

Recent Advances in Phishing Email Classification: A Comprehensive Review of Techniques, Datasets, and Emerging Trends

Parul Patel, Assistant Professor, Department of ICT, VNSGU, pjpatel@vnsgu.ac.in

Abstract—Phishing email remains one of the most prevalent cyber threats, continuously evolving through advanced social engineering techniques and the increasing use of artificial intelligence. Consequently, phishing email classification has become an active research area, with approaches progressing from traditional machine learning methods to deep learning, transformer-based architectures, and large language models (LLMs). This survey presents a comprehensive review of recent advances in phishing email classification by analyzing 29 research papers published between 2020 and 2026, along with industry phishing threat report. The selected studies are systematically examined with respect to benchmark datasets, feature engineering techniques, classification models, experimental setups, and performance evaluation metrics. The survey highlights the evolution of detection approaches, compares the strengths and limitations of various machine learning and deep learning models, and discusses the emerging role of transformer models, LLMs, and hybrid frameworks in improving phishing detection. Furthermore, key research challenges, including dataset limitations, model generalization, computational complexity, explainability, and evaluation inconsistencies, are identified. Finally, future research directions are outlined, emphasizing adaptive learning, explainable AI, lightweight architectures, standardized benchmarks, and robust detection against AI-generated phishing attacks. This survey provides researchers and practitioners with a structured overview of current developments and potential research opportunities in phishing email classification.

Keywords—*Phishing Email Detection, Large Language Model, Deep Learning, Cyber Security.*

I. INTRODUCTION

Email has become one of the most essential communication channels for individuals, businesses, educational institutions, and government organizations. However, its widespread adoption has also made it a primary target for cybercriminals. Among various email-based threats, phishing remains one of the most prevalent and damaging cyberattacks. By impersonating trusted entities, phishing emails aim to deceive recipients into revealing sensitive information, downloading malicious attachments, or performing unauthorized financial transactions. Phishing attacks are becoming more advanced, making them harder to detect.

Recent developments in AI have made phishing attacks more advanced. Attackers now leverage generative AI to create highly convincing, context-aware, and personalized phishing emails that closely resemble legitimate communications. These attacks exploit social engineering techniques such as urgency, authority, trust, and curiosity, making them increasingly difficult to detect using traditional security mechanisms. Recent industry reports also indicate a steady rise in AI-generated phishing, business email compromise (BEC), and targeted spear-phishing attacks,

emphasizing the need for more robust and intelligent detection approaches.

Early phishing detection systems mainly used predefined rules, blacklists, and manually created patterns to identify phishing emails. These methods worked well for known phishing attacks but could not detect new or unknown phishing emails because they depended on predefined rules and patterns.. To address these limitations, machine learning techniques were introduced to learn phishing patterns from historical data. Algorithms such as Support Vector Machines (SVM), Random Forest (RF), Naïve Bayes (NB), and Gradient Boosting improved phishing detection by using features such as email content, structure, and metadata.

As phishing attacks became more complex, deep learning approaches were introduced to improve detection accuracy. Unlike traditional methods, these techniques automatically learn useful features from raw email data. Deep learning models such as CNN, LSTM, Bi-LSTM, and GRU have achieved promising results by capturing the semantic meaning and sequence information present in phishing emails. Recently, transformer-based architectures such as BERT, RoBERTa, DistilBERT, and ALBERT have further enhanced phishing detection through contextual language understanding. The development of Large Language Models

(LLMs), such as GPT, Llama, and Qwen, has created new possibilities for phishing email detection. These models can understand the context of emails, explain their decisions, and analyze complete email content. Recently, hybrid approaches that combine different learning methods have become popular. These approaches improve phishing detection by providing better accuracy, reliability, and efficiency.

Existing studies use different datasets, features, learning algorithms, evaluation metrics, and experimental setups. This makes it difficult to compare their performance. In addition, many benchmark datasets are outdated and do not represent modern phishing attacks, especially AI-generated and multilingual phishing emails. The use of different preprocessing, validation, and evaluation methods makes it difficult to verify and compare the results of existing studies. While many review articles discuss phishing detection, only a few include recent transformer-based models, LLMs, experimental practices, and real-world deployment.

Motivated by these observations, this survey presents a comprehensive review of recent advances in phishing email classification. The survey synthesizes existing research by examining benchmark datasets, feature engineering techniques, machine learning, deep learning, transformer-based models, large language models, hybrid approaches, performance evaluation, and experimental practices. It also identifies current research challenges, emerging trends, and promising directions for future work.

The remainder of this paper is organized as follows. Section 2 describes the review methodology adopted in this survey. Section 3 presents a comparative analysis of phishing email datasets, followed by feature engineering techniques, machine learning, deep learning, transformer-based, large language model, and hybrid approaches for phishing email classification. It also discusses the reported performance of these approaches along with experimental setups and evaluation practices. Section 4 outlines the major research challenges, and Section 5 highlights future research directions. Finally, Section 6 concludes the survey.

II. REVIEW METHODOLOGY

A structured review methodology was adopted to ensure a comprehensive and systematic analysis of recent research on phishing email classification. Rather than providing a paper-by-paper description, the selected studies were synthesized to identify research trends, compare existing methodologies, evaluate their effectiveness, and highlight open research challenges. The review process comprised five major stages: defining research questions, conducting a systematic literature search, selecting relevant studies, extracting key information, and organizing the findings into a unified analytical framework. Figure 2 illustrates the overall review methodology adopted in this survey.

A. Literature Search Approach

A comprehensive literature search was conducted to identify recent studies on phishing email classification. The search focused on peer-reviewed journal articles, conference papers, and recent industry reports published in the field of cybersecurity. Relevant studies were identified using combinations of keywords such as phishing email detection, phishing email classification, machine learning, deep learning, transformer models, large language models, email security. To complement the academic literature, two recent cybersecurity threat reports published in 2026 were also considered to capture emerging phishing trends and real-world challenges. The collected literature was screened to ensure relevance and quality. Studies were included if they primarily focused on phishing email classification, proposed or evaluated detection techniques, reported experimental results, or provided valuable insights into datasets, feature engineering, or evaluation methodologies. Studies addressing unrelated phishing domains, duplicate publications, or lacking sufficient technical details were excluded. Following the screening process, 28 research papers and two industry reports were selected for detailed analysis.

B. Data Extraction

To facilitate a consistent comparison, information from each selected study was extracted using a standardized template. The extracted information included:

- ☐ Dataset characteristics
- ☐ Feature engineering techniques
- ☐ Classification models
- ☐ Experimental setup
- ☐ Performance metrics
- ☐ Advantages and limitations
- ☐ Future research directions

This standardized extraction process enabled a systematic comparison of different approaches while minimizing inconsistencies across studies.

C. Review Framework

After data extraction, the selected studies were organized into a structured review framework consisting of five analytical dimensions. Each dimension examines a specific aspect of phishing email classification to provide a comprehensive understanding of the current research landscape.

C.1 Phishing Email Datasets

The first stage analyzes the benchmark datasets used for model development and evaluation. Public datasets, hybrid datasets, and domain-specific corpora are compared in terms

of dataset size, class distribution, feature availability, and suitability for phishing email research. Their strengths and limitations are also discussed.

Table 1. Public Benchmark Datasets for Phishing Email Classification

Dataset	Total Samples	Phishing Samples	Legitimate Samples	Data Representation	Primary Purpose
Enron Email Dataset	517,431	0	517,431	Email content (Header, Subject, Body)	Legitimate (ham) email benchmark
SpamAssassin Public Corpus	9,324	1,897	7,427	Raw email (Header, Subject, Body)	Spam and legitimate email classification
Nazario Phishing Corpus	4,559	4,559	0	Raw email (Header, Subject, Body)	Phishing email benchmark
CEAS 2008 Spam Corpus	137,705	~82,000*	~55,700	Complete RFC-822 email	Spam/phishing email detection
PhishBench Benchmark	21,000	10,500	10,500	Raw email	Balanced benchmark for phishing email classification

Table 2. Website/URL Benchmark Datasets

Dataset	Total Samples	Phishing Samples	Legitimate Samples	Features	Primary Purpose
UCI Phishing Website Dataset	11,055	6,157	4,898	30 predictive features + 1 class label	Website phishing detection

Table 3. Online Phishing Repositories

Repository	Repository Statistics	Update Frequency	Available Information	Primary Use
PhishTank	4.14 million verified phishing URLs; ~51.8K active URLs	Continuously updated	Phishing URLs	Building phishing datasets and validating malicious URLs
OpenPhish	~5.17 million processed URLs	~8,600 new URLs added weekly	URL metadata (hostname, path, SSL, IP, ASN, country, brand, etc.)	Real-time phishing intelligence and URL validation

C.2 Feature Engineering Techniques

The second stage investigates the evolution of feature engineering methods. The review categorizes features into

header-based, URL-based, content-based, metadata-based, semantic, and contextual representations, highlighting the transition from handcrafted features to automatically learned representations.

Table 4. Comparison of Feature Categories Used in Phishing Email Classification

Feature Category	Description	Advantages	Limitations
Header Features	Sender information, SPF, DKIM, routing fields	Detect sender spoofing	Can be manipulated or missing
URL Features	URL length, domain, IP address, redirects	Strong phishing indicators	Less useful without embedded links
Lexical Features	Keywords, TF-IDF, n-grams	Simple and interpretable	Limited contextual understanding
Semantic Features	Word embeddings, contextual representations	Capture semantic meaning	Higher computational cost
HTML Features	HTML tags, forms, scripts	Detect malicious email structures	Not present in all emails
Metadata Features	Attachments, MIME type, timestamps	Complement textual features	Limited effectiveness individually
Learned Features	Transformer and LLM embeddings	Automatic contextual feature learning	Computationally intensive

The reviewed studies indicate a clear transition from handcrafted features to automated feature learning. Earlier research predominantly utilized lexical, header, and URL features because of their simplicity and effectiveness with traditional machine learning algorithms. As phishing attacks became more sophisticated, semantic representations and contextual embeddings gained prominence, enabling models to better capture subtle linguistic patterns. Recent transformer-based and LLM approaches enhance phishing detection by learning rich contextual representations directly from email content. However, many studies demonstrate that combining handcrafted features with learned representations

often yields more robust and reliable classification performance than relying on a single feature type.

Feature engineering has progressed from manual feature extraction to contextual representation learning. While transformer and LLM-based embedding have reduced the dependence on handcrafted features, traditional header and URL features continue to provide complementary information. Consequently, hybrid feature representations remain a promising approach for developing accurate and robust phishing email classification systems.

C.3 Phishing Email Classification Models

The third stage examines the evolution of classification techniques and categorizes the reviewed studies into the following learning paradigms:

Traditional machine learning remains an important baseline for phishing email classification due to its efficiency and ease of implementation. The studies primarily employ Support Vector Machines (SVM), Random Forest, Multinomial Naïve Bayes, LightGBM, and XGBoost.

Table 5. Comparison of Traditional Machine Learning Models

Authors	Model	Typical Dataset	Reported Performance	Advantages	Limitations
[15]	SVM	Unified Phishing Corpus	99.23% Accuracy, 99.43% F1	Effective for high-dimensional text classification	Computationally expensive on large datasets
[4],[20]	Random Forest	RF Phishing, SA-Nazario	Up to 99.70% Accuracy	Robust and handles heterogeneous features	May overfit with complex trees
[13],[19]	Multinomial Naïve Bayes	Kaggle, App River	98.78% Accuracy	Fast training and prediction	Assumes feature independence
[18],[21]	LightGBM	Integrated Corpora	96.81% Accuracy	Efficient with low inference latency	Sensitive to feature quality
[18],[21]	XGBoost	Integrated Corpora	95.98–96.76% Accuracy	High predictive performance	Requires careful hyperparameter tuning

C.4 Deep Learning Models

Deep learning approaches overcome the limitations of handcrafted feature engineering by learning hierarchical

representations directly from email content. CNN, LSTM, Bi-LSTM, GRU, and generative models such as PDGAN are the most frequently reported architectures.

Table 6. Comparison of Deep Learning Models

Authors	Model	Reported Performance	Advantages	Limitations
[3], [14], [27]	CNN	Up to 98.87% Accuracy	Automatic feature extraction	Limited sequential modeling
[3], [14], [19]	LSTM / Bi-LSTM	Up to 95.91% Accuracy	Captures long-term dependencies	Higher computational cost
[14], [27]	GRU / Bi-GRU	Up to 98.0% Accuracy	Faster than LSTM	Slightly lower accuracy in some studies
[26]	PDGAN	97.58% Accuracy	Improves robustness through synthetic data generation	Long training time and high memory usage

The reviewed studies indicate that deep learning models generally outperform traditional machine learning by learning discriminative representations directly from raw email data. Hybrid neural architectures, particularly CNN combined with recurrent networks, further enhance performance by simultaneously capturing local textual patterns and long-range contextual dependencies.

C.5 Transformer-based Models

Transformer architectures introduce attention mechanisms that enable contextual understanding across an entire email.

Table 7. Transformer-based Models

Autho	Model	Reported Performance	Advantages	Limitations
[11], [13]	BERT	91% Accuracy	Rich contextual understanding	Computationally intensive
[11]	RoBERTa	94% Accuracy	Improved contextual representation	High resource requirements
[13]	ALBERT	99.37% F1	Parameter efficient	Benefits depend on task complexity
[1],[12]	DistilBERT	92% Macro F1	Faster and lightweight	Slight reduction in accuracy

Compared with recurrent architectures, transformer models demonstrate superior contextual understanding and improved robustness against sophisticated phishing emails. However, their computational requirements remain significantly higher than those of traditional machine learning methods.

C.6 Large Language Models

Recent studies investigate large language models for phishing email classification, particularly to address AI-generated phishing attacks.

Table 8. Large Language Models

Author	Model	Advantages	Limitations
[6],[20]	Llama 3.1	Strong zero-shot reasoning	High computational requirements
[6],[20]	Qwen 2.5	Effective contextual analysis	Resource intensive
[1]	MultiPhishGuard	Multi-agent reasoning and multimodal analysis	Complex implementation
[7]	Flan-T5	Explainable predictions	Limited by model size

Large language models extend phishing detection beyond classification by providing contextual reasoning and natural language explanations. Although promising, their deployment remains challenging because of high computational costs and inference latency.

C.7 Hybrid Models

Several recent studies combine multiple learning paradigms to improve detection performance.

Table 9. Hybrid Models

Author	Model	Reported Performance	Strength
[6], [20]	LSTM-CNN	98.2% Accuracy	Learns spatial and sequential features
[6], [20]	1D-CNN + Bi-GRU	99.68% Accuracy	High detection accuracy
[1]	Dual-Layer Framework	98.74% Accuracy	Handles class imbalance effectively
[7]	Two-Stage Persuasion Model	96% F1	Improved interpretability

Hybrid models consistently report the highest performance by combining complementary feature learning capabilities, although they generally require greater computational resources and more complex training procedures. The reviewed studies reveal a gradual shift from traditional machine learning to transformer-based and LLM-driven approaches. Traditional algorithms remain attractive for resource-constrained environments because of their efficiency and interpretability. Deep learning models improve feature representation through automated learning, while transformer architectures enhance contextual understanding using attention mechanisms. Large language models further extend these capabilities by performing holistic reasoning and generating interpretable explanations. Hybrid models currently achieve the best overall

performance by integrating multiple architectures, although their computational demands may limit real-time deployment. The comparative analysis focuses on the strengths, limitations, and application scenarios of each category.

III. EXPERIMENTAL SETUP AND EVALUATION

The performance of phishing email classification models is influenced not only by the learning algorithm but also by the experimental setup, training strategy, and evaluation methodology. The reviewed studies adopted diverse configurations involving different optimizers, validation methods, hardware platforms, and performance metrics. Table 10 summarizes the common experimental practices reported across the selected studies.

Table 10. Summary of Experimental Setup and Evaluation Practices

Aspect	Common Practices Reported in the Literature
Optimizers	Adam was the most widely used optimizer, followed by AdamW for transformer-based models.
Learning Rate	Typically ranged from 2×10^{-5} to 1×10^{-3} , depending on the model architecture.
Batch Size	Most studies used batch sizes of 16–32.
Training Epochs	Transformer models generally required 3–5 epochs, while CNN/LSTM-based models used 40–100 epochs.
Loss Function	Binary Cross-Entropy was the most commonly adopted loss function for binary phishing classification.
Regularization	Dropout, Batch Normalization, Early Stopping, Gradient Clipping, and Learning Rate Scheduling.
Hyperparameter Tuning	Random Search, Talos optimization, Grid Search, and manual tuning.
Hardware Platforms	CPU systems for traditional ML; Google Colab (NVIDIA T4 GPU), cloud GPUs, and local GPU environments for deep learning and LLMs.
Validation Strategy	70:30 or 80:20 train-test split, stratified sampling, and 10-fold cross-validation.
Evaluation Metrics	Accuracy, Precision, Recall, F1-score, AUC, FPR, Training Time, and Inference Time.

Deep learning and transformer-based models predominantly employed the Adam optimizer because of its efficient convergence, while AdamW was commonly used for transformer fine-tuning. Learning rates typically ranged from 2×10^{-5} to 1×10^{-3} , with batch sizes between 16 and

32. Transformer models generally converged within 3–5 epochs, whereas CNN- and LSTM-based models required 40–100 epochs. To improve generalization and reduce overfitting, many studies incorporated techniques such as dropout, batch normalization, and early stopping.

Traditional machine learning models were generally trained on CPU-based systems, whereas deep learning and transformer-based models relied on GPU acceleration. Google Colab equipped with NVIDIA T4 GPUs emerged as the most frequently used experimental platform due to its accessibility and computational capability. Recent studies involving large language models also explored cloud-based GPU environments and lightweight fine-tuning techniques such as LoRA and QLoRA to reduce computational requirements.

The reviewed studies employed various validation strategies, including 70:30 and 80:20 train-test splits, stratified sampling, and 10-fold cross-validation. Traditional machine learning studies predominantly adopted cross-validation, whereas deep learning models commonly used a separate validation set for hyperparameter tuning and early stopping.

Model performance was primarily evaluated using accuracy, precision, recall, F1-score, and AUC, with several studies additionally reporting false positive rate (FPR), training time, and inference latency. Although Matthews Correlation Coefficient (MCC) has been recognized as an effective metric for evaluating imbalanced datasets, it was reported in relatively few studies.

Overall, the reviewed literature demonstrates an increasing emphasis on standardized evaluation and reproducible experimentation. Nevertheless, differences in dataset, preprocessing techniques, validation strategies, and reporting practices continue to make direct comparison across studies challenging. The adoption of standardized benchmark datasets and evaluation protocols would improve the consistency and reproducibility of future phishing email classification research.

IV. RESEARCH CHALLENGES

Despite significant advancements in phishing email classification, several challenges continue to limit the effectiveness and practical deployment of existing detection systems. A major concern is the dependence on outdated benchmark datasets, such as Enron, SpamAssassin, and Nazario, which do not adequately represent modern phishing campaigns involving AI-generated, multilingual, and highly personalized attacks. Additionally, class imbalance and the limited availability of large, publicly accessible datasets affect model robustness and reproducibility.

Recent deep learning, transformer-based, and large language models have achieved remarkable detection performance; however, these models require substantial computational resources, longer training times, and higher inference latency, making real-time deployment challenging. Another important limitation is the lack of model interpretability, as many high-performing approaches function as black boxes, reducing user trust and limiting adoption in security-critical environments.

Furthermore, phishing techniques continuously evolve through adversarial manipulation, social engineering, and generative AI, making it difficult for models trained on historical data to generalize to emerging threats. Finally, inconsistencies in dataset, preprocessing methods, validation strategies, and evaluation metrics hinder fair comparison across studies. Addressing these challenges requires the development of diverse benchmark datasets, lightweight and explainable models, standardized evaluation protocols, and adaptive learning techniques capable of responding to the rapidly evolving phishing landscape.

V. FUTURE RESEARCH DIRECTIONS

Future research should focus on developing adaptive, explainable, and computationally efficient phishing email detection systems. The creation of large, diverse, and up-to-date benchmark dataset containing multilingual and AI-generated phishing emails is essential for improving model generalization. Emerging approaches, including transformer-based models, large language models (LLMs), explainable AI (XAI), federated learning, and continual learning, offer promising directions for enhancing detection accuracy and robustness. Furthermore, standardized evaluation protocols and lightweight model optimization techniques will facilitate fair comparison and practical deployment in real-world environments.

VI. CONCLUSION

This survey examines recent developments in phishing email classification by reviewing relevant research studies and industry reports. The study examined datasets, feature engineering techniques, classification models, performance evaluation, and current research challenges. The findings indicate a clear transition from traditional machine learning to deep learning, transformer-based models, and LLMs, leading to improved detection performance. However, challenges related to dataset quality, explainability, computational complexity, and standardized evaluation remain. Addressing these issues will enable the development of more accurate, robust, and deployable phishing email detection systems capable of combating emerging cyber threats.

REFERENCES

- [1] Kushwaha, Chitra and Kumar, Prashant and Kiran Rajpoot, Abha, Advances in Phishing Detection: A Comprehensive Review of Machine Learning, Deep Learning, and Transformer-based Approaches (January 16, 2026). Proceedings of the International Conference on Innovative Computing and Communication (ICICC 2026), <http://dx.doi.org/10.2139/ssrn.6079506>
- [2] Anirudh S,P Radha Nishant,Sanjay Baitha, K Dinesh Kumar, An Ensemble Classification Model for Phishing Mail Detection, Procedia Computer Science, Volume 233, 2024, Pages 970-978, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2024.03.286>.
- [3] Rokunojjaman, Md & Malik, Anup & Sajeeb, MD & Yahiduzzaman, MD & Islam, MD. (2025). Multi-Feature Hybrid LSTM-CNN Framework for Phishing Email Detection. International Journal of

- Electronics and Communications Systems. 5. 93-103. 10.24042/ijecs.v5i1.26764.
- [4] Yasin, Adwan & Abuhasan, Abdelmunem. (2016). An Intelligent Classification Model for Phishing Email Detection. International Journal of Network Security & Its Applications. 8. 55-72. 10.5121/ijnsa.2016.8405.
- [5] Koide, Takashi & Fukushi, Naoki & Nakano, Hiroki & Chiba, Daiki. (2025). ChatSpamDetector: Leveraging Large Language Models for Effective Phishing Email Detection. 10.1007/978-3-031-94455-0_14.
- [6] Nuno Costa, José Cecilio, Ruben Salgueiro, Mauricio Rosa, and Dulce Domingos. 2026. LLM-Based Framework for Email Classification and Phishing Detection. Ada Lett. 45, 1 (June 2025), 70–74. <https://doi.org/10.1145/3784987.3784993>
- [7] SajadU, P. (2025). A Robust and Explainable Transformer-Based Framework for Phishing Email Detection.
- [8] Shah, A. (2021). Classification and Detection of email Phishing using random Forest supervised-unsupervised machine learning algorithms (Doctoral dissertation, Dublin, National College of Ireland).
- [9] S. Salloum, T. Gaber, S. Vadera and K. Shaalan, "A Systematic Literature Review on Phishing Email Detection Using Natural Language Processing Techniques," in IEEE Access, vol. 10, pp. 65703-65727, 2022, doi: 10.1109/ACCESS.2022.3183083.
- [10] Lukáš Halgaš, Ioannis Agrafiotis, and Jason R. C. Nurse. 2019. Catching the Phish: Detecting Phishing Attacks Using Recurrent Neural Networks (RNNs). In Information Security Applications: 20th International Conference, WISA 2019, Jeju Island, South Korea, August 21–24, 2019, Revised Selected Papers. Springer-Verlag, Berlin, Heidelberg, 219–233. https://doi.org/10.1007/978-3-030-39303-8_17
- [11] Ibrahim, M., & Elhafiz, R. (2026). Phishing Email Detection Using BERT and RoBERTa. Computation, 14(2), 46. <https://doi.org/10.3390/computation14020046>
- [12] Tooher, P., & Lallie, H. S. (2025). A Two-Stage Deep Learning Framework for AI-Driven Phishing Email Detection Based on Persuasion Principles. Computers, 14(12), 523.
- [13] Jay Doshi, Kunal Parmar, Raj Sanghavi, and Narendra Shekokar. 2023. A comprehensive dual-layer architecture for phishing and spam email detection. Comput. Secur. 133, C (Oct 2023). <https://doi.org/10.1016/j.cose.2023.103378>
- [14] Tusher, Ekramul & Ismail, Mohd Arfian & Farihan, Anis. (2025). Email Spam Classification Based on Deep Learning : A Review. Iraqi Journal for Computer Science and Mathematics. 06. 24-36.
- [15] Ahmed, Asif & Husnain, Esha & Eman, Hafza. (2026). Enhanced Phishing Email Classification through Multi-Level Feature Fusion.
- [16] Y. Sun, N. Chong and H. Ochiai, "Federated Phish Bowl: LSTM-Based Decentralized Phishing Email Detection," 2022 IEEE International Conference on Systems, Man, and Cybernetics (SMC), Prague, Czech Republic, 2022, pp. 20-25, doi: 10.1109/SMC53654.2022.9945584.
- [17] Ian Fette, Norman Sadeh, and Anthony Tomasic. 2007. Learning to detect phishing emails. In Proceedings of the 16th international conference on World Wide Web (WWW '07). Association for Computing Machinery, New York, NY, USA, 649–656. <https://doi.org/10.1145/1242572.1242660>
- [18] Jandaeng, C., Koad, P., Zolkipli, M. F., & Phuttharak, J. (2026). TERA: A Trade-Off Evaluation and Resource-Aware Framework for Spam and Phishing Email Detection. Informatics, 13(5), 72. <https://doi.org/10.3390/informatics13050072>
- [19] Bagui, S., Nandi, D. & White, R. J. (2021). Machine Learning and Deep Learning for Phishing Email Classification using One-Hot Encoding. Journal of Computer Science, 17(7), 610-623. <https://doi.org/10.3844/jcssp.2021.610.623>
- [20] Akinyelu, Andronicus A., Adewumi, Aderemi O., Classification of Phishing Email Using Random Forest Machine Learning Technique, Journal of Applied Mathematics, 2014, 425731, 6 pages, 2014. <https://doi.org/10.1155/2014/425731>
- [21] J. D. Duarte et al., "Machine Learning for Early Detection of Phishing URLs in Parked Domains: An Approach Applied to a Financial Institution," in IEEE Access, vol. 13, pp. 145736-145753, 2025, doi: 10.1109/ACCESS.2025.3599454.
- [22] Abdulla Al-Subaiey, Mohammed Al-Thani, Naser Abdullah Alam, Kaniz Fatema Antora, Amith Khandakar, SM Ashfaq Uz Zaman, Novel interpretable and robust web-based AI platform for phishing email detection, Computers and Electrical Engineering, Volume 120, Part A, 2024, 109625, ISSN 0045-7906, <https://doi.org/10.1016/j.compeleceng.2024.109625>.
- [23] Baskota, S. (2025). Phishing URL Detection using Bi-LSTM. ArXiv, abs/2504.21049.
- [24] Rathee, Dhruv & Mann, Suman. (2022). Detection of E-Mail Phishing Attacks – using Machine Learning and Deep Learning. International Journal of Computer Applications. 183. 1-7. 10.5120/ijca2022921868.
- [25] van, O., Stepan, L. & Vladislav, S. Analysis and testing of systems countering phishing and social engineering attacks at the corporate level. Int. J. Inf. Secur. 25, 69 (2026). <https://doi.org/10.1007/s10207-026-01238-w>
- [26] Kavya, S., Sumathi, D. Staying ahead of phishers: a review of recent advances and emerging methodologies in phishing detection. Artif Intell Rev 58, 50 (2025). <https://doi.org/10.1007/s10462-024-11055-z>
- [27] Altwaijry N, Al-Turaiki I, Alotaibi R, Alakeel F. Advancing Phishing Email Detection: A Comparative Study of Deep Learning Models. Sensors (Basel). 2024 Mar 24;24(7):2077. doi: 10.3390/s24072077. PMID: 38610289; PMCID: PMC11013960.
- [28] Anirudh S, P Radha Nishant, Sanjay Baitha, K Dinesh Kumar, An Ensemble Classification Model for Phishing Mail Detection, Procedia Computer Science, Volume 233, 2024, Pages 970-978, ISSN 1877-0509, <https://doi.org/10.1016/j.procs.2024.03.286>.
- [29] Anti-Phishing Working Group. (2026). Phishing activity trends report: 1st quarter 2026. <https://apwg.org/trendreports>