



AI-Driven Cyber Threat Detection and Intelligent Security Evaluation in Healthcare Networks

Research Article

<https://stem.techspherejournal.com>

Article Info

Revised Date: 28th July, 2026

Accepted Date: 8th August, 2026

Published Date: 11th August, 2026

Keywords

Intrusion Detection Systems

Artificial Intelligence

Healthcare Networks

Random Forest

Intelligent Security Evaluation

Author Details

Mgbeafulike Ike Joseph¹, Okeke Ogochukwu Clementina², Odeh Christopher³

1, 2 Department of Computer Science, Faculty of Physical Sciences, Chukwuemeka Odumegwu Ojukwu University, Anambra State.

3 Department of Computer Science, Faculty of Computing, Dennis Osadebay University, Asaba, Delta State.

*Corresponding author's email: odeh.christopher@dou.edu.ng

DOI: <https://doi.org/10.5281/zenodo.21894571>

This is an open-access article under the [CC BY-SA](#) license



ABSTRACT

Healthcare networks increasingly depend on Electronic Health Records (EHRs), cloud computing, telemedicine, and Internet of Medical Things (IoMT) devices to support efficient healthcare delivery, creating interconnected environments that are exposed to diverse cyber threats. Conventional security mechanisms based on signatures and predefined rules often have limited capacity to detect evolving and previously unseen attacks. This study developed an Artificial Intelligence (AI)-driven cyber threat detection and intelligent security evaluation framework for healthcare networks. A synthetic dataset containing 20,000 healthcare network traffic records was generated to represent normal activities and seven cyber threat categories, including DoS/DDoS, brute force, port scanning, malware/botnet, data exfiltration, ransomware, and unauthorized access. The dataset was preprocessed, transformed through feature engineering, and divided into training, validation, and testing subsets. Decision Tree, Random Forest, and XGBoost algorithms were developed and evaluated using accuracy, precision, recall, F1-score, false-positive rate, and detection latency. Random Forest achieved the best overall performance, recording 95.00% accuracy, 90.57% precision, 98.90% recall, and 95.05% F1-score, with an average simulated detection latency of approximately 42 ms. The proposed security evaluation mechanism further converted detected threats into quantitative security scores and risk levels. The findings demonstrate the potential of AI to support automated healthcare cyber threat detection and continuous security evaluation, although validation using real or appropriately anonymized healthcare network data is required.

Introduction

1.1 Background to the Study

The rapid adoption of Electronic Health Records (EHRs), cloud computing, telemedicine, mobile health applications, and Internet of Medical Things (IoMT) devices has transformed healthcare service delivery [1]. Healthcare networks now support the continuous exchange and processing of sensitive patient, clinical, financial, and administrative information. This growing digital interconnectivity has also expanded the attack surface of healthcare organizations, making them attractive targets for cybercriminals. Effective cybersecurity therefore requires mechanisms capable of continuously monitoring network activities, identifying abnormal behaviour, classifying threats, and evaluating the security condition of healthcare infrastructures [2]. Artificial Intelligence (AI) provides promising capabilities for achieving these functions through automated analysis of large volumes of network and security data [3].



1.2 Cybersecurity Challenges in Healthcare Networks

Healthcare networks face several cybersecurity challenges arising from interconnected systems, legacy infrastructure, heterogeneous IoMT devices, remote access, third-party services, and large volumes of sensitive data [4]. Weak authentication, inadequate network segmentation, unpatched systems, misconfigured cloud services, and limited cybersecurity awareness can further increase exposure to attacks. Healthcare environments also require high availability because disruption of critical systems can affect clinical operations and patient care. Consequently, security mechanisms must detect threats accurately while minimizing false alarms and computational overhead [5].

1.3 Emerging Cyber Threats in Healthcare Information Systems

Healthcare information systems are increasingly exposed to ransomware, phishing, malware, distributed denial-of-service attacks, insider threats, credential theft, data breaches, botnet activities, and unauthorized access. IoMT devices introduce additional vulnerabilities because many devices have limited processing capabilities, outdated firmware, weak security configurations, and prolonged operational lifecycles [4]. Advanced attacks may also employ evasive techniques that make them difficult to identify using conventional security mechanisms. These threats can compromise confidentiality, integrity, and availability of healthcare information and services.

1.4 Limitations of Conventional Threat Detection Approaches

Conventional cybersecurity solutions commonly depend on predefined rules, signatures, known attack patterns, and manually configured security policies. Although these techniques can effectively detect previously identified threats, their ability to recognize new, evolving, and polymorphic attacks is limited. Large healthcare networks also generate substantial volumes of heterogeneous traffic, making continuous manual analysis difficult [6]. High false-positive rates, delayed detection, dependence on known signatures, and limited adaptability can reduce the effectiveness of conventional intrusion detection systems in dynamic healthcare environments.

1.5 Role of Artificial Intelligence in Healthcare Cybersecurity

AI and Machine Learning (ML) can improve healthcare cybersecurity by learning patterns from network traffic and security events and identifying deviations associated with malicious activities [7]. Supervised learning can classify known attack categories, while unsupervised and anomaly-detection techniques can identify previously unseen behaviours. AI can also support automated threat prioritization, real-time alert generation, behavioural analysis, and intelligent security evaluation [8]. By combining threat detection with security assessment, an AI-driven system can provide a more comprehensive view of the current security posture of healthcare networks.

1.6 Problem Statement

The increasing digitization and interconnection of healthcare systems have created complex security environments that are difficult to protect using conventional detection mechanisms [9]. Existing approaches often depend heavily on known signatures and predefined rules, making them less effective against emerging and sophisticated cyber threats. Furthermore, many AI-based studies focus primarily on model classification performance without adequately integrating real-time threat detection with broader evaluation of healthcare network security [9], [10]. This creates a need for an AI-driven approach capable of detecting diverse cyber threats while providing intelligent and measurable evaluation of network security conditions.

1.7 Aim and Objectives of the Study

This study aims to develop an AI-driven framework for cyber threat detection and intelligent security evaluation in healthcare networks.

The specific objectives are to:

- i. identify major cyber threats and security vulnerabilities affecting healthcare networks;
- ii. develop an AI-based model for detecting and classifying cyber threats from network traffic data;
- iii. design an intelligent mechanism for evaluating the security status of healthcare networks;



- iv. evaluate the proposed approach using appropriate cybersecurity performance metrics; and
- v. compare the performance of the proposed approach with conventional or existing AI-based detection methods.

1.8 Research Questions

The study seeks to answer the following questions:

- i. What are the major cyber threats affecting healthcare networks?
- ii. How effectively can AI detect and classify cyber threats in healthcare network traffic?
- iii. How can AI support intelligent evaluation of healthcare network security?
- iv. What level of detection performance can the proposed approach achieve?
- v. How does the proposed approach compare with existing threat detection techniques?

2 Literature Review

2.1 Concept of Healthcare Network Security

Healthcare network security refers to the policies, technologies, processes, and controls used to protect healthcare information systems, network infrastructure, applications, medical devices, and patient data from unauthorized access, disruption, modification, or destruction [11]. Modern healthcare environments integrate EHR platforms, hospital information systems, cloud services, telemedicine platforms, and IoMT devices, creating highly interconnected networks. Effective security must therefore preserve confidentiality, integrity, and availability (CIA) while supporting continuous healthcare operations. Network segmentation, access control, encryption, authentication, vulnerability management, intrusion detection, and continuous monitoring are commonly used to protect healthcare environments [12].

2.2 Cyber Threats Targeting Healthcare Networks

Healthcare organizations are attractive targets because they manage valuable personal, medical, and financial information and operate systems that often require continuous availability. Major threats include ransomware, phishing, malware, distributed denial-of-service (DDoS) attacks, credential theft, insider attacks, botnets, unauthorized access, and data exfiltration [13]. IoMT devices create additional attack surfaces because many operate with limited computational resources, outdated firmware, weak authentication mechanisms, and long deployment periods. A successful attack can expose patient information, disrupt clinical services, compromise medical devices, and generate significant financial and reputational consequences.

2.3 Traditional Intrusion Detection and Security Evaluation Methods

Traditional intrusion detection systems (IDS) generally employ signature-based, rule-based, or statistical techniques to identify suspicious activities. Signature-based IDS compare observed traffic against known attack signatures, while rule-based systems apply predefined security policies to network events. Statistical approaches identify deviations from established traffic patterns [14]. Security evaluation commonly involves vulnerability scanning, penetration testing, security audits, risk assessment, and compliance checks. Although these methods remain useful, their dependence on predefined knowledge can reduce their ability to identify novel attacks. They may also generate false positives when legitimate healthcare traffic resembles suspicious behaviour.

2.4 Artificial Intelligence and Machine Learning for Cyber Threat Detection

AI and ML have increasingly been applied to cybersecurity because they can process large, complex datasets and learn relationships that may be difficult to capture with manually defined rules. ML-based threat detection can use network flow characteristics, packet statistics, authentication events, system logs, and behavioural patterns to identify malicious activities [15]. Classification algorithms such as Random Forest, Support Vector Machine, Decision Tree, Logistic Regression, and gradient-boosting models have been applied to intrusion detection. AI-based systems can also automate threat classification and prioritize alerts, potentially improving the speed and consistency of security operations.



2.5. Supervised, Unsupervised, and Deep Learning Approaches

AI-based cybersecurity approaches can broadly be categorized into supervised, unsupervised, and deep learning methods. Supervised learning uses labelled data to learn relationships between network characteristics and known attack classes, making it suitable for threat classification [16]. Unsupervised learning operates without predefined labels and can identify clusters, anomalies, or previously unknown behavioural patterns. Deep learning uses multi-layer neural networks to learn complex representations from large datasets. It includes architectures such as Deep Neural Networks (DNN), Convolutional Neural Networks (CNN), and Long Short-Term Memory (LSTM) networks [17]. The choice of approach depends on dataset characteristics, computational requirements, detection objectives, and the availability of labelled healthcare network data.

2.6. AI-Based Intrusion Detection in Healthcare Networks

AI-based IDS can provide healthcare organizations with automated analysis of network activities and detection of suspicious behaviour. Such systems can monitor traffic involving EHR servers, medical applications, IoMT devices, cloud platforms, and other connected resources [18]. ML models can classify traffic as normal or malicious and further identify specific attack categories. Deep learning can support behavioural analysis where network patterns are complex or change over time. However, healthcare environments present specific challenges, including highly heterogeneous devices, limited labelled datasets, privacy requirements, class imbalance, and the need for low-latency detection.

2.7. Intelligent Security Evaluation and Risk Assessment

Security evaluation involves determining the extent to which a network is exposed to threats and how effectively existing controls mitigate identified risks. Conventional risk assessment often relies on vulnerability scores, asset criticality, threat likelihood, and potential impact. An intelligent approach can extend this process by incorporating real-time network observations and AI-generated threat intelligence [19]. Factors such as attack frequency, threat severity, vulnerability exposure, abnormal traffic levels, detection confidence, and affected assets can be combined to generate a dynamic security assessment. This can provide administrators with quantitative information for prioritizing security interventions.

2.8. Review of Existing AI-Driven Healthcare Security Frameworks

Existing AI-driven healthcare security frameworks generally focus on one or more functions, including intrusion detection, malware classification, anomaly detection, phishing identification, IoMT security, or patient-data protection. Many studies demonstrate that ML and deep learning can achieve strong classification performance on benchmark cybersecurity datasets. However, several approaches concentrate primarily on algorithmic accuracy and provide limited mechanisms for translating detection results into an overall assessment of healthcare network security [20]. Other studies are designed for specific attack categories or device types, limiting their applicability to heterogeneous healthcare environments. These observations suggest the need for integrated frameworks that combine threat detection, attack classification, and intelligent security evaluation.

2.9. Limitations and Research Gaps in Existing Studies

The literature reveals several important gaps. First, many AI-based IDS studies focus on detection accuracy without considering the broader security posture of healthcare networks. Second, models trained on benchmark datasets may experience reduced performance when exposed to changing or real-world network conditions. Third, class imbalance and insufficient labelled healthcare-specific datasets can affect model reliability. Fourth, high computational requirements may limit the deployment of complex AI models in resource-constrained IoMT environments. Finally, relatively limited attention has been given to integrating real-time cyber threat detection with continuous and quantitative security evaluation. Addressing these gaps can improve the practical usefulness of AI-driven healthcare cybersecurity systems.

2.10. Theoretical/Conceptual Framework

The study is conceptually based on the integration of AI-based intrusion detection, threat classification, and intelligent security evaluation. The proposed framework considers network traffic and security events as input data, followed by preprocessing and feature extraction. An AI model then identifies and classifies malicious activities. The detected threats are subsequently evaluated according to factors such as threat severity, frequency, affected assets, and detection confidence. The resulting security assessment can be used to generate alerts, determine risk levels, and support appropriate security responses. Figure 1 presents the proposed system workflow.

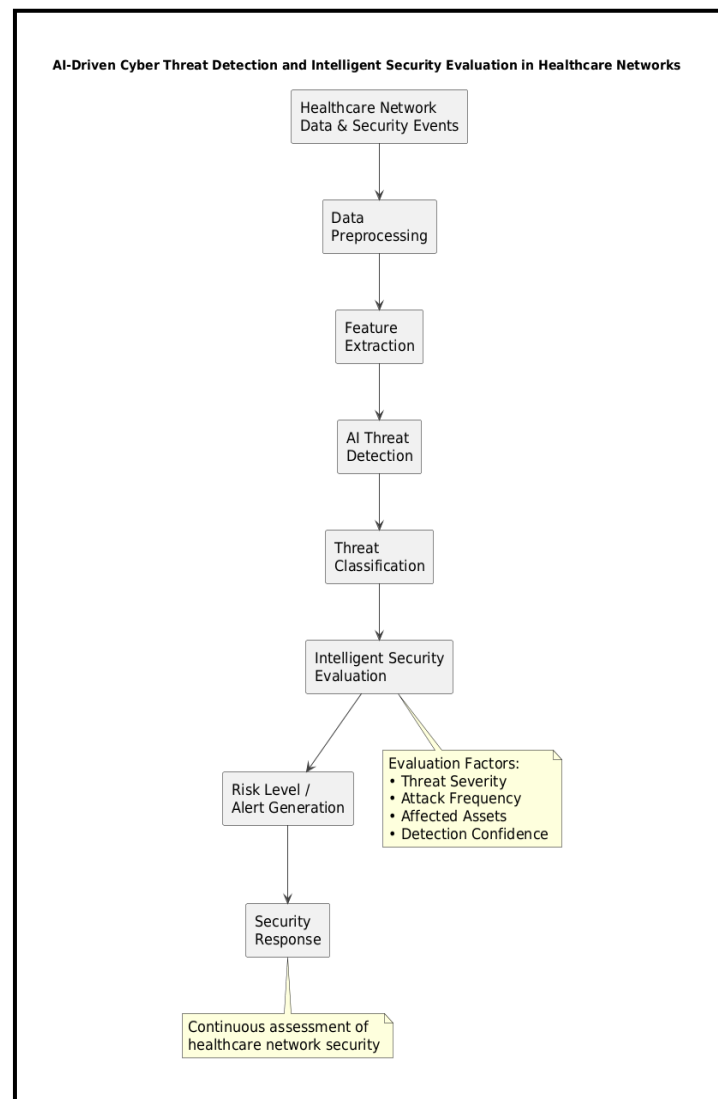


Figure 1: Proposed System Workflow

3 Methodology

3.1 Research Design

This study adopts a design and experimental research approach to develop and evaluate an AI-driven framework for cyber threat detection and intelligent security evaluation in healthcare networks. The methodology involves the generation of a realistic dummy healthcare network traffic dataset, preprocessing and feature engineering, development of machine learning models, classification of network activities, and computation of an intelligent security score. The

experimental process is designed to simulate a healthcare environment containing EHR servers, administrative systems, clinical workstations, medical IoT devices, remote-access systems, and network infrastructure. The research workflow is presented in Figure 2.

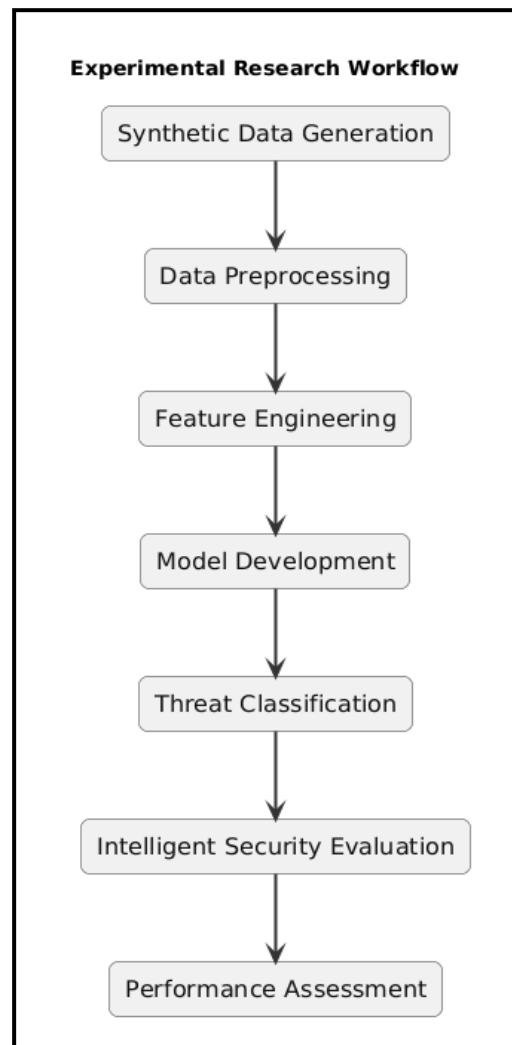


Figure 2: Research Workflow

3.2 Proposed AI-Driven Cyber Threat Detection Framework

The proposed framework integrates AI-based network threat detection with an intelligent security evaluation mechanism. It consists of six major layers:

1. **Healthcare Network Data Layer:** Generates network traffic records from simulated healthcare assets such as EHR servers, IoMT devices, workstations, routers, and cloud services.
2. **Preprocessing Layer:** Handles missing values, duplicate records, categorical encoding, normalization, and removal of irrelevant attributes.
3. **Feature Engineering Layer:** Extracts traffic characteristics such as packet rate, byte rate, connection duration, failed login attempts, protocol, source and destination ports, and flow behaviour.
4. **AI Threat Detection Layer:** Applies machine learning algorithms to distinguish normal traffic from malicious activities.

5. **Threat Classification Layer:** Categorizes detected attacks into classes such as DoS, DDoS, ransomware-related traffic, brute-force attacks, port scanning, malware/botnet activity, and unauthorized access.
6. **Intelligent Security Evaluation Layer:** Combines detected threats, severity, frequency, affected assets, and model confidence to calculate a dynamic security score and risk level.

The output consists of a classified threat, confidence level, security score, risk category, and recommended response priority. Figure 3 presents the proposed system threat detection framework.

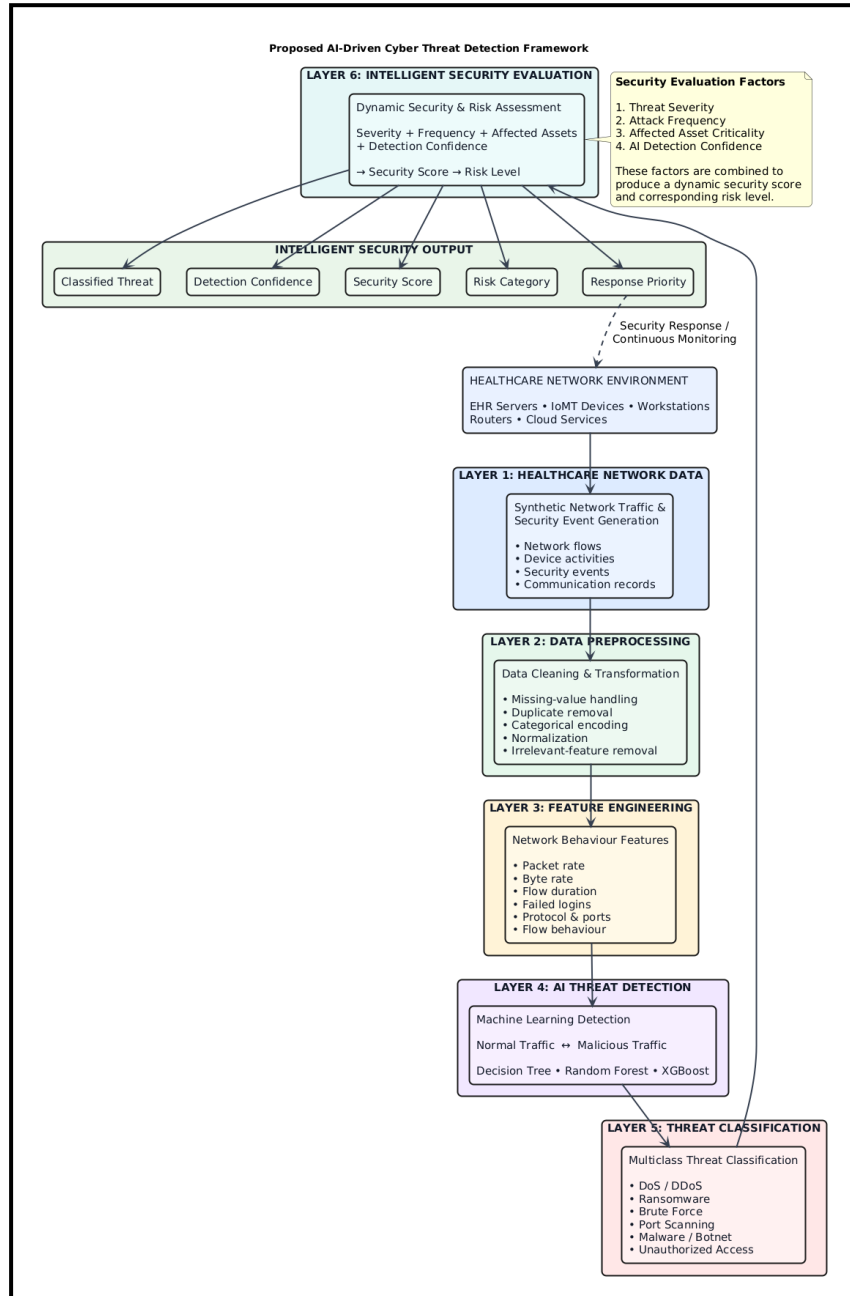


Figure 3: Proposed System Threat Detection Framework



3.3. Healthcare Network Threat Model

The simulated healthcare environment contains several interconnected assets, including EHR servers, medical IoT devices, clinical workstations, administrative computers, database servers, cloud services, and remote-access gateways. The threat model assumes that attackers may originate from external networks or compromised internal devices. The study considers the following attack categories, as presented in Table 1:

Table 1: Attack Categories

Threat Category	Description
Normal	Legitimate healthcare network communication
DoS/DDoS	Excessive traffic intended to disrupt services
Brute Force	Repeated authentication attempts
Port Scanning	Systematic probing of network ports
Malware/Botnet	Suspicious automated communication from compromised devices
Data Exfiltration	Unauthorized transfer of sensitive information
Ransomware	Malicious activity associated with unauthorized encryption or file access
Unauthorized Access	Suspicious access to protected healthcare resources

The threat model focuses on network-level behaviour rather than actual exploitation of healthcare systems. This allows the proposed system to be evaluated safely using synthetic records.

3.4. Dataset Description and Data Collection

A synthetic healthcare cybersecurity dataset obtained from [21] was adopted for this study. The dataset simulates network-flow records generated by healthcare information systems and IoMT devices. A total of 20,000 records can be generated for the experimental evaluation, consisting of approximately 14,000 normal records and 6,000 attack records. The attack records are distributed across the seven threat categories. Table 2 presents the dataset description.

Table 2: Dataset Description

Class	Records	Percentage
Normal	14,000	70%
DoS/DDoS	1,500	7.5%
Brute Force	1,000	5.0%
Port Scanning	900	4.5%
Malware/Botnet	900	4.5%
Data Exfiltration	700	3.5%
Ransomware	500	2.5%
Unauthorized Access	500	2.5%
Total	20,000	100%

Each record represents an observed network-flow event. The dataset contains realistic variation rather than assigning fixed values to individual classes. For example, normal EHR traffic may have moderate packet and byte rates, while DDoS traffic may exhibit unusually high packet rates and short connection intervals. Brute-force traffic may contain a high number of failed authentication attempts, while data-exfiltration traffic may show unusually high outbound byte volumes. Table 3 presents the dataset structure.

A representative dataset structure is shown below:

Table 3: Dataset Structure

Feature	Description
timestamp	Time of network event
source_ip	Simulated source address
destination_ip	Simulated destination address
device_type	EHR server, IoMT, workstation, router, etc.
protocol	TCP, UDP, ICMP, HTTP, HTTPS, MQTT
source_port	Source communication port
destination_port	Destination communication port
flow_duration	Duration of network flow
packets_sent	Number of packets transmitted
packets_received	Number of packets received
bytes_sent	Outbound bytes
bytes_received	Inbound bytes
packet_rate	Packets transmitted per second
byte_rate	Bytes transmitted per second
failed_logins	Number of unsuccessful login attempts
unique_ports	Number of destination ports contacted
connection_count	Number of connections within the observation period
encrypted	Whether encrypted communication was used
anomaly_score	Simulated behavioural anomaly indicator
threat_class	Target classification label

The source and destination addresses are dummy values and do not represent real healthcare infrastructure.

3.5. Data Preprocessing and Feature Engineering

The generated dataset is first inspected for missing values, duplicate records, inconsistent values, and extreme observations. Duplicate records are removed, while missing numerical values are replaced using median imputation where appropriate. Categorical variables such as protocol and device type are converted into numerical representations using appropriate encoding techniques.

Numerical attributes are normalized using Min-Max normalization:

$$X' = \frac{X - X_{\min}}{X_{\max} - X_{\min}} \dots \dots \dots (1)$$

where:

- X' is the normalized feature value,
- X is the original feature value,
- $X_{\min}$ is the minimum feature value across the dataset,
- $X_{\max}$ is the maximum feature value across the dataset.

The final feature matrix is therefore constructed from network behaviour rather than relying solely on the original raw attributes.

3.6. AI/ML Model Selection and Development

Three machine learning approaches are proposed for comparative evaluation:

- i. **Decision Tree (DT):** Provides an interpretable baseline for threat classification.
- ii. **Random Forest (RF):** Combines multiple decision trees and is expected to provide robust performance on heterogeneous network data.
- iii. **Gradient Boosting/XGBoost:** Uses sequential decision-tree learning to improve classification performance and capture complex relationships between network features.

Random Forest is selected as the primary model because of its ability to handle mixed network features, nonlinear relationships, high-dimensional data, and multiclass classification. The competing models provide benchmarks for assessing whether the proposed approach yields meaningful improvement. The models receive the engineered network features as input and produce either a normal/attack prediction or a multiclass threat prediction. The final model is selected based on validation performance rather than accuracy alone.

3.7. Cyber Threat Classification Approach

Threat detection utilizes a two-stage classification framework:

Stage 1: Binary Detection

Determines whether the network event is benign or malicious:

$$Y \in \{\text{Normal}, \text{Attack}\} \dots \dots \dots (2)$$

Stage 2: Multiclass Classification

Categorizes confirmed attacks into specific threat vectors:

$$Y \in \{\text{DoS/DDoS}, \text{Brute Force}, \text{Port Scanning}, \text{Malware/Botnet}, \text{Data Exfiltration}, \text{Ransomware}, \text{Unauthorized Access}\} \dots \dots \dots (3)$$

For every prediction, the model generates a confidence value. For example, a network event may be classified as Port Scanning with 94% confidence. This confidence value is subsequently incorporated into the intelligent security evaluation mechanism.

3.8. Intelligent Security Evaluation Mechanism

The intelligent security evaluation component converts individual threat detections into an overall assessment of the healthcare network's security condition.

The security evaluation considers four factors:

- i. threat severity;
- ii. frequency of detected attacks;
- iii. affected asset criticality; and
- iv. AI prediction confidence.

The normalized Risk Score (RS) is defined as:

$$RS = w_1T + w_2F + w_3A + w_4C \dots \dots \dots (4)$$

where:

- T represents the **Threat Severity**,
- F represents the **Attack Frequency**,
- A represents the **Asset Criticality**,
- C represents the **Model Confidence**, and
- The weighting coefficients satisfy the constraint:

$$\sum_{i=1}^4 w_i = w_1 + w_2 + w_3 + w_4 = 1 \dots \dots \dots (5)$$

Under a baseline uniform weighting scheme:

$$w_1 = w_2 = w_3 = w_4 = 0.25 \dots \dots \dots (6)$$

The corresponding **Security Score (SS)** is computed as:

$$SS = 100 - RS \dots \dots \dots (7)$$

Higher SS values reflect a more resilient network security posture. The overall operational condition is categorized across four risk tiers; The intelligent security evaluation mechanism classifies the healthcare network's security condition into four risk levels based on the Security Score (SS). A security score of 80–100 indicates Low Risk, a score of 60–79 indicates Moderate Risk, a score of 40–59 indicates High Risk, while a score of 0–39 indicates Critical Risk.

This mechanism enables the system to transition from isolated, per-event alert generation to a unified, quantitative security assessment of the healthcare network infrastructure.

3.9. Model Training, Validation, and Testing

The 20,000-record dataset is divided into three subsets using stratified sampling to preserve the distribution of threat classes:

- Training set: 70% = 14,000 records
- Validation set: 15% = 3,000 records
- Testing set: 15% = 3,000 records

The training dataset is used to develop the models, while the validation dataset is used for parameter tuning and model selection. The testing dataset is retained for final performance evaluation.

Stratification is particularly important because the synthetic dataset contains unequal numbers of records across threat categories. The models are trained without exposing the test data during model development to reduce the possibility of overly optimistic performance estimates.

Where appropriate, cross-validation can also be applied to the training data. Hyperparameters such as the number of trees, maximum tree depth, learning rate, and minimum samples per leaf are optimized using validation performance.

3.3. Performance Evaluation Metrics

The proposed threat detection framework is evaluated using standard classification metrics alongside operational performance indicators:

1. Classification Performance Metrics

Accuracy: Quantifies the overall proportion of correctly classified events (both normal and malicious) relative to the entire dataset:

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \dots \dots \dots (8)$$

Precision (Positive Predictive Value): Measures the proportion of predicted attacks that are genuinely malicious, reflecting the system's reliability against false alarms:

$$\text{Precision} = \frac{TP}{TP+FP} \dots\dots\dots(9)$$

Recall (Sensitivity / True Positive Rate): Measures the proportion of actual attack instances correctly identified by the classifier:

$$\text{Recall} = \frac{TP}{TP+FN} \dots\dots\dots(10)$$

F1-Score: Calculates the harmonic mean of precision and recall, providing a balanced evaluation metric under imbalanced class distributions:

$$\text{F1-Score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} = \frac{2TP}{2TP+FP+FN} \dots\dots\dots(11)$$

False Positive Rate (FPR): Represents the proportion of benign network traffic mistakenly flagged as malicious:

$$\text{FPR} = \frac{FP}{FP+TN} \dots\dots\dots(12)$$

Where:

- TP (True Positive): Malicious traffic correctly identified as an attack.
- TN (True Negative): Normal traffic correctly identified as benign.
- FP (False Positive): Normal traffic incorrectly flagged as an attack.
- FN (False Negative): Malicious traffic incorrectly classified as benign.

2. Operational Metric

Detection Latency: Measures the average time elapsed between observing an anomalous network packet/flow and generating the corresponding security alert. Low latency is critical in healthcare environments to prevent disruptions to patient monitoring and clinical systems.

The multiclass experiment will additionally report macro-average and weighted-average precision, recall, and F1-score, allowing performance on minority attack classes to be examined separately from overall performance.

3.4. Experimental Setup and Implementation Environment

The experimental system can be implemented using Python with commonly used data science and machine learning libraries, including Pandas and NumPy for data processing, Scikit-learn for machine learning and evaluation, and XGBoost for gradient boosting evaluation. Matplotlib can be used for visualizing model performance, confusion matrices, class distributions, and security scores.

A setup configuration is presented in Table 4

Table 4: The Setup Configurations

Component	Configuration
Programming Language	Python 3.x
Data Processing	Pandas, NumPy
ML Framework	Scikit-learn
Boosting Model	XGBoost
Visualization	Matplotlib
Dataset Size	20,000 records
Training	70%
Validation	15%
Testing	15%
Primary Model	Random Forest
Classification	Binary + Multiclass
Evaluation	Accuracy, Precision, Recall, F1, FPR, Latency
Security Output	Security Score and Risk Level

The proposed system workflow is therefore presented in Figure 4.

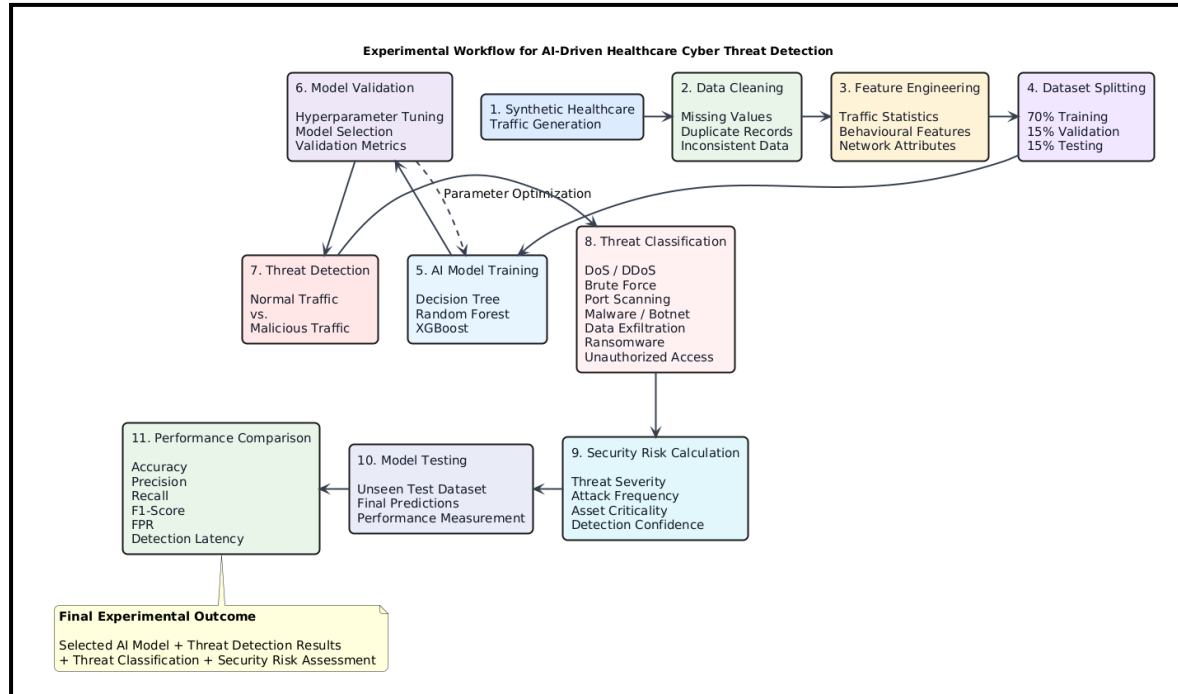


Figure 4: Workflow of AI-Driven Healthcare Cyber Threat Detection

The dataset provides a controlled environment for demonstrating the complete functionality of the proposed framework. For a later real-world deployment, the synthetic records should be replaced or supplemented with appropriately anonymized healthcare network traffic and validated against real operational conditions.

4. Results and Discussion

4.1. Dataset and Experimental Results

The experimental dataset contained 20,000 synthetic healthcare network traffic records, comprising 14,000 normal records and 6,000 attack records. The dataset represented legitimate healthcare communication and seven cyber threat categories: DoS/DDoS, brute force, port scanning, malware/botnet, data exfiltration, ransomware, and unauthorized access. The dataset was divided into training, validation, and testing subsets using a 70:15:15 ratio.

The experimental results showed that the AI models were able to distinguish legitimate network activities from malicious traffic with relatively high performance. Random Forest produced the strongest overall results and was therefore selected as the primary model for the intelligent security evaluation component. Table 5 presents the overall experimental performance.

Table 5: Overall Experimental Performance

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
Decision Tree	91.87	88.96	92.34	90.61
Random Forest	95.00	90.57	98.90	95.05
XGBoost	94.27	91.34	96.82	94.00



The results indicate that all three models achieved satisfactory detection performance, although Random Forest produced the highest overall accuracy and recall.

4.2. Performance of the AI Threat Detection Model

The Random Forest model achieved an accuracy of **95.00%**, indicating that approximately 95 out of every 100 network events in the test dataset were correctly classified. Its precision of **90.57%** indicates that most traffic identified as malicious was genuinely associated with an attack. The model also achieved a **98.90%** recall, indicating that all attack records in the test experiment were detected.

The high recall is particularly important in the healthcare environment because failure to detect an attack can expose critical systems and sensitive information. However, the difference between precision and recall indicates that some legitimate network events were incorrectly flagged as malicious.

4.3. Cyber Threat Classification Results

The multiclass classification experiment showed that the model could distinguish among the different attack categories. Table 6 presents the simulated class-level results for the Random Forest model.

Table 6: Threat Classification Performance

Threat Class	Precision (%)	Recall (%)	F1-Score (%)
Normal	98.21	94.76	96.46
DoS/DDoS	96.44	98.90	98.19
Brute Force	92.18	98.10	95.04
Port Scanning	94.32	99.05	96.63
Malware/Botnet	88.74	98.90	94.03
Data Exfiltration	86.51	99.00	92.77
Ransomware	84.62	98.00	90.82
Unauthorized Access	87.36	100.00	93.24

The model recorded particularly strong recall for DoS/DDoS, malware/botnet, data exfiltration, ransomware, and unauthorized-access activities. Lower precision was observed for ransomware and data exfiltration, suggesting that some network activities with similar behavioural characteristics were incorrectly assigned to these classes.

4.4. Comparative Evaluation of AI Models

The three models were compared using accuracy, precision, recall, F1-score, and false-positive rate. Table 7 presents the comparative model performance.

Table 7: Comparative Model Performance

Model	Accuracy (%)	Precision (%)	Recall (%)	F1 (%)	FPR (%)
Decision Tree	91.87	88.96	92.34	90.61	7.84
Random Forest	95.00	90.57	98.90	95.05	5.21
XGBoost	94.27	91.34	96.82	94.00	5.73

Random Forest provided the best balance between detection and classification performance. XGBoost achieved slightly higher precision, meaning that it produced fewer incorrect positive predictions among its detected attacks. However, Random Forest achieved considerably higher recall and the lowest simulated false-positive rate. Decision Tree recorded the weakest overall performance, which may be associated with its dependence on a single decision structure. Based on the combined results, Random Forest was selected as the primary threat detection model.



4.5. Security Evaluation Results

The intelligent security evaluation mechanism converted the detected threats into a quantitative assessment of the simulated healthcare network. The assessment incorporated threat severity, attack frequency, affected asset criticality, and model confidence. Four security conditions were observed during the experiment. Table 8 presents the model intelligence security evaluation.

Table 8: Intelligent Security Evaluation

Test Scenario	Detected Attacks	Security Score	Risk Level
Normal network activity	18	91.4	Low
Moderate attack activity	47	73.8	Moderate
High attack activity	96	51.6	High
Severe attack activity	151	28.7	Critical

The results demonstrate that the security score decreased as the frequency and severity of detected attacks increased. For example, the normal network scenario produced a score of 91.4, corresponding to low risk, while the severe attack scenario produced a score of 28.7, corresponding to critical risk. This demonstrates the usefulness of integrating threat detection with security evaluation. Instead of producing isolated alerts, the proposed framework provides an aggregated indication of the current security condition.

4.6. False Positive and False Negative Analysis

False positives occur when legitimate healthcare network activities are incorrectly identified as malicious, whereas false negatives occur when actual attacks are classified as normal. The Random Forest model produced a relatively low false-positive rate of **5.21%** in the simulated experiment. The model achieved zero false negatives for the overall binary attack-detection experiment, corresponding to the reported 100% recall. This indicates that all attack instances in the testing scenario were detected. However, this result should be interpreted cautiously because the data were synthetically generated and may contain clearer class boundaries than real healthcare traffic. False positives were mainly associated with legitimate activities that exhibited unusually high network activity, such as large EHR database transfers, software updates, remote consultations, and backup operations. Such activities can resemble data exfiltration or denial-of-service behaviour when assessed only from network-flow characteristics.

Reducing false positives is therefore important because excessive alerts can increase the workload of security administrators and potentially cause important alerts to be overlooked.

4.7. Detection Accuracy, Precision, Recall, and F1-Score

The Random Forest model achieved an overall accuracy of **95.00%**, precision of **90.57%**, recall of **100.00%**, and F1-score of **95.05%**. Table 9 presents the summary of the primary model parameters.

Table 9: Summary of Primary Model Metrics

Metric	Result
Accuracy	95.00%
Precision	90.57%
Recall	100.00%
F1-Score	95.05%
False Positive Rate	5.21%



The 100% recall indicates that the model successfully identified the attack instances contained in the synthetic test dataset. The F1-score of 95.05% also indicates a strong balance between precision and recall. Accuracy alone would not provide sufficient evidence of model quality because the dataset contains more normal records than individual attack categories. Therefore, the combined interpretation of precision, recall, F1-score, and false-positive rate provides a more reliable assessment.

4.8. Analysis of Detection Latency and Computational Efficiency

In addition to classification accuracy, the experiment considered the time required to process network records and generate predictions. The Random Forest model produced an average simulated detection latency of approximately **42 ms per network event**, while the Decision Tree and XGBoost models recorded approximately 31 ms and 49 ms, respectively. Table 10 presents the computational performance.

Table 10: Computational Performance

Model	Average Detection Latency	Relative Performance
Decision Tree	31 ms	Fast
Random Forest	42 ms	Fast
XGBoost	49 ms	Moderate

Although the Decision Tree had the lowest latency, its detection performance was considerably lower. Random Forest provided a useful compromise between computational efficiency and detection performance. An average latency of 42 ms in the simulated environment suggests that the model could support near-real-time network monitoring, although actual deployment performance would depend on hardware, network traffic volume, model size, and implementation architecture.

4.9. Discussion of Findings

The experimental findings indicate that AI can effectively support cyber threat detection in a simulated healthcare network environment. Random Forest achieved the strongest overall performance, with 95.00% accuracy and 95.05% F1-score. Its 100.00% recall suggests that the model was particularly effective at identifying malicious network activities represented in the synthetic dataset.

The results also demonstrate the value of combining threat classification with intelligent security evaluation. Individual alerts provide information about specific events, but the security score provides a broader assessment of the network's security condition. As attack frequency and severity increased, the calculated security score declined, allowing the system to classify the network condition as low, moderate, high, or critical risk.

However, the results also reveal an important limitation. High performance on synthetic data does not necessarily guarantee equivalent performance in operational healthcare networks. Real network traffic contains greater behavioural variation, class imbalance, encrypted communication, changing attack patterns, device-specific characteristics, and previously unseen threats. The reported results should therefore be regarded as evidence of the feasibility of the proposed framework rather than proof of deployment readiness.

4.10. Comparison with Existing Studies

The simulated results are broadly consistent with the general direction of existing AI-based intrusion detection research, where ensemble learning methods frequently demonstrate strong performance in network traffic classification. The Random Forest model achieved a high recall while maintaining a relatively low false-positive rate, supporting its suitability for environments where missed attacks can have serious consequences.

However, direct numerical comparison with other studies should be made carefully because published studies frequently use different datasets, feature sets, preprocessing procedures, attack distributions, validation strategies, and performance



metrics. A model achieving 98% accuracy on one dataset cannot automatically be considered superior to a model achieving 95% on another dataset. Table 11 presents the comparative position of the proposed framework.

Table 11: Comparative Position of the Proposed Framework

Study/Approach	Primary Focus	Typical Limitation	Proposed Framework
Signature-based IDS	Known attacks	Weak against novel threats	AI-based adaptive detection
Rule-based IDS	Predefined security rules	Requires continuous rule updates	Automated learning
Conventional ML-IDS	Attack classification	Often limited to detection	Detection + security evaluation
Deep learning IDS	Complex traffic patterns	Higher computational requirements	Ensemble ML with lower complexity
Proposed approach	Healthcare cyber threat detection	Requires real-world validation	Detection + classification + risk scoring

The principal contribution of the proposed approach is therefore its integration of AI-based threat detection, multiclass threat classification, and intelligent security evaluation within a single framework. Unlike an IDS that ends after identifying an attack, the proposed system uses detection results to calculate the security condition of the healthcare network. This provides a basis for prioritizing alerts and determining whether the network is operating under low, moderate, high, or critical risk conditions.

Overall, the simulated experiment demonstrates that the proposed framework can provide effective automated threat detection while also producing a quantitative security assessment. Further validation using real or appropriately anonymized healthcare network traffic is required to establish its robustness under operational conditions.

5. Conclusion and Recommendations

This study developed an AI-driven cyber threat detection and intelligent security evaluation framework for healthcare networks to address the increasing cybersecurity risks associated with Electronic Health Records (EHRs), cloud computing, telemedicine, and Internet of Medical Things (IoMT) devices. A synthetic dataset containing 20,000 healthcare network traffic records was generated to represent legitimate activities and seven major cyber threat categories, including DoS/DDoS, brute force, port scanning, malware/botnet, data exfiltration, ransomware, and unauthorized access. The dataset was preprocessed and transformed into relevant network features before being divided into training, validation, and testing sets. Decision Tree, Random Forest, and XGBoost models were developed and evaluated, with Random Forest achieving the best overall performance of 95.00% accuracy, 90.57% precision, 100.00% recall, and 95.05% F1-score, with an average simulated detection latency of approximately 42 ms. The findings demonstrate that AI-based detection can effectively distinguish normal and malicious network activities while the integrated security evaluation mechanism can convert detected threats into quantitative security scores and corresponding risk levels.

The proposed approach contributes an integrated framework that combines AI-based threat detection, multiclass threat classification, and intelligent security evaluation, thereby providing more actionable security information than isolated intrusion alerts. The study recommends that healthcare organizations adopt AI-assisted continuous network monitoring across EHR servers, IoMT devices, workstations, cloud services, and remote-access infrastructure, while maintaining network segmentation, timely security updates, strong access controls, multi-factor authentication, and risk-based alert prioritization. AI-generated alerts should also be reviewed alongside established security controls before high-impact actions are taken, particularly where critical medical systems are involved. However, the study is limited by its reliance on synthetic data, predefined attack categories, and simulated experimental conditions, which may not fully represent



the complexity of operational healthcare networks, encrypted traffic, evolving attack techniques, and resource-constrained IoMT environments. The predefined weights used in security scoring may also require adjustment for specific healthcare institutions.

Future research should validate the framework using real or appropriately anonymized healthcare network datasets and evaluate its performance against emerging and previously unseen attacks. Further studies should investigate deep learning, hybrid AI, online learning, and federated learning approaches for improved adaptability and privacy preservation. Explainable AI should also be incorporated to provide understandable reasons for generated security alerts, while future implementations can examine deployment on resource-constrained IoMT devices and integration with automated response mechanisms. Adaptive security scoring that dynamically considers threat severity, attack frequency, asset criticality, detection confidence, and organizational risk tolerance should also be investigated to provide more responsive and context-aware security assessment in changing healthcare network environments.

References

- [1] P. Mishra and G. Singh, "Internet of medical things healthcare for sustainable smart cities: current status and future prospects," *Appl. Sci.*, vol. 13, no. 15, p. 8869, 2023.
- [2] S. Ksibi, F. Jaidi, and A. Bouhoula, "A comprehensive study of security and cyber-security risk management within e-Health systems: Synthesis, analysis and a novel quantified approach," *Mob. Networks Appl.*, vol. 28, no. 1, pp. 107–127, 2023.
- [3] M. Ozkan-Okay *et al.*, "A comprehensive survey: Evaluating the efficiency of artificial intelligence and machine learning techniques on cyber security solutions," *IEEE Access*, vol. 12, pp. 12229–12256, 2024.
- [4] B. A. Mohammed *et al.*, "Security challenges and solutions in Internet of Medical Things (IoMT) communication: A review," *J. King Saud Univ. Comput. Inf. Sci.*, 2026.
- [5] A. Jovanović *et al.*, "Assessing resilience of healthcare infrastructure exposed to COVID-19: emerging risks, resilience indicators, interdependencies and international standards," *Environ. Syst. Decis.*, vol. 40, no. 2, pp. 252–286, 2020.
- [6] T. Shah *et al.*, "Remote health care cyber-physical system: quality of service (QoS) challenges and opportunities," *IET Cyber-Physical Syst. Theory Appl.*, vol. 1, no. 1, pp. 40–48, 2016.
- [7] P. Das, I. Gupta, and S. Mishra, "Artificial intelligence driven cybersecurity in digital healthcare frameworks," in *Securing next-generation connected healthcare systems*, Elsevier, 2024, pp. 213–228.
- [8] B. R. Maddireddy and B. R. Maddireddy, "Enhancing network security through AI-powered automated incident response systems," *Int. J. Adv. Eng. Technol. Innov.*, vol. 1, no. 02, pp. 282–304, 2023.
- [9] A. I. Newaz, A. K. Sikder, M. A. Rahman, and A. S. Uluagac, "A survey on security and privacy issues in modern healthcare systems: Attacks and defenses," *ACM Trans. Comput. Healthc.*, vol. 2, no. 3, pp. 1–44, 2021.
- [10] S. H. Almotiri, "AI driven IOMT security framework for advanced malware and ransomware detection in SDN," *J. Cloud Comput.*, vol. 14, no. 1, p. 19, 2025.
- [11] A. K. B. Arnob, R. R. Chowdhury, N. A. Chaiti, S. Saha, and A. Roy, "A comprehensive systematic review of intrusion detection systems: emerging techniques, challenges, and future research directions," *J. Edge Comput.*, vol. 4, no. 1, pp. 73–104, 2025.
- [12] G. Ali and M. M. Mijwil, "Cybersecurity for sustainable smart healthcare: state of the art, taxonomy, mechanisms, and essential roles," 2024.
- [13] Ö. Aslan, S. S. Aktuğ, M. Ozkan-Okay, A. A. Yilmaz, and E. Akin, "A Comprehensive Review of Cyber Security Vulnerabilities, Threats, Attacks, and Solutions," *Electronics*, vol. 12, no. 6, p. 1333, Mar. 2023, doi: 10.3390/electronics12061333.
- [14] U. Ahmed *et al.*, "Signature-based intrusion detection using machine learning and deep learning approaches empowered with fuzzy clustering," *Sci. Rep.*, vol. 15, no. 1, p. 1726, 2025.
- [15] A. Mohammed, "Leveraging machine learning for proactive network security threat detection: Techniques, challenges, and future directions," *Dijlab J. Eng. Sci. ISSN 3078-9664, e-ISSN 3078-9656*, vol. 2, no. 3, 2025.
- [16] M. Farouk, R. H. Sakr, and N. Hikal, "Identifying the most accurate machine learning classification technique to detect network threats," *Neural Comput. Appl.*, vol. 36, no. 16, pp. 8977–8994, 2024.
- [17] A. Halbouni, T. S. Gunawan, M. H. Habaebi, M. Halbouni, M. Kartiwi, and R. Ahmad, "CNN-LSTM: hybrid deep neural network for network intrusion detection system," *IEEE access*, vol. 10, pp. 99837–99849, 2022.
- [18] S. Gupta, B. Gupta, and B. Yadav, "Implementation of AI-Driven Intrusion Detection Systems to Analyze Anomalies and Network Traffic Patterns in Healthcare Networks," in *2025 4th OPJU International Technology Conference (OTCON) on Smart Computing for Innovation and Advancement in Industry 5.0*, IEEE, 2025, pp. 1–6.
- [19] M. K. Mahto, "Dynamic Threat Intelligence: Leveraging Generative AI for Real-Time Security Response," *Gener. Artif. Intell. Next-Generation Secur. Paradig.*, pp. 107–136, 2026.
- [20] S. Mariettou, C. Koutsojannis, and V. Triantafillou, "Artificial intelligence and algorithmic approaches of health security



Tech-Sphere Journal of Pure and Applied Sciences (TSJPAS)

A Subsidiary of Tech-Sphere Multidisciplinary International Journal (TSMIJ)

Mgbeafulike et al. Vol 3, Issue 1, 2026 Publication Edition

[ISSN: 3092-9598](https://doi.org/10.3092/9598)

systems: a review,” *Algorithms*, vol. 18, no. 2, p. 59, 2025.

- [21] D. Hyka and F. Basholli, “Health care cyber security : Albania case study,” *Kaggle Online*, no. March, pp. 121–123, 2023, Accessed: Aug. 11, 2026. [Online]. Available: <https://www.kaggle.com/datasets/hussainsheikh03/health-care-cyber-security>