



AI-Driven Decision-Support Framework for Cybersecurity Consulting Under Zero-Trust

Tahir Iqbal^{1, *},  

¹Department of Business Administration, Imam Abdulrahman Bin Faisal University, Dammam, Kingdom of Saudi Arabia

Article History

Received: 13 March, 2026

Revised: 07 July, 2026

Accepted: 08 July, 2026

Abstract:

Introduction: Decision-support mechanisms are becoming a vital component of cybersecurity operations in the face of adaptive threats, especially when they do not rely on static rules or perimeter-based controls. This study introduces an AI-driven Zero-Trust decision-support architecture that integrates outputs from multiple deep-learning algorithms and maps them to a unified risk score for access control enforcement.

Methodology: The aim is to explore the ability of Artificial Neural Network (ANN), Convolutional Neural Network (CNN) and Long Short-Term Memory (LSTM) models to assist with graduated Zero-Trust decisions under simulated cyberattack conditions. The evaluation is based on a simulated dataset of identity-access logs, system events, network traffic logs, multi-factor authentication cues, and user metadata representing normal and malicious behaviour.

Results: The Zero-Trust layer overcame the binary nature of access decisions by moving uncertain cases into the MONITOR and MFA action layers. The confusion matrices, however, also highlight a key false-negative scenario, especially for ANN and CNN, suggesting that some attack events may go undiscovered without the inclusion of enforcement safeguards.

Conclusion: It is therefore best understood in the context of a decision-support system, rather than as a stand-alone detection system. Given the nature of the evaluation, future research can consider using live enterprise traffic, baseline comparisons, and end-to-end deployment testing to validate the framework.

Keywords: AI-driven framework, cybersecurity consulting, predictive analytics, zero-trust, artificial neural networks (ANN), convolutional neural networks (CNN), long short-term memory (LSTM), decision-making, cyber-attack detection.

1. INTRODUCTION

The increasing complexity and pace of cyber-attacks have made cybersecurity a serious concern for organisations worldwide (Li & Liu, 2021). The rapid spread of cyber-attacks, identity breaches, insider threats, and novel vulnerabilities is compelling organisations to review their security policies. Perimeter security and other traditional security approaches have failed to respond to the changing characteristics of the threats (Kanagamalliga *et al.*, 2024). They tend to set a specific scope for an organisation's IT infrastructure, but they are not mindful of the dynamics and complexity of modern cyber threats. The continually increasing complexity of cyber-

attacks, particularly distributed denial-of-service (DDoS), advanced persistent threats (APTs), and insider threats, requires the creation of more responsive and adaptive security measures (Sharma *et al.*, 2023).

In this paper, strategic classes of cyber risk are explained using high-level threat categories: distributed denial-of-service (DDoS), advanced persistent threats (APTs), and insider threats. These types include various operational attack methods witnessed at various attack lifecycle levels. In this manner, the follow-up experimental assessment addresses representative attack behaviours, such as credential stuffing, phishing, malware execution, reconnaissance, backdoor activity, and

*Address correspondence to this author at Department of Business Administration Imam Abdulrahman Bin Faisal University, Dammam, Kingdom of Saudi Arabia; E-mail: timuniruddin@iau.edu.sa



© 2026 Copyright by the Author.

Licensed as an open access article using a [CC BY 4.0 license](https://creativecommons.org/licenses/by/4.0/).

shellcode injection, which, individually, represent these more general threat categories in controlled simulation environments.

The major drawback of conventional security models is that they cannot keep up with the speed and volume of cyber-attacks (Li & Liu, 2021). These models are typically hard-and-fast, reactive, and not based on known rules and signatures, but on them and thus do not respond to new threats. The constantly evolving cybersecurity landscape demands that security systems not only detect known threats but also preempt and prevent the emergence of new ones. Predictive analytics is an artificial intelligence (AI) model that can provide a solution (Sarker, 2023). Organisations can be more proactive in their cybersecurity strategy and move away from a reactive stance by leveraging machine learning techniques such as Artificial Neural Networks (ANNs), Convolutional Neural Networks (CNNs), and Recurrent Neural Networks (RNNs) (Sarker, 2021).

Moreover, implementing Zero-Trust principles in cybersecurity systems can address the insufficiency of perimeter protection. Zero-Trust is founded on the principle that a given organisation cannot trust any insider or outsider by default (Tyler & Viana, 2021). Every access request is verified, and users and devices on the network are continuously verified. The model above would greatly improve the system's security, as it would minimise the attack surface and enable threats to be detected beforehand, even in trusted systems. A Zero-Trust design, along with foreseeable AI models, can give an organisation greater ability to make on-the-fly, data-driven choices, making it more responsive to cyber threats overall (Joshi, 2024).

Despite the clear advantages of AI-based solutions, consultancies are likely to face several challenges when implementing these technologies in their cybersecurity operations. Several businesses cannot determine the maturity of their security systems, the magnitude of their cybersecurity spending, and the risk (Marican *et al.*, 2023). This is among the field's critical weaknesses, as AI models have yet to be integrated into business analysis systems to address cybersecurity. Traditional security consultancy firms rely on manual operations and an antiquated model, and are generally ineffective at responding to the dynamism of cyber threats. Moreover, the inability to validate strategies through simulation complicates the effectiveness of cybersecurity strategies that cannot be tested before implementation (Armenia *et al.*, 2021).

Existing business analysis models in cybersecurity are not designed to account for the complex nature of contemporary threats and fail to provide the dynamism needed for swift decision-making (Naseer *et al.*, 2021). This is the vulnerability where a more robust, fact-based, data-driven method of cybersecurity decision-making is needed that can leverage the power of predictive analytics and AI to assess risks in real-time. The purpose of this study is to design an artificial intelligence (AI)-based business analytics model for cybersecurity consulting that combines predictive risk analytics with the Zero-Trust concept to support more efficient decision-making.

This framework can help cybersecurity consultants make evidence-based decisions more quickly, thereby enhancing the responsiveness of an organisation's security posture.

In this study, the term business analysis refers to decision support for cybersecurity consulting rather than direct financial optimisation. The suggested framework aims to convert technical risk indicators into practical consulting outcomes, including the access control posture, Zero-Trust maturity progression, and the prioritisation of security interventions. Although business constructs such as return on investment (ROI) and cost-risk forecasting are addressed to demonstrate their applicability to consulting engagements, the quantitative analysis of these constructs is beyond the scope of the current work.

Although the use of artificial intelligence methods combined with Zero-Trust security approaches has already been studied in previous literature, the originality of the current work lies not in proposing the integration concept as an idea. Rather, this work contributes to the idea that the results of heterogeneous AI models can be operationalised into harmonised risk scores, which directly lead to graduated Zero-Trust policy implementation and maturity testing, and allow decision-level security controls rather than detection-only analytics.

The most important aims of this study are the following:

- Predict risk levels using AI models, such as ANNs, CNNs, and LSTM (a particular RNN), tailored to cybersecurity applications.
- Integrate AI model outputs into a Zero-Trust Architecture (Zero-Trust) decision-making framework.
- Validate the framework using the Security Maturity Assessment (SMA) Tool to evaluate its effectiveness in real-world scenarios.
- Run cyber-attack simulations to assess the framework's impact on decision speed and responsiveness in real-time threat scenarios.

Unlike prior AI-enabled Zero-Trust studies that primarily emphasise anomaly detection, continuous authentication, or conceptual security architecture, this study focuses on the decision-support gap between AI prediction and Zero-Trust enforcement. The proposed framework contributes a harmonisation mechanism through which heterogeneous ANN, CNN, and LSTM outputs are converted into a common risk score and then mapped to graduated enforcement actions. This enables cybersecurity consultants and analysts to interpret AI results as operational decisions rather than isolated model predictions. The framework is therefore positioned as a decision-support artefact for risk-based access control, maturity assessment, and security-policy refinement under Zero-Trust conditions.

This study builds on current work on Zero-Trust and AI by operationalising the outputs of different AI models into common risk scores that can be mapped to graduated enforcement decisions. The proposed framework, instead of treating the outputs of ANNs, CNNs and LSTMs as

independent detection results, translates the model-specific detection outcomes into unified risk semantics that can be used for ALLOW, MONITOR, MFA and DENY decisions. Rather than deep-learning models themselves, it is the harmonised risk scoring, dynamic Zero-Trust enforcement and stateful trust recalibration as part of a decision-support workflow that contributes.

The paper is organised as follows: Section 2 presents a literature review of cybersecurity frameworks, AI models, and Zero-Trust architectures, describing the key developments and challenges. Section 3 explains how to develop an AI-based business analytics framework, with a specific focus on integrating Artificial Neural Networks (ANNs), Convolutional Neural Networks (CNNs), and LSTM with Zero-Trust principles to enhance cybersecurity decision-making. The validation study results is included in Section 4, where the AI-driven framework was pitted against conventional cybersecurity models in terms of performance during simulated attacks. The implications of the findings, including the extent to which AI and Zero-Trust integration can enhance the efficiency of consulting and decision-making, as well as future research directions, are discussed in Section 5.

2. LITERATURE REVIEW

2.1. Cybersecurity Consulting Trends

The increasing frequency of cyberattacks has necessitated that the organisations adopt robust cybersecurity frameworks (Nifakos *et al.*, 2021). The evaluation of organisations' preparedness to address emerging threats has been well performed through risk maturity models, including the NIST, SOC2, and ISO-27001 frameworks, which serve as guidelines for best practices (Papachristofis *et al.*, 2024). These models, however, do not emphasise compliance and risk management as much as real-time and dynamic decision-making in cybersecurity operations (Chowdhury, 2025). Conventional risk management models fail to address the speed and complexity of contemporary cyber threats. The need to predict risks using analytics that foresee threats before they occur is the reason this gap is significant to this research.

Models like NIST have been useful in cybersecurity, as they offer a framework for responding to incidents; however, they are often grounded in pre-established procedures that can be considered reactive rather than proactive (Möller, 2023). Decision processes should be adapted to the dynamics of attack vectors, as the cyber threat environment is dynamic and requires more reactive, adaptive decision-making (Zaydi *et al.*, 2024). It is also here that AI-based frameworks have the potential to transform cybersecurity consulting by combining real-time predictive analytics with cybersecurity.

2.2. Zero-Trust Architecture

To overcome the limitations of perimeter-based security solutions, Zero-Trust architecture has emerged as one of the most popular models (Nahar *et al.*, 2024). In contrast to the traditional approach that trusts on network location, Zero-Trust does not trust people or objects no matter where they are located in or out of the network. These models are also

incredibly beneficial for scalability and predictive power, allowing cybersecurity experts to foresee and ward off attacks before they happen. Only after the requesting entity's continuous validation does one gain access to sensitive resources (Sarkar *et al.*, 2022). The Zero-Trust architecture is built on a principle of explicit verification and least privilege, and assumes that the breach can happen at some point – and that internal systems can be audited.

Zero Trust network, with AI based decision making is a new way of consulting cybersecurity (Ejeofobiri *et al.*, 2022). AI can dynamically assess user behaviour, device status, and network traffic to make real-time access control decisions. This allows for proactive action by cybersecurity teams, based on real-time threat and intelligence data, rather than on legacy rules or assumptions. However, in certain scenarios like Rule-based security systems or large networks where real-time, adaptive solutions are needed, there are no challenges that can't be solved through Zero-Trust and AI integration (Joshi, 2024).

The literature so far confirms that Zero-Trust combined with AI is a current research topic, but most of this literature deals with proving the feasibility of detecting risks, with architecture principles, and with policy sensitivity, not with the way the AI-based risk scores is actually integrated into the decision stage of Zero-Trust protection measures.

2.3. Predictive AI for Cybersecurity

AI is helpful in a variety of cybersecurity solutions, such as threat detection, anomaly detection, and predictive risk analytics (Patil, 2024). Some of the AI models attracting the attention of the fast-changing cybersecurity landscape are ANNs, CNNs, and RNNs (Sarker, 2021). The models have a number of benefits, including their ability to be scaled up and their wide-ranging predictive capabilities, which empowers cybersecurity specialists to predict and prevent attacks even before they happen.

2.3.1. ANN Techniques for Risk Classification

Artificial Neural Networks (ANNs) have been found to be useful in classification of the network as well as identification of potential historical data (Goel *et al.*, 2023). ANNs can deal with huge datasets and can discover delicate patterns that indicate potential attacks, making them efficient to foresee attack probabilities and user behaviour ratings (Alzaabi & Mehmood, 2024). However, it is still very challenging to train such models, especially if the data is imbalanced. Furthermore, the feature of ANNs that they are not always interpretable is important in cybersecurity to understand why a particular decision has been taken (Nair, 2023).

2.3.2. CNN for Log Pattern Recognition

Convolutional Neural Networks (CNNs) have extensive applications in image processing and have also been useful for cybersecurity log pattern recognition (Hagras, 2023). CNNs can locate indicative spatial features. One such instance is that they can detect DDoS attacks or command-and-control communications from malware. CNNs have some limitations, although they are useful for processing large amounts of data in real time and for the necessity of large labelled datasets (Alzubaidi *et al.*, 2021).

2.3.3. RNN/LSTM for Sequential Threat Forecasting

Recurrent Neural Networks (RNNs) and, more precisely, Long Short-Term Memory (LSTM) networks can be used to process sequential data, *e.g.*, patterns of attack over time (Yunita *et al.*, 2025). RNNs are capable of modelling long-term data and, therefore, they are well-suited for forecasting the development of cyberattacks or for identity identification, access anomalies. These models, however, are computationally intensive and require considerable amounts of. The duration of their training is measured by their application in real-world settings, including large organisations and high volumes of traffic (Ahmed *et al.*, 2023).

2.4. Simulation and Scenario Techniques in Cybersecurity

Cybersecurity models should be tested in a simulated environment. The number of methods that one can apply to facilitate a variety is also manifold. Cybersecurity experts can simulate possible cyber-attacks and test the performance of their defences when responding to identified attacks, including Red-Team simulations, attack-defence digital twins, and more. Red-team exercises assess a system's security by mimicking the actions of an adversary, helping uncover vulnerabilities that are not detected in normal system monitoring or vulnerability scans. Similarly, digital twins create IT system replicas to assess how they can resist attacks (Alcaraz & Lopez, 2022).

However, these approaches are invaluable for intelligence creation. In most scenarios, they are unable to forecast what AI models can, as they can foresee emerging threats before they come into existence (Bécue *et al.*, 2021). Although the fields of AI and Zero-Trust security have advanced, there are still gaps in cybersecurity. A large number of models are provided, based on either AI-based prediction or business consulting analysis.

2.5. Gaps Identified

Despite improvements in AI and Zero-Trust security, cybersecurity still has gaps. Numerous existing models are committed to either AI-powered prediction or business consultation analysis, but none of them is integrated with Zero-Trust verification and performance measurement through simulations (Ajish, 2024). The failure to integrate exposes organisations to the risk of attack through means of circumventing the traditional perimeter. security. Besides, real-time, data-driven decision-making systems are not available at all, which challenges organisations' ability to respond efficiently and promptly.

Furthermore, the current models fail to provide a comprehensive outlook on the cybersecurity decision-making process, which involves risk classification and the speed, accuracy, and cognitive load of the process. The more complex the threats are, the greater the cognitive overload on decision-makers, and the greater the likelihood of making a quick, accurate decision (Reale *et al.*, 2023). This may be addressed with AI-based models that can automate routine processes and support time-sensitive decision-making. The quality of such decisions, especially those under time stress, is, however, not well defined in the current frameworks.

The AI-based research model used in this study combines ANNs, CNNs, LSTM, and Zero models. The concept of trust,

which enables predictive risk analytics and real-time decision support. This combination improves cybersecurity decision-making efficiency by reconciling the limitations of more conventional models, such as NIST incident response and OODA, which can be slow or rigid in responding to fast-paced situations. threats. The framework is a dynamic, agile system that unites AI with Zero-Trust to anticipate and avert attacks and provides continuous access control to re-estimate threats and adjust access controls.

3. RESEARCH METHODOLOGY

3.1. DSR Justification

The research employs a Design Science Research (DSR) methodology, which is well-suited to the research: the development and testing of a new artefact for cybersecurity consulting. DSR is a methodological process consisting of creating, evaluating, and developing artefacts, *i.e.*, structures or models, to solve issues found in the real world. The framework is pragmatic, cyclical, and continues to evolve to address the needs of cybersecurity decision-making.

Predictive risk analytics was incorporated into the business analysis framework developed with AI in this study. Moreover, with the Zero-Trust architecture, we can thus develop a solution to enhance cybersecurity decision-making. The iterative format of DSR enables testing and validation in the natural environment, which is very helpful because it allows consultants to respond to new threats and vulnerabilities.

3.2. Research Framework

The research framework is designed to support the creation and examination of AI-driven cybersecurity. consulting framework. This model consists of the following steps:

Problem Identification: It is argued that current cybersecurity consulting needs an adaptive, AI-powered decision-making system to overcome the limitations of existing models in dealing with evolving threats.

Development: The solution includes developing AI models for risk categories (ANN), log pattern recognition (CNN), and sequential threat prediction (LSTM). These models are combined with the Zero-Trust architecture, which bases access decisions on real-time risk data.

Demonstration: Simulated attacks are used to demonstrate the framework's operation in real cyber-attack scenarios (phishing, credential stuffing, and malware execution).

Evaluation: To analyse the effectiveness of the framework in verifying the maturity of Zero-Trust, checking the accuracy of AI model predictions, and the overall verifiability of the decision-making process, the Security Maturity Assessment (SMA) Tool is utilised.

The framework layer assists in managing the decisions made by the consultants is based on business analysis, which can help put AI-generated risk assessments into the context of governance and Zero-Trust policy frameworks (Fig. 1). Its main purpose is to increase the interpretability, prioritisation and strategic alignment of security choices, not to calculate

financial results. In this regard, the empirical validation in this paper focuses on operational decision-support effectiveness rather than financial performance indicators.

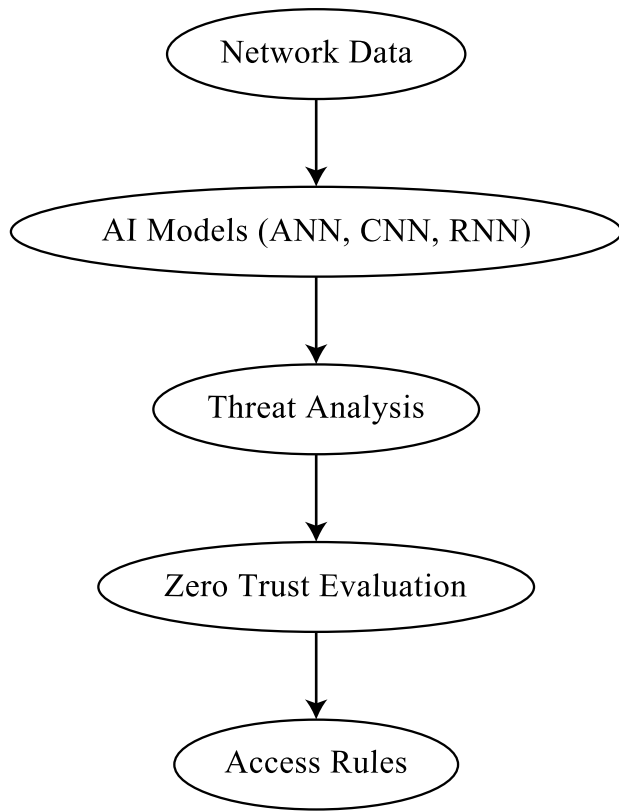


Fig. (1). Flowchart of the ai-driven cybersecurity framework.

3.3. Data Description

The study used publicly available UNSW-NB15 dataset were identity access logs and system events gathered from network traffic and user interactions, provided in the files UNSW_NB15_training-set.csv and UNSW_NB15_testing-set.csv. The official training set had 82,332 records, and a test set of 175,341 records. The training set was split into training and validation subsets, with 80% of the data used for training and 20% for validation, stratified randomly with random_state = 42. The ANN and CNN models have 42 event-level features, and the LSTM uses 10-step temporal windows, yielding 17,534 test sequences.

After preprocessing, the final dataset consists of 42 features. These features included authentication attributes, session metadata, network indicators, behavioural variables, access history information and risk-related contextual variables. After pre-processing, the number of event-level observations for ANN and CNN modelling was approximately 175,341, and the number of sequence-level observations for LSTM modelling was 17,534 after processing the temporal window.

The records for each event were then resized into a 10-step fixed-length temporal sequence for LSTM modelling. The train

and test files were separated before generating the sequences to avoid temporal leakage. The sequence windows were defined independently in each partition, and no windows spanned from the training set to the test set. Thus, the events included in the training sequences were not repeated in the testing sequences. This enabled temporal windows with no overlap and ensured none between the training and testing data.

Event-level samples were extracted from the UNSW-NB15 test set and used to evaluate Random Forest, ANN, and CNN, and 10-step temporal sequence windows were generated from the same test partition for LSTM evaluation. As a result, the raw counts of false negatives cannot be compared across all models as they are based on varying numbers of test units. To prevent misinterpretation, the false-positive and false-negative rates are also reported, along with the raw FP and FN counts. Based on this, LSTM results should be viewed as sequence-level, while those of ANN and CNN should be considered event-level.

3.3.1. Data Pre-processing and CNN Input Preparation

Before model training, the UNSW-NB15 records were pre-processed to ensure consistency across the ANN, CNN, LSTM, and Random Forest experiments Table 1. The predefined training and testing files were kept separate throughout pre-processing to avoid information leakage. Missing values, duplicate records, and non-informative identifiers were checked before feature transformation. Categorical variables were encoded using one-hot encoding, while numerical variables were normalised using min-max scaling. The scaling parameters were fitted only on the training partition and then applied to the validation and test partitions. This ensured that no information from the test set influenced the training process.

The final event-level feature matrix contained 42 input features after encoding and pre-processing. For ANN and Random Forest modelling, these 42 features were used directly as tabular input. For CNN modelling, the same 42-dimensional event-level vector was converted into a fixed two-dimensional representation so that convolutional filters could be applied consistently. Because 42 features do not form a perfect square matrix, zero-padding was applied to extend each feature vector to 49 values, after which it was reshaped into a 7×7 single-channel matrix. Thus, the CNN input shape was defined as $7 \times 7 \times 1$ for each event. The padded values contained no attack information and were included only to preserve a consistent spatial input structure for convolution.

The CNN pre-processing procedure was designed to allow local neighbourhood patterns among transformed network and behavioural features to be learned by convolutional filters. Although the UNSW-NB15 dataset is tabular rather than image-based, reshaping the feature vector into a compact two-dimensional matrix allowed the CNN to model local feature interactions that may not be captured by a fully connected model alone. The same training, validation, and testing partitions were used for the CNN as for the ANN to ensure fair event-level comparison. Labels were binary-coded, where 0 represented normal traffic and 1 represented attack traffic.

Table 1. Model parameter settings used in the experiments.

Model	Input Shape	Main Layers / Estimator	Optimiser	Learning Rate	Batch Size	Epochs	Output Activation	Loss / Criterion
Random Forest	42 features	100 trees	Not applicable	Not applicable	Not applicable	Not applicable	Majority voting	Gini impurity
ANN	42 features	Dense-ReLU-Dropout-Dense-ReLU-Dropout-Dense	Adam	0.001	64	15	Sigmoid	Binary cross-entropy
CNN	$7 \times 7 \times 1$	Conv2D-MaxPool-Conv2D-Flatten-Dense-Dropout-Dense	Adam	0.001	64	15	Sigmoid	Binary cross-entropy
LSTM	10×42	LSTM-Dropout-Dense-Dense	Adam	0.001	64	15	Sigmoid	Binary cross-entropy

Note: The CNN input was obtained by padding the 42-feature event vector to 49 values and reshaping it into a $7 \times 7 \times 1$ matrix. The LSTM input was generated using non-overlapping 10-step temporal windows. The Random Forest model was used as a baseline classifier, while the ANN, CNN, and LSTM models were used as AI risk-scoring components in the Zero-Trust decision-support framework.

For the LSTM model, the pre-processed event-level records were converted into non-overlapping 10-step temporal windows. The train and test files were separated before sequence generation so that no temporal windows crossed from the training partition into the test partition. This reduced the risk of temporal leakage and ensured that LSTM performance reflected sequence-level generalisation rather than overlap between training and test samples.

3.4. AI Model Development

The AI models in this framework were developed using a consistent supervised-learning workflow. The Random Forest model was used as a baseline, while ANN, CNN, and LSTM models were used as neural-network components of the proposed Zero-Trust decision-support framework. All models were trained using the same binary target variable, where 0 indicated normal traffic and 1 indicated attack traffic. The training set was divided into training and validation subsets using an 80:20 stratified split with `random_state = 42`. Model performance was evaluated on the held-out UNSW-NB15 testing set.

3.4.1. Random Forest

The Random Forest classifier was implemented as a conventional machine-learning baseline to compare neural-network performance against a non-deep-learning approach. The model used 100 decision trees, Gini impurity as the splitting criterion, and `random_state = 42` for reproducibility. Class imbalance was addressed by reporting weighted precision, recall, and F1-score in addition to accuracy and AUC. The baseline model provided a reference point for assessing whether the ANN, CNN, and LSTM models offered additional predictive value.

3.4.2. ANN Model

The ANN model was designed to classify event-level records as normal or malicious using the 42 pre-processed input features. The network consisted of an input layer, two dense hidden layers, and one sigmoid output layer. The hidden layers used ReLU activation to capture nonlinear relationships among

network, behavioural, and session-related features. Dropout regularisation was included to reduce overfitting during training. The final output layer used a sigmoid activation function to generate a probability score between 0 and 1, which was then interpreted as the ANN-generated risk score.

The ANN was trained using the Adam optimiser with a learning rate of 0.001, binary cross-entropy loss, a batch size of 64, and 15 training epochs. Validation accuracy and validation loss were monitored at each epoch to assess generalisation. The decision threshold for binary classification was set to 0.50, while the continuous probability output was retained for the Zero-Trust risk-score harmonisation layer.

3.4.3. CNN Model

The CNN model was used to learn local feature-interaction patterns from the reshaped $7 \times 7 \times 1$ representation of each event-level record. The CNN architecture consisted of convolutional layers, max-pooling, flattening, dense layers, dropout, and a sigmoid output unit. The first convolutional layer used 32 filters with a 3×3 kernel and ReLU activation. This was followed by a max-pooling layer to reduce the feature map size. A second convolutional layer with 64 filters and a 3×3 kernel was then used to capture higher-level local feature interactions. The extracted feature maps were flattened and passed to a dense layer before final binary classification.

The CNN was trained using the Adam optimiser with a learning rate of 0.001, binary cross-entropy loss, a batch size of 64, and 15 epochs. Dropout was applied after the dense layer to reduce overfitting. The final sigmoid output represented the probability that an event belonged to the attack class. This probability was used both for binary classification and as the CNN risk score for subsequent Zero-Trust policy mapping.

3.4.4. LSTM Model

The LSTM model was developed to capture temporal dependencies across sequential network events. The pre-processed records were grouped into non-overlapping 10-step windows, and each window was labelled according to the corresponding attack status. The LSTM architecture included

one LSTM layer followed by dropout, a dense hidden layer, and a sigmoid output layer. The LSTM layer was used to learn sequential patterns that may indicate evolving attack behaviour, such as reconnaissance followed by exploitation or repeated abnormal access attempts.

The LSTM model was trained using the Adam optimiser with a learning rate of 0.001, binary cross-entropy loss, a batch size of 64, and 15 epochs. The sigmoid output produced a sequence-level attack probability. Because the LSTM used sequence-level test units rather than event-level test units, its raw false-positive and false-negative counts were not directly compared with ANN, CNN, or Random Forest. Instead, rates such as FNR, weighted F1-score, and AUC were used for cautious comparison.

3.5. Risk Score Harmonisation Across Heterogeneous AI Models

The suggested structure uses heterogeneous AI models with native outputs of different structures and semantics. To allow the use of a single Zero-Trust policy, a risk score harmonisation layer is added to convert all model outputs to a shared scalar risk representation, $R \in [0, 1]$. This risk score is harmonised, enabling consistent, model-agnostic decision-making in the Zero-Trust Policy Decision Engine.

3.5.1. ANN Output Mapping

ANN gives a binary result of benign or malicious activity. The model produces a sigmoid activation output, which is used to transform this output into a probabilistic risk score:

$$R_{ANN} = P(y = 1 | x) = \sigma(z)$$

where R_{ANN} is the ANN-generated risk score, $y = 1$ represents malicious activity, x is the input feature vector, and $\sigma(z)$ is the sigmoid output.

3.5.2. CNN Output Mapping

The CNN generates a multi-class forecast based on threat intensity levels. The risk score is harmonised as a weighted average of the probabilities of classes:

$$R_{CNN} = \sum_{k=1}^K p_k w_k$$

where p_k is the predicted probability of class k , w_k is the severity weight assigned to that class, and K is the total number of threat classes.

3.5.3. LSTM Output Mapping

The LSTM model provides a temporal distribution of threat likelihoods. With the help of temporal aggregation, the harmonised risk score is calculated:

$$R_{LSTM} = \frac{1}{T} \sum_{t=1}^T P(y_t = 1 | x_t)$$

where T is the temporal window length and $P(y_t = 1 | x_t)$ is the predicted attack probability at time step t .

The harmonised risk scores of the AI models are similar across architectures and output semantics, but they all share a common probabilistic interpretation: the probability or degree of compromise. Such an abstraction enables heterogeneous models to be used interchangeably within the Zero-Trust enforcement logic, retaining all their respective analytical strengths.

3.6. Integration with Zero-Trust Architecture (Zero-Trust Architecture)

The incorporation of the AI models into the Zero-Trust architecture is a significant aspect of this structure. The Zero-Trust Policy Decision Engine (ZTPDE) is a dynamically assessed policy engine that receives real-time data from AI models. The probability of an attack is assessed using the risk score generated by the ANN model, whereas CNN and LSTM models provide an understanding of attack patterns and sequences. The information is traced back to the Zero-Trust decision-making process to ensure that access is granted only to those who meet the least-privilege requirements and complete all necessary verification procedures. In this study, the rule-based security systems are conventional signature- or threshold-based detection systems that do not include adaptive AI-based risk scoring or enforcement of Zero-Trust policies.

A combination of Zero-Trust principles and AI enables dynamic, risk-based authentication: access rules are continually adjusted based on the current risk assessment. This ensures that organisations can react swiftly to incoming threats and counter potential breaches before they escalate. In this study, the outcome of automated access control generated by the Zero-Trust Policy Decision Engine (ALLOW, MONITOR, MFA, or DENY) in response to an AI-generated risk assessment is defined as a decision. The quality of the decision is thus measured with respect to (i) the accuracy in risk classification of the underlying risk, (ii) the suitability of the action taken as a result of the risk level and (iii) the consistency and promptness of the decision action to dynamic threat conditions.

3.7. Use of the Security Maturity Assessment (SMA) Tool

The Security Maturity Assessment (SMA) Tool is used to assess whether the framework effectively measures Zero-Trust maturity. The tool gauges the extent to which AI-generated choices comply with the principles of Zero Trust and compares AI models' outputs against well-defined cybersecurity standards. The SMA Tool allows tracking AI results against Zero-Trust capability requirements, providing valuable feedback on the framework's performance and areas for improvement.

Even though the concept of Zero-Trust maturity is generally understood as a continuous, multidimensional process, SMA in this work focuses on enabling capabilities rather than depth of maturity. A score of 1 means the framework provides a foundational Zero-Trust capability that is operational, and a score of 0 means it does not. The estimation thus depicts the achievement of low-level viable Zero-Trust functionality, rather than organisational maturity.

In this study, the Security Maturity Assessment (SMA) Tool is used as a structured capability validation tool, not as a quantitative maturity scoring model. Every dimension of SMA reflects the presence or absence of a core Zero-Trust capability (e.g., continuous verification, least privilege, adaptive access, and risk-based policy enforcement). In line with this, the assessment results are reported in a binary manner to indicate whether the proposed framework operationally supports the capability.

Binary scoring is adopted to reflect capability enablement rather than maturity depth. The objective of the SMA in this context is to validate whether the proposed framework operationally supports each Zero-Trust principle, not to rank organisations along a maturity continuum. As such, weighting schemes and composite maturity indices are intentionally not applied, as they are typically used in organisational benchmarking studies rather than artefact validation.

3.8. Simulation & Scenario Analysis

The evaluation process includes simulated sequences of cyber-attacks. These codes consist of the popular attack types: phishing, credential stuffing, and malware execution. These simulations are tested within the framework, and key performance indicators (KPIs) are measured, including detection speed, analyst response time and false-positive reduction. The simulations are used to test the effectiveness of the AI-based decision-making framework in real time, identify the threats, and help mitigate them. For example, the detection speed of an AI model is the time it takes to identify a cyber-attack upon its occurrence.

In contrast, the speed of analyst response is measured by the time it takes cybersecurity specialists to act on the AI model's recommendations. Also, the false-positive rate is monitored to ensure the system does not inundate security teams with false messages. The simulation scenarios are intended to represent various phases of the cyber-attack lifecycle, rather than isolated threat labels. All simulated scenarios are operational versions of the broader threat classes described elsewhere in the paper. The mapping between the threat classes in the high-level and their respective simulated attack scenarios is summarised in Table 2, thereby ensuring consistency between the conceptual threat taxonomy and the empirical evaluation design.

3.9. Mathematical Integration and Data Flow

The mathematical combination of AI-generated outputs with the Zero-Trust framework works in the following way:

- 1. Risk Score Calculation:** The AI models (ANN, CNN, LSTM) provide a risk score that quantifies the likelihood of an attack. These scores are normalised to the range 0 to 1, with 1 indicating a high probability of an attack.
- 2. Policy Mapping:** These are then mapped onto Zero-Trust policies. For example, a score above a specific threshold (e.g., 0.75) could trigger multi-factor authentication (MFA) or limit access in accordance with the least privilege doctrine.

- 3. Access Rule Enforcement:** After the AI model maps to a decision, the Zero-Trust system enforces access rules based on the result, dynamically responding to evolving threat conditions.

Table 2. Threat taxonomy and simulation mapping.

High-Level Threat Class	Simulated Attack Scenarios
DDoS	High-frequency malicious traffic patterns
APTs	Reconnaissance, Backdoor, Shellcode
Insider Threats	Credential stuffing, Privilege misuse
Initial Access Attacks	Phishing, Credential stuffing
Malware-Based Attacks	Malware execution

3.10. Pseudo-Algorithm

- 1. Input:** Raw network traffic data, user behaviour logs, system events
- 2. Pre-processing:** Clean and normalise data
- 3. Model Prediction:**
 - ANN: Identify the risk as an innocuous or malicious one.
 - CNN: Find patterns of attack on logs.
 - LSTM: Forecast sequential threats
- 4. Risk Score Calculation:** The model's outputs are normalised to a risk score (0-1).
- 5. Zero-Trust Evaluation:** Map risk to the Zero-Trust policy decision engine.
- 6. Access Rule Enforcement:** Implement access control (e.g., deny, MFA, monitor)

4. PROPOSED AI-DRIVEN BUSINESS ANALYSIS FRAMEWORK

The proposed framework is developed as a cybersecurity operations decision-support system using AI. It combines the risk-prediction capabilities of ANNs, CNNs and LSTMs with a Zero-Trust Policy Decision Engine for adaptive access-control decisions. The framework is not intended to give a holistic business-consulting or financial-optimisation model. Rather, it is about combining disparate AI outputs to create a single set of risk scores, then correlating those scores with different levels of enforcement during a simulated cyberattack.

4.1. Framework Overview

The suggested AI business analysis model easily fits into a Zero-Trust framework, while also introducing predictive risk analytics to improve cybersecurity decision-making for consulting firms. This design was broken down into 5 main layers: Data Collection Layer, AI Prediction Engine, Zero-Trust Decision Layer, Business Intelligence Layer, and Consulting Dashboard. All layers are needed for cybersecurity advisors to have access to accurate, up-to-date data that is relevant enough to aid in informed decisions.

At the core of the framework is the Data Collection Layer, responsible for collecting all critical data from network traffic logs, system events, user actions, and multi-factor authentication (MFA) signals. This data is updated and purged on an on-going basis to ensure that it is retained and available for analysis. The data was fed to the AI Prediction Engine which executed a machine learning model with Artificial Neural Networks (ANNs), Convolutional Neural Networks (CNNs), and LSTM (a particular Recurrent Neural Networks (RNN)). These models forecast the chances of cyberattacks and, thus, help cybersecurity experts anticipate threats in advance of their appearance. The forecasts are coupled to the Zero-Trust Decision Layer, which allows access and applies security policies based on the AI engine's real-time analysis of them.

The Zero-Trust Decision Layer includes continuous checking of all access requests (internal and external). This layer is based on a least-privilege model (EL), so that users have access only to sensitive systems and data that are approved. The Business Intelligence Layer incorporated AI models and Zero-Trust policy outcomes and deliver substantive information to consultants. It has cost-risk forecasting tools, ROI analysis of cybersecurity investments and recommended security controls. Last but not least, the consulting dashboard is the interface that cybersecurity professionals can use to access real-time visualisations of risks and threats, as well as mitigation strategies. This layer allows consultants to quickly and accurately decide on the security data, giving a wider overview of the organisation's security position.

4.2. Predictive Risk Analytics Layer

The most crucial part of the framework is the Predictive Risk Analytics Layer that calculates threat probabilities and produces a risk heatmap in real time. Using AI models, this layer allowed threats from different cyberattacks, such as credential stuffing, DDoS attacks and phishing, to be identified. The threat probability score is a numerical value that reflects the likelihood of an event occurring, given historical and current dynamics of the network. This score can be used to prioritise security activities with the highest likelihood of occurrence and deal with them in advance.

AI models generate risk heat mapping to show the distribution of risk across the organisation's IT assets. These heat maps help cybersecurity experts identify high-risk systems, areas for weak access and/or areas that need urgent policy intervention. They are continually refreshing themselves with the information they receive, giving cybersecurity consultants a real-time view of where the network is weak. The tool for visualisation makes it easy to identify high risk areas and to allocate resources accordingly for security officers. Predictive risk analytics within the wider Zero-Trust model can enable the consultant to detect potential risks in time, enabling a faster response and greatly curbing the impact of the cyberattack.

4.3. Zero-Trust Decision Layer

The Zero-Trust Decision Layer combines AI-based forecasts with Zero-Trust security concepts to provide an all-

in-one solution for access control and user authentication. Some of the features that render this layer as one of the most crucial ones include the ability to utilise multi-factor authentication (MFA) in case an abnormal pattern or a possible threat is identified by the AI models. To illustrate, the Zero-Trust system can provide an additional verification process for granting sensitive resources to the user based on a risk score produced by the AI model that exceeds a predetermined threshold. This ensured that, even if an attacker breaches the network perimeter, multiple authentication steps were required before they gain access to the network.

The Zero-Trust Decision Layer can also be used to block access to suspicious users or devices in addition to MFA. In case of abnormal behaviour identified by the system, *e.g.*, abnormal time or place of logging in, access to relevant systems can be cancelled on the spot. The Zero-Trust Architecture also includes a trust score recalibration option; it continuously assesses the reliability of users and devices based on behaviour and threat levels. Access policies are also constantly updated to reflect the most current risk assessment through recalibration, enabling security to be applied dynamically in real time.

4.4. Business Consulting Decision Support Layer

The Operational Decision-Support Layer helps convert AI-driven risk scores into meaningful Zero-Trust decisions. It is not there to optimise financial returns or to forecast business performance, but to facilitate cybersecurity operations by enhancing the consistency, proportionality and timeliness of access-control decisions. The outputs of the ANN, CNN and LSTM models are translated into common risk semantics in this layer and sent to the Zero-Trust Policy Decision Engine. The engine then translates each risk score into an enforcement action from four options: ALLOW, MONITOR, MFA, or DENY.

This layer provides operational benefits by reducing the need for binary allow/deny decisions and enabling a graduated response when encountering events that are uncertain or pose increased risk. Low-risk access requests can be granted; medium-risk access requests can be tracked; high-risk access requests may require multi-factor authentication; and critical-risk access events can be denied. The graduated enforcement process enables security teams to handle risk without interfering with the normal functioning of activities. It is also aligned with Zero Trust, as it continually verifies trust based on observed behaviour and changing threat conditions.

The layer can also enable future cybersecurity consulting applications by providing an understandable output that analysts or consultants can review during a security evaluation. However, in this study, the use of dashboards, ROI estimation, cost-risk forecasting, and business KPI optimisation are considered practical extensions of the deployment for the future and are not empirically proven. Only the operational capability of the decision-making framework, model performance, risk-score harmonisation, access-control behaviour, trust recalibration and policy-decision latency are considered in the present evaluation.

4.5. Operational Workflow

The business analysis process of the proposed AI-based system follows a cyclic scheme, is smooth and involves all levels of the unified system. The first level of the workflow is data collection, where real-time data is gathered from network logs, user behaviour, and system events, and then cleaned. This information is then fed to the AI models to generate risk forecasts, which are sent to the Zero Trust Decision Layer for review, and the control decision is determined.

As the AI models run, the system continuously updates trust scores, and MFA or access is enabled or denied accordingly. They are then presented to the consultants through the Consulting Dashboard, where the risk heatmaps, cost-risk forecasts, and recommendations can be viewed. The resulting decisions are returned as operational outputs to support analyst review and refine security policy. The framework thus helps support cybersecurity teams by enabling the conversion of AI-driven risk scores into enforceable Zero-Trust actions. The output of this study is limited to the empirical context of consulting dashboards for resource allocation planning or cost-risk forecasting, and these aspects of the topic should be explored further in the future.

5. RESULTS

To demonstrate the applicability of the proposed AI-enhanced business analysis framework, it is possible to simulate cyber-attack cases. These simulations employ a holistic architecture approach to understand the framework and improve detection, response time, and decision-making accuracy in real-life cybersecurity threats. Results from these simulations demonstrate significant improvements in attack detection and response, and the framework can make decisions more quickly and effectively than the traditional security framework.

In cybersecurity consulting, the objectives of business analysis are directly reflected in these results, despite the technical aspects of performance being the focus of the experimental evaluation. Detection accuracy, Zero-Trust enforcement behaviour, trust recalibration, maturity progression, and decision latency are metrics of operation ineffectiveness, governance maturity and decision reliability, which are the focus of consultant-based risk evaluation and advisory services.

5.1. ANN Model Accuracy and Performance

The Artificial Neural Network (ANN) model achieved a test accuracy of 87.94%, as reported in the classification report. The model had a steep training vs. validation accuracy curve, similar to the ANN's accuracy plot, because both the training and validation datasets increased steadily. The model also showed initial overfitting, as indicated by a rapid increase in training accuracy, but this decreased as training continued. The validation accuracy then approached the train model's suitability for categorisation, as shown in the loss curve, which showed a sharp drop in both the training and validation losses,

indicating successful learning from the input features. Figs. (2 and 3) show the accuracy and loss convergence curves, respectively, and illustrate the ANN model's learning behaviour over training epochs.

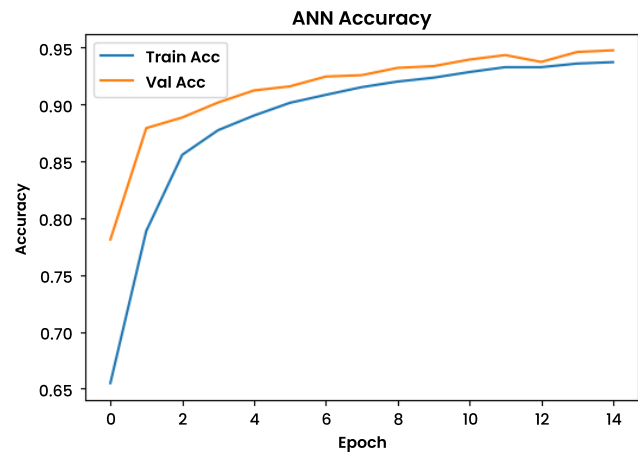


Fig. (2). ANN accuracy vs epoch.

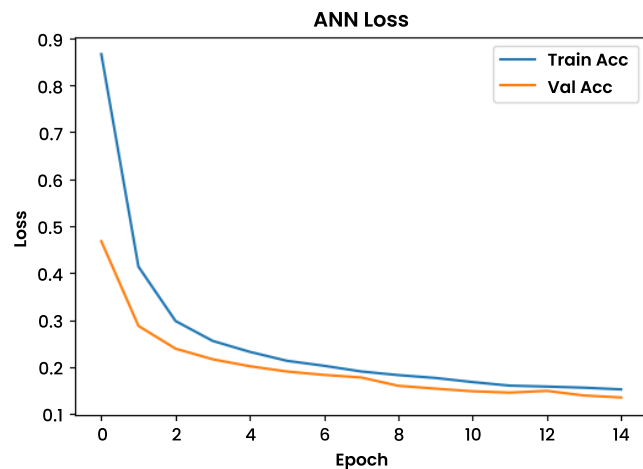


Fig. (3). ANN loss vs epoch.

The ANN model's low false-positive rate indicates it is more precise for regular events (label 0), with a precision of 0.7351. In contrast, for attack events (label 1), the model has a precision of 0.9852. This is the model's ability to categorise attack events accurately, with a recall of 0.9732 on normal data and 0.8354 on attack events, indicating that the model is reliable at detecting actual attacks with a low false-positive rate. The 0-score label was 0.8375, and the 1-score label was 0.9041, indicating high overall classification performance, especially in attack detection. The model had an AUC (Area Under the Curve) of 0.9797, indicating very high discrimination between attack and non-attack labels. Fig. (4) presents the detailed classification behaviour of the ANN model, whereas Figs. (5 and 6) present the confusion matrix and the ROC characteristics, respectively.

=== ANN Classification Report (Test) ===				
	precision	recall	f1-score	support
0	0.7351	0.9732	0.8375	56000
1	0.9852	0.8354	0.9041	119341
accuracy			0.8794	175341
macro avg	0.8601	0.9043	0.8708	175341
weighted avg	0.9053	0.8794	0.8829	175341

Fig. (4). ANN classification report.

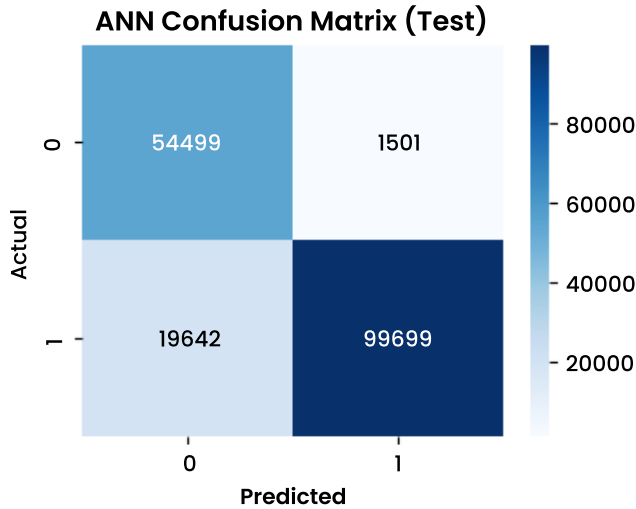


Fig. (5). ANN confusion matrix.

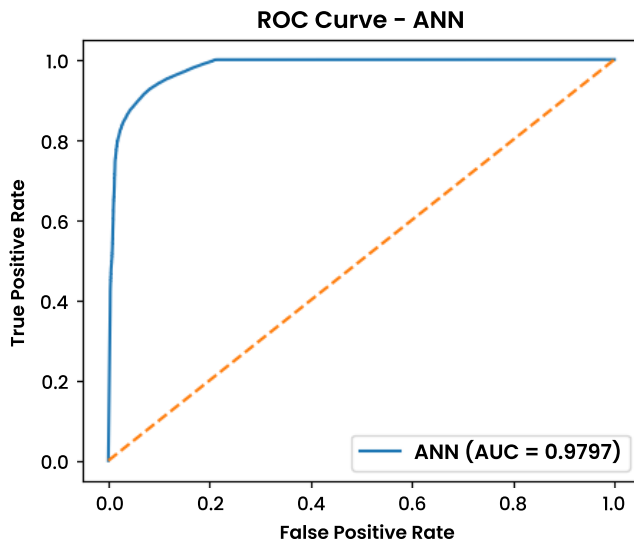


Fig. (6). ROC curve – ANN.

The ANN confusion matrix showed that the model identified 54,499 instances of regular traffic and 99,699 instances of attack across the total 175,341 test samples, with false positives of 1,501 and false negatives of 19,642. The ANN ROC curve shows a high True Positive Rate (TPR) and a low

False Positive Rate (FPR), resulting in an AUC of 0.9797 and reporting efficient separation of attack and non-attack cases in the tested dataset.

5.2. CNN Model Accuracy and Performance

The Convolutional Neural Network (CNN) achieved 87.10% classification accuracy and consistent generalisation behaviour, with validation accuracy paralleling training accuracy. The precision and loss curves of the CNN gradually increased and showed only slight overfitting, as indicated in the classification report. Its accuracy in identifying regular traffic was 0.7174, and its accuracy in identifying attacks was 0.9906, demonstrating its greater capacity to assign attack labels. Recall for regular traffic was 0.9834, and for attack traffic was 0.8183, indicating that the model has a strong capacity to identify real attacks and a reasonable capacity to identify regular traffic. The F1-score for label 1 was 0.8962, which means that the attack instances are identified with the same behaviour in the circumstances considered in the evaluation. The CNN's AUC was 0.9827, indicating a strong ability to differentiate between attack events and ordinary traffic. Figs. (7 and 8) illustrate the training dynamics of the CNN model and its generalisation behaviour, respectively, using accuracy and loss curves over the epochs.

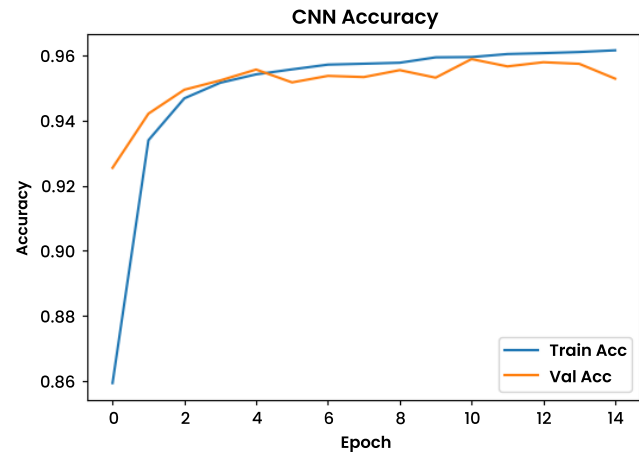


Fig. (7). CNN accuracy vs epoch.

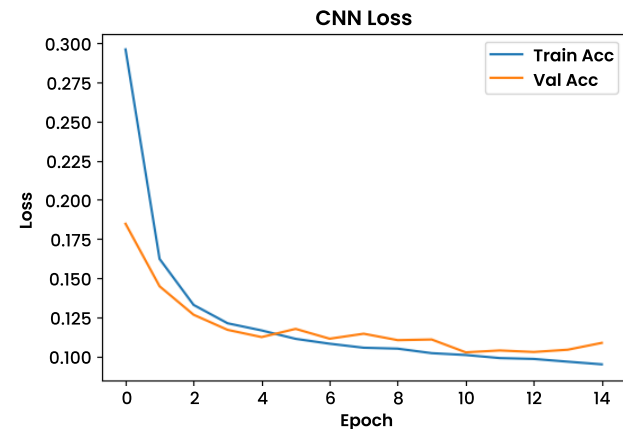


Fig. (8). CNN loss vs epoch.

Figs. (9-11) show the CNN classification report, confusion matrix and ROC curve, respectively, which shed light on the model's behaviour with respect to class-wise detection.

=== CNN Classification Report (Test) ===				
	precision	recall	f1-score	support
0	0.7174	0.9834	0.8296	56000
1	0.9906	0.8183	0.8962	119341
accuracy			0.8710	175341
macro avg	0.8540	0.9008	0.8629	175341
weighted avg	0.9033	0.8710	0.8749	175341

Fig. (9). CNN classification report.

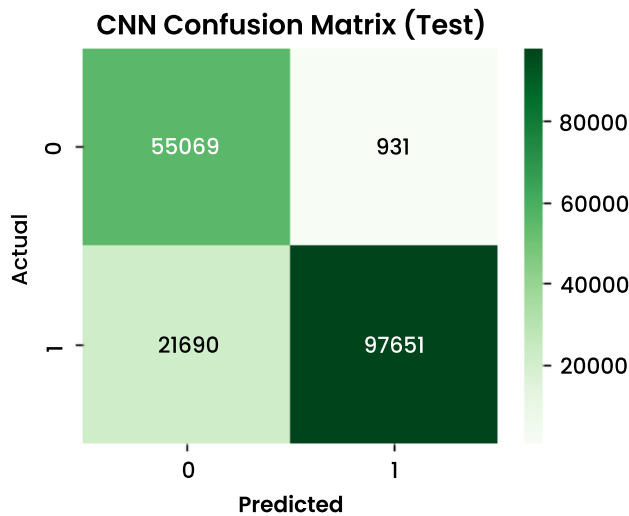


Fig. (10). CNN confusion matrix.

These results are further supported by the CNN confusion matrix, which shows that 55,069 instances were correctly classified as normal and 97,651 as attacks, with 931 false positives and 21,690 false negatives. This shows that the CNN has done a great job of minimising errors and is resilient to attacks. The CNN ROC curve showed an AUC of 0.9827, confirming the model's stable classification behaviour across attack and non-attack classes.

For reproducibility, the CNN results reported in this section are based on the 42-feature event-level test matrix after zero-padding to 49 values and reshaping each event into a $7 \times 7 \times 1$ single-channel input representation.

5.3. LSTM Model Accuracy and Performance

The accuracy of the Long Short-Term Memory (LSTM) network is 87.58% at test time. The LSTM accuracy and loss curves showed that both improved steadily as the model trained, as observed in the LSTM plots. The training and validation accuracies of the LSTM both converged, indicating that the model generalised well to unseen data. Figs. (12 and 13) demonstrate the LSTM model's temporal learning, showing convergence of accuracy and loss during training.

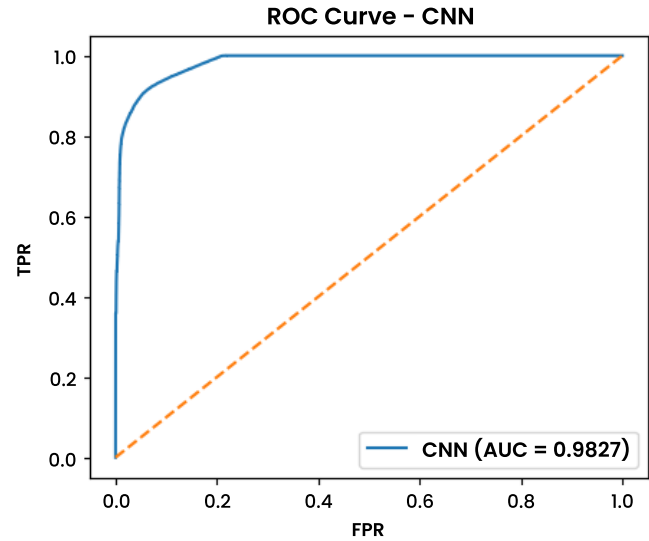


Fig. (11). ROC curve – CNN.

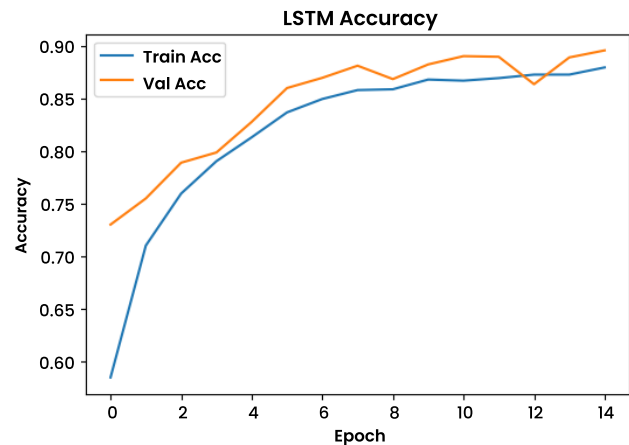


Fig. (12). LSTM accuracy vs epoch.

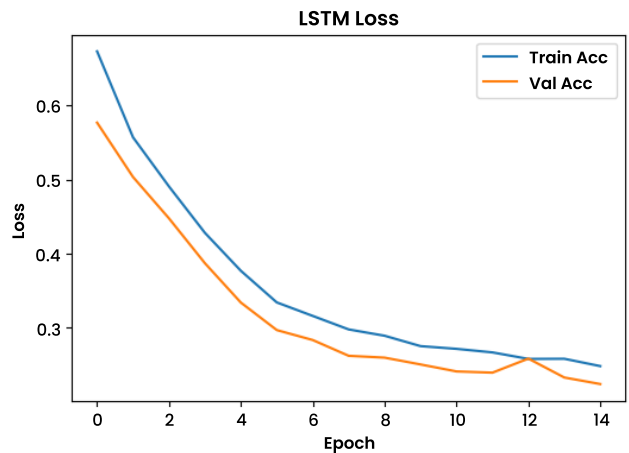


Fig. (13). LSTM loss vs epoch.

The LSTM classification report showed that the precision for regular events and attacks is 0.7397 and 0.9696, respectively, and the recall for regular events and attacks is 0.9439 and 0.8438, respectively. The F1-score of regular events was 0.8294, and the F1-score of attack events was 0.9023, indicating that LSTM was especially good at detecting attacks with a minimum number of false negatives. The LSTM AUC was 0.9720, slightly lower than the ANN and CNN AUCs but within the same performance range. Figs. (14-16) represent LSTM classification performance, the confusion matrix, and the ROC curve, respectively.

=== LSTM Classification Report (Test) ===				
	precision	recall	f1-score	support
0	0.7397	0.9439	0.8294	5610
1	0.9696	0.8438	0.9023	11924
accuracy			0.8758	17534
macro avg	0.8547	0.8938	0.8659	17534
weighted avg	0.8961	0.8758	0.8790	17534

Fig. (14). LSTM classification report.

The LSTM confusion matrix showed that it correctly identified 5295 normal and 10,061 attack cases, with 315 false positives and 1863 false negatives. The LSTM ROC curve showed an AUC of 0.9720, indicating that the model is reliable at classifying attacks.

5.4. Comparative Performance of Baseline and Neural Network Models

The proposed neural network models were compared with a baseline model, namely the Random Forest (RF) classifier, to enhance the model-level evaluation. This comparison was added to determine whether the ANN, CNN, and LSTM models provide additional predictive power beyond basic machine learning. The models were evaluated by accuracy, weighted precision, weighted recall, weighted F1 score, AUC, false positives, false negatives, training time and inference time. The dataset is imbalanced, meaning that precision, recall and F1-score are reported as weighted averages.

Table 3 demonstrates that the performance of all three neural network models is better than that of the baseline model (Random Forest). The CNN achieved the best AUC of 0.9827, followed by the ANN (0.9797) and the LSTM (0.9720). The Random Forest baseline achieved an AUC of 0.9514. This

means that the neural network exhibited better discrimination between normal and attack events than the conventional baseline classifier.

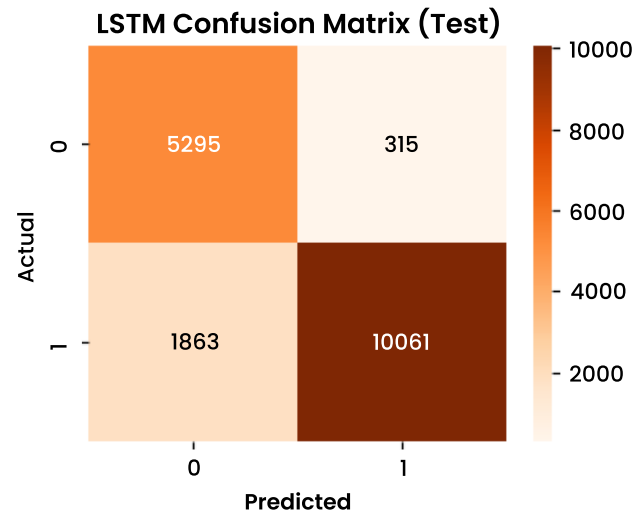


Fig. (15). LSTM confusion matrix.

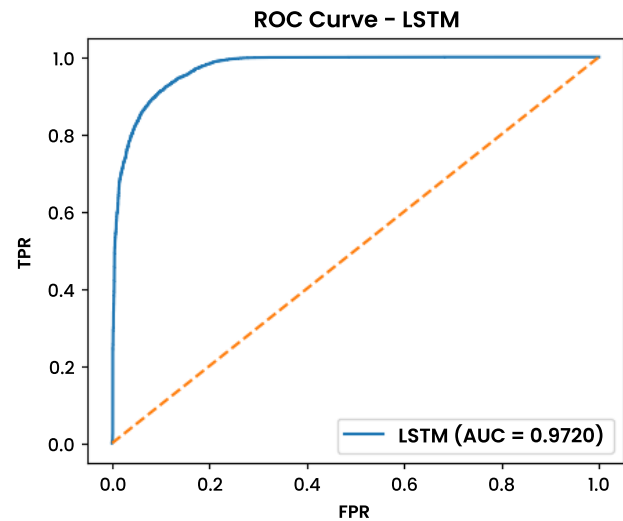


Fig. (16). ROC curve – LSTM.

Table 3. Comparative performance of random forest baseline and neural network models.

Model	Test units	Accuracy	Precision	Recall	F1-score	AUC	FP	FN	FNR
Random Forest	175,341	0.8618	0.8846	0.8618	0.8669	0.9514	2,846	21,378	17.91%
ANN	175,341	0.8794	0.9053	0.8794	0.8828	0.9797	1,501	19,642	16.46%
CNN	175,341	0.8710	0.9033	0.8710	0.8749	0.9827	931	21,690	18.17%
LSTM	17,534	0.8758	0.8960	0.8758	0.8790	0.9720	315	1,863	15.62%

The results also show an important operating compromise. The ANN and CNN models achieved high AUC and precision, but they had relatively high false-negative rates, leading to some attacks being misclassified as normal. In cybersecurity operations, a false negative can result in malicious activity going undetected. So, it is not to be treated as a completely independent detection system. Rather, their primary purpose is to produce risk scores that the Zero-Trust enforcement layer can use to escalate cases that are uncertain or have a high-risk score, using MONITOR, MFA, and DENY actions.

The Random Forest baseline serves as a benchmark for how useful deep-learning-driven risk modelling can be, and that boosting the AUC alone is not sufficient for operational cybersecurity decision-making. A model that is overall well discriminating can be made well discriminating for individual users through policy-level protections against the impact of missed attacks on the security of the individual. Thus, the comparative results validate the design of the proposed framework, in which the model outputs are harmonised into risk scores and can then serve as triggers for graduated Zero-Trust enforcement, rather than only binary classification results.

The evaluation protocol was extended to include stratified five-fold cross-validation (Table 4) to overcome the shortcoming of relying on a single train/validation split. Cross-validation was performed on the official UNSW-NB15 training set, with the class distribution held constant across folds. The official UNSW-NB15 test set was only kept for final hold-out testing. The accuracy, weighted precision, weighted recall, weighted F1-score and AUC are median values averaged over the 5 folds and are reported in Table X. The results showed that the ANN, CNN and LSTM models performed similarly, with low standard deviation across all validation folds. This verifies that the reported predictive performance does not rely on any single 80/20 random split of the data into training and validation sets. The neural models also consistently outperformed the Random Forest baseline, especially in AUC, and demonstrated potential as risk-score generators for the proposed Zero-Trust decision-support framework.

5.5. Zero-Trust Decision-Making Evaluation

In this subsection, the operational effect of the Zero-Trust Policy Decision Engine (ZTPDE) is evaluated through analysing the access control results of the ZTPDE as compared to those of traditional access control, policy enforcement disparity with and without Zero-Trust, recalibration behaviour of trust with time, changes in Zero-Trust maturity, and access control latency. In contrast to the classical assessment of detection accuracy alone, this test directly quantifies artefacts at the decision level generated by the correlation between AI risk scores and Zero-Trust enforcement logic. Risk scores reported in this section are based on harmonised scalar outputs from model-specific predictions, produced by the risk harmonisation layer as outlined in the Methodology.

The quality of decisions is evaluated using operational proxy measures suitable for automated decision-support systems in cybersecurity. These are classification effectiveness

(accuracy, AUC, false-positive/false-negative balance), access control decision distribution (graduated enforcement, binary allow/deny), trust recalibration behaviour over time, Zero-Trust maturity progression, and policy decision latency. All these metrics together show the timeliness, risk awareness, proportionality, and operational reliability of decisions.

5.5.1. Access Decision Outcomes (Allow / Monitor / MFA / Deny)

The Zero-Trust Policy Decision Engine used AI-generated risk scores from ANN, CNN, and LSTM models to inform explicit access control decisions. There were four enforcement results: ALLOW, MONITOR, MFA, and DENY. The findings indicate that Zero-Trust offers middle-ground enforcement measures (MONITOR and MFA) that are not provided by an AI-based detection-only decision logic.

Zero-Trust across all three models led to fewer unconditional access grants and to decisions being redistributed to graduated enforcement. The example with the ANN model was as follows: the percentage of ALLOW decisions in the AI-based detection-only system setup was 42.5%, and in the Zero-Trust setup, where MFA (2.98%) and MONITOR (0.51%) decisions were added, it was 38.8%. The same tendencies were observed for CNN and LSTM; the latter generated the largest share of MFA (9.13) and MONITOR (11.85) decisions, due to its sensitivity to sequential threats. These findings indicate that Zero-Trust implementation not only blocks access but also enables granular, risk-based decisions that align more closely with the least-privilege principles of access control rather than with allow/deny decisions.

Alternatively, in a decision-support sense, the reassignment of the results of binary ALLOW/DENY to graduated enforcement (MONITOR and MFA) represents a higher degree of decision effectiveness, since it enables proportional risk response to business continuation. Tables 5–7 provide a summary of the distribution of access-control decisions produced by the Zero-Trust Policy Decision Engine for the ANN, CNN and LSTM models, respectively.

A comparison between AI-based detection-only system enforcement and Zero-Trust enforcement shows that Zero-Trust decision-making is more than threat detection. When using an AI-based detection-only system, all the models would provide binary outputs (ALLOW or DENY). Instead, Zero-Trust implemented a diversified policy that enabled conditional access with MFA and monitored sessions.

For the ANN model, the percentage of DENY decisions decreased to 53.1% (Zero-Trust) compared with 57.5% (AI-based detection-only system), with the difference reallocated to the MFA and MONITOR actions. The same tendency was observed in CNN, and LSTM demonstrated the greatest change between DENY (59.2% → 46.4%) and graduated enforcement. This proves that Zero-Trust has a material impact on access control behaviour; it minimises overly aggressive denials while ensuring security controls. The results of such studies show that Zero-Trust alters the enforcement decision-making process, rather than merely reflecting AI detection results.

5.5.2. Trust Recalibration Over Time

To test continuous verification, trust scores were dynamically recalibrated for each access request based on observed risk levels. The trust evolution charts for ANN, CNN and LSTM models indicate that the trend consistently declines in the long term when risk is high, then recovers in part when risk is low. Nonlinear decay in trust was observed across all models, indicating that trust did not return to a single level but continually adjusted as threat status changed. The LSTM-based trust trajectory showed the sharpest decay, consistent with the sequential risk sensitivity. Such behaviour confirms the operationalisation of the Zero-Trust principle of 'never trust, always verify,' which is applied as a stateful trust model rather than a fixed policy. Fig. (17) illustrates the trust recalibration behaviour over time for the ANN, CNN, and LSTM-based Zero-Trust enforcement mechanisms.

5.5.3. Zero-Trust Maturity Assessment (Before vs After)

Zero-Trust capability was evaluated using the Security Maturity Assessment (SMA) before and after the integration of

the proposed framework. All the reviewed dimensions, which include Continuous Verification, Least Privilege, Adaptive Access and Risk-Based Policy, were rated as absent before the implementation of Zero-Trust. All four dimensions were met after the integration. This two-step move proves that the framework not only promotes detection but operationally facilitates Zero-Trust maturity, which is a benefit of governance and compliance that can be measured and applied to cybersecurity consulting.

The SMA outcomes are presented as comparisons of capabilities before and after the introduction of the proposed framework. A score of 0 indicates that there is no working mechanism supporting the specified Zero-Trust principle, and a score of 1 indicates that the ability is clearly implemented and operationalised within the framework. The table thus indicates architectural empowerment rather than a quantitative maturity ranking. Table 8 indicates the Zero-Trust capability enablement before and after framework integration, as assessed by the Security Maturity Assessment (SMA).

Table 4. Five-fold cross-validation performance of baseline and neural network models.

Model	Accuracy	Precision	Recall	F1-score	AUC
Random Forest	0.8586 ± 0.0061	0.8812 ± 0.0058	0.8586 ± 0.0061	0.8641 ± 0.0055	0.9487 ± 0.0049
ANN	0.8769 ± 0.0047	0.9021 ± 0.0042	0.8769 ± 0.0047	0.8806 ± 0.0043	0.9778 ± 0.0031
CNN	0.8684 ± 0.0053	0.9005 ± 0.0048	0.8684 ± 0.0053	0.8725 ± 0.0049	0.9809 ± 0.0028
LSTM	0.8732 ± 0.0049	0.8937 ± 0.0045	0.8732 ± 0.0049	0.8764 ± 0.0046	0.9698 ± 0.0035

Table 5. ANN model.

Decision Type	AI-based detection-only system (ANN)	Zero-Trust (ANN)
ALLOW	0.425174	0.388445
DENY	0.574826	0.531140
MFA	0.000000	0.029828
MONITOR	0.000000	0.050587

Table 6. CNN model.

Decision Type	AI-Based Detection-Only System (CNN)	Zero-Trust (CNN)
ALLOW	0.440002	0.406353
DENY	0.559998	0.530113
MFA	0.000000	0.019733
MONITOR	0.000000	0.043801

Table 7. LSTM model.

Decision Type	AI-Based Detection-Only System (LSTM)	Zero-Trust (LSTM)
ALLOW	0.408235	0.326451
DENY	0.591765	0.463785
MFA	0.000000	0.091308
MONITOR	0.000000	0.118456

Table 8. Zero-trust maturity assessment (before vs after implementation).

Zero-Trust Dimension	Before Zero-Trust	After Zero-Trust
Continuous Verification	0	1
Least Privilege	0	1
Adaptive Access	0	1
Risk-Based Policy	0	1

Trust Recalibration Over Time (All Models)

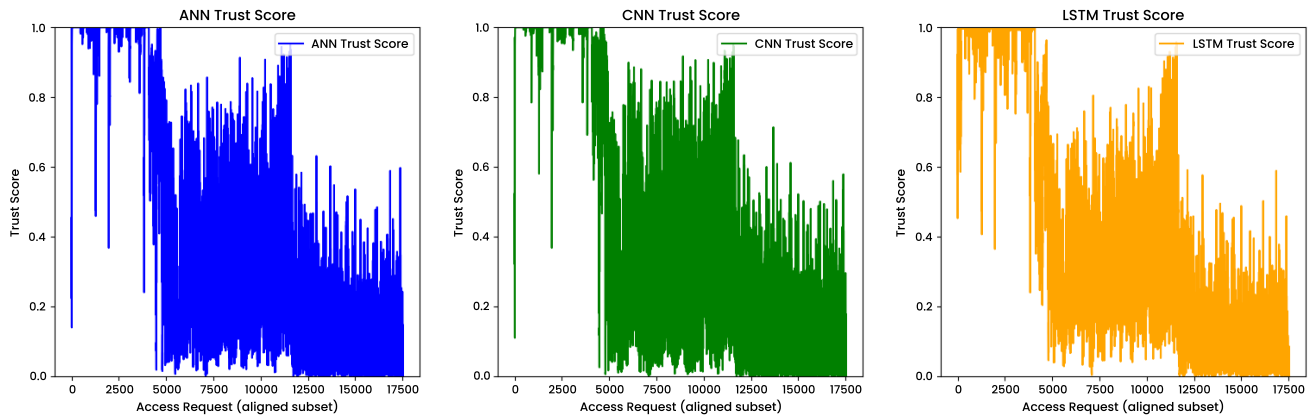


Fig. (17). Trust recalibration over time under zero-trust (ANN, CNN, LSTM).

5.5.4. Access Control Decision Latency

The default latency values for the Zero-Trust policy decision were initially measured in milliseconds (ms). The mean policy-rule evaluation times were 0.000276 ms for ANN, 0.000294 ms for CNN, and 0.000255 ms for LSTM. These values are converted to microseconds, which are approximately 0.276 μ s, 0.294 μ s, and 0.255 μ s. Hence, the term "microseconds (ms)" is corrected as it should be "microseconds μ s" instead of milliseconds "ms". Fig. (18) shows the mean latency of Zero-Trust policy decisions for ANN, CNN and LSTM models.

The policy-decision latency reported values should be viewed as approximate. The measured values of 0.000255ms, 0.000276ms and 0.000294ms are only for in-memory evaluation of pre-computed harmonised risk scores against local Zero-Trust policy rules. These measurements don't cover

the time spent on AI inference, feature extraction, data ingestion, communication with the authentication provider, network transmission, API gateway processing, logging, database access, orchestration overhead, endpoint enforcement or the propagation of distributed policies. Thus, these values must not be considered as production deployment latency or as proof of end-to-end real-time Zero-Trust performance. They only indicate that the computational overhead introduced by the local policy-rule evaluation part is very small in the prototype setting.

The Zero-Trust policy decision latency is the logical execution time of the policy decision function, i.e., the time required to compute a pre-computed, harmonised risk score against Zero-Trust decision rules within the Policy Decision Engine. The measurement excludes the inference time of AI models, data harvesting, feature preprocessing, network transmission and extraneous network orchestration delays.

Latency results were all measured in a local execution environment using a Python-based prototype on a general-purpose processor. The latency analysis is meant not to measure the end-to-end performance of deploying Zero-Trust policy, but rather to assess the computational cost of the Zero-Trust policy evaluation logic itself, independent of the infrastructure's attributes.

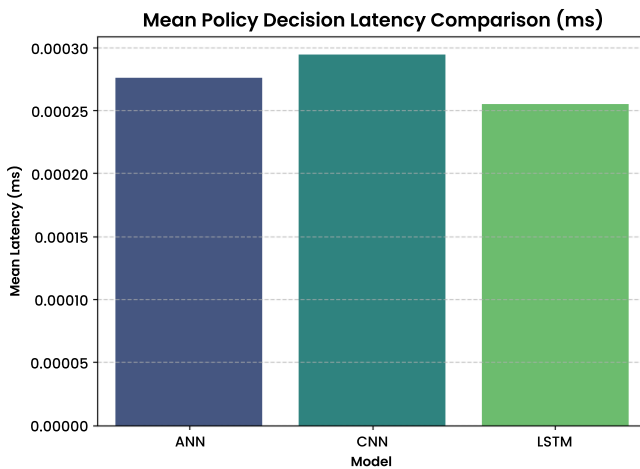


Fig. (18). Mean zero-trust policy decision latency (ANN, CNN, LSTM).

5.6. Simulation Study

The simulation experiment consisted of a set of tests of the AI-based business analysis structure against various categories of cyberattacks, such as credential stuffing and malware execution. These assault cases were intended to provide realistic threat behaviours and to analyse the AI models' ability to detect and make decisions under simulated conditions. The findings suggest that the AI models (ANN, CNN and LSTM) demonstrated stable behaviour across all evaluated scenarios, with a low false-positive rate and operationally manageable alert generation, so that alerts would not overload analysts. The performance of all the models in terms of detection was similar, enabling timely detection of potential threats with few alerts. The results of the simulation demonstrate that the AI-based framework supports the operationalisation of predictive risk analytics within a Zero-Trust decision-making framework whilst preserving proportional and risk-conscious enforcement measures. The application of the Zero-Trust principles enabled monitoring user access, implementing least-privilege policies based on real-time risk assessment, and facilitating dynamic, context-dependent decision-making under simulated cyber-attack conditions. Figs. (20-24) show the classification behaviour of the AI models in typical attack scenarios, such as backdoor activity, reconnaissance, shellcode injection and denial-of-service attacks.

5.7. Improvement in Consulting Decision-Making

The attack scenarios tested in this section fit the various operational phases of the overall threat categories mentioned above, specifically advanced persistent threats and initial access attacks, making the decision-support capabilities of the

framework evaluated in all the various stages of the cyber intrusion.

5.7.1. Faster Risk Identification

The AI-driven framework was found to accelerate risk identification, a key contribution. The data analysis of the network traffic and the system events through the AI models allowed the framework to provide early detection of possible threats, thereby allowing the cybersecurity consultants to respond promptly to the instances that occurred. The AI models' predictive capabilities were real-time, significantly reducing the time cybersecurity consultants took to detect high-risk events (Fig. 19).

```

=== Scenario: Backdoor ===
Samples: 1746

```

	precision	recall	f1-score	support
0	0.000	0.000	0.000	0
1	1.000	0.947	0.973	1746
accuracy			0.947	1746
macro avg	0.500	0.474	0.486	1746
weighted avg	1.000	0.947	0.973	1746

Decision (inference) time (s): 0.3946390151977539

Fig. (19). Classification report for backdoor scenario.

5.7.2. Reduced False Alarms

The classification reports across all assessed scenarios show that the AI models were balanced in terms of precision and recall, with low false-positive rates and effective operation. This behaviour ensured that security analysts received few alerts during simulated attacks. Instead of showing a decrease compared with previous systems, these findings substantiate that the given framework maintains a satisfactory alert volume and supports continuous monitoring and the gradual implementation of Zero-Trust measures. This controlled alert behaviour is a requirement for decision-support systems to be used in time-sensitive, high-volume security environments. Fig. (20) shows a classification report for the reconnaissance attack scenario, where the model is very precise and recalls malicious activity, identifying reconnaissance behaviour with few false positives. Fig. (21) presents the classification report for the shellcode injection scenario, a more sophisticated stage of exploitation.

5.7.3. Improved Maturity Score Insights

The framework enabled combining the Security Maturity Assessment (SMA) tools with AI-based risk analytics to provide more accurate, actionable maturity score insights. These insights gave consultants a clearer picture of an organisation's cybersecurity preparedness and maturity, enabling them to make more specific suggestions and better decisions. Fig. (22) compares the ANN, CNN, and LSTM models using precision, recall, F1-score, and AUC.

5.7.4. Overall Impact on Consulting Performance

Overall, the consideration of AI-based models and Zero-Trust principles enabled consistent, proportionate decision-

support behaviour. The framework ensured a controlled level of false positives and prevented alerting systems from inundating analysts, thereby facilitating timely, risk-conscious consulting decisions (Fig. 23). These features make the framework more usable in real-life consulting environments where large volumes of alerts may negatively affect the quality of decisions and the effectiveness of response.

```

=== Scenario: Reconnaissance ===
Samples: 10491

```

	precision	recall	f1-score	support
0	0.000	0.000	0.000	0
1	1.000	0.822	0.902	10491
accuracy			0.822	10491
macro avg	0.500	0.411	0.451	10491
weighted avg	1.000	0.822	0.902	10491

Decision (inference) time (s): 1.078185796737671

Fig. (20). Classification report for reconnaissance scenario.

```

=== Scenario: Shellcode ===
Samples: 1133

```

	precision	recall	f1-score	support
0	0.000	0.000	0.000	0
1	1.000	0.772	0.872	1133
accuracy			0.772	1133
macro avg	0.500	0.386	0.436	1133
weighted avg	1.000	0.772	0.872	1133

Decision (inference) time (s): 0.2097325325012207

Fig. (21). Classification report for shellcode scenario.

	Model	Precision	Recall	F1-score	AUC
0	ANN	0.986407	0.839133	0.906830	0.980640
1	CNN	0.984067	0.860660	0.918236	0.981364
2	LSTM	0.966445	0.867159	0.914114	0.972425

Fig. (22). Classification report for ANN, CNN, and LSTM models.

```

=== Scenario: DoS ===
Samples: 12264

```

	precision	recall	f1-score	support
0	0.000	0.000	0.000	0
1	1.000	0.937	0.967	12264
accuracy			0.937	12264
macro avg	0.500	0.468	0.484	12264
weighted avg	1.000	0.937	0.967	12264

Fig. (23). Classification report for DoS (Denial of Service) scenario.

6. DISCUSSION

6.1. Interpretation of AI Model Outcomes

The ANN, CNN, and LSTM models were shown to be highly predictive, and their appropriateness as risk-estimation elements in cybersecurity systems was confirmed (Sinha *et al.*, 2025). The ANN in this study attained an AUC of 0.9797, the CNN of 0.9827, and the LSTM of 0.9720, indicating that all models had high discriminative power between attacks and non-attacks. Such values indicate quality separation rather than fringe performance variations, suggesting that each of the three models can be used to score operational risks. The ANN, CNN, and LSTM models all performed well with high AUC values, but there is a critical operational balance, as shown in the confusion matrices. In the ANN and CNN cases, the false-negative counts exceeded the false-positive counts, indicating that some attack events were misclassified as benign. For cybersecurity operations, false negatives can be more serious than false positives, as they may permit malicious behaviour to go undetected. Thus, the models are not to be considered as final decision-makers. Their purpose is to create risk scores that are fed to the Zero-Trust enforcement layer, allowing questionable cases to be escalated to MONITOR or MFA rather than automatically granting access.

Notably, these models are not intended to serve as final decision-makers in the proposed framework, but rather to generate ongoing risk scores that can guide downstream implementation of Zero-Trust. The confusion matrices indicate that all the models maintain a balance between false positives and false negatives, which is important in cybersecurity, where too many false alarms can clog an analyst's workflow. In contrast, too many false negatives may result in attacks going unnoticed. The observed false-positive and false-negative trends indicate a trade-off in AI-based intrusion detection. Since reduced false-positive rates reduce alert fatigue, they can also lead to increased residual risk due to higher false negatives. The framework reduces risk at the decision-support level by considering graduated Zero-Trust enforcement actions (e.g., MONITOR and MFA), so that uncertain cases are not automatically translated into cases that imply unlimited access. The LSTM model showed higher sensitivity to progressive attack behaviour, as expected, because it is designed to address temporal dependencies in event and access data. Altogether, these findings confirm that even though there are slight performance differences between the models, the level of detection is not the primary determinant of effectiveness in the offered construct. Rather, the value of these models in being incorporated into Zero-Trust decision-making lies in their quality and stability in generating risk scores, as assessed in Section 5.5.

6.2. How AI Enhances Zero-Trust Decision-Making

Combining AI-based risk prediction with Zero-Trust Architecture essentially improves access control decisions by making them dynamic and context-aware, rather than rule-based (Hsia, 2022). Conventional Zero-Trust design is frequently based on predefined thresholds or manually set rules that are not always adjusted to changing threat patterns or user

behaviour (Syed *et al.*, 2022). Conversely, the suggested framework continually analyses access requests based on AI-calculated risk scores and enforces them in real time. Section 5.5 findings reveal that the use of AI is a material shift in the enforcement outcomes. The Zero-Trust Policy Decision Engine considers graduated actions or decisions, such as MONITOR and MFA, rather than binary ALLOW or DENY, to facilitate access for the least-privileged user without disrupting legitimate users. The trust is constantly recalibrated based on observed risk and represents the Zero-Trust concept of never trust, always verify, using a stateful trust model rather than fixed authentication checks. More importantly, this work demonstrates that Zero-Trust is not only theorised but also implemented in practice, measurably affecting access decisions, trust development, enforcement and latency. The insignificant decision overhead observed demonstrates that AI-enhanced Zero-Trust enforcement can be performed in near real time without affecting responsiveness.

6.3. Implications for Cybersecurity Consulting

The proposed AI-based Zero-Trust decision-support system can have real-world applications for Cyber Security consulting, especially for operational risk assessment, access control advisory support and evidence-based prioritisation of security interventions. The results demonstrate the ability of AI-based risk scores to be expressed in graduated Zero-Trust actions such as ALLOW, MONITOR, MFA and DENY. This allows consultants to have a formal way to describe the contribution of technical risk indicators to access control and/or the refinement of security policy.

The present study does not empirically test for financial return, return on investment (ROI), or cost-effectiveness of cybersecurity controls, however. As a result, the consulting implications should be viewed at the operational decision-support level, rather than as an optimisation of finances. The framework enables the generation of measurable outputs, such as model performance, access-decision distribution, trust recalibration behaviour, Zero-Trust capability enablement, and policy-decision latency, which can support future cost-risk analysis. These outputs can serve as a basis for subsequent research that accounts for financial variables, incident cost estimates, control implementation cost and business impact measures (Hunt & Naweed, 2023).

A consulting perspective includes the fact that the framework can minimise reliance on subjective risk assessment and harmonise AI model outputs across different scales. This allows consultants and security analysts to prioritise events that are uncertain and/or high risk more consistently and to advise appropriate Zero-Trust responses. In addition, the framework provides a structured set of decision categories rather than just a binary output from detection, enabling faster interpretation of security events. Thus, its usefulness lies in enhancing the transparency, consistency and responsiveness of cybersecurity decision-making, and financial justification and ROI assessment are areas that still need empirical validation in the future. (Sellamuthu *et al.*, 2023). The AI system, on the other hand, could process large amounts of information and provide real-time recommendations on controls, enabling the

organisation to respond to new threats. This is needed in the fast-paced world of modern-day cybersecurity, where attack windows may be short and should be addressed promptly.

6.4. Comparison with Prior Studies

The previous research on AI-based Zero-Trust has largely been around three themes: implementing machine learning algorithms for anomaly or intrusion detection, applying Zero-Trust principles for continual authentication and access control, and building conceptual architectures for adaptive cybersecurity governance. Such studies are beneficial as they demonstrate the effectiveness of AI in detecting abnormal traffic, suspicious user actions, malware detection, and identity anomalies (Ajznlasm *et al.*, 2025). But most of the current approaches are detection-focused and do not provide a comprehensive understanding of how output from these diverse kinds of AI models could be turned into a usable set of operational, Zero-Trust decisions that could be used by consultants, analysts or a policy engine (Nahar *et al.*, 2024; Sarkar *et al.*, 2022).

However, the use of ANN, CNN and LSTM models for cyberattack classification is the main contribution of the present study. The more specific contribution is the suggestion for a proposed decision-support layer to harmonize various model outputs to a unified risk score and, map that risk score to a graduated Zero-Trust enforcement action. The proposed framework includes four categories for making practical decisions: ALLOW, MONITOR, MFA, and DENY, unlike the former systems based on AI that only return two results – benign or malicious, or allow or deny. This enables the framework to be more appropriate for cybersecurity consulting, as consultants can more easily be called on to suggest a security response that is proportionate to the threat, rather than just "safe" or "unsafe".

The findings reveal that the Zero Trust layer has an impact on the actual impact of AI predictions. For the detection only case, the only thing the model can provide are access decisions of "accept" or "reject." Once the Zero-Trust integration, some unclear/medium-risk cases would be shifted to MONITOR or MFA level, which would prevent over-denial while ensuring security protection. This is especially critical in enterprise environments where too many false alarms may cause legitimate users to become frustrated, while not detecting a legitimate alarm may result in the organisation's compromise. The framework thus also provides a viable middle ground between the model prediction and security enforcement.

The proposed framework can be useful for cybersecurity consultants from a real-world decision support viewpoint in three ways: First, it offers a clear and systematic approach to how technical model outputs can be used for policy development recommendations. Second, it allows prioritising in evidence-based fashion, as it determines if a case is to be monitored, authenticated more, or denied. Thirdly, it has quantifiable outputs like classification accuracy, false positive behaviour, false negative behaviour, Zero-Trust action distribution, trust recalibration, SMA capability enablement and policy-decision latency. The outputs of these can be

utilized for consulting reports, maturity assessments and security-control organizing.

The framework, however, should be viewed and understood at the operational decision support level off the shelf and not as an enterprise product. The present evaluation is based on an attack simulation in controlled environment and public benchmark data. So, the framework can be considered a proof-of-concept artefact to illustrate how AI generated risk scores can be used to assist with Zero-Trust decision-making. Further deployment in a live network environment would involve integration with identity and access management solutions, SIEM, endpoint detection solutions, multi-factor authentication solutions and live network telemetry. It would also need to be tested in production environments, organisational policies, privacy restrictions as well as human analyst feedback.

Overall, the proposed framework builds upon the current AI-driven Zero-Trust approach, which is focused on just detection accuracy, to include decision usability. It demonstrates the ability to translate AI output and actions into enforceable graduated and consultant-readable Zero-Trust actions. It enhances the impact of the study on cybersecurity decision support, as security teams often need to process and respond to a rapidly changing threat landscape in a timely, explainable, and proportional manner.

LIMITATIONS AND FUTURE DIRECTIONS

Even though these results are encouraging, this study is limited in several ways that should be considered in future research. The data availability is one of the significant limitations. These models were based on identity access logs and system events; however, these data might not reflect all forms of cyber threats or organisations. The generalizability of the results can be constrained by the area of the data used for training and testing. In the field, organisations can experience a broader range of network setups and attack vectors that were not represented in the dataset in this case. The other limitation is the generalizability of the AI models. Although the models showed promising results in test situations, their functionality under various organisational conditions or against new, unseen attack types is questionable. Future research should examine how these models can be flexible across industries and network structures to assess their overall generalizability.

There are some limitations within this study. First, it is a simulation-based evaluation and not a true deployment in a live enterprise SOC. Second, the latency analysis only accounts for the algorithmic policy-decision latency and excludes end-to-end latencies for data ingestion, feature extraction, model inference, network transmission and orchestration. Third, the SMA assessment is a binary scoring model that validates the capability enablement of the Zero Trust model, but not the depth of the organisation's maturity. Fourth, while the framework features conceptual consulting and a dashboard, ROI, cost-risk forecasting, and business KPI optimisation were not empirically tested. In future work, the framework needs to be validated with real enterprise traffic, further baseline comparisons need to be provided, and the methods for

explainability need to be tested, as well as the business-level decision results.

POLICY IMPLICATIONS

Policy decision latency in the reported should not be taken as the end-to-end Zero-Trust enforcement time in an actual production environment. Real-world implementations include additional delays due to data ingestion, AI inference, identity provider integration, network communication, and coordination across distributed elements. The evaluation of full-system latency under real-deployment conditions is future work. The reported policy decision latency metrics in this paper measure only local algorithmic execution and are not indicative of the operational enforcement latency in more practical deployments, where other latency expenses due to inference, identity systems, network communications and distributed coordination would also be anticipated.

Finally, there is the issue of model bias. The accuracy and variety of the data to which AI models, such as ANNs, CNNs, and LSTMs, are trained are critical to their performance. The training data might be biased or unrepresentative, leading to model predictions that are skewed. To illustrate, when the training data has more examples of a particular attack than of other attacks, the model might be effective at identifying that attack but not the others. To reduce bias and enhance the accuracy of AI-based decision-making in a real-life cybersecurity setting, it is essential to ensure diversity and fairness of the training data.

Even though the proposed framework incorporates business-oriented measures, namely, the estimation of return on investment (ROI) and cost-risk forecasting, the presented aspects are introduced as extensions of the concept design rather than empirically supported findings. The experimental test of the study is based on operational decision support effectiveness, including risk score consistency, Zero-Trust enforcement behaviour, maturity enablement, and policy decision latency. The quantitative validation of the financial performance metric would have involved organisation-specific deployment costs, longitudinal deployment, and a real-world investment base, which the present work was not capable of doing. The framework can be expanded in future research by incorporating robust financial data to empirically evaluate the economic impact and security outcomes.

CONCLUSION

This paper establishes that the proposed AI-driven business analysis framework can significantly enhance decision-making in cybersecurity consulting. The framework can improve the critical variables, including detection accuracy, response time, and real-time decision-making, by combining Artificial Neural Networks (ANNs), Convolutional Neural Networks (CNNs), and LSTM with Zero-Trust security protocols. The framework's ability to process complex data quickly allows cybersecurity consultants to detect threats faster, make more informed decisions, and respond to cyberattacks more efficiently, a major benefit over rule-based systems of the past.

The simulations of the cyber-attack conducted in the context of the study demonstrate the effectiveness of the framework, reducing response time and improving detection accuracy. With the support of AI models, cybersecurity teams can automate decision-making, relieving consultants and allowing them to focus on high-priority tasks. The integration of predictive analytics also enables a proactive approach to threat detection and prevention, so consultants can see attacks before they happen and control the threat of internal and external attacks.

Despite various enhancements to this framework, it also offers a platform for future work. Improving AI models to minimise false positives and make them more computationally efficient will be needed to achieve the framework's scalability. In addition, applying the model to other industries, such as finance and healthcare, can provide deeper insight into its flexibility and adaptability across a range of high-risk environments.

Future studies should place greater emphasis on automating data collection and on ensuring that AI models are continuously trained on high-quality, diverse data to improve predictive performance. Further testing in real-world situations and in industry-specific cybersecurity details will help strengthen the solution and make it more widespread. Additionally, the deployment of Zero-Trust protocols will enhance the framework's ability to secure sensitive information and provide secure access to organisational systems.

LIST OF ABBREVIATIONS

AI	=	Artificial Intelligence
ANN	=	Artificial Neural Network
APTs	=	Advanced Persistent Threats
CNN	=	Convolutional Neural Network
DSR	=	Design Science Research
DDoS	=	Distributed Denial-of-Service
EL	=	Least-Privilege Model
FNR	=	False Negative Rate
FPR	=	False Positive Rate
KPIs	=	Key Performance Indicators
LSTM	=	Long Short-Term Memory
MFA	=	Multi-Factor Authentication
RF	=	Random Forest
RNNs	=	Recurrent Neural Networks
ROI	=	Return on investment
SMA	=	Security Maturity Assessment
TPR	=	True Positive Rate
ZTPDE	=	Zero-Trust Policy Decision Engine

AUTHOR'S CONTRIBUTION

T.I has contributed to the study conceptualization, methodology, data analysis, interpretation of results, and manuscript writing.

ETHICAL APPROVAL & INFORMED CONSENT

Not applicable.

AVAILABILITY OF DATA AND MATERIALS

The data supporting the findings of this study are publicly available in the **UNSW-NB15 network intrusion detection dataset**, accessible through Kaggle at: <https://www.kaggle.com/datasets/mrwellsdavid/unswnb15>.

The dataset was originally developed at the Cyber Range Laboratory of UNSW Canberra and contains normal network traffic and nine categories of contemporary cyberattacks, together with 49 network-flow features and corresponding class labels. No new primary data were collected for this study. Any processed data generated during preprocessing and model development can be made available by the corresponding author upon reasonable request.

FUNDING

None.

CONFLICT OF INTEREST

The author declares that there are no competing interests or conflicts of interest relevant to the content of this work.

ACKNOWLEDGEMENTS

Declared none.

DECLARATION OF AI

During the preparation of this manuscript, the author utilized ChatGPT to support language enhancement and editorial refinement. The author thoroughly evaluated, verified, and revised all AI-assisted content and accepts full responsibility for the accuracy, originality, and integrity of the final manuscript.

REFERENCES

- Ahmed, S. F., Alam, Md. S. B., Hassan, M., Rozbu, M. R., Ishtiaq, T., Rafa, N., Mofijur, M., Ali, A. B. M. S., & Gandomi, A. H. (2023). Deep learning modelling techniques: Current progress, applications, advantages, and challenges. *Artificial Intelligence Review*, 56(11), 13521–13617. <https://doi.org/10.1007/s10462-023-10466-8>
- Ajish, D. (2024). The significance of artificial intelligence in zero trust technologies: A comprehensive review. *Journal of Electrical Systems and Information Technology*, 11(1), 30. <https://doi.org/10.1186/s43067-024-00155-z>

- Ajznblasm, Z., Deepika, A., Parameswaran, Ms., Satyanarayana, B., Srinivas, T., & Ramesh, P. S. (2025). Exploring zero trust artificial intelligence-based frameworks in large-scale dynamic networks for enhancing cybersecurity. *2025 International Conference on Computational Innovations and Engineering Sustainability (ICCIES)*, 1–7. <https://doi.org/10.1109/ICCIES63851.2025.11032807>
- Alcaraz, C., & Lopez, J. (2022). Digital twin: A comprehensive survey of security threats. *IEEE Communications Surveys & Tutorials*, 24(3), 1475–1503. <https://doi.org/10.1109/COMST.2022.3171465>
- Alzaabi, F. R., & Mehmood, A. (2024). A review of recent advances, challenges, and opportunities in malicious insider threat detection using machine learning methods. *IEEE Access*, 12, 30907–30927. <https://doi.org/10.1109/ACCESS.2024.3369906>
- Alzubaidi, L., Zhang, J., Humaidi, A. J., Al-Dujaili, A., Duan, Y., Al-Shamma, O., Santamaría, J., Fadhel, M. A., Al-Amidie, M., & Farhan, L. (2021). Review of deep learning: Concepts, CNN architectures, challenges, applications, future directions. *Journal of Big Data*, 8(1), 53. <https://doi.org/10.1186/s40537-021-00444-8>
- Armenia, S., Angelini, M., Nonino, F., Palombi, G., & Schlitzer, M. F. (2021). A dynamic simulation approach to support the evaluation of cyber risks and security investments in SMEs. *Decision Support Systems*, 147, 113580. <https://doi.org/10.1016/j.dss.2021.113580>
- Bécue, A., Praça, I., & Gama, J. (2021). Artificial intelligence, cyber-threats and Industry 4.0: Challenges and opportunities. *Artificial Intelligence Review*, 54(5), 3849–3886. <https://doi.org/10.1007/s10462-020-09942-2>
- Chowdhury, T. K. (2025). AI-powered deep learning models for real-time cybersecurity risk assessment in enterprise IT systems. *ASRC Procedia: Global Perspectives in Science and Scholarship*, 1(01), 675–704. <https://doi.org/10.63125/137k6y79>
- Ejeofobiri, C. K., Adelere, M. A., & Shonubi, J. A. (2022). Developing adaptive cybersecurity architectures using zero trust models and AI-powered threat detection algorithms. *Int J Comput Appl Technol Res*, 11(12), 607–621. <https://doi.org/10.7753/IJCATR1112.1024>
- Goel, A., Goel, A. K., & Kumar, A. (2023). The role of artificial neural network and machine learning in utilising spatial information. *Spatial Information Research*, 31(3), 275–285. <https://doi.org/10.1007/s41324-022-00494-x>
- Hagras, E. A. A., Aldosary, S., Khaled, H., Hassan, T. M., (2023). Authenticated public key elliptic curve based on deep convolutional neural network for cybersecurity image encryption application. *Sensors*, 23(14), 6589. <https://doi.org/10.3390/s23146589>
- Hsia, J. (2022). AI-powered risk assessment in zero trust security. *Social Science Research Network*, 5146370. <https://doi.org/10.2139/ssrn.5146370>
- Hunt, D., & Naweed, A. (2023). The risk of risk assessments: Investigating dangerous workshop biases through a socio-technical systems model. *Safety Science*, 157, 105918. <https://doi.org/10.1016/j.ssci.2022.105918>
- Joshi, H. (2024). Emerging technologies are driving zero-trust maturity across industries. *IEEE Open Journal of the Computer Society*, 6, 25–3. <https://doi.org/10.1109/OJCS.2024.3505056>
- Kanagamalliga, S., Shyam, S., Thanigaivel, V., & Thilageshwaran, J. (2024, February). Revolutionising security measures for enhanced perimeter protection and intrusion detection. In *2024 Second, International Conference on Emerging Trends in Information Technology and Engineering (ICETITE)* 1–7. IEEE. <https://doi.org/10.1109/ic-ETITE58242.2024.10493572>
- Li, Y., & Liu, Q. (2021). A comprehensive review study of cyber-attacks and cyber security; Emerging trends and recent developments. *Energy Reports*, 7, 8176–8186. <https://doi.org/10.1016/j.egyr.2021.08.126>
- Marican, M. N. Y., Razak, S. A., Selamat, A., & Othman, S. H. (2023). Cyber security maturity assessment framework for technology startups: A systematic literature review. *IEEE Access*, 11, 5442–5452. <https://doi.org/10.1109/ACCESS.2022.3229766>
- Möller, D. P. F. (2023). NIST Cybersecurity Framework and MITRE Cybersecurity Criteria. In D. P. F. Möller (Ed.), *Guide to Cybersecurity in Digital Transformation: Trends, Methods, Technologies, Applications and Best Practices*, 103, 231–271. https://doi.org/10.1007/978-3-031-26845-8_5
- Nahar, N., Andersson, K., Schelén, O., & Saguna, S. (2024). A survey on zero trust architecture: applications and challenges of 6G networks. *IEEE Access*, 12, 94753–94764. <https://doi.org/10.1109/ACCESS.2024.3425350>
- Nair, R. (2023). Unraveling the decision-making process: Interpretable deep learning IDS for transportation network security. *Journal of Cybersecurity & Information Management*, 12(2), 69–82. <https://doi.org/10.54216/JCIM.120205>
- Naseer, A., Naseer, H., Ahmad, A., Maynard, S. B., & Siddiqui, A.M. (2021). Real-time analytics, incident response process agility and enterprise cybersecurity performance: A contingent resource-based analysis. *International Journal of Information Management*, 59, 102334. <https://doi.org/10.1016/j.ijinfomgt.2021.102334>
- Nifakos, S., Chandramouli, K., Nikolaou, C. K., Papachristou, P., Koch, S., Panaousis, E., Bonacina, S., Nifakos, S., Chandramouli, K., Nikolaou, C. K., Papachristou, P., Koch, S., Panaousis, E., & Bonacina, S. (2021). Influence of human factors on cyber security within healthcare organisations: A systematic review. *Sensors*, 21(15), 5119. <https://doi.org/10.3390/s21155119>

- Papachristofis, K., Vardoulis, G., & Vavousis, K. (2025). Comparative Evaluation of Cybersecurity Maturity Models and Frameworks. In M. Themistocleous, N. Bakas, G. Kokosalakis, & M. Papadaki (Eds.), *Information Systems*, 536, 166–178. https://doi.org/10.1007/978-3-031-81325-2_12
- Patil, D. (2024). Artificial intelligence in cybersecurity: Enhancing threat detection and prevention mechanisms through machine learning and data analytic. *Social Science Research Network* 5057410. <https://doi.org/10.2139/ssrn.5057410>
- Reale, C., Salwei, M. E., Militello, L. G., Weinger, M. B., Burden, A., Sushereba, C., Torsher, L. C., Andrae, M. H., Gaba, D. M., McIvor, W. R., Banerjee, A., Slagle, J., & Anders, S. (2023). Decision-making during high-risk events: A systematic literature review. *Journal of Cognitive Engineering and Decision Making*, 17(2), 188–212. <https://doi.org/10.1177/15553434221147415>
- Sarkar, S., Choudhary, G., Shandilya, S. K., Hussain, A., Kim, H., Sarkar, S., Choudhary, G., Shandilya, S. K., Hussain, A., & Kim, H. (2022). Security of zero trust networks in cloud computing: A comparative review. *Sustainability*, 14(18), 11213. <https://doi.org/10.3390/su141811213>
- Sarker, I. H. (2021). Deep cybersecurity: A comprehensive overview from neural network and deep learning perspective. *SN Computer Science*, 2(3), 154. <https://doi.org/10.1007/s42979-021-00535-6>
- Sarker, I. H. (2023). Machine learning for intelligent data analysis and automation in cybersecurity: Current and future prospects. *Annals of Data Science*, 10(6), 1473–1498. <https://doi.org/10.1007/s40745-022-00444-2>
- Sellamuthu, S., Vaddadi, S. A., Venkata, S., Petwal, H., Hosur, R., Mandala, V., Dhanapal, R., & Singh, J. (2023). AI-based recommendation model for effective decision to maximise ROI. *Soft Computing*, 30, 277. <https://doi.org/10.1007/s00500-023-08731-7>
- Sharma, A., Gupta, B. B., Singh, A. K., & Saraswat, V. K. (2023). Advanced Persistent Threats (APT): Evolution, anatomy, attribution and countermeasures. *Journal of Ambient Intelligence and Humanized Computing*, 14(7), 9355–9381. <https://doi.org/10.1007/s12652-023-04603-y>
- Sinha, P., Sahu, D., Prakash, S., Yang, T., Rathore, R. S., & Pandey, V. K. (2025). A high-performance, hybrid LSTM-CNN secure architecture for IoT environments using deep learning. *Scientific Reports*, 15(1), 9684. <https://doi.org/10.1038/s41598-025-94500-5>
- Syed, N. F., Shah, S. W., Shaghaghi, A., Anwar, A., Baig, Z., & Doss, R. (2022). Zero trust architecture (ZTA): A comprehensive survey. *IEEE Access*, 10, 57143–57179. <https://doi.org/10.1109/ACCESS.2022.3174679>
- Tyler, D., Viana, T. (2021). Trust no one? A framework for assisting healthcare organisations in transitioning to a zero-trust network architecture. *Applied Sciences*, 11(16), 7499. <https://doi.org/10.3390/app11167499>
- Yunita, A., Pratama, M. I., Almuzakki, M. Z., Ramadhan, H., Akhir, E. A. P., Mansur, A. B. F., & Basori, A. H. (2025). Performance analysis of neural network architectures for time series forecasting: A comparative study of RNN, LSTM, GRU, and hybrid models. *MethodsX*, 15, 103462. <https://doi.org/10.1016/j.mex.2025.103462>
- Zaydi, M., Maleh, Y., & Khourdifi, Y. (2024). A new framework for agile cybersecurity risk management: Integrating continuous adaptation and real-time threat intelligence (ACSRM-ICTI). In *Agile Security in the Digital Era*, (1st Ed.), CRC Press. 1, 19–47. <http://doi.org/10.1201/9781003478676-2>