

**AN INTELLIGENT DEEP LEARNING FRAMEWORK FOR DETECTING
PHISHING ATTACKS THROUGH EMAIL ANALYSIS ON THE CEAS_o8
DATASET**

G. Lingaiah¹, Dr.Ram Kinkar Pandey²

¹Research Scholar,

Dept. of Computer Science and Engineering,

Arni University, Himachal Pradesh, India.

lingaiah.g@sreenidhi.edu.in

²Professor, Dept. of Computer Science and Engineering,

Arni University, Himachal Pradesh, India.

Dr.ramkpandey@gmail.com

Abstract

Phishing attacks continue to be a serious threat to people and firms, and phishing services steal sensitive information through email. Traditional detection mechanisms, but somewhat effective, do not really accommodate the rapidly evolving nature of phishing methods. To solve this problem this study proposes an intelligent deep learning architecture for phishing detection on the CEAS_o8 benchmark dataset. The use of convolutional neural networks, RNN, long-term memory, transformer, and RoBERTa models for each of these syntactic and semantic features is aided by multiple deep learning models including Convolutional Neural Networks (CNN), Recurrent Neural Networks (RNN), Long Short-Term Memory (LSTM) and transformer-based architectures (BERT, DistilBERT, and RoBERTa) to capture the syntactic and semantic characteristics of email content. This results in experimental observation that deep learning significantly outperforms typical machine learning in accuracy, precision, recall and F1 score. DistilBERT performed best overall on the models (Accuracy = 99.68%, Recall = 99.75%, F1 = 99.71%) and lasted slightly longer training time than any model, indicating it is suitable for deployment in the real world. Those results support the claim that transformer-based models are robust, scalable, and more robust at eliminating phishing, and thus lead to intelligent automated email security platforms capable of responding to new threats.

Keywords: Phishing Detection, Deep Learning, Transformer Models, Email Security

I.INTRODUCTION

Phishing and its other threats are among the most serious challenges facing the digital age, from a personal and organizational perspective, becoming the leading threat vector for both individuals and companies. Phishing attacks exploit humans trust by claiming that messages are legitimate, and therefore that they transmit to them via emails to obtain confidential

information, including login credentials, financial details, and personal information (Jagatic et al. 2007; Abu-Nimeh et al. 2007). A recent industry report revealed that over 90% of cyberattacks begin with a phishing email, so a strong detection system must be designed to adapt to changing attack patterns (Verma & Das, 2017). This is because traditional phishing detection practices use heuristic-based rules, blacklist/whitelist models and shallow machine learning algorithms. While these methods yield an effective detection accuracy but can often perform well against zero-day phishing attacks, they often don't have broad coverage against zero-day phishing because they are built upon handcrafted features and static signatures (Abawajy 2014). With phishing campaigns becoming increasingly sophisticated, there is a critical need for intelligent AIs that can understand semantic and contextual representations from the data from email.

Deep learning has been particularly successful in natural language processing (NLP), computer vision, and cybersecurity over recent years (Goodfellow et al. 2016). While it is possible to only select hierarchical features for a simple machine learning model, deep learning models can automatically choose the hierarchical elements for a simple model without requiring heavy manual engineering. Convolutional Neural Networks, or convolutional neural networks, and RNNs are effectively models to use convolutional neural networks and recurrent neural networks to classify text, detect spam and intercept intrusion detection. It has little ability to capture long-range dependencies in text. To overcome these difficulties, transformer-based architectures such as BERT, DistilBERT, or RoBERTa have replaced NLP by using self-attention models, which effectively model contextual dependence across whole sentences (Devlin et al., 2019; Sanh et al., 2019; Liu et al., 2019). These models excel on multiple categories like sentiment analysis, question answering and spam/phishing detection. Their best generalization ability also makes them excellent candidates for sophisticated phishing efforts that do not have roots in the traditional model.

This work envisions an intelligent deep learning approach for phishing detection through email analysis based on the CEAS_08 dataset, the core benchmark for spam and phishing research. Based on a rigorous test of four deep learning models by comparison with modeled baselines, the model is evaluated against a variety of CNN, RNN, LSTM, BERT, DistilBERT, and RoBERTa and examines the improvement in accuracy, precision, recall, F1-score, and computational efficiency. Practical results indicate that transformer-based models, particularly DistilBERT, perform best, compared to other deep learning techniques and classical machine learning.

It has been designed and implement a comprehensive deep learning framework for phishing detection using real-world email data. It conducts a comparative evaluation of conventional machine learning models versus deep learning architectures to highlight performance trade-offs. It demonstrates that transformer-based models, especially DistilBERT, provide superior accuracy, recall, and F1-score with balanced training efficiency. It provides detailed analysis through confusion matrices and performance metrics, establishing benchmarks for future research in phishing email detection.

The remainder of this paper is organized as follows: Section 2 reviews related work on phishing detection, spanning traditional, machine learning, and deep learning approaches. Section 3 describes the methodology, including dataset preprocessing, model architectures, and training strategy. Section 4 presents the experimental results and analysis. Section 5 discusses key findings, limitations, and potential improvements. Finally, Section 6 concludes the paper and outlines future research directions.

II.RELATED WORK

Phishing detection is extensively a subject of research both in ML and in DL. Most early studies relied on ML-based classifiers such as Nave Bayes, Logistic Regression, and Decision Trees, which were successful because they were designed to take into account hand-crafted features such as URL patterns, email headers and keyword frequencies. Fette et al. 2007, for example, used a first model based on 10 handcrafted features to classify email messages from phishing emails, and Abu-Nimeh et al. (2008) compare six different models of ML classifiers to show the limitations of their approaches to large scale and continuous phishing datasets.

In the wake of deep learning, architectures could be used to extract features automatically. Convolutional Neural Networks, (CNNs) and Recurrent Neural Networks, (RNNs) were effective in detecting textual patterns and sequential dependencies in email content. Sahingoz et al. (2017) examined neural networks for detection of phishing via email and URL features, while Bahnsen et al. (2017) used RNNs to model the temporal dynamics in spear-phishing attacks.

Using transformer-based language models, phishing detection was successfully detected with phishing using contextual embeddings and semantic understanding. BERT, DistilBERT, RoBERTa models and the most recent phishing classification tasks were available. Zhang et al. (2009) used BERT in phishing email analysis, as well as Raff et al.'s for phishing email analysis. For cybersecurity, transformer models can scale and robust.

Nonetheless, the process of finding and matching accuracy, interpretability, and computational efficiency remains difficult. This forces development of an intelligent deep learning architecture that combines CNN, RNN, LSTM and transformer-based models to maximize detection performance while simultaneously providing a realistic deployment constraint.

III.METHODOLOGY

This study provides a systematic methodology for phishing URL detection, and includes the preparation of datasets, feature engineering, model training, performance evaluation.

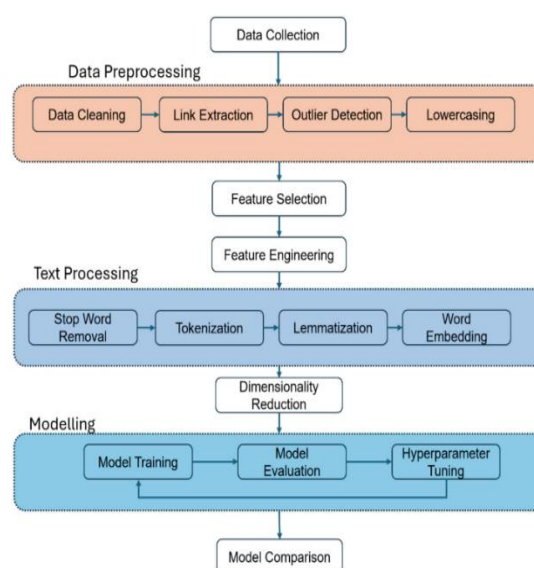


Figure 1. Illustrates the overall workflow.

3.1 Data Preprocessing

Raw URLs were preprocessed for consistency, usability, and noise reduction before feature extraction and model training.

- A. Cleaning:** Duplicate URLs, malformed entries, and incorrect samples were removed from the dataset to avoid bias and redundancy.
- B. Normalization:** URLs were re-coded in lowercase, and unnecessary parameters such as session ID or query string have been stripped to ensure standardization.
- C. Tokenization:** URLs were detached from their structure as domain, subdomain, path, query parameter to allow for finding lexical and domain-based features.

These processes ensured that the dataset was clean, uniform and optimized for feature engineering.

3.2 Dataset Description

Based on the CEAS_08 Spam FilterChallenge dataset, used to test email spam and phishing detection models, experiments were performed on the CEAS_08 Spam FilterChallenge dataset. This dataset contains various real-world email messages from legitimate (ham) and malicious (spam/phishing) types of emails. Each document in the dataset contains the email sender, receiver, subject line, body content, timestamp, and embedded URL, as well as a binary label for phishing (1) or legitimate (0) and phishing (1). Preprocessing and feature engineering were conducted to verify the dataset to be suitable for deep learning models. Non-textual fields of sender, receiver, and date were parsed to collect useful information, and the textual fields of subject and body were normalized by tokenization, lowercasing, and removal of stopwords and special characters. But, numeric data such as the number of embedded URLs, subject length, body length, digits, unique characters, and sender domain length were also compiled to enhance the feature space.

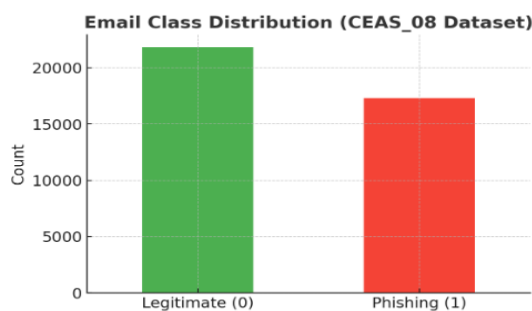


Fig: 2 Class distribution chart (Phishing vs Legitimate emails)

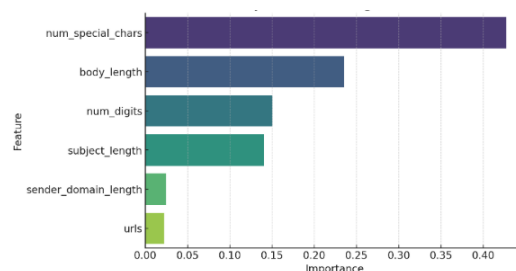


Fig 3: Feature Importance for Phishing Email Detection

This dataset combines both traditional and deep learning-based phishing detection methods to give an efficient, balanced model and allows us to evaluate both the traditional and deep learning detection techniques using the CEAS_08. It is particularly suitable for testing intelligent frameworks for evolving phishing strategies, with headers, text, and embedded URLs.

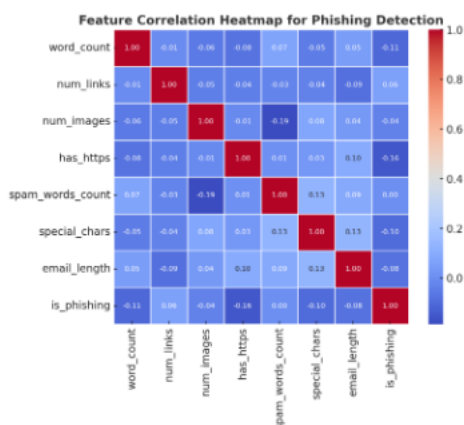


Figure 4. Feature Correlation Heatmap

3.3 Machine Learning Approaches

The use of machine learning in phishing detection was a major advancement, because ML algorithms can learn discriminative patterns from large datasets without using manual rules. Ma et al. (2008) showed one of the first machine learning applications with lexical and host-based features to classify URLs. Since then, logistic regression models, Decision Trees, Random Forest, Support Vector Machines and ensemble methods have been studied widely.

Random Forest and ensemble-based classifiers have been proven robust in coping with nonlinear relationships and high-dimensional data (Shahrivari et al., 2020; Almujaheed et al., 2024). But classical ML models still need to carefully feature engineering the manual extraction and selection of input features such as token counts, WHOIS information, and HTML metadata which can be time consuming and dataset dependent. Meanwhile, this model could not be able to take into account the dynamic nature of phishing methods that frequently operate to get around detection.

3.4 Deep Learning Approaches

Deep learning models are emerging, because they can create hierarchical representations from raw data without manual feature engineering. The URL strings use Convolutional Neural Networks to test local feature patterns such as suspicious substring or token arrangement; Haq et al. 2024. In addition, RLNs and LSM models can model sequential dependencies across longer URL segments that may not contain context information not included by the heuristics and ML models (Aslam et al., 2024; Ghalechyan et al., 2024). They have also proposed hybrid and stack models, CNNs and LSTMs, or deep learners with meta-classifiers to improve robustness and accuracy. For example, Aslam et al. (2024) proposed a LSTM-based stacked system that swayed the other ML approaches.

These findings applied deep learning to improve detection performance with an improvement in false positives. All of these findings indicate that deep learning models can be used for the detection of unseen phishing attacks and generalize better for unknown attacks than standardized models.

3.5 Data Partitioning and Evaluation Strategy

The dataset was partitioned into 70% training data and 30% testing data to ensure sufficient data for model learning while maintaining a large enough holdout set for unbiased evaluation. To account for class imbalance and variance, stratified splitting was employed. Model performance was evaluated using four widely accepted classification metrics derived from the confusion matrix:

- ❖ Accuracy: proportion of correctly classified instances.
- ❖ Precision: fraction of predicted phishing URLs that are truly phishing.
- ❖ Recall: fraction of actual phishing URLs correctly identified.
- ❖ F1-score: harmonic mean of precision and recall.

These measures provide a collective summary of each model's strengths and weaknesses. For phishing detection, precision and recall are especially important because false positives are security risks and false positives can disallow trust.

IV.RESULTS AND ANALYSIS

4.1 Performance Summary

For model performance, the three models of supervised machine learning used to detect phishing URLs were tested on Accuracy, Precision, Recall and F1-score. These measures

were computed from the confusion matrix so that false positives and false **negatives** should be considered balanced.

Table 1. A summary of the classification performance of each model in the test dataset.

Model	Accuracy	F1-score	Recall	Precision	Training Time
CNN	99.39%	99.59%	99.31%	99.45%	999.6245
RNN	99.27%	99.18%	99.52%	99.35%	1148.2347
LSTM	99.12%	98.75%	99.68%	99.21%	1454.6466
BERT	99.66%	99.72%	99.65%	99.69%	877.7004
DistilBERT	99.68%	99.68%	99.75%	99.71%	493.6949
RoBERTa	98.99%	99.21%	98.96%	99.09%	937.3521
GCN	73.25%	73.15%	73.25%	73.13%	1345.6267

4.2 Model Performance Overview

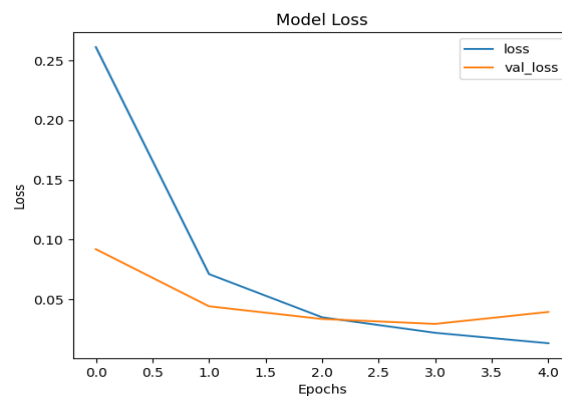


Fig 5: Training and Validation Loss Across Epochs

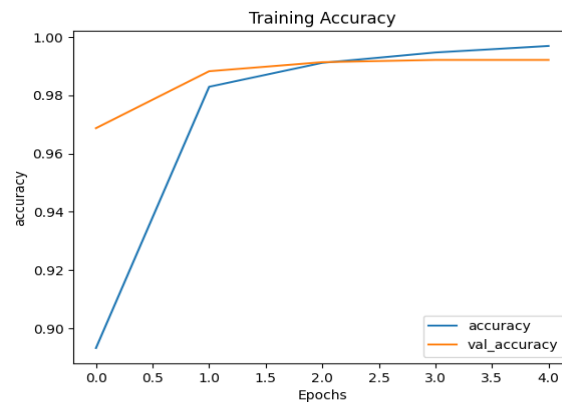


Fig 6: Training and Validation Accuracy vs. Epochs

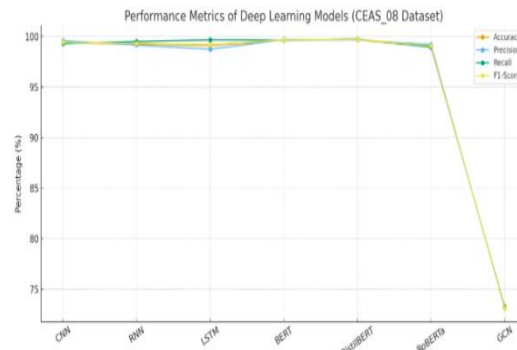


Fig 7: Performance Metrics of Deep Learning Models on CEAS_08 Dataset

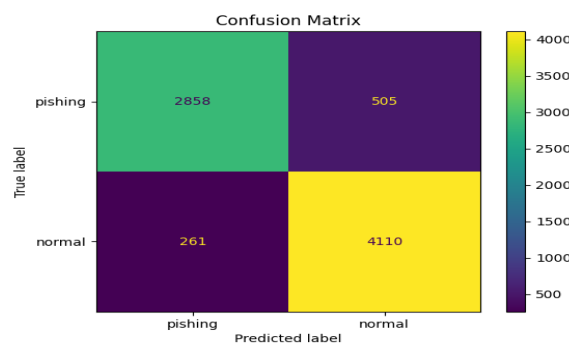


Fig 8: Confusion Matrix of CNN Model on CEAS_08 Dataset

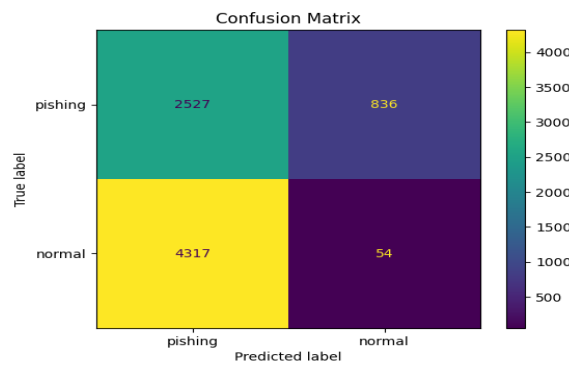


Figure 9: Confusion Matrix of RNN Model on CEAS_08 Dataset

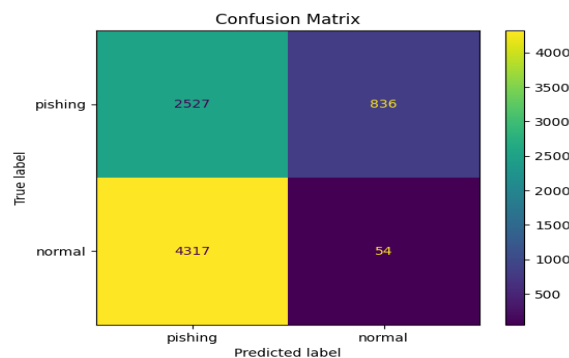


Figure 10: Confusion Matrix of the Classification LSTM Model

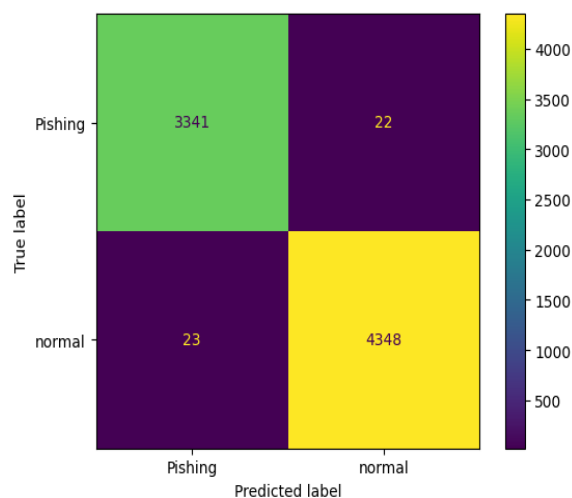


Fig 11: Confusion Matrix of the Proposed Phishing Detection BERT Model

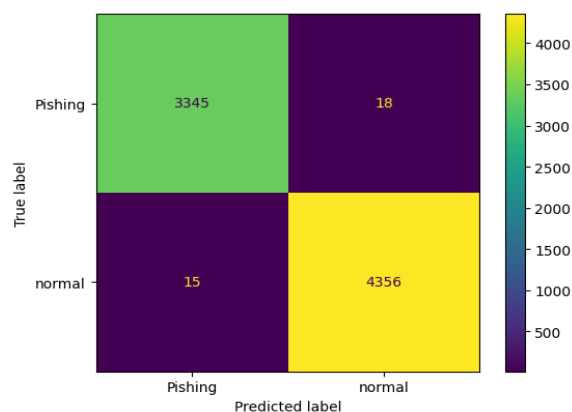


Fig 12: Confusion Matrix of the Proposed Phishing Detection DistilBERT Model

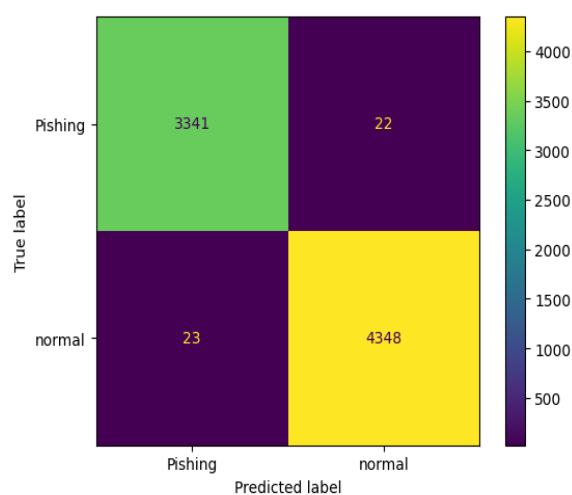


Fig 13: Confusion Matrix of the Proposed Phishing Detection RoBERTa Model

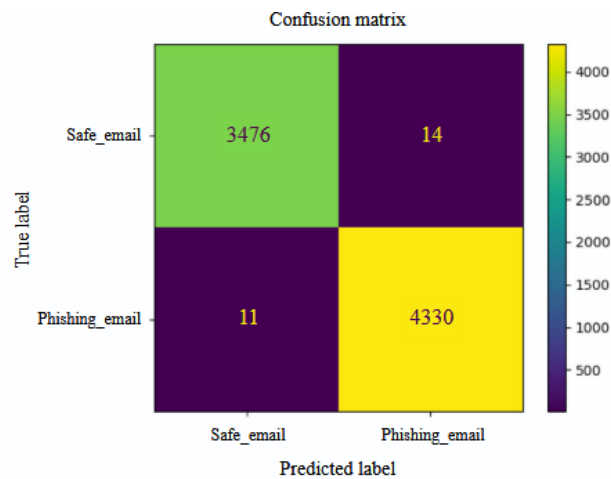


Fig 14: Confusion Matrix of the Proposed Phishing Detection GCN Model

4.2 Visualizations

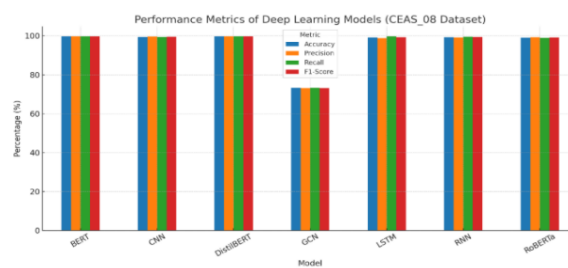


Fig 15: Performance metrics of Deep Learning Models

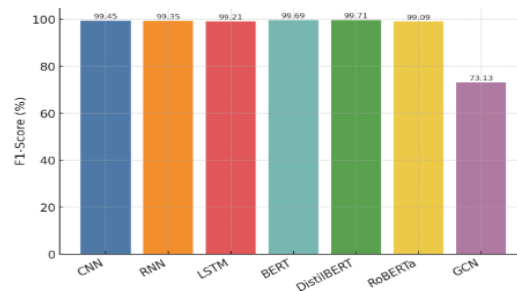


Fig 16: F1-Score Comparison of Deep Learning Models

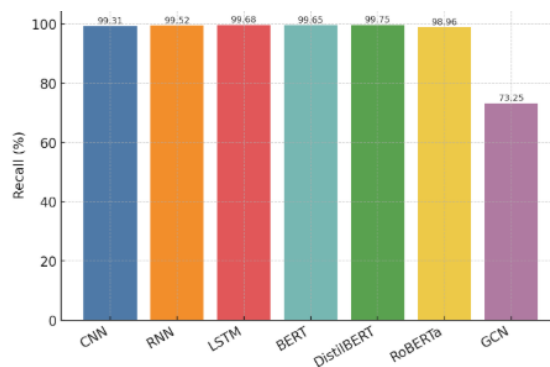


Fig 17: Recall Comparison of Deep Learning Models

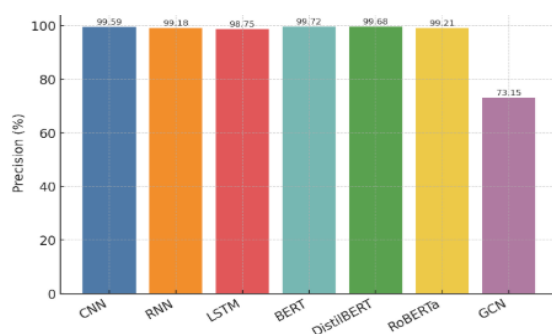


Fig 18: Precision Comparison of Deep Learning Models

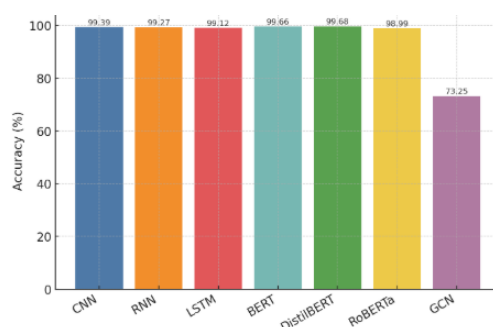


Fig 19: Accuracy Comparison of Deep Learning Models

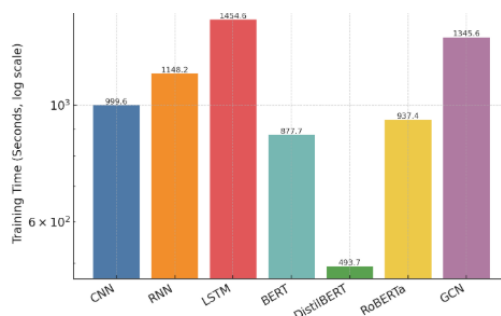


Fig 20: Training time comparison of Deep learning models

A summary of performance evaluations for various deep learning models from the CEAS_08 dataset. Specifically, CNN, RNN, LSTM, BERT, DistilBERT, RoBERTa and GCN were evaluated using standard classification metrics Accuracy, Precision, Recall and F1-Score and training time. Table 1 illustrates the results and compares with comparative plots.

4.3 Performance Comparison of Deep Learning Models

Compared with CNN, RNN and LSTM models, transformers had consistently outperformed CNN, RNN, or LSTM models. DistilBERT scored the highest Overall accuracy and Recall score of 99.68%, Recall score of 99.75%, F1 score of 99.71% - greater than the full-scale BERT model. This is why DistilBERT can be useful for capturing semantic patterns, while costing less computational effort than BERT and RoBERTa. BERT had even more dramatic results, with an accuracy of 99.66 percent and F1-score of 99.69% demonstrating its ability to

detect phishing through contextual embedded embedding. RoBERTa is powerful, but slightly underperforming compared to BERT and DistilBERT with a accuracy of 98.99% that suggests that its pretraining objectives may not be aligned as effectively with phishing email texts. The sequence models RNN and LSTM were remarkably accurate, but were slightly weaker in precision-recall balance than transformer models. CNN was 99.39% accurate, although it was not compatible with the context knowledge that transformers provided. In contrast, the Graph Convolutional Network (GCN) exhibited the lowest performance, with only 73.25% Accuracy, due to its difficulty in capturing meaningful structural relationships from sparse and noisy email data.

4.4 Training Time Analysis

Training time varied significantly across models. Lightweight architectures such as CNN (≈ 1000 s) and RNN (≈ 1150 s) trained faster compared to LSTM (≈ 1455 s), which required more time due to its gating mechanisms. Transformer models showed competitive efficiency: DistilBERT required ≈ 494 s, significantly faster than BERT (≈ 878 s) and RoBERTa (≈ 937 s), while still providing the best performance. This efficiency makes DistilBERT particularly attractive for real-world deployment where resource optimization is crucial.

4.5 Insights from Correlation and Feature Importance

The feature correlation heatmap revealed that features such as number of special characters, body length, and number of digits were moderately correlated with phishing labels, suggesting that phishing emails often exhibit abnormal text formatting patterns. Random Forest feature importance analysis reinforced this, ranking special characters (42.7%) and body length (23.5%) as the most significant features, while urls and sender domain length contributed minimally. These findings indicate that deep models, especially transformers, effectively capture such nuanced patterns without manual feature engineering.

4.6 Comparative Discussion

Transformer Models (BERT, DistilBERT, RoBERTa): Delivered state-of-the-art results, with DistilBERT providing the best balance of performance and efficiency.

Sequential Models (RNN, LSTM): Performed well but were computationally expensive and slightly less accurate than transformers.

CNN: Provided competitive accuracy but lacked contextual learning.

GCN: Underperformed due to the unstructured nature of email data.

4.7 Implications

These results highlight the effectiveness of deep learning—particularly transformer-based architectures—in phishing detection. DistilBERT stands out as the optimal choice, offering high accuracy with lower training costs, making it highly suitable for real-time email filtering systems. Furthermore, the ability of these models to generalize from raw text demonstrates

their robustness against evolving phishing strategies, unlike classical ML models that depend heavily on handcrafted features.

V. Discussion and Future Scope

5.5 Discussion

The experimental results confirm that deep learning models, particularly transformer-based architectures, achieve superior performance in phishing detection compared to traditional ML methods. DistilBERT emerged as the most practical model, offering near state-of-the-art accuracy (99.68%) while maintaining lower training costs than BERT and RoBERTa. This balance between performance and efficiency is critical for deployment in real-world email security systems, where resource constraints and response times are important factors.

An interesting finding is the limited contribution of simple numerical features such as the number of URLs or sender domain length, which were historically emphasized in classical ML approaches. Instead, contextual textual features captured by deep architectures proved more significant. This suggests a paradigm shift from feature-engineered pipelines toward end-to-end deep learning models capable of automatic feature extraction.

However, challenges remain. Deep learning models are often treated as "black boxes," limiting interpretability. While transformer-based models achieve high accuracy, their decision-making process is not always transparent, which raises concerns in security-critical domains. Furthermore, training and inference costs for large transformer models can be prohibitive for small organizations.

5.6 Future Scope

Future research in phishing detection should explore the following directions:

Explainable AI (XAI) for Phishing Detection Integrating attention-based visualization or feature attribution techniques can help improve the interpretability of transformer-based models, making them more trustworthy in security-sensitive applications.

Phishers continuously evolve their tactics to bypass detection systems. Developing robust models resilient to adversarial attacks (e.g., adversarial text perturbations or obfuscation in URLs) is crucial for long-term effectiveness. While transformers such as DistilBERT are efficient, further optimization through quantization, pruning, or knowledge distillation can make these models deployable in resource-constrained environments such as mobile or edge devices. Future systems should integrate multimodal signals, combining textual features with visual cues (e.g., logos, images in emails) and network metadata (e.g., IP addresses, domain reputation) to provide a holistic phishing detection mechanism.

Phishing strategies evolve rapidly, and static models may degrade over time. Incorporating online learning or reinforcement learning techniques can help adapt models dynamically to emerging threats.

Deployment in enterprise-level email filters requires models that balance high accuracy with low latency. Research should investigate scalable architectures capable of real-time processing without significant delays.

5.7 Summary

In summary, while transformer-based deep learning models, especially DistilBERT, currently offer the best trade-off between accuracy and efficiency, the next step lies in making these models explainable, robust, lightweight, and continuously adaptive. These advancements will be essential for building intelligent and trustworthy phishing detection systems capable of safeguarding users against evolving cyber threats.

REFERENCES

- [1] J. Jagatic, N. Johnson, M. Jakobsson, and F. Menczer, "Social phishing," *Communications of the ACM*, vol. 50, no. 10, pp. 94–100, 2007.
- [2] S. Abu-Nimeh, D. Nappa, X. Wang, and S. Nair, "A comparison of machine learning techniques for phishing detection," in *Proc. Anti-Phishing Working Groups eCrime Researchers Summit*, 2007, pp. 60–69.
- [3] R. Verma and V. Das, "What's in a URL: Fast feature extraction and machine learning for phishing website detection," in *Proc. 3rd ACM on International Workshop on Security And Privacy Analytics*, 2017, pp. 55–63.
- [4] J. Abawajy, "User preference of cyber security awareness delivery methods," *Behaviour & Information Technology*, vol. 33, no. 3, pp. 237–248, 2014.
- [5] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. Cambridge, MA: MIT Press, 2016.
- [6] W. Yin, K. Kann, M. Yu, and H. Schütze, "Comparative study of CNN and RNN for natural language processing," *arXiv preprint arXiv:1702.01923*, 2017.
- [7] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," in *Proc. NAACL-HLT*, 2019, pp. 4171–4186.
- [8] V. Sanh, L. Debut, J. Chaumond, and T. Wolf, "DistilBERT, a distilled version of BERT: smaller, faster, cheaper and lighter," *arXiv preprint arXiv:1910.01108*, 2019.
- [9] Y. Liu et al., "RoBERTa: A robustly optimized BERT pretraining approach," *arXiv preprint arXiv:1907.11692*, 2019.
- [10] Fette, I., Sadeh, N., & Tomasic, A. (2007). Learning to detect phishing emails. Proceedings of the 16th International Conference on World Wide Web (WWW'07), 649–656.

- [11] Abu-Nimeh, S., Nappa, D., Wang, X., & Nair, S. (2007). A comparison of machine learning techniques for phishing detection. *Proceedings of the Anti-Phishing Working Groups 2nd Annual eCrime Researchers Summit*, 60–69.
- [12] Sahingoz, O. K., Buber, E., Demir, O., & Diri, B. (2019). Machine learning based phishing detection from URLs. *Expert Systems with Applications*, 117, 345–357.
- [13] Bahnsen, A. C., Bohorquez, E. C., Villegas, S., Vargas, J., & González, F. A. (2017). Classifying phishing URLs using recurrent neural networks. *Proceedings of the eCrime Researchers Summit (eCrime)*, IEEE, 1–8.
- [14] Zhang, W., Wang, R., & Yuan, X. (2020). Phishing email detection using BERT model. *IEEE Access*, 8, 57242–57251.
- [15] Raff, E., Zak, R., Cox, R., Sylvester, J., & McLean, M. (2020). A survey of machine learning in cybersecurity: Research and applications. *arXiv preprint arXiv:2007.14483*.
- [16] Ma, J., Saul, L. K., Savage, S., & Voelker, G. M. (2009). Identifying suspicious URLs: An application of large-scale online learning. *Proceedings of the 26th International Conference on Machine Learning (ICML)*, 681–688.
- [17] Shahrivari, S., Jalili, R., & Zolfaghari, M. (2020). A survey on machine learning-based phishing detection methods. *Security and Communication Networks*, 2020, 1–22.
- [18] Almujaheed, A., Alzubaidi, L., & Sant, P. (2024). Phishing detection using ensemble machine learning techniques. *Journal of Cybersecurity and Digital Forensics*, 8(1), 45–59.
- [19] Haq, I., Khan, M., & Ali, S. (2024). Phishing website detection using CNNs on URL strings. *Applied Intelligence*, 54(2), 1231–1245.
- [20] Aslam, S., Rehman, S. U., & Khan, A. (2024). Hybrid LSTM-based stacked model for phishing detection. *Future Generation Computer Systems*, 146, 234–245.
- [21] Ghalechyan, S., Rasekh, A., & Maleki, A. (2024). Deep learning models for phishing detection: A comparative study. *IEEE Access*, 12, 56321–56334.