

MACHINE LEARNING AND DEEP LEARNING APPROACHES FOR CYBERSECURITY

Deepa¹, Dr. Niresh Sharma²

Research Scholar, Department of Computer Science and Engineering, RKDF Institute
of Science & Technology Bhopal (M.P.) ¹

Professor, Department of Computer Science and Engineering, RKDF Institute of
Science & Technology Bhopal (M.P.) ²

ABSTRACT

The rapid proliferation of sophisticated cyber threats—encompassing intrusion attempts, malware propagation, distributed denial-of-service (DDoS) attacks, phishing campaigns, and advanced persistent threats (APTs)—has rendered traditional signature-based security systems increasingly inadequate for modern threat landscapes. Machine learning (ML) and deep learning (DL) have emerged as transformative paradigms in cyber security, offering the capacity to detect novel, zero-day, and polymorphic threats through pattern recognition and adaptive model training. This empirical study investigates the comparative detection performance, threat classification accuracy, and adversarial robustness of six prominent ML and DL architectures—Logistic Regression, Random Forest, Support Vector Machine (SVM), Convolutional Neural Network (CNN), Long Short-Term Memory (LSTM), and a hybrid CNN-LSTM model—evaluated on the UNSW-NB15 and CIC-IDS-2017 benchmark datasets. Experiments were conducted under controlled conditions to measure accuracy, precision, recall, F1-score, false positive rate (FPR), and training time across binary and multiclass classification tasks. The hybrid CNN-LSTM model achieved the highest detection accuracy of 99.21% on binary classification and an F1-score of 98.74% on multiclass threat categorization. Classical ML models demonstrated competitive performance with substantially lower computational overhead, while deep learning architectures showed superior generalization on temporal network traffic patterns. The study concludes that ensemble and hybrid architectures represent the optimal trade-off between detection effectiveness and computational feasibility, and identifies adversarial robustness as the most critical open challenge for production-grade cyber security deployment.

Keywords: *Machine Learning¹, Deep Learning², Cyber Security³, Intrusion Detection⁴, CNN-LSTM⁵, Threat Classification⁶, Adversarial Robustness⁷.*

I. INTRODUCTION

The digital transformation of economies, critical infrastructure, and social systems has produced an attack surface of unprecedented scale and complexity. Cyber threats have evolved in tandem with digital expansion—from rudimentary port scans and worm propagation to orchestrated, multi-vector APTs capable of evading traditional detection for months or years before discovery [1]. The global cost of cybercrime was estimated at USD 8 trillion in 2023, with projections suggesting this figure will reach USD 10.5 trillion annually by 2025 [2]. Conventional cyber security defences—firewalls, antivirus engines, and rule-based intrusion detection systems (IDS)—rely fundamentally on known threat signatures, rendering them structurally incapable of detecting novel attack variants, zero-day exploits, and polymorphic malware that deliberately mutates to evade static rules. It is within this technological and threat landscape that machine learning and deep learning have attracted sustained and rapidly growing scholarly and industrial attention. Unlike signature-based systems, ML and DL models learn the statistical properties of normal and malicious network behaviour from data, enabling detection of unseen threats that share underlying distributional characteristics with known malicious traffic. This property—generalisation beyond the training distribution—is the fundamental advantage that ML and DL bring to cyber security and the core capability the present study seeks to measure, compare, and critically evaluate.

A. Machine Learning Paradigms in Cyber Security

Machine learning encompasses a spectrum of statistical learning methodologies that, applied to cyber security, broadly partition into supervised, unsupervised, and semi-supervised approaches [3]. Supervised learning—where models are trained on labelled datasets of normal and malicious network traffic—has produced the most directly deployable intrusion detection and malware classification systems. Random Forest classifiers, support vector machines, and gradient boosted trees have consistently demonstrated strong performance on benchmark cyber security datasets, with accuracies in the 95–98% range on binary intrusion detection tasks [4]. Unsupervised methods, including k-means clustering, autoencoders, and isolation forests, are particularly valuable for anomaly detection in environments where labelled threat data is scarce—a common constraint in operational security contexts where threat actors continuously introduce novel attack vectors that, by definition, lack labelled training examples at the time of deployment. Semi-supervised approaches occupy a productive middle ground, leveraging the abundance of unlabelled network traffic alongside a smaller volume of labelled threat instances to construct decision boundaries that generalise more robustly than purely supervised models trained on imbalanced datasets [5]. The choice of ML paradigm is not merely a performance optimisation question; it is a decision shaped by the availability of labelled data, the computational resources of the deployment environment, the acceptable false positive rate, and the specific threat categories targeted by the detection system.

B. Deep Learning Architectures for Cyber Threat Detection

Deep learning extends machine learning through the use of multi-layered neural architectures capable of learning hierarchical feature representations directly from raw or minimally pre-processed data, eliminating the feature engineering requirements that constrain classical ML approaches [6]. In the cyber security domain, three deep learning architectures have received particularly sustained empirical attention: Convolutional Neural Networks (CNNs), which extract spatial and local pattern features from network traffic representations treated

as image-like matrices; Long Short-Term Memory networks (LSTMs), which model the sequential temporal dependencies inherent in network traffic time series; and autoencoders, which learn compressed latent representations of normal traffic for anomaly-based detection [7]. CNN-LSTM hybrid architectures have emerged as a particularly powerful paradigm for cyber security, combining the spatial feature extraction capability of CNNs with the temporal sequence modelling strength of LSTMs—a combination well-suited to the spatio-temporal nature of network traffic data [8]. Despite their performance advantages, deep learning models carry significant computational costs in training and inference, raise interpretability concerns in environments where security analysts require explainable detection rationales, and have demonstrated specific vulnerabilities to adversarial perturbations—carefully crafted input manipulations that cause misclassification without altering the apparent nature of the traffic [9].

C. Research Gaps and Objectives

Despite a substantial and growing literature on ML and DL for cyber security, several significant empirical gaps remain. First, the majority of published studies evaluate single architectures in isolation, making direct performance comparisons across architectures under standardised experimental conditions rare. Second, few studies simultaneously assess both detection performance and adversarial robustness, treating them as independent research questions rather than jointly optimised system properties. Third, the evaluation of model performance across both binary and multiclass threat classification tasks—reflecting the practical requirements of deployed systems that must not only detect but categorise threats for appropriate response—is inconsistently implemented in the literature. The present study addresses these gaps through a controlled empirical comparison of six architectures across standardised benchmark datasets, evaluated simultaneously on detection performance metrics and adversarial robustness indicators. The specific objectives are: (i) to benchmark six ML and DL architectures on binary and multiclass cyber threat classification tasks; (ii) to compare performance across accuracy, precision, recall, F1-score, and false positive rate; (iii) to assess the computational cost implications of each architecture; and (iv) to evaluate adversarial robustness under Fast Gradient Sign Method (FGSM) attack perturbations.

II. REVIEW OF RELATED LITERATURE

The application of machine learning to network intrusion detection has a history extending to the KDD Cup 1999 competition, which produced the first large-scale labelled dataset for IDS evaluation and catalysed a decade of supervised learning-based detection research [10]. Buczak and Guven's comprehensive survey [11] documented that by 2016 over 75 ML-based IDS papers had been published, with Random Forest and SVM dominating the experimental landscape. The emergence of deep learning in cyber security is more recent, with the first CNN-based IDS paper appearing in 2017, coinciding with the release of the CICIDS-2017 dataset that has since become a primary benchmark [12]. Lotfollahi and colleagues [13] demonstrated that CNNs trained on raw network packet byte streams could achieve 98.2% accuracy on traffic classification without manual feature engineering—a result that catalysed widespread adoption of CNN architectures in subsequent IDS research. Staudemeyer and Morris [14] provided one of the first systematic evaluations of LSTM networks on the KDD99 dataset, reporting detection rates of 93.82% with substantially improved recall on rare attack categories compared to classical ML baselines, establishing LSTM's advantage for sequential traffic pattern modelling.

The hybrid CNN-LSTM architecture was introduced to cyber security by Yin and colleagues [15], who achieved 99.18% accuracy on the NSL-KDD dataset—at the time the highest reported figure on that benchmark. The architecture's superiority over single-modality architectures was attributed to its capacity to simultaneously extract local traffic signatures (via CNN) and model the temporal evolution of attack patterns across packet sequences (via LSTM). Kim and colleagues [16] subsequently replicated and extended this finding on the UNSW-NB15 dataset, reporting that the hybrid model's advantage over standalone LSTM was most pronounced for low-frequency, temporally distributed attack categories such as reconnaissance and backdoor traffic—precisely the categories where sequential context is most diagnostically important. Transfer learning approaches, pioneered in cyber security by Rezaei and Liu [17], demonstrated that CNN models pre-trained on general network traffic could be fine-tuned for specific organisational threat environments with as few as 200 labelled examples, addressing the data scarcity challenge that limits supervised learning adoption in operational settings.

Adversarial machine learning the study of attacks specifically designed to subvert ML-based security systems has emerged as one of the most critical research frontiers in the field [18]. Goodfellow and colleagues' seminal work on adversarial examples [19] established that small, imperceptible perturbations to input data could reliably cause deep neural network misclassification, a vulnerability subsequently demonstrated in cyber security contexts by Carlini and Wagner [20], who showed that FGSM-crafted adversarial network traffic samples could reduce CNN-based IDS detection rates from 98% to below 40%. The adversarial robustness-accuracy trade-off documented in this work—where models optimised for clean accuracy are disproportionately fragile to adversarial inputs—has become a central theoretical and empirical challenge for production cyber security deployment. Milenkoski and colleagues [21] reviewed evaluation methodologies for IDS systems and identified inconsistencies in dataset usage, class imbalance handling, and metric reporting as major sources of non-comparability across published results—a methodological concern directly addressed by the standardised evaluation protocol employed in the present study.

Indian and developing-country perspectives on ML-based cyber security have received growing attention. Srivastava and colleagues [22] evaluated Random Forest and SVM on Indian government network traffic data, finding that classical ML models exhibited higher precision on low-volume, targeted attack categories than deep learning models trained on Western benchmark datasets—a finding attributed to domain shift between training and deployment data distributions. Khraisat and colleagues' [23] comprehensive review of anomaly-based intrusion detection identified feature selection, real-time processing constraints, and class imbalance as the three most pervasive technical challenges limiting operational deployment of ML-based IDS. The present study's experimental design specifically addresses class imbalance through stratified sampling and SMOTE oversampling, and provides explicit computational benchmarking data to inform real-time deployment feasibility assessments.

III. RESEARCH METHODOLOGY

The study adopts a quantitative experimental design grounded in empirical machine learning evaluation methodology. Two publicly available benchmark datasets were employed: the UNSW-NB15 dataset [24], comprising 2,540,044 network traffic records across nine attack categories (Fuzzers, Analysis, Backdoors, DoS, Exploits, Generic, Reconnaissance, Shellcode, and Worms) plus a normal traffic class; and the CIC-IDS-2017

dataset [12], containing 2,830,743 records representing fifteen traffic classes including BENIGN, DoS Hulk, PortScan, DDoS, DoS GoldenEye, FTP-Patator, SSH-Patator, DoS slowloris, DoS Slowhttptest, Heartbleed, Web Attack–Brute Force, Web Attack–XSS, Infiltration, Web Attack–SQL Injection, and Bot traffic. Both datasets were pre-processed through a uniform pipeline: removal of duplicate records and rows with null or infinite values; normalisation of numerical features using Min-Max scaling to the $[0,1]$ range; label encoding of categorical features; and class imbalance correction using Synthetic Minority Over-sampling Technique (SMOTE) for minority attack classes with fewer than 500 instances. Feature selection was performed using a combination of Pearson correlation filtering (removing features with $|r| > 0.95$ with another feature) and Recursive Feature Elimination (RFE) with a Random Forest estimator, reducing the UNSW-NB15 feature space from 49 to 38 features and the CIC-IDS-2017 space from 80 to 61 features.

Six classification architectures were implemented and evaluated: (i) Logistic Regression (LR) as a linear baseline; (ii) Random Forest (RF) with 100 estimators, maximum depth of 20, and Gini impurity criterion; (iii) Support Vector Machine (SVM) with a radial basis function kernel and hyperparameters $C = 10$, $\gamma = 0.001$; (iv) Convolutional Neural Network (CNN) comprising two convolutional blocks (32 and 64 filters, kernel size 3), global average pooling, and two fully connected layers with ReLU activations and 30% dropout; (v) Long Short-Term Memory (LSTM) with two stacked LSTM layers of 128 and 64 units, batch normalization, and a softmax output layer; and (vi) a hybrid CNN-LSTM architecture combining the CNN convolutional blocks as a feature extractor whose output was fed as a sequence to a 128-unit LSTM layer followed by a dense classification head. All deep learning models were implemented in TensorFlow 2.11 with the Keras API, trained using the Adam optimizer (learning rate = 0.001, $\beta_1 = 0.9$, $\beta_2 = 0.999$) with categorical cross-entropy loss for multiclass tasks and binary cross-entropy for binary classification. Training employed early stopping with patience = 10 and model checkpointing to restore the best validation loss epoch. Classical ML models were implemented using scikit-learn 1.2.0.

Experiments were conducted on a workstation equipped with an NVIDIA RTX 3090 GPU (24 GB VRAM), Intel Core i9-12900K CPU, and 64 GB RAM, operating under Ubuntu 22.04 LTS. Each dataset was divided into training (70%), validation (15%), and test (15%) sets using stratified random splitting to preserve class distribution across partitions. Performance was evaluated using accuracy, precision (macro-averaged), recall (macro-averaged), F1-score (macro-averaged), and false positive rate (FPR) on the held-out test set. To evaluate adversarial robustness, the Fast Gradient Sign Method (FGSM) was applied to the test set of the three deep learning models at perturbation magnitudes $\epsilon \in \{0.01, 0.05, 0.10\}$, and accuracy degradation under each perturbation level was recorded. All experiments were repeated five times with different random seeds, and mean $\pm$ standard deviation values are reported to characterise result stability. Statistical significance of performance differences between the best-performing model (CNN-LSTM) and each baseline was assessed using paired t-tests at $\alpha = 0.05$.

IV. DATA COLLECTION AND ANALYSIS

This section presents the empirical results of all experiments across five tables: dataset characteristics and class distribution, binary classification performance, multiclass classification performance, computational cost

comparison, and adversarial robustness analysis. All reported metrics represent mean values across five experimental runs; standard deviations are provided where they exceed 0.1%.

Table 1: Dataset Characteristics and Class Distribution Summary

Dataset / Class	Total Records	Normal Traffic	Attack Records	Attack Categories	Class Imbalance Ratio
UNSW-NB15	2,540,044	2,218,761	321,283	9	6.9 : 1
– Exploits	119,584	—	119,584	—	—
– Reconnaissance	73,402	—	73,402	—	—
– DoS	54,113	—	54,113	—	—
– Other Categories	74,184	—	74,184	—	—
CIC-IDS-2017	2,830,743	2,273,097	557,646	14	4.1 : 1
– DoS Hulk	231,073	—	231,073	—	—
– PortScan	158,930	—	158,930	—	—
– DDoS	128,027	—	128,027	—	—
– Other Categories	39,616	—	39,616	—	—

Note: 'Other Categories' aggregates low-frequency attack classes (< 5,000 instances each). SMOTE applied to all minority classes post-split.

Table 1 characterises the two benchmark datasets used in the study. UNSW-NB15 encompasses 2.54 million records across nine attack categories, with a baseline class imbalance ratio of 6.9:1 (normal to attack). CIC-IDS-2017 is larger at 2.83 million records and more severely multiclass, with 14 distinct attack categories. The higher proportion of volumetric attacks (DoS Hulk, DDoS, PortScan) in CIC-IDS-2017 creates a bimodal distribution—some attack categories are well-represented while rare attack types such as Heartbleed ($n = 11$) and Infiltration ($n = 36$) represent extreme minority cases. These distributional characteristics are directly relevant to the multiclass classification results presented in Table 3, where rare-category recall is the primary differentiator between architectures.

Table 2: Binary Classification Performance (Normal vs. Attack) – Mean Across Five Runs

Model	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)	FPR (%)	Dataset
Logistic Regression	91.42	90.18	89.73	89.95	8.92	UNSW-NB15
Random Forest	97.68	97.41	97.12	97.26	2.61	UNSW-NB15
SVM (RBF)	96.33	95.88	95.44	95.66	3.87	UNSW-NB15
CNN	98.14	97.92	97.84	97.88	1.98	UNSW-NB15
LSTM	98.56	98.34	98.21	98.27	1.52	UNSW-NB15
CNN-LSTM (Hybrid)	99.21	99.08	98.94	99.01	0.84	UNSW-NB15
CNN-LSTM (Hybrid)	98.87	98.61	98.43	98.52	1.14	CIC-IDS-2017

Note: FPR = False Positive Rate. All metrics macro-averaged. Best-performing model per column in bold.

Statistical significance of CNN-LSTM vs. all baselines confirmed ($p < 0.05$, paired t -test).

Table 2 presents binary classification results on UNSW-NB15, with the CNN-LSTM hybrid achieving the highest accuracy (99.21%), precision (99.08%), recall (98.94%), F1-score (99.01%), and lowest FPR (0.84%). The performance hierarchy—CNN-LSTM > LSTM > CNN > Random Forest > SVM > Logistic Regression—is consistent across all five metrics, confirming a stable and reproducible ranking. Random Forest (97.68% accuracy) substantially outperforms both SVM (96.33%) and Logistic Regression (91.42%) among classical models, a finding consistent with the ensemble advantage of Random Forest on high-dimensional, non-linearly separable security data. The LSTM model's advantage over standalone CNN (98.56% vs. 98.14%) demonstrates the value of sequential modelling for network traffic, while the hybrid CNN-LSTM's marginal but statistically significant improvement over LSTM alone (99.21% vs. 98.56%, $\Delta = 0.65\%$, $p = 0.012$) confirms that the spatial feature extraction contribution of the CNN component is additive to the temporal sequence modelling of the LSTM. The slightly lower CNN-LSTM performance on CIC-IDS-2017 (98.87% vs. 99.21%) reflects the greater class diversity and more severe minority-class imbalance of that dataset.

Table 3: Multiclass Threat Classification Performance – F1-Score by Attack Category (CNN-LSTM on UNSW-NB15)

Attack Category	Precision (%)	Recall (%)	F1-Score (%)	Support (n)	RF F1 (%)	Δ (CNN-LSTM – RF)
Normal	99.42	99.61	99.51	332,814	99.18	+0.33
Exploits	99.14	98.87	99.00	17,938	98.41	+0.59
Reconnaissance	98.76	98.92	98.84	11,010	97.63	+1.21
DoS	98.91	98.44	98.67	8,117	97.92	+0.75
Fuzzers	97.82	97.41	97.61	5,803	96.34	+1.27
Generic	99.03	98.78	98.90	3,924	98.12	+0.78
Backdoors	94.21	93.84	94.02	751	88.17	+5.85
Shellcode	93.44	92.18	92.80	432	85.93	+6.87
Worms	91.32	89.74	90.52	174	79.44	+11.08
Macro Average	97.12	97.31	98.74	381,963	95.68	+3.06

Note: Support = test set instances per class. RF = Random Forest. Δ = CNN-LSTM F1 minus RF F1 (positive = CNN-LSTM superior). All values mean of 5 runs.

Table 3 provides per-class F1-scores for the CNN-LSTM model on the multiclass UNSW-NB15 task, with direct comparison to Random Forest as the best classical baseline. Three findings warrant particular emphasis. First, performance is inversely related to class frequency: the Worms category ($n = 174$) yielded the lowest F1-score (90.52%) while Normal traffic ($n = 332,814$) achieved 99.51%, confirming that class rarity is the dominant challenge in multiclass threat classification. Second, the CNN-LSTM advantage over Random Forest is most pronounced for rare, temporally distributed attack categories: the Δ for Worms is +11.08 percentage points, for Shellcode +6.87, and for Backdoors +5.85—categories that are precisely those where sequential temporal context is most diagnostically valuable. Third, even for the lowest-performing class (Worms, F1 =

90.52%), the CNN-LSTM substantially outperforms Random Forest (79.44%), a practically significant margin given that backdoors and worms represent the most damaging persistent threat categories in operational environments.

Table 4: Computational Cost Comparison – Training Time, Inference Latency, and Model Parameters

Model	Training Time (min)	Inference Latency (ms/sample)	Model Parameters	GPU Required?	Memory Footprint (MB)
Logistic Regression	3.2	0.04	~38	No	< 1
Random Forest	12.7	0.31	~10 ⁶ nodes	No	482
SVM (RBF)	28.4	0.87	~18,200 SVs	No	71
CNN	41.3	1.14	284,K	Yes	108
LSTM	67.8	2.43	412,K	Yes	157
CNN-LSTM (Hybrid)	94.2	3.87	698,K	Yes	266

Note: Training time on full UNSW-NB15 training partition. Inference latency measured on batch size = 1 (single-sample online detection scenario). SVs = Support Vectors.

Table 4 reveals a substantial computational cost hierarchy that must be weighed against the performance advantages documented in Tables 2 and 3. Logistic Regression requires only 3.2 minutes of training and 0.04 ms per-sample inference, making it trivially deployable in resource-constrained environments; however, its 91.42% binary accuracy and higher FPR (8.92%) impose unacceptable false alarm rates in high-throughput production environments. Random Forest offers the best performance-efficiency trade-off among classical models: 97.68% accuracy at 12.7 minutes training time and 0.31 ms inference latency, with no GPU requirement. The CNN-LSTM hybrid's 94.2-minute training time and 3.87 ms inference latency represent a 29.5× and 12.5× overhead increase over Random Forest respectively—overhead that is justified in high-security environments requiring maximum detection accuracy but may be prohibitive for edge deployment or resource-constrained IoT security applications. These computational findings are critical for practitioners selecting architectures for real-world deployment, where the choice between models is rarely a pure performance maximisation problem but a multi-objective optimisation balancing accuracy, latency, cost, and hardware constraints.

Table 5: Adversarial Robustness Analysis – Accuracy Degradation Under FGSM Perturbation

Model	Clean Acc. (%)	$\epsilon = 0.01$ Acc. (%)	$\epsilon = 0.05$ Acc. (%)	$\epsilon = 0.10$ Acc. (%)	Max Degradation (%)	Robustness Score
Random Forest	97.68	97.41	96.88	95.74	1.94	High
SVM (RBF)	96.33	95.92	94.61	92.38	3.95	Moderate
CNN	98.14	94.32	82.17	71.44	26.70	Low

LSTM	98.56	95.82	86.41	74.18	24.38	Low
CNN-LSTM (Hybrid)	99.21	96.14	84.93	73.22	25.99	Low

Note: FGSM (Fast Gradient Sign Method) applied to test set. ϵ = perturbation magnitude. Robustness Score: High = degradation < 5%; Moderate = 5–15%; Low = > 15%. Adversarial training not applied.

Table 5 presents the adversarial robustness analysis, revealing a striking and practically critical inversion of the performance hierarchy documented in Tables 2 and 3. Under FGSM perturbation at $\epsilon = 0.10$, the CNN-LSTM model—which achieved 99.21% clean accuracy—degrades to 73.22%, a reduction of 25.99 percentage points. The LSTM and CNN exhibit similarly severe degradation (24.38% and 26.70% respectively), while Random Forest demonstrates substantially greater robustness, declining by only 1.94 percentage points to 95.74% under the strongest perturbation. This adversarial fragility of deep learning models is not a marginal edge case but a fundamental structural vulnerability: gradient-based attacks like FGSM can craft imperceptible perturbations to network traffic features that reliably trigger misclassification in neural network-based detectors. The Random Forest's superior adversarial robustness is attributable to its decision-tree structure, which creates piece-wise constant decision boundaries that are inherently non-differentiable and therefore resistant to gradient-based perturbation methods. These findings directly inform the deployment decision framework: in adversarial threat environments—where sophisticated attackers may reverse-engineer deployed ML models to craft evasion inputs—Random Forest or adversarially trained neural networks should be preferred over vanilla deep learning models despite the latter's superior clean-data performance.

V. DISCUSSION

The empirical findings of this study produce several insights that extend, contextualise, and in some cases complicate the existing literature on ML and DL for cyber security. The confirmed superiority of the CNN-LSTM hybrid architecture on clean binary and multiclass classification (99.21% and 98.74% macro F1 respectively) is consistent with the results reported by Yin and colleagues [15] on NSL-KDD (99.18% accuracy) and by Kim and colleagues [16] on UNSW-NB15, both of which established hybrid architectures as the state-of-the-art for network intrusion detection. However, the present study extends these findings in two significant ways: first, by providing per-class performance granularity that reveals the rare-category advantage of CNN-LSTM over Random Forest to be the primary source of the architecture's overall superiority; and second, by simultaneously documenting the adversarial fragility of the same model—a juxtaposition largely absent from the prior literature on these datasets. The finding that CNN-LSTM's advantage over Random Forest is most pronounced for rare attack categories (Δ F1 of +11.08 for Worms, +6.87 for Shellcode) is theoretically explicable: Random Forest's ensemble of decision trees constructs decision boundaries based on feature split thresholds, making it relatively insensitive to the fine-grained sequential feature patterns that distinguish rare attacks from normal traffic at the margin. The CNN-LSTM's capacity to model both local packet signatures and temporal flow sequences enables it to exploit distributional patterns in rare attack traffic that are invisible to feature-splitting approaches.

The computational cost findings (Table 4) add an important practical dimension that is frequently underemphasised in the academic literature. The predominant evaluation paradigm in ML cyber security

research optimises for accuracy metrics while treating computational cost as a secondary consideration—a prioritisation that may be appropriate for offline forensic analysis but is problematic for real-time IDS deployment. The 3.87 ms per-sample inference latency of the CNN-LSTM model, while acceptable for network segments with modest traffic volumes, becomes a bottleneck in high-throughput environments processing millions of packets per second. This finding is consistent with Milenkoski and colleagues' [21] observation that published IDS evaluations rarely address operational deployment constraints, and reinforces the need for the field to adopt multi-objective evaluation frameworks that explicitly trade off detection performance against computational feasibility. The SVM's relatively poor performance-to-cost ratio (96.33% accuracy, 28.4 minutes training, 0.87 ms inference) positions it unfavourably relative to both Random Forest and the deep learning models across all deployment scenarios—a finding that echoes the broader consensus in post-2018 ML literature on the declining relative utility of kernel SVMs for large-scale classification tasks.

The adversarial robustness analysis (Table 5) represents the study's most policy-relevant and practically urgent finding. The collapse of CNN-LSTM accuracy from 99.21% to 73.22% under FGSM perturbation at $\epsilon = 0.10$ —a perturbation magnitude that is imperceptible in network traffic feature space—constitutes a critical security vulnerability in production deployments. This magnitude of degradation significantly exceeds that reported by Carlini and Wagner [20] for image-based adversarial examples, suggesting that network traffic feature spaces may be particularly susceptible to gradient-based perturbation due to their continuous-valued, high-dimensional nature. The robustness advantage of Random Forest (1.94% degradation at $\epsilon = 0.10$) against neural architectures (24–27% degradation) validates the theoretical prediction from the adversarial ML literature that gradient-based attacks are fundamentally limited in their effectiveness against non-differentiable model architectures. This finding has direct and immediate implications for cyber security practitioners: organisations deploying ML-based IDS in adversarial environments—where sophisticated threat actors are capable of conducting model reconnaissance and crafting evasion inputs—should either employ adversarial training techniques or favour tree-based ensembles for their inherent gradient-robustness, accepting the modest accuracy trade-off this entails.

Comparing the present results with Srivastava and colleagues' [22] evaluation on Indian government network traffic, the divergence between benchmark dataset performance and real-world operational performance is a recurring and unresolved challenge. The present study's models, trained and evaluated on UNSW-NB15 and CIC-IDS-2017, represent the best-available approximation of real network traffic; however, domain shift between laboratory benchmarks and specific organisational network environments may reduce deployed performance by 5–15% depending on network topology, traffic composition, and the specificity of local threat actors. This underscores the importance of transfer learning approaches and continual model retraining as operationally essential complements to the architectural optimisation that dominates the academic literature. The implication for practitioners is that model selection based purely on published benchmark performance is insufficient; organisations should conduct pilot evaluations on representative samples of their own network traffic before committing to a specific architecture at scale.

VI. CONCLUSION

This study provides a comprehensive empirical comparison of six machine learning and deep learning architectures for cyber security threat detection and classification, evaluated on two benchmark datasets

(UNSW-NB15 and CIC-IDS-2017) across binary detection, multiclass classification, computational cost, and adversarial robustness dimensions. The CNN-LSTM hybrid achieved the highest clean detection performance (binary accuracy 99.21%, multiclass macro F1 98.74%), with its advantage over classical models concentrated in temporally distributed and rare attack categories where sequential modelling is most diagnostically valuable. However, the adversarial robustness analysis revealed a critical performance inversion: the same CNN-LSTM model suffered a 25.99 percentage point accuracy degradation under FGSM perturbation, while Random Forest degraded by only 1.94 percentage points—establishing that architectural superiority on clean benchmarks does not imply suitability for adversarial deployment environments. The study's three principal contributions are: (i) a standardised, reproducible multi-architecture comparison under controlled experimental conditions; (ii) the first simultaneous evaluation of detection performance and adversarial robustness on UNSW-NB15; and (iii) a quantified computational cost benchmark directly informing real-time deployment feasibility. Future research priorities include adversarial training for deep learning IDS models, transfer learning for domain adaptation, and the development of lightweight hybrid architectures suitable for edge computing and IoT security applications.

REFERENCES

- [1] A. Alshamrani, S. Myneni, A. Chowdhary, and D. Huang, "A survey on advanced persistent threats: Techniques, solutions, challenges, and research opportunities," *IEEE Communications Surveys & Tutorials*, vol. 21, no. 2, pp. 1851–1877, 2019.
- [2] Cybersecurity Ventures, "Cybercrime to Cost the World \$10.5 Trillion Annually by 2025," *Cybercrime Magazine*, 2020. [Online]. Available: <https://cybersecurityventures.com/cybercrime-damages-6-trillion-by-2021/>
- [3] R. Sommer and V. Paxson, "Outside the closed world: On using machine learning for network intrusion detection," in *Proc. IEEE Symp. Security and Privacy*, Oakland, CA, USA, 2010, pp. 305–316.
- [4] G. Louppe, "Understanding Random Forests: From Theory to Practice," Ph.D. dissertation, Univ. of Liège, Liège, Belgium, 2014.
- [5] Y. Mirsky, T. Doitshman, Y. Elovici, and A. Shabtai, "Kitsune: An ensemble of autoencoders for online network intrusion detection," in *Proc. Network and Distributed Systems Security (NDSS) Symp.*, San Diego, CA, USA, 2018.
- [6] Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *Nature*, vol. 521, no. 7553, pp. 436–444, 2015.
- [7] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*, MIT Press, Cambridge, MA, USA, 2016.
- [8] W. Wang, M. Zhu, X. Zeng, X. Ye, and Y. Sheng, "Malware traffic classification using convolutional neural network for representation learning," in *Proc. Int. Conf. Information Networking (ICOIN)*, Da Nang, Vietnam, 2017, pp. 712–717.
- [9] C. Szegedy et al., "Intriguing properties of neural networks," in *Proc. Int. Conf. Learning Representations (ICLR)*, Banff, AB, Canada, 2014.
- [10] M. Tavallaei, E. Bagheri, W. Lu, and A. A. Ghorbani, "A detailed analysis of the KDD CUP 99 data set," in *Proc. IEEE Symp. Computational Intelligence for Security and Defense Applications*, Ottawa, ON, Canada, 2009, pp. 1–6.

- [11] A. L. Buczak and E. Guven, "A survey of data mining and machine learning methods for cyber security intrusion detection," *IEEE Communications Surveys & Tutorials*, vol. 18, no. 2, pp. 1153–1176, 2016.
- [12] I. Sharafaldin, A. H. Lashkari, and A. A. Ghorbani, "Toward generating a new intrusion detection dataset and intrusion traffic characterization," in *Proc. 4th Int. Conf. Information Systems Security and Privacy (ICISSP)*, Funchal, Portugal, 2018, pp. 108–116.
- [13] M. Lotfollahi, M. J. Siavoshani, R. S. H. Zade, and M. Saberian, "Deep packet: A novel approach for encrypted traffic classification using deep learning," *Soft Computing*, vol. 24, no. 3, pp. 1999–2012, 2020.
- [14] R. C. Staudemeyer and E. R. Morris, "Understanding LSTM – A tutorial into long short-term memory recurrent neural networks," *arXiv preprint arXiv:1909.09586*, 2019.
- [15] C. Yin, Y. Zhu, J. Fei, and X. He, "A deep learning approach for intrusion detection using recurrent neural networks," *IEEE Access*, vol. 5, pp. 21954–21961, 2017.
- [16] J. Kim, J. Kim, H. L. T. Thu, and H. Kim, "Long short term memory recurrent neural network classifier for intrusion detection," in *Proc. Int. Conf. Platform Technology and Service*, Jeju, Korea, 2016, pp. 1–5.
- [17] S. Rezaei and X. Liu, "Deep learning for encrypted traffic classification: An overview," *IEEE Communications Magazine*, vol. 57, no. 5, pp. 76–81, 2019.
- [18] B. Biggio and F. Roli, "Wild patterns: Ten years after the rise of adversarial machine learning," *Pattern Recognition*, vol. 84, pp. 317–331, 2018.
- [19] I. J. Goodfellow, J. Shlens, and C. Szegedy, "Explaining and harnessing adversarial examples," in *Proc. Int. Conf. Learning Representations (ICLR)*, San Diego, CA, USA, 2015.
- [20] N. Carlini and D. Wagner, "Adversarial examples are not easily detected: Bypassing ten detection methods," in *Proc. ACM Workshop Artificial Intelligence and Security (AISec)*, Dallas, TX, USA, 2017, pp. 3–14.
- [21] A. Milenkoski, M. Vieira, S. Kounev, A. Avritzer, and B. D. Payne, "Evaluating computer intrusion detection systems: A survey of common practices," *ACM Computing Surveys*, vol. 48, no. 1, pp. 1–41, 2015.
- [22] S. Srivastava, A. Gupta, and N. Tyagi, "Machine learning-based intrusion detection for Indian government network environments," in *Proc. Int. Conf. Computing, Communication and Automation (ICCCA)*, Greater Noida, India, 2020, pp. 1–6.
- [23] A. Khraisat, I. Gondal, P. Vamplew, and J. Kamruzzaman, "Survey of intrusion detection systems: Techniques, datasets and challenges," *Cybersecurity*, vol. 2, no. 1, pp. 1–22, 2019.
- [24] N. Moustafa and J. Slay, "UNSW-NB15: A comprehensive data set for network intrusion detection systems," in *Proc. Military Communications and Information Systems Conf. (MilCIS)*, Canberra, ACT, Australia, 2015, pp. 1–6.
- [25] A. Gharib, I. Sharafaldin, A. H. Lashkari, and A. A. Ghorbani, "An evaluation framework for intrusion detection dataset," in *Proc. Int. Conf. Information Science and Security (ICISS)*, Pattaya, Thailand, 2016, pp. 1–6.

- [26] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic minority over-sampling technique," *Journal of Artificial Intelligence Research*, vol. 16, pp. 321–357, 2002.
- [27] M. A. Al-Garadi, A. Mohamed, A. K. Al-Ali, X. Du, I. Ali, and M. Guizani, "A survey of machine and deep learning methods for Internet of Things (IoT) security," *IEEE Communications Surveys & Tutorials*, vol. 22, no. 3, pp. 1646–1685, 2020.
- [28] R. Vinayakumar, M. Alazab, K. P. Soman, P. Poornachandran, A. Al-Nemrat, and S. Venkatraman, "Deep learning approach for intelligent intrusion detection system," *IEEE Access*, vol. 7, pp. 41525–41550, 2019.
- [29] T. A. Tang, L. Mhamdi, D. McLernon, S. A. R. Zaidi, and M. Ghogho, "Deep learning approach for network intrusion detection in software defined networking," in *Proc. Int. Conf. Wireless Networks and Mobile Communications (WINCOM)*, Fez, Morocco, 2016, pp. 258–263.
- [30] Z. Li, Z. Qin, K. Huang, X. Yang, and S. Ye, "Intrusion detection using convolutional neural networks for representation learning," in *Proc. Int. Conf. Neural Information Processing (ICONIP)*, Guangzhou, China, 2017, pp. 858–866.