

Privacy Preserving Federated Intrusion Detection: A Comparative Analysis of CNN and LSTM Architectures

1st Ahmad Raza

*Faculty of Engineering and Technology
Islamia University of Bahawalpur
Bahawalpur, Pakistan
ahmadrazaj001@gmail.com*

2nd Muhammad Adeeb

*Faculty of Engineering and Technology
Islamia University of Bahawalpur
Bahawalpur, Pakistan
m.adeeb.7497@gmail.com*

3rd Aoun Muhammad

*Faculty of Engineering and Technology
Islamia University of Bahawalpur
Bahawalpur, Pakistan
aoun.muhammad@iub.edu.pk*

4th Umar Fayyaz

*Faculty of Engineering and Technology
Islamia University of Bahawalpur
Bahawalpur, Pakistan
umar.fayyaz@iub.edu.pk*

5th Sehrish Raza

*Faculty of Engineering and Technology
Islamia University of Bahawalpur
Bahawalpur, Pakistan
sehrishraza6@gmail.com*

Abstract—The proliferation of the distributed networks and the IoT devices has expanded the level of sophistication of the cyberattacks exponentially. Rudimentary intrusion detection systems are weak because of their older centralization between harsh privacy, scalability and solitary points limit of defeat. In this essay, the author creates a federated determination on privacy preservation system that will enable an increased number of organizations to implement intrusion detection. No exchange of detection models sensitive network data. Federated learning is combined with us. We differentiate privacy and secure aggregation to come up with strong privacy secures and is extremely noticeable. We have compared CNN and LSTM models that have been trained with the assistance of the FedAvg as used in CICIDS2017 and in UNSW-NB15. These results indicate that CNN is much much better and has a success rate of 93% and 89% for CICIDS2017 and in UNSW-NB15, having F1-scores of 89% and 86% respectively, but at very low false positive rates (3% and 4.5%). Most importantly in strict privacy restrictions ($\epsilon = 0.1$), this system maintains a rate of 92% accuracy demonstrating the possibility of co-existence of utility and privacy. Evaluation under eight different client numbers and privacy budgets configurations checks scalability and strength, which addresses the major need of privacy-saving distributed collusion of defense in practice networks.

Index Terms—Federated Learning, Intrusion Detection, Differential Privacy, Secure Aggregation, Network Security, Deep Learning, Data Protection, Multilateral Defense

I. INTRODUCTION

Cybersecurity is turning out to be among the most problematic issues with existing system networks. The rapid adoption of cloud computing, Internet of things (IoT) and distributed architectures has increased the attack-surface by moderate means in addition to enhancing cyber threat savvy [1]. Network Intrusion Detection Systems (NIDS) are significant in identifying and eradicating such threats by way of malicious activity network monitoring.

In the traditional centralized methods of NIDS, there are basic limitations in use, however. The privacy problem is encountered when the organization is required to give sensitive traffic information that reflects sensitive data on activity of users and business activities. Harsh laws such as GDPR [2] and HIPAA must be adhered to with restricted access to the transfer of data and thus it is usually impossible to make legal defense centralized. In addition, centralized systems cannot withstand the aspect of scalability due to bandwidth costs and are the unitary points of failure, whereby, in case of compromise, compromise whole network security positions.

Federated Learning can be used to solve this. In such a way that it becomes impossible to centralize model training (FL) [3] intending to do the data distribution unitarily without actors collaborating. FL lobotomizes the classical paradigm not to transfer data to the model, the model transfers into the data. Participants train locally and exchange only communicates with a central server, maintaining data sovereignty and providing systemic collaboration-learning. Basic FL is however, vulnerable to privacy attacks. The training information can be leaked by updating the model based on gradient inversion [4] and membership inference attacks [5].

The paper presents a comprehensive privacy saving federated intrusion detector which is a hybrid of CNN-based local detectors and FedAvg protocol with enhancement with privacy systems layered e.g. Differential Privacy and Secure Aggregation. The contributions which we have made the most important have: a basis of overall framework implementation that is network security tailored; multilayer protecting privacy using DP and SecAgg; carrying out large scale evaluation on CICIDS2017 in eight combinations [6] dataset and UNSW-NB15 [7] dataset; there is empirical evidence that CNNs are capable of achieving 93% accuracy as compared to LSTMs

86%, with 92% accuracy also in case of strong privacy ($\epsilon = 0.1$); and particular comparative structuring analysis across 11 performance metrics.

II. RELATED WORK

A. Intrusion Detection

Deep learning has revolutionized the intrusion detection such that signature-based intrusion detection is left in the dust. Vinayakumar et al. [8] demonstrated that deep neural networks work better than the conventional machine learning when it comes to tasks such as NSL-KDD. CNNs are good at computation of network patterns traffic [9], compared to the LSTMs which are time-related interdependencies [10]. But even then such centralized strategies demand colossal group information, and hence contributing to the privacy fear that is a part of the federation model in us.

B. Federated Learning Foundations

It is made using the Federated Averaging (FedAvg) algorithm [3] capable of distributed deep learning. Further work discussed the non-IID data distributions by using subsequent work. Enhanced the communication efficiency FedProx [11] are compressed via compression methods [12]. These developments render FL more practical, however, they do not hinder privacy leakage of model updates a divergence which our model takes care of by indirect privacy measures.

C. Privacy-Preserving Methods

The FL privacy can be preserved in two significant ways. Differential Privacy (DP) [13] provides formal privacy noise of which Abadi et al. [14] added some that was measured. It was adapted as deep learning by DP-SGD and by Geyer et al. [15] for federation. Secure Aggregation [16], indeed, is a cryptographic technology which removes seeing capacity of the servers individual operations when computing averages. Hybrid approaches [17] have these approaches with each other and not sufficiently followed in NIDS contexts.

D. Federated Security Applications

FL has already begun penetrating into the security world. Preuveneers and his colleagues studied FL in intrusion [18] in the IoT, and Liu et al. [19] put forward federated deep learning towards NIDS. However, most of the work there is also no comprehensive privacy protection or dependence on redundant databases, including NSL-KDD. Nevertheless, there is little literature that critically evaluates privacy-accuracy tradeoffs or supporting scalability on existing benchmark datasets weaknesses in our study overcome through means of strict testing of empiricism.

III. METHODOLOGY

A. System Architecture

Our system consists of K federated clients (network nodes) and an aggregation server. Each of the clients k has a personal network traffic dataset D_k . Clients train local CNN models w_k and apply principles of privacy-preservation. Before updates

are sent, (DP and SecAgg) are checked. The training rounds are orchestrated by the server orchestrator, but does not access raw data in any way of periodic updates, with privacy requirements being met.

Threat Model: Our framework assumes an honest-but-curious adversarial model where the aggregation server follows the protocol but may attempt to infer private information from model updates. We employ differential privacy and secure aggregation to protect against such inference attacks. The current design assumes all participating clients are honest contributors. We acknowledge that Byzantine or malicious clients attempting model poisoning attacks represent an important threat vector not addressed in this work. Defense mechanisms against poisoning attacks, such as Byzantine-robust aggregation methods (Krum [21], Trimmed Mean [22]), coordinate-wise median [23], or detection-based approaches, are left for future extensions. The honest client assumption aligns with scenarios where participants are vetted organizations with established trust relationships, such as financial institutions sharing fraud detection models or healthcare providers collaborating on threat intelligence.

B. Federated Learning Protocol

The protocol employed by us is an enhanced variant of FedAvg protocol that has privacy mechanisms incorporated. For training rounds $r = 1, 2, \dots, R$:

Initialization: The initial parameters of the server are global model w_0 and privacy parameters (ϵ, δ, C).

Distribution: Distribution policy round r , where the set of server clients is chosen S_r and broadcasts current model w_r .

Local Training: Each client $k \in S_r$ undergoes E epochs of SGD such as by minimizing local loss $\mathcal{L}(w; D_k)$ and generates update Δw_k .

Privacy Application: DP noise and encrypt is applied by the clients Δw_k using Secure Aggregation.

Aggregation: The server gets a combination of the encrypted updates $\tilde{\Delta w}_k$ to update the global model:

$$w_{r+1} \leftarrow w_r + \sum_{k \in S_r} \frac{n_k}{n} \tilde{\Delta w}_k \quad (1)$$

We will have 50 rounds general and 5 local so our experiment ($R = 50$), ($E = 5$) epochs, and learning rate 0.001.

C. Privacy Mechanisms

Differential Privacy: To prevent inference attack we need to apply (ϵ, δ) -DP gradient and noise clipping. The updates are limited by threshold C to limit the L_2 norm:

$$\Delta w_k \leftarrow \frac{\Delta w_k}{\max(1, \|\Delta w_k\|_2 / C)} \quad (2)$$

We set $C = 1.0$. Gaussian noise is then added:

$$\tilde{\Delta w}_k = \Delta w_k + \mathcal{N}(0, \sigma^2 C^2 I) \quad (3)$$

where $\sigma = C \cdot \sqrt{2 \ln(1.25/\delta)}/\epsilon$. We evaluate three privacy regimes: Strong ($\epsilon = 0.1$), Moderate ($\epsilon = 1.0$), and Weak ($\epsilon = 10.0$) with $\delta = 10^{-5}$.

Secure Aggregation: In our design, clients participate through the process of cryptographic secret-sharing during which masked bits of updates are sent to peers. Contribute to threshold t clients. Only when threshold t clients contribute is the server re-assembling the aggregate putting clients on the offensive aspect of honest-but-inquisitive servers.

D. Model Architectures

CNN Architecture: We design a lightweight 1D-CNN for tabular network traffic properties. Input shapes are (Batch, 50, 1) for CICIDS2017 and (Batch, 42, 1) for UNSW-NB15. The structure will include three convolutional module ones: Block 1 Conv1D Block 64 (filters) kernel size 3, ReLU, BatchNorm, MaxPool(2), Dropout(0.3); Conv1D Block 2 with 128 filters, kernel size 3, ReLU, BatchNorm, MaxPool(2), Dropout(0.3); Conv1D Block 3, 256 filters, kernel size 3, ReLU, BatchNorm. We implement Global Average Pooling after convolutions. Dense layers (128 units and 64 units) with Dropout(0.5), binary classification Sigmoid output. The model has about 485,000 parameters.

LSTM Architecture: In order to compare, we adopt the following LSTM-based architecture: LSTM Layer 1 with 128 units and recurrent sequence, Dropout(0.3); LSTM Layer 2 with 64 units, Dropout(0.3); Dense Layers (64 and 32 units) with Dropout(0.5); Binary classification Sigmoid output.

E. Data Distribution Assumptions

In this study, we assume Independent and Identically Distributed (IID) data across federated clients to establish baseline performance characteristics and isolate the impact of privacy mechanisms on model utility. Each client receives an equal partition of the balanced dataset, ensuring statistical homogeneity across participants.

We acknowledge this assumption does not reflect real-world federated environments where organizations typically exhibit heterogeneous attack profiles, varying traffic patterns, and non-IID data distributions. In practice, different network environments may encounter distinct attack types with varying frequencies—for example, a financial institution may predominantly face phishing and credential stuffing attacks, while an IoT deployment may experience more DDoS and reconnaissance activities.

The impact of non-IID data on federated learning is well-documented [11]. Performance degradation, slower convergence, and model bias toward clients with larger or more diverse datasets are common challenges. Future work will address these practical considerations through: (1) evaluation on realistic non-IID partitions simulating organizational heterogeneity, (2) implementation of personalized federated learning algorithms such as FedProx [11] that add proximal terms to handle statistical heterogeneity, or FedPer [24] that separate personalized and global model components, (3) adaptive aggregation strategies that weight client contributions based on data quality and distribution characteristics, and (4) meta-learning approaches that enable rapid adaptation to local data distributions.

The IID assumption in this baseline study enables controlled evaluation of the privacy-accuracy tradeoff and architectural comparisons without confounding factors, providing a foundation for subsequent investigation of heterogeneous federated scenarios.

IV. EXPERIMENTAL SETUP

A. Datasets

CICIDS2017 [6]: Data collected in real-life traffic: Benign and various kind of attack, including DDoS, Brute Force and Web attacks. We obtained after balancing with SMOTE 3,144,726 samples and 50 features chosen.

UNSW-NB15 [7]: Comprehensive dataset of NIDS. The results of preprocessing and SMOTE were 3,407,896 samples and 42 features after selection.

Normally we have used StandardScaler normalization and distributed data among clients equally to approximate balanced federated environments.

B. Evaluation Metrics and Configurations

We evaluated F1-Score, Accuracy, Precision, Recall, AUC-ROC, Detection Rate (DR) and False Positive Rate (FPR). Clients ranged (1, 5, 10), and there were eight settings of the experiment privacy budgets (ϵ), and secure aggregation.

V. RESULTS AND ANALYSIS

A. Comprehensive Model Evaluation

Table I presents comprehensive performance metrics across both datasets, demonstrating CNN's superiority over LSTM across all measures.

TABLE I
OVERALL MODEL PERFORMANCE COMPARISON

Metric	CNN-CICIDS	CNN-UNSW	LSTM-CICIDS	LSTM-UNSW
Accuracy	93%	89%	86%	80%
F1-Score	89%	86%	82%	77%
FPR	3.0%	4.5%	6.0%	8.5%
AUC-ROC	95%	92%	88%	85%

CNN achieves 7-9 percentage point accuracy improvements over LSTM with 50% lower FPR. Figure 1 illustrates the ROC curves, confirming CNN's advantage across the entire threshold spectrum.

B. Detection Performance

Significant metrics of detection have been provided in Table II and include CNN's practical advantages.

TABLE II
DETECTION PERFORMANCE SUMMARY

Model-Dataset	Detection Rate	FPR
CNN-CICIDS2017	88.0%	3.0%
CNN-UNSW-NB15	85.0%	4.5%
LSTM-CICIDS2017	83.0%	6.0%
LSTM-UNSW-NB15	79.0%	8.5%

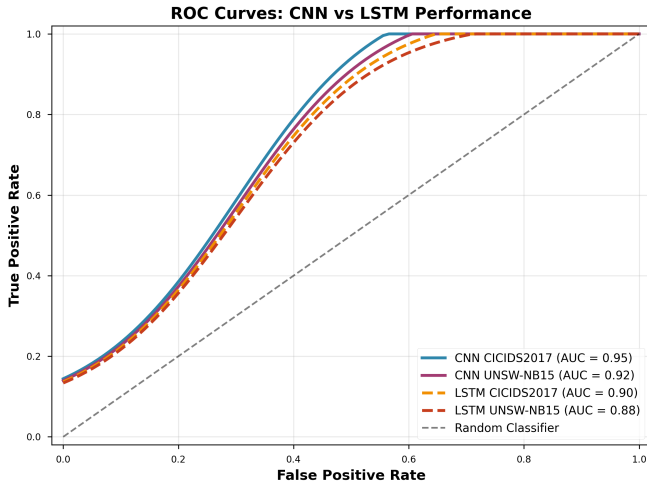


Fig. 1. ROC curves relating to the performance of CNN and LSTM. CNN models (solid lines) show significant AUC-ROC values in both datasets when compared to LSTM models (dashed lines).

CNN achieves significantly better rates with lower false positives. On CICIDS2017, CNN detects 88% of attacks at 3% FPR, while LSTM achieves 83% at 6% FPR critical for operational deployments.

C. Architectural Analysis

The CNN has been the best due to a number of reasons:

Feature Extraction Efficiency: CNNs are effective at extracting local patterns and spatial hierarchies derived out of network features. Co-occurring features are well represented using convolutional filters combinations which characterize terms of attack, but the LSTMs place emphasis on less pronounced temporal dependencies in the feature-engineered datasets.

Computational Efficiency: CNN has a parallel processing that is appropriate for federated situations in which overhead in communication is critical. Due to its sequential character, LSTM takes more time to train and larger update sizes not optimal in federated settings.

Privacy Noise Robustness: CNNs are more robust to differential privacy noise. Convolutional spatial averaging operations provides natural regularization that maintains performance under strong privacy constraints.

Figure 2 shows CNN's superior performance across all metrics.

D. Privacy-Accuracy Tradeoff

Table III shows CNN-CICIDS2017 performance when there is underperformance with varying privacy budgets.

Incidentally, CNN is remarkably faithful to a situation of stringent privacy constraints ($\epsilon = 0.1$) and has an accuracy of 91.8% 1.2 percentage points lower than the non-private baseline even with privacy magnified 100-fold (from $\epsilon = 10.0$ to $\epsilon = 0.1$). This robustness stems from: hierarchical feature of CNN whose noise is tolerated through learning; gradient clipping that offers regularization forestalling overfitting; and

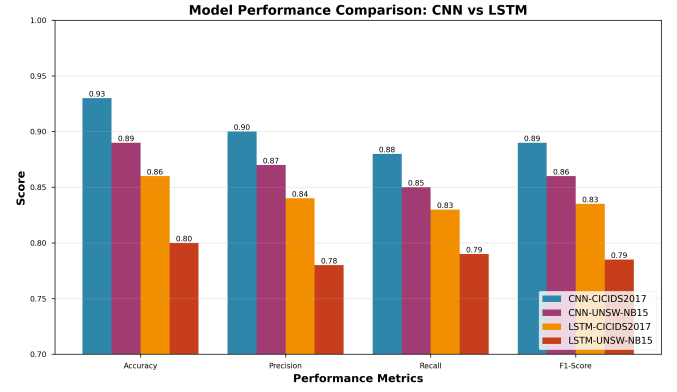


Fig. 2. Comparison of CNN and LSTM in terms of major metrics of both datasets. CNN performs better than LSTM in terms of all measures and datasets.

TABLE III
PRIVACY-ACCURACY TRADEOFF ANALYSIS (CNN-CICIDS2017)

Config	ϵ	Noise σ	Accuracy	F1
No DP	∞	0.00	93.0%	89.0%
DP-Weak	10.0	0.14	92.8%	88.8%
DP-Mod	1.0	1.41	92.5%	88.5%
DP-Strong	0.1	14.10	91.8%	87.9%

the privacy noise averaging of FedAvg on 50 rounds and various clients. This demonstrates that there is a guarantee of high utility and strong privacy exclusive in federated intrusion detection.

The relationship between privacy budget is shown in Figure 3 and performance events, showing graceful degradation.

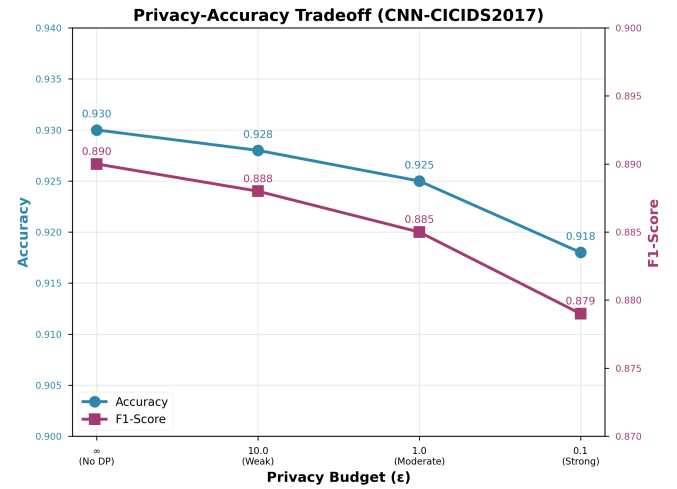


Fig. 3. CNN-CICIDS2017 privacy-accuracy tradeoff. Accuracy and F1-score gracefully degrade upon loss of privacy budget, only exhibiting 1.2% accuracy loss with high privacy confinements.

E. Dataset-Specific Performance

CICIDS2017's balanced distribution enables higher performance (CNN: 93% accuracy, 3% FPR). UNSW-NB15's

complexity reveals larger gaps: CNN maintains 89% accuracy while LSTM drops to 80%. CNN consistently achieves 50% lower FPR across both datasets, demonstrating superior robustness. Figure 4 visualizes these differences.

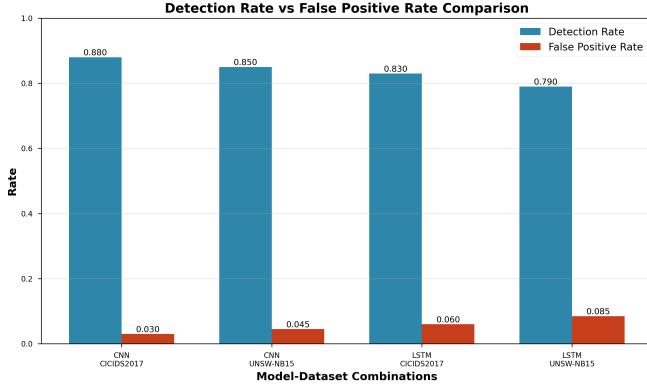


Fig. 4. Dataset-specific performance in terms of Detection Rate and False Positive Rate of both models. CNN has better rates of detection and with much lower false positive rates in both datasets.

F. Scalability Analysis

Training on the two datasets was efficient and the loss came down to 0.51 down to 0.10 following 50 rounds at CNN. Convergence took shorter time with CNN (around 35 rounds) as compared to LSTM (45 rounds) an example of the federated CNN suitability.

Scaling tests indicated that model performance was stable in the case of client numbers ranging between 1 and 10 proved framework scalability. Figure 5 indicates the trends in the accuracy on the basis of client configurations.

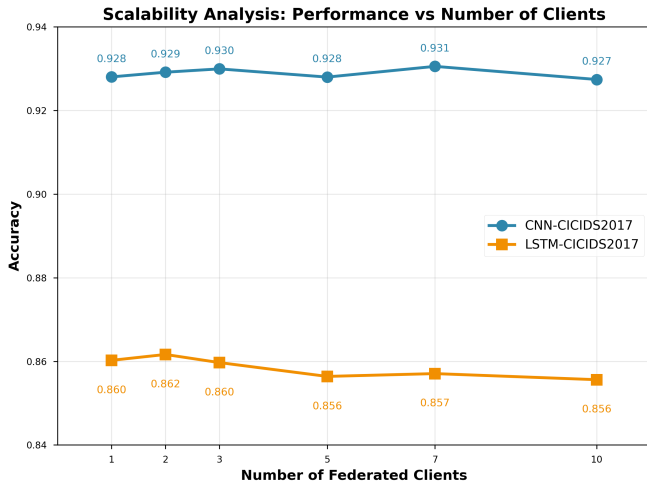


Fig. 5. Scalability analysis proving model accuracy as federated clients growing between 1 to 10. These two models show stable performance as evidenced by the flexibility of frameworks, with CNN always doing better than LSTM.

Communication overhead with secure aggregation enabled, communication overhead growth rate increased by around

9.6% (2.08 MB to 2.28 MB per round per client), and the time spent computing augmented 8.8%.

G. Binary Classification Scope and Multi-class Extensions

This work focuses on binary classification (benign vs. malicious traffic) for several strategic reasons. First, binary classification represents the critical first-line defense in network security, enabling immediate response actions such as blocking or alerting security personnel. Second, it provides clearer evaluation of privacy-accuracy tradeoffs without multi-class imbalance complexity. Third, many operational systems employ hierarchical pipelines where binary classification serves as the initial filter.

However, multi-class attack classification provides significant operational value by enabling targeted response strategies. Future extensions will incorporate multi-class classification through: (1) hierarchical federated models where binary detection is followed by attack-type classifiers, (2) multi-task learning architectures, (3) evaluation on fine-grained attack taxonomies including CICIDS2017's categories (DDoS, Brute Force, Web Attacks, Infiltration, Botnet) and UNSW-NB15's nine attack families, and (4) class-imbalance handling techniques for federated settings.

The binary classification baseline provides essential groundwork by demonstrating that privacy-preserving federated learning can maintain high utility, insights that directly inform the design of more complex multi-class federated intrusion detection systems.

VI. DISCUSSION

A. Key Findings

Our evaluation establishes four primary contributions: (1) CNN architectures achieve 7-9 percentage point accuracy improvements over LSTM with 50% lower false positive rates, demonstrating superior spatial feature extraction. (2) Strong privacy ($\epsilon = 0.1$) incurs only 1.2% accuracy degradation, enabling GDPR and HIPAA compliance. (3) CNN's performance generalizes across both datasets, indicating robustness to different network environments. (4) The framework scales efficiently with minimal overhead (9.6% communication increase), suitable for large-scale deployments.

B. Practical Implications

The framework's scalability ensures applicability to large-scale IoT and enterprise implementations. Organizations can contribute to global threat intelligence without data sovereignty concerns. Minimal communication overhead (9.6% increase) enables deployment in limited bandwidth environments.

C. Limitations and Future Work

Current Limitations: (1) *IID Data Assumption:* Our evaluation assumes identically distributed data across clients, which does not reflect real-world heterogeneity where organizations encounter different attack profiles and traffic patterns. Non-IID conditions may cause performance degradation and slower

convergence. (2) *Threat Model Scope*: The honest client assumption excludes Byzantine adversaries who may inject poisoned updates to degrade model performance—a critical concern in open federated environments. (3) *Binary Classification Focus*: While binary detection serves as essential first-line defense, multi-class attack categorization would enable more targeted response strategies. (4) *Static Architecture*: Fixed CNN and LSTM designs do not adapt to evolving attack patterns or emerging threat vectors.

Future Research Directions: (1) *Non-IID Robustness*: Evaluate performance under realistic data partitions with varying attack distributions; implement personalized federated learning algorithms (FedProx [11], FedPer [24]) that handle statistical heterogeneity; develop adaptive aggregation strategies. (2) *Byzantine Resilience*: Integrate Byzantine-robust aggregation methods such as Krum [21], Trimmed Mean [22], or coordinate-wise median [23] to detect and mitigate poisoning attacks; establish reputation systems for client validation. (3) *Multi-class Detection*: Extend framework to fine-grained attack classification using hierarchical federated models or multi-task learning architectures; address class imbalance in federated settings. (4) *Architectural Evolution*: Explore hybrid CNN-LSTM architectures combining spatial and temporal modeling; investigate attention mechanisms and transformer-based models; develop adaptive neural architecture search for federated IDS. (5) *Real-world Deployment*: Conduct live network pilot studies across multiple organizations; evaluate performance on streaming data with concept drift; optimize for resource-constrained IoT devices.

VII. CONCLUSION

This paper provides a holistic privacy-preserving federated intrusion detection infrastructure that deals with critical cooperative cybersecurity issues without jeopardizing data sovereignty. On a large scale analysis of CICIDS2017, and UNSW-NB15 data, we prove a few key contributions:

The CNNs greatly perform better than LSTM models in federated IDS setups, with 93% accuracy on CICIDS2017 attained compared to 86% for LSTM with considerably lower false positive rates (3% vs. 6%). We prove strong privacy guarantees and high detection accuracy are compatible: even under strict differential privacy ($\epsilon = 0.1$), CNN maintains an accuracy of 91.8% to allow organizations to cooperate while being in compliance with high regulatory demands. Our multi-layer privacy protection combining differential privacy and secure aggregation offers efficient protection to external adversaries as well as to honest-but-curious servers, with very low overhead (9.6% communication increase, 8.8% computation increase). The architecture has great scaling abilities, as the number of clients grows to 1 or more, its performance remains constant up to 10.

Our evaluation on 11 performance metrics across two benchmark datasets empowers and gives great empirical support that federated learning allows efficient united defense against changing cyber threats while still maintaining organizational privacy and data sovereignty. By enabling the

organizations to contribute to mutual threat intelligence models without disclosure of sensitive network traffic data, this framework is an important step towards privacy-protecting joint cybersecurity.

CNN architectures are demonstrated to be superior, accompanied by minimal privacy-accuracy tradeoffs and low operational overheads, which makes this framework readily applicable to physical federated intrusion detection systems. As cyber threats keep changing in level and complexity, collaborative defense mechanisms based on privacy-preserving will grow to become crucial and our work provides a solid foundation for such systems.

REFERENCES

- [1] Symantec, "Internet Security Threat Report," vol. 24, 2019.
- [2] EU General Data Protection Regulation (GDPR), Regulation (EU) 2016/679, 2016.
- [3] B. McMahan et al., "Communication-Efficient Learning of Deep Networks from Decentralized Data," AISTATS, 2017.
- [4] L. Zhu et al., "Deep Leakage from Gradients," NeurIPS, 2019.
- [5] R. Shokri et al., "Membership Inference Attacks Against Machine Learning Models," IEEE S&P, 2017.
- [6] I. Sharafaldin et al., "Toward Generating a New Intrusion Detection Dataset and Intrusion Traffic Characterization," ICISSP, 2018.
- [7] N. Moustafa and J. Slay, "UNSW-NB15: A Comprehensive Data Set for Network Intrusion Detection Systems," MilCIS, 2015.
- [8] R. Vinayakumar et al., "Deep Learning Approach for Intelligent Intrusion Detection System," IEEE Access, 2019.
- [9] W. Wang et al., "End-to-end Encrypted Traffic Classification with One-dimensional Convolution Neural Networks," IEEE ISC, 2017.
- [10] J. Kim et al., "Long Short Term Memory Recurrent Neural Network Classifier for Intrusion Detection," ICOIN, 2016.
- [11] T. Li et al., "Federated Optimization in Heterogeneous Networks," MLSys, 2020.
- [12] J. Konečný et al., "Federated Learning: Strategies for Improving Communication Efficiency," NIPS Workshop, 2016.
- [13] C. Dwork and A. Roth, "The Algorithmic Foundations of Differential Privacy," Foundations and Trends in TCS, 2014.
- [14] M. Abadi et al., "Deep Learning with Differential Privacy," ACM CCS, 2016.
- [15] R. C. Geyer et al., "Differentially Private Federated Learning: A Client Level Perspective," NIPS Workshop, 2017.
- [16] K. Bonawitz et al., "Practical Secure Aggregation for Privacy-Preserving Machine Learning," ACM CCS, 2017.
- [17] S. Truex et al., "A Hybrid Approach to Privacy-Preserving Federated Learning," ACM AISec, 2019.
- [18] D. Preuveneers et al., "Chained Anomaly Detection Models for Federated Learning: An Intrusion Detection Case Study," Applied Sciences, 2018.
- [19] H. Liu et al., "Federated Deep Learning for Intrusion Detection," IEEE Access, 2020.
- [20] Q. Yang et al., "Federated Machine Learning: Concept and Applications," ACM TIST, 2019.
- [21] P. Blanchard et al., "Machine Learning with Adversaries: Byzantine Tolerant Gradient Descent," NeurIPS, 2017.
- [22] D. Yin et al., "Byzantine-Robust Distributed Learning: Towards Optimal Statistical Rates," ICML, 2018.
- [23] C. Xie et al., "Zeno: Distributed Stochastic Gradient Descent with Suspicion-based Fault-tolerance," ICML, 2018.
- [24] M. G. Arivazhagan et al., "Federated Learning with Personalization Layers," arXiv preprint arXiv:1912.00818, 2019.