

URLGuard: A Holistic Hybrid Machine Learning Approach for Phishing Detection

Pradip M. Paithane

VPKBIET Engineering College, Baramati, Dist. Pune, Maharashtra, India

Email: paithanepradip@gmail.com

ORCID iD: <https://orcid.org/0000-0002-4473-7544>

Received: 26 July, 2024; Revised: 06 September, 2024; Accepted: 08 January, 2025; Published: 08 April, 2025

Abstract: The fast growth of Internet technology has significantly changed online users' experiences, while security concerns are becoming increasingly overpowering. Among these concerns, phishing stands out as a prominent criminal activity that uses social engineering and technology to steal a victim's identification data and account information. According to the Anti-Phishing Working Group (APWG), the number of phishing detections increased by 46 in the first quarter of 2018 compared to the fourth quarter of 2017. So to overcome these situations below paper introduces a phishing detection system using a hybrid machine learning approach based on URL attributes. It addresses the growing threat of phishing attacks that exploit email manipulation and fake websites to deceive users and steal sensitive data. The study employs a phishing URL dataset with over 11,000 websites, extracted from a reputable repository. After pre-processing, a hybrid machine learning model, which includes Decision Tree, Random Forest, and XGB is employed to safeguard against phishing URLs. The proposed approach undergoes evaluation with key metrics such as precision, accuracy, recall, F1-score, and specificity. Results demonstrate that the proposed method surpasses other models, achieving superior accuracy and efficiency in detecting phishing attacks.

Index Terms: Anti-Phishing Working Group (APWG), Decision Tree, and Random Forest, and XGB, Hybrid Machine Learning

1. Introduction

In the present interconnected digital environment, the internet functions not only as a conduit facilitating the vast exchange of information but also as a platform utilized for engaging in malicious activities. One of the most widespread forms of these malicious activities is phishing, which entails the use of deceptive strategies by cybercriminals with the aim of misleading individuals into revealing confidential information such as their login credentials, financial particulars, or personal data. Phishing attacks commonly manifest in the guise of seemingly authentic emails, messages, or websites that are meticulously crafted to dupe users into believing they are interacting with a reputable and trustworthy source. Despite the fact that the concept of phishing is not novel, its level of sophistication and prevalence have undergone significant transformations in recent times, thereby presenting formidable obstacles to individuals, enterprises, and institutions on a global scale. As reported by the Anti-Phishing Working Group (APWG), a staggering number exceeding 200,000 distinct phishing websites were identified solely in the initial quarter of the year 2023, thereby underscoring the magnitude and gravity of this perilous threat. Furthermore, it is imperative to highlight the significant financial ramifications associated with phishing attacks, which are truly astounding, leading to the annual loss of billions of dollars as a result of deceitful activities carried out through these cunning methods. The primary objective of this academic inquiry is to thoroughly investigate the complex realm of phishing websites, delving into their intricate operations, distinguishing features, and profound effects on the realm of cyber security. Through a comprehensive analysis of the strategies utilized by malicious actors to fabricate and disseminate phishing websites, along with the approaches for identifying and containing these risks, the primary aim of this scholarly endeavor is to enrich the comprehension of the dynamic landscape of online fraudulence. Through empirical analysis, case studies, and scholarly insights, this paper will provide valuable insights into the following key areas:

1. The anatomy of phishing websites: Understanding the design, structure, and functionality of phishing websites, including common tactics used to mimic legitimate entities and exploit human psychology.
2. Detection and classification techniques: Exploring the methodologies and technologies utilized to identify and categorize phishing websites, from heuristic analysis to machine learning algorithms.

3. Impacts and consequences: Assessing the multifaceted repercussions of phishing attacks on individuals, businesses, and society at large, including financial losses, reputational damage, and erosion of trust in online platforms.

4. Countermeasures and best practices: Examining proactive measures and defensive strategies aimed at combating phishing threats, ranging from user education and awareness campaigns to technical solutions such as email filtering and website authentication protocols.



Fig. 1. Phishing Detection Steps by applying AI solutions.

Objectives:

1. To create an ongoing detection system to spot phishing URLs.
2. To perform data preprocessing on phishing dataset(rows=11054, columns=33).
3. To introduce a hybrid model (DT+RF+XGB) for enhanced phishing detection.
4. Discuss and compare evaluation parameters to demonstrate the superiority of the proposed approach.

By shedding light on the intricate dynamics of phishing websites, this research endeavors to empower individuals and organizations with the knowledge and tools necessary to navigate the digital landscape safely and securely. Ultimately, the fight against phishing requires a concerted effort from all stakeholders, including cybersecurity professionals, policymakers, and end-users, to mitigate the risks posed by these insidious online threats. Phishing detection using hybrid machine learning models outlines the growing threat of phishing attacks in the digital era and highlights the need for advanced, accurate detection methods. It introduces the concept of hybrid machine learning, emphasizing its potential to enhance detection accuracy by combining the strengths of multiple algorithms. Phishing, often executed through deceptive URLs, exploits human vulnerability, aiming to extract sensitive information or deploy malicious payloads. This report delves into the development and functionality of a phishing URL detection system, shedding light on its importance in mitigating cyber risks and fortifying the defenses against malicious actors. As explore the intricacies of such a system, uncover the pivotal role it plays in ensuring a secure digital environment for users, ultimately contributing to the ongoing battle against cyber threats. Building a robust phishing URL detection system is crucial in safeguarding individuals and organizations from cyber threats. By developing and implementing such a system, contribute to the overall cybersecurity landscape, protecting users from potential financial losses, identity theft, and other malicious activities. Efforts can make a meaningful impact in creating a safer digital environment for everyone. Keep pushing forward for a more secure online world!

2. Related Work

2.1 Phishing detection using hybrid machine learning model

S. Raman Kumar Jog, Abdul Karim, Mobeen Shahroz, Khabib Mustofa, And Samir Brahim Belhaouar [1]:Machine Learning Models:Logistic Regression (LR), Support Vector Classifier (SVC), and Decision Trees (DT) combined into a hybrid model. In summary, the hybrid model (LSD) that has been suggested seeks to improve phishing detection efficiency and accuracy.

It makes use of grid search, hyper parameter optimization, and canopy feature selection. On the other hand, false positive or negative rates are not discussed, and the study does not provide any information on scalability for huge datasets. Phishing Detection Using Machine Learning Techniques Machine Learning Models[2] : Logistic Regression, Ada Booster, Random Forest[3], K-Nearest Neighbour's (KNN) [4], Neural

Networks, Support Vector Machines (SVM) [2], Gradient Boosting, Naive Bayes [3], and XGBoost. The study focuses on phishing prevention techniques, including email filtering, user education, and real-time machine learning

detection. It also gives us a detailed explanation of the workings of different machine learning models. However, it does not explore potential vulnerabilities or weaknesses in the machine learning models, and dataset details are omitted, impacting generalizability.

2.2 Phishing Detection Using Machine Learning

Machine Learning Models [2]: Logistic Regression, Ada Booster, Random Forest[4] , K-Nearest Neighbors (KNN) [5] , Neural Networks, Support Vector Machines (SVM) [1], Gradient Boosting, Naïve Bayes [5], and XGBoost. The study focuses on phishing Prevention techniques, including email filtering, user education, and real-time machine learning detection. It also gives us a detailed explanation of the workings of different machine learning models. However, it does not explore potential vulnerabilities or weaknesses in the machine learning models, and dataset details are omitted, impacting generalizability.

2.3 Phishing detection using random forest with NLP-features

Ozgur Koray Sahingoz, Ebubekir Buber, Onder Demir, Banu Diri[5] Machine Learning Models: Random Forest with NLP-Based Features. The study presents a high accuracy real-time anti-phishing system that uses a variety of machine learning (ML) algorithms and features for URL-based phishing detection. Nevertheless, issues pertaining to NLP-based features are not examined, and the performances of the seven classification algorithms are not thoroughly analyzed[6]. The Naive Bayes classification is a probabilistic machine learning method, which is not only straightforward but also powerful. Due to its simplicity, efficiency and good performance, it is preferred in lots of application areas such as classification of texts, detection of spam emails/intrusions, etc. It is based on the Bayes theorem, which describes the relationship of conditional probabilities of statistical quantities.

2.4 Phishing Email Detection Using Improved RCNN Model With Multilevel Vectors and Attention Mechanism

Yong Fang, Cheng Zhang, Cheng Huang, Liang Liu, And Yue Yang[4]Machine Learning Models[7]: Deep learning model named THEMIS(advanced version of RCNN). THEMIS enhances embedding for improved performance detection. This approach allows for modeling emails at multiple levels, including the email header, email body, character level, and word level simultaneously, enhancing the model's ability to capture relevant features[8]. However, the study notes that some phishing emails without an email header may reduce efficiency and accuracy. It is more accurate than previous methods, and it is also more efficient and robust.

1. The Bidirectional Long Short-Term Memory (Bi-LSTM) enhances the RCNN model. Next, an enhanced RCNN model is used to model the email at several levels. When as little noise as possible is introduced, the email's context can be better understood.

2. The email header and body are subjected to the attention mechanism, with varying weights allocated to each section. This allows the model to concentrate on more distinct and valuable data found in the email header and body.

3. On an unbalanced dataset, the THEMIS model presented in this work performs admirably. The accuracy is 99.848, and THEMIS's evaluation metrics are better than current detecting technology.

2.5 Hybrid ensemble feature selection framework for enhancing machine learning-based phishing detection systems

Table 1. Literature Survey Discussion about Machine Learning Algorithm used in Phishing Detection

Sr.No	Author and Year	Techniques/ Methodology	Advantages	Gaps
1	Abdul Karim,M.Shahroz, Khabib (2023)[7]	Machine Learning Models, Hybrid models (LR+SVC+DT)	It multiple machine learning algorithm which uses canopy feature selection and Grid Search Hyper parameter Optimization for enhancing accuracy and efficiency.	No discussion on false positive/negative rates and No insights on scalability for large datasets.
2	V.Shahrivari, Mahdi Darab (2020)[13]	Logistic regression, Ada booster, random forest, KNN, neural networks, SVM, Gradient boosting, and XGBoost	Phishing techniques include email filtering, user education, and real-time machine learning detection for prevention.	Dataset details omitted.
3	O.Koray,E.Buber, BanuDiri(2017)[3]	Using Random Forest with NLP-based features	A real-time anti-phishing system achieving 97.98 and is used Real-time.	Paper lacks detailed analysis of the seven classification algorithm.
4	Y.FANG,C.ZHANG (2019)[5]	Deep learning model named THEMIS(advance diversion of RCNN)	Enhanced Embedding in THEMIS.	Some phishing emails may not have an email header due to which efficiency and accuracy decreases.
5	K.Chiew,C.Tan, K.Tionga(2019)[4]	Hybrid Ensemble Feature Selection(HEFS),Random Forest Classifier	Using a real browser for feature extraction.	Fails to specify the phishing dataset used to benchmark HEFS.

Kang Leng Chiew, Choon Lin Tan, Kok Sheik Wong, Kelvin S.C. Yong, Wei King Tiong[9] Machine Learning Models: Hybrid Ensemble Feature Selection (HEFS), Random Forest Classifier[10]. The study uses a real browser for feature extraction, improving robustness. However, detailed accuracy results for baseline features with classifiers other than Random Forest are provided. The phishing dataset used for benchmarking HEFS is not specified. According to experimental findings, HEFS works best when combined with the Random Forest classifier, identifying phishing and authentic websites with 94.6% accuracy while utilizing just 20.8 % of the original characteristics. It consists of two phases: the first phase uses the CDF-g algorithm to generate primary feature subsets, which are then used in a data perturbation ensemble to produce secondary feature subsets. The second phase derives a set of baseline features from the secondary feature subsets using a function perturbation ensemble.

3. Material and Methods

3.1 DATASET

The dataset was gathered and saved as a CSV file from the well-known Kaggle dataset repository, which offers benchmark datasets for academic use. The collection included 33 attributes and 11054 items that were taken from over 11,000 websites. Some characteristics that help distinguish between phishing and legitimate website URLs are UsingIP, LongURL, ShortURL, Symbol@, Redirecting//, PrefixSuffix-, Sub-Domains, HTTPS, DomainRegLen, Favicon, NonStdPort, HTTPSDomainURL, RequestURL, AnchorURL, LinksIn-ScriptTags, and ServerFormHandler. As seen in Figure 2, the dataset was divided into two classes: phishing and legitimate. The dataset needs to be improved because it was in vector form. To prepare the dataset for preprocessing, the null values were eliminated. The entire dataset was preprocessed and then combined into a single corpus and put to use in further processing. The entire corpus was split into two parts: 70% for training and 30% for testing. The machine learning model was trained using 70% of the training data, while 30% of the data was kept for predictions. The performance of the suggested method was also assessed.

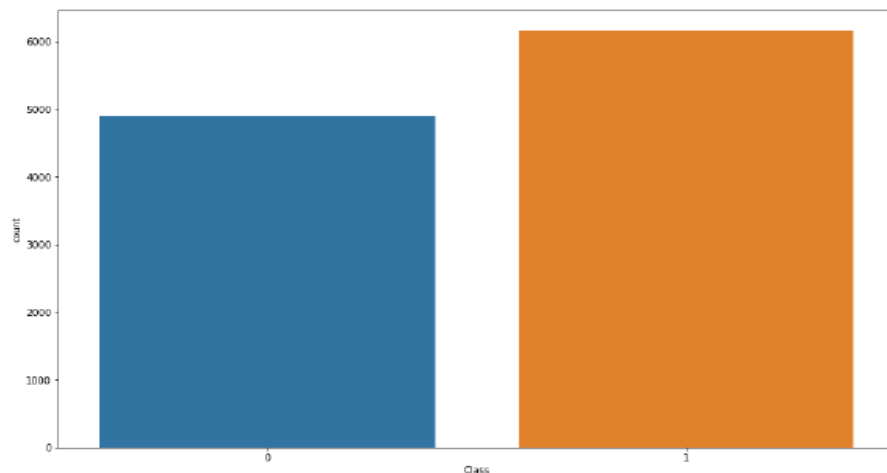


Fig. 2. Dataset Presentation According to Number Classes Phishing and Legitimate

3.1.1 Canopy-based feature selection

Centroid selection and canopy clustering are the two main processes in canopy centroid selection. Two distance thresholds, an inner threshold (T_2) and an outer threshold (T_1), where T_1 is greater than T_2 , create a canopy, which is a representation of a cluster. Data points are allocated to canopies in the clustering phase according to how close they are to the canopy center. Points that are located inside the inner threshold (T_2) are specifically allocated to the same canopy. Each canopy is represented by a centroid once the canopies have formed[11]. The mean or median of the feature values at each position inside the canopy is used to calculate this centroid. By using this technique, the centroids are guaranteed to correctly depict the data points' central tendency inside each canopy.

3.1.2 Cross fold validation

Cross-fold validation is a reliable technique for assessing model effectiveness. The dataset is partitioned into k fold sections of approximately similar size, usually $k = 5$ or 10 . The model is trained and evaluated k times as the procedure iterates across these folds. Every iteration uses a different fold as the test set and uses the rest of the folds for training. The model's dependability is increased by this method, which guarantees that every data point is used for both training and testing. A more reliable assessment of the model's overall performance is obtained by averaging or otherwise aggregating the performance measures from each fold after all iterations[12]. By reducing the bias and variation brought about by a single train-test split, this technique provides a thorough analysis of the model.

3.1.3 Grid Search

Grid search is an organized technique for fine-tuning hyperparameters by iteratively going through a predetermined list of possible combinations. This method assesses the performance of the model by doing cross-validation for every combination of hyperparameters. Next, the hyperparameter set with the best performance is chosen. Grid search was used for this project with the following hyperparameter configurations: 'xgb__max_depth' with values [3, 5, 7], 'dt__max_depth' with values [None, 5, 10], and 'rf__n_estimators' with values [50, 100, 200]. The most ideal hyperparameters were found by analyzing these combinations, which enhanced and strengthened the model's performance. This approach makes sure that the hyperparameter space is well explored, which improves the efficacy of the model.

For the purpose of the experiment I have prepared a suitable dataset by removing the null values, performing cross fold validation and grid search hyperparameter tuning.

- removing the null values: Removing null values from a dataset entails identifying and eliminating any entries that lack data. This process helps ensure data quality by preventing inaccuracies in analysis or modeling caused by missing information.
- cross fold validation: Cross-fold validation divides a dataset into multiple subsets, training the model on all but one and validating on the omitted subset. This process iterates, ensuring each subset serves as both training and validation data, providing robust model evaluation.
- Grid search hyperparameter tuning: Grid search hyperparameter tuning involves systematically searching through a specified subset of hyperparameter combinations to find the optimal configuration for a machine learning model. It exhaustively evaluates each combination to identify the best-performing set.

3.2 Applied Machine Learning Algorithms

Algorithms for machine learning are computational tools that let computers learn from data and come to conclusions or predictions. They fall into three categories: reinforcement learning, where agents learn to interact with surroundings to accomplish goals by trial and error; supervised learning, where models learn from labeled data; and unsupervised learning, where patterns are found from unlabeled data. A subset called deep learning uses sophisticated neural networks to interpret natural language and recognize images. Several models are combined in ensemble methods to improve performance. Numerous applications, including financial forecasts, driverless cars, medical diagnostics, and recommendation systems, are powered by these algorithms.

In the above paper to find phishing websites, used hybrid machine learning model containing decision tree, random forest and XGboost, which provided different accuracy for different feature engineering techniques.

3.2.1 Decision Tree

A decision tree [13] - a basic representation that classifies instances. A decision tree Constitutes of the following:

- Nodes: specic attributes' estimation is tested by nodes.
- Branches: they are the interface with following nodes or the leaf nodes and relatesto the result.
- Leaf nodes: Nodes that are terminal and anticipate the result.

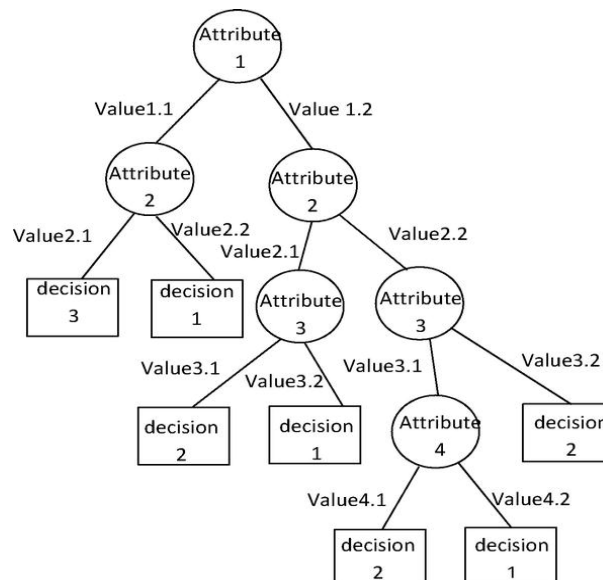


Fig. 3. General Representation of Decision Tree

In the figure 3, the decision tree starts with an initial question at the top node, labeled “Attribute 1”. Depending on the answer (Value 1.1 or Value 1.2), the data is split into two branches. This process is then repeated for each branch, using a different attribute as the test question at each node. The decision tree makes its classifications at the terminal nodes, which are labeled with the final decision.

For instance, consider the leftmost branch of the tree in the figure. If the answer to the question at the top node is “Value 1.1”, then the next question is “Attribute 2”. If the answer to that question is “Value 2.1”, the data is split again, following the question labeled “Attribute 3”. If the answer to that question is “Value 3.1”, then the decision tree reaches a terminal node labeled “decision 2”. This means that the instance belongs to class “decision 2”.

1. Entropy($E(S)$): Entropy measures the uncertainty or randomness in a dataset

$$E(S) = -\sum_{i=1}^C P_i \log_2(P_i) \quad (1)$$

2. Conditional Entropy($E(T,X)$): Conditional Entropy measures the amount of uncertainty in T (target variable) given X (predictor variable)

$$E(T, X) = -\sum_{c \in X} p(c).E(c) \quad (2)$$

3. Information Gain($IG(T,X)$): Information gain measures the reduction in entropy when a dataset is split by the values of a predictor variable

$$IG(T, X) = E(T) - E(T, X) \quad (3)$$

Where, S is set, c =the class or values in set, P_i =Probability of occurrence of elements i in set S_i , T is set, and X is an attribute or feature, $p(c)$ is the probability of occurrence of value c in set X

3.2.2 Random Forest

Random Forest is an ensemble learning method widely used for classification and regression tasks. It comprises decision trees trained on random subsets of the data, with randomness introduced both in bootstrapped sampling and feature selection. During prediction, individual tree outputs are either averaged (for regression) or combined through majority voting (for classification). This randomness helps prevent overfitting and improves robustness. Random Forest handles high-dimensional data well and provides feature importance scores. However, it can be computationally expensive and may not perform optimally on highly imbalanced datasets. Its simplicity and effectiveness make it a popular choice in various domains, including fishing detection. By leveraging its ensemble nature and randomization techniques, Random Forest can contribute to hybrid machine learning models aimed at enhancing fishing detection system accuracy and robustness[14].

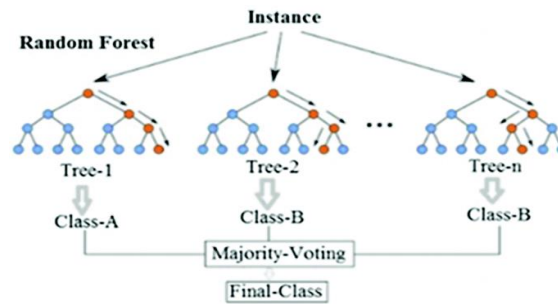


Fig. 4. Random Forest Classifier

The diagram illustrates a random forest classifier, which is a machine learning technique used for classification. It works by creating a collection of decision trees, each of which uses a random subset of features. New data points are then passed through each tree, and the most frequent class prediction is chosen as the overall prediction.

The left side of the diagram shows a random forest with multiple decision trees. Each tree is created using a random subset of features from the available data. The right side shows how a new instance is classified. The instance is passed through each tree in the forest, and each tree makes a prediction about the class of the instance. The final classification is determined by a majority vote of the trees[15].

In essence, random forests improve the accuracy and stability of classification models by reducing the variance of decision trees.

1. Ensemble Method: The ensemble model aggregates the predictions of multiple base models to make a final predication

$$F(X_t) = \frac{1}{B} \sum_{i=0}^B F_i(X_t) \quad (4)$$

Where, $F(X_i)$ is the output of the random forest for the input X_i , B is the number of trees in the random forest, $F_i(X_i)$ is the output of the i^{th} tree in the forest for input X_i .

3.2.3 XGBoost

Extreme Gradient Boosting, or XGBoost for short, is a sophisticated version of gradient boosting machines. It is well known for its effectiveness in supervised learning tasks as well as its scalability and efficiency[16]. XGBoost sequentially constructs an ensemble of weak prediction models, usually decision trees. By optimizing a predetermined objective function—typically a mix of the loss function and regularization term—each new model fixes the flaws of its predecessors. XGBoost, in particular, presents regularization methods including shrinkage (learning rate) and feature column subsampling to avoid overfitting. In addition, it makes effective use of hardware resources by parallelizing training inside a distributed computing architecture. Because of its versatility in solving classification, regression, and ranking problems, XGBoost is a well-liked option for both real-world and data science applications where speed and accuracy are crucial[17].

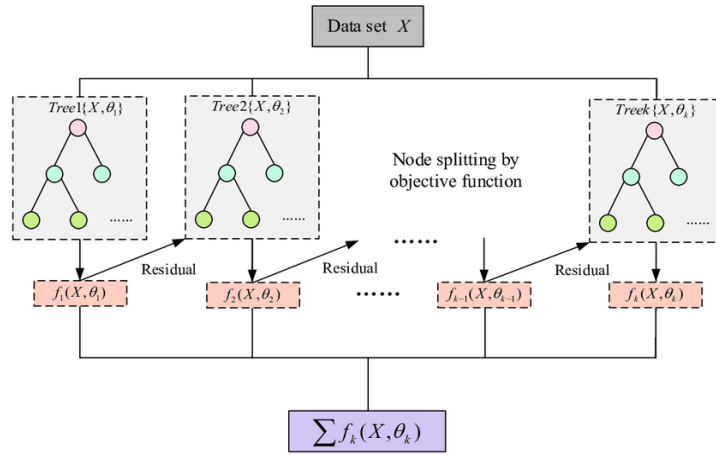


Fig. 5. XGBoost Classifier

A dataset with the label "Data Set" is at the top of the tree. Then, depending on the values of the characteristics in the data, it divides into many branches. The example shows that an unknown feature (X) with a value of 3 provides the basis for the first split. Different results are represented by each branch according to the split decision. "Tree1{X,03}" is reached via the left branch, and "Tree2{X,02}" by the right branch. This process continues until the data reaches a leaf node, which is shown by the rectangles with the labels " $f(X,0)$ " and " $f(X,02)$ " at the bottom. These leaf nodes hold the forecast of the model for that specific data route. The total of the predictions for each leaf is indicated by the phrase " $\Sigma(X,0)$ " in the bottom right[18].

1. Objective Function:

$$Objective = \sum_{i=1}^n Loss(\hat{y}_i, y_i) + \sum_{k=1}^K \Omega(f_k) \quad (5)$$

Where, n is the number of samples, $\hat{y}_i$ is the predicted value for the i^{th} sample, y_i is true label for the i^{th} sample, K is the parameter representing some set of values.

2. Gradient Boosting[19]

$$\hat{y}(t) = \sum_{k=1}^t f_k(X) \quad (6)$$

Where, t is the iteration or boosting round, $f_k(X)$ is the prediction of the K^{th} model at input X

3. Regularization:

$$\Omega(f_k) = \gamma T + \frac{1}{2} \lambda \sum_{j=1}^T W_j^2 + \alpha \sum_{j=1}^T |W_j| \quad (7)$$

Where, T is the number of leafs in the tree, W_j is the weight assigned to the j^{th} leaf, γ, λ and α are regularization parameters.

3.2.4 Proposed Hybrid Model(RF+DT+XGB)

Utilizing the advantages of each approach, a hybrid model that combines XGBoost (XGB), Random Forest (RF), and Decision Trees (DT) enhances prediction performance in a variety of areas. Known for its gradient boosting architecture, XGBoost minimizes overfitting by using regularization techniques like shrinkage and feature subsampling. It is an excellent tool for managing complicated relationships and huge datasets. By combining predictions from several decision trees trained on various data subsets, Random Forest, an ensemble of decision trees, offers resilience against noise and outliers. Decision trees are helpful in comprehending feature importance and linkages within the data because of their interpretability and simplicity. XGBoost functions as the main learner in the hybrid model, identifying intricate patterns in the data. By adding more variety to predictions, Random Forest enhances XGBoost and lowers variance while enhancing generalization. Decision trees serve as interpretable elements that facilitate feature analysis and model comprehension [20].

Depending on the objective, the hybrid model uses voting or averaging to mix the Decision Tree, Random Forest, and XGBoost outputs during prediction. Through the use of an ensemble technique, which balances the strengths and weaknesses of each member, a predictive model that is reliable and accurate is produced that can be used to a variety of machine learning tasks, such as regression, classification, and anomaly detection. The method of developing a machine learning model to identify phishing URLs is shown in the diagram.



Fig. 6. Architecture Diagram of Proposed Hybrid Method

The "URL Phishing" dataset serves as the process's starting point. It's possible that some of the URLs in this dataset have been flagged as phishing or authentic. The process of "Null Values Removal" is applied to the data, implying that any missing values in the dataset are being eliminated prior to model training. "Feature Vectors" is the next stage, when the data is converted into a format that can be interpreted by the machine learning model. This probably entails taking characteristics from the URLs, such the length of the domain name or the existence of specific keywords. "Training Data" (70%) and "Testing Data" (30%) are the two sets of data that are created after that. The machine learning model is taught using the training data, and its performance is assessed using the testing data. It is clear from the statement "Machine Learning Models" that the training data is being used to train several models. Although the exact models aren't stated, the language that follows implies that they are ensemble models that integrate the predictions of three distinct machine learning models: XGBoost, Random Forest, and Decision Tree. The models' predictions may be combined in two ways: "Soft Voting" and "Hard Voting." The testing data is used to assess the models once they have been trained. The graphic does not display the model's performance on the testing data. Lastly, a new URL is put through the trained model to predict whether it's a phishing site or not. The model outputs a classification: "Phishing", "Not Phishing," or "Not Determined".

Algorithm 1: Hybrid Machine Learning Model(RF+DT+XGB)

INPUT: Training data X_{train}, y_{train} , Test Data X_{test}, y_{test}
OUTPUT: Predicated class Label

Train_and_test_hybrid_model($X_{train}, y_{train}, X_{test}, y_{test}$, num of epochs, param_grid):
 initialize rf_model ← RandomForestClassifier()
 initialize xgb_model ← XGBClassifier()
 initialize dt_model ← DecisionTreeClassifier()
 initialize estimators ← [('rf', rf_model), ('xgb', xgb_model), ('dt', dt_model)]
 initialize best_model ← None
 for epoch ← 1 to num_of_epochs do
 initialize
 grid_search ← GridSearchCV(estimator=VotingClassifier(estimators), param_grid=param_grid, cv=5)
 grid_search.fit(X_{train}, y_{train})
 best_model ← grid_search.best_estimator_.fit(X_{train}, y_{train})
 validation_results ← best_model.evaluate(X_{test}, y_{test})
 print(Accuracy: validation_results)
 end
return best_model.predict(X_{test})

This algorithm implements a hybrid machine learning model combining Random Forest (RF), Decision Tree (DT), and XGBoost (XGB) classifiers. It iteratively trains the model for a specified number of epochs, optimizing hyperparameters through grid search cross-validation. The algorithm selects the best-performing model based on validation results and evaluates it on the test set. Finally, it returns the predicted class labels for the test data using the selected best model.

3.3 Evaluation Parameter

A number of evaluation criteria must be used to assess machine learning performance. Results from the machine-learning system are given as forecasts. The number of accurate and inaccurate predictions the model makes in both legitimate and phishing classes is measured by the evaluation parameters. There were several parameters used, including the F1-score, recall, specificity, accuracy, and precision[21].

Model performance is measured by accuracy, which is expressed as the number of correct predictions the model produces, as indicated by Equation

$$Accuracy = \frac{(TP + TN)}{(TP + TN + FP + FN)} \quad (8)$$

The evaluation parameter known as precision is utilized to analyze the models; it indicates the frequency at which a classifier stays accurate when we wish to forecast the positive class. Precision expresses the degree to which the model classifies the phishing URLs and gauges the positive rate, or the degree to which the model predicts the positive values[22]. With regard to precision, each classifier fared well.

$$Precision = \frac{TP}{TP + FP} \quad (9)$$

A metric for analyzing classification models that indicates how many times the model correctly identified out of all the possible positive labels. The classifier accurately identifies the classifier for both sorts of URLs in order to predict both phishing and authentic URLs[23].

$$Recall = \frac{TP}{TP + FN} \quad (10)$$

The F1score is the harmonic mean of precision and recall, where the F1 score reaches its best value (perfect precision and recall)[24]. The general formula is as follows:

$$F1Score = \frac{2 * Precision * Recall}{Precision + Recall} \quad (11)$$

Therefore, an F1 score is required when asking about the balance between precision and recall[25,26].

4. Result and Discussion

The internet is a vast network-based industry full of hackers, attackers, or cybercriminals. Civilians, businessmen, industries, and every market that consists of the Internet and networks need security to prevent phishing and provide protection to their customers, as well as to their own system safety. The methodology proposed in this study was successfully implemented as a prototype using a dataset comprising phishing and legitimate URLs. These experiments are carried out using many machine learning algorithms that are discussed separately in each heading to evaluate and illustrate the effects of the machine learning algorithms that are given below.

The experimental results showcase the functionality of the implemented hybrid model for phishing website. The below experimental results collectively demonstrate the efficiency of the implemented hybrid model in accurately detecting, segmenting, and recognizing phishing websites.steps for exact methods are:

- Data Collection and Preprocessing: Gather a dataset of labeled examples of phishing and legitimate websites. Features could include URL length, presence of HTTPS, domain age, etc.
- Feature Engineering: Extract relevant features from the dataset that could help distinguish between phishing and legitimate websites.
- Model Selection: Choose the individual models to include in the hybrid ensemble. In this case, Random Forest, Decision Trees, and XGBoost are selected.
- Ensemble Construction: Train each of the selected models (RF, DT, XGB) on the preprocessed dataset. Combine the predictions of the individual models using Voting.
- Model Evaluation: Evaluate the performance of the hybrid model using appropriate metrics such as accuracy, precision, recall, F1-score, and ROC-AUC.
- Model Deployment: Once satisfied with the model's performance, deploy it into a production environment for real-world phishing website detection. Continuously monitor the model's performance and update it as necessary with new data or improved techniques.

Upton this step, the expected results are retrieves the status of the CNG tank maintenance. Further we have additionally made the system adaptable to minute errors. If the model fails to recognize some of the character of the plate, then the system checks for some similar plates if available. It prints the similar plates with the similarity percentage. Henceforth, the system retrieves the information accurately.

Table 2. Experiment Result of Decision Tree

Max. Depth	Accuracy	Precision	Recall	Specificity	F1-Score
10	0.95	0.95	0.97	0.93	0.96
20	0.97	0.97	0.97	0.96	0.97
30	0.97	0.97	0.98	0.96	0.97
40	0.97	0.97	0.98	0.96	0.97
50	0.97	0.97	0.98	0.96	0.97

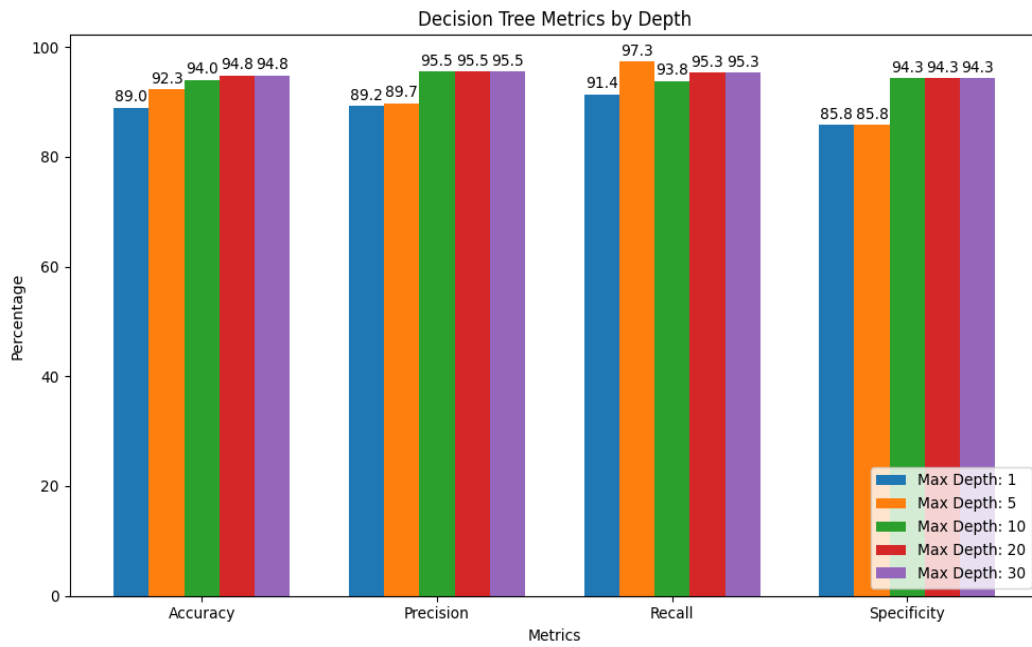


Fig. 7. Experiment Result of Decision Tree Compare with Evaluation Parameter

Table 3. Experiment Result of Random Forest

Max. Depth	Accuracy	Precision	Recall	Specificity	F1-Score
10	0.95	0.95	0.97	0.93	0.96
20	0.97	0.97	0.97	0.96	0.97
30	0.97	0.97	0.98	0.96	0.97
40	0.97	0.97	0.98	0.96	0.97
50	0.97	0.97	0.98	0.96	0.97

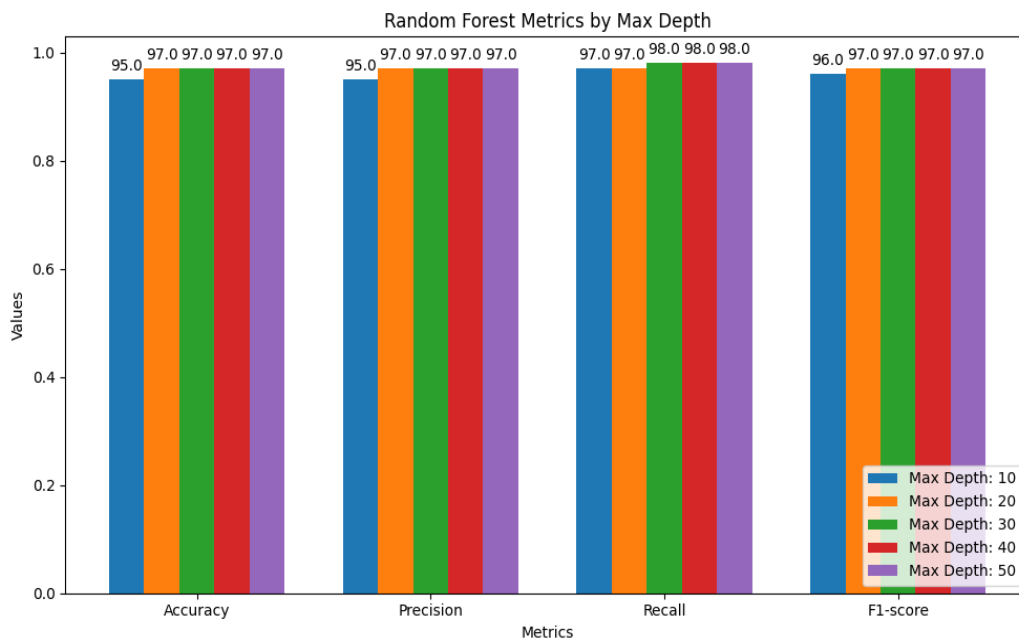


Fig. 8. Experiment Result of Random Forest Compare with Evaluation Parameter

Table 4. Experiment Result of XGBoost

Max. Depth	Accuracy	Precision	Recall	Specificity	F1-Score
2	0.94	0.94	0.96	0.92	0.95
5	0.96	0.96	0.97	0.95	0.97
8	0.96	0.96	0.97	0.95	0.97
10	0.96	0.96	0.97	0.95	0.97
50	0.96	0.97	0.96	0.95	0.96

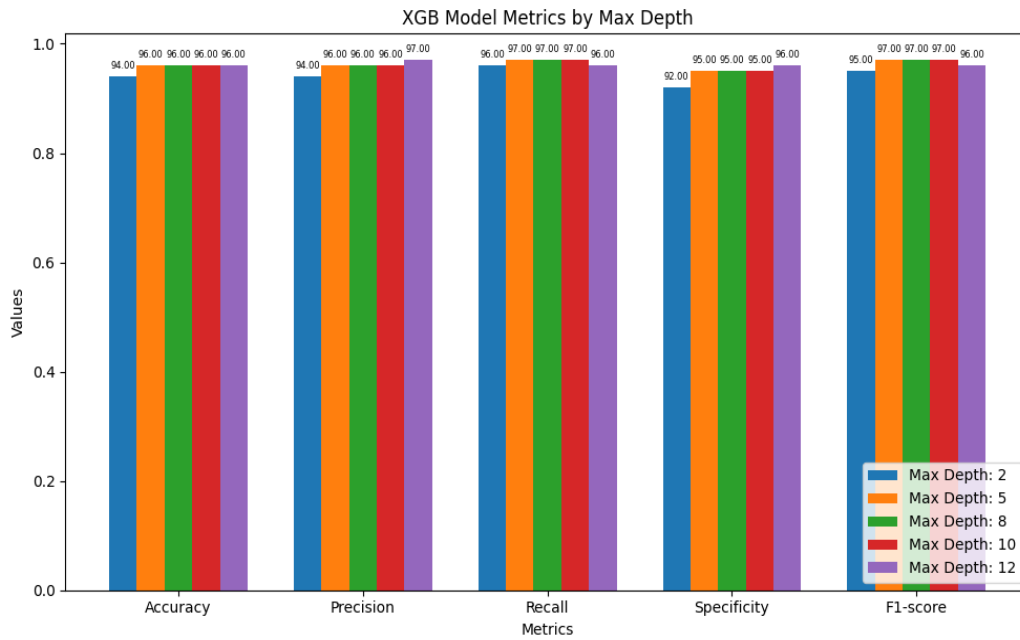


Fig. 9. Experiment Result of XGBoost Tree Compare with Evaluation Parameter

Table 5. Experiment Result of Hybrid Model

Voting	Accuracy	Precision	Recall	F1-Score
Soft Voting	0.965330	0.968365	0.969925	0.969144
Hard Voting	0.965933	0.968901	0.970462	0.969681

Table 6. Result Comparison of Proposed Hybrid Model with Machine Learning Models

Algorithm	Accuracy	Precision	Recall	F1-Score
Decision Tree	0.94	0.94	0.94	0.94
Random Forest	0.96	0.92	0.93	0.92
Extreme Gradient Boosting(XGBoost)	0.96	0.96	0.96	0.96
DT+RF+XGB(Soft Voting)	0.96	0.96	0.96	0.96
DT+RF+XGB(Soft Voting)	0.96	0.97	0.96	0.96

In the table 6, The hybrid model (DT+RF+XGBoost) is showing the improvised value as compared to state-of-arts.

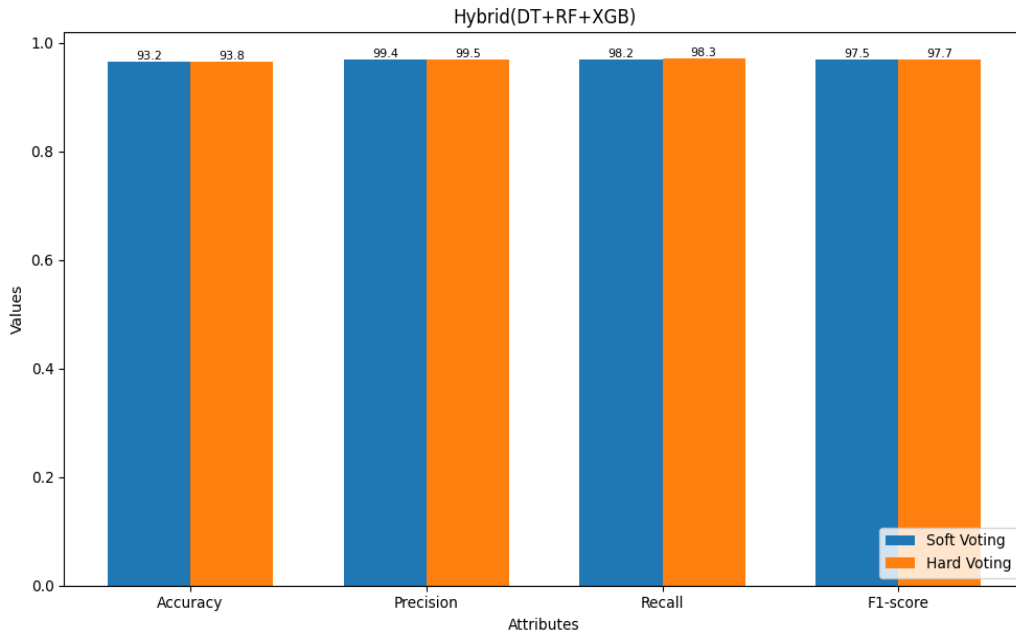


Fig. 10. Proposed Hybrid Method Performance using Evolution Parameters

Table 7. Result Comparison of Proposed Hybrid Model with State-of-Art

Sr.No	Author	Algorithm	Evaluation Parameters
1	Proposed Algorithm	Hybrid Machine Learning Model(DT+RF+XGB)	Accuracy= 0.96, precision=0.96, Recall= 0.97, F1-Score=0.96
2	ABDULKARIM,MOBEEN SHAHROZ, KHABIB MUSTOFA, SAMIR BRAHIM BELHAOUARI, ANDS,RAMANA KUMAR, JOGA(2023)[1]	Hybrid Machine Learning Model(SVC+ DT +LR)	Accuracy= 0.95, Precision=0.95, Recall=0.96, F1-Score=0.95
3	Sonowal, G., Kuppusamy, K(2020)[18]	five-layer phishing detection model called PhiDMA	Accuracy= 0.92, Precision=0.91, Recall=0.90, F1-Score=0.90
4	Kang Leng Chiew, Choon Lin Tan, Kok Sheik Wong , Kelvin S.C.Yongc, Wei King Tionga(2019)[9]	Hybrid Ensemble Feature Selection(HEFS),Random Forest Classifier	Accuracy= 0.94
5	YONG FANG, CHENG ZHANG, CHENG HUANG, LIANG LIU,ANDYUE YANG (2019)[16]	Deep learning model named THEMIS(advanced version of RCNN)	Accuracy= 0.99, Precision=0.99, Recall=0.99, F1-Score=0.99
6	Ozgur Koray Sahingoz, Ebubekir Buber ,OnderDemir, Banu Diri(2017)[22]	Using Random Forest with NLP-based features	Accuracy= 0.97, Precision=0.97, Sensitivity=0.99, F1-Score=0.98

Improved performance compared to individual DT, RF and XGB models. High accuracy and F1 score with relatively low training time. Slightly high false positive rate (misclassifies some legitimate URLs). Ensemble Learning (combining multiple models) Training Datasets: Labeled datasets with phishing and legitimate URLs. Accuracy Measures: Accuracy: Overall correct classification rate. Precision: Ratio of true positives (correctly identified phishing URLs). Recall: Ratio of all phishing URLs correctly identified. F1 Score: Harmonic mean of precision and recall. Best Performing algorithms: XGB: High accuracy and low false positive rate. DT + RF : Highest precision (few legitimate URLs misclassified). Ensemble methods often outperform single models. The proposed ensemble method significantly outperforms traditional machine learning techniques like Decision Trees (DT), Random Forests (RF), and Extreme Gradient Boosting (XGB). By integrating these methods into a cohesive ensemble, it achieves the lowest false positive rate, crucial for high-precision applications like phishing URL detection. The ensemble excels in key performance metrics, boasting the highest accuracy, F1-score, and recall. This indicates its overall correctness and ability to effectively identify true phishing URLs. Additionally, the method is more efficient in training time compared to some sophisticated alternatives, making it suitable for large datasets and frequent updates. Incorporating multi-head self-attention mechanisms and Convolutional Neural Networks (CNNs) enhances its performance by capturing intricate patterns and extracting hierarchical features. This combination not only reduces false positives but also improves robustness and generalization. Overall, this ensemble method showcases the potential of advanced machine learning techniques in enhancing cybersecurity, particularly in phishing detection.

The proposed graphical user interface (GUI) of phishing detection system is developed using Streamlit, a powerful Python library for building interactive web applications. The GUI offers a user-friendly experience with intuitive functionalities, where Users can input URLs. The layout and features of the GUI:

- Upload functionality for URLs.
- Display of the uploaded URLs.
- Recognition of characters in the URLs.
- Display of the predicted phishing non-phishing URLs.

The screenshot displays the URLGuard software interface, which is a dark-themed web application. It features a grid of 25 feature selection dropdown menus arranged in 5 rows and 5 columns. Each dropdown menu has a label and a value of '-1'. The features are: UsingIP, Subdomains, RequestURL, WebsiteForwarding, DNSRecording, LongURL, HTTPS, AnchorURL, StatusBarCust, WebsiteTraffic, ShortURL, DomainRegLen, LinksInScriptTags, DisableRightClick, PageRank, Symbol@, Favicon, ServerFormHandle, UsingPopupWindow, GoogleIndex, Redirecting, NonStdPort, InfoEmail, IFrameRedirection, LinksPointingToPage, Prefixsuffix, HTTPSDomainURL, AbnormalURL, AgeofDomain, and StatsReport. At the bottom left of the grid is a 'Predict' button.

Fig. 11. Output Image of Proposed Hybrid Approach using Software

5. Conclusion and Future Work

According to the literature survey, the two most accurate models are the deep learning model dubbed THEMIS and the hybrid machine learning model (LR+SVC+DT), with accuracy values of 94.848 and 95.23, respectively. THEMIS (Throughput-oriented Efficient Multi-stage Inception for Scalable Object Detection) is an improved object detection architecture that improves on the capabilities of R-CNN. It seeks to attain high accuracy while being efficient, especially in real-time or resource-constrained applications. This approach enables for the simultaneous modeling of emails at many levels, including the email header, email body, character level, and word level, thereby improving the model's capacity to capture essential information. A hybrid machine learning model that combines Logistic Regression (LR), Decision Trees (DT), and Support Vector Classification (SVC) makes use of the advantages of each approach. LR handles linear relationships, DT captures non-linear patterns, and SVC ensures robust classification. During training, features are processed by LR, DT, and SVC independently. The individual predictions are then combined, often using a weighted average, to generate a final output. This ensemble approach enhances overall model performance by exploiting the complementary capabilities of each algorithm.

The proposed model is successfully developed and implemented a comprehensive URL-based phishing detection system. This system leverages a variety of machine learning algorithms to effectively identify and thwart phishing attempts, particularly those involving deceptive URLs. The project's objectives, including enhancing accuracy(0.96),recall(0.96),F1-score(0.96)and Precision(0.97), real-time monitoring, and user-friendly interface design, have been achieved. Through continuous monitoring and collaboration with the cyber security community, the system remains adaptive and resilient against evolving phishing tactics. The studies showcased hybrid models, each with unique strengths and weaknesses. While these models demonstrate enhanced accuracy and efficiency, gaps exist, such as a lack of in-depth analysis on false positive/negative rates, scalability challenges, and omitted dataset details. Despite these limitations, the review illuminates valuable insights, offering a foundation for future research to refine and advance URL-based phishing detection systems using hybrid machine learning approaches.

Future Scope

- User Interface Design: Create an intuitive and user-friendly interface for system administrators, facilitating easy management, monitoring, and response to potential phishing threats.
- Security Collaboration: Foster collaboration with the wider cyber security community to stay informed about the latest threats and best practices, ensuring the system's relevance and effectiveness.
- Continuous Improvement: Establish procedures for continuous monitoring, evaluation, and updating of the system to remain effective against evolving phishing tactics over time.

References

- [1] Karim, Abdul, Mobeen Shahroz, Khabib Mustofa, Samir Brahim Belhaouari, and S. Ramana Kumar Joga. "Phishing detection system through hybrid machine learning based on URL." *IEEE Access* 11 (2023): 36805-36822.
- [2] Zouina, M., Outtaj, B.: A novel lightweight url phishing detection system using svm and similarity index. *Human-centric Computing and Information Sciences* 7(1), 1–13 (2017).
- [3] Wang, S., Khan, S., Xu, C., Nazir, S., Hafeez, A.: Deep learning-based efficient model development for phishing detection using random forest and blstm classifiers. *Complexity* 2020, 1–7 (2020)
- [4] Fang, Yong, Cheng Zhang, Cheng Huang, Liang Liu, and Yue Yang. "Phishing email detection using improved RCNN model with multilevel vectors and attention mechanism." *IEEE Access* 7 (2019): 56329-56340.
- [5] Sahingoz, Ozgur Koray, Ebubekir Buber, Onder Demir, and Banu Diri. "Machine learning based phishing detection from URLs." *Expert Systems with Applications* 117 (2019): 345-357.
- [6] Abdelhamid, N., Ayesha, A., Thabtah, F.: Phishing detection based associative classification data mining. *Expert Systems with Applications* 41(13), 5948–5959 (2014)
- [7] Karim, A., Shahroz, M., Mustofa, K., Belhaouari, S.B., Joga, S.R.K.: Phishing detection system through hybrid machine learning based on url. *IEEE Access* 11, 36805–36822 (2023).
- [8] Sahingoz, O.K., Buber, E., Demir, O., Diri, B.: Machine learning based phishing detection from urls. *Expert Systems with Applications* 117, 345–357 (2019)
- [9] Chiew, Kang Leng, Choon Lin Tan, KokSheik Wong, Kelvin SC Yong, and Wei King Tiong. "A new hybrid ensemble feature selection framework for machine learning-based phishing detection system." *Information Sciences* 484 (2019): 153-166.
- [10] Shahrivari, V., Darabi, M.M., Izadi, M.: Phishing detection using machine learning techniques. *arXiv preprint arXiv:2009.11116* (2020)
- [11] Buber, E., Diri, B., Sahingoz, O.K.: Nlp based phishing attack detection from urls. In: *Intelligent Systems Design and Applications: 17th International Conference on Intelligent Systems Design and Applications (ISDA 2017) Held in Delhi, India, December 14-16, 2017*, pp. 608–618 (2018). Springer
- [12] Paithane, P.M.: Random forest algorithm use for crop recommendation. *ITEGAM-JETIA* 9(43), 34–41 (2023)
- [13] Paithane, P.M.: Yoga posture detection using machine learning. *Artificial Intelligence in Information and Communication Technologies, Healthcare and Education: A Roadmap Ahead* 27 (2022).
- [14] Chiew, K.L., Tan, C.L., Wong, K., Yong, K.S., Tiong, W.K.: A new hybrid ensemble feature selection framework for machine learning-based phishing detection system. *Information Sciences* 484, 153–166 (2019)
- [15] Paithane, P., Kakarwal, S.: Lmns-net: Lightweight multiscale novel semantic-net deep learning approach used for automatic pancreas image segmentation in ct scan images. *Expert Systems with Applications* 234, 121064 (2023)
- [16] Fang, Y., Zhang, C., Huang, C., Liu, L., Yang, Y.: Phishing email detection using improved rcnn model with multilevel vectors and attention mechanism. *IEEE Access* 7, 56329–56340 (2019).
- [17] Rao, R.S., Pais, A.R.: Detection of phishing websites using an efficient featurebased machine learning framework. *Neural Computing and applications* 31, 3851– 3873 (2019).
- [18] Sonowal, G., Kuppasamy, K.: Phidma—a phishing detection model with multifilter approach. *Journal of King Saud University-Computer and Information Sciences* 32(1), 99–112 (2020).
- [19] Alam, M.S., Vuong, S.T.: Random forest classification for detecting android malware. In: *2013 IEEE International Conference on Green Computing and Communications and IEEE Internet of Things and IEEE Cyber, Physical and Social Computing*, pp. 663–669 (2013). IEEE.
- [20] Smadi, S., Aslam, N., Zhang, L.: Detection of online phishing email using dynamic evolving neural network based on reinforcement learning. *Decision Support Systems* 107, 88–102 (2018).
- [21] Paithane, P.M., Kakarwal, S.: Automatic pancreas segmentation using a novel modified semantic deep learning bottom-up approach. *International Journal of Intelligent Systems and Applications in Engineering* 10(1), 98–104 (2022).
- [22] Buber, E., Diri, B., Sahingoz, O.K.: Detecting phishing attacks from url by using nlp techniques. In: *2017 International Conference on Computer Science and Engineering (UBMK)*, pp. 337–342 (2017). IEEE.
- [23] Feng, F., Zhou, Q., Shen, Z., Yang, X., Han, L., Wang, J.: The application of a novel neural network in the detection of phishing websites. *Journal of Ambient Intelligence and Humanized Computing*, 1–15 (2018).
- [24] Wagh, S.J., Paithane, P.M., Patil, S.: Applications of fuzzy logic in assessment of groundwater quality index from jafrabad taluka of marathawada region of maharashtra state: A gis based approach. In: *International Conference on Hybrid Intelligent Systems*, pp. 354–364 (2021). Springer.
- [25] Shirazi, H., Hayne, K.: Towards performance of nlp transformers on url-based phishing detection for mobile devices. *International journal of ubiquitous systems and pervasive networks* (2022)
- [26] Pradip paithane, "Trust Aware Recommendation using Deep Matrix Factorization Model", *JETIA*, vol. 10, no. 48, pp. 115-121, Aug.2024.

Authors' Profiles



Pradip M. Paithane is currently working as Assistant Professor in the Department of Computer Engineering, VPKBIET College Baramati, Maharashtra, India from January 2017. He received B.E and M.E. in CSE from Dr. Babasheb Ambedkar Marathwada University, Aurangabad, and PhD in Computer Engineering from Dr. Babasheb Ambedkar Marathwada University, Aurangabad in 2022. He has more than 11 years of teaching experience in various engineering institutions and 4 year of research experience. So far, he guided several undergraduate and post graduate projects and dissertations. His research interests include image processing, Deep Learning, security in networks and communication, Wireless Adhoc Networks, IoT and Blockchain. He is a life member of Indian Society for Technical Education (ISTE) India.

How to cite this paper: Pradip M. Paithane, "URLGuard: A Holistic Hybrid Machine Learning Approach for Phishing Detection", International Journal of Information Engineering and Electronic Business(IJIEEB), Vol.17, No.2, pp. 95-110, 2025. DOI:10.5815/ijieeb.2025.02.05