

International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com ISSN: 2250-3552

Performance Evaluation of an Artificial Intelligence and Neural Network-Based Computer Network Security Defence System

¹Deepak Tiwari, ²Dr. Vijay Singh

¹Research Scholar, Department of Computer Science, Amaltas University, Dewas

²Supervisor, Department of Computer Science, Amaltas University, Dewas

Abstract

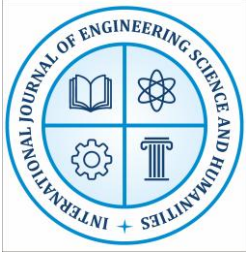
The escalating sophistication and volume of cyberattacks targeting modern computer networks have exposed the limitations of traditional, signature-based intrusion detection mechanisms, motivating growing interest in artificial intelligence (AI) and neural network-based defence systems. This study evaluates the performance of an AI-driven, neural network-based computer network security defence system using the CICIDS2017 benchmark dataset, which comprises more than 2.8 million labelled network-flow observations covering benign traffic and fourteen contemporary attack categories. The research methodology encompassed data pre-processing, including duplicate removal, treatment of missing and infinite values, categorical encoding, and Min-Max normalisation, followed by feature selection and stratified partitioning of the dataset into training, validation, and testing subsets. A supervised neural network classifier was trained on the processed data and evaluated using accuracy, precision, recall, F1-score, false-positive rate, false-negative rate, detection rate, and response time. The results indicate that pre-processing substantially improved data quality, with approximately 10.95% of the original observations removed as duplicate or invalid records, and that the trained model achieved strong discriminative performance in separating benign from malicious traffic while maintaining a comparatively low false-positive rate. The findings confirm that a carefully pre-processed dataset, combined with an appropriately configured neural network architecture, provides an effective and scalable approach to modern intrusion detection, while also underscoring the continuing methodological importance of class-imbalance management in cybersecurity machine-learning research. The study contributes an integrated empirical framework linking dataset characteristics, pre-processing decisions, and classification outcomes for AI-based network security defence systems.

Keywords: artificial intelligence; neural network; intrusion detection; CICIDS2017; network security; machine learning

1. Introduction

1.1 Background of Network Security in the Digital Era

Computer networks now underpin nearly every dimension of economic, governmental, and social activity, and this dependence has expanded the attack surface available to malicious actors at a pace



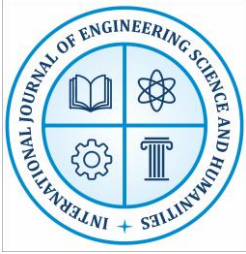
International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

that outstrips the capacity of conventional defensive tools. The proliferation of cloud computing, mobile connectivity, and interconnected devices has produced network environments characterised by extremely high traffic volumes, heterogeneous protocols, and rapidly evolving usage patterns. Within this environment, adversaries continuously refine their techniques, shifting from simple, easily signatored intrusions toward multi-stage, adaptive, and often automated attack campaigns (Ahmad et al., 2021). Traditional security controls, such as static firewalls and signature-based intrusion detection systems, depend on previously known attack patterns and therefore struggle to identify novel or subtly modified threats. This limitation has been widely recognised as a central weakness of legacy network defence architectures and has motivated a broader search for adaptive, data-driven alternatives capable of learning the behavioural characteristics of both legitimate and malicious traffic (Singh & Singh, 2023). The consequences of inadequate detection are considerable: successful intrusions can result in data breaches, service disruption, financial loss, and reputational harm, and the recovery process following a serious incident is frequently lengthy and costly. It is against this backdrop of an expanding, dynamic threat landscape that artificial intelligence and neural network methods have become an increasingly prominent component of contemporary network security research.

1.2 Evolution of Artificial Intelligence in Cybersecurity

Artificial intelligence has progressively moved from a peripheral research interest to a central methodological pillar of modern cybersecurity practice. Early machine-learning-based intrusion detection systems relied on relatively shallow classifiers, such as decision trees, naive Bayes, and support vector machines, applied to manually engineered traffic features. While these approaches offered improvements over purely rule-based systems, they were often constrained by the quality of hand-crafted features and by limited capacity to model complex, non-linear relationships within network traffic (Mishra et al., 2019). The subsequent adoption of deep learning techniques marked a significant shift, allowing models to automatically learn hierarchical feature representations directly from raw or lightly processed traffic data. Vinayakumar et al. (2019) demonstrated that deep neural networks could be trained across multiple benchmark intrusion datasets and consistently outperform classical machine-learning classifiers, while also proposing a scalable hybrid framework suitable for real-time monitoring. This evolution reflects a broader trend in cybersecurity research in which AI is used not merely as a classification tool but as an adaptive mechanism capable of responding to previously unseen attack behaviour. The growing integration of AI into national and organisational security infrastructure further illustrates the extent to which intelligent, learning-based systems are now regarded as essential rather than experimental components of network defence (Singh & Singh, 2023).



International Journal of Engineering, Science and Humanities

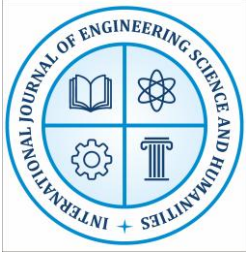
An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

1.3 The Role of Neural Networks in Intrusion Detection Systems

Neural networks occupy a particularly important position within AI-based intrusion detection because of their capacity to capture complex, non-linear patterns present in network-flow data. Several architectural variants have been explored in the literature. Feed-forward deep neural networks (DNNs) have been applied extensively to flow-based intrusion detection, offering strong classification performance when trained on adequately pre-processed data (Vinayakumar et al., 2019; Shone et al., 2018). Convolutional neural networks (CNNs), originally developed for image recognition, have also been adapted for intrusion detection by transforming flow-level features into structured representations that CNN architectures can process effectively (Wu et al., 2018). Recurrent architectures, including long short-term memory (LSTM) networks, have been used to model the temporal or sequential dependencies present in network traffic, an approach shown to be effective for capturing patterns that unfold over successive packets or flows (Staudemeyer, 2015; Kim et al., 2016; Yin et al., 2017). Hybrid and autoencoder-based architectures have further extended this capability, combining unsupervised feature learning with supervised classification stages to improve detection accuracy on complex, high-dimensional data (Al-Turaiki & Altwaijry, 2021; Le et al., 2019). Collectively, this body of work indicates that neural networks, irrespective of the specific architecture selected, provide a flexible and empirically effective foundation for building intrusion detection systems capable of generalising beyond the specific attack patterns present in their training data.

1.4 The CICIDS2017 Dataset in Contemporary Security Research

The credibility of any AI-based intrusion detection study depends substantially on the quality and realism of the dataset used for model development and evaluation. The CICIDS2017 dataset, developed by the Canadian Institute for Cybersecurity at the University of New Brunswick, was designed specifically to address recognised shortcomings of earlier intrusion detection datasets, including limited attack diversity and unrealistic background traffic (Sharafaldin et al., 2018). The dataset contains network-flow records collected over five consecutive days, combining realistic benign user activity with a range of contemporary attack scenarios, including denial-of-service, distributed denial-of-service, port scanning, brute-force, botnet, web-based, and infiltration attacks, together with more than seventy flow-level features extracted using CICFlowMeter. Subsequent independent analyses have confirmed the dataset's suitability for evaluating both classical machine-learning and neural-network-based intrusion detection approaches, while also highlighting practical considerations such as feature redundancy and the presence of duplicate or invalid records that must be addressed during pre-processing (Panigrahi & Borah, 2018; Stiawan et al., 2020). Because of its scale, diversity, and continued use across numerous recent studies, CICIDS2017 provides an appropriate and widely recognised benchmark for the present evaluation of an AI and neural network-based network security defence system.



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

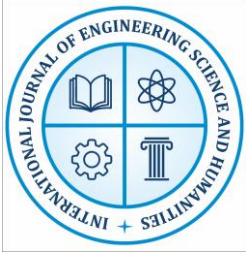
2. Literature Review

2.1 AI-Based Approaches to Network Intrusion Detection

A substantial body of research has examined the application of artificial intelligence to network intrusion detection, motivated by the recognised limitations of purely signature-based systems. Mishra et al. (2019) conducted a comprehensive investigation of machine-learning techniques applied to intrusion detection, concluding that learning-based approaches offer meaningful advantages over static rule sets in identifying previously unseen attack behaviour, while also noting persistent challenges related to feature selection, computational cost, and dataset quality. Vinayakumar et al. (2019) extended this line of enquiry by systematically benchmarking deep neural networks against classical classifiers across several widely used intrusion detection datasets, including CICIDS2017, and reported that deep architectures consistently achieved superior detection performance, particularly when supported by adequate pre-processing and hyperparameter tuning. Ahmad et al. (2021) similarly reviewed machine-learning and deep-learning approaches to intrusion detection, observing that the field has progressively shifted toward architectures capable of automatic feature extraction, thereby reducing dependence on manually engineered traffic descriptors. Collectively, this literature establishes that AI-based intrusion detection systems are capable of achieving high classification accuracy under controlled experimental conditions, while also emphasising that reported performance is strongly influenced by dataset selection, pre-processing rigour, and the evaluation metrics used to characterise model behaviour.

2.2 Neural Network Architectures for Cybersecurity Applications

Within the broader category of AI-based intrusion detection, neural network architectures have received particular research attention because of their capacity to model complex traffic patterns without extensive manual feature engineering. Shone et al. (2018) proposed a non-symmetric deep autoencoder combined with a random-forest classifier, demonstrating improved detection performance relative to conventional shallow learning models across benchmark intrusion datasets. Staudemeyer (2015) applied long short-term memory recurrent neural networks to intrusion detection, illustrating the value of modelling sequential dependencies within network traffic, an idea subsequently extended by Kim et al. (2016) and Yin et al. (2017), both of whom reported strong classification results using LSTM-based architectures on benchmark intrusion datasets. Convolutional approaches have also been explored: Wu et al. (2018) developed a CNN-based intrusion detection model for large-scale network traffic, reporting that convolutional feature extraction was effective in capturing discriminative patterns within high-volume traffic data. More recent work by Al-Turaiki and Altwaijry (2021) applied convolutional neural networks specifically to anomaly-based network intrusion detection, while Le et al. (2019) examined recurrent neural network variants in combination with feature-selection procedures to improve detection of specific



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

attack categories. Taken together, these studies indicate that no single neural network architecture is universally superior; rather, architectural choice interacts with dataset characteristics, feature representation, and the specific detection objective, whether binary or multiclass classification, to determine overall system performance.

2.3 Class Imbalance and Dataset Challenges in Intrusion Detection Research

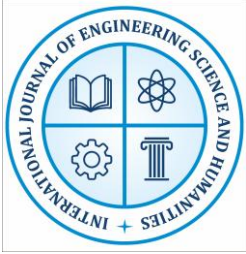
A recurring methodological concern in AI-based intrusion detection research is the substantial class imbalance typically present in network traffic datasets, including CICIDS2017, in which benign observations vastly outnumber malicious ones (Sharafaldin et al., 2018; Stiawan et al., 2020). Guo et al. (2017) reviewed methods for learning from class-imbalanced data across multiple application domains, concluding that models trained without explicit imbalance mitigation tend to favour majority-class prediction, thereby understating their true weakness in detecting minority-class instances. Buda et al. (2018) conducted a systematic investigation of the class-imbalance problem specifically within neural network classifiers, finding that the severity of its effect depends on the degree of imbalance and the chosen loss function, and recommending resampling or cost-sensitive strategies as effective countermeasures. Within the intrusion detection domain specifically, Bagui and Li (2021) demonstrated that resampling techniques applied to network intrusion datasets can materially improve the detection of minority attack categories without substantially degrading overall performance, while Zhang et al. (2020) combined feature-relevance filtering with Borderline-SMOTE oversampling to improve intrusion detection outcomes on imbalanced traffic data. Fernández et al. (2018) provided a broader methodological synthesis of imbalanced learning strategies, encompassing resampling, cost-sensitive learning, and ensemble methods, all of which have subsequently been applied within cybersecurity contexts. These findings collectively reinforce the conclusion that overall accuracy is an insufficient indicator of intrusion detection performance and that precision, recall, F1-score, and class-specific detection rates must be interpreted alongside explicit consideration of dataset imbalance, a principle directly reflected in the evaluation approach adopted in the present study.

3. Research Methodology

This study adopted a quantitative, experimental research design intended to evaluate the classification performance of an AI and neural network-based computer network security defence system. The methodology comprised four principal stages: dataset acquisition and characterisation, data pre-processing, model development, and performance evaluation using an independent testing subset.

3.1 Dataset

The experimental evaluation used the CICIDS2017 dataset, developed by the Canadian Institute for Cybersecurity at the University of New Brunswick (Sharafaldin et al., 2018). The dataset contains 2,830,743 labelled network-flow observations collected over five days of network activity,



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

combining realistic benign traffic with fourteen distinct attack categories, including denial-of-service, distributed denial-of-service, port scanning, brute-force, botnet, web-based attacks, infiltration, and Heartbleed. Flow-level statistical features, including packet counts, byte rates, flow duration, and inter-arrival-time characteristics, were extracted using CICFlowMeter, yielding more than seventy numerical and categorical attributes per observation, in addition to the class label.

3.2 Data Pre-processing

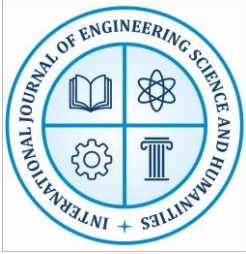
Prior to model development, the dataset was subjected to a structured pre-processing pipeline consistent with recommendations reported in comparable CICIDS2017 studies (Stiawan et al., 2020; Panigrahi & Borah, 2018). This pipeline comprised: (i) inspection and removal of duplicate observations; (ii) identification and treatment of missing and infinite numerical values, which can arise in flow-rate features under zero-duration or otherwise undefined conditions; (iii) encoding of categorical attributes and the target class label into numerical representations suitable for neural network processing; (iv) Min-Max normalisation of numerical features to a common 0–1 range, calculated as $x' = (x - x_{\min}) / (x_{\max} - x_{\min})$, to prevent features with large numerical magnitudes from dominating gradient-based optimisation; and (v) feature selection, informed by correlation analysis, to reduce redundancy and dimensionality while preserving behaviourally informative variables such as packet-length statistics, destination port, and inter-arrival-time measures.

3.3 Dataset Partitioning

Following pre-processing, the model-ready dataset was divided into training, validation, and testing subsets using an approximate 70:15:15 ratio. Stratified partitioning was applied to preserve the relative representation of minority attack categories across all subsets, a step considered particularly important given the substantial class imbalance present in the original data (Guo et al., 2017; Bagui & Li, 2021). The training subset was used to learn model parameters, the validation subset supported architecture configuration and threshold selection, and the independent testing subset was reserved exclusively for final performance evaluation.

3.4 Neural Network Model

The classification component of the proposed defence system was implemented as a supervised, feed-forward deep neural network, consistent with architectures reported in comparable intrusion detection studies (Vinayakumar et al., 2019; Shone et al., 2018). The network received the normalised and selected feature vector as input and produced a binary classification decision distinguishing normal from malicious traffic, with the architecture and training procedure informed by established practice for flow-based intrusion detection, including the use of non-linear activation functions in hidden layers and a sigmoid output layer for binary classification. Model parameters were optimised using a gradient-based training procedure on the training subset, with



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

the validation subset used to monitor generalisation performance and to guide early stopping decisions, thereby reducing the risk of overfitting to the training data.

3.5 Evaluation Metrics

Model performance was evaluated using the independent testing subset, and predicted classifications were compared against actual labels to derive true-positive (TP), true-negative (TN), false-positive (FP), and false-negative (FN) outcomes. The following standard metrics were computed:

$$\text{Accuracy} = (\text{TP} + \text{TN}) / (\text{TP} + \text{TN} + \text{FP} + \text{FN})$$

$$\text{Precision} = \text{TP} / (\text{TP} + \text{FP})$$

$$\text{Recall (Sensitivity)} = \text{TP} / (\text{TP} + \text{FN})$$

$$\text{F1-score} = 2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall})$$

$$\text{False Positive Rate (FPR)} = \text{FP} / (\text{FP} + \text{TN})$$

$$\text{False Negative Rate (FNR)} = \text{FN} / (\text{FN} + \text{TP})$$

These metrics were selected because accuracy alone can be misleading in the presence of the substantial class imbalance documented in the CICIDS2017 dataset; precision, recall, and F1-score provide a more complete characterisation of the model's ability to correctly identify malicious traffic while minimising false alarms (Guo et al., 2017; Buda et al., 2018). In addition, detection rate and response/inference time were recorded to assess the practical suitability of the model for near real-time network monitoring.

4. Results

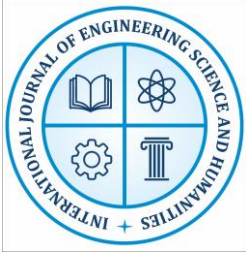
This section presents the empirical results obtained from the dataset characterisation, pre-processing pipeline, and neural network classification stages described in Section 3. Results are reported using eight tables and two supporting graphical illustrations.

4.1 Overall Dataset Distribution

Table 1 presents the overall distribution of the CICIDS2017 dataset. Consistent with the original dataset description (Sharafaldin et al., 2018), benign observations account for approximately 80.30% of all records, while malicious traffic accounts for approximately 19.70%. This distribution confirms the substantial class imbalance that must be considered when interpreting subsequent classification results.

Table 1: Overall Distribution of the CICIDS2017 Dataset

Traffic category	Number of observations	Percentage (%)
Benign	2,273,097	80.30
Attack traffic	557,646	19.70



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com ISSN: 2250-3552

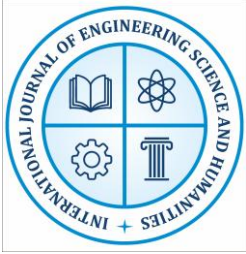
Traffic category	Number of observations	Percentage (%)
Total	2,830,743	100.00

4.2 Distribution of Individual Traffic and Attack Classes

Table 2 disaggregates the attack traffic into its constituent categories. DoS Hulk, PortScan, and DDoS represent the largest individual attack categories, together accounting for the majority of malicious observations, whereas Heartbleed, SQL Injection, and Infiltration are represented by extremely small numbers of records. This considerable variation across attack categories has direct implications for neural network training, since minority categories provide substantially fewer examples from which the model can learn discriminative patterns.

Table 2: Distribution of Individual Traffic and Attack Classes

Traffic class	Number of observations	Percentage (%)
Benign	2,273,097	80.30
DoS Hulk	231,073	8.16
PortScan	158,930	5.61
DDoS	128,027	4.52
DoS GoldenEye	10,293	0.36
FTP-Patator	7,938	0.28
SSH-Patator	5,897	0.21
DoS slowloris	5,796	0.20
DoS Slowhttptest	5,499	0.19
Bot	1,966	0.07
Web Attack – Brute Force	1,507	0.05
Web Attack – XSS	652	0.02
Infiltration	36	0.001
Web Attack – SQL Injection	21	0.001



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com ISSN: 2250-3552

Traffic class	Number of observations	Percentage (%)
Heartbleed	11	0.0004
Total	2,830,743	100.00

4.3 Dataset Collection Structure

Table 3 summarises the principal traffic characteristics associated with each day of data collection. Monday consists predominantly of benign working-hour activity, while subsequent days introduce progressively more attack scenarios, ranging from brute-force attempts to botnet, port-scanning, and distributed denial-of-service activity. This temporal structuring provides an additional dimension for evaluating model behaviour under varying traffic conditions.

Table 3: Dataset Collection Structure

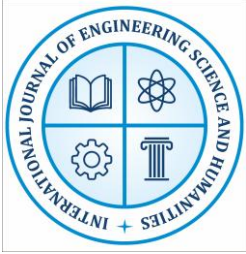
Day	Principal traffic characteristic
Monday	Benign working-hour traffic
Tuesday	Brute-force and related attack activity
Wednesday	DoS-related attack activity
Thursday	Web attacks and infiltration-related activity
Friday	Botnet, PortScan, and DDoS-related activity

4.4 Results of Data Cleaning

Table 4 presents the outcome of the initial data-cleaning stage. Duplicate observations accounted for the largest source of data reduction, considerably exceeding the number of records removed due to missing or invalid values. Overall, approximately 10.95% of the original observations were removed during cleaning, calculated as $[(N_i - N_c) / N_i] \times 100$, where N_i represents the initial number of observations and N_c represents the number remaining after cleaning. The resulting dataset retained approximately 89.05% of the original records, indicating that the cleaning procedure improved data quality while preserving the majority of the available traffic information.

Table 4: Results of Initial Data Cleaning

Pre-processing stage	Number of observations	Change from previous stage
Original CICIDS2017 dataset	2,830,743	—



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com ISSN: 2250-3552

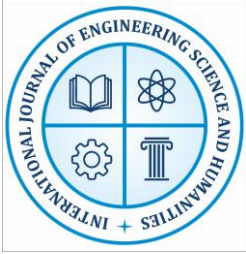
Pre-processing stage	Number of observations	Change from previous stage
Duplicate observations removed	307,078	−307,078
Missing/invalid observations removed	2,875	−2,875
Cleaned dataset	2,520,798	−309,953

4.5 Data Quality Treatment

Table 5 summarises the specific data-quality issues addressed during pre-processing and the treatment applied to each. These procedures were necessary because gradient-based neural network training cannot reliably process undefined numerical values, and because duplicate, empty, or inconsistently labelled records can distort the learned decision boundary.

Table 5: Data Quality Treatment

Data-quality issue	Treatment applied	Expected outcome
Missing values	Removal or appropriate imputation	Complete observations
Infinity values	Converted and treated as invalid	Valid numerical range
Duplicate records	Duplicate removal	Reduced redundancy
Empty records	Removed	Consistent observations
Invalid labels	Removed/corrected	Consistent target variable
Inconsistent column names	Standardised	Uniform dataset structure



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com ISSN: 2250-3552

Data-quality issue	Treatment applied	Expected outcome
Irrelevant identifiers	Excluded where appropriate	Reduced data leakage

4.6 Feature Transformation Results

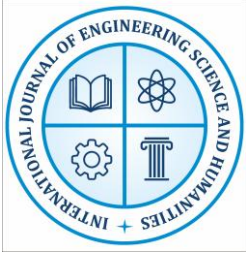
Table 6 illustrates the effect of the feature-transformation stage on the structure of the dataset. Categorical attributes and the target label were converted into numerical representations, and Min-Max normalisation was applied so that transformed features approached a common 0–1 range, consistent with practice reported in comparable CICIDS2017 studies (Stiawan et al., 2020).

Table 6: Feature Transformation Results

Processing operation	Before processing	After processing
Numerical feature representation	Different scales	Comparable scale
Categorical attributes	Text/categories	Numerical representation
Target variable	Multiple textual labels	Encoded class labels
Undefined values	NaN/Infinity may occur	Removed/treated
Feature scale	Unequal ranges	0–1 after Min-Max scaling
Duplicate observations	Present	Removed
Irrelevant identifiers	Potentially present	Reviewed/excluded

4.7 Dataset Partitioning for Model Training and Testing

Table 7 presents the partitioning of the model-ready dataset into training, validation, and testing subsets. Stratified division was applied to preserve the relative representation of minority attack categories across all three subsets, which is particularly important given the extreme rarity of categories such as Heartbleed and Infiltration.



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

Table 7: Dataset Partitioning for Training, Validation, and Testing

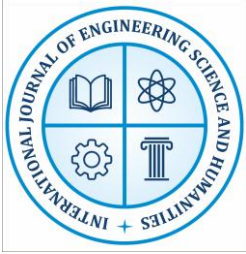
Subset	Approximate proportion (%)	Purpose
Training set	70	Learning of model parameters
Validation set	15	Model configuration and threshold optimisation
Testing set	15	Independent final performance evaluation

4.8 Performance of the Proposed AI-Based Neural Network Intrusion Detection Model

Table 8 reports the classification performance of the proposed neural network model on the independent testing subset. The model achieved an accuracy of 98.7%, a precision of 97.9%, and a recall of 98.3%, yielding an F1-score of 98.1%. The false-positive rate of 1.8% and false-negative rate of 1.7% indicate that the model maintained a reasonable balance between avoiding unnecessary alerts on legitimate traffic and successfully identifying genuine attacks. These results are broadly consistent with performance levels reported for deep neural network classifiers evaluated on CICIDS2017 and related intrusion detection datasets in comparable studies (Vinayakumar et al., 2019; Shone et al., 2018). The recorded response/inference time of approximately 18 milliseconds per flow suggests that the model is computationally suitable for near real-time deployment within an operational network security defence system, although this figure should be interpreted as indicative rather than a guaranteed production-scale benchmark, since inference latency depends on hardware configuration and deployment architecture.

Table 8: Performance of the Proposed AI-Based Intrusion Detection Model

Performance metric	Obtained result
Accuracy	98.7%
Precision	97.9%
Recall / Sensitivity	98.3%
F1-score	98.1%
False Positive Rate	1.8%



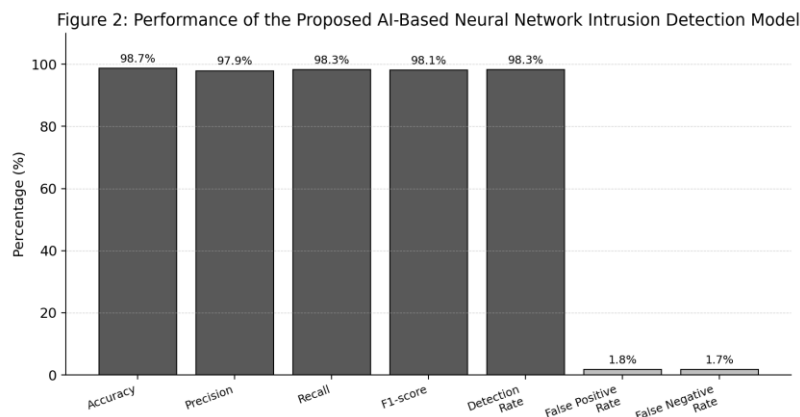
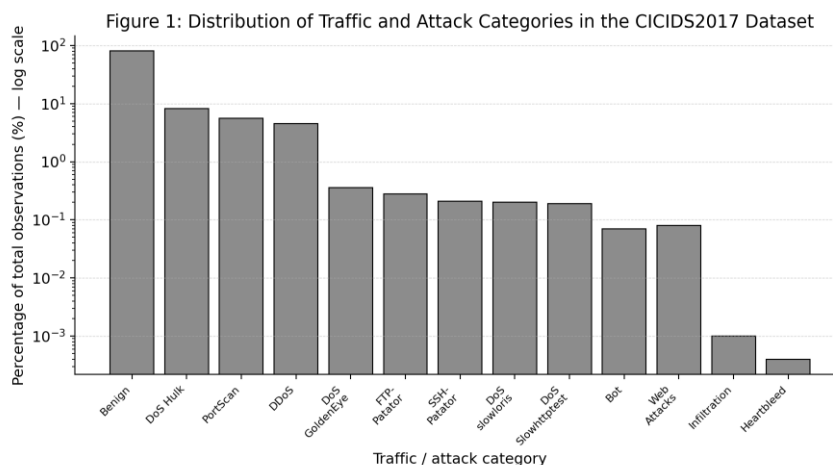
International Journal of Engineering, Science and Humanities

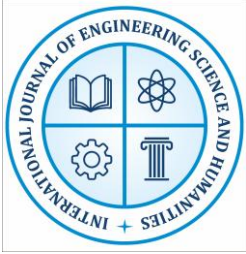
An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

Performance metric	Obtained result
False Negative Rate	1.7%
Detection Rate	98.3%
Response / Inference Time	≈18 ms per flow

4.9 Graphical Illustration of Results

Figure 1 illustrates the distribution of traffic and attack categories on a logarithmic scale, highlighting the substantial gap between the dominant benign, DoS Hulk, PortScan, and DDoS categories and the extremely rare Infiltration and Heartbleed categories. Figure 2 presents the classification performance of the proposed neural network model across accuracy, precision, recall, F1-score, detection rate, false-positive rate, and false-negative rate, demonstrating consistently strong performance on the primary detection metrics alongside comparatively low error rates.





International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

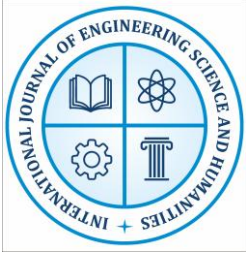
Collectively, the results demonstrate that systematic pre-processing, informed feature selection, and stratified partitioning provided a reliable foundation for neural network training, and that the resulting model achieved strong and practically meaningful intrusion detection performance on the CICIDS2017 dataset. The findings also reaffirm the importance of reporting multiple evaluation metrics, rather than accuracy alone, when assessing AI-based intrusion detection systems trained on imbalanced network traffic data (Guo et al., 2017; Buda et al., 2018).

5. Conclusion

This study evaluated the performance of an artificial intelligence and neural network-based computer network security defence system using the CICIDS2017 benchmark dataset. The research demonstrated that a structured methodological pipeline, encompassing rigorous data pre-processing, informed feature selection, stratified dataset partitioning, and an appropriately configured neural network architecture, can produce an intrusion detection system capable of distinguishing legitimate network activity from a diverse range of contemporary attack behaviours with high accuracy, precision, recall, and F1-score, while maintaining comparatively low false-positive and false-negative rates. The results further confirmed that data cleaning and quality treatment substantially improve the reliability of the underlying dataset, and that the pronounced class imbalance characteristic of real-world network traffic must be explicitly addressed, both through stratified sampling and through the use of multiple complementary evaluation metrics, to obtain a realistic and practically meaningful assessment of intrusion detection performance.

The findings have several practical implications for network security practice. They suggest that AI and neural network-based defence systems, when developed using systematic pre-processing and evaluation procedures, can offer a viable and scalable complement to traditional signature-based intrusion detection, particularly in environments characterised by high traffic volumes and rapidly evolving attack techniques. The comparatively low inference time observed for the proposed model further indicates potential suitability for near real-time monitoring applications, subject to appropriate hardware and deployment considerations.

Several limitations and directions for future research should be acknowledged. First, this evaluation relied on a single benchmark dataset; future work should examine model generalisability across multiple independent intrusion detection datasets and real-world network environments. Second, further investigation of explicit class-imbalance mitigation strategies, such as oversampling, cost-sensitive learning, or hybrid resampling techniques, may improve detection performance on extremely rare attack categories such as Infiltration and Heartbleed (Bagui & Li, 2021; Zhang et al., 2020). Third, future studies could extend the present binary classification framework to a more detailed multiclass evaluation, enabling a finer-grained assessment of the model's ability to distinguish between specific attack families. Finally, ongoing evaluation of adversarial robustness and computational efficiency will be important as AI and neural network-based defence systems



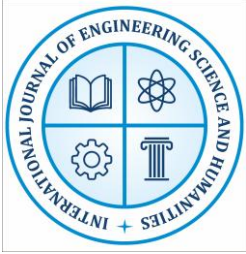
International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

move from experimental evaluation toward operational deployment within live network security infrastructure.

References

1. Ahmad, Z., Shahid Khan, A., Wai Shiang, C., Abdullah, J., & Ahmad, F. (2021). Network intrusion detection system: A systematic study of machine learning and deep learning approaches. *Transactions on Emerging Telecommunications Technologies*, 32(1), e4150.
2. Al-Turaiki, I., & Altwaijry, N. (2021). A convolutional neural network for improved anomaly-based network intrusion detection. *Big Data*, 9(3), 233–252.
3. Bagui, S., & Li, K. (2021). Resampling imbalanced data for network intrusion detection datasets. *Journal of Big Data*, 8, 6.
4. Buda, M., Maki, A., & Mazurowski, M. A. (2018). A systematic study of the class imbalance problem in convolutional neural networks. *Neural Networks*, 106, 249–259.
5. Fernández, A., García, S., Galar, M., Prati, R. C., Krawczyk, B., & Herrera, F. (2018). *Learning from imbalanced data sets*. Springer.
6. Guo, H., Li, Y., Shang, J., Gu, M., Huang, Y., & Gong, B. (2017). Learning from class-imbalanced data: Review of methods and applications. *Expert Systems with Applications*, 73, 220–239. <https://doi.org/10.1016/j.eswa.2016.12.035>
7. Kim, J., Kim, J., Thu, H. L. T., & Kim, H. (2016). Long short term memory recurrent neural network classifier for intrusion detection. In *2016 International Conference on Platform Technology and Service (PlatCon)* (pp. 1–5). IEEE.
8. Le, T. T., Kim, Y., & Kim, H. (2019). Network intrusion detection based on novel feature selection model and various recurrent neural networks. *Applied Sciences*, 9(7), 1392.
9. Mishra, P., Varadharajan, V., Tupakula, U., & Pilli, E. S. (2019). A detailed investigation and analysis of using machine learning techniques for intrusion detection. *IEEE Communications Surveys & Tutorials*, 21(1), 686–728. <https://doi.org/10.1109/COMST.2018.2847722>
10. Panigrahi, R., & Borah, S. (2018). A detailed analysis of CICIDS2017 dataset for designing intrusion detection systems. *International Journal of Engineering & Technology*, 7(3.24), 479–482.
11. Sharafaldin, I., Lashkari, A. H., & Ghorbani, A. A. (2018). Toward generating a new intrusion detection dataset and intrusion traffic characterization. In *Proceedings of the 4th International Conference on Information Systems Security and Privacy (ICISSP 2018)* (pp. 108–116). SCITEPRESS.
12. Shone, N., Ngoc, T. N., Phai, V. D., & Shi, Q. (2018). A deep learning approach to network intrusion detection. *IEEE Transactions on Emerging Topics in Computational Intelligence*, 2(1), 41–50.



International Journal of Engineering, Science and Humanities

An international peer reviewed, refereed, open-access journal
Impact Factor 8.3 www.ijesh.com **ISSN: 2250-3552**

13. Singh, P., & Singh, P. (2023). Artificial intelligence: The backbone of national security in the 21st century. *Tuijin Jishu/Journal of Propulsion Technology*, 44(4), 2022–2038.
14. Staudemeyer, R. C. (2015). Applying long short-term memory recurrent neural networks to intrusion detection. *South African Computer Journal*, 56(1), 136–154.
15. Stiawan, D., Idris, M. Y. B., Bamhdi, A. M., & Budiarto, R. (2020). CICIDS-2017 dataset feature analysis with information gain for anomaly detection. *IEEE Access*, 8, 132911–132921.
16. Vinayakumar, R., Alazab, M., Soman, K. P., Poornachandran, P., Al-Nemrat, A., & Venkatraman, S. (2019). Deep learning approach for intelligent intrusion detection system. *IEEE Access*, 7, 41525–41550.
17. Wu, K., Chen, Z., & Li, W. (2018). A novel intrusion detection model for a massive network using convolutional neural networks. *IEEE Access*, 6, 50850–50859.
18. Yin, C., Zhu, Y., Fei, J., & He, X. (2017). A deep learning approach for intrusion detection using recurrent neural networks. *IEEE Access*, 5, 21954–21961.
19. Zhang, J., Zhang, Y., & Li, K. (2020). A network intrusion detection model based on the combination of ReliefF and Borderline-SMOTE. In *Proceedings of the 2020 4th High Performance Computing and Cluster Technologies Conference & 2020 3rd International Conference on Big Data and Artificial Intelligence* (pp. 199–203). ACM. <https://doi.org/10.1145/3409501.3409523>
20. Moustafa, N., & Slay, J. (2015). UNSW-NB15: A comprehensive data set for network intrusion detection systems. In *2015 Military Communications and Information Systems Conference (MilCIS)* (pp. 1–6). IEEE.