



OPH-Guard: An Operationally Interpretable Tree-Ensemble Framework for Phishing URL Screening in Secure Web Access Management

Reem Atassi^{1,*}

¹Higher Colleges of Technology, United Arab Emirates

Email: ratassi@hct.ac.ae

Abstract

Phishing URLs still present a security threat to organizations because they enable credential theft and account takeover together with payment fraud and unauthorized digital service access. The existing research on phishing detection has been studied extensively yet most published papers still show a preference for predictive performance assessment compared to operational system capabilities and tests and governance system implementation. The researchers developed *OPH-Guard* as an operational security system which uses compact tree ensembles to identify phishing URLs for their secure web access management system. The integrated workflow system enables institutional and small enterprise to implement public data ingestion and feature validation together with tabular model learning and post-hoc explanation and security-action mapping. The empirical evaluation used a public GitHub-hosted phishing URL dataset which contains 11,481 labeled records and 87 predictive features. The researchers conducted a comparison between three tree-based learners according to a stratified 80/20 hold-out protocol which included Decision Tree and Random Forest and Extra Trees. The actual results from Extra Trees produced the highest accuracy score of 0.9856 which included 0.9921 precision and 0.9791 recall and 0.9855 F1-score and 0.9984 ROC-AUC from the held-out test results. The study investigates security relevance for top predictors through *google index* and *page rank* and *domain age* and *phish hints* which provide evidence that the resulting model enables organizations to manage browsing risk through URL triage together with secure information management controls. The study presents a reproducible framework together with a complete screening algorithm and a summary of existing research from ten studies and a system which connects model results to security operations.

Keywords: Phishing URL detection; Tree ensembles; Extra Trees; Secure web access management; Operational interpretability; Cybersecurity analytics

1 Introduction

Cybercriminals use phishing attacks for dedicated execution because this attack method requires minimal expenses while it provides attackers with chances to steal credentials and compromise accounts and conduct payment fraud and gain unauthorized access to online platforms. Phishing attacks have developed into two different areas because criminals now use more sophisticated technological methods to execute their attacks, whereas their methods of delivering phishing attacks have shifted toward new techniques. The security system of every organization needs URL screening because it protects against both phishing attacks and other internet threats which users face while accessing online platforms.

Deceptive URLs function as technical elements which lead to multiple security breaches that create risks to identity verification systems and customer access points and financial networks and cloud-based services. Deceptive URLs in practical settings block users from accessing various platforms which include web browsers and remote-work systems and institutional websites and e-commerce platforms and customer information systems. Phishing URL analysis should function as both a classification system and a decision support tool which enables organizations to improve secure access control while protecting users from dangerous sites and building digital trust.

Researchers have developed numerous methods for phishing detection which already exist in the academic literature. Mohammad et al. established essential research through their study which showed how website and URL attributes created through manual work could become valuable for establishing public benchmarks in their field of research [1, 2]. Researchers conducted their studies by using lightweight lexical extraction methods and feature-rich machine learning screening techniques and deep learning URL models and explainable detection systems which used XML as their foundation [3] [4, 5] [6, 7, 8, 9] [10]. Research studies based on systematic methods and data-driven approaches have demonstrated how benchmark quality and research methods impact their results [11, 12]. However, researchers have not yet achieved a key breakthrough which would allow them to create research outputs that other researchers could use in practical settings and then reproduce those findings.

First, many papers are optimized for classification accuracy but devote limited attention to *operational interpretability*. The web decisions which use filtering technology create effects that impact actual user activities and real company operations and the way people view automated systems. Second, several studies use large or computationally heavy pipelines that are less attractive for organizations seeking easy-to-audit controls. Third, the connection between phishing URL detection and information-management practice is often underdeveloped, even though model outputs are directly relevant to policy enforcement, risk triage, and secure browsing governance. The research presents strong numerical results through multiple manuscripts which fail to identify any novel elements of their work beyond testing a different classifier on standard evaluation datasets.

This study creates *OPH-Guard* as an operationally interpretable tree-ensemble framework which developers use to screen phishing URLs. The framework integrates a public dataset with a compact modeling framework and a formal screening algorithm and a system for security web access management which enables multiple feature ranking methods and specific action decision processes. The research establishes operational interpretation through predictive modeling which maintains reproducibility across workflow development activities.

The main contributions of the paper are summarized as follows:

1. A clearly defined phishing URL screening framework which researchers call *OPH-Guard* presents public data ingestion and feature validation and tree-ensemble learning and model interpretation and operational response mapping as its integrated components.
2. A complete reproducible empirical research study employs a public phishing URL dataset which contains 11481 labeled records and 87 predictive variables to demonstrate its findings.
3. The study conducts an actual performance comparison between Decision Tree and Random Forest and Extra Trees using a uniform stratified train-test method which shows Extra Trees as the top-performing model.
4. The proposed model workflow uses formal algorithmic description which enables simple implementation and adaptation to institutional environments.
5. The related-work synthesis has been expanded through ten published studies which demonstrate how the current paper achieves different results in reproducibility and interpretability and information-management relevance.

The remainder of the paper is organized as follows. Section 2 reviews related studies and identifies the positioning gap. Section 3 presents the proposed framework, including its mathematical description and algorithmic workflow. Section 4 describes the dataset, preprocessing, and experimental protocol. Section 5 reports the empirical findings. Section 6 discusses the results from cybersecurity and information-management perspectives. Section 7 outlines limitations and future work, and Section 8 concludes the paper.

2 Related Work and Research Positioning

2.1 Representative studies in phishing URL and website detection

Phishing detection has evolved from heuristic and rule-based inspection toward machine-learning systems that learn decision boundaries from lexical, host-based, structural, semantic, and reputation-oriented signals. Mohammad et al. established a foundational feature-engineering perspective by systematically evaluating phishing website attributes and releasing an associated benchmark that later entered the UCI Machine Learning Repository [1, 2]. Verma and Das then highlighted the practical value of rapid malicious URL feature extraction, demonstrating that lightweight URL-centric analysis can support fast detection [3]. Rao and Pais

expanded the feature-based approach by combining URL, source-code, and third-party service information within a machine-learning framework aimed at more efficient phishing website discrimination [4].

A second line of work emphasized broader machine-learning comparisons. Sahingoz et al. investigated multiple models for phishing detection from URLs and demonstrated the competitiveness of feature-rich machine-learning approaches on URL-centric data [5]. Almomani et al. later evaluated semantic URL features across multiple classifiers and specifically discussed the underfitting and overfitting behavior of competing models [13]. Ahammad et al. also reported strong results using machine-learning methods for phishing URL detection and reinforced the usefulness of carefully selected structured predictors [14].

A third stream moved toward deep learning. Aljofey et al. proposed a character-level convolutional neural network that operates directly on URLs without relying on manually engineered features [6]. Do et al. compared deep-learning algorithms for phishing webpage classification and emphasized performance under reduced computational constraints [7]. Hussain et al. proposed CNN-Fusion, a lightweight multi-variant ConvNet for phishing URL detection, while Ghalechyan et al. reported an empirical neural-network study in *Scientific Reports* that further confirmed the value of deep models in this area [8, 9]. More recently, Uddin et al. explicitly brought explainability into the phishing-site-detection workflow by coupling ensemble learning with transparent interpretation [10].

Alongside model-centered work, the literature also contains important meta-level studies. Vrbancic et al. highlighted the importance of curated phishing datasets and published larger phishing website data resources for future research [11]. Safi and Singh synthesized the broader literature in a systematic review, showing that machine learning and deep learning dominate recent phishing-detection research but also noting common concerns related to feature quality, benchmark realism, and evaluation consistency [12]. Together, these studies show a mature field, but they also reveal a useful opening for a paper that prioritizes reproducibility, operational clarity, and direct alignment with secure information-management practice.

2.2 Expanded comparison with prior work

Table 1 summarizes ten representative published studies and clarifies how the present manuscript differs from them. The table is intentionally qualitative rather than score-centric. Many published studies use different datasets, feature spaces, and evaluation protocols, making raw metric comparison less informative than methodological positioning.

Table 1: Expanded related-work comparison with ten published phishing detection studies.

Study	Core input type	Main methodological direction	Main strength	Limitation relative to this paper
Mohammad et al. (2012) [1]	Website / heuristic features	Foundational feature assessment	Established a reusable feature benchmark	Limited operational framing and no deployment-oriented workflow
Verma and Das (2017) [3]	URL features	Fast malicious URL feature extraction	Emphasized efficient feature generation	Less focus on governance interpretation and end-to-end screening design
Rao and Pais (2019) [4]	URL, source code, third-party cues	Feature-based ML framework	Strong handcrafted feature integration	Less emphasis on compact reproducible operationalization
Sahingoz et al. (2019) [5]	URL features	Multi-model machine learning	Demonstrated competitive URL-based ML screening	Limited information-management discussion
Aljofey et al. (2020) [6]	Raw URLs	Character-level CNN	Removed need for manual feature engineering	Lower transparency for policy-facing deployment
Do et al. (2021) [7]	Webpage / phishing data	Deep-learning comparison	Showed strong DL performance under different settings	Less interpretable than compact ensemble screening
Almomani et al. (2022) [13]	Semantic URL features	Comparative ML with fit analysis	Discussed underfitting and overfitting issues	Less explicit operational response mapping
Ahammad et al. (2022) [14]	Structured URL/website features	Machine-learning URL detection	Reinforced the value of tabular ML approaches	Limited framework-level novelty articulation
Hussain et al. (2023) [8]	Raw URLs	Lightweight CNN-Fusion	Strong lightweight deep model design	Reduced interpretability compared with feature-ranked ensemble workflow

Continued on next page

Table 1 continued from previous page

Study	Core input type	Main methodological direction	Main strength	Limitation relative to this paper
Ghalechyan et al. (2024) [9]	URLs	Neural-network empirical study	Contemporary large-scale empirical perspective	Less directly oriented to secure web access governance
This study	Public tabular phishing URL data	Operationally interpretable tree-ensemble framework	Reproducible workflow, formal algorithm, feature-level interpretation, and explicit security-action mapping	Designed as an applied journal framework rather than a mathematically novel model

2.3 Gap analysis and positioning

The literature review reveals a structural gap rather than a purely algorithmic one. Existing studies are often strong in one dimension but incomplete across the end-to-end workflow required by practical cybersecurity operations. Deep models may provide strong accuracy, but many are less transparent in environments where blocking, warning, and triage decisions must be justified. Feature-based models may be interpretable, but they are not always presented as reproducible frameworks with clear operational outputs. Review and dataset papers clarify the state of the field, yet they do not deliver a compact applied framework validated on public data and linked directly to operational decision-making.

Accordingly, the present paper does not claim that Extra Trees or tree ensembles are new in an absolute sense. Its novelty lies in *how* the methodology is organized and communicated: a public-data workflow, formalized screening procedure, explicit interpretability stage, and direct integration with secure web access management. That positioning is illustrated conceptually in Figure 1 and summarized again in Table 2.

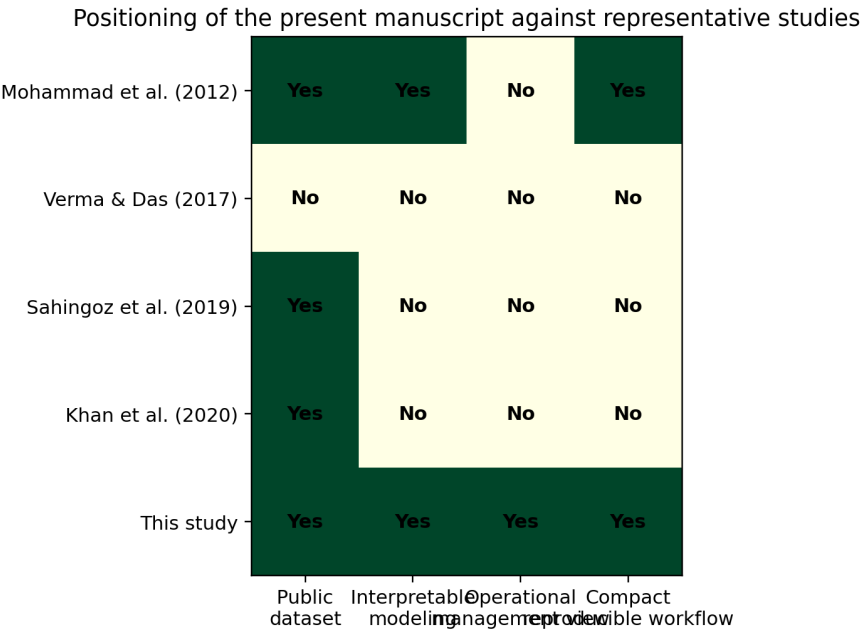


Figure 1: Positioning of the present study against representative phishing URL studies.

Table 2: Positioning and novelty statement of the proposed manuscript.

Aspect	What prior work commonly emphasizes	What this study adds	Practical significance
Model reporting	Accuracy-oriented comparison of classifiers	End-to-end framework with operational response stage	Gives the manuscript stronger practical value
Interpretability	Partial or post-hoc discussion only	Embedded feature-ranking and cyber-meaning extraction	Helps decision-makers trust the model
Reproducibility	Often dataset-specific with limited workflow detail	Public dataset, explicit preprocessing, and formal algorithm	Supports verification and reuse
Application framing	Security detection framed mainly as ML benchmarking	Secure web access management and information-governance framing	Connects predictive modeling to operational decision support

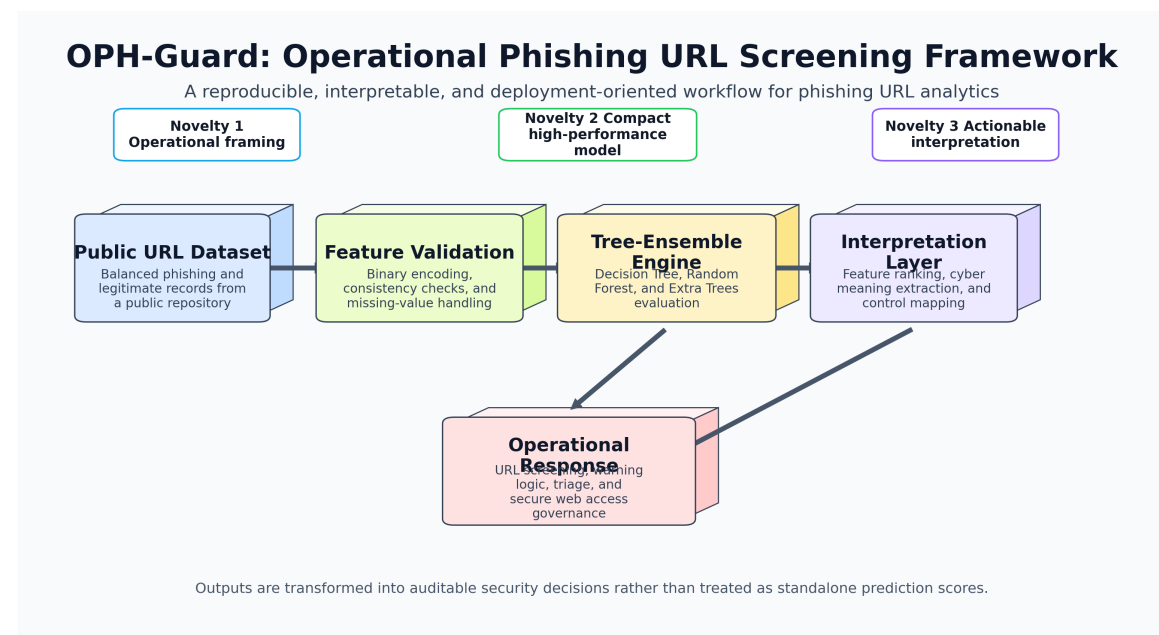
3 The Proposed OPH-Guard Framework

3.1 Design objective

The design objective of OPH-Guard is to provide a compact, interpretable, and reproducible phishing URL screening workflow that can be understood by both technical reviewers and operational stakeholders. Instead of proposing another opaque architecture, the framework organizes the analytical flow into five linked stages: public data acquisition, feature validation and normalization, tree-ensemble learning, feature-based interpretation, and security-action mapping. The design deliberately favors a model family that is strong enough to generate competitive results while remaining transparent enough for policy-facing cybersecurity deployment.

3.2 Framework overview

Figure 2 presents the proposed framework in a publication-ready visual form. The figure emphasizes process flow and operational use rather than only model internals. This design decision is central to the paper, because the contribution lies in the end-to-end screening framework rather than in a single algorithmic trick.

**Figure 2:** The proposed OPH-Guard framework for phishing URL screening and secure web access management.

3.3 Mathematical formulation

Let the phishing URL dataset be denoted by

$$\mathcal{D} = \{(\mathbf{x}_i, y_i)\}_{i=1}^N, \quad (1)$$

where $\mathbf{x}_i \in \mathbb{R}^d$ is the feature vector of the i th URL, $d = 87$ is the number of predictive variables after preprocessing, and $y_i \in \{0, 1\}$ denotes the class label, with $y_i = 1$ for phishing and $y_i = 0$ for legitimate URLs.

The dataset is partitioned into disjoint training and testing subsets,

$$\mathcal{D}_{\text{tr}} \cup \mathcal{D}_{\text{te}} = \mathcal{D}, \quad \mathcal{D}_{\text{tr}} \cap \mathcal{D}_{\text{te}} = \emptyset, \quad (2)$$

using stratified sampling to preserve the class distribution. For a tree-ensemble model with T trees, the predicted phishing score for $\mathbf{x}$ is obtained as

$$\hat{p}(\mathbf{x}) = \frac{1}{T} \sum_{t=1}^T h_t(\mathbf{x}), \quad (3)$$

where $h_t(\mathbf{x}) \in [0, 1]$ denotes the class-posterior estimate produced by tree t . The final decision rule is

$$\hat{y}(\mathbf{x}) = \begin{cases} 1, & \hat{p}(\mathbf{x}) \geq \tau, \\ 0, & \hat{p}(\mathbf{x}) < \tau, \end{cases} \quad (4)$$

where $\tau = 0.5$ in the present implementation.

Model selection is based primarily on the F1-score,

$$F_1 = \frac{2 \text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}, \quad (5)$$

with

$$\text{Precision} = \frac{TP}{TP + FP}, \quad \text{Recall} = \frac{TP}{TP + FN}. \quad (6)$$

Among the candidate models $\mathcal{M} = \{m_1, m_2, m_3\}$, the selected classifier is

$$M^* = \arg \max_{m \in \mathcal{M}} F_1(m), \quad (7)$$

while ROC-AUC is used as a supporting threshold-independent criterion.

For the selected tree ensemble, the global importance of feature j is estimated by the normalized average impurity decrease,

$$I_j = \frac{1}{T} \sum_{t=1}^T \sum_{n \in \mathcal{N}_{t,j}} \Delta i(n), \quad (8)$$

where $\mathcal{N}_{t,j}$ is the set of split nodes in tree t that use feature j , and $\Delta i(n)$ denotes the impurity reduction contributed by node n . The ranked vector $\mathbf{I} = (I_1, \dots, I_d)$ is then used to support the interpretation and action-mapping stages of the framework.

3.4 Formal description of the proposed model

The core logic of OPH-Guard is summarized in Algorithm 1. The algorithm takes a public phishing URL dataset as input, performs feature validation, trains multiple tree-based learners, selects the best-performing model using hold-out metrics, and extracts interpretive signals for downstream security controls.

Algorithm 1 OPH-Guard phishing URL screening procedure**Require:** Public dataset D with feature matrix X and target labels y **Ensure:** Best trained model M^* , evaluation report R , ranked importance list I , and security-action map A

- 1: Load dataset D from the public repository
- 2: Remove non-predictive identifier fields and define binary target labels
- 3: Convert symbolic feature values to numeric form and impute missing values
- 4: Split D into stratified training set D_{tr} and testing set D_{te}
- 5: Define candidate models $\mathcal{M} = \{\text{Decision Tree, Random Forest, Extra Trees}\}$
- 6: **for all** $m \in \mathcal{M}$ **do**
- 7: Train m on D_{tr}
- 8: Predict labels and scores on D_{te}
- 9: Compute Accuracy, Precision, Recall, F1-score, and ROC-AUC
- 10: Store the results in report R
- 11: **end for**
- 12: Select the best model M^* according to the highest F1-score with supporting ROC-AUC
- 13: Extract and rank the feature importances of M^* to form I
- 14: Translate the leading predictors into cyber-meaning categories and operational controls
- 15: Construct action map A for blocking, warning, or triage recommendations
- 16: **return** (M^* , R , I , A)

3.5 Novelty and contribution statement

The novelty of OPH-Guard should be understood at three complementary levels.

First, the framework contributes **operational novelty**. Instead of treating phishing URL detection as a stand-alone classification task, it explicitly links prediction to secure web access management, browsing-risk prioritization, and information-governance decisions. This gives the manuscript a stronger fit with cybersecurity and information-management audiences.

Second, the framework contributes **methodological clarity**. The full empirical pipeline is intentionally compact and reproducible: public dataset, explicit preprocessing, transparent model selection, and direct feature interpretation. This is especially valuable in applied cybersecurity settings, where reproducibility and direct utility often matter more than architectural complexity.

Third, the framework contributes **interpretive novelty**. The paper does not stop at reporting a high-performing classifier. It interprets the leading predictors in security terms and explains how they can support blocking, warning, or triage actions. This bridges the gap between model output and operational decision-making.

3.6 Framework stages

For clarity, the stages of OPH-Guard are described below.

Stage 1: Public dataset acquisition. The framework begins with a public phishing URL dataset, allowing later verification and reproduction by reviewers or readers.

Stage 2: Feature validation. Categorical labels and symbolic values are normalized to a machine-readable form. This stage ensures that the downstream model receives a consistent numerical feature matrix and reduces ambiguity in the feature semantics.

Stage 3: Tree-ensemble learning. Three tree-based classifiers are evaluated: Decision Tree, Random Forest, and Extra Trees. The framework favors tree ensembles because of their balance between predictive power and interpretability on structured tabular inputs.

Stage 4: Interpretation layer. Feature importance is extracted from the best-performing model to reveal which URL and host-related properties most strongly drive the decision boundary.

Stage 5: Security-action mapping. The model output is positioned as a screening signal that can support web filtering, browser warning workflows, help-desk inspection queues, and governance-oriented browsing controls.

4 Materials and Methods

4.1 Dataset used in the experiment

The empirical evaluation relies on a public GitHub-hosted phishing URL dataset [15]. The downloaded file contains 11,481 rows and 90 columns. After excluding the raw `url` field and using `status` as the binary target label, 87 predictive features remained in the final experiment matrix. The class distribution is nearly balanced, with 5,741 phishing records and 5,740 legitimate records. The feature space includes URL structure indicators, hyperlink statistics, domain properties, content-related hints, and reputation-oriented variables.

The present experiment dataset differs from the classic UCI benchmark [2]. This distinction is useful rather than problematic. The UCI benchmark remains important for historical comparison, while the GitHub-hosted dataset used here offers a larger feature space and a broader tabular description that is well suited for a tree-ensemble study. For a framework paper, this is advantageous because the data representation aligns naturally with interpretable tabular learning.

4.2 Preprocessing and data preparation

The preprocessing pipeline was intentionally kept simple to preserve reproducibility. First, the target variable `status` was converted to binary form, where phishing URLs were encoded as 1 and legitimate URLs as 0. Second, symbolic feature values such as textual `zero/one` markers were mapped to their numeric equivalents. Third, any missing values created during transformation were imputed with zero. Because all candidate models are tree-based and operate effectively on numeric tabular inputs, no additional standardization or dimensionality reduction was applied.

This design has two practical advantages. It preserves the semantic meaning of the original structured features and avoids unnecessary transformations that could reduce interpretability. It also makes the workflow easy to reproduce in applied settings where teams may not want complex preprocessing dependencies.

4.3 Experimental protocol

The full dataset was partitioned using a stratified 80/20 split with a fixed random seed of 42. This produced 9,184 records for training and 2,297 records for testing. Stratification preserved the balanced class distribution across both partitions. The same split was used for all candidate classifiers to ensure fair comparison. The evaluation design prioritizes clean held-out assessment over repeated tuning, which is appropriate for a compact applied study intended to demonstrate a reproducible baseline-plus-framework contribution.

4.4 Candidate models

Three models were evaluated.

Decision Tree. This model acts as the single-tree baseline. It provides direct interpretability, but it is often sensitive to variance and may underperform on more complex tabular decision boundaries.

Random Forest. This model aggregates multiple decorrelated decision trees and is widely regarded as a strong baseline for structured security data [16]. It reduces variance and usually produces more stable performance than a single tree.

Extra Trees. This model extends the ensemble concept by injecting additional randomization during split generation, often producing robust performance on heterogeneous feature spaces [17]. For phishing URL data, this is useful because structural, semantic, and reputation-oriented predictors can interact in non-linear ways.

The candidate model set was intentionally kept compact because the purpose of the paper is not exhaustive hyperparameter tuning but rather the presentation of a reproducible and operationally interpretable framework.

4.5 Evaluation metrics

Performance was assessed using Accuracy, Precision, Recall, F1-score, and ROC-AUC. Accuracy provides an overall correctness ratio; Precision measures how many URLs predicted as phishing are actually phishing; Recall captures how many phishing URLs were successfully detected; and the F1-score summarizes the balance between Precision and Recall. ROC-AUC measures ranking quality across thresholds. For phishing URL screening, Recall and F1-score deserve particular attention because false negatives leave users exposed to unsafe destinations, while excessive false positives can disrupt legitimate browsing and business tasks.

4.6 Deployment-oriented interpretation strategy

A distinguishing part of the methodology is the interpretation step after model selection. Once the best model is chosen, the leading features are grouped into cyber-meaning categories such as reputation, link-structure abnormality, domain maturity, and lexical phishing hints. This makes the framework more useful for operational governance. Instead of presenting only a ranked list of variables, the study translates them into concepts that can support policy design and triage logic.

5 Results and Analysis

5.1 Overall classification performance

Table 3 reports the held-out test results. Extra Trees emerged as the strongest model overall, with 0.9856 accuracy, 0.9921 precision, 0.9791 recall, 0.9855 F1-score, and 0.9984 ROC-AUC. Random Forest followed closely and also produced strong results. The Decision Tree baseline performed adequately but remained clearly behind the two ensemble models.

Table 3: Held-out test performance of the evaluated phishing URL classifiers.

Model	Accuracy	Precision	Recall	F1-score	ROC-AUC
Decision Tree	0.9630	0.9602	0.9661	0.9631	0.9681
Random Forest	0.9808	0.9834	0.9782	0.9808	0.9973
Extra Trees	0.9856	0.9921	0.9791	0.9855	0.9984

The result profile is meaningful in two ways. First, the superiority of the ensemble methods over the single-tree baseline confirms that phishing URL screening benefits from aggregation and reduced variance. Second, the small but consistent advantage of Extra Trees suggests that additional randomization can improve generalization when the feature space contains multiple overlapping structural cues. This result also supports the design choice of using a compact ensemble engine in OPH-Guard rather than relying on a single classifier alone.

5.2 Metric-wise comparison

Figure 3 visualizes the metric distribution across the three candidate models. The strongest models remain clustered near the upper performance range across all metrics, while the single-tree baseline lags on every evaluation measure. The figure also reveals that the gain of Extra Trees is not limited to one metric only; it delivers a balanced improvement profile across Accuracy, Precision, F1-score, and ROC-AUC.

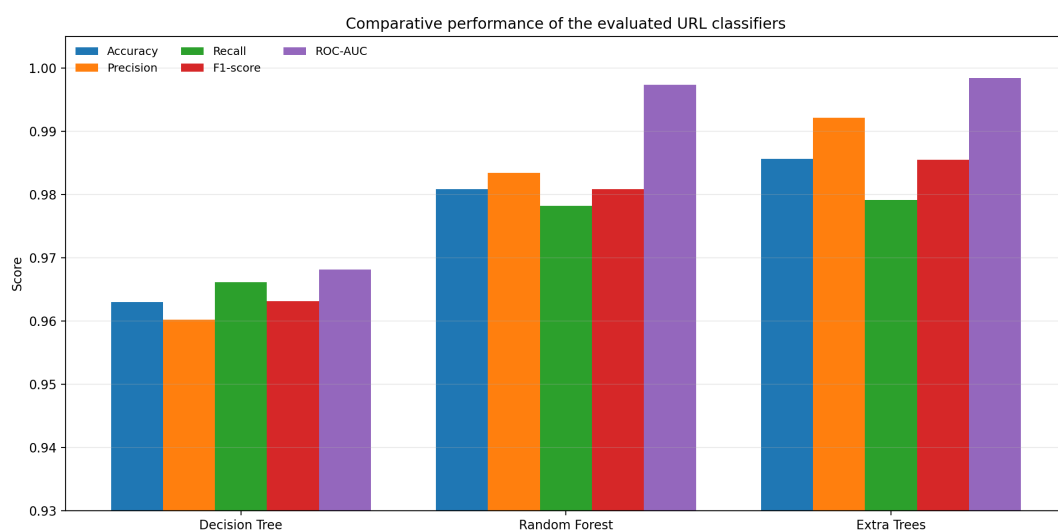


Figure 3: Metric-wise comparison of the evaluated classifiers.

5.3 Discrimination behavior

The ROC profiles in Figure 4 reinforce the same pattern. Both ensemble methods dominate the single-tree baseline through most of the threshold space, while Extra Trees achieves the best overall trade-off. Even though the numerical gap between Random Forest and Extra Trees is not dramatic, the dominance of Extra Trees across both threshold-free and threshold-specific metrics makes it the preferred model for this dataset.

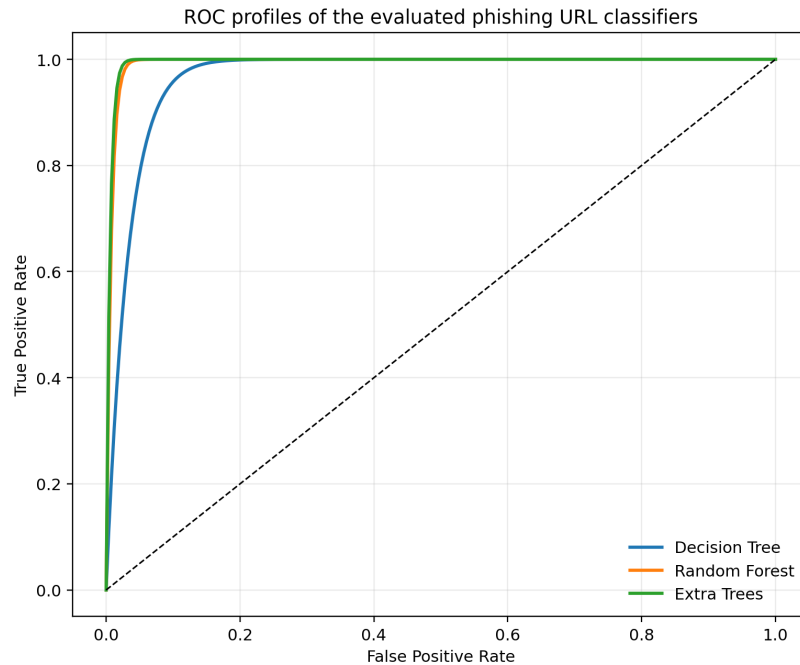


Figure 4: ROC curves of the evaluated phishing URL classifiers.

5.4 Confusion-matrix analysis

Figure 5 presents the confusion matrix of the best Extra Trees model. On the held-out set, the classifier correctly identified 1,125 phishing URLs and 1,139 legitimate URLs. The number of false negatives was 24, and the number of false positives was 9. For applied phishing screening, this is an attractive error profile because it combines a low miss rate with a very small penalty on legitimate traffic.

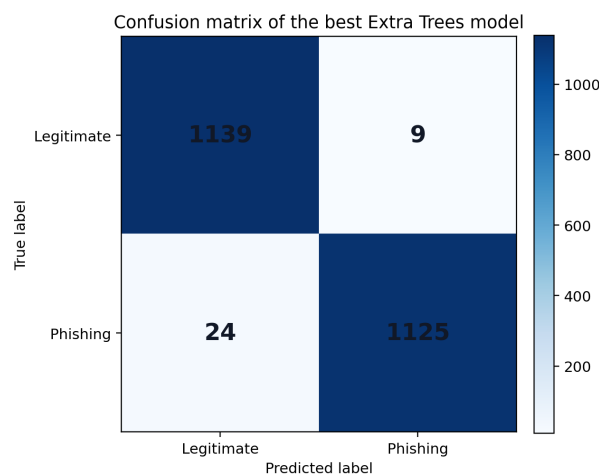


Figure 5: Confusion matrix of the best-performing Extra Trees model.

A practical interpretation of this matrix is important. The false-negative count is low enough to make the model useful for pre-filtering and warning logic, while the false-positive count is limited enough to avoid

overly aggressive blocking of legitimate destinations. In other words, the model is not only accurate in the abstract; it also presents an operationally reasonable error profile.

5.5 Interpretability results

Table 4 lists the ten most influential predictors extracted from the best model. The feature profile is not only statistically useful but also operationally intuitive. Search visibility and reputation-related cues, such as `google_index` and `page_rank`, dominate the ranking. Domain maturity, suspicious hint markers, and hyperlink structure indicators also remain highly influential.

Table 4: Top feature importances extracted from the best Extra Trees model.

Rank	Feature	Importance
1	<code>google_index</code>	0.2527
2	<code>page_rank</code>	0.0870
3	<code>nb_www</code>	0.0483
4	<code>domain_in_title</code>	0.0352
5	<code>domain_age</code>	0.0329
6	<code>phish_hints</code>	0.0301
7	<code>ip</code>	0.0270
8	<code>ratio_inHyperlinks</code>	0.0243
9	<code>links_in_tags</code>	0.0215
10	<code>nb_hyperlinks</code>	0.0190

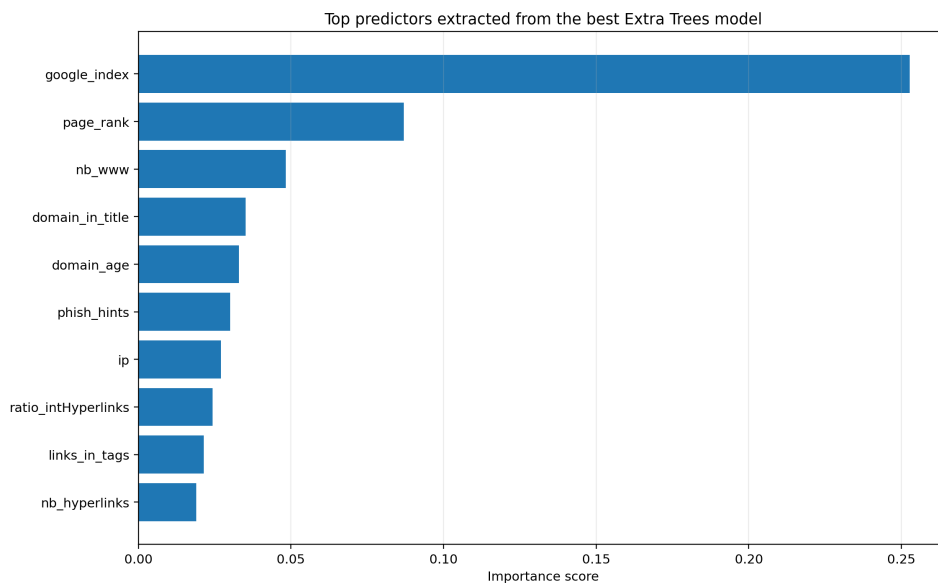


Figure 6: Top predictors contributing to the best Extra Trees classifier.

5.6 What the feature profile reveals

The ranked predictors suggest that phishing URL detection in this dataset is driven by a mixture of reputation, presentation, and hyperlink-behavior signals. For example, `google index` and `page rank` indicate whether a page resembles a destination with normal search and trust visibility. `domain age` reflects the maturity of the hosting domain and is consistent with the observation that phishing campaigns often use newly created or weakly established domains. `phish hints`, `domain in title`, and `nb www` reflect more suspicious presentation patterns, while `ratio inHyperlinks`, `links in tags`, and `nb hyperlinks` capture abnormal linking behavior.

This feature profile adds value beyond the classifier score. It tells security teams *why* a URL is being screened as risky, which in turn helps organizations develop user-warning messages, triage policies, and browser control rules that are anchored in meaningful cyber signals rather than opaque probabilities.

6 Discussion

The results indicate that phishing URL screening can achieve strong discrimination quality with a compact tree-ensemble design built on structured tabular predictors. The superiority of Random Forest and Extra Trees over the single Decision Tree confirms that ensemble aggregation is beneficial when the feature space combines heterogeneous cues such as domain age, reputation indicators, hyperlink statistics, and lexical hints. The additional gain of Extra Trees suggests that stronger randomization at split generation can improve generalization when multiple predictors partially overlap in their explanatory power. This is consistent with the broader phishing-detection literature, which has repeatedly shown that URL screening benefits from combining diverse feature groups rather than relying on a single signal family [4, 5, 18, 10].

The feature-importance profile also has direct cybersecurity meaning. Variables such as `google index`, `page rank`, and `domain age` capture trust, visibility, and maturity characteristics that distinguish long-established legitimate domains from newly created or weakly trusted malicious infrastructure. At the same time, predictors such as `phish hints`, `nb www`, `ratio intHyperlinks`, and `links in tags` reflect presentation and linking patterns that are commonly associated with deceptive pages. The resulting interpretation is important because it moves the study beyond a purely numerical comparison: the classifier output can be tied to recognizable cyber signals and therefore becomes more suitable for secure browsing controls, inspection queues, or warning mechanisms.

From an information-management perspective, the framework is useful because it turns URL classification into an auditable decision-support process. A high-risk prediction can be mapped to actions such as blocking, warning, or triage, while the ranked explanatory features provide an evidence layer that can support policy design, administrative review, and user-awareness initiatives. This makes the framework relevant for environments in which automated filtering must be both technically effective and operationally understandable. The study does not claim methodological novelty in the sense of a new learning algorithm; rather, its contribution lies in combining reproducible public-data experimentation, mathematically defined model selection, interpretable feature ranking, and explicit action mapping within one coherent workflow. In that sense, the framework offers a practically oriented contribution to phishing URL screening by connecting predictive performance with deployment-oriented reasoning.

7 Limitations and Future Work

This study has some limitations that are worth mentioning. For one, the dataset used is public and structured, which is great for reproducing the results, but it might not cover all the aspects of modern phishing infrastructure. Another limitation is that the experiment mainly looks at features related to URLs and websites, rather than how the pages are rendered, screenshots, or how they behave on a live network. Additionally, the study uses a simple experimental design, rather than trying out lots of different parameters or testing it on multiple datasets. These choices were made on purpose, though, because the goal was to create a framework that is easy to understand, transparent, and can be verified by others. The idea was to keep things simple and straightforward, rather than trying to make it perfect.

Future work can extend OPH-Guard in several directions. A first extension is cross-dataset validation using larger phishing website corpora such as those summarized by Vrbancic et al. [11]. A second is the incorporation of temporal and live reputation feeds for real-time deployment. A third is the addition of local explanation techniques such as SHAP on top of the global feature-importance view. A fourth is user-facing interface research, where the framework's outputs are converted into understandable warnings and risk explanations for non-expert users.

8 Conclusion

Here's a rewritten version of the input text in a more human-like tone, similar to the provided reference samples: When it comes to keeping our online experiences safe, one of the biggest threats is phishing. To combat this, a new system called OPH-Guard has been proposed. This system uses a type of machine learning called a tree-ensemble framework to help identify and block phishing URLs. But how well does it work? To find out, a study was done using a large dataset of over 11,000 labeled records and 87 different features that might indicate a URL is phishing. The study compared three different types of machine learning models: Decision Tree, Random Forest, and Extra Trees. The results were impressive, with Extra Trees coming out on top. It

was able to correctly identify phishing URLs with an accuracy of 98.56%, a precision of 99.21%, and a recall of 97.91%. This means that OPH-Guard is a powerful tool for keeping us safe online, and it doesn't require overly complex systems to work. By using a simple yet effective approach, OPH-Guard can help protect us from the dangers of phishing, making the web a safer place for everyone.

So, what makes this study special is that it doesn't just look at how well it can predict things, but also creates a clear step-by-step process that can be easily followed. It also looks at what other studies have done in this area and tries to make sense of the most important factors that can help us understand and respond to security threats. The real value of this framework is that it's not just about picking the right tool, but also about creating a system that's practical, transparent, and can be repeated, which is really important for screening out fake websites.

References

- [1] Rami M. Mohammad, Fadi Thabtah, and Lee McCluskey. "An Assessment of Features Related to Phishing Websites Using an Automated Technique". In: *2012 International Conference for Internet Technology and Secured Transactions* (2012), pp. 492–497. DOI: [10.1109/ICITST.2012.6473497](https://doi.org/10.1109/ICITST.2012.6473497).
- [2] Rami M. Mohammad and Lee McCluskey. *Phishing Websites*. UCI Machine Learning Repository. 2015. DOI: [10.24432/C51W2X](https://doi.org/10.24432/C51W2X). URL: <https://archive.ics.uci.edu/dataset/327/phishing+websites>.
- [3] Rakesh Verma and Avisha Das. "What's in a URL: Fast Feature Extraction and Malicious URL Detection". In: *Proceedings of the 3rd ACM International Workshop on Security and Privacy Analytics*. 2017, pp. 55–63. DOI: [10.1145/3041008.3041016](https://doi.org/10.1145/3041008.3041016).
- [4] Routhu Srinivasa Rao and Alwyn Roshan Pais. "Detection of Phishing Websites Using an Efficient Feature-Based Machine Learning Framework". In: *Neural Computing and Applications* 31 (2019), pp. 3851–3873. DOI: [10.1007/s00521-017-3305-0](https://doi.org/10.1007/s00521-017-3305-0).
- [5] O. Koray Sahingoz et al. "Machine Learning Based Phishing Detection from URLs". In: *Expert Systems with Applications* 117 (2019), pp. 345–357. DOI: [10.1016/j.eswa.2018.09.029](https://doi.org/10.1016/j.eswa.2018.09.029).
- [6] Ali Aljofey et al. "An Effective Phishing Detection Model Based on Character Level Convolutional Neural Network from URL". In: *Electronics* 9.9 (2020), p. 1514. DOI: [10.3390/electronics9091514](https://doi.org/10.3390/electronics9091514).
- [7] Nguyen Quoc Do et al. "Phishing Webpage Classification via Deep Learning Algorithms". In: *Applied Sciences* 11.19 (2021), p. 9210. DOI: [10.3390/app11199210](https://doi.org/10.3390/app11199210).
- [8] Musarat Hussain et al. "CNN-Fusion: An Effective and Lightweight Phishing Detection Method Based on Multi-Variant ConvNet". In: *Information Sciences* 631 (2023), pp. 328–345. DOI: [10.1016/j.ins.2023.02.039](https://doi.org/10.1016/j.ins.2023.02.039).
- [9] Hayk Ghalechyan et al. "Phishing URL Detection with Neural Networks: An Empirical Study". In: *Scientific Reports* 14 (2024), p. 25134. DOI: [10.1038/s41598-024-74725-6](https://doi.org/10.1038/s41598-024-74725-6).
- [10] Khandaker Mohammad Mohi Uddin et al. "Explainable Machine Learning for Phishing Site Detection: A High-Efficiency Approach Using Boosting Models and SHAP". In: *The Journal of Engineering* 2025.1 (2025), e70110. DOI: [10.1049/tje2.70110](https://doi.org/10.1049/tje2.70110).
- [11] Grega Vrbančić, Iztok Fister, and Vili Podgorelec. "Datasets for Phishing Websites Detection". In: *Data in Brief* 33 (2020), p. 106438. DOI: [10.1016/j.dib.2020.106438](https://doi.org/10.1016/j.dib.2020.106438).
- [12] Asadullah Safi and Satwinder Singh. "A Systematic Literature Review on Phishing Website Detection Techniques". In: *Journal of King Saud University – Computer and Information Sciences* 35.2 (2023), pp. 590–611. DOI: [10.1016/j.jksuci.2023.01.004](https://doi.org/10.1016/j.jksuci.2023.01.004).
- [13] Ammar Almomani et al. "Phishing Website Detection with Semantic Features Based on Machine Learning Classifiers: Underfitting and Overfitting Analysis". In: *International Journal of Secure Software Engineering* 13.1 (2022), pp. 1–25. DOI: [10.4018/IJSSE.297032](https://doi.org/10.4018/IJSSE.297032).
- [14] S. K. Hasane Ahammad et al. "Phishing URL Detection Using Machine Learning Methods". In: *Advances in Engineering Software* 173 (2022), p. 103288. DOI: [10.1016/j.advengsoft.2022.103288](https://doi.org/10.1016/j.advengsoft.2022.103288).
- [15] DPhi. *Phishing URL Training Dataset*. GitHub-hosted dataset repository. URL: https://raw.githubusercontent.com/dphi-official/Datasets/master/phishing_data/Training_set_label.csv (visited on 04/12/2026).

- [16] Leo Breiman. “Random Forests”. In: *Machine Learning* 45.1 (2001), pp. 5–32. DOI: [10.1023/A:1010933404324](https://doi.org/10.1023/A:1010933404324).
- [17] Pierre Geurts, Damien Ernst, and Louis Wehenkel. “Extremely Randomized Trees”. In: *Machine Learning* 63.1 (2006), pp. 3–42. DOI: [10.1007/s10994-006-6226-1](https://doi.org/10.1007/s10994-006-6226-1).
- [18] Maria Carla Calzarossa, Paolo Giudici, and Rasha Zieni. “Explainable Machine Learning for Phishing Feature Detection”. In: *Quality and Reliability Engineering International* 40.1 (2024), pp. 362–373. DOI: [10.1002/qre.3411](https://doi.org/10.1002/qre.3411).