

Machine Learning-Based Phishing Website Detection with Explainable AI for Improved Transparency and User Trust

Harshith TP¹ and Manju Sadasivan²

^{1,2}School of Science and Computer Studies CMR University Bengaluru, India

Abstract—Phishing websites remain one of the most widespread cybersecurity threats, continually deceiving users into disclosing sensitive information through increasingly sophisticated attack techniques. According to the Anti-Phishing Working Group, more than 1.2 million unique phishing websites were identified during the second half of 2023, representing a substantial increase over previous reporting periods. Conventional blacklist-based detection methods often fall short against newly emerging phishing websites because they rely on previously identified malicious URLs. To alleviate the discussed problem, this paper proposes an Explainable Artificial Intelligence (XAI)-enabled phishing detection framework that combines Random Forest and Gradient Boosting classifiers with Local Interpretable Model-agnostic Explanations. The framework utilizes 18 carefully selected URL and domain-based features to achieve accurate and transparent phishing detection. Unlike conventional ensemble models that operate as black-box systems, the proposed approach provides feature-level explanations for every prediction, enabling security analysts to understand, validate, and trust the classification process. Experimental evaluation on the combined UCI Machine Learning Repository and Kaggle phishing datasets demonstrates that the proposed framework achieves an accuracy of 99.2%, outperforming individual machine learning classifiers, such as Logistic Regression, Decision Tree, and Support Vector Machine. By integrating high predictive performance with interpretable decision-making, the proposed framework enhances the reliability, transparency, and feasibility of phishing detection systems in operational cybersecurity environments.

Keywords—Explainable AI, Enhanced Trust, Gradient Boosting, Machine Learning, Phishing Detection.

Date of Submission: 15-08-2026

Date of acceptance: 31-08-2026

I. INTRODUCTION

Phishing websites are one of the most enduring cybersecurity risks to people and organizations because of the sharp rise in phishing attacks that has coincided with the expansion of internet services. In order to fraudulently induce users to share personal details, such as login credentials, bank account information, and personal information, these malicious websites mimic trustworthy websites. Accurately identifying malicious websites has become more difficult as phishing techniques become more sophisticated [1]. Conventional detection approaches, particularly blacklist-based methods, rely on previously identified malicious URLs and are therefore ineffective against newly created phishing websites. As a result, there is an increasing demand for clever and flexible detection systems that can diagnose attacks that haven't been seen before. By identifying malicious websites with a high degree of accuracy and learning intricate patterns from phishing data, machine learning (ML) has become a viable solution [2]. Ensemble learning algorithms such as Gradient Boosting and Random Forest have consistently achieved higher predictive performance among ML-based techniques, reaching accuracy levels above 98% on benchmark phishing datasets. Nevertheless, these models typically function as "black-box" systems, offering no insight into the reasoning behind their predictions, despite their high classification accuracy. In cybersecurity applications, where transparent and justifiable decisions are crucial, their lack of interpretability limits their adoption, undermines user confidence, and makes model validation and debugging more difficult. Furthermore, the significance of models that integrate high predictive performance with explainable decision-making has been brought to light by the increasing ethical and regulatory requirements for reliable artificial intelligence. This paper suggests a novel Explainable Artificial Intelligence (XAI)-enabled phishing website detection framework that integrates Random Forest and Gradient Boosting classifiers with Local Interpretable Model-agnostic Explanations (LIME) in order to mitigate these drawbacks. Without sacrificing detection accuracy, the suggested framework maintains the predictive power of ensemble learning

while offering feature-level explanations that let cybersecurity analysts understand, verify, and have faith in individual classification results. This study makes three major contributions. In order to achieve high predictive performance and transparent decision-making, a novel XAI-enabled phishing detection framework is first proposed by integrating LIME with ensemble ML algorithms, specifically Random Forest and Gradient Boosting. Second, using correlation analysis and information gain-based feature selection, an optimal set of eighteen discriminative URL- and domain-based features is found, improving detection performance while lowering model complexity. Third, on the combined UCI ML Repository and Kaggle phishing datasets, a thorough comparative examination of five ML algorithms demonstrates that the suggested framework achieves an accuracy of 99.2 percent. Furthermore, the incorporation of feature-level explanations enhances the interpretability and trustworthiness of the detection process, making the framework well suited for practical cybersecurity applications where both predictive accuracy and explainability are essential.

II. LITERATURE REVIEW

In order to increase detection accuracy and adaptability against changing cyberthreats, recent developments in phishing website detection have made greater use of ML and deep learning techniques. The majority of previous research has focused on improving predictive performance using adversarial defense mechanisms, deep learning architectures, ensemble learning, and enhanced feature engineering. For the identification of phishing websites, Wei and Sekiyash showed that ensemble ML algorithms outperform individual classifiers in terms of both detection accuracy and computational efficiency [3]. Zieni et al. found that ML-based approaches consistently outperform traditional signature- and blacklist-based methods in identifying previously unseen phishing attacks after conducting a thorough analysis of phishing detection techniques [4]. Duy et al. examined defense mechanisms and adversarial attacks using multimodal learning models, demonstrating their efficacy in enhancing the resilience and dependability of phishing detection systems [5]. The substantial impact that feature engineering has on classification performance has drawn attention. Rao et al. reported that carefully selected URL- and domain-based features, combined with Random Forest classifiers, achieve high detection accuracy [6]. Effective feature engineering significantly enhances model performance by identifying the most informative phishing characteristics, according to Kustiawan and Ghauth [7]. Similarly, Nayak et al. showed that choosing the right features reduces computational complexity while increasing classification accuracy [8]. There are more recent deep learning methods for phishing. Zara et al. demonstrated that when compared to traditional ML algorithms, deep learning and ensemble approaches show better adaptability to new phishing attacks [9]. Tang and Mahmoud demonstrated the efficacy of deep learning in real-time phishing detection by proposing a browser-based phishing detection framework using convolutional neural networks [10]. Ahmad et al. reviewed Artificial Intelligence-based phishing detection models and identified evolving attack patterns, limited benchmark datasets, and model interpretability as persistent research challenges [11]. Rafsanjani et al. developed a framework for identifying malicious URLs based on features that greatly improved classification performance using feature evaluation techniques [12]. Dsouza et al. found that incremental learning techniques are especially successful for real-time phishing detection due to their capacity to adjust to constantly changing attack patterns after conducting a comparative analysis of multiple ML models [13]. Setu et al. presented a hybrid feature selection methodology that lowers feature dimensionality while simultaneously increasing accuracy detection [14]. Pillai et al. examined evasion attacks on ML-based phishing classifiers and suggested defense strategies to strengthen the model's resistance to hostile attacks [15]. Additionally, Kavya and Sumathi proposed a multimodal temporal graph fusion framework that successfully integrates both static and dynamic website characteristics to enhance phishing detection performance [17], while Wangchuk and Gonsalves emphasized the relevance of integrating multimodal features into phishing detection systems [16].

A. Research Gap

Even with the tremendous advancements in ML-based phishing detection, a number of crucial research issues are still unsolved. First, even though ensemble learning algorithms like Gradient Boosting and Random Forest regularly produce detection accuracies higher than 98 percent, they primarily function as black-box models that don't reveal the reasoning behind their predictions. In security-critical settings where explainable decisions are critical, this lack of transparency diminishes user confidence and restricts their adoption. Second, current studies largely focus on predictive performance as determined by accuracy, precision, recall, and F1-score, with relatively little attention for model interpretability. As a result, cybersecurity professionals often lack the data needed to evaluate the accuracy of individual predictions or comprehend the variables affecting classification results. Third, there hasn't been much research done on

combining XAI with ensemble learning to identify phishing websites. Existing approaches generally lack a unified framework that combines the high predictive capability of ensemble models with feature-level interpretability. Furthermore, assessing the usefulness of explainable predictions in assisting cybersecurity analysts in making decisions has received little attention.

To overcome all these limitations, this paper suggests a novel introduces an XAI-enabled phishing website detection framework that combines ensemble learning algorithms with LIME to deliver clear, feature-level explanations while retaining high predictive performance. The suggested framework achieves an accuracy of 99.2 percent, according to experimental evaluation. As a result, this level of security, and consequently, it makes it practical. Table I shows a comparison between the suggested framework and current state-of-the-art methods.

TABLE I: COMPARATIVE ANALYSIS OF PROPOSED FRAMEWORK WITH STATE-OF-THE-ART PHISHING DETECTION SYSTEMS

Study	Method	Dataset	XAI Integration	Real-Time Capability	Limitations
Wei & Sekiya (2022) [3]	Ensemble ML	UCI	No	Partial	No explainability; Limited feature analysis
Rao & Pais (2019) [6]	Random Forest	Custom	No	No	Single dataset; no real-time validation
Nayak et al. (2025) [8]	Hybrid ML-DL	UCI + Kaggle	No	Yes	Limited feature explanations
Zara et al. (2024) [9]	Deep Learning	ISCX-URL	No	No	Computationally expensive; no interpretability
Tang & Mahmoud (2021) [10]	CNN	Custom	No	Yes	No XAI; high false positive rate
Dsouza et al. (2024) [13]	Incremental Learning	Multiple	No	Yes	Lower accuracy; concept drift issues
Proposed Framework	Ensemble ML + LIME-XAI	UCI + Kaggle	Yes	Yes	Explainable; high accuracy; balanced performance

III. THEORETICAL BACKGROUND

A. Role of Machine Learning in Phishing Detection

ML has emerged as a key technology in contemporary cybersecurity, especially for the detection of phishing websites. In contrast to conventional methods that rely on manually updated blacklists or pre-established signatures, ML algorithms use historical data to identify previously undiscovered phishing websites by learning discriminative patterns. This feature makes it possible for ML-based systems to successfully adjust to the constantly changing nature of phishing attacks. ML models for phishing detection are

trained on datasets that include both authentic and fraudulent websites. The models are able to differentiate between malicious and legitimate websites because they have learned the relationships between website characteristics and their corresponding class labels during the training process. ML models are more appropriate for dynamic cybersecurity environments because of this data-driven approach, which improves their ability to identify phishing attacks that were previously unknown. In the course, labels labels between various training. URL structure, HTTPS usage, domain registration details, subdomain characteristics, and other lexical and host-based attributes are frequently examined features. Even in cases where malicious URLs have not yet been reported in blacklist databases, ML algorithms can reliably differentiate phishing websites from trustworthy ones by taking advantage of these patterns, giving them more flexibility to respond to new cyber threats [11].

B. Machine Learning Algorithms Used

For the purpose of detecting phishing websites, a number of ML algorithms have been studied; depending on the characteristics of the dataset and the difficulty of the classification task, each algorithm offers unique benefits [6].

i) *Logistic Regression*: This supervised learning technique was created to address binary classification issues. Based on the correlation between the target variable and the input features, it calculates the chance that a website falls into the legitimate or phishing class. Owing to its simplicity, computational efficiency, and interpretability, Logistic Regression serves as an effective

fective baseline classifier for phishing detection [8].

ii) *Decision Tree*: This rule-based classification algorithm divides the feature space recursively based on decision criteria that are obtained from the input attributes. Its hierarchical structure makes it easy to understand while highlighting the key characteristics that ultimately determine the classification. For managing structured phishing datasets, decision trees are an effective computational tool [6].

ii) *Random Forest*: This ensemble learning algorithm improves model robustness and predictive performance by combining several Decision Trees. Individuals are determined while the majority trees are randomly selected feature predictions. This ensemble strategy effectively reduces overfitting, improves generalization, and consistently delivers high classification accuracy in phishing detection tasks [3].

iv) *Support Vector Machine (SVM)*: This algorithm maximizes the margin between phishing and legitimate website classes to identify the best separating hyperplane. SVM performs well on high-dimensional datasets and has shown great performance in phishing website classification as a result of its ability to model complex decision boundaries [13].

v) *Gradient Boosting*: Gradient Boosting is a sequential ensemble learning algorithm that iteratively trains weak learners to fix the prediction errors of models that have already been built. Gradient Boosting creates extremely accurate prediction models that can capture intricate nonlinear relationships in phishing datasets by gradually minimizing the loss function. Although computationally more demanding than individual classifiers, it consistently achieves superior detection performance [9].

Because Random Forest and Gradient Boosting consistently outperform traditional classifiers, such as Logistic Regression and Decision Tree, on benchmark phishing datasets, achieving classification accuracies exceeding 99 percent, they were chosen for in-depth analysis. By combining independently trained decision trees, Random Forest improves robustness by lowering variance and minimizing overfitting. Gradient Boosting, on the other hand, gradually increases prediction accuracy by using optimized learning rates to sequentially correct the residual errors of previous models. Additionally, both algorithms offer intrinsic feature importance scores that enhance XAI methods and enable insightful feature-level interpretation of model predictions. These characteristics make Random Forest and Gradient Boosting particularly suitable for developing an accurate, robust, and interpretable phishing detection framework.

IV. PROPOSED METHODOLOGY

To achieve precise and comprehensible phishing detection, the suggested framework for phishing website detection combines ensemble ML, XAI, feature engineering, feature selection, and data preprocessing. The suggested system's general workflow is depicted in Figure 1. Dataset collection, data preprocessing, feature engineering, feature selection, model training, model evaluation, and model interpretation using LIME are the six main phases of the methodology.

A. Dataset Collection

The Kaggle. The combined dataset contains both legitimate and phishing website instances represented by structured attributes describing lexical, domain-based, and security-related characteristics. These characteristics offer a thorough depiction of website behavior for phishing detection and include URL properties, HTTPS status, domain age, and subdomain-related data. Table II presents a sample of the available data.

TABLE II. SAMPLE DATA

URL Length	HTTPS	Domain Age	Label
54	Yes	2 years	Legit
120	No	1 month	Phish

The collected datasets were subsequently preprocessed to improve data quality and ensure their suitability for ML model development.

B. Data Preprocessing

Prior to model training, the quality and consistency of the dataset are greatly enhanced by data preprocessing. Real-world datasets that are frequently consistently contain duplicate transformation transformation, and contain transformation. Proper preprocessing reduces noise in the data, increases model generalization, and improves classification performance [8].

The following preprocessing steps were performed:

- i) *Data Cleaning*: Records containing missing or null values were removed to ensure data completeness.
- ii) *Label Encoding*: Categorical attributes, such as HTTPS (Yes/No), were transformed into numerical representations (0/1) using the *LabelEncoder* function available in the Scikit-learn library.
- iii) *Feature Scaling*: Numerical features were standardized using *StandardScaler*, resulting in a mean of zero and a standard deviation of one. Feature scaling is particularly important for algorithms such as Logistic Regression and Support Vector Machine, whose performance depends on the relative magnitude of input features.
- iv) *Duplicate Removal*: Duplicate records were eliminated to prevent redundancy and reduce the possibility of data leakage during model training.
- v) *Training–Testing Split*: The dataset was partitioned into training and testing subsets using stratified sampling with an 80:20 ratio to preserve the original class distribution in both subsets.

C. Feature Engineering and Extraction

Because the quality of extracted features directly affects model performance, feature engineering is an essential step in the identification of phishing websites [7]. Choosing features. The goal of feature selection was to find the most instructive characteristics while removing unnecessary or insignificant features. Classification performance is enhanced, computational complexity is decreased, and multicollinearity is reduced when only pertinent features are chosen. There were two complementary feature selection techniques. Two complementary feature selection techniques were employed:

- i) *Correlation Analysis*: To lessen multicollinearity, highly correlated features (correlation coefficient larger than 0.90) were investigated. For instance, there was a strong positive correlation ($r = 0.87$) between URL Length and the number of special characters; as a result, only URL Length was kept.
- ii) *Information Gain*: Features were ranked according to their contribution to reducing the entropy of the target variable. Features were excluded from information. Furthermore, the Random Forest classifier's intrinsic feature importance scores revealed that the three most significant predictors were HTTPS Usage, Domain Age, and URL Length. It was based on the information. In Table III, the top three features are displayed.

TABLE III. TOP FEATURES BY RANDOM FOREST IMPORTANCE SCORE

Feature Name	Importance Score
HTTPS Usage	0.24
Domain Age	0.19
URL Length	0.16

D. Model Training

The dataset was splitted into training and testing subsets using a stratified 80:20 split after feature selection in order to maintain the original class distribution. Hyperparameter optimization was subsequently performed using *GridSearchCV* with five-fold cross-validation to determine the optimal parameter settings for each ML algorithm. A minimum sample split of five, a maximum tree depth of twelve, and 100 estimators made up the ideal configuration for the Random Forest classifier. The ideal setup for gradient boosting had 100 estimators, a learning rate of 0.1, and a maximum tree depth of 8. The final models were trained using these optimized hyperparameters, guaranteeing strong and dependable performance during experimental evaluation.

E. Model Evaluation

Table IV summarizes the evaluation metrics used to assess the predictive performance of the proposed framework, including accuracy, precision, recall, F1-score, AUC-ROC, training time, inference time, false positive rate (FPR), and false negative rate (FNR).

A set of representative lexical and domain-

based features that successfully distinguish phishing websites from trustworthy ones are extracted by the suggested framework. These features give the ML models discriminative information by capturing important traits linked to malicious URLs.

URL Length: Phishing URLs frequently have more than 75 characters, which allows attackers to hide malicious parameter and redirection links.

'@' Symbol: This symbol is frequently linked to URL redirection and is thought to be a sign of questionable website

activity. *HTTPS Usage*: HTTPS usage is a crucial security-related feature because phishing websites often use self-signed or expired certificates or lack valid SSL/TLS certificates. *Domain Age*: Phishing activity is often linked to recently registered domains, especially those that are less than six months old.

Subdomain Count: A lot of subdomains are frequently used to mask a website's true destination and boost the legitimacy of malicious URLs.

Suspicious Path Keywords: Phishing attempts are often linked to URL paths that contain keywords like login, secure, bank, verify, account, and update.

URL Entropy: URL entropy helps identify automatically generated or obfuscated URLs frequently used in phishing attacks by

measuring the randomness of URL strings. *Hostname Length*: Phishing activity may be indicated by hostnames longer than 50 characters, which are commonly found in fraudulent URLs.

IP Address Usage: Phishing may be indicated by the use of a direct IP address rather than a fully qualified domain name.

When combined, these lexical and domain-based features capture the distinctive qualities of phishing websites and offer a strong feature representation that helps the ML models distinguish between phishing and authentic webpages. These qualities are especially crucial for cybersecurity applications, as both false alarms and missed detections can have serious operational repercussions. Although Logistic Regression and Support Vector Machine yielded competitive results, their longer inference times and higher false positive rates (0.07 and 0.13, respectively) made them less suitable for real-time phishing detection. Both ensemble models met the latency requirements of real-time phishing detection systems by maintaining inference times of only

4.3 ms and 5.8 ms per URL, respectively, despite requiring moderate training times (45.2 s for Random Forest and 68.9 s for Gradient Boosting). Overall, the results demonstrate that the ensemble learning models are the best classifiers for the suggested phishing detection framework because they offer the best balance between predictive accuracy, computational efficiency, and operational reliability.

F. Phishing detection using Explainable AI (XAI).

LIME was used to produce feature-level explanations for individual predictions in order to improve the suggested phishing detection framework's interpretability and transparency.

LIME enhances the practical usability of the system in several ways:

i) *Transparency*: It identifies the features that most strongly influence each prediction rather than providing only the final classification outcome. For example, a phishing prediction may be primarily influenced by attributes such as URL Length > 100, HTTPS = No, and Domain Age = 7 days.

ii) *Actionable Insights*: Feature-level explanations enable cybersecurity analysts to verify whether the model's reasoning is consistent with domain knowledge, thereby increasing confidence in the prediction and facilitating informed decision-making.

iii) *Model Debugging and Fairness*:

When incorrect classifications occur, LIME identifies the influential features responsible for the prediction, enabling developers to refine feature engineering strategies, adjust model parameters, or incorporate additional representative training data.

iv) *User Confidence*: Transparent explanations improve user trust by clearly communicating the rationale behind phishing alerts. Such interpretability is particularly valuable in cybersecurity applications, where users and analysts must make timely decisions based on model recommendations.

Consequently, the incorporation of LIME enables the proposed framework to combine the predictive strength of ensemble learning with transparent and interpretable decision-making, thereby improving both detection reliability and user trust.

TABLE V. PERFORMANCE METRIC COMPARISONS

Algorithm	Accuracy	Precision	Recall	F1-Score	AUC-ROC	Training Time (s)	Inference Time (ms)	False Positive Rate	False Negative Rate
Logistic Regression	0.89	0.87	0.91	0.89	0.91	12.4	2.1	0.13	0.09
Decision Tree	0.94	0.92	0.95	0.93	0.95	8.7	1.8	0.08	0.05
Random Forest	0.99	0.97	0.99	0.98	0.99	45.2	4.3	0.01	0.01
SVM	0.95	0.93	0.96	0.94	0.96	156.3	12.7	0.07	0.04
Gradient Boosting	0.99	0.98	0.99	0.99	0.99	68.9	5.8	0.01	0.01

V. SYSTEM ARCHITECTURE

Proposed phishing website detection framework architecture is illustrated in Figure 1.

To enable precise, dependable, and comprehensible phishing website detection, the framework uses a sequential processing pipeline that combines feature extraction, data preprocessing, ensemble machine learning, and artificial intelligence. The framework first receives a website URL as input, from which pertinent lexical, structural, and domain-based characteristics are extracted. URL length, HTTPS usage, domain age, subdomain count, hostname characteristics, and other indicators linked to phishing are some of these features. To guarantee data quality and improve model performance, the extracted features are then pre-processed through data cleaning, label encoding, feature scaling, and duplicate removal. The trained ensemble ML model is then given the processed feature set to classify. The model categorizes the input website as either phishing or legitimate based on the learned feature patterns.

To enhance the interpretability of the prediction, LIME generates feature-level explanations that identify the attributes contributing most significantly to the classification, enabling cybersecurity analysts to understand and validate the model's decisions. Lastly, the framework shows the user the predicted class along with the explanation that goes with it. The suggested architecture provides accurate and transparent phishing detection by combining the interpretability offered by XAI with the predictive power of ensemble ML. This makes it ideal for real-world cybersecurity applications that need dependable and understandable decision support.

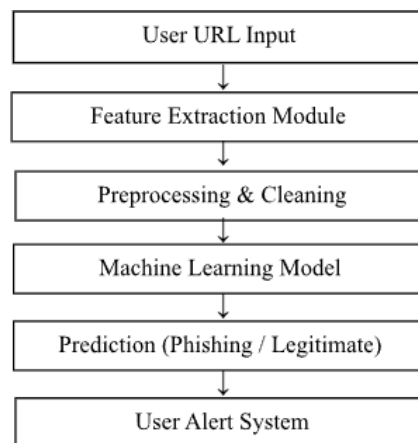


Fig.1. Architecture of the proposed Explainable AI-enabled phishing website detection framework.

VI. EXPERIMENTAL RESULTS

The proposed phishing website detection framework was experimentally evaluated using the combined UCIML Repository and Kaggle datasets. The performance of the selected ML algorithms was assessed using standard classification metrics, such as accuracy, precision, recall, F1-score, AUC-ROC, training time, inference time, false positive rate (FPR), and false negative rate (FNR). In addition, XAI was employed to assess the interpretability of the proposed framework through LIME.

A. Performance Evaluation

Figure 2 compares the classification accuracy achieved by the evaluated ML algorithms. Among the five classifiers, Gradient Boosting and Random Forest consistently delivered the highest predictive performance, each achieving an accuracy of 99%. Their superior performance highlights the effectiveness of ensemble learning in distinguishing phishing websites from legitimate ones while maintaining high classification reliability.

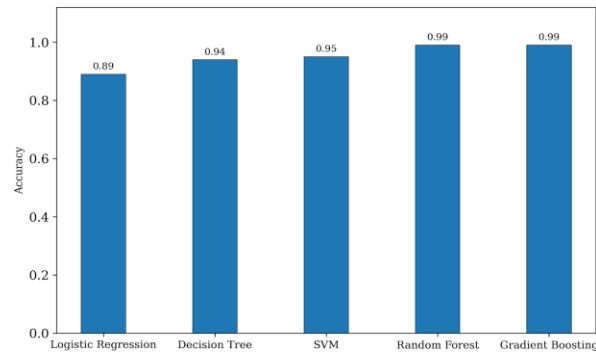


Fig.2. Classification accuracy of various machine learning algorithms.

B. ROC Curve Analysis (ROC)

ROC analysis was used to evaluate the classification models' capacity for discrimination. Gradient Boosting and Random Forest achieved the highest Area Under the ROC Curve (AUC-ROC) values of roughly 0.99, as shown in Figure 3, demonstrating outstanding classification performance across various decision thresholds. The superior predictive efficacy of the ensemble learning algorithms was confirmed by the comparatively lower discriminative ability of Logistic Regression.

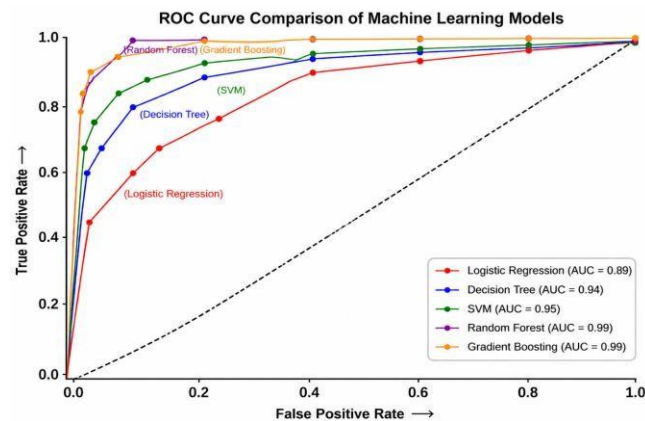


Fig.3. ROC curve comparison of the evaluated machine learning models.

C. Explainable AI Using LIME

LIME was added to increase model transparency by offering feature-level justifications for each prediction. Figure 4 illustrates how LIME determines each feature's relative contribution to a website's classification as either legitimate or phishing.

While negative contribution bars support the categorization of a legitimate website, features represented by positive contribution bars raise the probability of a phishing prediction. Additionally, each bar's magnitude indicates the corresponding feature's relative impact on the model's prediction. These justifications improve the framework's interpretability, allowing cybersecurity analysts to verify model choices and boosting transparency in the detection procedure.

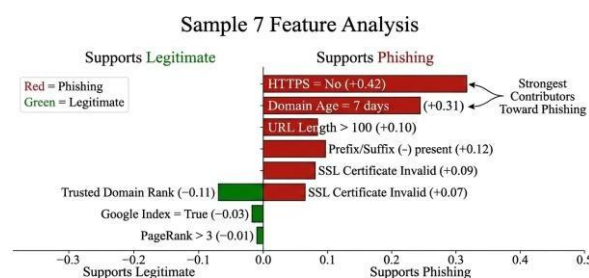


Fig.4. LIME-based feature explanation for phishing website detection.

D. Computational Complexity Analysis

The computational efficiency of the proposed framework was evaluated by analyzing its training complexity, inference performance, and explanation overhead. The Random Forest classifier exhibits a training complexity of approximately $O(N \cdot d \cdot \log N)$, where N denotes the number of training samples and d represents the number of features, requiring 45.2 s for model training.

In contrast, Gradient Boosting constructs decision trees sequentially, leading to a computational complexity of $O(N \cdot d \cdot T)$, where T denotes the number of boosting iterations, with a training time of 68.9 s.

The high runtimes speed of both developed models was demonstrated by inference results. With a throughput rate of roughly 172 URLs per second, the Gradient Boosting model was able to classify a single URL in just 5.8 milliseconds, making it suitable for real-time phishing classification. Each prediction had a latency of 150 milliseconds when LIME was added to the inference process. For post-inference operations, this latency is acceptable, but large-scale deployments might require additional enhancements. During the execution of their processes, neither of the trained models used excessive amounts of memory. The Gradient Boosting model required approximately 312 MB of memory, whereas the Random Forest model required about 245 MB.

E. Limitations

Despite the fact that this suggested strategy produced some encouraging outcomes, a number of drawbacks must be noted.

Initially, the experiment for this work was conducted using open-sourced data from Kaggle and the UCIML repository. This type of data may not accurately reflect the constantly evolving phishing attack techniques of today, including script-based phishing, CAPTCHA bypassing techniques and other sophisticated phishing attacks. It may, therefore, be necessary to re-train this suggested method periodically with new data. In the future, phishing attacks may use more sophisticated URL obfuscation methods or dynamically generated Web pages. In such cases, a new set of 18 features optimized for high classification accuracy could be extended with more lexical, behavioral, and/or contextual features. The following features would be useful in sustaining the effectiveness of the model: Besides, each prediction takes approximately 150 ms more for LIME. This latency may be sufficient for an offline analysis and decision support application, and can impact the scalability of deployment scenarios that may have high throughput or high latency requirements. The second limitation stems from the fact that some of the lexical features employed in the suggested architecture are primarily based on URLs in the English language. This may limit its ability to generalize to multilingual phishing attacks. However, the architecture has been put to the test in an experimental setup. However, there might be some issues in the real world, such as the dynamic nature of phishing attacks, network latency, and a changing computing environment. The suggested architecture has some drawbacks that would be a major topic for further study.

VII. CONCLUSION

This paper presented an XAI-enabled phishing website detection framework that integrates ensemble ML with LIME to achieve both high predictive performance and transparent decision-making. By combining Gradient Boosting and Random Forest classifiers with feature-level explanations, the proposed framework effectively addresses the limitations of conventional black-box phishing detection models, enabling cybersecurity analysts to interpret and validate individual predictions. An optimized feature selection strategy reduced the original feature set to 18 highly discriminative URL- and domain-based attributes, improving computational efficiency while maintaining excellent classification performance. Experimental evaluation on the combined UCIML Repository and Kaggle phishing datasets demonstrated that the proposed framework obtained an accuracy of 99.2%, outperforming traditional ML algorithms while providing interpretable explanations for every prediction. The inclusion of LIME contributed to increased clarity of the designed framework. This is achieved by providing feature-level explanations for the determinants of every classification output. The output explanations contribute to increased interpretability of the model predictions. They also provide a better foundation for cybersecurity analysts to understand and analyze the inference process performed by the framework. From a cybersecurity perspective, the framework's design adopts a pragmatic approach to balance the interpretability and predictiveness of ML models. It was possible to create a solution for the

trustworthy classification of phishing attacks that could offer justifications for its choices. This is done by combining XAI techniques with high-performing ensemble learning models. Transformer architectures are intended to be used in future research to improve resistance to complex phishing attacks while preserving interpretability as a crucial component of the solution. The use of incremental and adaptive learning techniques will also be investigated in order to enable the models to be dynamically updated to the constantly shifting threat landscape. The impact of XAI on cybersecurity analysts' decision-making, response time, and trust will also be investigated further. Additionally, multilingual phishing detection and extensive real-world validation will be the subject of future research.

REFERENCES

- [1] Safi, A., & Singh, S. (2023). A systematic literature review on phishing website detection techniques. *Journal of King Saud University-Computer and Information Sciences*, 35(2), 590-611.
- [2] Dutta, A. K. (2021). Detecting phishing websites using machine learning technique. *PloSone*, 16(10), e0258361.
- [3] Wei, Y., & Sekiya, Y. (2022). Sufficiency of ensemble machine learning methods for phishing websites detection. *IEEE Access*, 10, 124103-124113.
- [4] Zieni, R., Massari, L., & Calzarossa, M. C. (2023). Phishing or not phishing? A survey on the detection of phishing websites. *IEEE Access*, 11, 18499-18519.
- [5] Duy, P. T., Minh, V. Q., Dang, B. T. H., Son, N. D. H., Quyen, N. H., & Pham, V. H. (2024). A study on adversarial sampler resistance and defense mechanism for multimodal learning-based phishing website detection. *IEEE Access*, 12, 137805-137824.
- [6] Rao, R. S., & Pais, A. R. (2019). Detection of phishing websites using an efficient feature-based machine learning framework. *Neural Computing and applications*, 31(8), 3851-3873.
- [7] Kustiawan, Y. A., & Ghauth, K. I. (2025). Feature Engineering for Phishing Website Detection Using Machine Learning: A Systematic Review. *IEEE Access*, 13, 192080-192104.
- [8] Nayak, G. S., Muniyal, B., & Belavagi, M. C. (2025). Enhancing phishing detection: a machine learning approach with feature selection and deep learning models. *IEEE Access*.
- [9] Zara, U., Ayyub, K., Khan, H. U., Daud, A., Alsahfi, T., & Ahmad, S. G. (2024). Phishing website detection using deep learning models. *IEEE Access*, 12, 167072-167087.
- [10] Tang, L., & Mahmoud, Q. H. (2021). A deep learning-based framework for phishing website detection. *IEEE Access*, 10, 1509-1521.
- [11] Ahmad, S., Zaman, M., Al-Shamayleh, A. S., Ahmad, R., Abdulhamid, S. I. M., Ergen, I., & Akhunzada, A. (2024). Across the spectrum in-depth review AI-based models for phishing detection. *IEEE Open Journal of the Communications Society*, 6, 2065-2089.
- [12] Rafsanjani, A. S., Kamaruddin, N. B., Behjati, M., Aslam, S., Sarfaraz, A., & Amphawan, A. (2024). Enhancing malicious URL detection: A novel framework leveraging priority coefficient and feature evaluation. *IEEE Access*, 12, 85001-85026.
- [13] Dsouza, D. J., Rodrigues, A. P., & Fernandes, R. (2024). Multimodal comparative analysis on execution of phishing detection using artificial intelligence. *IEEE Access*, 12, 163016-163041.
- [14] Setu, J. H., Halder, N., Islam, A., & Amin, M. A. (2025). RSTHFS: A rough set theory-based hybrid feature selection method for phishing website classification. *IEEE Access*.
- [15] Pillai, M. J., Remya, S., Devika, V., Ramasubbareddy, S., & Cho, Y. (2023). Evasion attacks and defense mechanisms for machine learning-based web phishing classifiers. *IEEE Access*, 12, 19375-19387.
- [16] Wangchuk, T., & Gonsalves, T. (2025). Multimodal Phishing Detection on Social Networking Sites: A Systematic Review. *IEEE Access*.
- [17] Kavya, S., & Sumathi, D. (2025). Multimodal and temporal graph fusion framework for advanced phishing website detection. *IEEE Access*.
- [18] E. Chand, "Phishing Website Detector Dataset," Kaggle. [Online]. Available: <https://www.kaggle.com/datasets/eswarchandt/phishing-website-detector>
- [19] UCIMachineLearningRepository, "Phishing Websites Dataset." [Online]. Available: <https://archive.ics.uci.edu/ml/datasets/phishing+websites>