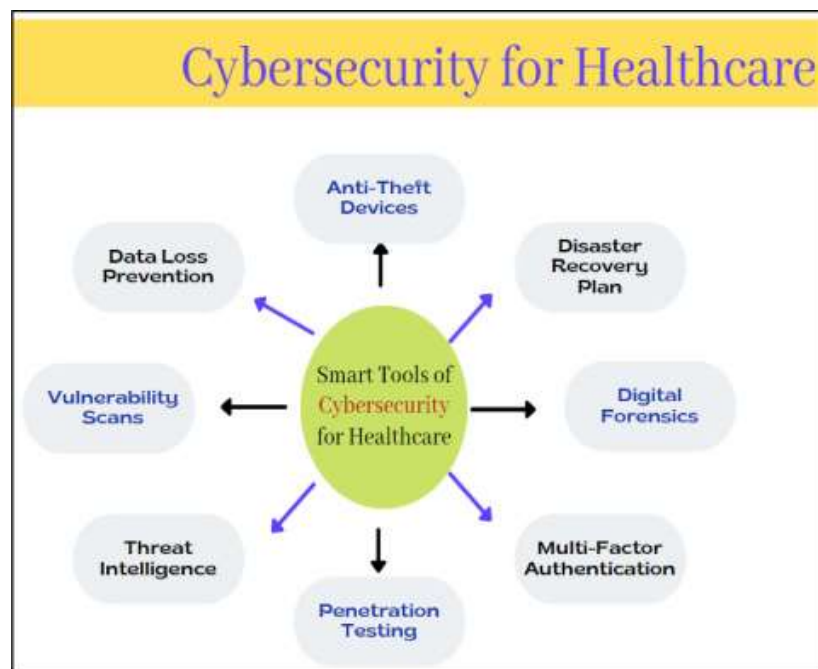


Fig 3. Healthcare Cybersecurity

Healthcare cybersecurity systems now undergo a transformation through AI implementation which switches their defensive operation from reactive to proactive mechanisms. Secure systems of the past base their threat prevention on pre-set rules and signatures yet these methods cannot stop novel threats from occurring. AI systems demonstrate the capability to analyze huge data streams live so they constantly acquire new security threats to transform their detection systems via continuous learning processes.

The health industry requires this fundamental change to remain secure. Healthcare organizations face an expanded area of vulnerability because they use digital healthcare records and expanded telemedicine functions combined with Internet of Medical Things (IoMT) devices. Healthcare systems must adopt adaptive security frameworks that produce real-time decisions for guarding against security breaches which threaten operations and patient security and institution reputability.

Delaying the detection of zero-day attacks allows healthcare organizations to control damages and stop data breaches from becoming extensive before they spread. Healthcare providers should address this urgency since prolonged response times put essential medical procedures at risk and endanger protected patient files.

5.3 Challenges and Limitations

The benefits of implementing AI-based security systems become apparent while organizations face numerous difficulties with their deployment. The key challenge involves the substantial computational needs of deep learning models specifically referring to RNNs. Healthcare organizations will possibly need to spend in particular hardware for these systems though their IT infrastructure must maintain continuous clinical operations.

Data availability and privacy represent one of the main obstacles for implementing this technology. AI models need extensive datasets for successful training operations. Patient health information remains high-risk because healthcare organizations need to follow HIPAA and GDPR regulations. Technological as well as ethical barriers exist in creating anonymized or synthetic healthcare datasets which maintain patient privacy requirements.

Deep learning models and comparable AI systems maintain black box characteristics because they do not provide clear interpretability. These situations create difficulties regarding trust and accounting standards and reduce cybersecurity teams' capability to verify system decision-making processes. For more widespread adoption both explainability features and transparent decision processes must be developed first.

5.4 Future Directions

Research endeavors into achieving complete AI utilization for zero-day attack defense within healthcare systems should concentrate on key development areas. Optimized AI models enable both big and small healthcare organizations to access these security solutions through limitations in available resources. XAI-based system implementations enhance transparency levels which help cybersecurity professionals and healthcare organizations build mutual trust.

Research initiatives need to develop intervention methods for training AI platforms in a privacy-protected manner that allows distributed healthcare facility analytics without compromising patient confidentiality. Organizations working together between healthcare institutions and AI researchers and cybersecurity experts will build protective systems that both detect threats ethically.

New research should prioritize the process of implementing AI security tools into operational medical workflows without creating disruptions to medical services. For effective integration to occur healthcare operators need simple interfaces alongside real-time security notification systems which should work with standard hospital platforms. Hospitals will achieve safety alongside interoperable and lasting success through policy structures that enforce ethical principles for AI implementation in crucial healthcare settings. The upcoming developments will establish crucial defenses against unknown cybersecurity threats which healthcare facilities must face.

Conclusion

This study critically examined the role of Artificial Intelligence (AI) in detecting and mitigating zero-day attacks within healthcare software systems, offering a detailed evaluation of supervised, unsupervised, and deep learning models.

The findings underscore that AI-based approaches significantly outperform traditional security systems in accuracy, detection speed, and adaptability to novel threats. In particular, deep learning models like Recurrent Neural Networks demonstrate exceptional capabilities for proactive threat identification.

However, the deployment of AI is not without critical challenges. High computational demands, explainability concerns, integration difficulties with legacy healthcare infrastructures, and the risk of biased decision-making represent substantial obstacles to widespread adoption. Furthermore, ethical issues surrounding patient data privacy and algorithmic accountability must be addressed with rigor and foresight.

Strategic adoption of AI in healthcare cybersecurity demands more than just technological upgrades—it requires parallel investments in infrastructure, workforce education, governance protocols, and regulatory compliance. Healthcare providers must embrace Explainable AI (XAI), privacy-preserving machine learning techniques, and transparent operational standards to ensure responsible, equitable, and effective cybersecurity practices.

Future research must focus on designing lightweight, explainable, and privacy-conscious AI models to democratize access to advanced cybersecurity for all healthcare institutions, regardless of size or resources.

Ultimately, AI offers a transformative opportunity to shift healthcare cybersecurity from a reactive posture to a proactive, adaptive defense mechanism. However, realizing this potential will depend not on technology alone, but on the ethical and strategic choices made by healthcare leaders, policymakers, and researchers alike.

References

1. Ahmad, R., Alsmadi, I., Alhamdani, W., & Tawalbeh, L. (2023). Zero-day attack detection: a systematic literature review. *Artificial Intelligence Review*, 56(10), 10733–10811. <https://doi.org/10.1007/s10462-023-10437-z>
2. Ali, S., Rehman, S. U., Imran, A., Adeem, G., Iqbal, Z., & Kim, K. I. (2022). Comparative Evaluation of AI-Based Techniques for Zero-Day Attacks Detection. *Electronics* (Switzerland), 11(23). <https://doi.org/10.3390/electronics11233934>
3. Al-Mhiqani, M. N., Ahmad, R., Abidin, Z. Z., Yassin, W., Hassan, A., Abdulkareem, K. H., ... Yunus, Z. (2020, August 1). A review of insider threat detection: Classification, machine learning techniques, datasets, open challenges, and recommendations. *Applied Sciences* (Switzerland). MDPI AG. <https://doi.org/10.3390/app10155208>
4. Almotiri, S. H., Nadeem, M., Al Ghamdi, M. A., & Khan, R. A. (2023). Analytic Review of Healthcare Software by Using Quantum Computing Security Techniques. *International Journal of Fuzzy Logic and Intelligent Systems*. Korean Institute of Intelligent Systems. <https://doi.org/10.5391/IJFIS.2023.23.3.336>
5. Fernandez De Arroyabe, I., Arranz, C. F. A., Arroyabe, M. F., & Fernandez de Arroyabe, J. C. (2023). Cybersecurity capabilities and cyber-attacks as drivers of investment in cybersecurity systems: A UK survey for 2018 and 2019. *Computers and Security*, 124. <https://doi.org/10.1016/j.cose.2022.102954>
6. Guo, Y. (2023, January 15). A review of Machine Learning-based zero-day attack detection: Challenges and future directions. *Computer Communications*. Elsevier B.V. <https://doi.org/10.1016/j.comcom.2022.11.001>
7. Kim, A., Oh, J., Ryu, J., & Lee, K. (2020). A review of insider threat detection approaches with IoT perspective. *IEEE Access*, 8, 78847–78867. <https://doi.org/10.1109/ACCESS.2020.2990195>
8. Kim, J., Park, M., Kim, H., Cho, S., & Kang, P. (2019). Insider threat detection based on user behavior modeling and anomaly detection algorithms. *Applied Sciences* (Switzerland), 9(19). <https://doi.org/10.3390/app9194018>
9. Kumar, A., Maskara, R., Maskara, S., & Chiang, I. J. (2014). Conceptualization and application of an approach for designing healthcare software interfaces. *Journal of Biomedical Informatics*, 49, 171–186. <https://doi.org/10.1016/j.jbi.2014.02.007>
10. Nyakasoka, L., & Naidoo, R. (2022). Barriers to dynamic cybersecurity capabilities in healthcare software services. In *EPiC Series in Computing* (Vol. 85, pp. 231–242). EasyChair. <https://doi.org/10.29007/8g47>

11. Ronchieri, E., & Canaparo, M. (2023). Assessing the impact of software quality models in healthcare software systems. *Health Systems*, 12(1), 85–97. <https://doi.org/10.1080/20476965.2022.2162445>
12. Sarhan, M., Layeghy, S., Gallagher, M., & Portmann, M. (2023). From zero-shot machine learning to zero-day attack detection. *International Journal of Information Security*, 22(4), 947–959. <https://doi.org/10.1007/s10207-023-00676-0>
13. Sarker, I. H., Kayes, A. S. M., Badsha, S., Alqahtani, H., Watters, P., & Ng, A. (2020). Cybersecurity data science: an overview from machine learning perspective. *Journal of Big Data*, 7(1). <https://doi.org/10.1186/s40537-020-00318-5>
14. Slapničar, S., Vuko, T., Čular, M., & Drašček, M. (2022). Effectiveness of cybersecurity audit. *International Journal of Accounting Information Systems*, 44. <https://doi.org/10.1016/j.accinf.2021.100548>
15. Taherdoost, H. (2022, July 1). Understanding Cybersecurity Frameworks and Information Security Standards—A Review and Comprehensive Overview. *Electronics* (Switzerland). MDPI. <https://doi.org/10.3390/electronics11142181>

