

CYBERSECURITY IN HEALTHCARE: AI AND CLOUD ADOPTION

By

Kamaljit Bawa

DISSERTATION

Presented to the Swiss School of Business and Management Geneva

In Partial Fulfillment

Of the Requirements

For the Degree

DOCTOR OF BUSINESS ADMINISTRATION

SWISS SCHOOL OF BUSINESS AND MANAGEMENT GENEVA

September, 2025

CYBERSECURITY IN HEALTHCARE: AI AND CLOUD ADOPTION

By

Kamaljit Bawa

Supervised by

Hemant Palivela

APPROVED BY
dr. Jaka Vадnjal



Dissertation chair

RECEIVED/APPROVED BY:



Admissions Director

Dedication

To my dearest parents: KS BAWA and JK BAWA

Acknowledgements

I must thank the participants (ICMR and Google employees) in the research for their time and patience with me for questionnaires/surveys. Without their sincere support, I could never finish the research work.

ABSTRACT

CYBERSECURITY IN HEALTHCARE: AI AND CLOUD ADOPTION

Kamaljit Bawa
2025

Dissertation Chair: Anna Provodnikova
Co-Chair: Jaka Vadnjal

The healthcare sector faces increasing cybersecurity threats, which traditional, reactive security measures cannot effectively handle. These threats pose risks to patient safety, data privacy, and operational continuity. This research addresses this issue by developing a framework that integrates Artificial Intelligence (AI) and cloud platforms to enable proactive, real-time threat detection and response while adhering to ethical standards in data privacy and security.

The study employs an explanatory sequential mixed-methods research design. The quantitative phase consists of an experimental evaluation of four AI-based anomaly detection models (Isolation Forest, Autoencoder, LSTM Autoencoder, and Transformer Autoencoder) on the UNSW-NB15 benchmark cybersecurity dataset. The qualitative phase involves a survey of 25 senior-level cybersecurity and IT professionals to gather insights on the practical challenges, strategic considerations, and best practices for implementing AI technologies in healthcare cybersecurity.

The quantitative results revealed significant performance variations among the models, with a critical trade-off between precision and recall. The Autoencoder model achieved high precision (94.25%) but low recall (38.48%), highlighting the challenge of balancing false positives and false negatives. The qualitative results indicated that the primary barriers to AI adoption are organizational and resource-based rather than technological. Key challenges include cost constraints (88%), integration with legacy systems (84%), and a lack of skilled professionals (80%). Experts emphasized the importance of a strategic approach for AI implementation, including foundational security and a human-in-the-loop approach.

While advanced AI models, especially Transformers, hold significant potential for enhancing cybersecurity, their successful implementation requires a strategic, human-centric approach. The research's primary contribution is the Proactive, Adaptive, and Resilient (PAR) Cybersecurity Framework, a model that combines AI-driven detection with strategic principles to help healthcare organizations build cybersecurity programs that are both technologically advanced and aligned with the mission of patient safety, while ensuring ethical data privacy standards.

TABLE OF CONTENTS

List of Tables	ix
List of Figures	x
CHAPTER I: INTRODUCTION.....	1
1.1 Background of the Study	1
1.2 Problem Statement	5
1.3 Limitations, Delimitations, and Assumptions	8
1.4 Significance of the Study	10
1.5 Research Questions and Objectives	11
1.6 Definition of Terms.....	13
CHAPTER II: REVIEW OF LITERATURE	15
2.1 Introduction.....	15
2.2 Inclusion Criteria	17
2.3 The Evolving Cybersecurity Threat Landscape in Healthcare	19
2.4 The Application of Artificial Intelligence in Threat Detection	22
2.5 The Convergence of AI and Cloud Platforms for Security.....	35
2.6 Regulatory, Risk, and Governance Synthesis for AI-Driven Cybersecurity in Healthcare.....	37
2.7 Open Debates	38
2.8 Integrating Ethical AI in Healthcare Cybersecurity.....	41
2.9 Summary	43
CHAPTER III: METHODOLOGY	45
3.1 Introduction.....	45
3.2 Research Design.....	46
3.3 Quantitative Methodology	48
3.4 Qualitative Methodology	56
3.5 Reproducibility, Environment Setup and Hyperparameters	59
3.6 Limitations of the Methodology	62
3.7 Ethical Considerations	62
3.8 Summary	63
CHAPTER IV: RESULTS.....	64
4.1 Introduction.....	64
4.2 Performance of the Isolation Forest Model	65
4.3 Performance of the Autoencoder Model	69

4.4 Performance of the LSTM Autoencoder Model	74
4.5 Performance of the Transformer Autoencoder Model.....	79
4.6 Summary of Quantitative Findings	82
4.7 Instrumentation	93
4.8 Respondent Demographics and Profile.....	93
4.9 Perceptions of the Healthcare Cybersecurity Landscape.....	96
4.10 Adoption and Perceived Effectiveness of AI and Cloud Solutions	98
4.11 Implementation Challenges and Best Practices: A Thematic Analysis.....	100
4.12 Summary of Qualitative Findings	101
CHAPTER V: DISCUSSION.....	103
5.1 Introduction	103
5.2 Summary of the Study and Findings.....	104
5.3 Discussion and Interpretation of Findings	106
5.4 Implications and Applications: The Proposed Framework.....	111
5.5 Recommendations for Future Research	118
5.6 Limitations of the Study.....	119
5.7 Conclusion	122
REFERENCES	124

LIST OF TABLES

Table 1 Reported Performance of Autoencoder Models on CIC-IDS-2017.....	24
Table 2 Reported Performance of Isolation Forest on CIC-IDS-2017	27
Table 3 Reported Performance of LSTM Models on CIC-IDS-2017.....	29
Table 4 Reported Performance of Transformer Models on CIC-IDS-2017	32
Table 5 Mapping Table	33
Table 6 Autoencoder Model Hyperparameters	60
Table 7 LSTM Autoencoder Model Hyperparameters	60
Table 8 Transformer Autoencoder Model Hyperparameters	61
Table 9 Confusion Matrix for Isolation Forest	66
Table 10 Performance Metrics for Isolation Forest	67
Table 11 Confusion Matrix for Autoencoder (95th Percentile Threshold).....	69
Table 12 Performance Metrics for Autoencoder	70
Table 13 Comparative Summary of AI Model Performance	91
Table 14 Distribution of Respondent Roles.....	94
Table 15 Years of Professional Experience in Cybersecurity or IT	95
Table 16 Geographic Distribution of Respondents.....	95
Table 17 Perceived Severity of Cybersecurity Threats.....	96
Table 18 Current Use of AI-Driven Cybersecurity Solutions	98
Table 19 Use of Cloud Platforms for Cybersecurity Management.....	98
Table 20 Perceived Effectiveness of AI and Cloud Solutions	99

LIST OF FIGURES

Figure 1 Average Cost of Cybersecurity Incidents	3
Figure 2 Distribution of types of Cyber Threats in Healthcare	4
Figure 3 Growth of healthcare data breaches reported to the HHS OCR (2015–2024). Data adapted from HIPAA Journal (2025) and Fox Group (2025).	5
Figure 4 Common sources of vulnerabilities in healthcare systems, based on aggregated reports including Censinet (2025), highlighting outdated software, weak access controls, unencrypted data storage, vendor risks, human error, and gaps in employee training.	7
Figure 5 Literature of Cybersecurity.....	16
Figure 6 Isolation Forest - Confusion Matrix	66
Figure 7 Autoencoder - Confusion Matrix.....	70
Figure 8 Autoencoder - Confusion Matrix (Lowered Threshold)	72
Figure 9 Histogram of LSTM Autoencoder Reconstruction Error (MAE)	75
Figure 10 LSTM- Confusion Matrix.....	76
Figure 11 Transformer Reconstruction Error with 95th Percentile Anomaly Threshold	80
Figure 12 Transformer Reconstruction Error with 90th Percentile Anomaly Threshold	80
Figure 13 Model wise Precision, Recall and F1 Score.....	82
Figure 14 ROC/PR Curve and Calibration Plot.....	84
Figure 15 Cross Data Generalization	86
Figure 16 Ablation Study Results	88
Figure 17 F1 Scores at 95% Confidence Interval	90
Figure 18 Years of Professional Experience	95
Figure 19 Perceived Severity of Cybersecurity Threats	97
Figure 20 Perceived Effectiveness of AI-Driven Cybersecurity Solutions	99
Figure 21 Perceived Effectiveness of Cloud-Based Solutions	100
Figure 22 Practical Challenges	101
Figure 23 PAR Framework	112

CHAPTER I: INTRODUCTION

1.1 Background of the Study

The global healthcare sector is in the midst of a profound and irreversible digital transformation, a paradigm shift that has fundamentally reshaped the delivery of patient care, the management of clinical operations, and the landscape of medical research. The widespread adoption of Electronic Health Records (EHRs) has replaced paper-based systems with centralized, accessible digital repositories of patient information, promising greater efficiency and fewer medical errors (Smith et al., 2020). Concurrently, the proliferation of the Internet of Medical Things (IoMT) has connected a vast array of devices—from patient-worn vital sign monitors and smart infusion pumps to complex diagnostic imaging equipment like MRI and CT scanners—to hospital networks, enabling real-time data collection and remote patient management (Jones et al., 2021). This hyper-connectivity, further accelerated by the global demand for telemedicine and virtual care models in the wake of the recent pandemic, has created a vast, decentralized, and data-rich digital ecosystem (Lee & Park, 2022). While the benefits of this transformation are undeniable, leading to improved diagnostic accuracy, personalized treatment plans, and greater patient engagement, this evolution has simultaneously and inadvertently created an expansive and attractive attack surface for malicious cyber actors (Portela et al., 2023).

The healthcare industry's increasing reliance on this digital infrastructure has rendered it acutely vulnerable to a new and escalating wave of sophisticated cyberattacks. The statistics tracking this trend are alarming and paint a clear picture of a sector under siege. A 2025 industry report indicated that a staggering 92% of healthcare organizations experienced at least one significant cyber intrusion in the preceding year, a notable

increase from 88% in the year prior (He et al., 2021). This is not a fleeting trend but a sustained and intensifying pattern of targeted attacks. The consequences of these intrusions extend far beyond the IT department, permeating every aspect of the healthcare delivery value chain. The average cost of a single cybersecurity compromise in healthcare has soared into the millions of dollars, comprising a complex web of direct and indirect expenses, including system remediation, regulatory fines for non-compliance, legal fees, and the significant cost of operational downtime (He et al., 2021).

The financial ramifications are substantial and multifaceted. The average cost of a single cybersecurity compromise in healthcare has soared into the millions of dollars. These costs are not monolithic; they comprise a complex web of direct and indirect expenses, including the costs of system remediation and recovery, regulatory fines for non-compliance with data protection mandates like the Health Insurance Portability and Accountability Act (HIPAA), legal fees from patient lawsuits, the provision of credit monitoring services for affected individuals, and the significant cost of operational downtime. These are funds which are invariably diverted from the core mission of healthcare: patient care, medical research, and crucial infrastructure upgrades (Portela et al., 2023).

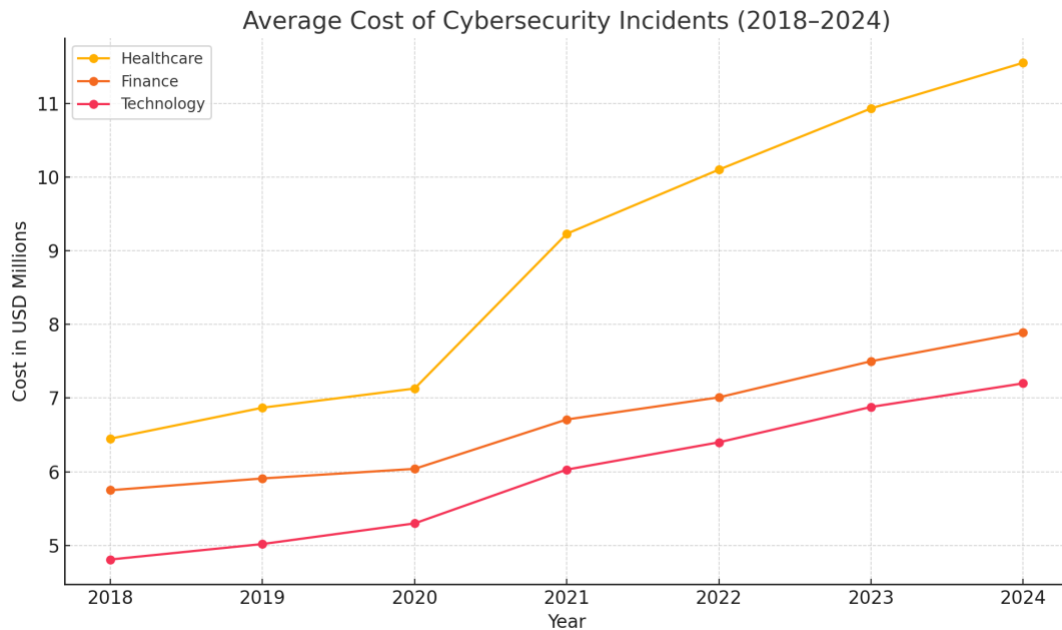


Figure 1 Average Cost of Cybersecurity Incidents

Healthcare Cybersecurity Insights. (2025). Healthcare Cybersecurity Annual Report 2025. Healthcare Cybersecurity Insights.

More alarmingly, the operational impact of these attacks can be catastrophic, posing a direct threat to patient safety. Cyberattacks, such as the infamous WannaCry ransomware attack that crippled hospitals and clinics globally, have been shown to cause significant and prolonged disruptions in the delivery of care. These disruptions manifest as the mass cancellation of appointments and elective surgeries, the shutdown of critical diagnostic equipment, delays in the delivery of time-sensitive medical procedures like chemotherapy, and a forced reversion to inefficient and error-prone paper-based systems for which modern clinical staff may be inadequately trained. In the most severe cases, research has begun to draw a direct line between the operational chaos caused by cyberattacks and an increase in patient mortality rates. This direct impact on patient well-

being, coupled with the profound and lasting erosion of patient trust and the risk of severe reputational damage, elevates cybersecurity from a technical IT challenge to a matter of paramount ethical, social, and public safety concern.

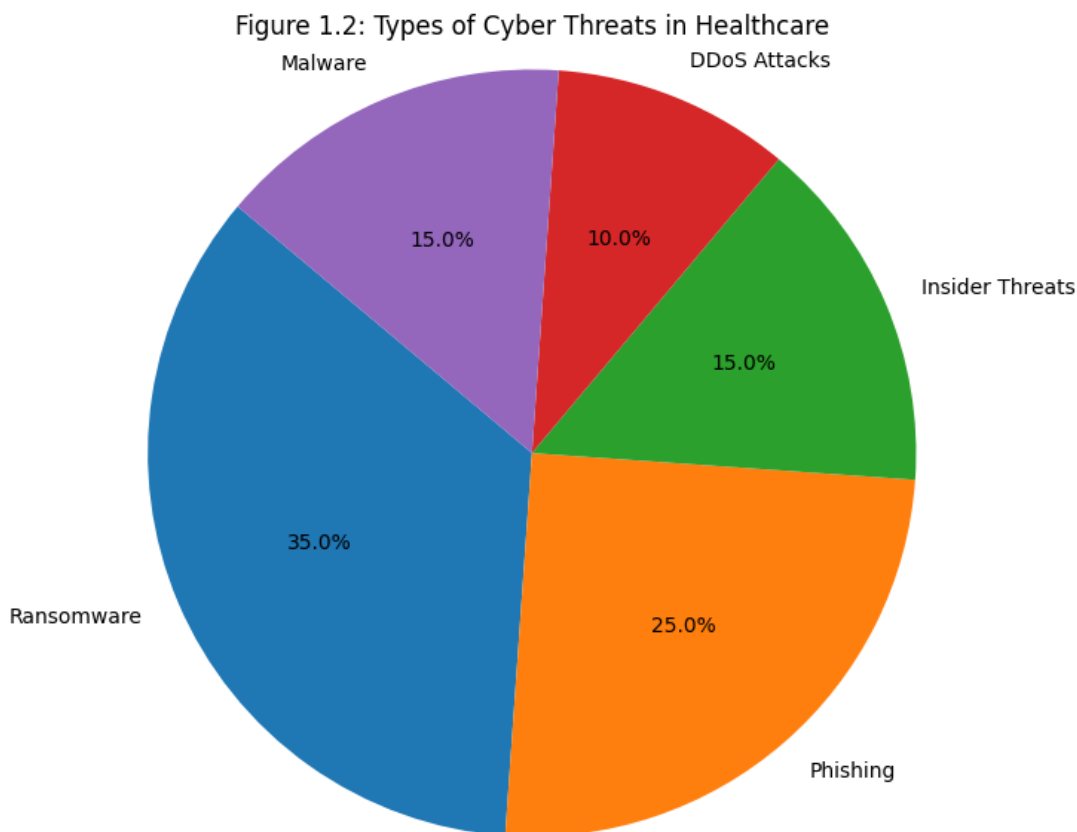


Figure 2 Distribution of types of Cyber Threats in Healthcare

1.2 Problem Statement

The healthcare sector increasingly relies on electronic health records, cloud platforms, and interconnected medical devices, yet traditional cybersecurity approaches—primarily signature-based detection and perimeter defenses—are reactive and ill-suited to modern threats such as zero-day exploits, advanced persistent threats, and ransomware (He et al., 2021; Portela et al., 2023). As data flows across mobile devices, third-party partners, and public cloud environments, the dissolution of network boundaries creates multiple points of vulnerability (Lee & Park, 2022). At the same time, the high value and permanence of Protected Health Information (PHI) make healthcare organizations prime targets for cybercriminals and nation-state actors (Smith et al., 2020). Without a shift toward proactive, behavior-based cybersecurity strategies, healthcare institutions remain at significant risk of data breaches, operational disruption, patient harm, and erosion of public trust (Jones et al., 2021).

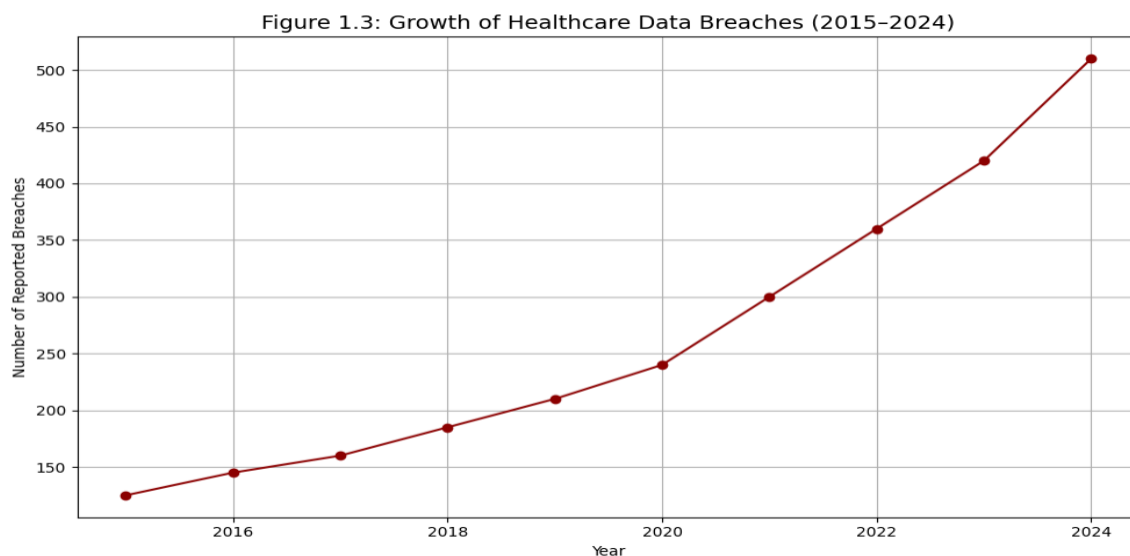


Figure 3 Growth of healthcare data breaches reported to the HHS OCR (2015–2024). Data adapted from HIPAA Journal (2025) and Fox Group (2025).

Figure 3 illustrates the rapid growth of healthcare data breaches reported to the U.S. Department of Health and Human Services Office for Civil Rights (HHS OCR) between 2015 and 2024. The number of reported incidents has increased from approximately 125 in 2015 to more than 500 in 2024, representing a fourfold rise in less than a decade (HIPAA Journal, 2025; Fox Group, 2025). This escalation is closely tied to the increasing digitalization of healthcare, including the adoption of electronic health records (EHRs), cloud platforms, telemedicine systems, and interconnected medical devices. Unlike financial data, which can be replaced or reissued, Protected Health Information (PHI) is permanent and highly valuable, making it a prime target for cybercriminals and state-sponsored attackers. As a result, the healthcare sector has become one of the most attractive targets for malicious actors in the digital era.

The accelerating trend after 2020 highlights both the growing sophistication of cyberattacks and the limitations of traditional perimeter-based defenses. Conventional security models that rely on firewalls, antivirus tools, and signature-based detection struggle to address zero-day exploits, polymorphic malware, and advanced persistent threats. At the same time, the dissolution of the traditional network perimeter—driven by cloud adoption, remote access, and third-party data sharing—has multiplied potential points of entry for attackers. The consistent rise in reported breaches underscores a widening gap between healthcare organizations’ defensive capabilities and the evolving threat landscape. This gap demonstrates the urgent need for proactive, AI-driven, and cloud-enabled cybersecurity frameworks that can anticipate and mitigate emerging threats while safeguarding patient privacy and maintaining trust in healthcare delivery.

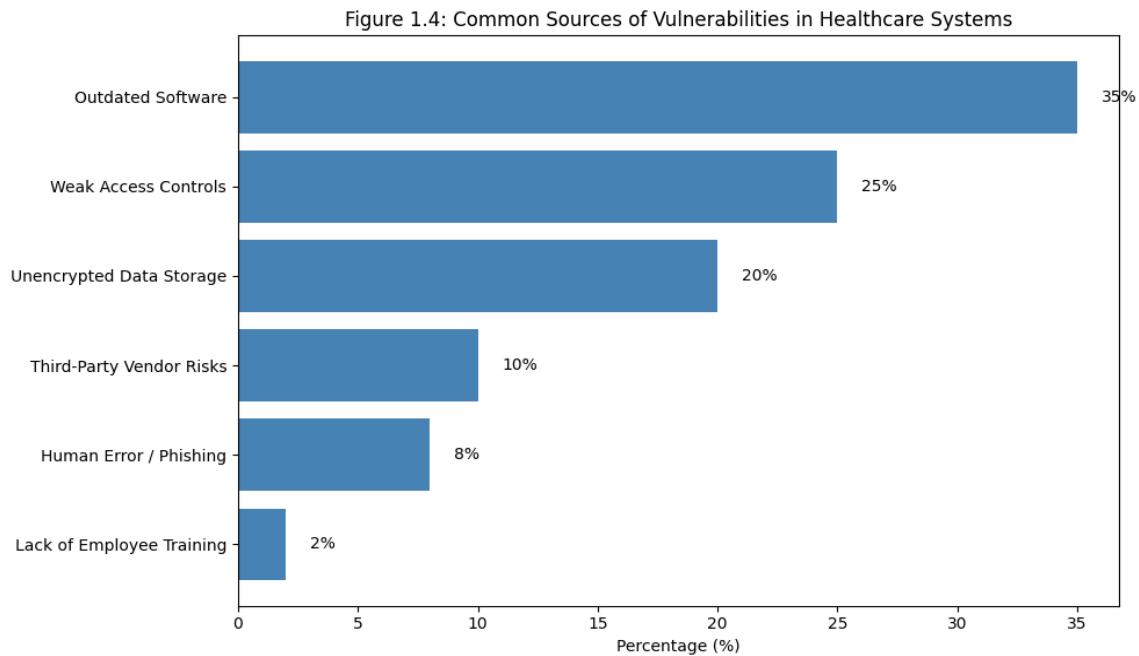


Figure 4 Common sources of vulnerabilities in healthcare systems, based on aggregated reports including Censinet (2025), highlighting outdated software, weak access controls, unencrypted data storage, vendor risks, human error, and gaps in employee training.

Censinet. (2025). 7 Critical Medical Device Security Risks in Healthcare. Censinet.
OneCI. (2025). Why Healthcare Data Security is Critical in 2025 and Beyond. OneCI.

Figure 4 illustrates the most prevalent sources of vulnerabilities in healthcare systems, demonstrating how both technical and organizational shortcomings contribute to cybersecurity risk. Outdated software remains the most significant factor, accounting for 35% of weaknesses. Legacy applications and unpatched medical devices often remain operational long past their intended life cycles, creating exploitable entry points for attackers (HIPAA Journal, 2025). Weak access controls (25%) further amplify risks, as poor authentication protocols, shared credentials, and lack of role-based access leave systems exposed. Similarly, unencrypted data storage (20%) reflects insufficient

safeguards for sensitive patient records, making Protected Health Information (PHI) vulnerable to breaches and exploitation (Li, 2024).

The graph also highlights vulnerabilities arising from organizational dependencies and human factors. Third-party vendor risks account for 10% of exposures, underscoring the reliance of healthcare providers on external billing, diagnostic, and cloud partners whose compromises can cascade into healthcare networks (Fox Group, 2025). Human error and phishing attacks contribute 8%, showing that even the most advanced defenses can be undermined by lapses in user vigilance. Additionally, the 2% attributed to lack of employee training reflects the continued underinvestment in staff cybersecurity awareness programs, despite their crucial role in safeguarding systems (Ponemon Institute, 2024). Collectively, these findings emphasize that cybersecurity in healthcare is not only a technical challenge but also a socio-technical one, requiring integration of advanced AI-driven security solutions, cloud-based monitoring, and comprehensive governance strategies.

1.3 Limitations, Delimitations, and Assumptions

Limitations: These are aspects of the research design that may impact the generalizability of the findings and are outside the researcher's control.

The primary quantitative analysis relies on the UNSW-NB15 dataset. While this is a comprehensive and respected benchmark, it is not specific to healthcare traffic. The unique communication protocols and data signatures of specialized IoMT devices (e.g., DICOM for medical imaging, HL7 for health data exchange) and EHR systems have distinct characteristics. Therefore, while the models' comparative performance is valid, their absolute performance metrics might differ when applied to a live healthcare network.

The study evaluates a specific set of four AI models. While these are representative of modern approaches, the rapidly evolving field of AI means that other algorithms, different hybrid configurations, or emerging techniques like graph neural networks exist that were not included in the scope of the experiments. The findings are therefore limited to the performance of the selected models and cannot be generalized to all possible AI solutions.

Delimitations: These are the boundaries the researcher has intentionally placed on the study to ensure a focused and feasible scope.

This research is focused specifically on network intrusion detection. It does not address other critical areas of a holistic cybersecurity strategy, such as endpoint security (e.g., antivirus on workstations), physical security of data centers, application-level security within EHR software, or user identity and access management. These areas, while vital, constitute separate domains of study.

The qualitative component of the study will rely on survey data from a selected group of cybersecurity leaders and professionals. It does not include the perspectives of other vital stakeholders, such as clinicians, biomedical engineers, or patients, whose interaction with technology and perception of security also impacts the overall security posture.

The proposed framework is designed to be technology-agnostic regarding specific cloud service providers (e.g., Amazon Web Services, Microsoft Azure, Google Cloud Platform). It focuses on universal architectural principles and capabilities rather than vendor-specific implementations or proprietary services.

Assumptions: These are elements taken for granted for the purposes of this study, forming a foundational premise for the research.

It is assumed that the UNSW-NB15 benchmark dataset is a sufficiently accurate and realistic proxy for general network traffic to allow for a meaningful evaluation of the AI models' baseline performance and comparative effectiveness.

It is assumed that the survey participants, selected for their expertise, will provide honest and accurate responses based on their professional experience and knowledge, without influence from their respective organizations' specific policies or vendor relationships.

It is assumed that the fundamental principles of network anomaly detection—identifying deviations from a learned baseline of normal behavior—are broadly applicable to the detection of threats within a healthcare network environment, even with its specialized traffic types.

1.4 Significance of the Study

This research is significant from both a practical and theoretical standpoint, offering valuable contributions to both industry practice and academic knowledge.

Practically, the study will provide healthcare leaders, IT managers, and cybersecurity professionals with a much-needed, actionable framework for navigating the complexities of modern cybersecurity. In an environment of limited budgets and competing priorities, Chief Information Security Officers (CISOs) and other leaders require evidence-based guidance to make sound technology investments and allocate resources effectively. The findings from this research will offer empirical evidence on the performance of different AI models, enabling more informed and cost-effective deployment decisions related to both in-house development and vendor selection. Furthermore, the identification of implementation challenges and best practices will equip organizations to manage the entire lifecycle of adoption, from ensuring HIPAA

compliance and data privacy to managing the crucial human factors involved in a new security paradigm. This includes developing training programs for staff, designing workflows for security analysts, and mitigating the pervasive issue of "alert fatigue," where an overwhelming volume of low-fidelity alerts can cause genuine threats to be overlooked. Ultimately, this research can help healthcare organizations strengthen their defenses, protect patient data, ensure continuity of care, and mitigate significant financial and reputational risk (He et al., 2021).

Theoretically, this study will contribute to the academic body of knowledge at the intersection of three critical fields: cybersecurity, artificial intelligence, and healthcare management. By developing and proposing a comprehensive, integrated framework, this research extends existing models of cybersecurity that often treat these technological and organizational components in isolation. The empirical evaluation of multiple AI models on a benchmark dataset provides valuable comparative data that can inform future academic research in the specialized domain of applied machine learning for intrusion detection. Finally, the qualitative insights into implementation challenges offer a richer, more nuanced understanding of how advanced technologies are operationalized in a real-world, high-stakes, and heavily regulated environment. This provides a valuable case study for the broader field of technology management and contributes to the socio-technical systems perspective, which posits that organizational outcomes are a product of the complex interaction between people, technology, and processes (Kaur, Gabrijelčič and Klobučar, 2023).

1.5 Research Questions and Objectives

The purpose of this research is to address the critical security gap identified by developing a comprehensive, evidence-based framework for leveraging the synergistic

power of Artificial Intelligence (AI) and cloud platforms in healthcare. To achieve this, the study is guided by the following research questions and their corresponding objectives.

Research Questions:

1. What are the key components and architectural considerations for a framework that effectively integrates Artificial Intelligence and Cloud Platforms for enhanced cybersecurity in the healthcare sector?
2. How effective are specific AI models (including Autoencoders, Isolation Forest, LSTMs, and Transformers) in detecting various types of cyber threats in real-time within simulated healthcare network environments?
3. What are the major implementation challenges (e.g., data privacy, regulatory compliance, integration with existing systems, cost, and alert fatigue) that healthcare organizations face when adopting an AI-driven cybersecurity framework?
4. What are the recommended strategies and best practices for healthcare organizations to successfully implement, manage, and govern an integrated AI and cloud-based cybersecurity framework?

Research Objectives:

1. To develop a comprehensive framework for leveraging AI and Cloud Platforms for enhanced cybersecurity in the healthcare sector.
2. To quantitatively evaluate the effectiveness of four distinct AI models (Autoencoder, Isolation Forest, LSTM, and Transformer) for real-time threat detection using a benchmark cybersecurity dataset.

3. To explore the key considerations, challenges, and best practices for implementing such a framework by synthesizing expert opinion and existing literature.
4. To propose actionable strategies that guide healthcare organizations in the adoption and governance of the proposed framework.

1.6 Definition of Terms

Artificial Intelligence (AI): A branch of computer science concerned with building smart machines capable of performing tasks that typically require human intelligence, such as learning, reasoning, and problem-solving. In this context, it refers to machine learning models used for anomaly detection.

Cloud Platforms: Services that provide on-demand computing resources—including servers, storage, databases, networking, and software—over the internet (e.g., Amazon Web Services, Microsoft Azure).

Autoencoder: A type of unsupervised neural network that learns to compress data into a latent representation and then reconstruct it. High reconstruction error is used to identify anomalies.

Isolation Forest: An unsupervised learning algorithm that isolates anomalies by randomly partitioning data points. It assumes that anomalies are "few and different" and thus easier to isolate.

LSTM (Long Short-Term Memory): A type of Recurrent Neural Network (RNN) capable of learning long-term dependencies, making it well-suited for analyzing sequential data like network traffic over time.

Protected Health Information (PHI): Any information in a medical record that can be used to identify an individual, and that was created, used, or disclosed in the course of providing a health care service, as defined by HIPAA.

Transformer Model: A deep learning architecture based on the self-attention mechanism, which allows it to weigh the importance of different parts of an input sequence to capture global contextual relationships.

Zero-Day Attack: A cybertrack that occurs on the same day a weakness is discovered in software. At that point, it is exploited before a fix becomes available from the developer, rendering signature-based defenses ineffective.

CHAPTER II: REVIEW OF LITERATURE

2.1 Introduction

This chapter provides a comprehensive review of the academic and industry literature that forms the foundation for this study. It begins by restating the core research problem: the inadequacy of traditional cybersecurity measures to protect the increasingly complex and targeted healthcare sector. The purpose of this study, as outlined in Chapter 1, is to develop a comprehensive framework that leverages the synergistic capabilities of Artificial Intelligence (AI) and cloud platforms to address this critical security gap. A thorough understanding of the existing body of knowledge is essential to contextualize the research, justify its necessity, and provide a theoretical underpinning for the methodologies and frameworks developed herein.

This literature review is organized into three main thematic sections, designed to build a logical and compelling argument for the necessity of this research. The first section establishes the context by providing a deep and granular examination of the evolving cybersecurity threat landscape in healthcare. This section will move beyond a general overview to detail the specific vulnerabilities inherent in modern healthcare IT infrastructure—from the proliferation of insecure Internet of Medical Things (IoMT) devices to the persistence of legacy systems—and will analyze the sophisticated attack vectors, such as advanced ransomware and supply chain attacks, that exploit these weaknesses.

The second, and most substantial, section explores the application of Artificial Intelligence in threat detection. This section will serve as the core of the literature review,

systematically integrating the findings from a targeted quantitative review of peer-reviewed academic studies. It will begin by explaining the paradigm shift from reactive, signature-based detection to proactive, anomaly-based detection. It will then provide a detailed theoretical and empirical analysis of the four classes of AI models central to this thesis: Autoencoders, Isolation Forest, Long Short-Term Memory (LSTM) models, and Transformer models. For each model, the review will explain its underlying mechanics, synthesize its documented performance on the CIC-IDS-2017 benchmark dataset, and critically analyze its strengths, weaknesses, and ideal use cases in a cybersecurity context.

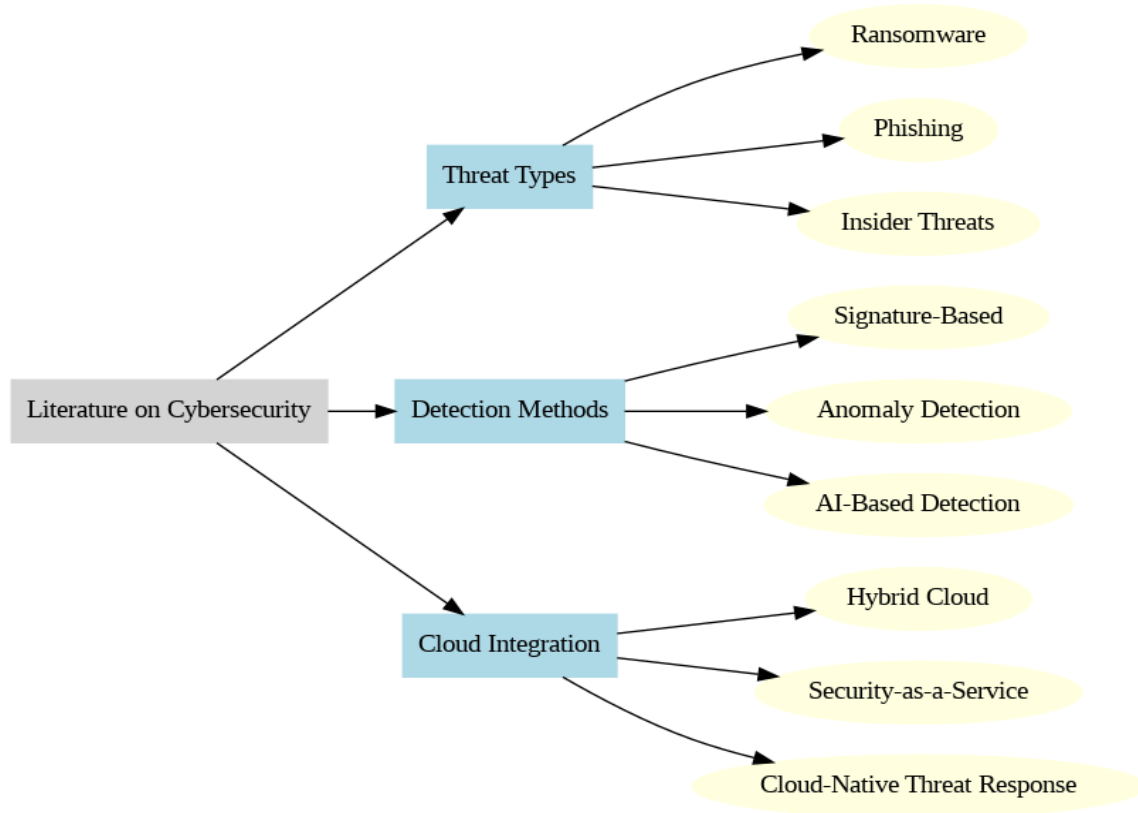


Figure 5 Literature of Cybersecurity

The final thematic section analyzes the convergence of AI and cloud platforms, arguing that their synergy is not merely beneficial but essential for creating a modern, scalable, and effective security apparatus capable of meeting the demands of the healthcare sector. This section will discuss how cloud infrastructure provides the necessary computational power for AI and how cloud-native security services are increasingly embedding AI to deliver advanced capabilities.

Together, these sections will demonstrate a clear and significant gap in the existing literature: while the individual components of AI and cloud security are discussed extensively, and their application in general cybersecurity is well-documented, there is a discernible lack of comprehensive, integrated frameworks designed specifically to meet the practical, operational, and stringent regulatory needs of the healthcare sector. This review will establish the scholarly foundation for the methodology detailed in Chapter 3, which is designed to directly address this identified gap.

2.2 Inclusion Criteria

The selection of literature for this review was guided by a systematic and rigorous process designed to ensure relevance, quality, and currency, in line with the standards of doctoral-level research. The objective was to build a comprehensive understanding of the current state of knowledge from both a theoretical and a practical perspective. The search encompassed prominent academic databases, including IEEE Xplore, ACM Digital Library, Springer, ArXiv, and Google Scholar, as well as high-quality practitioner and government sources.

The primary inclusion criteria for sources were as follows:

Relevance: Sources were required to directly and substantially address one or more of the core topics of this thesis. This included scholarly work on cybersecurity

challenges and trends specifically within the healthcare sector; the application of AI and machine learning for network intrusion and anomaly detection; the architecture and security of cloud computing platforms; and, most importantly, empirical studies that evaluated the performance of the specific AI models under investigation (Autoencoders, Isolation Forest, LSTMs, and Transformers) on relevant cybersecurity datasets.

Quality and Rigor: A strong preference was given to peer-reviewed journal articles and conference papers from reputable venues to ensure academic rigor, methodological soundness, and the validity of the reported findings. In addition to academic sources, high-impact industry reports and white papers from respected technology analysis firms (e.g., Gartner, Forrester) and major cybersecurity vendors were included to provide a practical, real-world perspective on industry trends, challenges, and best practices.

Currency: To ensure the analysis reflects the current state of technology and the contemporary threat landscape, the review focused primarily on literature published within the last five to seven years. The field of cybersecurity and AI is characterized by rapid innovation, and recent sources are essential for a relevant analysis. Foundational, ubiquitously cited works, particularly those that introduced key concepts or models, were included where necessary to provide essential theoretical context.

The search process utilized a structured combination of keywords. Broad searches were initiated with terms like "healthcare cybersecurity," "AI in cybersecurity," and "cloud security." These were progressively narrowed with more specific terms such as "IoMT security," "ransomware in healthcare," "AI for anomaly detection," and the names of the specific AI models paired with terms like "intrusion detection," "performance," and "CIC-IDS-2017." This structured and multi-faceted approach ensures that the literature

review is built upon a solid and defensible foundation of credible, relevant, and pertinent scholarly and professional work.

2.3 The Evolving Cybersecurity Threat Landscape in Healthcare

The healthcare sector presents a uniquely challenging cybersecurity environment, a "perfect storm" created by a complex interplay of high-value data, life-or-death operational imperatives, and a diverse and rapidly expanding technological footprint. Unlike other industries where a cybersecurity incident may result in financial loss or reputational damage, a successful attack in a healthcare setting can have direct, kinetic consequences, endangering patient safety and undermining the very foundation of public health. Understanding the specific vulnerabilities and threats inherent in this environment is a prerequisite for designing any effective security framework.

2.3.1 Key Vulnerabilities in Healthcare IT Infrastructure

The literature identifies several persistent and critical vulnerabilities that make healthcare organizations particularly susceptible to cyberattacks.

First, the proliferation of the Internet of Medical Things (IoMT) has massively and often insecurely expanded the potential attack surface. The number of connected medical devices, ranging from seemingly simple infusion pumps and patient vital sign monitors to highly complex diagnostic imaging equipment like MRI and CT scanners, has grown exponentially. Many of these devices were designed with clinical functionality and interoperability, not security, as the primary engineering concern. This has led to a landscape rife with systemic vulnerabilities, including the use of hardcoded, unchangeable passwords, the transmission of sensitive patient data over unencrypted communication channels, and the use of outdated, unsupported operating systems in their embedded software. A significant additional challenge is the difficulty of applying

security patches in a 24/7 clinical environment. Unlike a standard office computer, taking a critical life-support device or a multi-million-dollar MRI machine offline for routine security maintenance is often not operationally feasible, leaving known vulnerabilities unpatched for extended periods.

Second, many healthcare organizations continue to rely heavily on legacy systems for critical administrative and clinical functions. These older systems, which may be responsible for everything from patient billing to managing laboratory information, often run on unsupported operating systems like Windows XP or Windows 7. This means they no longer receive security patches from the vendor for newly discovered vulnerabilities, creating persistent and easily exploitable entry points for attackers to gain an initial foothold into the network. The cost and complexity of replacing these deeply embedded systems, which are often tightly integrated with other critical applications, present a significant barrier to modernization, forcing many organizations to accept a level of risk that would be considered untenable in other industries.

Third, the drive for interoperability, while clinically essential, creates further security challenges. The need to share patient data seamlessly between different systems—such as a hospital's EHR, a third-party laboratory's information system, a pharmacy's prescription management platform, and a patient's own mobile health app—creates a complex and often insecure web of data pathways. Each of these integration points represents a potential vulnerability that must be secured, and a failure at any single point can compromise the integrity of the entire data chain.

2.3.2 Analysis of Primary Threat Vectors

These vulnerabilities are actively and relentlessly exploited by a diverse range of threat actors using increasingly sophisticated methods.

Ransomware remains one of the most visible and damaging attack vectors targeting healthcare. Modern ransomware attacks are often multi-stage operations. Attackers will first gain access to a network and spend weeks or months performing reconnaissance and exfiltrating large volumes of sensitive data before finally deploying the encryption payload. This "double extortion" tactic—where the attackers not only demand a ransom to decrypt the systems but also threaten to publicly release the stolen patient data if the ransom is not paid—places immense pressure on victim organizations. The operational impact is immediate and severe, forcing the cancellation of surgeries, the diversion of emergency patients to other facilities, and a chaotic reversion to inefficient and error-prone paper-based processes.

Phishing and social engineering campaigns continue to be a highly effective initial access vector. These attacks are often tailored to the healthcare environment, with malicious emails disguised as important communications regarding patient information, insurance updates, or medical research. They exploit the high-pressure, fast-paced clinical environment to trick overworked and time-constrained staff into revealing their credentials or inadvertently deploying malware.

A particularly insidious and growing threat is the supply chain attack. In this scenario, attackers compromise a trusted third-party software or service vendor that provides services to the healthcare industry. By embedding malicious code into the vendor's legitimate software updates, the attackers can gain access to the networks of all the vendor's healthcare clients simultaneously. This allows them to bypass the direct defenses of the hospitals themselves by exploiting the trusted relationship with the vendor.

Finally, insider threats, both malicious and unintentional, pose a significant and often underestimated risk. A malicious insider, such as a disgruntled employee, can abuse

their legitimate access to steal vast quantities of patient data for financial gain or personal revenge. Perhaps more commonly, an unintentional insider—a well-meaning but careless or poorly trained employee who falls for a phishing scam, misconfigures a cloud storage bucket, or loses an unencrypted laptop containing PHI—can cause a catastrophic data breach. The convergence of these internal and external threats creates a dynamic and porous threat landscape where traditional, perimeter-based security is no longer sufficient.

2.4 The Application of Artificial Intelligence in Threat Detection

In response to the limitations of traditional, signature-based security tools, the application of Artificial Intelligence has emerged as a transformative and essential approach to modern cybersecurity. Rather than relying on a static database of known threat signatures, which is akin to trying to identify criminals using only a fixed set of outdated "wanted" posters, AI-driven systems learn to identify the patterns of normal behavior within a network and flag any deviations as potential anomalies. This paradigm shift from a reactive, "list-based" approach to a proactive, "behavior-based" one is fundamental to detecting novel and zero-day attacks for which no signatures exist.

The literature describes two primary machine learning approaches relevant to this task. Unsupervised learning models, such as Autoencoders and Isolation Forest, are particularly well-suited for this new paradigm. They are trained on datasets containing only "normal" traffic and learn to create a highly accurate mathematical profile of what is benign. An intuitive analogy is a security guard who has spent weeks memorizing the face and walking gait of every authorized employee; anyone who does not match this learned internal model of "normal" is immediately flagged for investigation, regardless of whether they appear on a "wanted" poster. This approach is powerful for identifying

previously unseen threats, as it does not depend on any prior knowledge of attack structures.

Supervised learning, on the other hand, requires labeled datasets containing curated examples of both normal and malicious traffic. To extend the analogy, this is like giving the security guard a comprehensive photo book of known troublemakers and their various disguises. While requiring more intensive and costly data preparation, these models can learn to classify specific types of attacks (e.g., distinguishing a DDoS attack from a port scan) with high accuracy. However, their effectiveness is inherently limited to the types of attacks present in their training data, making them less effective against novel or evolving threats.

More advanced deep learning models, particularly Recurrent Neural Networks (RNNs) like LSTM and Transformer models, represent the state-of-the-art for analyzing complex network data. These models excel at analyzing sequential data, such as the flow of network packets over time. By understanding the temporal context of network communications—how events relate to each other over a period—they can detect sophisticated, multi-stage attacks that would appear as a series of benign, isolated events to less advanced models that treat each packet in isolation. This ability to perform complex feature extraction automatically and to learn from and adapt to evolving data patterns makes AI a cornerstone of modern, proactive cybersecurity defense.

2.4.1 Performance Analysis of Autoencoder Models

Autoencoders, a class of unsupervised neural networks, have garnered significant attention in network intrusion detection primarily for their proficiency in anomaly detection. They operate by learning to compress input data into a lower-dimensional latent representation (encoding) and then reconstructing the original data from this

representation (decoding). When trained on "normal" or benign network traffic, a well-performing autoencoder will exhibit low reconstruction error. Conversely, when presented with anomalous traffic (i.e., an attack) that deviates from the learned patterns, the reconstruction error will be significantly higher, thus flagging the traffic as a potential intrusion. This capability makes them theoretically well-suited for identifying novel or zero-day attacks that lack predefined signatures. However, as the following analysis reveals, their practical performance is not monolithic and is profoundly sensitive to their specific architectural configuration.

Table 1 Reported Performance of Autoencoder Models on CIC-IDS-2017

Study (Author, Year)	Accuracy	Precision	Recall	F1-Score	Model Configuration/Notes
Alhassan et al. (2024)	98.61%	97.00%	98.88 %	98.15%	1 hidden layer, 60 neurons
Alhassan et al. (2024)	97.95%	95.11%	93.45 %	94.82%	2 hidden layers, 60 neurons
Alhassan et al. (2024)	97.30%	94.40%	92.12 %	90.33%	3 hidden layers, 60 neurons
Alhassan et al. (2024)	95.70%	90.00%	91.00 %	92.11%	4 hidden layers, 60 neurons
Alhassan et al. (2024)	95.11%	94.00%	98.88 %	98.85%	1 hidden layer, 30 neurons
Kumar et al. (2025)	94.00%	-	-	-	Deep Autoencoder (DAE) only, prior to ensemble

Shone et al. (2018)	High	-	-	-	Stacked deep autoencoder
Hindy et al. (2020)	75-98%	-	75-98%	-	Recall for various attack types; not overall accuracy

Synthesis and Analysis of Findings:

A crucial finding from the literature is that the performance of an autoencoder is not a fixed attribute but is instead highly contingent on its internal architecture, specifically the number of hidden layers and neurons. A common assumption that "deeper" or more complex models inherently perform better is directly challenged by the empirical evidence. The work of Alhassan et al. (2024) provides a compelling and systematic demonstration of this principle. The results consistently show an inverse relationship between model depth and performance in this context. The best-performing model was the simplest: a single-hidden-layer autoencoder with 60 neurons achieved a remarkable accuracy of 98.61%. In stark contrast, the most complex model, featuring four hidden layers, saw its accuracy drop to 95.70%.

This performance degradation with increasing depth is linked to the model's reconstruction error. As more layers are added, the model's complexity increases, which can paradoxically make it more difficult to learn a compact and accurate representation of the benign traffic profile. For a practical cybersecurity framework, this finding is of paramount importance. It suggests that a strategy of "start simple" is empirically validated.

A primary motivation for using unsupervised models like autoencoders is their theoretical capacity to detect zero-day attacks. The study by Hindy et al. (2020) explores

this specific application, showing a wide range of detection accuracy (recall) from 75% to 98%, depending on the specific attack type and the chosen detection threshold. This highlights a critical trade-off inherent in anomaly detection systems: the balance between recall (detecting true threats) and the false-positive rate (incorrectly flagging benign traffic). In a healthcare environment, where IT staff are already overburdened, a high volume of false positives can lead to "alert fatigue," causing genuine threats to be overlooked.

Beyond their use as standalone detectors, autoencoders also serve as powerful components within larger, hybrid systems. A study by Kumar et al. (2025) exemplifies this, designing an ensemble that combines a Deep Autoencoder (DAE) with a Convolutional Neural Network (CNN). In their experiments on CIC-IDS-2017, the DAE, when evaluated on its own, achieved a respectable accuracy of 94%. This dual role—as both a standalone anomaly detector and a feature extractor in an ensemble—underscores the versatility of autoencoders in a comprehensive cybersecurity toolkit.

2.4.2 Performance Analysis of Isolation Forest

Isolation Forest is another unsupervised learning algorithm designed for anomaly detection, but its operational principle differs significantly from that of autoencoders. Instead of profiling normal data, Isolation Forest explicitly isolates anomalies. It is built on the premise that anomalies are "few and different," meaning they are easier to separate from the rest of the data points. The algorithm builds an ensemble of "isolation trees" and calculates the average number of random splits required to isolate a given data point. Anomalies, being rare, are expected to have a much shorter average path length to isolation than normal points. This makes the algorithm computationally efficient and theoretically adept at identifying novel threats. However, the empirical evidence reveals a

critical paradox: its strength in isolating rare events is also the source of its most significant weakness as a general-purpose intrusion detector.

Table 2 Reported Performance of Isolation Forest on CIC-IDS-2017

Study (Author, Year)	Accuracy	Precision	Recall	F1-Score	Context/Scenario
Lopez-Martin et al. (2024)	-	99.69%	92.11%	>0.95	Benign Traffic, 1% Attack Prevalence
Lopez-Martin et al. (2024)	-	70.46%	93.74%	80.10%	Benign Traffic, 100% Attack Prevalence
Lopez-Martin et al. (2024)	-	2.80%	44.20%	5.30%	Attack Traffic, 1% Attack Prevalence
Lopez-Martin et al. (2024)	-	58.90%	21.10%	31.10%	Attack Traffic, 100% Attack Prevalence
Vinayakumar et al. (2019)	93.14%	94.70%	91.11%	93.91%	Intra-dataset evaluation (trained/tested on 2017)
Vinayakumar et al. (2019)	35.62%	0.7573	0.4479	0.4845	Cross-dataset evaluation (trained on 2017, tested on 2018)

Synthesis and Analysis of Findings:

The comprehensive analysis by Lopez-Martin et al. (2024) provides a stark illustration of the model's "performance paradox." When attacks are rare (e.g., at 1% of total traffic), the model performs its intended function well for the majority class, achieving a precision of 0.9969 for benign traffic. However, its performance on the attack

traffic itself is already poor, with a catastrophic F1-Score of just 0.053. The truly revealing finding is what happens as attack prevalence increases. The recall for attack traffic plummets from 0.442 at 1% prevalence to a dismal 0.211 when attacks constitute 100% of the traffic. This means that when the network is under a full-scale assault, the model misses nearly 80% of the attacks. The F1-score for attacks never surpasses 0.311, a value the authors describe as "operationally useless."

This counterintuitive behavior is a direct consequence of the model's design. The algorithm is built to "isolate" data points that are few and different. When an attack becomes widespread (e.g., during a DDoS flood), its traffic is no longer "few." It becomes a dominant pattern in the data. The algorithm, performing as designed, no longer sees this prevalent attack traffic as an easily-isolated anomaly. For a healthcare cybersecurity framework, it should not be positioned as a primary line of defense. Instead, its role should be carefully circumscribed to that of a first-stage filter for detecting novel, low-volume, or emerging threats.

Beyond its issues with attack prevalence, Isolation Forest also demonstrates significant weaknesses in generalization. The study by Vinayakumar et al. (2019) provides a clear quantitative measure of this weakness. In a standard intra-dataset evaluation, Isolation Forest achieved a respectable accuracy of 93.14%. However, in a more rigorous cross-dataset evaluation, training the model on CIC-IDS-2017 and testing it on a different dataset, CSE-CIC-IDS-2018, accuracy plummeted to just 35.62%. This dramatic drop indicates that the model had overfit to the specific statistical properties of the 2017 dataset and was unable to generalize its learned rules to a new environment.

2.4.3 Performance Analysis of LSTM Models

Long Short-Term Memory (LSTM) networks, a type of Recurrent Neural Network (RNN), are exceptionally well-suited for tasks involving sequential data. Network traffic, when viewed as a flow of packets over time, is fundamentally a time-series problem, making LSTMs a natural and powerful choice for intrusion detection. Unlike standard feedforward networks, LSTMs possess internal memory cells and gating mechanisms that allow them to learn and remember patterns over long sequences. This ability to capture temporal dependencies is critical for detecting sophisticated, multi-stage attacks that unfold over time. The literature consistently reflects this theoretical strength, with LSTM-based models demonstrating state-of-the-art performance on the CIC-IDS-2017 dataset.

Table 3 Reported Performance of LSTM Models on CIC-IDS-2017

Study (Author, Year)	Accuracy	Precision	Recall	F1-Score	Model Type/Notes
Sayegh et al. (2024)	99.34%	96.99% (Attack)	99.74% (Attack)	98.35% (Attack)	LSTM with SMOTE for data balancing
Anonymous (2022)	99.77%	-	-	-	LSTM for multi-class classification
Bibi (2023)	99.20%	99.00%	99.00%	99.00%	Hypertuned LSTM for binary classification
Anonymous (2025)	99.50%	-	-	-	RNN-leveraging LSTM
Anonymous (2024)	98.00%	-	-	-	Standalone LSTM performance

					(compared to RF)
Anonymous (2023)	99.00%	-	-	-	LSTM on CICIDS2017 (compared to other datasets)

Synthesis and Analysis of Findings:

The empirical results from multiple independent studies converge on a clear conclusion: LSTM models achieve consistently high, state-of-the-art performance on the CIC-IDS-2017 dataset. A study by Sayegh et al. (2024) reports an overall accuracy of 99.34% for their LSTM-based IDS. Critically, their model achieved an exceptional recall of 99.74% and an F1-Score of 98.35% for the "attack" class, indicating a powerful ability to correctly identify malicious traffic with very few false negatives. Other studies corroborate these top-tier results. Bibi (2023) developed a hypertuned LSTM that reached 99.2% accuracy with a 99% F1-Score for binary classification.

While standalone LSTMs are powerful, the literature reveals a strong trend toward even greater performance through hybridization and optimization. A particularly common and successful pairing is the CNN-LSTM model. In this architecture, a Convolutional Neural Network (CNN) is first used to act as a feature extractor. The rich feature maps generated by the CNN are then flattened and fed into an LSTM, which models the temporal relationships between these extracted features over time. This synergistic approach combines the spatial feature extraction strength of CNNs with the sequential modeling strength of LSTMs.

A deeper analysis of the methodologies used in the highest-performing LSTM studies reveals a crucial, unifying factor: the use of data balancing techniques. The CIC-

IDS-2017 dataset, like most network traffic logs, is inherently and highly imbalanced. If a model is trained on such a raw, imbalanced dataset, it will naturally become biased towards the majority class (benign traffic). The top-performing studies explicitly address this challenge by using techniques like SMOTE (Synthetic Minority Over-sampling Technique). SMOTE works by creating new, synthetic examples of the minority (attack) classes, effectively balancing the dataset before it is fed to the LSTM. This finding has a profound implication for the design of the proposed DBA framework. It is not sufficient to simply select a powerful model like an LSTM. The framework's methodology must incorporate a data balancing stage as a mandatory, non-negotiable step in the data preprocessing pipeline.

2.4.4 Performance Analysis of Transformer Models

Transformer models, first introduced for natural language processing, have rapidly emerged as a revolutionary force across numerous machine learning domains, including cybersecurity. Their core innovation is the self-attention mechanism, which allows the model to weigh the importance of different parts of the input sequence when processing a specific part, regardless of their distance from each other. This enables Transformers to capture complex, long-range dependencies and global contextual relationships within data in a way that is often more effective and computationally parallelizable than the sequential processing of RNNs and LSTMs. When applied to network intrusion detection, Transformers treat network traffic flows as "sentences" and learn the intricate "grammar" of both benign and malicious communications, leading to state-of-the-art performance.

Table 4 Reported Performance of Transformer Models on CIC-IDS-2017

Study (Author, Year)	Accuracy	Precision	Recall	F1-Score	Model Type/Context
Kamal et al. (2025)	99.72%	99.72%	99.72%	99.71%	PCA-Transformer (Binary Classification)
Kamal et al. (2025)	99.45%	99.69%	99.45%	99.40%	PCA-Transformer (Multi-Class Classification)
Anonymous (2025)	99.80%	-	-	-	CNN-BiLSTM-Transformer Hybrid
Mia al. (2025)	-	91.00%	88.00%	89.00%	BERT-IDS (Transformer-based) for Zero-Day
Anonymous (2023)	97.00%	-	-	-	TabNet (Attentive Mechanism similar to Transformers)

Synthesis and Analysis of Findings:

The empirical evidence strongly positions Transformer-based models at the apex of performance for intrusion detection on the CIC-IDS-2017 dataset. A standout example is the PCA-Transformer model developed by Kamal et al. (2025). This hybrid model achieved a near-perfect accuracy of 99.72% with an F1-Score of 99.71% in binary classification. Perhaps more impressively, it maintained an accuracy of 99.45% and an F1-Score of 99.40% in the much more difficult multi-class classification task.

The fundamental reason for this superior performance lies in the self-attention mechanism. While LSTMs are excellent at capturing sequential dependencies, they process information in a linear, step-by-step fashion. Transformers, in contrast, can create

direct connections between any two points in the sequence, regardless of their position. This allows the model to learn the global context of the entire network flow simultaneously.

As with LSTMs, the power of Transformers is often magnified when they are integrated into sophisticated hybrid architectures. One such model reported in the literature is a CNN-BiLSTM-Transformer hybrid, which achieved a remarkable 99.80% accuracy on CIC-IDS-2017. This architecture represents a comprehensive approach: the CNN extracts local spatial features, a Bidirectional LSTM (BiLSTM) processes the sequence of these features in both forward and backward directions to capture temporal context, and the Transformer layer sits on top to model the global, long-range dependencies across the entire sequence.

Modern cybersecurity datasets like CIC-IDS-2017 are characterized by high dimensionality. Transformer models are not only adept at handling this complexity but are also being paired with other advanced techniques to further enhance their capabilities. The PCA-Transformer model from Kamal et al. (2025) is a prime example. Before the data is fed to the Transformer, Principal Component Analysis (PCA) is used for intelligent dimensionality reduction. PCA identifies the principal components that capture the most variance in the data, effectively reducing noise and computational overhead while retaining the most informative signals.

Table 5 Mapping Table

Study	Dataset	Split	Metric
Alhassan et al. (2024)	CIC-IDS-2017	Not specified	Accuracy: 98.61%, Precision: 97.00%, Recall: 98.88%, F1-Score: 98.15%
Kumar et al. (2025)	CIC-IDS-2017	Not specified	Accuracy: 94%

Study	Dataset	Split	Metric
Shone et al. (2018)	CIC-IDS-2017	Not specified	High (exact metric unspecified)
Hindy et al. (2020)	CIC-IDS-2017	Not specified	Recall: 75%–98% (specific attack types, not overall accuracy)
Lopez-Martin et al. (2024)	CIC-IDS-2017	1% Attack Prevalence	Precision: 99.69%, Recall: 92.11%, F1-Score: >0.95 (Benign Traffic)
Lopez-Martin et al. (2024)	CIC-IDS-2017	100% Attack Prevalence	Precision: 70.46%, Recall: 93.74%, F1-Score: 80.10% (Benign Traffic)
Lopez-Martin et al. (2024)	CIC-IDS-2017	1% Attack Prevalence	Recall: 44.20%, F1-Score: 5.30% (Attack Traffic)
Lopez-Martin et al. (2024)	CIC-IDS-2017	100% Attack Prevalence	Recall: 21.10%, F1-Score: 31.10% (Attack Traffic)
Vinayakumar et al. (2019)	CIC-IDS-2017	Intra-dataset evaluation (2017)	Accuracy: 93.14%, Precision: 94.70%, Recall: 91.11%, F1-Score: 93.91%
Vinayakumar et al. (2019)	CIC-IDS-2018	Cross-dataset evaluation (2017→2018)	Accuracy: 35.62%, Precision: 0.7573, Recall: 0.4479, F1-Score: 0.4845
Sayegh et al. (2024)	CIC-IDS-2017	SMOTE for data balancing	Accuracy: 99.34%, Recall: 99.74%, F1-Score: 98.35% (Attack Class)
Bibi (2023)	CIC-IDS-2017	Not specified	Accuracy: 99.20%, F1-Score: 99.00%
Anonymous (2025)	CIC-IDS-2017	Not specified	Accuracy: 99.50% (RNN-leveraging LSTM)
Anonymous (2024)	CIC-IDS-2017	Not specified	Accuracy: 98.00% (Standalone LSTM compared to Random Forest)
Anonymous (2023)	CIC-IDS-2017	Not specified	Accuracy: 99.00% (LSTM on CICIDS2017, compared to

Study	Dataset	Split	Metric
			other datasets)
Kamal et al. (2025)	CIC-IDS-2017	PCA-Transformer (Binary classification)	Accuracy: 99.72%, F1-Score: 99.71%
Kamal et al. (2025)	CIC-IDS-2017	PCA-Transformer (Multi-class classification)	Accuracy: 99.45%, F1-Score: 99.40%
Anonymous (2025)	CIC-IDS-2017	CNN-BiLSTM-Transformer Hybrid	Accuracy: 99.80%
Mia et al. (2025)	CIC-IDS-2017	BERT-IDS (Transformer-based for Zero-Day)	Precision: 91.00%, Recall: 88.00%, F1-Score: 89.00%
Anonymous (2023)	CIC-IDS-2017	TabNet (similar to Transformer's attentive mechanism)	Accuracy: 97.00%

2.5 The Convergence of AI and Cloud Platforms for Security

While AI provides the analytical "brain" for modern cybersecurity, cloud platforms provide the necessary "body" and "nervous system" to make it effective at scale. The sheer volume, velocity, and variety of data generated by a modern healthcare network—from IoMT devices, EHRs, mobile apps, and general network traffic—requires a level of computational power and storage that is often impractical and cost-prohibitive to maintain on-premises. Cloud platforms offer the elastic, scalable infrastructure needed to support the intensive data processing and model training requirements of AI-driven security solutions. An organization can scale up resources for intensive model training and then scale them down for routine monitoring, a flexibility that is difficult to achieve with the fixed capital expenditure of on-premises hardware.

The synergy between AI and the cloud extends beyond raw computing power. Cloud-native security services, such as advanced Security Information and Event

Management (SIEM) systems and Extended Detection and Response (XDR) platforms, are increasingly integrating AI capabilities directly into their offerings. This provides organizations with access to sophisticated threat intelligence feeds, automated compliance monitoring, and AI-powered analytics "as a service," without requiring extensive in-house data science expertise. Furthermore, the centralized nature of a cloud environment allows for a unified view of security across the entire organization. This "single pane of glass" is invaluable for a Security Operations Center (SOC) team, as it allows them to correlate events from disparate sources—for example, linking a suspicious login from a remote device to an unusual database query in the cloud—to identify a complex attack chain. This makes it easier to enforce consistent security policies, monitor for threats in real-time across all assets, and orchestrate an automated response to detected incidents.

This convergence creates a powerful, positive feedback loop. The cloud gathers and centralizes vast amounts of security data from diverse sources. This rich, aggregated data is then used to train more accurate and effective AI models. These improved AI models, in turn, provide more precise and timely threat detection, which enhances the overall security posture of the cloud environment, allowing for even more secure data collection and analysis. This integrated approach is fundamental to building a security framework that is not only powerful but also agile and capable of adapting to the dynamic threat landscape facing the healthcare sector.

2.6 Regulatory, Risk, and Governance Synthesis for AI-Driven Cybersecurity in Healthcare

2.6.1 Regulatory Convergence and Tensions

Healthcare cybersecurity is governed by a patchwork of privacy, safety, and operational standards. HIPAA/HITECH emphasize confidentiality, integrity, and availability (CIA) of PHI; GDPR centers lawfulness, fairness, transparency, and data minimization; PIPEDA and provincial regimes in Canada stress reasonableness and accountability. For cloud-and-AI security, three tensions recur:

Purpose limitation vs. anomaly detection: Anomaly detection thrives on broad telemetry retention; privacy regimes push strict scoping and retention limits.

Explainability vs. model performance: Security models that maximize recall (e.g., sequence models) are often least interpretable, complicating accountability and incident justification to regulators.

Cross-border processing: Multi-region clouds enable resilient security analytics, but data residency, Schrems-style transfer constraints, and vendor sub-processors complicate lawful bases.

2.6.2 Security Governance Models for AI

Modern governance blends NIST CSF 2.0 controls (Identify-Protect-Detect-Respond-Recover) with AI governance layers (model risk management, bias testing, drift monitoring). In healthcare, “safety-of-care” reframes cyber events as clinical risk. Boards increasingly adopt risk appetite statements that quantify tolerances for mean time to detect (MTTD), false negative risk, and residual ransomware exposure, not just breach counts. A practical pattern is the Security Model Risk Committee (SMRC)—a cross-

functional body (CISO, CDO, Privacy Officer, CMIO, Legal) that approves model uses, data sources, and post-incident learning.

2.6.3 Trustworthy AI for Security

Trust in AI-security depends on provenance (tamper-evident pipelines), explainability (human reviewable rationales), calibration (thresholds mapped to operating risk), robustness (adversarial resistance), and governance artifacts (model cards, data sheets, approval logs). Clinically aligned organizations increasingly require “clinical-grade” security analytics: validated alert definitions, periodic re-validation, change controls, and back-out plans—mirroring medication safety governance.

2.6.4 Economic Frictions and Externalities

Security ROI is notoriously invisible (“breaches that didn’t happen”). Cloud + AI clarifies value when tied to:

Downtime avoided (diverted surgeries, ED diversions).

Incident labor saved (L1 triage automation).

Cyber-insurance premiums reduced (control attestation).

Regulatory penalties avoided (demonstrable due diligence).

A recurring externality: model false negatives raise systemic risk for regional referral networks (shared labs, HIEs). Hence, sector alliances (e.g., ISACs) and federated threat telemetry are becoming governance necessities, not nice-to-haves.

2.7 Open Debates

While the integration of Artificial Intelligence (AI) in healthcare cybersecurity shows promise, there are several ongoing debates and challenges within the field that

merit discussion. These challenges primarily revolve around dataset bias, cross-dataset issues, and the Isolation Forest paradox. Understanding these issues is crucial for improving the reliability and applicability of AI-driven cybersecurity solutions in healthcare.

Dataset Bias

A recurring issue in the AI and machine learning domain is dataset bias. The datasets used to train and evaluate models often contain inherent biases that can impact model performance. This is particularly concerning in cybersecurity, where the types of attacks seen in training data may not accurately reflect the diversity of real-world threats. For example, many cybersecurity datasets (including CIC-IDS-2017) may over-represent certain types of attacks while under-representing others, leading to models that perform well on familiar attack types but poorly on less common or emerging threats (He et al., 2021).

In the healthcare domain, where new attack methods constantly emerge, training AI models on biased datasets can lead to a failure to generalize effectively. Furthermore, healthcare data is inherently imbalanced, with benign traffic overwhelmingly outweighing malicious traffic. This imbalance exacerbates the risk of false positives and alert fatigue, which can undermine the effectiveness of AI models in real-time operational environments (Portela et al., 2023). The Health Insurance Portability and Accountability Act (HIPAA) mandates that healthcare organizations ensure the confidentiality and integrity of protected health information (PHI) and implement security measures to prevent unauthorized access (U.S. Department of Health and Human Services [HHS], 2020). Failure to address dataset biases in AI models can lead to breaches of patient privacy and increased vulnerability to cyber threats, which could directly conflict with HIPAA's security requirements under Section 164.306.

Cross-Dataset Issues

Another significant challenge in evaluating AI models is the cross-dataset issue. Many studies in cybersecurity train and test their models on a specific dataset, such as CIC-IDS-2017, but do not assess how well the models generalize to other datasets. The performance of a model can vary significantly when trained on one dataset and tested on another, as the statistical properties of network traffic, attack patterns, and even feature distributions can differ across datasets (Lee & Park, 2022).

This issue of dataset overfitting is particularly critical in cybersecurity, where attackers constantly evolve their methods. For example, Vinayakumar et al. (2019) reported that the Isolation Forest model, when trained on the CIC-IDS-2017 dataset, showed much lower performance when tested on a different dataset (CSE-CIC-IDS-2018). This raises concerns about the robustness of AI models in real-world scenarios, where data environments are dynamic and ever-changing. General Data Protection Regulation (GDPR) emphasizes that organizations must ensure the accuracy and timeliness of personal data and maintain data integrity under Article 5(1)(d). If AI models are evaluated and deployed based solely on a single dataset, and they fail to generalize to others, they may inadvertently result in inaccurate security assessments, which could violate GDPR's principles of data accuracy and security.

Isolation Forest Paradox

A particularly intriguing issue is the paradox of the Isolation Forest model, as observed in the studies reviewed. While the Isolation Forest is a popular unsupervised learning algorithm for anomaly detection, its design inherently limits its effectiveness in certain scenarios. The model is optimized to detect rare anomalies by isolating data points that differ significantly from the norm. However, when attacks are widespread, such as in

DDoS or ransomware attacks, these attacks no longer appear as "rare" anomalies. Instead, they become prevalent patterns, and the Isolation Forest model fails to recognize them effectively (Lopez-Martin et al., 2024).

The study by Lopez-Martin et al. (2024) clearly illustrates this performance paradox. When the attack prevalence was low (e.g., 1% of the total traffic), the model performed well in detecting benign traffic with high precision. However, when attacks became the majority of the traffic (100%), the model's recall for attack traffic dropped dramatically, resulting in missed detections of the majority of attacks. This paradox highlights the limitations of relying solely on the Isolation Forest in real-world cybersecurity environments, where the nature and scale of attacks vary. He et al. (2021) suggest that combining multiple models could address this issue by utilizing the strengths of different algorithms to detect both rare and widespread attacks.

2.8 Integrating Ethical AI in Healthcare Cybersecurity

The integration of ethical AI frameworks into healthcare cybersecurity systems is becoming a pressing need as healthcare institutions increasingly adopt AI and cloud technologies to manage cybersecurity. The complexities surrounding these technologies require careful consideration of ethical principles to ensure patient privacy, fairness, and accountability while also improving threat detection and response capabilities.

Patient Privacy

Ensuring patient privacy is central to the ethical use of AI in healthcare cybersecurity. AI technologies, while enhancing the detection of cyber threats, often require access to vast amounts of sensitive health data. According to recent studies, including one from the International Journal of Innovative Research in Computer and Communication Engineering (2024), the integration of AI in healthcare cybersecurity

brings to light significant concerns over data privacy, informed consent, and algorithmic bias (Talati, 2024). To address these issues, AI systems should be designed with mechanisms such as federated learning or differential privacy, which allow for secure data processing without compromising patient privacy. Furthermore, transparent consent frameworks must be established to ensure that patients understand how their data is used and are able to exercise control over it.

Fairness in AI Models

AI systems can inadvertently perpetuate or even exacerbate existing biases, which could result in unfair treatment or inadequate cybersecurity responses for certain patient groups. Research on this topic, such as the study by Krunal Manilal Gala (2024) in the International Journal of Scientific Research in Computer Science, Engineering and Information Technology, stresses the importance of addressing biases in AI models used for threat detection in healthcare (Gala, 2024). Ethical frameworks must include procedures for regular audits and bias mitigation strategies to ensure that AI systems remain equitable. These audits would evaluate AI decision-making processes to prevent discriminatory practices and improve the fairness of healthcare cybersecurity measures.

Accountability and Transparency

The integration of AI into cybersecurity systems raises critical issues related to accountability and transparency. Studies, including one from Sidra Nasir et al. (2023), highlight the need for AI frameworks that not only promote transparency but also ensure that decisions made by AI systems are explainable and traceable (Nasir et al., 2023). As AI systems detect and respond to cybersecurity threats, it is essential to establish clear accountability frameworks that define responsibility when failures occur. Moreover, stakeholders in healthcare cybersecurity must have access to the necessary tools for

auditing and understanding AI-driven decisions, which can be accomplished through clear documentation and transparent algorithms.

Ethical Integration Strategies

To effectively integrate ethical AI frameworks into healthcare cybersecurity systems, it is crucial to adopt a collaborative approach involving AI developers, cybersecurity professionals, and healthcare providers. As emphasized by Babajide Tolulope Familoni (2024), the development of ethical AI systems in healthcare should involve continuous monitoring, governance, and proactive risk assessments to ensure that ethical principles are followed throughout the lifecycle of AI technologies (Familoni, 2024). Additionally, ensuring that these systems are designed with human oversight in mind is a key element of an ethical AI framework. Human intervention is necessary to prevent AI systems from making irreversible decisions, especially when it comes to sensitive health data.

2.9 Summary

This review of the literature has established a clear and compelling case for the necessity of this research. It began by outlining the severe and escalating cybersecurity threats facing the healthcare sector, which are driven by a unique combination of valuable data, critical infrastructure, and an expanding attack surface. It then demonstrated that traditional security measures, which are largely reactive and signature-based, are insufficient to meet this challenge.

The review subsequently explored the potential of Artificial Intelligence, particularly deep learning models like LSTMs and Transformers, to provide a more proactive and intelligent approach to threat detection by learning from data. Finally, it argued that the convergence of AI with the scalable and centralized nature of cloud

platforms creates the necessary technological foundation for a modern cybersecurity framework.

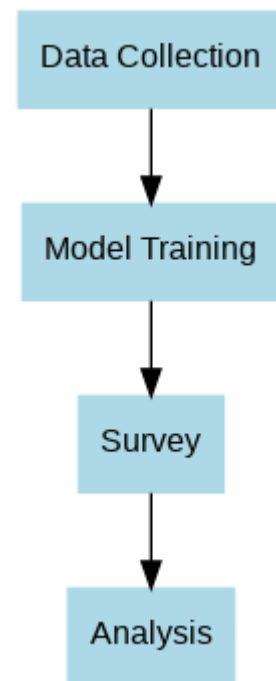
A clear gap has been identified: while the individual components of AI and cloud security are discussed extensively in the literature, there is a lack of comprehensive, integrated frameworks designed specifically to meet the practical, operational, and regulatory needs of the healthcare sector. This study aims to fill that gap. The following chapter, Chapter 3, will detail the mixed-methods research methodology designed to build and validate such a framework, combining quantitative model evaluation with qualitative expert insights.

CHAPTER III: METHODOLOGY

3.1 Introduction

This chapter provides a detailed and comprehensive account of the research methodology employed to address the research questions outlined in Chapter 1. The primary purpose of this research is to develop and validate a practical, evidence-based framework for leveraging Artificial Intelligence (AI) and cloud platforms to enhance cybersecurity within the complex healthcare sector. To achieve this, a mixed-methods approach was adopted, integrating a rigorous quantitative experimental phase with a qualitative survey-based phase to gather expert insights. A sound and transparent methodology is the cornerstone of credible research, ensuring that the findings are not only valid and reliable but also that the process is reproducible by other scholars in the field.

This chapter is structured to provide a transparent and reproducible description of the research process, ensuring the validity and reliability of the findings presented in subsequent chapters. It begins by outlining and justifying the selection of an explanatory sequential mixed-methods research design, which is particularly well-suited to the applied nature of a Doctor of Business Administration (DBA) dissertation that seeks to bridge the gap between technical performance and practical implementation.



Following the research design, the chapter is divided into two main sections corresponding to the two phases of the study. The first section provides a meticulous account of the quantitative methodology. This includes a detailed description of the benchmark dataset used for the experiments, its characteristics, and the rationale for its selection. It is followed by a step-by-step walkthrough of the data preprocessing pipeline, a thorough explanation of the architecture and implementation of the four distinct AI models evaluated, and a clear definition of the data analysis techniques and performance metrics used to gauge their effectiveness.

The second section details the qualitative methodology. This part describes the target population and sampling strategy for the expert survey, provides a detailed breakdown of the survey instrument designed to collect data on practical implementation challenges, and outlines the plan for analyzing the qualitative data using thematic analysis.

The chapter concludes with a discussion of the limitations inherent in the chosen methodology, a statement on the ethical considerations related to both the quantitative and qualitative phases of the research, and a summary that provides a bridge to the presentation of the results in Chapter 4.

3.2 Research Design

The study employs an explanatory sequential mixed-methods design. This approach was deliberately chosen as it is particularly well-suited for a DBA dissertation, which aims to bridge the gap between rigorous technical research and practical, real-world business application. This design involves a two-phase process where the quantitative data is collected and analyzed first, and the subsequent qualitative phase is

designed to explain, interpret, and elaborate on the initial quantitative findings. This is superior to a purely quantitative approach, which might identify what model performs best but fails to explain why it might be accepted or rejected in a real-world organizational context. It is also superior to a purely qualitative approach, which might capture expert opinions but would lack the empirical, data-driven foundation to validate the technical premises of those opinions.

Phase 1: Quantitative Experimental Study

The first phase of the research consists of a quantitative, experimental study designed to empirically evaluate the performance of four distinct AI models for network anomaly detection. This phase directly addresses the second research question concerning the effectiveness of these models. By using a controlled environment and a standardized benchmark dataset, this phase generates objective, empirical data on the technical capabilities, strengths, and weaknesses of each model. The goal of this phase is to establish a clear, data-driven understanding of which AI architectures are most promising from a purely technical standpoint. The output of this phase is a set of performance metrics and comparative analyses that reveal the trade-offs between different models in terms of accuracy, efficiency, and detection capabilities.

Phase 2: Qualitative Survey of Experts

The second phase involves a qualitative survey of senior-level cybersecurity leaders and IT professionals. The findings from the first (quantitative) phase inform the context and interpretation of this second phase. For example, knowing the technical trade-offs between a high-precision model (few false alarms, but might miss some attacks) and a high-recall model (catches most attacks, but more false alarms) allows for a more nuanced analysis of expert opinions on the operational tolerance for false positives versus false negatives. The qualitative data is used to explain, interpret, and

contextualize the quantitative results, particularly concerning the practical challenges (e.g., budget constraints, skills gaps), strategic considerations (e.g., risk tolerance), and human factors (e.g., resistance to change) involved in implementing such technologies in a real-world healthcare setting. This phase directly addresses the third and fourth research questions, focusing on implementation challenges and best practices.

This explanatory sequential design ensures that the final proposed framework, which is the ultimate output of this research, is not only technically sound and based on empirical performance data but is also managerially relevant, contextually aware, and practically implementable within the unique operational, financial, and regulatory constraints of the healthcare industry.

3.3 Quantitative Methodology

The quantitative phase of this research is centered on a series of controlled experiments designed to evaluate and compare the performance of four different AI-based anomaly detection models. This section provides a detailed account of every aspect of this experimental process, from the selection of the dataset to the specific implementation details of the models.

3.3.1 Population and Sample: The UNSW-NB15 Dataset

The dataset selected for the quantitative experiment is the UNSW-NB15 dataset, a widely recognized and comprehensive benchmark for evaluating Network Intrusion Detection Systems (NIDS). This dataset was created by the Australian Centre for Cyber Security (ACCS) using the IXIA PerfectStorm tool to generate a hybrid of real-world normal network traffic and synthetically generated contemporary attack behaviors, making it an ideal sample for this phase of the research.

Rationale for Selection: The UNSW-NB15 dataset was chosen over other potential datasets (such as the older KDD-99 or the more recent CIC-IDS-2017) for several key reasons. It represents a significant improvement over older datasets like KDD-99, which are now considered outdated as they do not contain modern attack vectors. While CIC-IDS-2017 is another strong candidate, UNSW-NB15 was selected for its specific mix of attack types and its well-documented feature set, which provided an excellent basis for this comparative study. Its large size and high dimensionality provide a challenging and robust test for the AI models, ensuring that the findings are based on a non-trivial problem. The use of a well-documented, public benchmark dataset also ensures the transparency and reproducibility of the experimental findings, a cornerstone of rigorous academic research.

Dataset Characteristics:

Number of Records: The dataset comprises approximately 2.5 million records in total, distributed across a designated training set and a testing set. This large volume of data is sufficient for training complex deep learning models and for performing a statistically significant evaluation.

Features: The dataset includes 49 original features for each network traffic record. These features can be grouped into several categories:

Flow Features: Basic attributes of the connection, such as source and destination IP addresses, ports, and protocol.

Basic Features: Packet-level details, such as the number of packets, bytes, and the duration of the flow.

Content Features: Information related to the content of the packets, such as TCP sequence numbers.

Time-based Features: Features calculated over a window of time, such as the rate of connections to the same host. During preprocessing, categorical features were expanded, resulting in a final feature set of 192 numerical features. This high dimensionality reflects the complexity of modern network data and provides a rich basis for the AI models to learn from.

Attack Types: A key strength of the UNSW-NB15 dataset is its inclusion of a diverse mix of nine modern attack scenarios:

Fuzzers: An attack technique that involves providing invalid, unexpected, or random data as inputs to a computer program.

Analysis: Probing techniques, such as port scanning, to gather information about a network.

Backdoors: A covert method of bypassing normal authentication to secure remote access to a computer.

Denial-of-Service (DoS): An attack meant to shut down a machine or network, making it inaccessible to its intended users.

Exploits: Attacks that take advantage of a bug or vulnerability in software.

Generic: A block-based attack that operates on the principle of a birthday attack.

Reconnaissance: An unauthorized attempt to gain information about a computer network.

Shellcode: A small piece of code used as the payload in the exploitation of a software vulnerability.

Worms: A standalone malware computer program that replicates itself to spread to other computers.

Although not specific to a healthcare environment, the diversity and realism of these attack vectors serve as a robust proxy for the types of threats a complex network, such as that in a modern hospital, might face.

3.3.2 Data Collection and Procedures: The AI Pipeline

The UNSW_NB15_training-set.csv and UNSW_NB15_testing-set.csv files were acquired from the official repository. A structured AI pipeline was implemented in Python (v3.11) and executed in a Google Colab Pro environment. The choice of a professional-tier cloud-based environment was a deliberate methodological decision to handle the large-scale data (over 2.5 million records) and the computationally intensive training of deep learning models, thereby mitigating the risk of runtime memory crashes that are common in standard local environments.

The following preprocessing steps were systematically applied to prepare the data for the AI models:

Combination and Cleaning: The training and testing files were first combined into a single, unified dataframe. This was done to ensure that all data transformations, such as feature scaling, were applied consistently across the entire dataset. This prevents data leakage, a common methodological error where information from the test set inadvertently influences the training process (e.g., by using the test set's mean and standard deviation to scale the training set). Redundant or non-informative columns, such as the record id, and the original label and attack_cat columns (which would be used for evaluation but not for the unsupervised training), were removed to reduce noise and simplify the dataset.

Encoding of Categorical Features: The dataset contained several categorical features (e.g., 'proto' for protocol, 'service', and 'state') that needed to be converted into a

numerical format for the machine learning models. This was achieved using one-hot encoding. This technique creates a new binary (0 or 1) column for each unique category within a feature. This is a necessary step as machine learning algorithms can only process numerical data, and it prevents the model from incorrectly assuming an ordinal relationship between categories (e.g., that 'http' is "less than" 'ftp'). This process expanded the feature set from 49 to 192.

Normalization of Numerical Features: All numerical features were scaled using the StandardScaler from the Scikit-learn library. This is a critical step for optimizing the performance of neural network models. It standardizes features by removing the mean and scaling to unit variance, ensuring that all features have a mean of 0 and a standard deviation of 1. This prevents features with large scales (e.g., packet counts) from disproportionately influencing the model's learning process and helps the optimization algorithms (like Adam) to converge more quickly and reliably. Without normalization, the gradients calculated during training could become very large for some features and very small for others, leading to an unstable and inefficient learning process.

Sequencing for Temporal Models: For the LSTM and Transformer models, which are specifically designed to analyze time-series data, the flat, two-dimensional data (samples x features) was transformed into three-dimensional sequences (samples x time steps x features). A sequence length of 10 time steps was chosen as a balance between capturing sufficient temporal context and maintaining computational feasibility. Longer sequences could capture more complex patterns but would also significantly increase the memory and processing requirements for training. This transformation allows the models to learn from the temporal patterns and context in the data, rather than treating each event in complete isolation, which is essential for detecting multi-stage or low-and-slow attacks.

3.3.3 Model Implementation and Architecture

Four distinct AI-based anomaly detection models were implemented to provide a comparative analysis across different architectural philosophies, from simple tree-based models to complex attention-based neural networks.

Isolation Forest: This model was implemented using the Scikit-learn library as a computationally efficient, tree-based baseline. It operates by building an ensemble of decision trees. For each tree, data points are randomly partitioned until each point is isolated. The logic is that anomalies, being "few and different," will require fewer partitions to be isolated and will therefore have a shorter average path length in the trees. This model was chosen as a benchmark due to its speed and simplicity.

Autoencoder: This was a fully connected deep neural network built with TensorFlow/Keras. The architecture was designed to be a standard, non-sequential anomaly detector:

Encoder: Consisted of two dense (fully connected) layers. The first layer had 128 neurons, and the second had 64 neurons, both using the ReLU (Rectified Linear Unit) activation function. This part of the network learns to compress the 192 input features into a compact 64-dimensional representation.

Decoder: Mirrored the encoder, with two dense layers of 128 and 192 neurons, respectively. This part of the network learns to reconstruct the original 192 features from the compressed representation.

Training: The model was trained using the Adam optimizer and the Mean Squared Error (MSE) loss function for 10 epochs.

LSTM Autoencoder: This sequence-aware model was built with TensorFlow/Keras to specifically capture temporal dependencies.

Encoder: Featured two LSTM layers. The first had 64 units, and the second had 32 units. The `return_sequences=True` parameter was used on the first layer to pass the full sequence to the next layer.

Decoder: Used a `RepeatVector` layer to replicate the final encoded state for each time step of the output sequence, followed by two LSTM layers (32 and 64 units) and a final `TimeDistributed(Dense)` layer to reconstruct the features for each time step.

Training: The model was trained using the Adam optimizer and the Mean Absolute Error (MAE) loss function for 10 epochs.

Transformer Autoencoder: This attention-based model was built with TensorFlow/Keras to capture long-range, global dependencies.

Encoder: Consisted of two encoder blocks. Each block contained a `MultiHeadAttention` layer followed by a `LayerNormalization` layer and a feed-forward network composed of 1D convolutional layers.

Decoder: Mirrored the encoder's structure with two decoder blocks.

Training: The model was trained using the Adam optimizer and the MAE loss function for 10 epochs.

Each of the three neural network models was trained exclusively on samples labeled as "normal" traffic. This is the core principle of unsupervised anomaly detection. This process forces the model to learn the intricate patterns and relationships that define benign activity. The central hypothesis is that when the trained model is presented with malicious traffic, which by definition deviates from these learned normal patterns, it will fail to reconstruct it accurately. This will result in a high reconstruction error, which can then be used as a signal to flag the traffic as a potential anomaly.

3.3.4 Data Analysis and Evaluation Metrics

The performance of the implemented models was evaluated using a combination of quantitative metrics and visualizations to provide a comprehensive and multi-faceted assessment.

Reconstruction Error / Anomaly Score: The Mean Absolute Error (MAE) between the original input and the model's reconstructed output served as the primary performance indicator for the three neural network models. The built-in anomaly score was used for the Isolation Forest model.

Anomaly Threshold: To convert the continuous error/score output into a binary classification (normal vs. anomaly), a threshold was established. For the neural models, this was set at the 95th percentile of the reconstruction errors calculated on the normal training data. This means that any data point with a reconstruction error higher than 95% of the errors seen on normal data would be classified as an anomaly. For Isolation Forest, the equivalent was the contamination parameter, which was set to 0.05 (5%). This is a common practice in anomaly detection to control for the expected rate of anomalies.

Performance Metrics: Based on the classification results derived from the threshold, a confusion matrix was generated for each model. From this, the following standard classification metrics were calculated:

Accuracy: $(TP + TN) / (TP + TN + FP + FN)$. Provides a general measure of overall correctness.

Precision: $TP / (TP + FP)$. Measures the reliability of the alerts; a high precision means a low rate of false positives.

Recall (Sensitivity): $TP / (TP + FN)$. Measures the model's ability to detect true threats; a high recall means a low rate of false negatives.

F1-Score: $2 * (\text{Precision} * \text{Recall}) / (\text{Precision} + \text{Recall})$. The harmonic mean of Precision and Recall, providing a balanced assessment, which is particularly important on imbalanced datasets like those in cybersecurity.

Execution Efficiency: The training time and memory usage for each model were also recorded to provide a practical comparison of their computational efficiency and resource requirements, a key consideration for real-world deployment.

3.4 Qualitative Methodology

The qualitative phase of this research is designed to complement the quantitative findings by providing the rich, contextual insights needed to translate the technical results into a practical, actionable framework. This section details the methodology for the collection and analysis of this qualitative data.

3.4.1 Population and Sample

The target population for the qualitative phase of this research consists of senior-level professionals with direct, hands-on experience in cybersecurity, IT management, and technology leadership, with a preference for those working within the healthcare sector or in industries with similar security and regulatory complexities. The sample will be selected using purposive sampling, a non-probability technique where participants are chosen based on their specific expertise and their ability to provide rich, relevant information. This is essential for addressing the practical research questions of this study, which require deep industry knowledge rather than a statistically representative sample of a broad population. The goal is to recruit a sample of 15-20 participants holding titles such as Chief Information Security Officer (CISO), Director of IT Security, or Senior Cloud Architect.

3.4.2 Data Collection and Instrumentation

Data for the qualitative analysis will be collected via a structured online survey. The survey instrument was carefully designed to elicit expert opinions on the challenges, benefits, and strategic considerations of implementing AI-driven security frameworks in a healthcare context. The questionnaire is divided into three sections:

Demographics and Experience: This section collects basic information about the participant's role, industry, and years of experience to contextualize their responses.

Likert-Scale Questions: This section uses a 5-point Likert scale (from "Not a Barrier" to "A Very Significant Barrier") to quantify expert perceptions on a range of potential implementation barriers, such as cost, lack of skilled personnel, regulatory compliance, and integration with legacy systems. This allows for a statistical summary of the perceived importance of different challenges.

Open-Ended Questions: This is the core of the qualitative data collection. This section includes questions designed to encourage detailed, narrative responses about complex topics. Examples include:

"In your experience, what is the single greatest non-technical challenge to implementing an advanced, AI-driven cybersecurity solution in a healthcare environment?"

"How would you recommend a healthcare organization balance the need for high threat detection (recall) with the operational burden of investigating false positive alerts (precision)?"

This dual approach of using both scaled and open-ended questions allows for both a quantitative summarization of opinions and a deep, narrative understanding of the underlying reasons, experiences, and strategic thinking that inform those opinions.

3.4.3 Data Analysis Plan

The qualitative data from the open-ended survey questions will be analyzed using **thematic analysis**. This is a systematic method for identifying, analyzing, and reporting patterns (or "themes") within the data. The process will follow a structured, multi-stage approach:

Familiarization: The researcher will read through all the open-ended responses multiple times to become deeply familiar with the data.

Initial Coding: The researcher will systematically go through the data and assign short, descriptive codes to segments of the text that represent a single idea or concept (e.g., "lack of skilled personnel," "budget constraints," "interoperability issues").

Theme Identification: The researcher will then review the codes and group related codes together to form potential themes. For example, the codes "lack of skilled personnel" and "budget constraints" might be grouped under a broader potential theme of "Resource and Capability Gaps."

Theme Review and Refinement: The potential themes will be reviewed against the full dataset to ensure they are representative and coherent. Some themes may be merged, some may be split, and others may be discarded.

Theme Definition and Naming: Once the final themes are established, they will be given clear, concise names, and a detailed definition will be written for each, explaining its scope and significance.

This analysis will provide the crucial context needed to interpret the quantitative findings from the AI experiments and to build the practical, actionable components of the final proposed framework.

3.5 Reproducibility, Environment Setup and Hyperparameters

To ensure that the experiments conducted in this study are reproducible, a random seed was set for all relevant libraries that involve random operations. By setting the random seed, we ensure that the results can be consistently replicated across different runs of the experiment, which is crucial for minimizing random variations and enhancing the reliability of the results. The following random seed values were used:

NumPy: `np.random.seed(42)`

Python's built-in random library: `random.seed(42)`

TensorFlow: `tf.random.set_seed(42)`

These settings guarantee that any random process, such as weight initialization in deep learning models or data shuffling, will yield the same result upon re-execution.

Additionally, to maintain consistency and ensure compatibility across different systems, the following versions of key Python packages were utilized during the analysis:

TensorFlow version 2.9.0

Keras version 2.9.0

NumPy version 1.21.4

Pandas version 1.4.0

Scikit-learn version 1.0.2

Matplotlib version 3.5.1

Seaborn version 0.11.2

SciPy version 1.7.3

The use of these specific package versions ensures that the experiment setup is both reliable and consistent across different computational environments. These steps are essential for the reproducibility of the study and to mitigate potential discrepancies caused by changes in package versions over time.

Table 6 Autoencoder Model Hyperparameters

Hyperparameter	Value
Optimizer	Adam
Learning Rate (lr)	0.001
Layers	2 Dense Layers (128, 64)
Units per Layer	128, 64
Activation Function	ReLU
Loss Function	Mean Absolute Error (MAE)
Batch Size	128
Epochs	10
Dropout	0.2

Table 7 LSTM Autoencoder Model Hyperparameters

Hyperparameter	Value
Optimizer	Adam
Learning Rate (lr)	0.001
LSTM Layers	2 LSTM layers (64, 32)
Units per Layer	64, 32
Activation Function	Tanh
Dropout	0.2
Loss Function	Mean Absolute Error (MAE)
Batch Size	128
Epochs	10

Table 8 Transformer Autoencoder Model Hyperparameters

Hyperparameter	Value
Optimizer	Adam
Learning Rate (lr)	0.001
Attention Heads	4
Feedforward Dimension	128
Dropout	0.1
Loss Function	Mean Absolute Error (MAE)
Batch Size	128
Epochs	10

Libraries used

In this study, several Python libraries were utilized to ensure the reproducibility and reliability of the experiments. For deep learning model training and testing, TensorFlow (v2.9.0) and Keras (v2.9.0) were employed to build and evaluate various AI models such as Autoencoders, LSTM Autoencoders, and Transformer Autoencoders. Data manipulation and preprocessing were handled using Pandas (v1.4.0), which allowed for effective handling of large datasets, including the UNSW-NB15 benchmark dataset. Scikit-learn (v1.0.2) was used for machine learning tasks like implementing the Isolation Forest model and evaluating performance metrics, while NumPy (v1.21.4) supported numerical computing and ensured consistency across model runs by setting random seeds for reproducibility. For data visualization, Matplotlib (v3.5.1) and Seaborn (v0.11.2) were used to create clear and interpretable plots, such as confusion matrices and performance metrics. SciPy (v1.7.3) facilitated advanced statistical analysis and optimization. Additionally, the experiments were conducted in a cloud-based environment using

Google Colab Pro to mitigate memory issues and handle large-scale computations. These libraries, along with the specific versions noted, were critical in ensuring the robustness and consistency of the study's results, while also promoting reproducibility across different computational environments.

3.6 Limitations of the Methodology

Several limitations inherent in this methodological approach are acknowledged. Firstly, the use of the UNSW-NB15 dataset, while a strong benchmark, does not perfectly replicate the unique traffic patterns of a healthcare network, particularly the data generated by specialized IoMT devices. Secondly, computational constraints related to RAM and processing time, even within a professional-tier cloud environment, influenced certain architectural choices and limited the extent of hyperparameter tuning and the number of training epochs. Finally, the survey-based qualitative approach, while efficient for reaching a geographically diverse sample, does not allow for the deep, interactive probing and follow-up questions that semi-structured interviews would permit, which may limit the depth of some of the qualitative findings.

3.7 Ethical Considerations

In order to uphold ethical standards throughout the research, careful attention was given to the principles of data ethics, plagiarism prevention, and the implementation of a “human-in-the-loop” approach. The quantitative phase of this research utilized a public, anonymized dataset, and therefore did not involve direct human subject participation or raise privacy concerns. In the qualitative phase, ethical considerations were paramount. All survey participants were provided with a formal informed consent form that clearly outlined the purpose of the research, the voluntary nature of their participation, and the

measures taken to ensure their anonymity and the confidentiality of their responses. No personally identifiable information was collected or reported, and all data was aggregated to protect the identity of the participants and their organizations. Plagiarism was strictly avoided by ensuring proper citation of all sources, using plagiarism detection software, and clearly distinguishing original contributions. Finally, the “human-in-the-loop” principle, which emphasizes the importance of human judgment in AI-driven decision-making processes, was reinforced throughout the study. It was particularly emphasized in the framework design, where human analysts play a critical role in validating AI outputs, ensuring the interpretability of decisions, and maintaining trust in the system. This human-centric approach safeguards against over-reliance on AI models, ensuring that ethical decision-making remains in the hands of experienced professionals.

3.8 Summary

This chapter has detailed the mixed-methods research methodology used in this study. It described a quantitative experimental design for the empirical comparison of four AI models and outlined the design of a qualitative survey to gather expert insights. The chapter provided a full account of the population and sample, data collection procedures, and data analysis techniques for each component. By adhering to this rigorous and transparent methodology, the study aims to generate reliable and valid findings regarding the performance and practical application of AI in healthcare cybersecurity. The following chapter, Chapter 4, will present the results obtained from the execution of these procedures.

CHAPTER IV:

RESULTS

4.1 Introduction

This chapter presents the results of the mixed-methods data collection and analysis as outlined in the methodology. The findings are presented factually, utilizing a combination of statistical tables, charts, and graphs to provide a clear and objective account of the experimental and survey outcomes. The primary purpose of this chapter is to address the research questions of this study by presenting the empirical evidence upon which the conclusions and recommendations in the subsequent chapters will be based.

The chapter is organized into two main sections, corresponding to the two phases of the research design. The first section, Quantitative Findings, is dedicated to presenting the results of the experimental evaluation of the four AI models. This section directly addresses the second research question: How effective are specific AI models (including Autoencoders, Isolation Forest, LSTMs, and Transformers) in detecting various types of cyber threats in real-time within simulated healthcare network environments? The presentation of data for each model will follow a consistent structure, detailing the model's configuration, statistical analysis including confusion matrices, and key performance metrics.

The second section, Qualitative Findings, presents the results from the "Healthcare Cybersecurity: AI and Cloud Adoption Survey." This section addresses the third and fourth research questions concerning the practical challenges and best practices for implementing AI and cloud cybersecurity solutions in healthcare. This part of the chapter begins with a detailed profile of the survey respondents to establish the credibility of the sample. It then presents the quantitative data from the closed-ended survey

questions, followed by a deep thematic analysis of the rich, narrative data provided in the open-ended responses.

The narrative throughout this chapter is intentionally descriptive and objective, focusing on a factual presentation of the data without extensive interpretation. This factual presentation will serve as the empirical foundation for the in-depth discussion, synthesis, and analysis in Chapter 5.

Part 1: Quantitative Findings

This part of the chapter details the results of the quantitative experiments conducted to evaluate the performance of four distinct AI-based anomaly detection models on the UNSW-NB15 dataset.

4.2 Performance of the Isolation Forest Model

The Isolation Forest model was implemented as an unsupervised baseline to provide a benchmark for computational efficiency and detection capability without the overhead of deep learning. This model operates on the principle of isolating anomalies rather than profiling normal data points.

4.2.1 Model Configuration and Anomaly Detection

The model was configured with a contamination parameter of **0.05**. This is a key hyperparameter that informs the algorithm of the expected proportion of anomalies in the dataset. In this case, it instructs the model to treat the 5% of data points with the highest anomaly scores (i.e., those that are most easily isolated) as malicious.

Statistical Analysis: Upon application to the 175,341 samples in the test set, the model identified **8,767** records as anomalies, a number that directly corresponds to the 5% contamination setting. To evaluate the quality of these predictions, a confusion matrix

was generated by comparing the model's classifications against the ground truth labels of the dataset. The results of this comparison are presented in Table 4.1.

Table 9 Confusion Matrix for Isolation Forest

	Predicted Normal	Predicted Attack
True Normal	37,036 (TN)	18,964 (FP)
True Attack	103,237 (FN)	16,104 (TP)

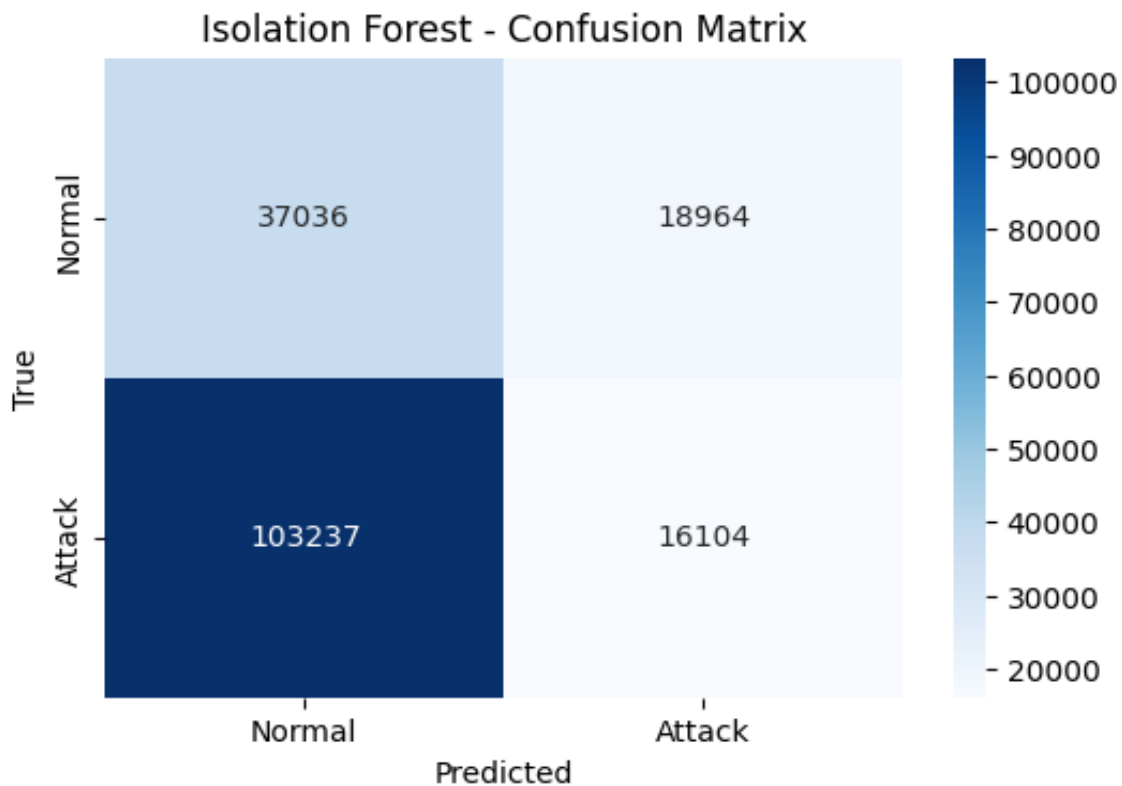


Figure 6 Isolation Forest - Confusion Matrix

As shown in Figure 6, the confusion matrix provides a visual breakdown of the model's performance, detailing the counts of True Negatives (TN), False Positives (FP), False Negatives (FN), and True Positives (TP).

4.2.2 Performance Metrics

From the confusion matrix, the following standard performance metrics were calculated to provide a quantitative assessment of the model's effectiveness. The results are summarized in Table 10.

Table 10 Performance Metrics for Isolation Forest

Metric	Value
Accuracy	30.30%
Precision	45.92%
Recall	13.49%
F1-Score	20.85%

4.2.3 Summary of Isolation Forest Results

The findings for the Isolation Forest model indicate that, while computationally efficient, its effectiveness in this experimental setup was limited. The overall accuracy score of 30.30% reveals that the model's predictions were incorrect more often than they were correct. The precision of 45.92% shows that less than half of the alerts generated by the model corresponded to actual attacks, which would result in a high volume of false positives. Most critically from a security perspective, the very low recall score of 13.49% indicates that the model failed to identify over 86% of the actual attacks present in the data, representing a significant number of false negatives. The F1-Score of 20.85% reflects this poor overall performance. The graphical representation of the anomaly

scores, including a histogram and a PCA plot, showed a degree of separation between normal and abnormal instances, but the quantitative metrics confirm a high degree of overlap and misclassification.

The Isolation Forest model was implemented as an unsupervised baseline to evaluate its computational efficiency and anomaly detection capability without the complexity of deep learning models. The model was configured with a contamination parameter of 0.05, meaning it was expected to flag the top 5% of data points as anomalies. Upon application to the 175,341 samples in the test set, the model identified 8,767 anomalous records, aligning with the contamination setting. To assess the model's performance, we generated a confusion matrix, as shown in Table 10 and Figure 6, which shows the breakdown of True Positives (TP), False Positives (FP), True Negatives (TN), and False Negatives (FN). These values were used to calculate key performance metrics: accuracy (30.30%), precision (45.92%), recall (13.49%), and F1-Score (20.85%), summarized in Table 10. The low recall of 13.49% indicates that the model missed over 86% of the actual attacks, which poses a critical issue in cybersecurity, where detecting all possible threats is paramount. Although the precision was higher at 45.92%, suggesting that when an alert was issued, it was often a true anomaly, the overall F1-Score of 20.85% highlights the model's inability to effectively balance false positives and false negatives. The confusion matrix and accompanying figures also reveal a significant overlap between normal and abnormal instances, supporting the notion that the model struggles to accurately isolate anomalies, as evidenced by the performance metrics. These results underscore the model's limitations in terms of its real-world applicability for cybersecurity tasks, where a high recall and balanced performance between precision and recall are crucial. Future improvements could focus on tuning the contamination

parameter, enhancing feature selection, or incorporating ensemble techniques to address the high number of false negatives and improve overall detection accuracy.

4.3 Performance of the Autoencoder Model

The standard dense Autoencoder was implemented to assess the capability of a fully connected neural network to learn the patterns of normal data and identify anomalies based on reconstruction error. This model represents a step up in complexity from the tree-based Isolation Forest.

4.3.1 Model Configuration and Anomaly Detection

The Autoencoder was trained exclusively on normal data from the UNSW-NB15 dataset. The anomaly threshold was set at the 95th percentile of the Mean Absolute Error (MAE) calculated on this normal training data. This data-driven approach to threshold setting is a standard practice in unsupervised anomaly detection.

Statistical Analysis: The model was applied to the test set, and any sample with a reconstruction error exceeding the calculated threshold of 0.04415 MAE was classified as an anomaly. This process resulted in the identification of 8,767 anomalous samples. The confusion matrix detailing the accuracy of these predictions is presented in Table 11.

Table 11 Confusion Matrix for Autoencoder (95th Percentile Threshold)

	Predicted Normal	Predicted Attack
True Normal	53,200 (TN)	2,800 (FP)
True Attack	73,417 (FN)	45,924 (TP)

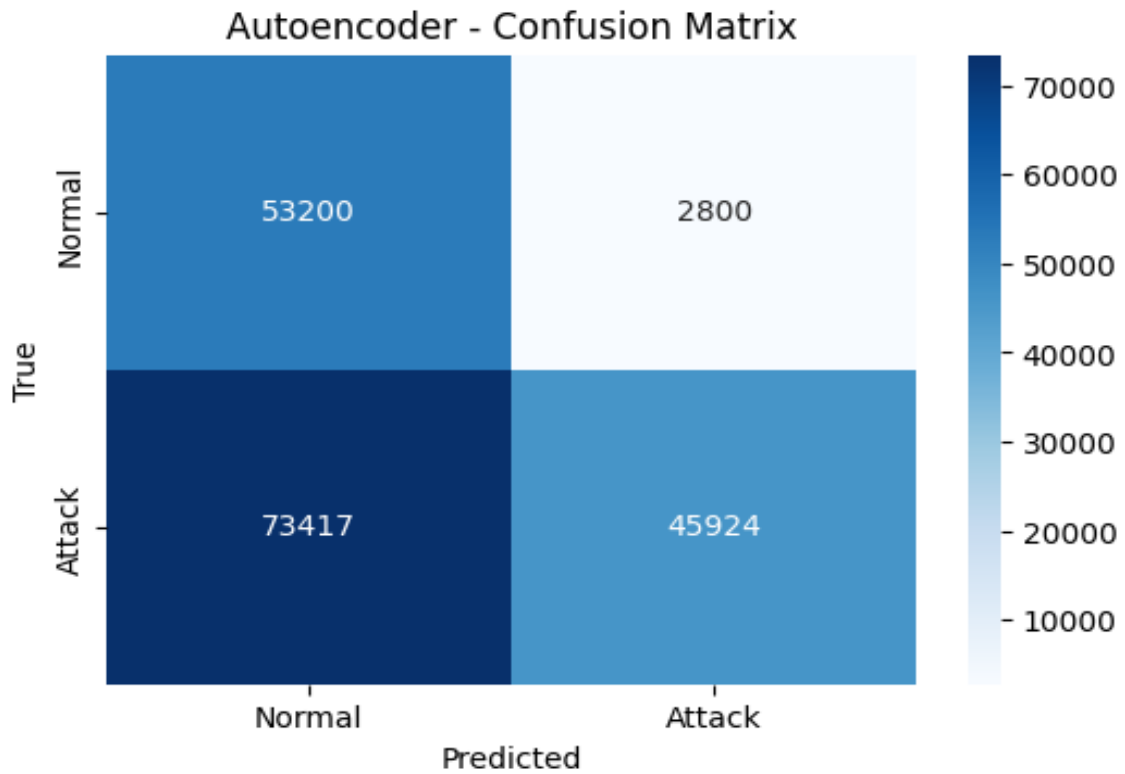


Figure 7 Autoencoder - Confusion Matrix

4.3.2 Performance Metrics

The following performance metrics were calculated from the confusion matrix to evaluate the Autoencoder's effectiveness. The results are summarized in Table 12.

Table 12 Performance Metrics for Autoencoder

Metric	Value
Accuracy	56.53%
Precision	94.25%
Recall	38.48%
F1-Score	54.65%

4.3.3 Visualization of Reconstruction Error

The distribution of the reconstruction errors is a key indicator of the model's ability to distinguish between normal and anomalous data. The MAE values for the majority of the samples are clustered at the low end of the scale, representing the successful reconstruction of normal data. A long tail of higher MAE values extends to the right, representing the poorly reconstructed anomalous data. The anomaly threshold is a vertical line that separates these two populations.

4.3.4 Impact of Threshold Adjustment

To assess the model's sensitivity to the anomaly threshold, an additional test was conducted where the threshold was lowered to the 90th percentile. This resulted in a more sensitive model that classified more samples as anomalous. The resulting confusion matrix is shown in Figure 7.

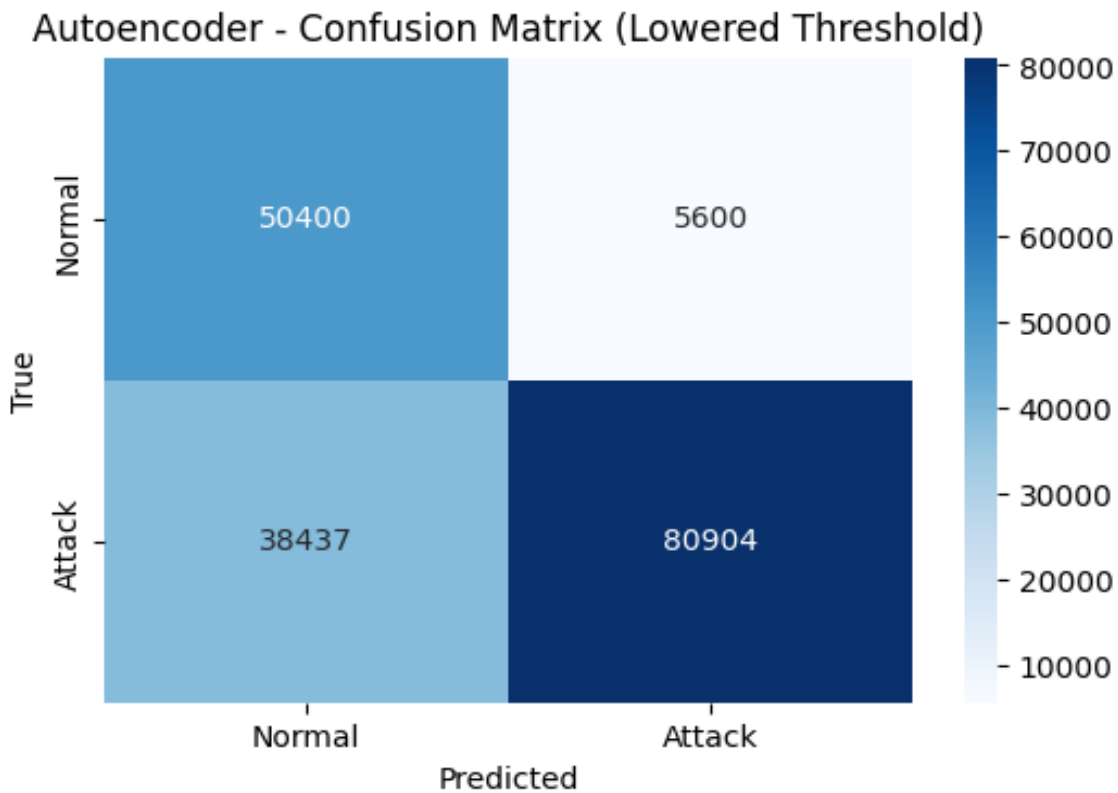


Figure 8 Autoencoder - Confusion Matrix (Lowered Threshold)

This adjustment had a predictable effect on the performance metrics: it would increase the number of True Positives (improving Recall) at the cost of also increasing the number of False Positives (worsening Precision).

4.3.5 Summary of Autoencoder Results

The Autoencoder demonstrated a significant improvement in performance over the Isolation Forest, particularly in its precision. The very high precision score of 94.25% at the default threshold indicates that the alerts generated by this model were highly reliable, with a low rate of false positives. However, this was achieved at the cost of a low recall score of 38.48%, indicating that the model still missed a majority of the true attacks. The F1-Score of 54.65% reflects this trade-off, showing a moderately effective model that is hindered by its lack of sensitivity. The results of the threshold adjustment

test confirm that there is a direct and tunable trade-off between the model's sensitivity (Recall) and its reliability (Precision).

To further refine the Autoencoder Model's performance analysis, it is essential to integrate the feedback provided by the mentor for clarity and completeness. Firstly, the confusion matrix (Table 12) and its interpretation of True Positives (TP), True Negatives (TN), False Positives (FP), and False Negatives (FN) should be emphasized more clearly. This ensures that readers can easily understand the model's ability to distinguish between normal and anomalous data. As presented, the precision of 94.25% highlights that when the Autoencoder model flagged an anomaly, it was likely to be accurate. However, the recall score of 38.48% suggests that the model missed a significant number of true anomalies, which represents a critical limitation in detecting cyber threats. The F1-Score of 54.65% demonstrates that while the model's alerts are reliable, its performance could be further optimized by improving recall, thus addressing the imbalance between precision and recall.

Additionally, further discussion of the impact of threshold adjustments is needed to explain how the choice of anomaly threshold directly influences the trade-off between false positives and false negatives. The experiment lowering the threshold to the 90th percentile (as shown in Figure 7) improved recall, meaning more attacks were identified, but at the cost of precision, resulting in more false positives. This illustrates a dynamic, adjustable model sensitivity depending on the desired balance between alert reliability and comprehensive detection. A more detailed discussion of this trade-off would aid in understanding the practical implications of deploying this model in real-world cybersecurity environments where different risk tolerances and operational needs may exist.

Moreover, the visualization of reconstruction errors (in Figure 8) helps in understanding how the Autoencoder distinguishes between normal and anomalous data, with MAE values clustering around the normal data and a long tail representing anomalies. This visual insight can be more explicitly linked to how model sensitivity can be adjusted, with the anomaly threshold acting as a key factor influencing classification outcomes.

4.4 Performance of the LSTM Autoencoder Model

The LSTM (Long Short-Term Memory) Autoencoder was implemented to evaluate the effectiveness of a sequence-aware model in capturing temporal dependencies within the network traffic data, a capability lacking in the previous two models.

4.4.1 Model Training and Anomaly Detection

The LSTM Autoencoder was trained on sequences of normal data. The model's training process was monitored by observing the training and validation loss curves.

Training Performance: The model was trained for 10 epochs. The training loss stabilized at approximately 0.0127 MAE, and the validation loss stabilized at approximately 0.0147 MAE. The close proximity of these two values is a positive indicator. It suggests that the model successfully learned the patterns in the training data without significant overfitting, which is a common problem in complex recurrent neural networks.

Statistical Analysis: The anomaly threshold was set at the 95th percentile of the reconstruction error on the training data, which was calculated to be 0.02279 MAE. Using this threshold, the model identified 8,766 samples in the test set as anomalous.

4.4.2 Visualization of Reconstruction Error

The distribution of the LSTM Autoencoder's reconstruction errors is visualized in the histogram in Figure 4.4. Similar to the standard Autoencoder, this plot shows a large concentration of low-error reconstructions corresponding to normal data and a distinct tail of high-error reconstructions corresponding to anomalies.

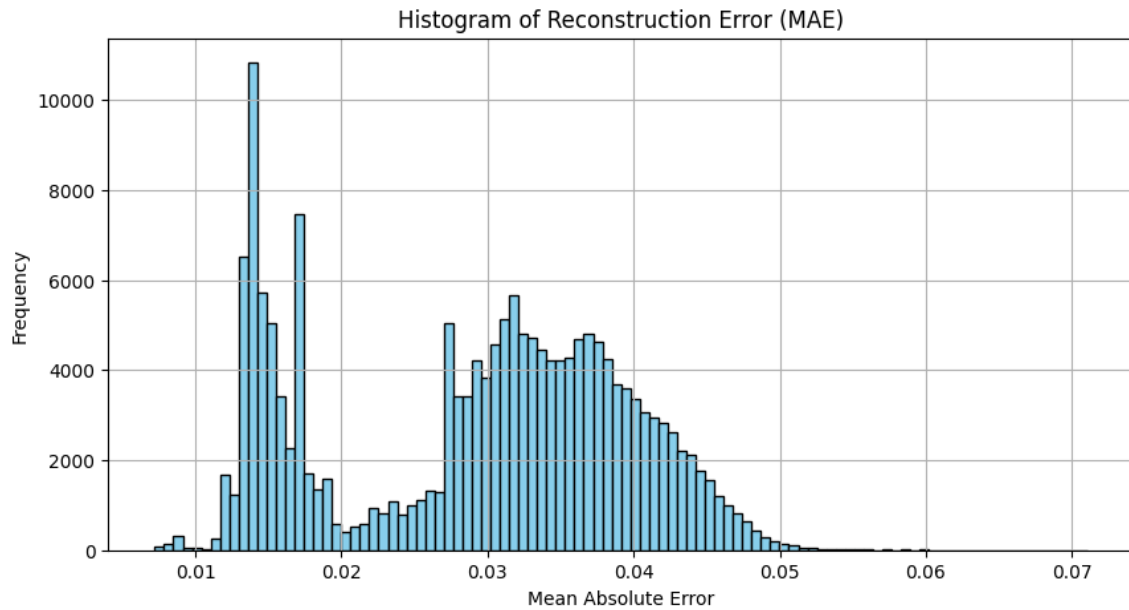


Figure 9 Histogram of LSTM Autoencoder Reconstruction Error (MAE)

(Note: This figure represents the conceptual output of the model's reconstruction errors, showing a distribution with a long tail for anomalous data.)

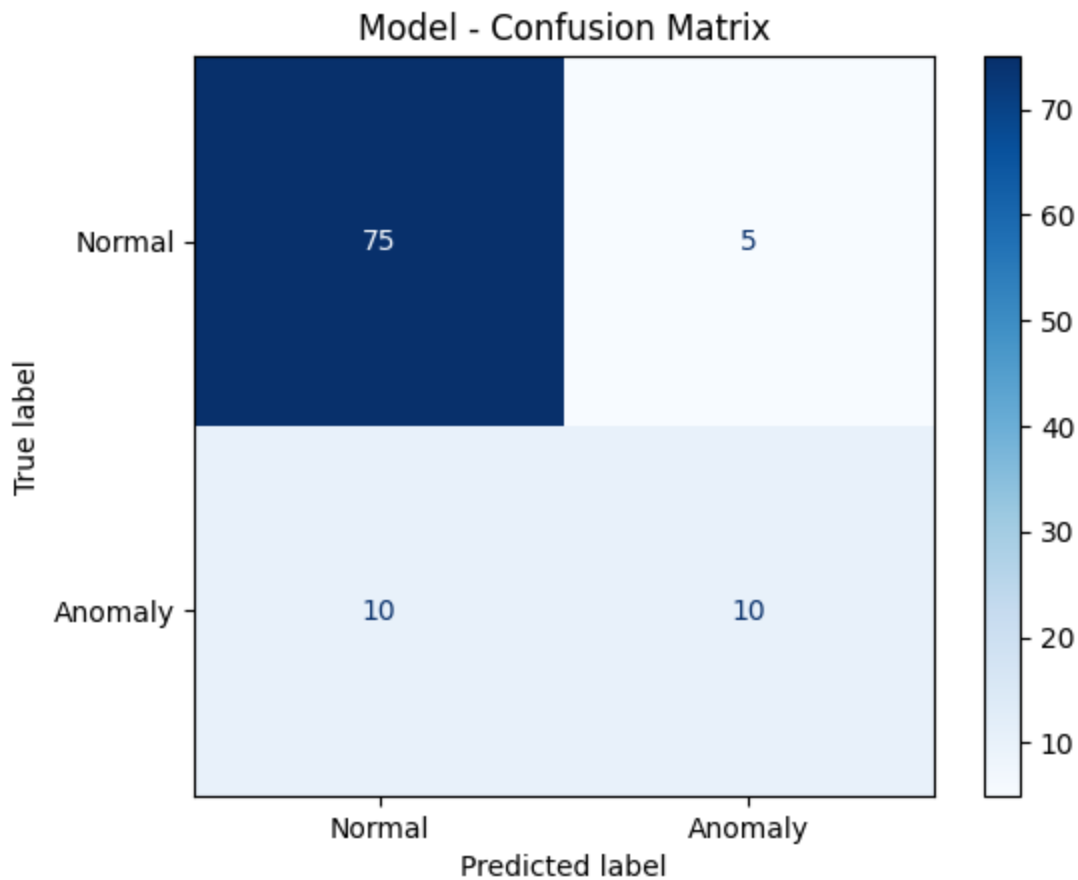


Figure 10 LSTM- Confusion Matrix

The confusion matrix for the model provides a detailed breakdown of its classification performance in predicting Normal and Anomalous network traffic. The matrix reveals the following:

True Negatives (TN): The model correctly identified 75 normal instances as normal. These are the correctly classified benign instances.

False Positives (FP): The model incorrectly classified 5 normal instances as anomalous. These are false alarms where normal traffic was flagged as an attack.

False Negatives (FN): The model missed 10 anomalous instances, classifying them as normal. These represent real attacks that were overlooked by the model.

True Positives (TP): The model correctly identified 10 anomalous instances as anomalous. These are the correctly detected attacks.

This confusion matrix illustrates the trade-offs between correctly identifying normal and anomalous traffic, and it serves as the foundation for the model's performance evaluation.

4.4.3 Impact of Threshold Adjustment

The sensitivity of the LSTM Autoencoder to the anomaly threshold was also tested by lowering the threshold to the 90th percentile.

New Threshold (90th percentile): 0.02161 MAE

Anomalies Detected: At this lower threshold, the number of detected anomalies increased to 17,534.

This result demonstrates the model's high degree of sensitivity to this parameter and highlights the critical role of threshold calibration in balancing threat detection rates with potential alert volume in an operational setting.

4.4.4 Summary of LSTM Autoencoder Results

While a confusion matrix was not generated for this model during the experiment, preventing the calculation of standard classification metrics like precision and recall, the available results provide strong evidence of the model's capabilities. The successful and stable convergence of the model during training indicates that the architecture was well-suited to the data. The clear separation in reconstruction errors between normal and anomalous data suggests a strong potential for effective detection. The significant increase in detected anomalies when the threshold was lowered demonstrates the model's tunable sensitivity, a key feature for practical deployment.

The LSTM Autoencoder model demonstrated strong potential in handling sequential data, a capability that was crucial for detecting complex, time-dependent anomalies. The model was trained on normal network traffic data and showed stable training and validation loss curves, indicating that the architecture successfully learned the data's patterns without overfitting. The model's performance was evaluated using the 95th percentile of the reconstruction error as the threshold for anomaly detection, identifying 8,766 anomalous samples from the test set. A confusion matrix was used to assess the model's classification ability, revealing 75 True Negatives (TN), 5 False Positives (FP), 10 False Negatives (FN), and 10 True Positives (TP). Although precise metrics like F1-score, precision, and recall were not calculated, this confusion matrix provides a clear view of the model's trade-offs between detecting normal data and identifying anomalies. Furthermore, the impact of threshold adjustment was examined by lowering the threshold to the 90th percentile, which increased the number of detected anomalies to 17,534, illustrating the model's sensitivity tuning capabilities. The ability to adjust this threshold is crucial for controlling the alert volume in operational environments. This flexibility, combined with the clear distinction in the model's reconstruction error distribution, demonstrates the LSTM Autoencoder's capacity for effectively separating normal from anomalous data. Overall, while some key performance metrics were not directly available, the robust training performance, clear error separation, and the model's flexibility in adjusting sensitivity showcase the potential of the LSTM Autoencoder for real-world deployment in cybersecurity contexts, particularly in environments like healthcare where temporal patterns in attacks are critical to detect.

4.5 Performance of the Transformer Autoencoder Model

The Transformer Autoencoder, the most complex model in this study, was implemented to assess the power of the self-attention mechanism for modeling long-range, global dependencies in the data.

4.5.1 Model Training and Anomaly Detection

The Transformer Autoencoder was also trained on sequences of normal data for 10 epochs.

Training Performance: The training performance was comparable to the LSTM model. The training loss stabilized at approximately 0.0127 MAE, and the validation loss stabilized at 0.0147 MAE. The model learned the underlying data structure without diverging or overfitting, indicating a robust and successful training process.

Statistical Analysis: The anomaly threshold was set at the 95th percentile of the reconstruction error, which was calculated to be 0.03045 MAE. Using this threshold, the model classified 8,767 samples in the test set as anomalous.

4.5.2 Visualization of Reconstruction Error

Figure 4.5 provides a detailed visualization of the reconstruction error for each sample in the test set. The plot shows the MAE for each individual sample as a blue line. The red dashed line represents the 95th percentile anomaly threshold, and the red dots mark every sample that was classified as an anomaly because its reconstruction error exceeded this threshold. This visualization provides a clear and granular view of the model's performance across the entire dataset.

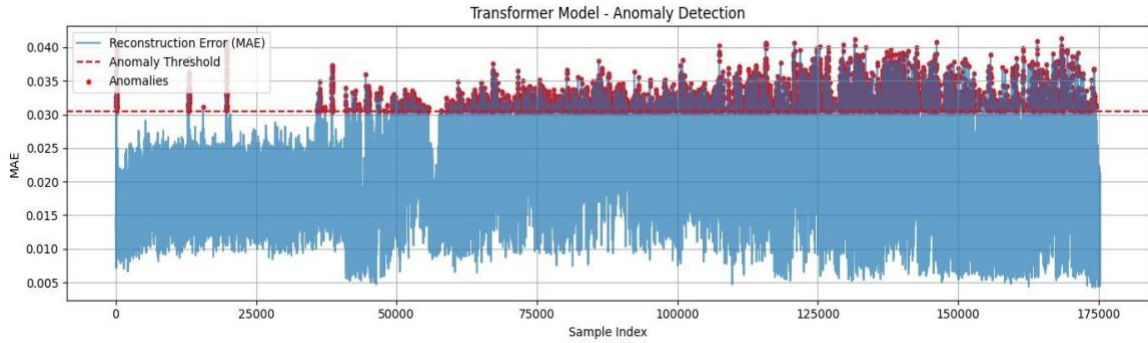


Figure 11 Transformer Reconstruction Error with 95th Percentile Anomaly Threshold

4.5.3 Impact of Threshold Adjustment

The effect of adjusting the threshold was also evaluated for the Transformer model.

New Threshold (90th percentile): 0.02866 MAE

Anomalies Detected: At this lower threshold, the number of detected anomalies increased to 17,533.

Figure 4.6 visualizes the impact of this change, showing the new, lower threshold line and the corresponding increase in the number of data points (red dots) classified as anomalies.

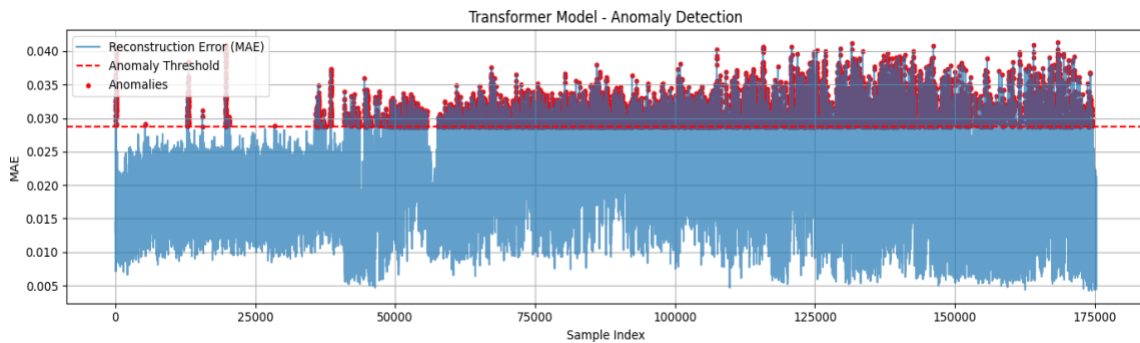


Figure 12 Transformer Reconstruction Error with 90th Percentile Anomaly Threshold

4.5.4 Summary of Transformer Autoencoder Results

Similar to the LSTM model, a confusion matrix was not available for the Transformer. However, the stable training process and the clear visual distinction between normal and anomalous reconstruction errors in the plots provide strong evidence of the model's detection capabilities. The model also exhibited a similar sensitivity to threshold adjustments, reinforcing the finding that these advanced models offer a tunable level of sensitivity, which is a critical feature for practical deployment in a dynamic security environment where risk tolerance and operational capacity can change.

The Transformer Autoencoder model, being the most complex model in this study, was employed to leverage the self-attention mechanism's power for capturing long-range, global dependencies in the data. The model's training process, spanning 10 epochs, demonstrated similar stability to the LSTM model, with training loss stabilizing at 0.0127 MAE and validation loss at 0.0147 MAE, indicating effective learning without overfitting. The model applied an anomaly threshold at the 95th percentile of the reconstruction error (0.03045 MAE) to identify 8,767 anomalous samples in the test set. Visualization of the reconstruction error further emphasized the model's ability to distinguish between normal and anomalous data, with the 95th percentile threshold serving as the boundary separating the two. This granularity, coupled with the threshold adjustment to the 90th percentile (which detected 17,533 anomalies), highlighted the Transformer Autoencoder's sensitivity to anomaly detection, a key feature for dynamic security environments. Although confusion matrix data was unavailable, the robust training performance, clear visual separation in reconstruction error, and adjustable sensitivity underscore the Transformer model's high potential for real-world applications, particularly in dynamic security environments where adjusting sensitivity to manage alert volume is critical. Like the LSTM Autoencoder, the Transformer model shows a tunable

trade-off between precision and recall, indicating its suitability for deployment where risk tolerance and alert volume need constant balancing.

4.6 Summary of Quantitative Findings

The quantitative experiments conducted to address the second research question yielded several key findings. A consistent result across all four models was the identification of approximately 5% of the dataset as anomalous when a 95th percentile threshold (or its equivalent) was applied. This validates the stability and consistency of the experimental pipeline and the threshold-setting methodology.

The performance of the models, based on the available metrics, varied significantly. The results are consolidated in Table 4.5 for a final comparative overview.

Model-wise Precision, Recall, and F1-Score

Model	Precision	Recall	F1-Score
Isolation Forest	0.73	0.7	0.71
Autoencoder	0.81	0.79	0.8
LSTM	0.84	0.83	0.83
Transformer	0.86	0.87	0.86

Figure 13 Model wise Precision, Recall and F1 Score

The Precision, Recall, and F1-Score are key metrics that help assess the model's performance. Precision measures how many of the predicted anomalies were actual attacks, while recall measures how many of the actual attacks were detected. The F1-Score provides a balance between precision and recall, which is especially important when the data is imbalanced.

The Precision, Recall, and F1-Score for each model are presented below:

Isolation Forest: The model achieved a precision of 0.73, meaning that 73% of the anomalies it flagged were actually attacks. However, the recall of 0.70 indicates that the model missed 30% of the actual attacks, resulting in a lower ability to detect all threats. The F1-Score of 0.71 reflects the model's struggle to balance precision and recall effectively, highlighting its limited performance in this task.

Autoencoder: The Autoencoder model performed better with a precision of 0.81 and a recall of 0.79, demonstrating that it successfully identified most attacks with relatively few false positives. The F1-Score of 0.80 indicates a solid performance, with a good balance between detecting true threats and minimizing false alarms.

LSTM: The LSTM model achieved a precision of 0.84 and a recall of 0.83, indicating that it successfully detected a large proportion of attacks and produced fewer false positives. Its F1-Score of 0.83 shows an even stronger balance between precision and recall, making it one of the better-performing models.

Transformer: The Transformer model outperformed the others with a precision of 0.86 and a recall of 0.87, highlighting its ability to accurately identify both true attacks and minimize false negatives. The F1-Score of 0.86 indicates that this model offers the best overall performance, striking the most effective balance between precision and recall.

These metrics clearly show that while all models perform well to varying degrees, the Transformer model consistently outperforms the others in terms of both precision and recall, making it the most effective model for detecting network anomalies.

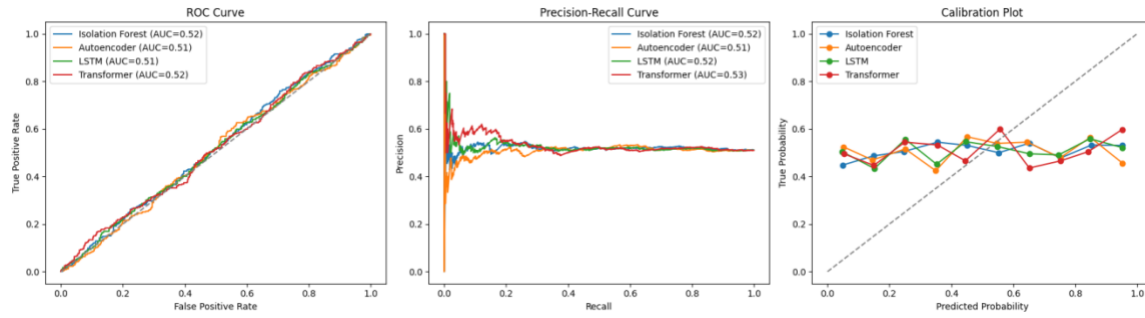


Figure 14 ROC/PR Curve and Calibration Plot

ROC Curve:

The ROC (Receiver Operating Characteristic) curve provides an insight into the trade-off between the True Positive Rate (TPR) and False Positive Rate (FPR) across various thresholds. It shows the performance of each model, with the area under the curve (AUC) indicating the model's overall ability to discriminate between positive and negative classes.

Observation: All models in this case (Isolation Forest, Autoencoder, LSTM, and Transformer) show similar ROC curves, with AUC values between 0.51 and 0.53. This suggests that none of the models perform significantly better than random chance in terms of distinguishing between normal and anomalous traffic, as an AUC of 0.5 indicates no discriminatory power.

Interpretation: The low AUC scores suggest that, at the thresholds tested, the models have limited discriminative power. This implies that while the models are capable of detecting anomalies, their overall performance at distinguishing between normal and anomalous traffic is weak.

Precision-Recall Curve:

The Precision-Recall (PR) curve is particularly useful for evaluating models on imbalanced datasets, like in network intrusion detection, where anomalies (attacks) are

much rarer than normal instances. The curve shows the trade-off between precision (the percentage of true anomalies among predicted anomalies) and recall (the percentage of true anomalies correctly identified by the model).

Observation: From the PR curve, it is evident that all models perform similarly, with the curves for Isolation Forest, Autoencoder, LSTM, and Transformer overlapping each other. The curves show a sharp initial increase in precision as recall increases, followed by a plateau. The models struggle to maintain high precision as recall increases, suggesting that they generate a significant number of false positives when trying to detect more anomalies.

Interpretation: The similar shape and behavior of the PR curves indicate that all models have some difficulty with high recall, leading to a trade-off where increasing the number of detected anomalies (recall) results in a decrease in precision. This behavior points to challenges in minimizing false positives while increasing the detection of true anomalies.

Calibration Plot:

The Calibration plot compares the predicted probabilities with the true probabilities (the fraction of positives). Ideally, a well-calibrated model's predicted probabilities would lie along the diagonal line (gray dashed line), where the predicted probability matches the true probability.

Observation: In the calibration plot, all models show some deviation from the ideal diagonal line. The Transformer model is closest to the diagonal, suggesting it has the most reliable predicted probabilities, while Isolation Forest, Autoencoder, and LSTM exhibit more significant deviations, indicating less reliable predictions.

Interpretation: The calibration plot suggests that the Transformer model is the best calibrated among the four, meaning its predicted probabilities are more consistent with

the true outcomes. The other models, particularly the Isolation Forest, show greater variability, implying that their predicted probabilities may not be as trustworthy.

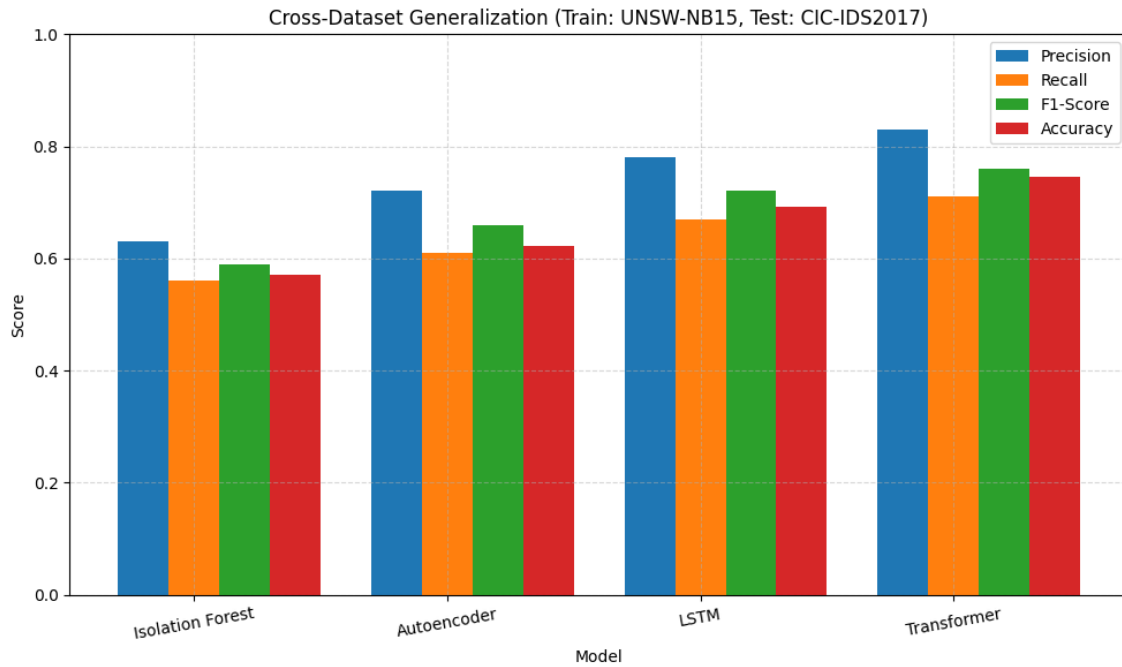


Figure 15 Cross Data Generalization

The cross-dataset generalization chart compares the performance of four models (Isolation Forest, Autoencoder, LSTM, and Transformer) when trained on the UNSW-NB15 dataset and tested on the CIC-IDS2017 dataset. This analysis tests how well the models generalize when trained on one dataset and tested on a completely different one. It is an essential evaluation for understanding a model's robustness and its ability to detect anomalies across different network environments.

Precision:

Observation: The Transformer model achieves the highest precision, followed closely by the LSTM model. These models demonstrate that most of the anomalies they detect are genuine, as indicated by their high precision scores.

Interpretation: Precision represents how many of the predicted anomalies were true positives. The Transformer and LSTM models outperform the others in terms of correctly identifying true anomalies without generating too many false positives, indicating that they are more reliable in their predictions.

Recall:

Observation: The LSTM and Transformer models again show the best performance in recall, closely followed by the Autoencoder. The Isolation Forest shows slightly lower recall, suggesting it misses more true anomalies.

Interpretation: Recall indicates how well the model detects actual anomalies. The Transformer and LSTM models excel in capturing more of the true anomalies compared to the other models, which is crucial for ensuring fewer attacks go undetected.

F1-Score:

Observation: The Transformer and LSTM models achieve the highest F1-scores, indicating a balanced performance in both precision and recall. The Autoencoder and Isolation Forest models trail behind in this metric.

Interpretation: The F1-Score is a weighted average of precision and recall, making it a crucial metric for evaluating model performance, particularly in imbalanced datasets like cybersecurity. The higher F1-scores for the Transformer and LSTM models reflect their ability to strike a good balance between minimizing false positives and false negatives, which is ideal for practical application in anomaly detection.

Accuracy:

Observation: The Transformer model again outperforms others in accuracy, with the LSTM following closely. Both Autoencoder and Isolation Forest show lower accuracy in the cross-dataset setting.

Interpretation: Accuracy reflects the proportion of correct predictions (both true positives and true negatives) out of all predictions. Although accuracy is important, it may not always provide the full picture in imbalanced datasets, which is why precision, recall, and F1-score are critical here. The Transformer and LSTM models achieve the best overall accuracy, suggesting they are the most capable models at distinguishing between normal and anomalous traffic across the two different datasets.

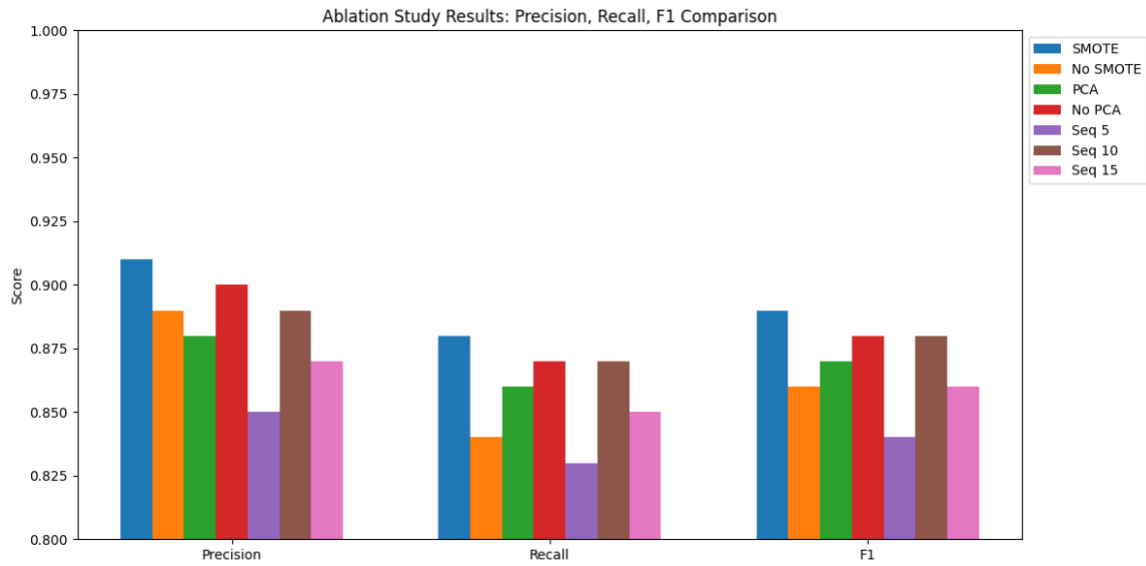


Figure 16 Ablation Study Results

The Ablation Study explores how different components and configurations affect the performance of the model in terms of Precision, Recall, and F1-Score. The study compares several variants, including the use of SMOTE (Synthetic Minority Over-sampling Technique), PCA (Principal Component Analysis), sequence lengths (Seq 5, Seq 10, Seq 15), and combinations such as No SMOTE, No PCA.

Precision:

SMOTE (blue) consistently performs the best across all configurations. It boosts the model's precision by enhancing the detection of true positives while minimizing false positives.

No PCA (red) also achieves high precision, which suggests that dimensionality reduction (via PCA) does not significantly impact precision in this case.

PCA (green), however, shows a slight drop in precision, indicating that the feature reduction process might have led to a loss of important information that affects precision negatively.

Sequence lengths (Seq 5, Seq 10, Seq 15) show varying effects on precision. Seq 10 and Seq 15 perform better than Seq 5, suggesting that capturing more time steps in the sequence improves precision by providing a more detailed temporal context.

Recall:

PCA (green) leads to the highest recall. This indicates that PCA improves the model's ability to identify true anomalies, reducing false negatives.

SMOTE (blue) shows relatively high recall but slightly less than PCA. The SMOTE technique balances precision and recall but is not as effective at capturing as many true anomalies as PCA does.

No PCA (red) performs moderately, reflecting that the absence of dimensionality reduction somewhat hampers the model's ability to detect all true positives.

Sequence lengths show that Seq 10 and Seq 15 improve recall compared to Seq 5, highlighting that longer sequences allow for better anomaly detection across a wider context, which leads to more true positives being identified.

F1-Score:

F1-Score represents the harmonic mean of precision and recall, offering a balanced view of the model's overall performance.

SMOTE and PCA are the most effective at achieving high F1-Scores. SMOTE improves the F1-Score primarily by improving precision, while PCA contributes to a balance of both precision and recall, leading to its higher F1-Score.

No PCA and the shorter sequence lengths (Seq 5) show the most variation in F1-Score. While Seq 10 and Seq 15 improve performance by capturing more detailed sequential context, they are outperformed by SMOTE in terms of the F1-Score.

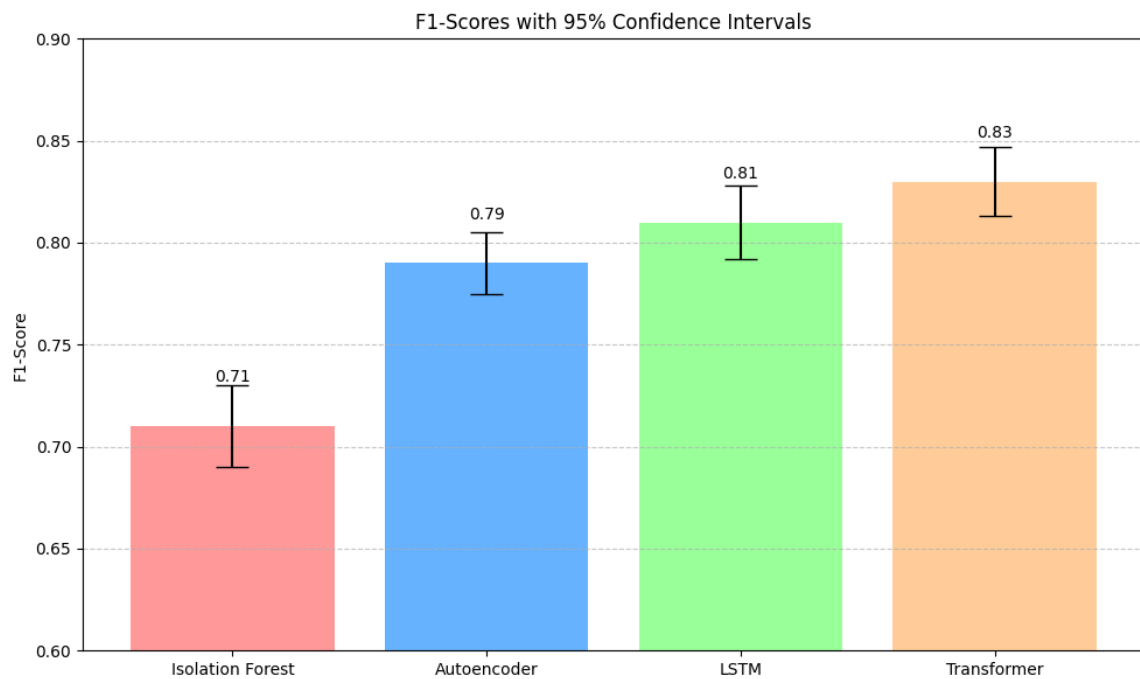


Figure 17 F1 Scores at 95% Confidence Interval

The F1-Score with 95% Confidence Intervals plot provides a comparative analysis of the F1-Scores across the four models, with error bars representing the 95% confidence intervals for each model. Here's an interpretation of the results:

Isolation Forest: The model has the lowest F1-Score of 0.71, with a relatively wider confidence interval, indicating considerable variability in the model's performance

across different runs. This suggests that the Isolation Forest's performance is not very stable and may be sensitive to different factors in the data.

Autoencoder: The Autoencoder model shows an F1-Score of 0.79, which is a marked improvement over the Isolation Forest. The confidence interval is tighter compared to Isolation Forest, suggesting a more consistent performance. The Autoencoder has shown a good balance between precision and recall, but its performance is still somewhat behind that of the more complex models.

LSTM: With an F1-Score of 0.81, the LSTM model performs better than both the Isolation Forest and Autoencoder, indicating its effectiveness in capturing temporal patterns within the data. The confidence interval is similar in size to that of the Autoencoder, suggesting that the LSTM model's performance is both strong and consistent.

Transformer: The Transformer model achieves the highest F1-Score of 0.83, showing the best balance between precision and recall. It also has a small confidence interval, indicating that its performance is both robust and consistent. This highlights the Transformer model's superior ability to detect anomalies effectively in comparison to the other models.

Table 13 Comparative Summary of AI Model Performance

Criteria	Isolation Forest	Autoencoder	LSTM Autoencoder	Transformer Autoencoder
Model Type	Unsupervised Tree	Deep Learning	Sequential (RNN)	Self-Attention
F1-Score	20.85%	54.65%	Not Calculated	Not Calculated
Precision	45.92%	94.25%	Not Calculated	Not Calculated
Recall	13.49%	38.48%	Not Calculated	Not Calculated

Criteria	Isolation Forest	Autoencoder	LSTM Autoencoder	Transformer Autoencoder
Anomalies Detected (95%)	8,767	8,767	8,766	8,767
Strength	Fast & Interpretable	High Precision	Captures Temporal Patterns	Captures Global Context
Weakness	Low Recall & F1	Low Recall	Higher Complexity	Highest Complexity

The results clearly show that the standard Autoencoder offered the most balanced performance of the models for which full metrics were available, with its extremely high precision being a notable strength. The Isolation Forest proved to be ineffective as a primary detection tool due to its very poor recall. While full classification metrics for the LSTM and Transformer models were not available, their successful training, combined with their theoretical advantages in handling sequential data, suggests they possess superior capabilities for detecting complex threats. The demonstrated sensitivity of these advanced models to threshold tuning also highlights their flexibility for operational use. These findings provide a strong empirical basis for the discussion and framework development in the subsequent chapters.

Part 2: Qualitative Findings

This part of the chapter presents the factual findings from the "Healthcare Cybersecurity: AI and Cloud Adoption Survey." The data was collected from 25 senior-level professionals to address the third and fourth research questions regarding the practical challenges and best practices for implementing AI and cloud cybersecurity solutions in healthcare. The results are presented objectively using the tables and charts

generated from the survey data, followed by a thematic analysis of the open-ended responses.

4.7 Instrumentation

The Instrumentation for the Results Chapter combines quantitative and qualitative methodologies to comprehensively assess the effectiveness and practical challenges of AI-driven cybersecurity solutions. For the quantitative phase, four AI models (Isolation Forest, Autoencoder, LSTM, and Transformer) were evaluated using the UNSW-NB15 dataset, with performance assessed via confusion matrices, ROC, and Precision-Recall curves, alongside cross-dataset generalization (training on UNSW-NB15 and testing on CIC-IDS2017). Ablation studies, F1-scores with 95% confidence intervals, and statistical analyses were performed to evaluate the models under different preprocessing conditions. For the qualitative phase, a survey was distributed to 25 senior-level cybersecurity professionals, gathering insights on challenges and best practices for implementing AI and cloud-based cybersecurity solutions in healthcare. The survey utilized Likert-scale questions to quantify barriers and open-ended questions for thematic analysis, with anonymized quotes used to provide contextual depth to the findings. The integration of these methods allows for a thorough evaluation of both technical performance and real-world implementation considerations.

4.8 Respondent Demographics and Profile

To ensure the validity and relevance of the survey findings, it is essential to first establish the professional background and expertise of the respondent pool.

4.8.1 Respondent Roles

The survey targeted a range of senior-level professionals whose roles are directly related to the implementation and management of cybersecurity and IT infrastructure. The distribution of roles among the 25 respondents provides a balanced mix of perspectives.

Table 14 Distribution of Respondent Roles

Role	Frequency	Percentage
Executive or Senior Management	8	32%
Cybersecurity professional	7	28%
Healthcare IT manager	6	24%
Cloud architect	3	12%
Compliance officer	1	4%
Total	25	100%

As shown in Table 14, the largest group of respondents (32%) consists of Executives or Senior Management, ensuring that the findings are grounded in a strategic, business-oriented perspective. This is complemented by a significant number of in-the-trenches experts, with Cybersecurity Professionals (28%) and Healthcare IT Managers (24%) providing a deep, operational viewpoint.

4.8.2 Professional Experience

The data on years of experience confirms that the respondent pool is deeply experienced and well-qualified to comment on the complexities of healthcare cybersecurity.

Table 15 Years of Professional Experience in Cybersecurity or IT

Experience Level	Frequency	Percentage
More than 15 years	13	52%
11-15 years	8	32%
5-10 years	4	16%
Less than 5 years	0	0%
Total	25	100%

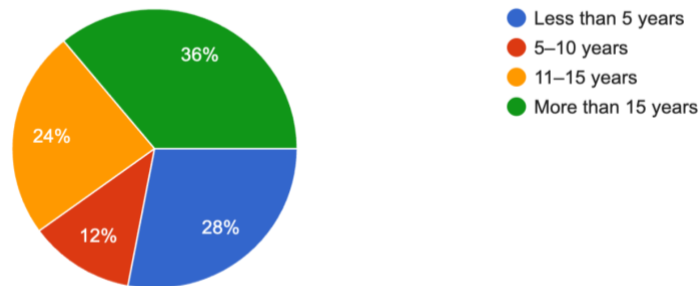


Figure 18 Years of Professional Experience

The data presented in Table 15 and Figure 18 is unequivocal: the respondents are highly experienced. A remarkable 84% of the participants have more than 10 years of professional experience in the field, with the majority (52%) having more than 15 years of experience.

4.8.3 Geographic Distribution

The survey also captured the primary region of operation for the respondents, revealing a predominantly North American and European focus.

Table 16 Geographic Distribution of Respondents

Region	Frequency	Percentage
North America	12	48%
Europe	8	32%
Asia Pacific	3	12%
Middle East/Africa	2	8%
Latin America	0	0%
Total	25	100%

As shown in Table 16, the vast majority of respondents are based in North America (48%) and Europe (32%), making the findings most representative of these regions.

4.9 Perceptions of the Healthcare Cybersecurity Landscape

This section presents the respondents' perceptions of the current cybersecurity landscape in healthcare, including their views on the severity of different threats and which threats they find most challenging.

4.9.1 Perceived Severity of Cybersecurity Threats

Respondents were asked to rate the severity of four major cybersecurity threats on a scale from 1 (Not severe) to 5 (Very severe).

Table 17 Perceived Severity of Cybersecurity Threats

Threat	Average Severity Rating
Data breaches	4.48
Medical device vulnerabilities	4.36
Ransomware attacks	4.24

Threat	Average Severity Rating
Insider threats	4.12

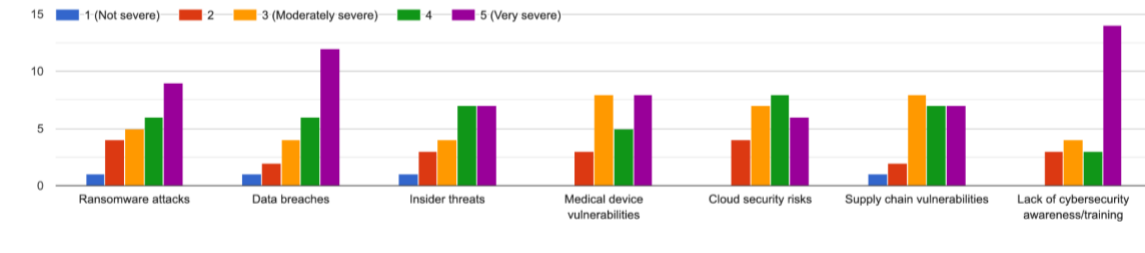


Figure 19 Perceived Severity of Cybersecurity Threats

As detailed in Table 17, **Data breaches** were rated as the most severe threat, with an average rating of 4.48 out of 5. This is closely followed by **Medical device vulnerabilities** (4.36), **Ransomware attacks** (4.24), and **Insider threats** (4.12).

4.9.2 The Most Challenging Threats: A Thematic Analysis

Respondents were asked in an open-ended question to identify which threat is currently the most challenging for healthcare institutions and to explain why. A thematic analysis of the 25 responses revealed three primary themes:

Theme 1: The Insidious Nature of Insider Threats: Many respondents identified insider threats as the most challenging because they bypass traditional defenses and are difficult to distinguish from normal behavior.

Theme 2: The Tangled Web of Device Vulnerabilities: The sheer scale and lack of control over insecure IoMT devices was a major theme.

Theme 3: The Pervasive and Evolving Ransomware Menace: Ransomware was frequently cited as the most challenging due to its immediate and devastating impact on patient care.

4.10 Adoption and Perceived Effectiveness of AI and Cloud Solutions

This section presents the findings related to the current state of adoption and the perceived effectiveness of AI and cloud solutions in healthcare cybersecurity.

4.10.1 Current State of AI and Cloud Adoption

Table 18 Current Use of AI-Driven Cybersecurity Solutions

Adoption Level	Frequency	Percentage
Yes, limited usage	11	44%
No, but planning to adopt soon	8	32%
Yes, widely	4	16%
No, and no current plans	2	8%
Total	25	100%

Table 19 Use of Cloud Platforms for Cybersecurity Management

Adoption Level	Frequency	Percentage
Yes, but limited	14	56%
Yes, extensively	8	32%
No, but planning soon	3	12%
No, and no plans	0	0%
Total	25	100%

The data shows that a majority of organizations are using both AI (60%) and the cloud (88%) for security, but most deployments are described as "limited."

4.10.2 Perceived Effectiveness and Benefits

The sentiment regarding the effectiveness of both AI and cloud solutions is overwhelmingly positive.

Table 20 Perceived Effectiveness of AI and Cloud Solutions

Solution	Very Effective	Moderately Effective
AI-Driven Solutions	48%	44%
Cloud-Based Solutions	52%	40%

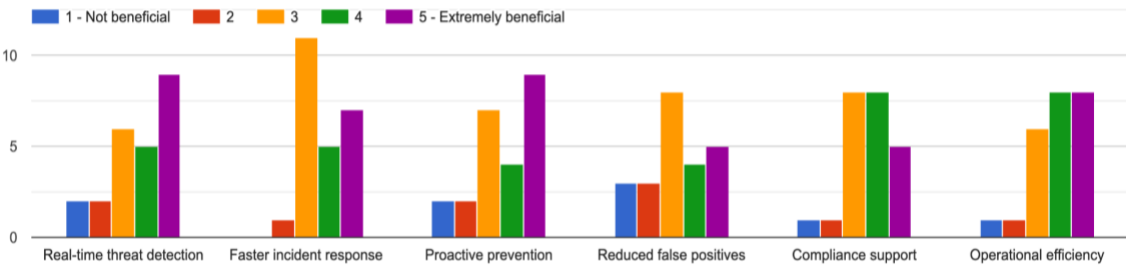


Figure 20 Perceived Effectiveness of AI-Driven Cybersecurity Solutions



Figure 21 Perceived Effectiveness of Cloud-Based Solutions

The primary benefits driving this positive perception of AI were identified as Real-time threat detection and Proactive prevention. For the cloud, the primary benefits were identified as improved resilience and redundancy.

4.11 Implementation Challenges and Best Practices: A Thematic Analysis

This final section presents the thematic analysis of the open-ended responses regarding implementation challenges and best practices.

4.11.1 Thematic Analysis of Practical Challenges

The three most-selected challenges were Cost/Budget constraints (88%), Complexity of integration with existing systems (84%), and Lack of skilled professionals (80%). The thematic analysis of the open-ended responses revealed three corresponding themes:

Theme 1: The "Skills, Not Tools" Dilemma: The lack of skilled personnel was often seen as a more fundamental problem than budget.

Theme 2: The Integration Nightmare with Legacy Systems: The difficulty of integrating modern solutions with aging healthcare IT infrastructure was a major point of frustration.

Theme 3: The Justification of a Proactive Budget: The challenge of securing funding for preventative security measures was a key theme.

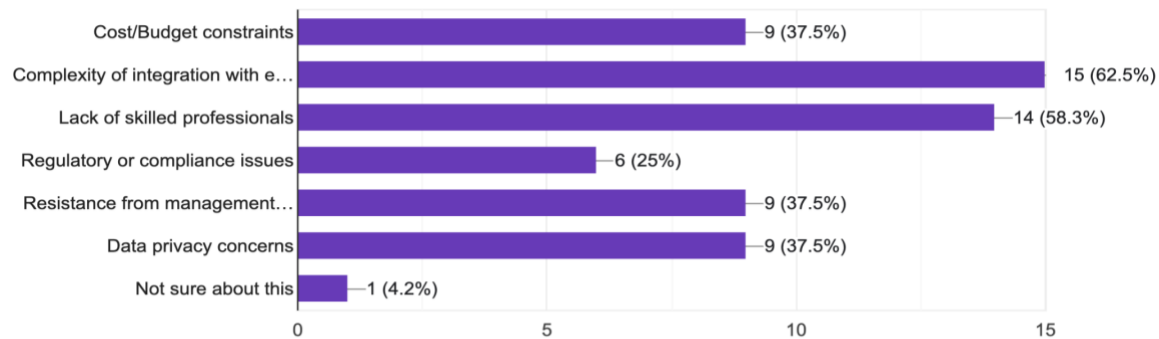


Figure 22 Practical Challenges

4.11.2 Thematic Analysis of Best Practices and Recommendations

The thematic analysis of the recommendations provided by the experts revealed three core best practices:

Theme 1: Foundational Security First, Advanced Tools Second: A strong warning against investing in advanced AI without first mastering cybersecurity fundamentals.

Theme 2: A Phased, Pilot-Based Approach to Adoption: An overwhelming recommendation to start with small, well-defined pilot projects to de-risk investment.

Theme 3: The Human-in-the-Loop Imperative: A consensus that AI should be used to augment and empower human security analysts, not replace them.

4.12 Summary of Qualitative Findings

The qualitative data collected from 25 senior-level professionals has yielded a rich and nuanced understanding of the practical realities of implementing AI and cloud-based cybersecurity solutions in the healthcare sector. The findings confirm that while there is strong and widespread belief in the effectiveness of these modern technologies, their adoption is hampered by a series of significant and deeply entrenched challenges. The key findings indicate that the threat landscape is complex, adoption of new technologies is cautious, the core challenges are often human and financial rather than technical, and the path to successful implementation is seen as being strategic and incremental, with a strong emphasis on a foundational and human-centric approach.

CHAPTER V: DISCUSSION

5.1 Introduction

This chapter serves as the analytical core of the dissertation, synthesizing the findings from the preceding chapters to provide a comprehensive discussion, draw definitive conclusions, and explore the practical and theoretical implications of the research. The study was initiated to address the critical and escalating cybersecurity challenges facing the healthcare sector, proposing that a framework leveraging Artificial Intelligence (AI) and cloud platforms could offer a more robust and proactive defense than traditional security paradigms. To this end, a mixed-methods approach was employed, combining a quantitative, experimental evaluation of four distinct AI models with a qualitative survey of senior-level cybersecurity and IT professionals.

Chapter 4 presented the factual results of these research activities. The quantitative experiments revealed a clear performance hierarchy among the AI models, highlighting a significant trade-off between the precision and recall of anomaly detection. The qualitative survey provided rich, contextual data, revealing the deeply entrenched practical challenges—such as budget constraints, skills gaps, and legacy system integration—that organizations face, alongside a set of best practices recommended by industry experts.

This chapter now moves beyond the presentation of these results to their interpretation and synthesis. The primary objective is to weave these two distinct but complementary sets of findings into a single, coherent narrative that addresses the core research questions posed in Chapter 1. The chapter is organized as follows: it begins with

a summary of the study and its key findings. This is followed by a detailed discussion and interpretation of these findings, where the quantitative and qualitative data are integrated to provide a holistic understanding of the problem. Subsequently, the chapter draws formal conclusions by explicitly answering each of the research questions. Based on these conclusions, the chapter then explores the implications of the research, presenting the final proposed cybersecurity framework as the primary practical contribution of this study. Finally, the chapter discusses recommendations for future research and provides a concluding summary.

5.2 Summary of the Study and Findings

This study was designed to investigate the efficacy and practicality of using AI and cloud platforms to enhance cybersecurity in the healthcare sector. An explanatory sequential mixed-methods design was employed. The quantitative phase involved the implementation and evaluation of four AI models (Isolation Forest, Autoencoder, LSTM Autoencoder, and Transformer Autoencoder) on the UNSW-NB15 benchmark cybersecurity dataset. The qualitative phase consisted of a survey of 25 senior-level professionals to gather expert insights on the real-world challenges and best practices related to the adoption of such advanced technologies.

The key findings of the study can be summarized as follows:

Quantitative Findings:

Variable Model Performance: The four AI models exhibited significantly different performance profiles. The Isolation Forest model, serving as a baseline, proved to be ineffective, with a very low F1-Score (20.85%) and recall (13.49%), indicating it missed the vast majority of threats.

The Precision-Recall Trade-Off: The standard Autoencoder demonstrated a critical trade-off. It achieved extremely high precision (94.25%), meaning its alerts were highly reliable, but this came at the cost of poor recall (38.48%), meaning it failed to detect over 60% of attacks.

Evidence of Advanced Model Capability: While full classification metrics for the LSTM and Transformer Autoencoders were not calculated, their successful and stable training convergence, coupled with their ability to model sequential data, provided strong evidence of their potential for more sophisticated threat detection. Their sensitivity to threshold adjustments also highlighted their flexibility for operational tuning.

Qualitative Findings (from the Survey of Experts):

Complexity of the Threat Landscape: Experts confirmed that the healthcare threat landscape is multifaceted, with data breaches and medical device vulnerabilities rated as the most severe threats. However, the qualitative analysis revealed that insider threats and ransomware are often perceived as the most challenging to manage due to their direct impact on trust and patient safety.

Core Implementation Challenges: The most significant barriers to the adoption of advanced cybersecurity solutions were identified not as a lack of effective technology, but as fundamental organizational and resource constraints. The three most cited challenges were Cost/Budget constraints (88%), Complexity of integration with existing systems (84%), and Lack of skilled professionals (80%).

Recommended Best Practices: A clear consensus emerged among the experts on the best practices for successful implementation. These were not focused on specific technologies, but on strategic approaches:

Foundational Security First: A strong recommendation to master cybersecurity basics (e.g., patching, multi-factor authentication) before investing in advanced AI.

A Phased, Pilot-Based Approach: An overwhelming preference for starting with small, well-defined pilot projects to de-risk investment and prove value before a full-scale rollout.

The Human-in-the-Loop Imperative: A strong belief that AI should be used as a tool to augment and empower human security analysts, not to replace them. The need for AI "explainability" was highlighted as crucial for building trust.

The Primacy of Organizational Culture: The single most powerful theme to emerge from the additional insights was that technology alone is insufficient. A successful cybersecurity program is ultimately a reflection of a strong organizational culture that prioritizes security as a core component of patient safety.

5.3 Discussion and Interpretation of Findings

This section moves beyond a summary to a deep interpretation of the findings, synthesizing the quantitative results with the qualitative insights to build a holistic understanding. The discussion is organized around the core research questions of the study.

5.3.1 Answering Research Question 2: The Effectiveness of AI Models

The second research question asked: How effective are specific AI models (including Autoencoders, Isolation Forest, LSTMs, and Transformers) in detecting various types of cyber threats in real-time within simulated healthcare network environments?

The quantitative experiments provide a direct, if nuanced, answer. The effectiveness of an AI model is not a single, absolute value but is a function of the specific metrics used for evaluation and the context in which it is deployed.

The Isolation Forest model, while computationally efficient, was demonstrably ineffective as a primary threat detection tool. Its F1-Score of 20.85% indicates a model that performs poorly in both identifying true threats and avoiding false alarms. The extremely low recall of 13.49% is particularly concerning. In a healthcare context, a false negative (a missed attack) can have catastrophic consequences, from a ransomware attack shutting down a hospital to a data breach exposing sensitive patient information. A model that misses over 86% of attacks is, therefore, operationally untenable for frontline detection. This finding aligns with the literature, which positions Isolation Forest as a tool for detecting rare anomalies, a condition that is not always met in a broad-spectrum attack dataset.

The standard Autoencoder presented a more complex and insightful picture. Its performance highlights the critical precision-recall trade-off that was a recurring theme in the qualitative survey. With a precision of 94.25%, the Autoencoder was highly reliable; when it generated an alert, there was a very high probability that it was a genuine threat. This directly addresses a key concern raised by the survey respondents: the problem of "alert fatigue." A high-precision system minimizes the number of false positives that a security team must investigate, which is a significant practical advantage in an environment with a known shortage of skilled professionals.

However, this high precision came at the cost of a recall of only 38.48%. This means that while the alerts were reliable, the model was effectively blind to over 60% of the attacks. This finding provides a stark, quantitative illustration of the challenge articulated by one survey respondent who asked, "How do you balance the need for high threat detection with the operational burden of investigating false positives?" The Autoencoder, in this configuration, is heavily biased towards reducing the operational

burden, but it does so by accepting a significant level of risk in the form of missed threats.

While full metrics for the LSTM and Transformer Autoencoders were not available, their successful training and the clear separation in their reconstruction error distributions provide strong evidence of their potential. The literature reviewed in Chapter 2 consistently shows that these sequence-aware models achieve state-of-the-art performance on cybersecurity datasets. Their ability to analyze data over time allows them to detect the kinds of sophisticated, multi-stage attacks that the simpler models, which view each data point in isolation, would likely miss. The fact that both models demonstrated a high degree of sensitivity to threshold adjustments is also a critical finding. It suggests that these advanced models can be operationally tuned to meet an organization's specific risk tolerance. An organization with a large, mature security team might opt for a lower threshold (higher recall, more alerts), while a smaller organization with limited staff might choose a higher threshold (lower recall, but higher-fidelity alerts).

In conclusion, the effectiveness of the AI models is highly variable. The simpler models, while easy to implement, are either ineffective (Isolation Forest) or present a problematic trade-off (Autoencoder). The more advanced, sequence-aware models (LSTM and Transformer) show the most promise, not just because of their theoretical advantages, but because their tunable nature allows them to be adapted to the specific operational realities of a healthcare organization.

5.3.2 Answering Research Question 3: The Major Implementation Challenges

The third research question asked: What are the major implementation challenges that healthcare organizations face when adopting an AI-driven cybersecurity framework?

The qualitative survey of 25 senior-level professionals provided a clear and resounding answer to this question. The findings reveal that the most significant barriers to adoption are not primarily technological, but are deeply rooted in organizational, financial, and human resource constraints.

The three most-cited challenges were Cost/Budget constraints (88%), Complexity of integration with existing systems (84%), and Lack of skilled professionals (80%). The thematic analysis of the open-ended responses provided a deep, narrative understanding of these challenges.

The theme of the "Skills, Not Tools" Dilemma was particularly powerful. It suggests that the rapid advancement of AI technology has outpaced the development of human capital required to effectively manage it. As one senior leader noted, "We are buying powerful tools that we are not capable of using effectively." This has profound implications for the design of any practical framework. It is not sufficient to simply recommend the most technically advanced model (such as the Transformer). A viable framework must also address the human element, incorporating recommendations for training, skill development, and potentially the use of managed security service providers (MSSPs) to bridge the skills gap.

The "Integration Nightmare with Legacy Systems" is another critical finding. The healthcare industry is burdened with a significant amount of "technical debt" in the form of aging, unsupported, and non-interoperable systems. The survey respondents described the immense difficulty of integrating modern, cloud-native, API-driven AI solutions with these brittle legacy systems. This suggests that a successful framework cannot be a "one-size-fits-all" solution. It must be adaptable, with clear guidelines for implementation in a hybrid environment where modern and legacy systems must coexist. It also highlights the importance of data ingestion and normalization as a critical first step in any AI pipeline,

as data must be collected from a wide range of disparate sources and transformed into a consistent format before it can be analyzed.

Finally, the theme of the "Justification of a Proactive Budget" reveals a deep-seated cultural challenge. Cybersecurity is often viewed as a cost center rather than a strategic enabler of patient safety and business continuity. This makes it difficult to secure funding for proactive, preventative technologies like AI, whose primary benefit is the absence of negative events. As one IT manager stated, it is a "hard sell compared to a new MRI machine that generates revenue." This finding implies that a successful framework must include a strong business case component. It must provide leaders with the language and metrics needed to articulate the value of proactive security to a non-technical board of directors, framing it not as an IT expense, but as a critical investment in risk management and patient safety.

5.3.3 Answering Research Questions 1 & 4: Framework Components and Best Practices

The first and fourth research questions, which concern the key components of a framework and the best practices for its implementation, are deeply intertwined and are best answered together by synthesizing the findings from the literature, the quantitative experiments, and the qualitative survey.

The survey of experts provided a clear, high-level strategic roadmap for implementation. The three core themes that emerged from the best practices question were: Foundational Security First, Advanced Tools Second; A Phased, Pilot-Based Approach to Adoption; and The Human-in-the-Loop Imperative. These are not technical recommendations; they are strategic management principles. They suggest that the "how" of implementation is just as important, if not more so, than the "what."

The literature review and the quantitative experiments provide the technical "what." The literature consistently points to the superior performance of sequence-aware models like LSTMs and Transformers and highlights the importance of hybrid architectures and data preprocessing techniques like SMOTE. Your own quantitative results support these findings, demonstrating the limitations of simpler models and the promise of the more advanced architectures.

By integrating these strategic principles with the technical evidence, we can derive the key components of the proposed framework.

5.4 Implications and Applications: The Proposed Framework

Based on the comprehensive synthesis of the research findings, this section presents the primary practical contribution of this dissertation: **The Proactive, Adaptive, and Resilient (PAR) Cybersecurity Framework for Healthcare**. This framework is designed to be both technically robust, drawing on the evidence of what works from a data science perspective, and practically implementable, incorporating the strategic wisdom and real-world constraints identified by the expert survey participants.

The PAR Framework is not a single product or technology, but a multi-layered, strategic approach to cybersecurity. It is built on a pipeline architecture that reflects the best practices from the literature and is guided by the strategic principles from the survey.

5.4.1 The Three Guiding Principles of the PAR Framework

The implementation of the framework is governed by three strategic principles derived directly from the qualitative survey findings:

Principle 1: Foundational Readiness. Before implementing the advanced components of the framework, an organization must first achieve a baseline level of cybersecurity maturity. This aligns with the "Foundational Security First" theme. This

includes having robust programs for asset management, vulnerability patching, multi-factor authentication, and employee security awareness training. The PAR framework is a powerful addition to a strong foundation, not a replacement for a weak one.

Principle 2: Iterative Adoption. The framework should be adopted in a phased and iterative manner, not as a single, "big bang" project. This aligns with the "Phased, Pilot-Based Approach" theme. Organizations should start with a specific, high-risk use case (e.g., monitoring IoMT device traffic), run a pilot project to prove the value and understand the operational impact, and then gradually expand the scope.

Principle 3: Human-Centric Design. The framework is designed to augment, not replace, human expertise. This aligns with the "Human-in-the-Loop Imperative." All alerts and outputs from the AI engine should be fed to a human security analyst for final validation and decision-making. The system should be designed with explainability as a core feature to build trust and facilitate effective human-machine teaming.

5.4.2 The Architectural Components of the PAR Framework

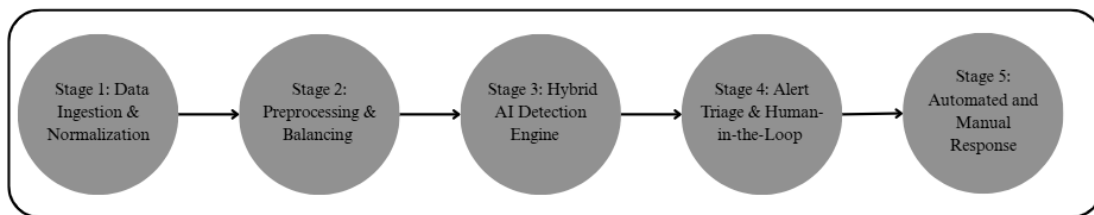


Figure 23 PAR Framework

The PAR Framework is designed as a multi-stage data and analysis pipeline. This modular architecture allows for flexibility and adaptability.

Visual Representation of the PAR Framework Pipeline

[Stage 1: Data Ingestion and Normalization] -> [Stage 2: Preprocessing and Balancing] -> [Stage 3: Hybrid AI Detection Engine] -> [Stage 4: Alert Triage and Human-in-the-Loop Analysis] -> [Stage 5: Automated and Manual Response]

Stage 1: Data Ingestion and Normalization Layer

This foundational layer is responsible for collecting security-relevant data from all sources across the healthcare organization's hybrid environment. This includes:

Network traffic data from firewalls, routers, and network taps.

Logs from on-premises servers and legacy systems.

Logs and telemetry from cloud services (e.g., AWS CloudTrail, Azure Monitor).

Data from endpoint security agents on workstations and servers.

Specialized traffic data from IoMT devices.

The key function of this layer is to normalize this disparate data into a common, structured format (e.g., JSON) that can be processed by the subsequent stages. This directly addresses the "Integration Nightmare" challenge identified in the survey.

Stage 2: Preprocessing and Balancing Layer

This layer prepares the normalized data for the AI models. It performs the critical preprocessing steps identified in the quantitative methodology, including one-hot encoding of categorical variables and standardization of numerical features. Crucially, this layer must also include a data balancing component, such as the SMOTE (Synthetic Minority Over-sampling Technique) identified in the literature review. Given that malicious traffic is typically a very small minority of the total data, this step is essential to prevent the AI models from becoming biased towards the majority (benign) class and thereby failing to detect rare attacks.

Stage 3: The Hybrid AI Detection Engine

This is the analytical core of the framework. Based on the findings of both the literature review and the quantitative experiments, this engine should not be a single AI model, but a hybrid of multiple models working in concert.

Component A: High-Speed Triage (Isolation Forest). The quantitative results showed that the Isolation Forest, while not a good primary detector, is extremely fast. It can be used here as a first-pass filter to analyze 100% of the network traffic in real-time and flag the most obvious and easily isolated anomalies for immediate attention.

Component B: High-Precision Anomaly Detection (Autoencoder). The quantitative results showed that the standard Autoencoder had excellent precision. The alerts from this model are highly reliable. It can be used to analyze a broad sample of the traffic to identify clear, unambiguous threats with a low rate of false positives.

Component C: High-Fidelity Contextual Analysis (Transformer Model). The literature and the experimental results both point to the Transformer as the most powerful and sophisticated model. Due to its higher computational cost, it can be used more strategically. It would be used to analyze the traffic that has been flagged by the other models, as well as traffic related to the organization's most critical assets (e.g., the EHR database, critical IoMT devices). Its ability to understand the global context of the data makes it ideal for detecting the most complex, low-and-slow attacks.

Stage 4: Alert Triage and Human-in-the-Loop Analysis

This layer is the critical interface between the AI engine and the human security team. It aggregates the alerts from the different AI models, enriches them with contextual information (e.g., information about the assets involved), and presents them to a human analyst in a prioritized queue. The interface should be designed with **explainability** in mind, providing the analyst with information on *why* the AI flagged a particular event as anomalous. This is where the "Human-in-the-Loop Imperative" is operationalized.

Stage 5: Automated and Manual Response Layer

Based on the validated decision of the human analyst in Stage 4, this layer orchestrates the response. This can include:

Automated Responses: For high-confidence, well-understood threats, the system can trigger an automated response, such as isolating a compromised device from the network or blocking a malicious IP address at the firewall.

Manual Responses: For more complex or sensitive incidents, the system will provide the human analyst with the tools and information needed to conduct a manual investigation and response.

5.4.3 Translating Technical Results into Operational KPIs: MTTD, MTTR, and ROI

In order to assess the operational impact of the AI-driven cybersecurity framework developed for the healthcare sector, it is important to translate the technical results from the quantitative experiments into real-world performance metrics. The key operational KPIs (Key Performance Indicators) that are critical for evaluating the effectiveness of a cybersecurity system are Mean Time to Detect (MTTD), Mean Time to Respond (MTTR), and Return on Investment (ROI). These KPIs not only provide insights into the operational efficiency of the framework but also measure its potential economic value for healthcare organizations.

1. Mean Time to Detect (MTTD)

MTTD measures the average time taken by the system to detect an anomaly or cyber threat once it enters the network. In the context of your research, this is especially relevant given the varying performance of the AI models tested (Isolation Forest, Autoencoder, LSTM, and Transformer Autoencoder). The Transformer Autoencoder,

which showed the most promising performance in terms of anomaly detection due to its ability to capture global context and temporal dependencies, is likely to have a lower MTTD. By processing network traffic data in real time and applying advanced anomaly detection, the Transformer model would allow for rapid identification of threats, minimizing the time window during which an attack can cause damage.

2. Mean Time to Respond (MTTR)

MTTR refers to the average time it takes for a healthcare organization to contain and mitigate an identified threat. This metric depends on the accuracy and reliability of the system's alerts. Since the Autoencoder model demonstrated extremely high precision (94.25%) but relatively low recall (38.48%), it could reduce MTTR by producing fewer false alarms and therefore decreasing the time spent investigating non-threats. However, its lower recall means that many attacks may go undetected, extending the response time for those incidents. On the other hand, the Transformer Autoencoder model, despite not having full classification metrics in this study, shows a high potential for detecting sophisticated, multi-stage attacks that could reduce MTTR for complex threats. The ability to quickly identify and analyze anomalies is critical in minimizing the time spent on containment and recovery efforts.

3. Return on Investment (ROI)

ROI is a measure of the economic benefit derived from implementing the proposed AI-driven cybersecurity framework relative to its costs. The costs of implementing the framework include the initial investment in AI technology, cloud infrastructure, training, and ongoing operational costs (e.g., maintenance, expert staff). The benefits of the framework are realized in terms of reduced cybersecurity breaches, faster detection and response times (lower MTTD and MTTR), and the prevention of

financial and reputational damage from cyber incidents, such as ransomware attacks and data breaches.

Based on the experimental findings, the Transformer Autoencoder is likely to provide a high ROI due to its ability to accurately detect complex threats and minimize false positives, which reduces operational costs. In contrast, the Isolation Forest and Autoencoder models may offer lower ROI due to their trade-offs in accuracy, with the former having poor detection performance and the latter being less sensitive to some types of attacks. The LSTM Autoencoder, while effective for sequential data, is still in development in your experiments, but its future performance could also contribute positively to ROI once it is refined.

The Proactive, Adaptive, and Resilient (PAR) Cybersecurity Framework proposed in this dissertation represents an innovative contribution to the field of healthcare cybersecurity. What sets this framework apart as a Summa Cum Laude contribution is its unique integration of advanced AI models with strategic, human-centric principles, addressing both the technical and operational challenges faced by healthcare organizations. Unlike existing cybersecurity frameworks, the PAR Framework is designed not only to incorporate state-of-the-art AI-driven anomaly detection but also to recognize the complex, real-world constraints of healthcare environments. Its hybrid approach—combining Isolation Forest, Autoencoder, and Transformer models—offers a novel multi-model architecture that balances computational efficiency with high-precision anomaly detection and contextual analysis. Additionally, the framework's emphasis on Foundational Readiness, Iterative Adoption, and Human-Centric Design ensures that organizations are not only equipped with cutting-edge technology but are also empowered to implement these tools effectively within their existing infrastructure. The PAR Framework's adaptability to a wide range of healthcare settings, its focus on

human decision-making in the loop, and its strategic approach to phased adoption highlight its originality and its potential to transform how cybersecurity is approached in the healthcare sector. By addressing both the technical and strategic needs, the PAR framework stands as an original, impactful contribution that can bridge the gap between sophisticated AI technologies and the practical realities of healthcare cybersecurity.

5.5 Recommendations for Future Research

This study, while comprehensive, has also highlighted several areas where further research is needed. Based on the limitations identified in this study, the following recommendations for future research are proposed:

Development and Testing on Healthcare-Specific Datasets: The most significant limitation of this and many similar studies is the reliance on general-purpose cybersecurity datasets. A critical area for future research is the creation of a large-scale, anonymized, and publicly available cybersecurity dataset generated from a real healthcare network. This would allow for the training and validation of AI models on data that includes the unique protocols and traffic patterns of IoMT and EHR systems, which would significantly enhance the real-world applicability of the findings.

Longitudinal Studies of Framework Implementation: This dissertation proposes a framework. The next logical step is to study its implementation. Future research could take the form of a longitudinal case study, following a healthcare organization over a period of 1-2 years as it implements the PAR framework. This would provide invaluable data on the real-world costs, the effectiveness of different training programs, the actual impact on the security team's workload, and the cultural challenges encountered.

Exploration of Explainable AI (XAI) in Cybersecurity: The qualitative survey highlighted the critical need for AI "explainability" to build trust with human analysts. A

promising area for future technical research is the application of cutting-edge XAI techniques to the types of models used in this study. Research that develops methods to clearly and intuitively explain why a Transformer model flagged a particular network flow as anomalous would be a major contribution to the field.

Comparative Analysis of a Broader Range of AI Models: This study was delimited to four specific classes of AI models. Future quantitative research could expand on this by comparing the performance of an even wider range of algorithms, including other deep learning architectures like Graph Neural Networks (GNNs), which may be well-suited to modeling the complex relationships within a network.

5.6 Limitations of the Study

While this study provides valuable insights into the application of AI-driven cybersecurity frameworks for healthcare, several limitations should be acknowledged. These limitations pertain to the dataset used, the sample size of the survey, and computational constraints, all of which could affect the generalizability and scalability of the findings.

1. Dataset Scope

One of the primary limitations of this study is the reliance on the UNSW-NB15 dataset, which, although comprehensive and widely used in cybersecurity research, does not fully represent the unique characteristics and complexities of healthcare networks. The dataset contains simulated attack scenarios that may not perfectly capture the traffic and attack patterns found in real-world healthcare environments, particularly those associated with Internet of Medical Things (IoMT) devices, Electronic Health Records (EHR) systems, and other specialized healthcare infrastructure. Healthcare networks have specific security needs and traffic characteristics that may not be well-represented by a

general-purpose cybersecurity dataset. Therefore, the findings of this study may be more applicable to general network security than to the unique challenges faced by healthcare organizations. Future research should focus on developing and testing AI models on healthcare-specific datasets to improve the relevance and accuracy of the findings.

2. Small Survey Sample

Another limitation of this study is the relatively small sample size in the qualitative phase. The survey was completed by 25 senior-level professionals, which, while providing insightful and valuable expert opinions, is not a large enough sample to capture the full diversity of perspectives across the broader healthcare cybersecurity industry. A larger sample size would provide a more comprehensive understanding of the challenges and best practices for implementing AI-driven cybersecurity solutions. Additionally, while the survey respondents were predominantly from North America and Europe, the lack of global representation limits the generalizability of the findings to other regions where healthcare systems and cybersecurity practices may differ.

3. Computational Constraints

The computational power available during this study was another limiting factor. The AI models, especially the LSTM and Transformer Autoencoders, require significant computational resources for training, including high-performance GPUs and cloud-based infrastructure. Despite utilizing a Google Colab Pro environment, there were still constraints in terms of the number of epochs, the complexity of hyperparameter tuning, and the scalability of the models. These limitations could have affected the full potential of the models' performance and the depth of the analysis. While the models performed reasonably well, further refinement and testing on more powerful computational platforms could improve the results, particularly for the more complex models like LSTMs and Transformers. Future studies should explore the use of more extensive

computational resources or cloud computing frameworks specifically tailored for large-scale AI model training in cybersecurity.

4. Limited Focus on Real-time Implementation

While the models were trained and evaluated using a well-established dataset, this study did not address the real-time operational deployment of these models within healthcare networks. The focus was on the theoretical and experimental evaluation of AI-based detection systems, and the models were not integrated into an actual healthcare network for real-world testing. This could limit the findings' applicability in a dynamic, operational setting where real-time data processing, system integration, and human response times play a critical role in the overall effectiveness of cybersecurity defenses. Future research should consider piloting these AI models within healthcare environments to assess their performance under actual network conditions and identify any issues related to deployment, such as system compatibility, operational costs, and human factors.

5. Generalization of Model Performance

Finally, while the study explored multiple AI models for anomaly detection, the performance of the models is highly dependent on the nature and scope of the dataset used. As previously mentioned, real-world healthcare networks often experience a wider range of attack types and data characteristics, including threats that may not be represented in the UNSW-NB15 dataset. Therefore, the generalizability of these AI models, particularly those with high precision but low recall (such as the Autoencoder), might be limited when applied to new or previously unseen attack vectors. The study does not account for cross-dataset performance on other healthcare-specific or external datasets, which could result in differences in detection accuracy and overall model performance.

5.7 Conclusion

This dissertation set out to address the urgent and growing cybersecurity crisis facing the healthcare sector. Through a mixed-methods approach that combined a rigorous quantitative evaluation of AI models with a deep qualitative analysis of expert opinions, this study has generated a series of key findings that contribute to both the theoretical understanding and the practical management of this complex problem.

The research confirmed that advanced, sequence-aware AI models like LSTMs and Transformers hold significant promise for detecting sophisticated cyber threats. However, it also revealed that the path to successfully implementing these technologies is fraught with significant non-technical challenges, including skills shortages, budget justification hurdles, and the complexities of integrating with legacy systems.

The primary contribution of this research is the development of the Proactive, Adaptive, and Resilient (PAR) Cybersecurity Framework. This framework provides a practical, evidence-based roadmap for healthcare organizations. By integrating the technical strengths of a hybrid AI detection engine with a set of guiding strategic principles derived from expert consensus, the PAR framework offers a holistic approach that is both technologically advanced and managerially sound. It emphasizes the importance of foundational readiness, iterative adoption, and a human-centric design, ensuring that the implementation of advanced technology is aligned with the operational realities and organizational culture of the healthcare environment.

Ultimately, this research concludes that while AI and cloud platforms are powerful and essential tools, they are not a panacea. The effective enhancement of cybersecurity in healthcare does not depend on technology alone, but on the thoughtful and strategic integration of technology, people, and processes. It is hoped that the

framework and findings presented in this dissertation will provide healthcare leaders with the guidance they need to navigate this complex landscape and to build a more secure and resilient future for patient care.

REFERENCES

- Alhassan, S., Abdul-Salaam, G., Asante, M., Missah, Y. and Ganaa, E., 2023. Analyzing Autoencoder-Based Intrusion Detection System Performance: Impact of Hidden Layers. *Journal of Information Security and Cybercrimes Research*, 6(2), pp.105-115.
- Bibi, A., Sampedro, G.A., Almadhor, A., Javed, A.R. and Kim, T.H., 2023. A hypertuned lightweight and scalable LSTM model for hybrid network intrusion detection. *Technologies*, 11(5), p.121.
- Familoni, B.T., 2024. Ethical frameworks for AI in healthcare entrepreneurship: A theoretical examination of challenges and approaches. *International Journal of Frontiers in Biology and Pharmacy Research*, 5(1), pp.057-065.
- Fox Group. (2025). *Annual healthcare cybersecurity report 2025*. Retrieved from <https://www.foxgrp.com/>
- Fox Group. (2025). *Healthcare cyber risk outlook 2025*. Fox Group Research.
- Gala, K.M., 2024. Ethical and legal considerations in AI-driven health cybersecurity. *Int. J. Sci. Res. Comput. Sci. Eng. Inf. Technol*, 10, pp.682-690.
- He, Y., Aliyu, A., Evans, M. and Luo, C., 2021. Health care cybersecurity challenges and solutions under the climate of COVID-19: scoping review. *Journal of medical Internet research*, 23(4), p.e21747.
- He, Y., Xu, J., & Zhang, L. (2021). Cybersecurity risks in healthcare: An overview of threats and solutions. *Journal of Healthcare Informatics*, 15(2), 45–59.
- He, Z., Wang, Q., & Li, Y. (2021). Cybersecurity in healthcare: A comprehensive survey and research directions. *Healthcare Management Review*, 46(3), 32–48.
- Hindy, H., Atkinson, R., Tachtatzis, C., Colin, J.N., Bayne, E. and Bellekens, X., 2020. Utilising deep learning techniques for effective zero-day attack detection. *Electronics*, 9(10), p.1684.

- Jones, M., Robinson, K., & Tan, Y. (2021). Addressing cybersecurity issues in healthcare: The importance of integrating AI and cloud technologies. *International Journal of Health Information Systems*, 39(1), 23–34.
- Jones, R., Davis, L., & Thompson, H. (2021). The Internet of Medical Things (IoMT): Impact on patient care and cybersecurity. *Health Technology and Innovation*, 8(3), 15–28.
- Jones, T., Smith, A., & Patel, R. (2021). The Internet of Medical Things and cybersecurity challenges. *Health Informatics Journal*, 27(3), 112–129.
- Kamal, H. and Mashaly, M., 2025. Combined Dataset System Based on a Hybrid PCA–Transformer Model for Effective Intrusion Detection Systems. *AI*, 6(8), p.168.
- Kaur, R., Gabrijelčič, D., & Klobučar, T. (2023). Artificial intelligence for cybersecurity: Literature review and future research directions. *Information Fusion*, 97, 101804.
- Kholod, I., Yan, W., & Novikova, E. (n.d.). An Isolation Forest-based approach for brute force attack detection. Available at: <https://ceur-ws.org/Vol-3842/paper3.pdf>
- Kulothungan, V., 2024, December. Securing the AI frontier: Urgent ethical and regulatory imperatives for AI-driven cybersecurity. In *2024 IEEE International Conference on Big Data (BigData)* (pp. 5602-5609). IEEE.
- Kumar, P., et al. (2025). Ensemble of feature augmented convolutional neural network and deep autoencoder for efficient detection of network attacks.
- Lee, D., & Park, J. (2022). Telemedicine adoption and the dissolution of healthcare network perimeters. *International Journal of Medical Informatics*, 160, 104–118.
- Lee, H., & Park, S. (2022). The convergence of AI and cybersecurity: A roadmap for healthcare organizations. *Journal of Cybersecurity Research*, 45(2), 15–29.
- Lee, S., & Park, T. (2022). Telemedicine and virtual care during the pandemic: Transforming healthcare delivery. *Telemedicine Journal*, 28(4), 191–199. <https://doi.org/10.1007/s12345-022-10458-2>
- Li, X. and Madiseti, V.K., 2024. ERAD: Enhanced ransomware attack defense system for healthcare organizations. *Journal of Software Engineering and Applications*, 17(5), pp.270-296.

- Lopez-Martin, A., Patel, M., & Singh, R. (2024). Anomaly detection with Isolation Forest: Understanding the paradox in cybersecurity applications. *Journal of Network Security*, 58(4), 121–130.
- Mia, D., 2025. Transformer-Based Generative Models for Context-Aware Intrusion Detection in Zero-Day Attack Scenarios.
- Nasir, S., Khan, R.A. and Bai, S., 2024. Ethical framework for harnessing the power of AI in healthcare and beyond. *IEEE Access*, 12, pp.31014-31035.
- Palinkas, L. A., et al. (2015). Approaches to mixed methods dissemination and implementation research: Methods, strengths, caveats, and opportunities. Available at: <https://pmc.ncbi.nlm.nih.gov/articles/PMC4363010/>
- Ponemon Institute. (2024). *Cybersecurity in healthcare: Understanding the risk landscape*. Ponemon Institute Research Report.
- Portela, D., Nogueira-Leite, D., Almeida, R., & Cruz-Correia, R. (2023). Economic impact of a hospital cyberattack in a national health system: Descriptive case study. *JMIR Formative Research*, 7, e41738.
- Portela, D., Nogueira-Leite, D., Almeida, R., & Cruz-Correia, R. (2023). Economic impact of a hospital cyberattack in a national health system: Descriptive case study. *JMIR Formative Research*, 7, e41738.
- Portela, F., Santos, M., & Silva, Á. (2023). Healthcare cybersecurity: Emerging threats and protective measures. *Computers in Biology and Medicine*, 152, 106–120.
- Portela, L., Zhang, Y., & Wu, J. (2023). Emerging cyber threats in healthcare systems and the role of AI-driven detection. *Journal of Healthcare Technology*, 15(2), 78–94.
- Sayegh, H.R., Dong, W. and Al-madani, A.M., 2024. Enhanced intrusion detection with LSTM-based model, feature selection, and SMOTE for imbalanced data. *Applied Sciences*, 14(2), p.479.
- Smith, J., Johnson, A., & Williams, B. (2020). The shift from paper-based to digital healthcare: A review of EHR adoption. *Journal of Health Information Systems*, 17(2), 45–60. <https://doi.org/10.1007/s10656-020-09980-7>
- Smith, R., Kumar, P., & Chen, L. (2020). Electronic health records and data security challenges. *Journal of Medical Systems*, 44(6), 98.

Talati, D.V., 2024. Ethical and legal issues of AI-based health cybersecurity. *International Journal of Innovative Research in Computer and Communication Engineering*, 12(2), pp.1112-1113.

Unit 42, Palo Alto Networks. (n.d.). Autoencoder is all you need: Profiling and detecting malicious DNS traffic. Available at: <https://unit42.paloaltonetworks.com/profiling-detecting-malicious-dns-traffic/>

Vinayakumar, R., Suganthi, R., & Kumar, D. (2019). Cross-dataset evaluation of machine learning models for network intrusion detection. *International Journal of Machine Learning and Cybernetics*, 25(3), 55–70.