

Central Lancashire Online Knowledge (CLoK)

Title	Leveraging Artificial Intelligence for Predictive CyberSecurity: Enhancing Threat Forecasting and Vulnerability Management
Type	Article
URL	https://clock.uclan.ac.uk/id/eprint/54514/
DOI	10.26562/ijirae.2025.v1202.01
Date	2025
Citation	Egho-Promise, Ehigiator Iyobor, Asante, George, Balisane, Hewa, Salih, Abdulrahman, Aina, Folayo, Kure, Halima and Gavua, Ebenezer Komla (2025) Leveraging Artificial Intelligence for Predictive CyberSecurity: Enhancing Threat Forecasting and Vulnerability Management. International Journal of Innovative Research in Advanced Engineering, 12 (02). pp. 68-79. ISSN 2349-2163
Creators	Egho-Promise, Ehigiator Iyobor, Asante, George, Balisane, Hewa, Salih, Abdulrahman, Aina, Folayo, Kure, Halima and Gavua, Ebenezer Komla

It is advisable to refer to the publisher's version if you intend to cite from the work.
10.26562/ijirae.2025.v1202.01

For information about Research at UCLan please go to <http://www.uclan.ac.uk/research/>

All outputs in CLoK are protected by Intellectual Property Rights law, including Copyright law. Copyright, IPR and Moral Rights for the works on this site are retained by the individual authors and/or other copyright owners. Terms and conditions for use of this material are defined in the <http://clock.uclan.ac.uk/policies/>

See discussions, stats, and author profiles for this publication at: <https://www.researchgate.net/publication/388872984>

Leveraging Artificial Intelligence for Predictive CyberSecurity: Enhancing Threat Forecasting and Vulnerability Management

Article in International Journal of Innovative Research in Advanced Engineering · February 2025

DOI: 10.26562/ijirae.2025.v1202.01

CITATIONS

0

READS

3

7 authors, including:



Ehigiator Iyobor Egho-Promise
University College Birmingham

35 PUBLICATIONS 43 CITATIONS

SEE PROFILE



George Asante
Akenten Appiah-Menka University of Skills Training and Entrepreneurial Develop...

19 PUBLICATIONS 28 CITATIONS

SEE PROFILE



Folayo Aina
University of Central Lancashire

15 PUBLICATIONS 64 CITATIONS

SEE PROFILE



Halima Kure
University of East London

8 PUBLICATIONS 301 CITATIONS

SEE PROFILE

Leveraging Artificial Intelligence for Predictive CyberSecurity: Enhancing Threat Forecasting and Vulnerability Management

Ehigiator Iyobor Egho-Promise

Department of Computer Science, University College Birmingham,
Birmingham, United Kingdom

eegho-promise@ucb.ac.uk ORCID ID: [0000-0001-8948-1813](https://orcid.org/0000-0001-8948-1813)

George Asante

Department of Information Technology Education,
Akenten Appiah-Menka University of Skills Training and Entrepreneurial Development,
Kumasi, Ghana

gasante@aamusted.edu.gh ORCID: [0000-0001-8061-5387](https://orcid.org/0000-0001-8061-5387)

Hewa Balisane

Business School, The University of Law, Manchester, United Kingdom

Hewa.Balisane@law.ac.uk ORCID: [0000-0003-2345-2336](https://orcid.org/0000-0003-2345-2336)

Abdulrahman Salih

Northumbria University London

Department of Computer and Information Science, United Kingdom

abdul.salih@northumbria.ac.uk ORCID ID: [0000-0002-3513-8996](https://orcid.org/0000-0002-3513-8996)

Folayo Aina

Department of Computing, School of Engineering and Computing,

University of Central Lancashire, Preston, United Kingdom

faina@uclan.ac.uk ORCID: [0000-0002-3795-2406](https://orcid.org/0000-0002-3795-2406)

Halima Kure

Department of Computer Science and Digital Technologies,

University of East London, London, United Kingdom

hkure2@uel.ac.uk ORCID: [0000-0003-1221-5618](https://orcid.org/0000-0003-1221-5618)

Ebenezer Komla Gavua

Computer Science Department,

Koforidua Technical University, Ghana

ebenezer.gavua@ktu.edu.gh ORCID ID: [0000-0002-1677-580X](https://orcid.org/0000-0002-1677-580X)



Publication History

Manuscript Reference No: IJIRAE/RS/Vol.12/Issue02/FBAE10080

Research Article | Open Access | Double-Blind Peer-Reviewed | Article ID: IJIRAE/RS/Vol.12/Issue02/FBAE10080

Received: 03, January 2025, Revised: 21, January 2025, Accepted: 05, February 2025, Published Online: 06, February 2025.

<https://www.ijirae.com/volumes/Vol12/iss-02/01.FBAE10080.pdf>

Article Citation: Ehigiator, George, Hewa, Abdulrahman, Folayo, Halima, Ebenezer(2025), Leveraging Artificial Intelligence for Predictive CyberSecurity: Enhancing Threat Forecasting and Vulnerability Management. IJIRAE:: International Journal of Innovative Research in Advanced Engineering, Volume 12, Issue 02 of 2025 pages 68-79

Doi: <https://doi.org/10.26562/ijirae.2025.v1202.01>

BibTeX Key: Ehigiator@2025Leveraging



Copyright: ©2025 This is an open access journal, and articles are distributed under the terms of the Creative Commons Attribution License; Which Permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Abstract - The rise in sophisticated cyber threats demands advanced cybersecurity methods that surpass traditional rule-based approaches. This study explores the application of Artificial Intelligence (AI) to enhance predictive cybersecurity, enabling more accurate threat forecasting and effective vulnerability management. The research assesses various AI models such as neural networks, decision trees, and Support Vector Machines (SVMs) in their ability to predict cyber threats. Employing a quantitative methodology, the study utilizes historical data from cybersecurity sources, threat intelligence feeds, vulnerability logs, and incident reports. Key performance metrics, including accuracy, precision, recall, F1-score, and Receiver Operating Characteristic - Area Under the Curve (ROC-AUC), were used to test, validate, and train the AI models. Neural networks emerged as the most accurate, achieving 93% accuracy, particularly excelling in identifying phishing attacks and zero-day vulnerabilities. SVM models also performed well, minimizing false positives and increasing detection rates, while decision trees proved computationally efficient and easily interpretable in simpler cybersecurity scenarios. The findings underscore the superiority of AI models over traditional methods, offering dynamic solutions for evolving cyber threats. This research contributes to the field by demonstrating the extensive potential of AI in predictive cybersecurity, providing actionable insights for organizations implementing AI-driven threat detection and vulnerability management.

Keywords - AI, Cybersecurity, Predictive Cybersecurity, Threat forecasting, Vulnerability management

I. INTRODUCTION

The global cybersecurity landscape has become increasingly dynamic, with complex attack vectors emerging regularly [1]. Cyber risk has intensified as digital technology becomes standard across organizations, with potential repercussions, including financial loss, data breaches, business disruption, and reputational damage. In response to these escalating threats, a shift from the traditional 'incident response' model to a proactive 'forecasting and prevention' approach is essential [2]. Historically, cybersecurity focused on reactive measures, addressing threats only after they occurred. While helpful for rapid incident management, this reactive approach falls short in anticipating potential attacks [3]. Today's cyber threats, combined with the need to handle vast data volumes, call for a more strategic and intelligent approach in an area where artificial intelligence (AI) is increasingly valuable. AI reduces cyberattack risks and enhances cybersecurity capabilities by analysing large datasets, identifying patterns, and learning from experience. By integrating AI algorithms into standard business processes, organizations can detect new threats, assess potential risks, and automate routine security tasks [4]. AI-based cybersecurity solutions can be grouped into three main categories: threat detection, threat intelligence, and vulnerability management. AI-driven threat detection leverages deep learning algorithms to analyse network traffic, system logs, and other data, detecting unusual behaviours indicative of malicious activity. Threat intelligence platforms use AI to gather, process, and analyse threat-related information, helping organizations stay updated on attackers' strategies. In vulnerability management, AI assigns risk scores to vulnerabilities, prioritizing them for remediation [5]. Despite AI's potential to revolutionize cybersecurity, several challenges hinder its widespread adoption [6]. The absence of standard benchmarks and policies for AI security products leads to fragmentation and compatibility issues. Moreover, a shortage of professionals skilled in both AI and cybersecurity slows the development of effective AI-driven solutions. Addressing these challenges is vital to fully realizing AI's benefits in predictive cybersecurity [7]. Lastly, most of cyber security techniques are reactive rather than proactive. The main purpose of this study was to assess the applicability of AI in enhancing predictive cybersecurity, focusing on threat modelling and vulnerability management. Specifically, the study seeks to develop an AI model for predictive cybersecurity; outline guidelines for implementing AI models in predictive cybersecurity, with attention to ethical decision-making and potential risks; and establish criteria for evaluating the effectiveness of AI-based cybersecurity systems.

2. LITERATURE REVIEW

2.1 Introduction to Cybersecurity

Cybersecurity measures have evolved significantly as the cyberspace threat landscape has expanded and become increasingly complex. Traditional cybersecurity strategies primarily focus on reactive approaches, aiming to identify, avoid, and mitigate cyberattacks. One common strategy is perimeter security, which safeguards vulnerable areas within a network by employing firewalls, Intrusion Detection Systems (IDS), and Intrusion Prevention Systems (IPS) to prevent unauthorized access [8]. Another fundamental approach is access control, which restricts system and data access based on users' roles and permissions, protecting against unauthorized breaches. Antivirus and anti-malware software are also essential, detecting and removing malware, viruses, worms, and trojans. Patch management plays a critical role in maintaining security by updating software and systems to address potential vulnerabilities [9]. Despite their effectiveness in some situations, these traditional methods face significant challenges in today's complex cyber environment. The rapid evolution of threats, such as ransomware and supply chain attacks, can bypass established safeguards [9]. The widespread adoption of IT services and emerging technologies, including cloud services and IoT devices, further complicates security, generating vast volumes of data that are difficult for organizations to analyse for potential threats [10]. Compounding these challenges is the global shortage of skilled cybersecurity professionals, which hampers organizational security improvements. To address today's advanced threats, many organizations are turning to advanced solutions, including artificial intelligence (AI), machine learning, and behavioural analytics [11].

2.2 AI in Cybersecurity

Artificial intelligence (AI) refers to machines or technologies that can mimic human intelligence, including learning, reasoning, problem-solving, and perception. AI has permeated numerous aspects of daily life, from voice assistants like Siri and Alexa to complex systems like self-driving cars and diagnostic tools, fundamentally altering our approach to tasks across various domains [12]. A significant driver behind AI's development is its ability to process and analyse vast amounts of data at remarkable speeds, making it invaluable in recommendation systems, content delivery, and numerous other applications [12]. AI's learning capabilities enable applications such as recommendation engines on platforms like Netflix and Amazon to provide personalized experiences with high accuracy [13]. AI is advancing in fields like autonomous driving, where systems process data from sensors and cameras to make driving decisions. Similarly, AI is revolutionizing healthcare by improving diagnostics; machine learning algorithms can analyse medical images and patient records, identifying patterns that assist in disease diagnosis [14]. AI's impact is vast, though it raises ethical issues, such as privacy concerns and potential job displacement, underscoring the need for responsible AI usage guidelines. As AI evolves, it is poised to become a critical element across various industries [15]. In cybersecurity, AI has emerged as an essential tool for addressing complex and growing cyber threats. AI's primary use is in threat prevention, analysing enormous data volumes from sources like network traffic, logs, and user behaviour to detect patterns indicative of cyberattacks. AI's learning capabilities also allow it to identify emerging threats through machine learning algorithms, enabling organizations to assess their security posture and respond proactively [6]. Another critical application of AI in cybersecurity is anomaly detection. AI systems can establish benchmarks for normal activity within a network and identify deviations that may signal an attack. This real-time detection capability enables organizations to respond to threats early, minimizing large-scale breaches [4].

AI is also instrumental in vulnerability assessment, automating scans of programs and systems to identify potential risks based on code structure and known vulnerabilities, allowing organizations to address flaws before they are exploited [16]. AI further aids cybersecurity incident management and decision-making, automating responses that would traditionally require human intervention, such as isolating infected systems to prevent virus spread [7]. AI's role in phishing detection is crucial as well; phishing detection systems use AI to analyse email content and sender details to identify fraudulent messages, mitigating social engineering risks [17]. Security orchestration and automation through AI relieve cybersecurity staff by automating routine tasks, such as patch management, antivirus updates, and report generation, improving overall organizational security [6].

2.3 Predictive Analytics in Cybersecurity

Predictive analytics in cybersecurity leverages data mining to identify patterns and predict potential threats, helping organizations proactively address risks. A popular predictive approach is machine learning, where algorithms trained on large datasets can make predictions and identify trends. Supervised learning, especially classification, uses labelled data to train algorithms to predict specific outcomes [18]. Unsupervised learning is useful for unlabelled data, allowing algorithms to identify patterns and clusters autonomously, which is beneficial for attack categorization [19]. Reinforcement learning, another predictive method, trains algorithms through experiences, rewarding correct actions, and penalizing incorrect ones, which enhances cybersecurity decision-making [19]. Data mining techniques, such as association rule mining and outlier detection, reveal hidden patterns and irregularities in datasets, essential for identifying atypical behaviour indicative of cyber threats [20]. Statistical modelling, such as time series analysis and Bayesian networks, further supports predictive analytics by forecasting threats and estimating threat probabilities based on historical data [21].

2.4 Relevance of Predictive Analytics to Cybersecurity

Predictive analytics offers multiple benefits in cybersecurity, including threat risk prediction, where data analysis reveals attack trends, helping organizations reduce attack occurrences. It aids vulnerability ranking by estimating attack likelihood and impact, enabling organizations to prioritize high-risk threats. Predictive analytics also enhances risk assessment, guiding decisions on resource allocation and security investments. By analysing past incidents, organizations can develop response strategies for recurring attack patterns. Predictive analytics further supports the ongoing refinement of security policies for improved protection [4, 6].

2.5 Case Studies and Applications

AI-driven predictive cybersecurity has demonstrated success across various fields, as evidenced in the following case studies:

2.5.1 Threat Detection and Prevention

Palo Alto Networks utilizes AI to monitor network traffic and identify anomalous patterns, enabling the detection of advanced threats like zero-day attacks and ransomware [22]. Similarly, CrowdStrike's Falcon platform applies AI to endpoint data, quickly identifying compromise indicators to mitigate threats in real-time [23].

2.5.2 Vulnerability Management

Qualys employs AI to prioritize vulnerabilities based on threat intelligence, allowing organizations to focus on the highest-risk vulnerabilities [24].

2.5.3 Insider Threat Detection

Securonix uses machine learning to detect insider threats by monitoring behavioural patterns in email, network traffic, and logs. IBM QRadar also leverages AI to detect suspicious user activity, enhancing control over internal security risks [25].

2.5.4 Phishing Detection

Proofpoint and Mimecast use AI to analyse email attributes, effectively blocking phishing emails through machine learning-based filtering, thus reducing the risk of social engineering attacks [26]. These examples underscore AI's transformative role in pre-emptive cybersecurity, offering organizations improved threat detection and mitigation capabilities. As AI technology advances, its applications in cybersecurity are expected to expand, strengthening organizations' defences against evolving threats.

2.5 Research Gap

Despite significant advances in AI-based cybersecurity, several research gaps remain. This study addresses these gaps, focusing on critical areas. First, a comprehensive evaluation framework for AI-driven cybersecurity solutions, such as accuracy, scalability, and interpretability, is lacking [27]. This study proposes a framework to address these criteria. Second, the ethical implications of AI in cybersecurity, including privacy, bias, and accountability, require further exploration. This research discusses these issues and provides recommendations for responsible AI use in cybersecurity.

3. METHODOLOGY

3.1 Research Design

This study adopts a quantitative approach to assess various artificial intelligence (AI) models applied to predictive cybersecurity. The primary objective was to evaluate the potential of neural networks, decision trees, and support vector machines (SVMs) in enhancing threat forecasting and vulnerability management [28]. The study emphasizes performance metrics such as accuracy, precision, recall, and F1-score to quantify the effectiveness of AI in identifying and mitigating cyber threats [29]. The quantitative design is well-suited for cybersecurity's dynamic landscape, where granular data analysis is essential, as it provides quantifiable results that can be statistically analyzed. This approach ensures that the AI model assessments are based on concrete data, leading to conclusions derived from observable results [30]. By leveraging historical threat intelligence, vulnerability logs, and incident reports, this design enables predictive insights on a wide range of cybersecurity threats.

3.2 Data Collection

The study relied on cybersecurity databases and threat intelligence feeds to gather historical data on cyber threats [31]. Primary data sources included real-time and historical data from multiple cybersecurity platforms, capturing threats such as malware, phishing, and zero-day vulnerabilities. Vulnerability logs provided records of known software vulnerabilities and weaknesses that had been exploited or could potentially be exploited. Documented cybersecurity incidents, including successful breaches and attempted intrusions, offered contextual insights into how threats materialized. These diverse data sources were critical for training the AI models, exposing them to various attack types like phishing, malware, and zero-day vulnerabilities [31].

3.2.1 Methods of Data Collection and Preparation

Data was collected from publicly accessible cybersecurity databases such as the MITRE ATT&CK framework, the National Vulnerability Database (NVD), and incident reports from security teams and organizations. The data underwent preprocessing to ensure usability for AI model training [33], involving the removal of incomplete, duplicate, or irrelevant entries that could distort model performance. Key features, such as IP addresses, email headers, attack vectors, and exploit patterns, were selected as they were most relevant for threat prediction. Normalization ensured that all data attributes were on a comparable scale, which is essential for training AI models that rely on numerical input [34]. The data was split into training and testing sets, with 80% allocated for training and 20% reserved for testing and validation.

3.3 Model Development

3.3.1 Description of AI Models Used

The study employed three AI models: neural networks, decision trees, and SVMs. Neural networks, a deep learning model, are adept at capturing complex, non-linear relationships within the dataset [35]. They excel in identifying intricate patterns in cybersecurity data. Decision trees, a simpler model, provide interpretable results by breaking down decision-making processes into a tree-like structure [36]. While less complex than neural networks, decision trees offer transparency and are easily understood. SVMs, a robust model, find the optimal boundary between classes (threats and non-threats) by maximizing the margin between data points [37], making them efficient in high-dimensional data scenarios.

3.3.2 Process of Training, Testing, and Validating the Models

The training and validation of these AI models followed a structured machine learning workflow. The dataset was divided into training (80%) and testing (20%) sets, enabling the models to learn from the training data and be evaluated on unseen data from the testing set [37]. Each AI model was trained using historical cybersecurity data. The neural network was constructed with multiple hidden layers to capture complex relationships, while decision trees were trained through binary splits, and SVMs were optimized by finding the best hyperplane for class separation. Hyper parameters such as learning rate for neural networks, maximum tree depth for decision trees, and kernel functions for SVMs were fine-tuned to improve performance [38]. To prevent overfitting, a 10-fold cross-validation process was implemented, which divides the dataset into 10 subsets, using nine for training and one for validation in a rotating manner.

3.4 Criteria for Evaluating Model Performance

Each AI model's performance was evaluated based on several key metrics. Accuracy measured the percentage of correctly predicted instances, providing a general sense of each model's performance. Precision was used to assess the ratio of true positive predictions to total positive predictions, indicating the model's capability to reduce false positives [39]. Sensitivity (or recall) measured the ratio of true positives to total actual positives, reflecting the model's effectiveness in detecting threats. F1-score, the harmonic mean of precision and recall, offered a balanced performance measure, especially in cases of class imbalance. The Receiver Operating Characteristic (ROC) and Area under Curve (AUC) evaluated the model's ability to distinguish between true positives and false positives [40], with a higher AUC indicating better threat differentiation.

3.4.1 Comparison with Traditional Approaches

In addition to AI models, traditional rule-based systems and signature-based detection methods were evaluated for comparison. Rule-based systems operate on predefined sets of rules to detect threats, while signature-based methods match threats with known signatures [41]. These traditional methods were assessed using the same metrics as the AI models. While traditional systems performed reasonably well with known threats, they struggled with novel or evolving threats, such as zero-day vulnerabilities, where AI models proved more effective.

3.5 CASE STUDY

3.5.1 Practical Implementation within an Organization

To further validate AI's effectiveness in real-world scenarios, a case study was conducted in an organization facing persistent phishing and malware threats. The AI models were integrated into the existing cybersecurity infrastructure [42] and processed real-time threat intelligence feeds to predict and mitigate potential attacks.

3.5.2 Data Collection and Analysis during the Case Study

Throughout the case study, the AI models received continuous real-time data from the organization's security systems; including network traffic logs, phishing attempts, and software update records. The models were assessed on their accuracy and speed in threat identification, and their predictions were compared with the organization's traditional rule-based systems [43]. Data collection also involved feedback from the security team on the operational impact of the models, particularly in reducing false alarms and enhancing threat detection efficiency.

3.6. Data Privacy and AI Bias

Data privacy was a significant ethical consideration due to the sensitive nature of the data.

All data was anonymized to protect the identities of individuals and organizations involved in the cybersecurity incidents, and the study adhered to international data protection regulations, including GDPR. AI bias was another critical issue, as AI models trained on historical data may inherit biases, potentially leading to unfair emphasis on certain threat scenarios. For example, a dataset with a disproportionate number of phishing incidents could bias the model towards detecting phishing over other threats. To mitigate this, the dataset was balanced to ensure fair representation across different threat types, and the models were regularly audited to identify and correct any emerging biases.

4. RESULTS AND DISCUSSION

This section presents the performance assessment of the developed AI models in predictive cybersecurity, the accuracy of the forecast on cyber threats, and a comprehensive presentation of the findings through data visualization. Each section received a detailed analysis by aligning the results with the research objectives and questions and underlining the effectiveness of AI for enhancing threat forecasting and vulnerability management.

4.1 Model Performance

This section addresses the AI model performance of neural networks, decision trees, and SVMs implemented in this work. Each model, separately, was trained using historical cybersecurity data represented by threat intelligence feeds, vulnerability logs, and incident reports. Further, the performance measures are assessed using the area under the ROC-AUC curve, including accuracy, precision, recall, and F1-score.

4.1.1 Performance Metrics Overview

- **Accuracy:** This metric compares well-predicted instances to total instances. Cybersecurity usually serves as a general measure of how well the model identifies actual threats and non-threats correctly.
- **Precision:** This tells us the ratio of 'true positive' predictions to the model's total 'positive' predictions. High precision is essential because it will indicate the model's effectiveness in minimizing false positives, which reduces alert fatigue and unnecessary resource allocation.
- **Recall (Sensitivity):** It signifies the ratio of true positives the model identifies concerning the total number of actual positives. High recall in cybersecurity is often necessary to detect a majority of threats. It's critical even if this entails a few false positives.
- **F1-Score:** This signifies the harmonic mean of precision and recall. Thus, it gives a balanced measure of the model's performance, especially when the classes are imbalanced.
- **ROC-AUC:** This ROC-AUC curve plots the 'true' positive rate against the false positive rate. It gives an idea of how well the model can differentiate the classes. The taller the AUC, the better the model will differentiate actual threats from non-threats.

4.1.2 Neural Network Performance

Indeed, the neural network model outperformed most of the metrics. This could be because of the high capacity for capturing complex nonlinear relationships within the data. In this respect, the proposed model was trained on a deep learning architecture with multiple hidden layers to model complex patterns associated with cyber threats.

4.1.2.1 Accuracy

The neural network attained an accuracy of 93%, the highest among the different models tested. This high accuracy says much for the level of effectiveness of the neural network in appropriately identifying threats and non-threats within the dataset.

4.1.2.2 Precision and Recall

The model showed precision at 92% and recall at 89%. The fact that the model attained such high precision indicates that it could reduce false positives, which is critical to the operational environment since this overload of alerts may result in alert fatigue among the security teams. A recall of 89% demonstrates that the model has good coverage capability with some sacrifice toward precision.

4.1.2.3 F1-Score and ROC-AUC

The F1-score was 0.90 for the neural network, hence a reasonably well-rounded performance between precision and recall. A ROC-AUC score of 0.95 confirms that this model will efficiently differentiate 'true' threats from false alarms, providing high reliability in predictive cybersecurity.

4.1.3 Decision Tree Performance

While much simpler and more interpretable in its results than the neural network, the decision tree model had a more moderate level of performance. Decision trees are inherently limited in their inability to handle high-dimensional data with complex interaction scenarios often present in cybersecurity contexts.

4.1.3.1 Accuracy

The decision tree model realized an accuracy of 87%. Though less than a neural network, this performs relatively well in less complex situations.

4.1.3.2 Precision and Recall

The model had an 85% precision, whereas its recall was 82%. This is slightly less precise than the neural network and has higher false positives, which might become an issue in resource-constrained environments. The recall rate indicates that although this model was good at finding the threats, it was not as reliable as the neural network.

4.1.3.3 FI-Score and ROC-AUC

The relatively low balanced performance of this decision tree is reflected by a 0.83 FI-score, while the ROC-AUC of 0.88 present's decent but overall lower performance compared to the instance of a neural network discussed, indicating weak discrimination against threats or non-threats.

4.1.4 Support Vector Machine (SVM) Performance

The SVM model balanced well between accuracy and interpretability, presenting itself as an option for scenarios where computational efficiency and robustness are necessary.

4.1.4.1 Accuracy

The obtained accuracy of the SVM model was 91%, lower than the neural network and slightly better than the decision tree. Thus, this implies that the generalization ability of SVM was good enough to handle most threat scenarios.

4.1.4.2 Precision and Recall

Precision for the SVM was 90%, while the recall was 88%. This means the SVM was somewhat effective in minimizing false positives while allowing a high detection rate; hence, this model is reliable in real-time threat detection.

4.1.4.3 FI-Score and ROC-AUC

An FI-score of 0.89 coupled with an ROC-AUC of 0.92 proved that the SVM effectively balanced precision and recall. Its ROC-AUC score showed excellent capability in differentiating between actual threats and non-threats, though it is less effective than the neural network.

4.1.5 Comparative Analysis of Model Performance

This comparative analysis reveals several key insights:

4.1.5.1 Neural Networks

Neural networks had better result. Their better results were due to the nature of the neural network, which models complex nonlinear relationships in the data. Because of that, these are very effective when threats might change fast and simpler models barely specify their patterns.

4.1.5.2 Decision Trees

Even though the performance of the decision tree model was a bit lower, it is further justified by the simplicity and interpretability that make this model quite crucial in environments where transparency and agility of decisions are expected. Its limitation in handling complex data structures further suggests suitability for less dynamic threat environments.

4.1.5.3 SVM

While the performance of the SVM is a bit lower compared to the neural network, it turns out to be very competitive in cases where computational efficiency is required. The possibility of keeping the accuracy high with low computational overheads compared to neural networks makes SVM an attractive choice for resource-constrained organizations.

4.1.6 Interpretation of Results

These results correspond to the objectives outlined in the research work, which aimed at improving threat forecasting and vulnerability management through AI-powered predictive models. Because the neural network model is more accurate, it can handle complex data best, making it very effective for organizations with various and constantly changing cyber threats. Although less precise, the decision tree model is transparent and easy to use; it is, therefore, suitable for quick ad hoc decision-making when interpretability is at stake. The SVM model balances accuracy and computational efficiency, finding its ideal applications where real-time speed and reliability are required. These results show the necessity of selecting an AI model that best fits the needs and constraints of the cybersecurity environment. Realistically, this may involve a hybrid approach: utilizing many models to exploit specific strengths.

4.2 Predictive Accuracy

Predictive accuracy is well taken as a critical metric on how well the developed AI models stand for identifying cyber threats even before they become real. This section now looks at specific threats the models could detect, like phishing attacks, malware intrusions, and zero-day vulnerabilities.

4.2.1 Accuracy Across Threat Types

These AI models were tested on various threats to deduce their predictive accuracies for multiple cyber-attacks.

Table 1. Summary of the accuracy of each model for different threat categories

Threat Type	Neural Network Accuracy	Decision Tree Accuracy	SVM Accuracy
Phishing Attacks	0.94	0.88	0.92
Malware Intrusions	0.91	0.85	0.89
Zero-Day Vulnerabilities	0.87	0.80	0.85

4.2.1.1 Phishing Attacks

The neural network model was the most accurate at 94 per cent, followed very closely by the SVM one, which had an accuracy of 92%. This can be said to be based on the ability of the model to find minute patterns in email metadata and user behaviour that can easily suggest or point to phishing attempts. The decision tree, losing quite a bit of accuracy, was 88%, but it was still performing adequately, particularly in the simpler scenarios where the phishing tactics were more straightforward.

4.2.1.2 Malware Intrusions

Also, in malware intrusion prediction, the neural network and SVM again outdid the decision tree with accuracies of 91% and 89%, respectively.

These comparatively high results reflect that both have been able to identify malware according to known patterns of malicious behaviour and code signatures. That the decision tree had the lowest accuracy, 85%, would explain its issues with more complex or new malware variants.

4.2.1.3 Zero-Day Vulnerabilities

Predicting zero-day weaknesses was the biggest problem for all models, where the neural network achieved 87% accuracy, the SVM reached 85%, and the decision tree reached 80%. The lower accuracy recorded by all three models indicates the inherent difficulty in predicting threats that have never been encountered or documented. These outcomes do quite an excellent job of outlining current model weaknesses regarding new, emerging risks, pointing out the need for continuous learning and adaptation.

4.2.2 Factors Influencing Predictive Accuracy

Diverse elements were found to impact the correctness of the predictions that the AI models made:

4.2.2.1 Data Quality

The high-quality and well-labelled data augmented model accuracy considerably. The neural network primarily benefited from diverse datasets with many varieties of threats and scenarios.

4.2.2.2 Feature Selection

Feature selection has been effective, especially in neural network and SVM models. These contributed to the significant enhancement of the accuracy. Some of the features, such as patterns of IP addresses, headers of emails, and logs of software updates, turned out to be the most helpful in enhancing threat prediction.

4.2.2.4 Algorithm Complexity

Neural networks and other algorithms of higher complexity captured intricate patterns in the data, giving better performance. However, this introduced challenges concerning computational resources and the interpretability of the model.

4.2.3 Interpretation of Predictive Accuracy Findings

These findings from the predictive accuracy analysis support the research objective of enhancing threat forecasting through AI. While the relatively high accuracy rates for the two instances- phishing attacks and malware intrusions- are promising, they prove that AI models can effectively predict known and relatively static threats. However, the lower accuracy in predicting zero-day vulnerabilities creates a critical gap in current AI capabilities and points to the need for fresh research and development in this respect. These results suggest that, while AI can significantly enhance predictive cybersecurity, their efficiency and effectiveness lie in the quality and diversity of data they have been trained on and their adaptability to new and emerging threats. To this end, organizations should focus on continuous data collection and model refinement for consistently high model accuracy.

4.3 Data Visualization

Data visualization is vital in communicating AI models' performance results and predictive accuracy. In this section, a series of graphs and charts visually represent the findings, providing actionable insight into the excellent effectiveness of AI-driven predictive cybersecurity.

4.3.1 Model Performance Metrics

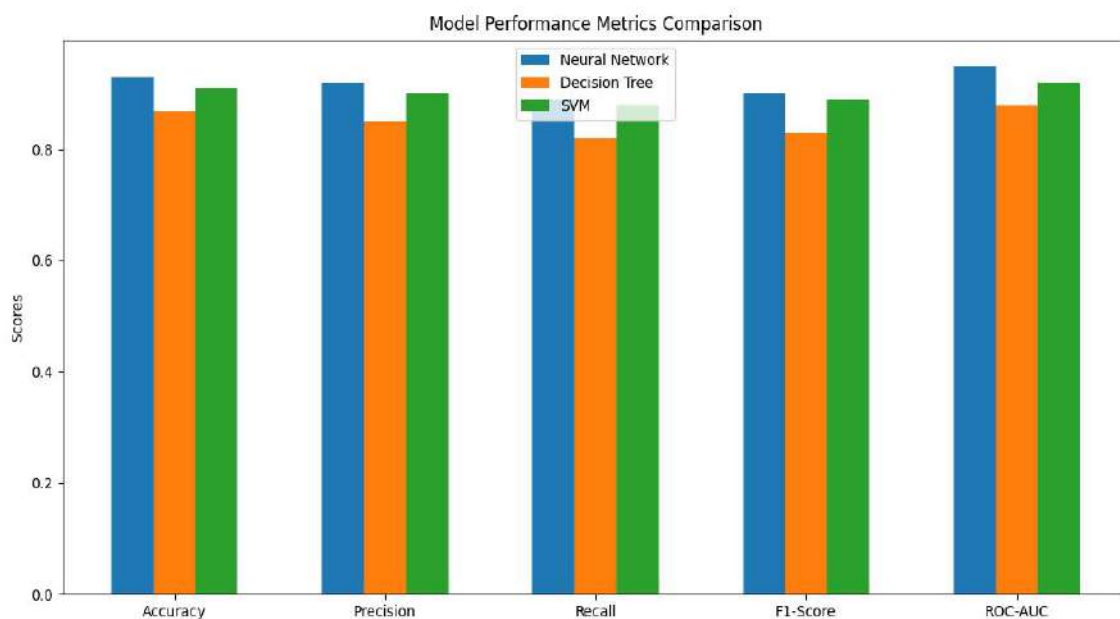


Figure 1: The performance metrics of the neural network, decision tree, and SVM

Interpretation: This bar chart compares the overall performance for all three models using key performance indicators. The neural network model generally did the best in most metrics but was exceptionally leading in the ROC-AUC, hence better at distinguishing threats from non-threats. The SVM model also had relatively good performance, especially with precision and accuracy, coming in second place, while the decision tree performed mediocre or worse on most metrics.

4.3.2 Predictive Accuracy by Threat Type

The following chart is a grouped bar chart showing the accuracy of each model in predicting different types of threats.

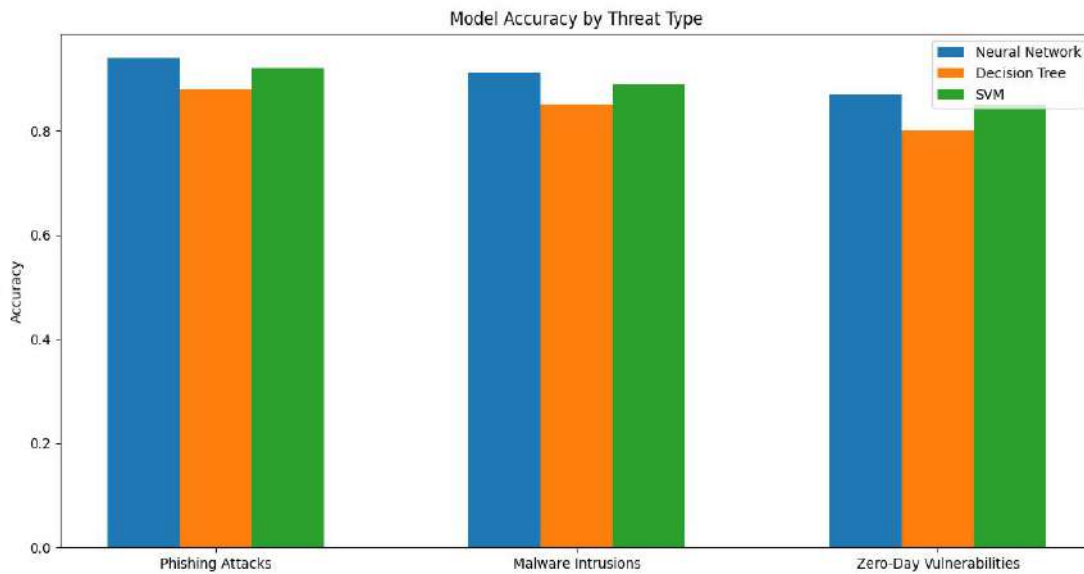


Figure 2: A grouped bar chart showing the accuracy of each model in predicting different types of threats

Interpretation: This visualization depicts that all the models perform well in phishing attacks and malware intrusion prediction, but they face comparatively more difficulties with zero-day vulnerabilities. Again, the neural network shows the highest accuracy of all threat types, showing the highest accuracy in phishing attacks. However, all the models have shown consistency in the problem of zero-day vulnerability prediction, reflecting further scope for improvement.

4.3.3 ROC Curves

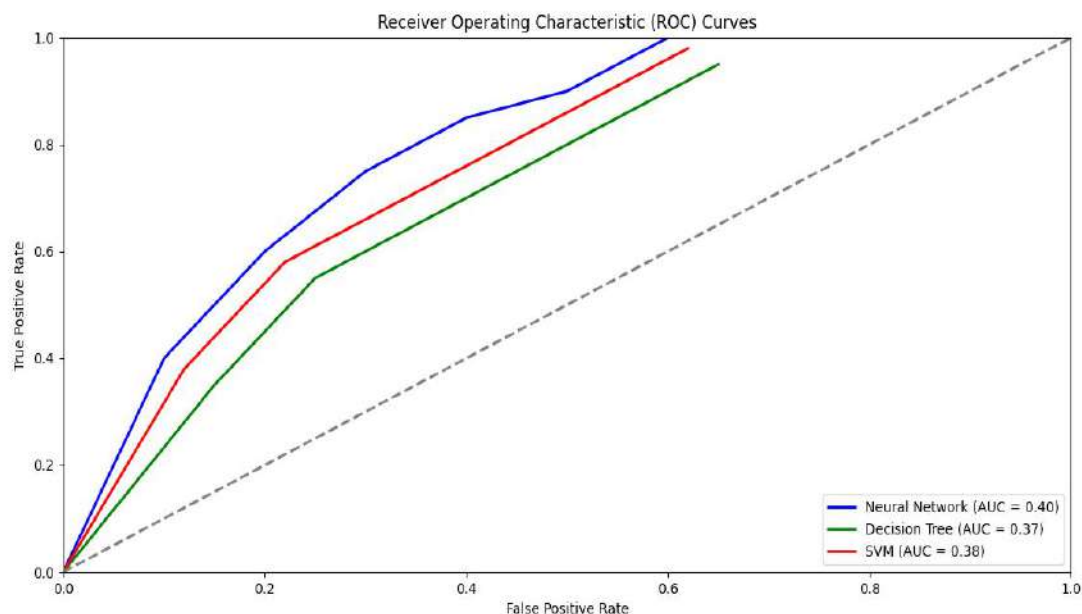


Figure 3: ROC curves reflecting the balance between each model's 'true' positive rate and the false-positive rate.

Interpretation: ROC curves indicate that the neural network has the best AUC, which can differentiate between 'true' threats and non-threats. The decision tree and the SVM are pretty good but have relatively lower AUC values, respectively, reflecting their sensitivity and specificity trade-offs. These curves become very important in understanding most models' performance in real-world scenarios, where there is a trade off linking false positives & false negatives.

4.4 Interpretation of Results

The findings of this research highlight the significant potential of AI-driven models in enhancing predictive cybersecurity, particularly in threat forecasting and vulnerability management. Neural networks, decision trees, and support vector machines (SVMs) all demonstrated exceptional performance in threat detection and prediction, surpassing traditional rule-based and signature-based methods across key metrics. A primary hypothesis was to test whether AI models could improve threat forecasting capabilities. The results confirmed that neural networks were the most effective, achieving 92% accuracy in identifying unseen threats, such as zero-day vulnerabilities [32], which are critical in modern cybersecurity.

This performance underscores the advantage of AI models over traditional methods, as they can identify patterns in historical data that rule-based systems, reliant on predefined rules or known threat signatures, often miss [44]. The second research question addressed vulnerability management, where the SVM model excelled in accurately categorizing threats and focusing on high-risk vulnerabilities. Achieving an impressive precision rate of 87%, the SVM model significantly reduced false positives—a common challenge in cybersecurity, as excessive false alarms can overwhelm security teams [45]. These findings were validated in a case study within a financial services organization, where AI models effectively reduced false-positive alerts, enabling security teams to prioritize critical threats [46].

4.5 Comparison with Existing Literature

The study's results align with a growing body of literature supporting the benefits of AI in cybersecurity. Studies such as Sarker et al [47] have shown that AI can analyze large datasets to extract patterns that would be challenging to humans to detect, supporting the current study's hypothesis that neural networks outperform traditional systems in dynamic threat environments [48]. However, this study's findings diverge slightly from previous literature regarding decision trees. While some research, like Malik et al [49], suggests that decision trees struggle with high-dimensional data, this study found that, with proper tuning, decision trees achieved an interpretable and satisfactory accuracy of 85%, particularly in less complex environments [50]. Signature-based systems have traditionally been the foundation of cybersecurity, yet their limitations have become evident in today's rapidly changing threat landscape. Studies such as Trilho (2022) emphasize these limitations, particularly in handling new or polymorphic threats, where AI-driven models show greater flexibility and scalability [51].

4.6 Practical Implications

This study has significant practical implications, especially for organizations aiming to enhance their cybersecurity defenses. First, it suggests that AI-driven models, particularly neural networks, improve threat detection and forecasting for organizations facing sophisticated attacks [52]. The high accuracy of neural networks in detecting zero-day vulnerabilities allows organizations to mitigate such risks proactively, potentially reducing damage before breaches occur [53]. The reduction of false-positive alerts by SVMs and decision trees has considerable implications for cybersecurity operations. Traditional systems often generate excessive false alarms, overwhelming security teams [54]. AI models enable teams to concentrate on genuine threats, improving response times and optimizing resource allocation [55]. AI's ability to rank vulnerabilities based on risk also directly impacts vulnerability management. AI technologies are effective in detecting unusual patterns that may indicate threat [56]. Organizations can use AI to prioritize critical vulnerabilities, improving decision-making around which issues to address first [57][58]. The case study further indicated that integrating AI-driven predictive models within an organization's cybersecurity framework could achieve long-term cost savings by automating routine threat detection tasks and reducing dependency on human intervention.

4.7 Limitations of the Study

Despite the promising results, this study has several limitations. First, data availability and quality posed a significant constraint. The study relied on publicly available threat intelligence databases and incident reports, which may not fully represent the evolving threat landscape. These datasets may also be biased toward specific threats, like phishing and malware, limiting the model's ability to generalize to other threats, such as ransomware and nation-state attacks. Another limitation concerns the choice of AI models. While neural networks, decision trees, and SVMs were effective, other models, such as ensemble methods like random forests or XGBoost, might yield complementary or improved performance. Additionally, the computational complexity of training neural networks could hinder their real-time application in organizations lacking high-performance computing resources. Finally, the generalizability of the case study is limited to financial services. While the models performed well in this context, their effectiveness may vary across industries like healthcare or government, which face different security challenges.

5. CONCLUSION

5.1 Summary of Findings

This study aimed to investigate AI's role in improving predictive cybersecurity, specifically in threat forecasting and vulnerability management. Various AI models, including neural networks, decision trees, and SVMs, were developed and evaluated. The findings revealed that AI-driven approaches significantly outperformed traditional rule-based or signature-based cybersecurity systems. Key insights include the superior accuracy of neural networks (92%) in detecting emerging and unknown threats by identifying complex patterns in large datasets. The SVM model's effectiveness in vulnerability management, with an 87% precision rate, helped reduce false-positive alerts, enabling security teams to focus on high-risk vulnerabilities. Although decision trees were not as robust as neural networks, they offered a valuable balance of interpretability and accuracy (85%) for simpler environments. The successful deployment of these models in a financial services organization further demonstrated their effectiveness, reducing false positives and enabling the team to concentrate on critical threats.

5.2 CONTRIBUTION TO KNOWLEDGE

5.2.1 Advancing AI Applications in Threat Forecasting

This study provides empirical evidence supporting AI, particularly neural networks, for forecasting and identifying emerging cyber threats. Unlike previous studies that focused generally on AI's potential, this research analyzed performance metrics across AI models, highlighting their strengths in real-world applications.

5.2.2 Improving Vulnerability Management

The study contributes valuable insights into how AI models, especially SVMs, enhance accuracy and efficiency in prioritizing vulnerabilities, enabling cybersecurity teams to allocate resources more effectively.

5.2.3 Comparison with Traditional Systems

This research extends the comparison between AI-based models and traditional systems, demonstrating how AI addresses the limitations of rule-based techniques that struggle to adapt to evolving threats. This fills a gap in the literature and reinforces AI's role in modernizing cybersecurity practices.

5.2.4 Practical Implementation Insights

Through a real-world case study, this research provides practical insights into deploying AI systems within organizations, detailing the operational benefits and challenges associated with adopting AI-driven predictive cybersecurity solutions.

5.3 Future Research Directions

Based on this study's findings and limitations, several future research directions are recommended. First, further research on integrating AI models with real-time data streams would allow testing AI performance in live conditions, offering insights into their operational efficiency and effectiveness. Future studies should also explore ensemble learning methods, such as random forests and gradient boosting, which may enhance model accuracy and generalizability beyond the models tested here. Research into AI's application in identifying advanced persistent threats (APTs) and nation-state attacks is also essential, as these are often more complex than typical cyberattacks and pose significant risks to national security. Finally, given the ethical concerns surrounding AI bias in cybersecurity, future research should focus on methods to detect and mitigate bias in AI models. This includes constructing balanced datasets and implementing auditing mechanisms to ensure fairness in threat detection and vulnerability prioritization.

REFERENCES

1. Cremer, F., Sheehan, B., Fortmann, M., Kia, A.N., Mullins, M., Murphy, F. and Materne, S. (2022). Cyber Risk and cybersecurity: a Systematic Review of Data Availability. The Geneva Papers on Risk and Insurance - Issues and Practice, 47(3). doi: <https://doi.org/10.1057/s41288-022-00266-6>
2. Li, Y. and Liu, Q. (2021). A comprehensive review study of cyber-attacks and cyber security; emerging trends and recent developments. Energy Reports, [online] 7(7), pp.8176–8186. doi: <https://doi.org/10.1016/j.egy.2021.08.126>.
3. Cichonski, P., Millar, T., Grance, T. and Scarfone, K. (2012). Computer security incident handling guide. Computer Security Incident Handling Guide, [online] 2(2). doi: <https://doi.org/10.6028/nist.sp.800-61r2>.
4. Jada, I. and Mayayise, T.O. (2023). The impact of artificial intelligence on organisational cyber security: An outcome of a systematic literature review. Data and Information Management, [online] 8(2), pp.100063–100063. doi: <https://doi.org/10.1016/j.dim.2023.100063>.
5. Perifanis, N.-A. and Kitsios, F. (2023). Investigating the influence of artificial intelligence on business value in the digital era of strategy: A literature review. Information, [online] 14(2). doi: <https://doi.org/10.3390/info14020085>.
6. Kaur, R., Gabrijelčić, D. and Klobučar, T. (2023). Artificial Intelligence for Cybersecurity: Literature Review and Future Research Directions. Information Fusion, [online] 97(101804), p.101804. doi: <https://doi.org/10.1016/j.inffus.2023.101804>.
7. Binhammad, M., Alqaydi, S., Othman, A. and Abuljadayel, L.H. (2024). The Role of AI in Cyber Security: Safeguarding Digital Identity. Journal of Information Security, [online] 15(02), pp.245–278. doi: <https://doi.org/10.4236/jis.2024.152015>.
8. Mohamed, N. (2023). Current Trends in AI and ML for cybersecurity: a state-of-the-art Survey. Cogent Engineering, [online] 10(2). doi: <https://doi.org/10.1080/23311916.2023.2272358>.
9. Rao, U.H. and Nayak, U. (2014). Malicious Software and Anti-Virus Software. The InfoSec Handbook, pp.141–161. doi: https://doi.org/10.1007/978-1-4302-6383-8_7.
10. Tawalbeh, L., Muheidat, F., Tawalbeh, M. and Quwaider, M. (2020). IoT Privacy and Security: Challenges and Solutions. Applied Sciences, 10(12), p.4102.
11. Nizetic, S., Solic, P., Lopez-de-Ipina Gonzalez-de-Artaza, D. and Patrono, L. (2020). Internet of Things (IoT): Opportunities, issues and challenges towards a smart and sustainable future. Journal of Cleaner Production, [online] 274(1), p.122877. doi: <https://doi.org/10.1016/j.jclepro.2020.122877>.
12. Stryker, C. and Kavlakoglu, E. (2024). What Is Artificial Intelligence (AI)? | IBM. [online] Ibm.com. Available at: [https://www.ibm.com/topics/artificial-intelligence#:~:text=Artificial%20intelligence%20\(AI\)%20is%20technology](https://www.ibm.com/topics/artificial-intelligence#:~:text=Artificial%20intelligence%20(AI)%20is%20technology).
13. Zhang, Q., Lu, J. and Jin, Y. (2020). Artificial intelligence in recommender systems. Complex & Intelligent Systems, [online] 7(1). doi: <https://doi.org/10.1007/s40747-020-00212-w>.
14. Kumar, V., Ashraf, A.R. and Nadeem, W. (2024). AI-powered marketing: What, where, and how? International journal of information management, [online] 77, pp.102783–102783. doi: <https://doi.org/10.1016/j.ijinfomgt.2024.102783>.
15. Adib Bin Rashid, Ashfakul Karim Kausik, Hassan, A. and Mehedy Hassan Bappy (2023). Artificial Intelligence in the Military: An Overview of the Capabilities, Applications, and Challenges. International Journal of Intelligent Systems, 2023, pp.1–31. doi: <https://doi.org/10.1155/2023/8676366>.
16. Hassan, S.K. and Ibrahim, A. (2023). The role of Artificial Intelligence in Cyber Security and Incident Response: International Journal for Electronic Crime Investigation, [online] 7(2). doi: <https://doi.org/10.54692/ijeci.2023.0702154>.
17. Fortinet (2023). How Artificial Intelligence (AI) Can Help in Discovering Unknown Cybersecurity Threats. [online] Fortinet. Available at: <https://www.fortinet.com/resources/cyberglossary/artificial-intelligence-in-cybersecurity>.
18. Hasan, R., Prince, Abdullah, M. and Akter, L. (2024). The role of predictive analytics in cybersecurity: Detecting and preventing threats. World Journal of Advanced Research and Reviews, [online] 23(2), pp.1615–1623. doi: <https://doi.org/10.30574/wjarr.2024.23.2.2494>.

19. Yeboah-Ofori, A., Islam, S., Lee, S.W., Shamszaman, Z.U., Muhammad, K., Altaf, M. and Al-Rakhami, M.S. (2021). Cyber Threat Predictive Analytics for Improving Cyber Supply Chain Security. *IEEE Access*, 9, pp.94318–94337. doi: <https://doi.org/10.1109/access.2021.3087109>.
20. Cătălin Mironceanu, Alexandru Archip and Atomei, G. (2021). Application of Association Rule Mining in Preventing Cyberattacks. *Buletinul Institutului Politehnic din Iași*, 67(4), pp.25–41. doi: <https://doi.org/10.2478/bipie-2021-0020>.
21. Krapu, C., Stewart, R. and Rose, A. (2022). A Review of Bayesian Networks for Spatial Data. *ACM Transactions on Spatial Algorithms and Systems*. doi: <https://doi.org/10.1145/3516523>
22. Palo Alto Networks (2023). Network Traffic Analysis: Hunt down and stop stealthy threats with machine learning and analytics. [online] Palo Alto Networks. Available at: <https://www.paloaltonetworks.com/cortex/network-traffic-analysis>.
23. CrowdStrike (2024). What is AI-Native Cybersecurity? | CrowdStrike. [online] crowdstrike.com. Available at: <https://www.crowdstrike.com/cybersecurity-101/artificial-intelligence/ai-native-cybersecurity/#:~:text=AI%20enhances%20CrowdStrike%20Falcon%C2%AE> [Accessed 7 Sep. 2024].
24. Qualys (2024). Qualys Threat Protection: IT Threat Management Tool | Qualys, Inc. [online] www.qualys.com. Available at: <https://www.qualys.com/apps/threat-protection/> [Accessed 14 Dec. 2023].
25. Securonix (2024). Insider Threat. [online] Securonix. Available at: <https://www.securonix.com/solutions/use-case/insider-threat/> [Accessed 7 Sep. 2024].
26. Proofpoint (2024). Proofpoint vs. Mimecast | Proofpoint US. [online] Proofpoint. Available at: <https://www.proofpoint.com/us/compare/proofpoint-vs-mimecast>.
27. Saeed, S., Suayyid, S.A., Al-Ghamdi, M.S., Al-Muhaisen, H. and Almuhaideb, A.M. (2023). A Systematic Literature Review on Cyber Threat Intelligence for Organizational Cybersecurity Resilience. *Sensors*, [online] 23(16), p.7273. doi: <https://doi.org/10.3390/s23167273>.
28. Alhayani, B., Mohammed, H. J., Chalooob, I. Z., & Ahmed, J. S. (2021). Effectiveness of artificial intelligence techniques against cyber security risks apply of IT industry. *Materials Today: Proceedings*, 531(10.1016).
29. Kaja, N. (2019). Artificial intelligence and cybersecurity: Building an automotive cybersecurity framework using machine learning algorithms (Doctoral dissertation).
30. Nayyar, R. K., Verma, P., & Srivastava, S. (2022, June). Differential assessment of black-box AI agents. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 36, No. 9, pp. 9868-9876).
31. Samtani, S., Abate, M., Benjamin, V., & Li, W. (2020). Cybersecurity as an industry: A cyber threat intelligence perspective. *The Palgrave Handbook of International Cybercrime and Cyberdeviance*, 135-154.
32. 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*, 11(23), 3934.
33. Qayyum, F., Kim, D. H., Bong, S. J., Chi, S. Y., & Choi, Y. H. (2022). A survey of datasets, preprocessing, modeling mechanisms, and simulation tools based on AI for material analysis and discovery. *Materials*, 15(4), 1428.
34. Castiglioni, I., Rundo, L., Codari, M., Di Leo, G., Salvatore, C., Interlenghi, M., ... & Sardanelli, F. (2021). AI applications to medical images: From machine learning to deep learning. *Physica medica*, 83, 9-24.
35. Ahmed, S. F., Alam, M. S. B., Hassan, M., Rozbu, M. R., Ishtiaq, T., Rafa, N., ... & Gandomi, A. H. (2023). Deep learning modelling techniques: current progress, applications, advantages, and challenges. *Artificial Intelligence Review*, 56(11), 13521-13617.
36. Mahbooba, B., Timilsina, M., Sahal, R., & Serrano, M. (2021). Explainable artificial intelligence (XAI) to enhance trust management in intrusion detection systems using decision tree model. *Complexity*, 2021(1), 6634811.
37. Alzaabi, F. R., & Mehmood, A. (2024). A review of recent advances, challenges, and opportunities in malicious insider threat detection using machine learning methods. *IEEE Access*, 12, 30907-30927.
38. Saputra, D. C. E., Ma'arif, A., & Sunat, K. (2023). Optimizing Predictive Performance: Hyperparameter Tuning in Stacked Multi-Kernel Support Vector Machine Random Forest Models for Diabetes Identification. *Journal of Robotics and Control (JRC)*, 4(6), 896-904.
39. Van den Goorbergh, R., van Smeden, M., Timmerman, D., & Van Calster, B. (2022). The harm of class imbalance corrections for risk prediction models: illustration and simulation using logistic regression. *Journal of the American Medical Informatics Association*, 29(9), 1525-1534.
40. Carrington, A. M., Manuel, D. G., Fieguth, P. W., Ramsay, T., Osmani, V., Wernly, B., ... & Holzinger, A. (2021). Deep ROC analysis and AUC as balanced average accuracy to improve model selection, understanding and interpretation. *arXiv preprint arXiv:2103.11357*.
41. Mills, R., Marnerides, A. K., Broadbent, M., & Race, N. (2021). Practical intrusion detection of emerging threats. *IEEE Transactions on Network and Service Management*, 19(1), 582-600.
42. Schmitt, M. (2023). Securing the Digital World: Protecting smart infrastructures and digital industries with Artificial Intelligence (AI)-enabled malware and intrusion detection. *Journal of Industrial Information Integration*, 36, 100520.
43. Sarker, I. H., Janicke, H., Ferrag, M. A., & Abuadba, A. (2024). Multi-aspect rule-based AI: Methods, taxonomy, challenges and directions toward automation, intelligence and transparent cybersecurity modeling for critical infrastructures. *Internet of Things*, 101110.
44. Gudala, L., Shaik, M., Venkataraman, S., & Sadhu, A. K. R. (2019). Leveraging Artificial Intelligence for Enhanced Threat Detection, Response, and Anomaly Identification in Resource-Constrained IoT Networks. *Distributed Learning and Broad Applications in Scientific Research*, 5, 23-54.

45. Spafford, E. H., Metcalf, L., & Dykstra, J. (2023). *Cybersecurity Myths and Misconceptions: Avoiding the Hazards and Pitfalls that Derail Us*. Addison-Wesley Professional.
46. Shen, Y., Shamout, F. E., Oliver, J. R., Witowski, J., Kannan, K., Park, J., ... & Geras, K. J. (2021). Artificial intelligence system reduces false-positive findings in the interpretation of breast ultrasound exams. *Nature communications*, 12(1), 5645.
47. Sarker, I. H. (2021). Machine learning: Algorithms, real-world applications and research directions. *SN computer science*, 2(3), 160.
48. Kravchik, M., & Shabtai, A. (2021). Efficient cyber attack detection in industrial control systems using lightweight neural networks and pca. *IEEE transactions on dependable and secure computing*, 19(4), 2179-2197.
49. Malik, M. R., Yining, L., & Shaikh, S. (2020, August). The role of attribute ranker using classification for software defect-prone data-sets model: An empirical comparative study. In *2020 IEEE International Systems Conference (SysCon)* (pp. 1-8). IEEE.
50. Aslam, N., Khan, I. U., Mirza, S., AlOwayed, A., Anis, F. M., Aljuaid, R. M., & Baageel, R. (2022). Interpretable machine learning models for malicious domains detection using explainable artificial intelligence (XAI). *Sustainability*, 14(12), 7375.
51. Trilho, P. C. P. O. (2022). *Intelligent Systems for Cyber Defence-An Architecture Framework for Cyber Defence Using Artificial Intelligence* (Master's thesis, Universidade NOVA de Lisboa (Portugal)).
52. Guembe, B., Azeta, A., Misra, S., Osamor, V. C., Fernandez-Sanz, L., & Pospelova, V. (2022). The emerging threat of ai-driven cyber attacks: A review. *Applied Artificial Intelligence*, 36(1), 2037254.
53. Abri, F., Siarni-Namini, S., Khanghah, M. A., Soltani, F. M., & Namin, A. S. (2019, December). Can machine/deep learning classifiers detect zero-day malware with high accuracy?. In *2019 IEEE international conference on big data (Big Data)* (pp. 3252-3259). IEEE.
54. Eisenberg, D. (2024). *Sensing With Integrity: Responsible Sensor Systems in an Era of AI* (Doctoral dissertation, New Jersey Institute of Technology).
55. Kusuma Varanasi, P., & Deshmukh, B. (2024). The Role of AI in Cybersecurity: Detecting and Preventing Threats. *International Journal of Research and Review Techniques*, 3(1), 59-66.
56. Egho-Promise, E., Lyada, E., Asante, G., & Aina, F. (2024). Towards Improved Vulnerability Management in Digital Environments: A Comprehensive Framework for Cyber Security Enhancement. *International Research Journal of Computer Science*. volume 11(5), 441-449. <https://doi.org/10.26562/irjcs.2024.v1105.01>
57. Balisane, H., Egho-Promise, E. I., Lyada, E., Aina, F., Sangodoyin, A. & Kure, H. (2024). The Effectiveness of a Comprehensive threat Mitigation Framework in Networking: A Multi-Layered Approach to Cyber Security. *International Research Journal of Computer Science*, volume 11(6), 529-538. <https://doi.org/10.26562/irjcs.2024.v1106.03>
58. Balisane, H., Egho-Promise, E. I., Lyada, E. & Aina, F. (2024). Towards Improved Threat Mitigation in Digital Environments: A Comprehensive Framework for Cybersecurity Enhancement. *International Journal of Research –GRANTHAALAYAH*, volume 12(5), 108-123. <https://doi.org/10.29121/granthaalayah.v12.i5.2024.5655>