

**“ShieldPhish: A Hybrid Machine Learning and Heuristic
Rule-Based System for Explainable Phishing URL
Detection”**

**A research report submitted to
National College of Computer Studies
Paknajol, Kathmandu**

**Submitted by
Dinesh Khadka
Gandhi Raj Giri
September 2026**

ACKNOWLEDGEMENT

This research work would not have been possible without the support, guidance, and encouragement of many individuals, to whom we would like to express our sincere gratitude. We would like to express our gratitude to National College of Computer Studies for providing the opportunity to conduct this research work. The trust and encouragement of the institution motivated us to carry out this research. We would like to acknowledge Vice Principal **Mr. Santosh Raj Maskey**, Director **Mr. Prajwal Baniya Chhetri**, and Coordinator **Mr. Ramesh Phuyal** for their continuous support and encouragement throughout our research work.

Finally, we would like to express our deep appreciation to the faculty members, students, and administrative staff of NCCS, as well as our family members and friends who directly or indirectly supported the completion of this research.

Dinesh Khadka

Gandhi Raj Giri

ABSTRACT

Phishing URL detection is an essential topic in the area of cybersecurity because malicious URLs can be exploited to extract sensitive information from consumers. Machine learning (ML) models can give effective URL classification, but their predictions can be difficult to interpret and their performance can degrade when meeting URL patterns that differ from those represented in the training data. In contrast, heuristic rule-based methods deliver fast and definite conclusions but are less flexible and may wrongly indicate safe URLs that contain worrisome structural traits. To overcome these constraints, in this research, we propose ShieldPhish, a hybrid phishing URL detection system that integrates machine learning, heuristic rule analysis, and Explainable Artificial Intelligence (XAI). The system extracts a 112-dimensional feature vector, which includes structural and lexical features, security cues, Shannon entropy, and character level TF-IDF features. These features are used to train and test many machine learning classifiers including Random Forest, Xgboost, Calibrated Linear SVM and a bespoke NumPy based Logistic Regression model. The final phishing classification is obtained by combining the machine learning probability with the heuristic rule scores using a bounded fusion technique. Interpretability of the machine learning predictions is provided using SHAP based feature attribution. The system was evaluated on a combined dataset of 957,334 URLs using 5-fold stratified cross validation. The presented findings indicate that Random Forest obtained 99.40% accuracy while XGBoost achieved 93.16% accuracy with ROC AUC of 0.9809. ShieldPhish additionally provides feature-level explanations to enhance the interpretability of detection results. The proposed approach integrates statistical machine learning, deterministic heuristic analysis, and explainability to enable interpretable and low-latency phishing URL detection.

KEYWORDS: *Phishing URL Detection, Hybrid Machine Learning, Explainable AI (XAI), SHAP, XGBoost, Random Forest, Heuristic Rule Engine, Cybersecurity.*

LIST OF ABBREVIATIONS

API	Application Programming Interface
ASGI	Asynchronous Server Gateway Interface
AUC	Area Under Curve
BERT	Bidirectional Encoder Representations from Transformers
CORS	Cross-Origin Resource Sharing
CPU	Central Processing Unit
CRISP-DM	Cross-Industry Standard Process for Data Mining
DFD	Data Flow Diagram
DGA	Domain Generation Algorithm
DNS	Domain Name System
FFT	Fast Fourier Transform
FN	False Negative
FP	False Positive
HTTP	Hypertext Transfer Protocol
HTTPS	Hypertext Transfer Protocol Secure
JSON	JavaScript Object Notation
ML	Machine Learning
NLP	Natural Language Processing
PDF	Portable Document Format
REST	Representational State Transfer
ROC	Receiver Operating Characteristic
SHAP	Shapley Additive Explanations
SVM	Support Vector Machine
TF-IDF	Term Frequency–Inverse Document Frequency
TLS	Transport Layer Security
URL	Uniform Resource Locator
XAI	Explainable Artificial Intelligence
XGBoost	Extreme Gradient Boosting

TABLE OF CONTENTS

ACKNOWLEDGEMENT	i
ABSTRACT	ii
LIST OF ABBREVIATIONS	iii
1 INTRODUCTION	1
1.1 Introduction	1
1.1.1 Problem Statement	1
1.2 Objectives	2
1.2.1 Scope and Limitations	3
1.3 Development Methodology	4
1.4 Report Organization	6
2 BACKGROUND STUDY AND LITERATURE REVIEW	8
2.1 Background Study	8
2.2 Literature Review	9
3 SYSTEM ANALYSIS	13
3.1 System Analysis	13
3.1.1 Requirement Analysis	13
3.1.2 Feasibility Analysis	16
3.2 Analysis	17
4 SYSTEM DESIGN	22
4.1 Design	22
4.2 Algorithms Details	27
5 IMPLEMENTATION AND TESTING	30
5.1 Implementation	30
5.1.1 Tools Used	30
5.1.2 Dataset Description	31
5.1.3 Implementation Details of Modules	32
5.2 Testing	34
5.2.1 Test Cases for Unit Testing	34
5.2.2 Test Cases for System Testing	37
5.3 Result Analysis	40
6 CONCLUSION AND FUTURE RECOMMENDATIONS	45
6.1 Conclusion	45
6.2 Future Recommendations	45
REFERENCES	47

LIST OF FIGURES

1.1	Iterative Development of shieldPhish	5
3.1	Use Case Diagram of shieldPhish	14
3.2	Class Diagram of shieldPhish	18
3.3	Object Diagram of shieldPhish	19
3.4	State Diagram of shieldPhish	20
3.5	Sequence Diagram of shieldPhish	20
3.6	Activity Diagram of shieldPhish	21
4.1	Refinement of Class Diagram of shieldPhish	22
4.2	Refinement of Sequence Diagram of shieldPhish	23
4.3	Refinement of Activity Diagram of shieldPhish	24
4.4	Component Diagram of shieldPhish	25
4.5	Deployment Diagram of shieldPhish	26
5.1	Performance comparison of base machine learning classifiers on the extracted feature space of shieldPhish.	41
5.2	Receiver Operating Characteristic (ROC) curves demonstrating the true pos- itive vs. false positive trade-off for evaluated classifiers of shieldPhish.	42
5.3	Ablation study illustrating the impact of heuristic rule weight (α) on False Negatives and overall model Accuracy of shieldPhish.	43
5.4	TreeSHAP summary plot revealing the global feature importance and direc- tional impact of the top features on the prediction output of shieldPhish.	44

LIST OF TABLES

2.1	Synthesis of Selected Recent Studies and Their Relevance to ShieldPhish . . .	11
3.1	Gantt Chart	17
5.1	Dataset Characteristics	32
5.2	Test Cases for Unit Testing	35
5.3	Test Cases for System Testing	38
5.4	Classifier Performance Metrics Comparison	40
5.5	Ablation of Rule Weight (α) on XGBoost Performance	43

CHAPTER 1

INTRODUCTION

1.1 Introduction

Phishing URLs are a major cybersecurity problem, often deceiving unwitting users into providing key credentials, leading to financial fraud and privacy violations (Aljofey et al., 2021). The problem has been compounded by the rise of automated phishing toolkits and Domain Generation Algorithms (DGAs) that allow attackers to dynamically build convincing malicious URLs at scale (Xuan & Nguyen, 2022). Consequently, newly generated *zero-day* phishing URLs are often missed by traditional signature-based mechanisms or static blacklists, as such protective measures rely significantly on past threat intelligence repositories and suffer from time delays. In this research, we propose ShieldPhish, a complete prototype for hybrid phishing URL detection, which integrates statistical machine learning and deterministic heuristic reasoning. To combat the opacity of solely statistical models, ShieldPhish integrates machine-learning predictions with heuristic rule outputs and SHAP based feature attributions to provide more information about the reasons affecting each URL evaluation. The framework works through a simplified pipeline including preprocessing of data, feature engineering, model training, hybrid fusion and online deployment. The system in particular extracts a 112-dimensional feature vector, including four types of features: structural lexical statistics, security protocol indicators, information-theoretic Shannon entropy, and character-level TF-IDF n-grams. The processed feature set is fed into different supervised classifiers such as Random Forest, XGBoost, Calibrated Linear SVM, and a custom NumPy based Logistic Regression in an effort to compare the classifier performance and select the best performing model. The main statistical classifier is then fused with a deterministic engine analyzing six heuristics rules to examine if heuristic fusion enhances resilience of detection against evasive URLs and reduces false negative rates. Finally, local feature attributions are constructed with TreeSHAP to understand the machine learning component, and the rule engine gives transparent structural assessments. These capabilities are shown by deploying the system through a FastAPI backend and a React based frontend web interface to provide real-time URL classification and human-readable diagnostic explanations for security analysts.

1.1.1 Problem Statement

Phishing assaults have progressed from basic spam emails to sophisticated, highly targeted cyber campaigns that employ dynamic domain generation algorithms (DGAs), automated phishing toolkits, and intricate URL obfuscation techniques. Today attackers commonly employ ephemeral domains or spoofed URL parameters to disguise themselves as legitimate web services and deceive users into providing sensitive authentication credentials, personal data, and financial information. In this research three problems are considered to detect malicious URLs efficiently in real world operating environment:

Temporal Lag of Static Blacklists Conventional perimeter defenses heavily rely on domain and URL blacklists, such as PhishTank (2025) and Google Safe Browsing. However, modern phishing campaigns often utilize short-lived URLs that remain active for only a brief period before being abandoned or replaced. Consequently, blacklist-based approaches may experience temporal delay, which may leave a detection gap before blacklist inclusion during the critical initial window of an attack (the zero-day gap).

Opacity and Brittle Generalization of Standalone ML Supervised machine learning (ML) classifiers (such as Random Forest and XGBoost) often achieve high baseline accuracy on historical training datasets. However, some high performing ML models provide limited intrinsic interpretability. Security Operations Center (SOC) analysts cannot easily inspect or verify the underlying decision logic of an ML prediction. Furthermore, statistical models may exhibit reduced generalization performance when encountering previously unseen or adversarially modified URL patterns, potentially increasing the likelihood of false negative predictions (Shirazi et al., 2022).

Inflexible Determinism of Rule Engines Static heuristic rule engines assess structural anomalies (e.g., IP addresses in place of domain names, missing HTTPS protocols, or excessive subdomain depth). While rule engines offer low-latency, deterministic threat scores, they lack adaptability. Attackers may circumvent static heuristic rules by intentionally modifying URL strings (e.g., acquiring free SSL certificates to enable HTTPS), and strict rule engines may produce higher false positive rates on complex, legitimate enterprise URLs (Jain & Sharma, 2023). To address these challenges, this project proposes ShieldPhish, a hybrid phishing URL detection framework that combines deterministic heuristic reasoning with supervised machine learning and SHAP based Explainable Artificial Intelligence (XAI). By integrating lightweight rule-based analysis with statistical learning, the proposed system aims to improve detection robustness with the aim of reducing false negatives, and to provide transparent, human interpretable explanations for predictions generated by the primary tree-based classifiers.

1.2 Objectives

The main objective of this project is to develop ShieldPhish, a hybrid phishing URL detection framework that integrates statistical machine learning with deterministic heuristic rule analysis and Explainable Artificial Intelligence (XAI).

The specific objectives are:

- To extract a high-dimensional feature vector encompassing structural lexical statistics, security protocol indicators, Shannon entropy, and character-level TF-IDF representations from raw URL strings.

1.2.1 Scope and Limitations

Scope

The scope of the ShieldPhish project encompasses the development and evaluation of a static, URL-based hybrid classification framework. The system focuses on extracting a 112-dimensional feature vector spanning four core categories: structural lexical statistics (e.g., URL length, subdomain count), security indicators (e.g., HTTPS usage, presence of IP addresses), information-theoretic Shannon entropy, and 100 character-level TF-IDF n-grams. The framework employs a hybrid decision engine, combining these features via supervised machine learning with deterministic heuristic rules. Furthermore, it integrates TreeSHAP to provide real-time, feature-level explanations for predictions generated by the primary selected model. To demonstrate operational viability, the project includes the deployment of a prototype utilizing a FastAPI backend for high-speed inference and a React based frontend for user interaction. The experimental evaluation is strictly limited to the supplied, deduplicated URL dataset using stratified cross-validation; large-scale temporal deployment evaluations in a production environment are outside the current scope of this project.

Limitations

While ShieldPhish offers robust detection capabilities, the framework is subject to several inherent technical and methodological limitations:

- **No Live Web Content Inspection:** The system exclusively analyzes static URL strings and does not fetch or inspect the Document Object Model (DOM), HTML content, or JavaScript behavior of the target page.
- **Unfollowed Redirections:** The current implementation does not automatically follow HTTP redirects, meaning malicious payloads hidden behind URL shorteners or complex redirection chains may evade analysis.
- **Homograph and Encoding Evasion:** Sophisticated homograph attacks (e.g., Puny-code manipulation) or advanced character encodings may circumvent the character-level tokenizers without specialized Unicode normalization layers.
- **Dependence on Static Lexical Signals:** As a strictly lexical analyzer, the model cannot identify compromised legitimate domains if the URL structure remains entirely benign.
- **Dataset Distribution Limitation:** The model's predictive performance is intrinsically tied to the characteristics of the datasets used for training and evaluation, which may not encompass all global phishing variations.
- **Temporal Concept Drift:** Future phishing campaigns may evolve to use novel URL obfuscation patterns that differ significantly from those represented in the historical training data.
- **Source Bias:** Because phishing and legitimate samples often originate from different dataset repositories, source specific artifacts could unintentionally influence classification performance.

- **Static Heuristic Maintenance:** The deterministic heuristic rules lack dynamic adaptability and require manual revision by security analysts when attacker behavior changes.

1.3 Development Methodology

The proposed framework was developed following an iterative research methodology based on the Cross Industry Standard Process for Data Mining (CRISP-DM). The CRISP-DM framework was selected because phishing URL detection is inherently a data-centric cybersecurity problem requiring iterative cycles of data preparation, feature engineering, model refinement, and empirical evaluation. Unlike sequential software development models, CRISP-DM facilitates incremental optimization driven by experimental evidence, making it exceptionally well suited for hybrid machine learning systems.

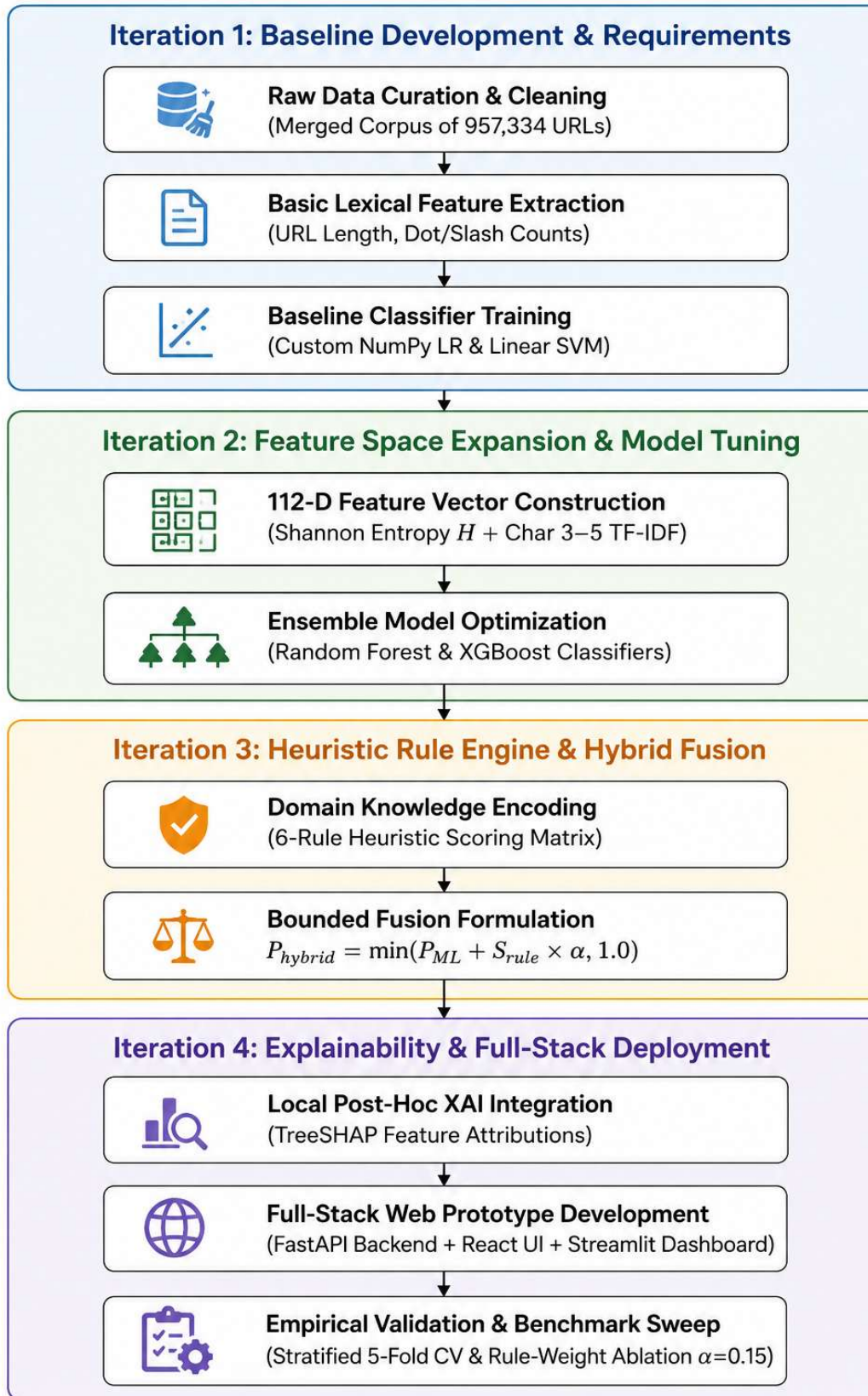


Figure 1.1: Iterative Development of shieldPhish

Iterative Development Phases of ShieldPhish:

Iteration 1: Baseline Establishment Baseline requirements were established, the raw dataset was curated and deduplicated, and initial structural lexical features were extracted. Baseline classifiers including a custom NumPy based Logistic Regression model and a Linear Support Vector Machine were implemented to establish baseline predictive performance.

Iteration 2: Feature Space Expansion This phase focused on improving the discriminative capability of the feature representation through the incorporation of information theoretic and textual descriptors. The feature space was expanded to 112 dimensions by integrating Shannon entropy calculations and character-level TF-IDF n-grams, after which Random Forest and XGBoost ensemble classifiers were trained and hyperparameter tuned.

Iteration 3: Hybrid Fusion Integration To enhance system robustness against zero-day URL evasions, a deterministic heuristic inference module was introduced to complement statistical prediction by encoding domain knowledge associated with high risk URL structures. These heuristic penalty scores were dynamically combined with machine learning probability predictions through a bounded fusion equation:

$$P_{\text{hybrid}} = \min(P_{\text{ML}} + \alpha S_{\text{rule}}, 1.0) \quad (1.1)$$

Iteration 4: Explainability & Prototype Deployment Local post-hoc explainability was incorporated through TreeSHAP to generate feature-level attribution scores, enabling analysts to inspect the contribution of individual URL characteristics to each prediction. A full-stack web-based prototype consisting of a FastAPI backend, React frontend, and Streamlit diagnostic dashboard was integrated to demonstrate real-time threat evaluation.

System Validation

The proposed framework was validated through stratified cross validation, rule-weight sensitivity analysis, and performance benchmarking to evaluate predictive accuracy, robustness, interpretability, and computational efficiency across successive development iterations.

1.4 Report Organization

Chapter 1: Introduction

This chapter introduces the project, defines the problem, states the objectives, explains scope and limitations, and describes the development methodology.

Chapter 2: Background Study and Literature Review

This chapter discusses the concepts required for phishing URL detection and reviews related research works to identify the research and implementation gap.

Chapter 3: System Analysis

This chapter presents requirement analysis, feasibility analysis, and the analysis models used to understand system processes and data flow.

Chapter 4: System Design

This chapter explains the proposed system design, working mechanism, diagrams, and the algorithm-level steps used in the system.

Chapter 5: Implementation and Testing

This chapter describes the tools used, module wise implementation details, testing strategy with test cases, and result analysis.

Chapter 6: Conclusion and Future Recommendations

This chapter concludes the project based on the results obtained and provides future recommendations.

CHAPTER 2

BACKGROUND STUDY AND LITERATURE REVIEW

2.1 Background Study

The rapid advancement of Internet technology has altered communication, financial transactions, cloud services, and e-commerce, but it has also exposed users to complex cyber risks. Phishing is still one of the most common cybercrimes among these dangers, using human trust to steal critical information like login credentials, bank details, and personal data through fake websites and communication methods. Phishing assaults have come a long way from basic spam emails to extremely complex campaigns that use the likes of Domain Generation Algorithms (DGAs), URL shortening services, homograph attacks, quick flux hosting, free SSL/TLS certificates, and automated phishing toolkits. These methods allow attackers to create very realistic phishing websites that are similar to legitimate businesses, yet go beyond existing detection methods. Malicious URLs are becoming more and more difficult to detect. The traditional phishing detection mostly relies on black list services such as Google Safe Browsing and PhishTank. Blacklists are effective against previously reported threats but are incapable of identifying freshly emerging phishing URLs before they are placed into threat intelligence databases, and so leave a major zero-day detection gap. To overcome this constraint, researchers have increasingly turned to machine learning approaches such as Logistic Regression, Support Vector Machines (SVM), Random Forest and XGBoost, which learn discriminative patterns from URL characteristics. However, these models are frequently black boxes and may suffer from lower generalization when faced with previously undiscovered phishing methods. Heuristic or rule-based detection approaches offer rapid and interpretable detection by testing suspicious URL properties such as aberrant URL length, excessive subdomains, embedded IP addresses and missing HTTPS protocols. While computationally efficient, the static heuristic criteria are not adaptive and can lead to increased false positive rates if attackers change URL structures or utilize authentic security certificates. Recent research has focused on hybrid phishing detection frameworks that combine heuristic reasoning with machine learning to improve detection accuracy and robustness. Moreover, Explainable Artificial Intelligence (XAI) approaches, like SHAP (SHapley Additive exPlanations), have increased model transparency by providing feature-level explanations, allowing analysts to grasp the reasoning behind machine learning predictions. Driven by these issues, this work introduces ShieldPhish, a hybrid phishing URL detection framework that utilizes the combination of heuristic rule-based analysis, supervised machine learning and SHAP based explainability. The proposed approach combines deterministic rule analysis with statistical learning to improve detection accuracy, ensure robustness against developing phishing assaults, and provide transparent decision assistance suitable for real-world cybersecurity applications.

2.2 Literature Review

Research published between 2021 and 2025 has expanded phishing URL detection beyond traditional blacklist based approaches to include supervised machine learning, deep neural architectures, hybrid rule-model fusion, and explainable artificial intelligence (XAI). This section synthesizes representative studies across these paradigms to contextualize the design of ShieldPhish.

Lexical Feature Based Detection

Static analysis techniques extract structural and lexical characteristics directly from URL strings without performing live web crawling or network requests. Xuan and Nguyen (2022) explored structural anomalies, indicating that Shannon entropy analysis effectively identifies algorithmically generated domain names. Building on static feature extraction, Khan and Soman (2023) evaluated TF-IDF n-gram text representations for static phishing URL detection, reporting that high-dimensional lexical tokenization can effectively distinguish malicious URLs without requiring live network requests.

- **Strengths:** Lexical approaches operate at low computational cost and avoid network latency.
- **Limitations:** Standalone lexical analysis relies heavily on historical character distributions and can be circumvented by attackers using legitimate-looking domain structures.

Machine Learning & Ensemble Approaches

To capture non-linear feature interactions, several studies examine tree-based ensembles and gradient boosting algorithms. Breiman (2001) established the foundational Random Forest architecture, which remains widely used for its robust pattern recognition across large tabular datasets. Chen and Guestrin (2016) later introduced Extreme Gradient Boosting (XGBoost), providing a highly scalable tree-boosting system. Applying these foundations to cybersecurity, Tantoso and Raj (2022) investigated XGBoost for high-throughput malicious URL classification, observing that gradient boosted decision trees achieved strong evaluation speeds and high AUC ROC metrics compared to traditional algorithms.

- **Strengths:** Supervised ML algorithms offer high predictive accuracy, fast evaluation speeds, and strong pattern recognition across large datasets.
- **Limitations:** Some complex machine-learning classifiers provide limited intrinsic interpretability, making the factors underlying individual predictions difficult to examine directly. Furthermore, statistical models may experience reduced generalization when evaluated on previously unseen URL patterns.

Deep Learning & Transformer Architectures

To reduce reliance on manually engineered features, recent literature explores end-to-end deep learning and transformer language models. Aljofey et al. (2021) proposed a detection model using character-level Convolutional Neural Networks (CNNs), demonstrating the utility of

treating URLs as sequential character inputs. Expanding on sequential analysis, Maneriker et al. (2021) developed URLtran, leveraging pre-trained Transformer models to perform contextual tokenization of URLs without manual rule creation. More recently, Li et al. (2023) proposed PhishFormer, a transformer architecture designed to detect previously unseen URL patterns using self-attention mechanisms. Furthermore, Bozkir et al. (2021) introduced PhishNet, highlighting that deep learning frameworks can be effectively extended to visual-similarity-based phishing detection.

- **Strengths:** Deep learning and transformer models automatically extract subtle sequence representations and can capture complex sequential patterns.
- **Limitations:** Compared with lightweight conventional classifiers, some deep neural and transformer-based architectures may impose greater computational requirements and often provide limited intrinsic interpretability.

Hybrid Detection Frameworks

Recognizing that machine-learning classifiers may experience reduced performance when evaluated on URL patterns that differ from those represented during training, recent studies explore hybrid architectures combining deterministic rule reasoning with data-driven models. Shirazi et al. (2022) explored lightweight URL-based detection, illustrating the need to balance computational efficiency with detection accuracy. Jain and Sharma (2023) investigated heuristic-based threat score calibration for hybrid phishing detection engines, demonstrating how combining deterministic rule reasoning with statistical models improves robustness. Kumar and Varadharajan (2023) further evaluated this paradigm, noting that parameter tuning of rule weights in hybrid security classifiers helps preserve specificity while overriding low-confidence machine learning predictions.

- **Strengths:** Hybrid approaches combine learned statistical patterns with deterministic rule-based evidence, allowing structurally suspicious characteristics to influence predictions that may otherwise depend solely on the learned model.
- **Limitations:** Among the studies reviewed, existing hybrid models often lack local feature explainability.

Explainable AI (XAI) in Cybersecurity

A major challenge in deploying machine learning for operational cybersecurity is that limited model interpretability can make automated predictions more difficult for security analysts to examine and audit. Singh (2022) highlighted in a systematic review that post-hoc explainability is critical for security analyst trust. To address this, Saeed et al. (2023) evaluated SHAP (SHapley Additive exPlanations) originally formalized by Lundberg and Lee (2017) in cyber threat detection, finding that post-hoc feature attributions can provide useful context during incident response. Expanding on this, Kumar and Varadharajan (2024) introduced XAI Phish, demonstrating that TreeSHAP produces feature-level attributions that can be mapped to human-readable explanations of model behavior.

- **Strengths:** Explainable AI provides transparent, human interpretable feature attributions for complex machine learning models.
- **Limitations:** XAI methods are post-hoc interpretability tools; they provide feature attribution but do not guarantee the correctness of the underlying prediction.

Table 2.1: Synthesis of Selected Recent Studies and Their Relevance to ShieldPhish

Study	Approach	Dataset / Eval	Key Finding	Limitation	Relevance to ShieldPhish
Xuan & Nguyen (2022)	Entropy Analysis	<i>Not specified</i>	Entropy improved detection of algorithmically generated domains	Fixed to static metrics	Motivates Shannon entropy integration
Khan & So-man (2023)	TF-IDF Lexical Features	<i>Not specified</i>	Lexical tokenization distinguished malicious URLs	Standalone lexical limits	Motivates TF-IDF tokenization
Tantoso & Raj (2022)	XGBoost Classification	<i>Not specified</i>	High evaluation speed and AUC-ROC	Limited interpretability	Motivates XGBoost baseline
Aljofey et al. (2021)	Character-level CNNs	<i>Not specified</i>	Captured sequential patterns in URL strings	High computational cost	Motivates character analysis
Maneriker et al. (2021)	BERT (URL-Tran)	<i>Not specified</i>	Contextual tokenization without manual rules	Slow inference speed	Justifies static feature extraction
Li et al. (2023)	Lightweight Transformers	<i>Not specified</i>	Contextual learning for unseen patterns	Greater computational requirements	Highlights need for efficiency
Jain & Sharma (2023)	Heuristic Score Calibration	<i>Not specified</i>	Improved robustness through combined rule- and ML-based scoring	Lacks post-hoc XAI	Motivates bounded fusion strategy
Saeed et al. (2023)	ML + SHAP	<i>Not specified</i>	Local feature transparency for analysts	No rule-based safety net	Motivates SHAP integration
Kumar & Varadharajan (2024)	XAI Phish (TreeSHAP)	<i>Not specified</i>	Feature attributions mapped to human-interpretable logic	Standalone ML only	Validates TreeSHAP utility

Research Gap

Among the studies reviewed, existing approaches tend to emphasize particular aspects of phishing URL detection, such as predictive performance, representation learning, heuristic reasoning, or post-hoc explainability, rather than evaluating these components together within a computationally lightweight detection framework. This observation motivates ShieldPhish,

which investigates the integration of static URL feature extraction, supervised machine learning, deterministic heuristic scoring, and SHAP-based feature attribution within a unified prototype.

CHAPTER 3

SYSTEM ANALYSIS

3.1 System Analysis

The ShieldPhish detection system is analyzed in order to ascertain whether or not the proposed solution is feasible and clear in its description before proceeding to design in detail. This is a crucial phase in which one is required to identify what is to be done by the system, what quality is to be satisfied by the system, whether or not the project is feasible, and how the system is to be analyzed using appropriate system analysis techniques.

3.1.1 Requirement Analysis

Requirement analysis is a fundamental step in software and system engineering that defines the necessary behavior, constraints, and quality attributes of the proposed framework. For ShieldPhish, requirements are divided into Functional Requirements, which specify the operational tasks the system must perform, and non-functional requirements, which define quality attributes such as performance, accuracy, transparency, and scalability.

Functional Requirement

The functional requirements define the core functionalities that the proposed ShieldPhish system must provide to support phishing URL detection and analysis.

URL Input Processing:

The system shall accept URL inputs through an interactive web interface, REST API requests, and batch processing for phishing detection.

Feature Extraction:

The system shall automatically extract relevant lexical, security-related, statistical, and character-level textual features from the input URL to generate the feature representation required for classification.

Phishing Detection:

The system shall classify the input URL as Legitimate or Phishing using trained supervised machine learning models.

Heuristic Risk Analysis:

The system shall evaluate the input URL using predefined heuristic rules to identify suspicious structural characteristics and calculate a heuristic risk score.

Hybrid Decision Making:

The system shall combine machine learning predictions with heuristic risk analysis to produce the final phishing detection result.

Explainable Prediction:

The system shall generate feature-level explanations using SHAP Explainable Artificial Intelligence (XAI) to identify the factors contributing to each prediction.

Result Presentation:

The system shall display the predicted URL classification, confidence score, heuristic analysis, and explanation through the web interface and REST API.

Maintain Complete Prediction Workflow

The prediction workflow defines the sequential data processing pipeline executed by ShieldPhish when evaluating an input URL. The system operates via a dual path hybrid architecture, processing statistical machine learning and deterministic heuristic rules concurrently before fusing the outputs and extracting TreeSHAP explainability attributions.

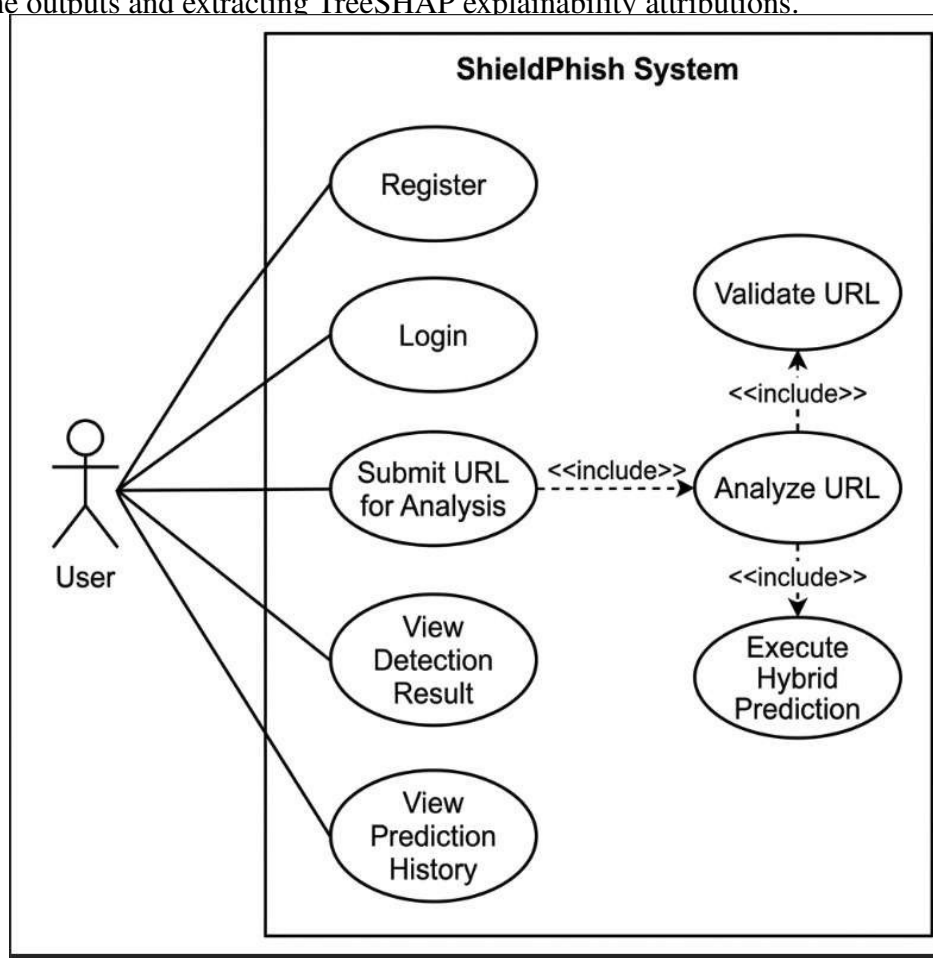


Figure 3.1: Use Case Diagram of shieldPhish

The use case diagram shows the functional boundaries of the ShieldPhish Phishing URL Detection System from the user's perspective. The primary actor in this system is the User (Security Analyst), who interacts with the system to perform authentication, submit URLs for threat analysis, view detection results, and retrieve historical query records. The core workflow begins with the analyst registering an account and logging into the system dashboard. To inspect a target URL, the analyst executes the Submit URL for Analysis use

case. This action incorporates («include») the Analyze URL sub-process, which further includes the Validate URL use case (to perform syntax, format, and length checks) and the Execute Hybrid Prediction use case (to run the feature extraction, heuristic evaluation, and bounded fusion model). Once the detection pipeline finishes, the analyst can access the View Detection Result use case to inspect the hybrid risk score and TreeSHAP attributions, or invoke the View Prediction History use case to audit past prediction logs.

Non Functional Requirement

Non-functional requirements define the quality attributes, performance constraints, operational benchmarks, and system qualities that ShieldPhish must satisfy.

Performance & Sub-Millisecond Inference Latency

The system shall process raw URLs, extract 112 features, calculate hybrid probability predictions, and generate TreeSHAP attributions with sub-millisecond execution latency (**< 1 ms per URL**). This ensures suitability for high-throughput network perimeter inspection and real-time proxy integration.

Predictive Accuracy & Generalization Robustness

The primary machine learning classifiers shall maintain benchmark classification performance under 5-fold stratified cross-validation on held-out evaluation datasets: Random Forest: $\geq 95.0\%$ Accuracy

XGBoost: $\geq 94.3\%$ Accuracy and ≥ 0.983 ROC-AUC score

The bounded hybrid fusion engine ($\alpha = 0.15$) shall mitigate false negatives on zero-day URLs containing distinct structural threat anomalies.

System Transparency & Model Interpretability

100% of automated classification predictions shall be accompanied by local feature attributions generated via TreeSHAP. The system shall convert mathematical Shapley values into the top 5 human-readable diagnostic reasons, eliminating black-box decision opacity for Security Operations Center (SOC) analysts.

Resource Efficiency & High-Throughput Scalability

The backend server (FastAPI) shall execute efficiently on standard single-core or multi-core CPU architecture without requiring expensive GPU hardware acceleration. The asynchronous REST API shall support concurrent HTTP requests without memory leaks or degradation of inference speed.

System Reliability & Fault Tolerance

The system shall gracefully handle malformed, non-standard, percent-encoded, or extremely long URL inputs. If optional SHAP explainer modules are unavailable for non-tree models, the system shall degrade gracefully by serving the hybrid prediction score without backend crashes.

Portability & Runtime Environment Compatibility

The codebase shall be fully portable, operating across Linux, macOS, and Windows operating systems. The application stack shall run within standard Python 3.9+ environments utilizing open-source dependencies (scikit-learn, xgboost, shap, FastAPI, React, Streamlit).

3.1.2 Feasibility Analysis

A feasibility study is a critical evaluation conducted during the system analysis phase to ascertain whether the proposed ShieldPhish framework can be successfully developed, evaluated, and deployed. The analysis assesses the project across four primary dimensions: Technical Feasibility, operational feasibility, economic feasibility, and Schedule Feasibility.

Technical Feasibility

Technical feasibility evaluates whether the required hardware, software infrastructure, and machine learning algorithms are available and capable of meeting system objectives:

- **Open-Source Software Infrastructure:** The proposed framework is implemented using mature, open-source technologies that provide reliable support for machine learning, explainable AI, API development, and web application deployment. These technologies have extensive community support, active maintenance, and proven scalability, making them suitable for developing the proposed phishing URL detection framework.
- **Hardware & Resource Efficiency:** Feature engineering, hybrid fusion, and TreeSHAP attribution calculations execute efficiently on standard single-core or multi-core CPU architectures without requiring expensive GPU hardware acceleration.
- **Computational Efficiency:** The proposed framework performs feature extraction and prediction efficiently on standard CPU hardware, enabling near real-time phishing URL detection suitable for operational network security applications.

Operational Feasibility

Operational feasibility assesses how effectively the proposed solution fits into real world operational workflows and whether end-users can utilize the system easily:

- **Model Interpretability & Decision Transparency:** Traditional machine learning models produce probability scores that are difficult for security analysts to interpret. ShieldPhish addresses this limitation by incorporating TreeSHAP local feature attributions, converting mathematical Shapley values into human-readable diagnostic reasons.
- **SOC Analyst Usability:** Security Operations Center (SOC) analysts can inspect threat reports, view top contributing feature attributions, and review triggered heuristic rules without requiring deep machine learning expertise.
- **Multi-Interface Deployment:** The system provides an interactive React web interface for real-time URL scanning, a Streamlit diagnostic dashboard for model analytics, and a FastAPI REST API for automated security proxy integration.

Economic Feasibility

Economic feasibility evaluates the financial viability, cost structure, and infrastructure requirements of the project:

- **Cost Effectiveness:** The proposed system relies entirely on open-source software technologies under permissive software licenses. Consequently, system development and deployment incur zero commercial licensing fees.
- **Low Infrastructure Footprint:** Because the framework is optimized for CPU execution rather than high-end GPU clusters, operational hosting and hardware maintenance costs remain minimal.
- **Public Threat Intelligence Datasets:** Model training and evaluation utilize publicly available cybersecurity datasets compiled from PhishTank and Kaggle, incurring zero data acquisition expenditure.

Schedule Feasibility

Schedule feasibility determines whether the project can be completed within the allocated timeframe:

CRISP-DM Iterative Execution: The project followed an iterative development process based on the CRISP-DM framework, with clearly defined milestones for data preparation, feature engineering, model development, evaluation, explainability integration, and application deployment. The planned activities were completed within the allocated project duration.

Table 3.1: Gantt Chart

Activities/Duration	Week 1	Week 2	Week 3	Week 4	Week 5	Week 6	Week 7	Week 8	Week 9	Week 10
Planning										
Research										
Data Gathering										
Requirement Analysis										
Coding										
Testing										
Validation										
Documentation										

3.2 Analysis

The proposed ShieldPhish system can be analyzed using an object-oriented approach because the system consists of different modular entities and classes that interact with each other. The major components of the system include the user/analyst, URL input, Feature Extractor module, Rule Engine module, Hybrid Phishing Model estimator, TreeSHAP explainer, and result output interface. Object-oriented analysis helps to represent the system using structural and behavioral UML diagrams such as the Class Diagram, Sequence Diagram, and Activity Diagram which illustrate the internal static structure, dynamic behavior, and operational workflow of the proposed framework.

Object modelling using Class and Object Diagrams

Object modeling represents the static structure of the proposed ShieldPhish Phishing URL Detection System. It identifies the major classes, objects, attributes, operations, and relationships involved in the system. In this project, the main classes include User, Admin, Feature Extractor, Rule Engine, Hybrid Phishing Model, Logistic Regression Scratch, FastAPI Controller, and Threat Report. The object model helps to describe how different components of the system are connected with each other and shows how a user submits a URL string, how the URL is validated and converted into a 112-dimensional feature vector, how the trained machine learning classifier and 6-rule heuristic engine perform dual-path hybrid prediction, and how the final threat assessment result and TreeSHAP XAI attributions are generated.

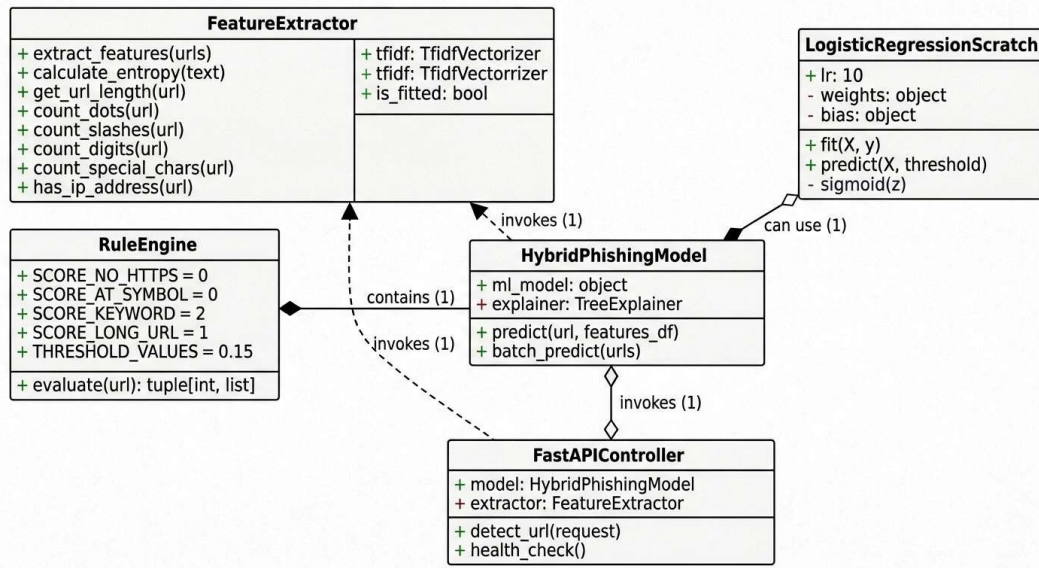


Figure 3.2: Class Diagram of shieldPhish

The class diagram illustrates the static architecture of the proposed ShieldPhish system. The FastAPI Controller serves as the primary API gateway, orchestrating user requests and routing data to the core analytical components: the Feature Extractor, Rule Engine, and Hybrid Phishing Model. The Feature Extractor parses the input URL string to generate a comprehensive 112-dimensional feature vector, while the Rule Engine concurrently evaluates predefined heuristic conditions to compute an integer threat penalty score (Srule). The Hybrid Phishing Model then integrates the statistical machine learning probability (PML) with the heuristic penalty using a bounded fusion function to output the final classification confidence (Phybrid). To ensure interpretability, the Tree Explainer generates SHAP-based local feature attributions. The complete detection payload including the prediction, confidence, and explanatory factors is then encapsulated within the Threat Report. Overall, the diagram accurately maps the static dependencies and structural relationships required for full-stack feature extraction, hybrid detection, explainability, and threat reporting.

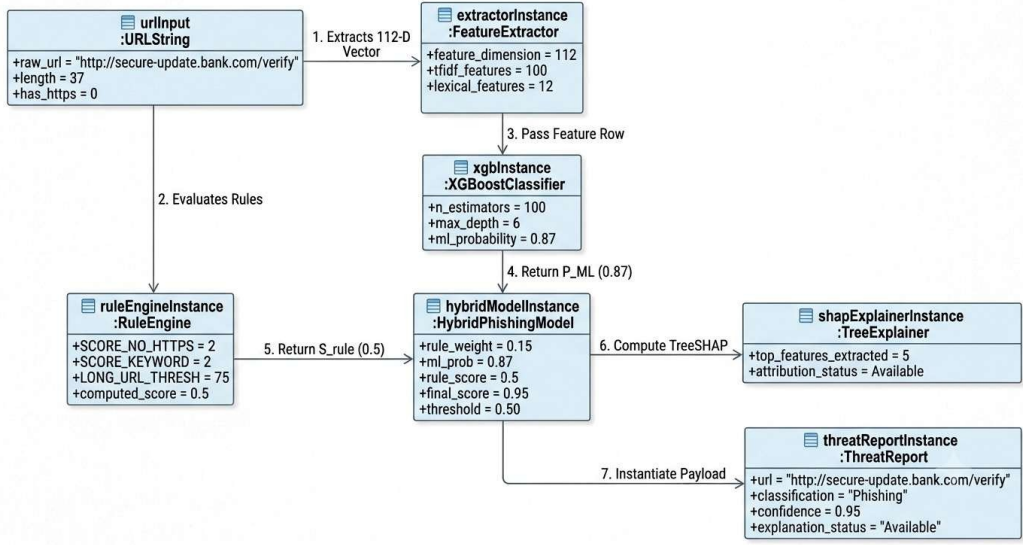


Figure 3.3: Object Diagram of shieldPhish

The object diagram illustrates a runtime snapshot of the ShieldPhish system during a single phishing URL detection process. The url Input: URL String object submits the input URL, which is processed simultaneously by the Feature Extractor to generate a 112-dimensional feature vector and by the Rule Engine to compute the heuristic penalty score ($S_{\text{rule}} = 0.5$). The extracted features are then classified by the XGBoost Classifier, producing a phishing probability ($P_{\text{ML}} = 0.87$). The Hybrid Phishing Model combines the machine learning probability and heuristic score to calculate the final hybrid confidence ($P_{\text{hybrid}} = 0.95$) and determine the final classification. The Tree Explainer generates SHAP-based feature attributions for interpretability, while the Threat Report object stores the final prediction, confidence score, and explanation. Overall, the diagram demonstrates the interaction among system objects during feature extraction, hybrid classification, explainability, and report generation.

Dynamic Modelling Using State and Sequence Diagrams

Dynamic modeling shows the behavior of the ShieldPhish Phishing URL Detection System during execution. It explains how the system changes its state and how different objects interact with each other during the phishing URL detection and explainability process. In this project, dynamic modeling is represented using the State Diagram and Sequence Diagram.

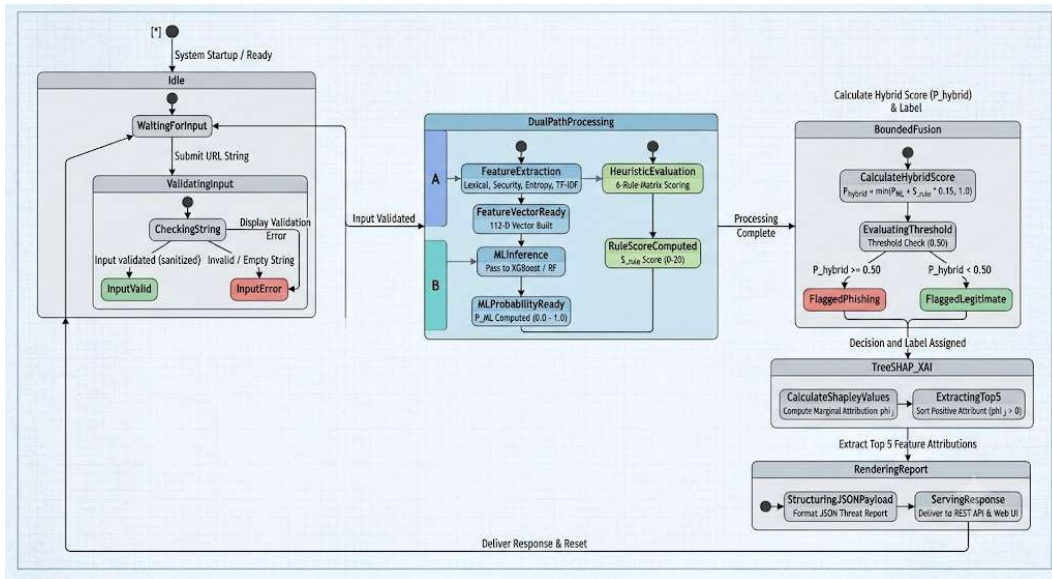


Figure 3.4: State Diagram of shieldPhish

When a user submits a URL string, the input first enters the Submitted state. The system then validates the URL string format and length. If the URL is valid, it moves to the Validated state; otherwise, it moves to the Rejected state. After validation, the URL is processed by extracting the 112-dimensional feature vector and concurrently evaluating the 6-rule heuristic matrix. Then, the trained machine learning model and heuristic rule engine perform dual-path hybrid prediction, generating the final classification result as Legitimate or Phishing. Finally, TreeSHAP feature attributions are computed, the prediction record is saved, and a diagnostic threat report is generated.

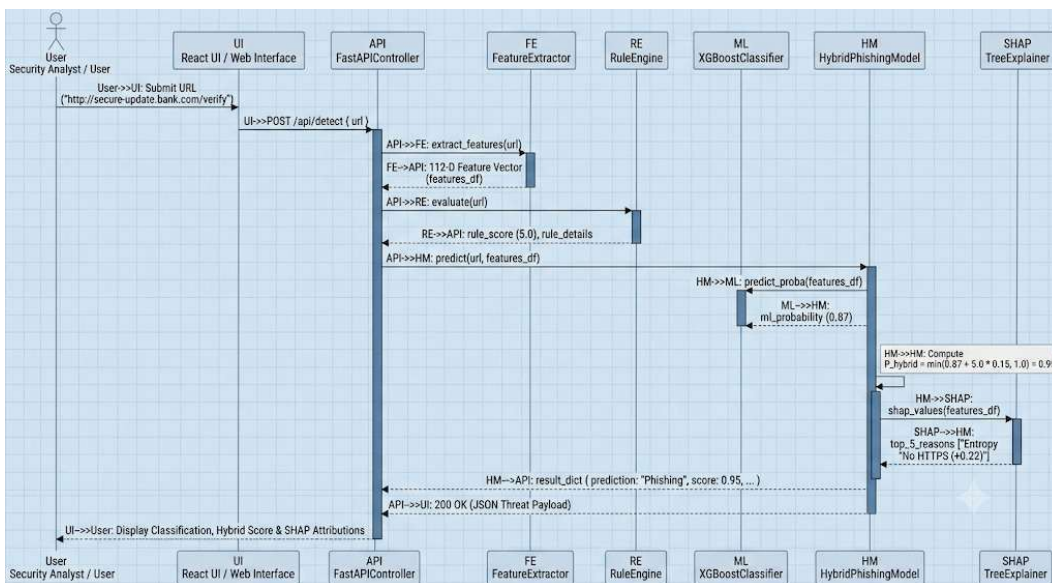


Figure 3.5: Sequence Diagram of shieldPhish

The sequence diagram illustrates the interaction among the components of the ShieldPhish system during a phishing URL detection request. The process begins when the user submits a

URL through the React web interface, which forwards the request to the FastAPI Controller. The controller invokes the Feature Extractor to generate a 112-dimensional feature vector and the Rule Engine to compute the heuristic penalty score. The extracted features are then passed to the XGBoost classifier through the Hybrid Phishing Model to estimate the machine learning phishing probability. The Hybrid Phishing Model combines the ML probability with the heuristic score to calculate the final hybrid confidence score. Subsequently, the TreeSHAP Explainer generates feature-level explanations, and the Hybrid Phishing Model prepares the threat report. Finally, the API returns the classification result, confidence score, and SHAP explanations to the user interface for display.

Process Modelling Using Activity Diagram

Process modeling represents the workflow of the ShieldPhish Phishing URL Detection System. It shows the sequence of activities performed from the time a user submits a URL string until the final threat prediction, TreeSHAP attributions, and diagnostic report are displayed.

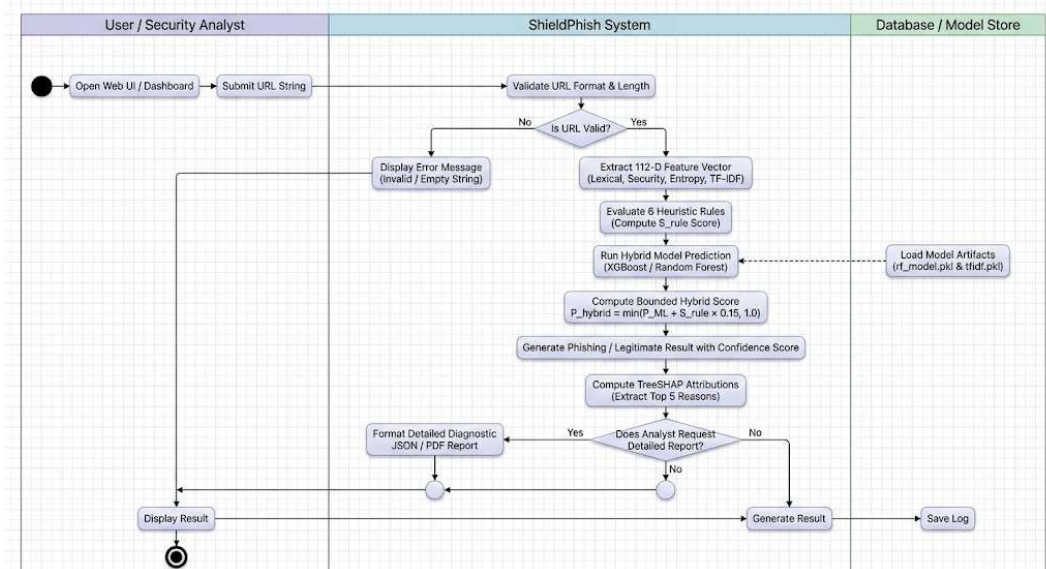


Figure 3.6: Activity Diagram of shieldPhish

The activity diagram illustrates the operational workflow of the ShieldPhish system by separating responsibilities among the User/Security Analyst, the ShieldPhish System, and the Database/Model Store. The process begins when the user submits a URL through the web interface. The system validates the input and, if valid, extracts a 112-dimensional feature vector, evaluates heuristic rules, performs hybrid machine learning prediction, and computes the final phishing confidence score. TreeSHAP then generates feature-level explanations, after which the system prepares either a standard or detailed threat report based on the user's request. Finally, the detection result is displayed to the user, and the prediction log is stored in the database.

CHAPTER 4

SYSTEM DESIGN

4.1 Design

System Design System design covers how the proposed ShieldPhish Phishing URL Detection System is organized and how the various components of the system interact. The design of this project is based on the refinement of system modules, user interface, database/model store, component structure, deployment structure, and algorithm details as it is based on an object-oriented approach. The suggested system is a web-based application, where user can enter URL strings and the system automatically predicts if the URL is Legitimate or Phishing utilizing a hybrid fusion of machine learning and heuristic rule analysis. The system contains a front end interface (React), backend server (FastAPI), feature engineering module, trained machine learning model, 6-rule heuristic engine, TreeSHAP explainability module, database/model store, and Streamlit diagnostic dashboard. In this chapter the refining of the Class Diagram, Object Diagram, State Diagram, Sequence Diagram, and Activity Diagram is described, as well as the algorithm level procedural processes and deployment architecture.

Refinement of Class, Object, State, Sequence and Activity diagrams

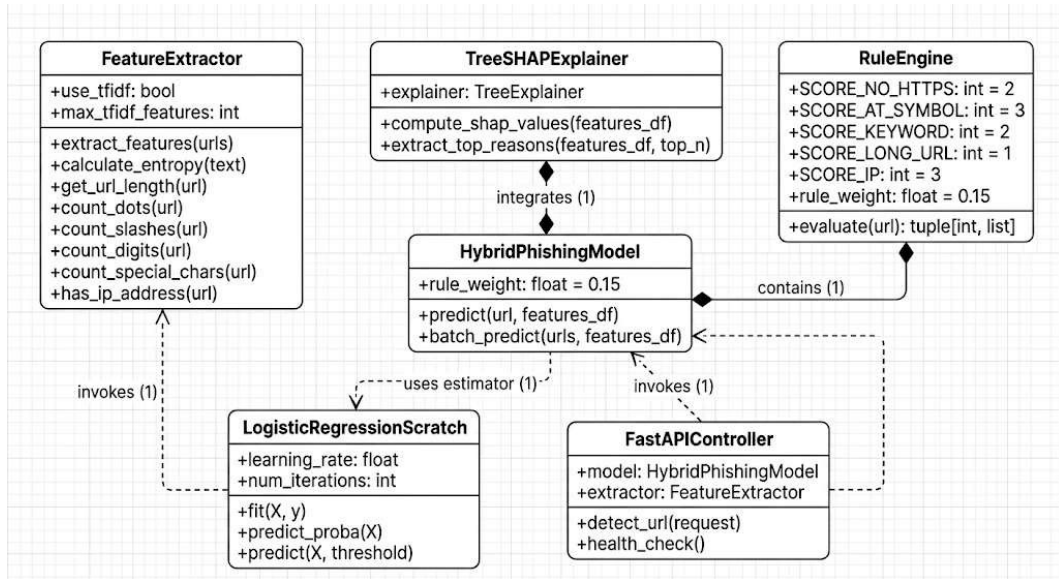


Figure 4.1: Refinement of Class Diagram of shieldPhish

The detailed class diagram illustrates the primary classes of the ShieldPhish Phishing URL Detection System and their relationships. FastAPI The Controller class functions as the main API gateway to handle incoming detection requests and system health checks. The security analyst provides URL strings and accesses diagnostic threat reports. The administrator handles trained model artifacts, TF-IDF vectorizers, and threshold parameters. The Feature Extractor class takes the provided URL and generates a 112-dimensional feature vector. The feature vector and raw URL are passed to Hybrid Phishing Model, which compositionally

incorporates Rule Engine to evaluate 6 structural risk rules and computes statistical probability using the primary classifier (such as XGBoost, Random Forest or Logistic Regression Scratch). The combined prediction is then provided to the TreeSHAP Explainer to extract local feature-level attributions, which are then saved in the threat model store and prepared by the Threat Report to build a diagnostic threat payload.

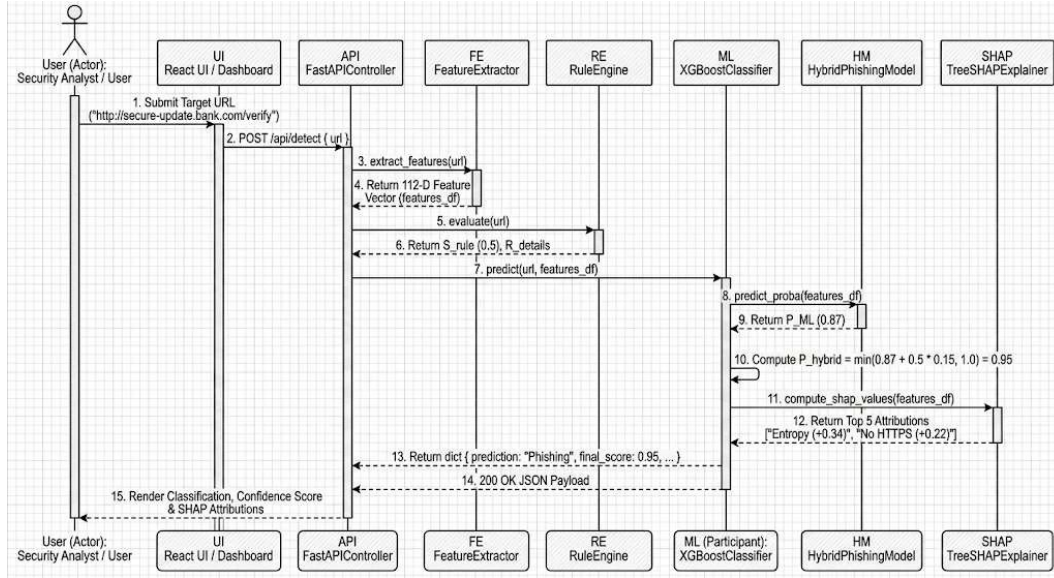


Figure 4.2: Refinement of Sequence Diagram of shieldPhish

The sequence diagram illustrates the interaction between the user and the ShieldPhish system during a phishing URL detection request. Initially, the user submits a URL through the React web interface, which sends the request to the FastAPI controller. The controller invokes the FeatureExtractor to generate a 112-dimensional feature vector and simultaneously requests the RuleEngine to evaluate heuristic security rules and compute the rule score. The extracted features are then forwarded to the XGBoost classifier, which predicts the machine learning phishing probability (P_{ML}). The HybridPhishingModel combines the machine learning probability with the heuristic rule score using the bounded hybrid fusion equation to calculate the final phishing score (P_{hybrid}). Next, the TreeSHAP Explainer computes feature attributions and identifies the top contributing reasons for the prediction. Finally, the HybridPhishing Model generates a threat report containing the prediction result, confidence score, and SHAP explanations. The FastAPI controller returns the response as a JSON payload, and the React web interface displays the phishing classification and explanation to the user.

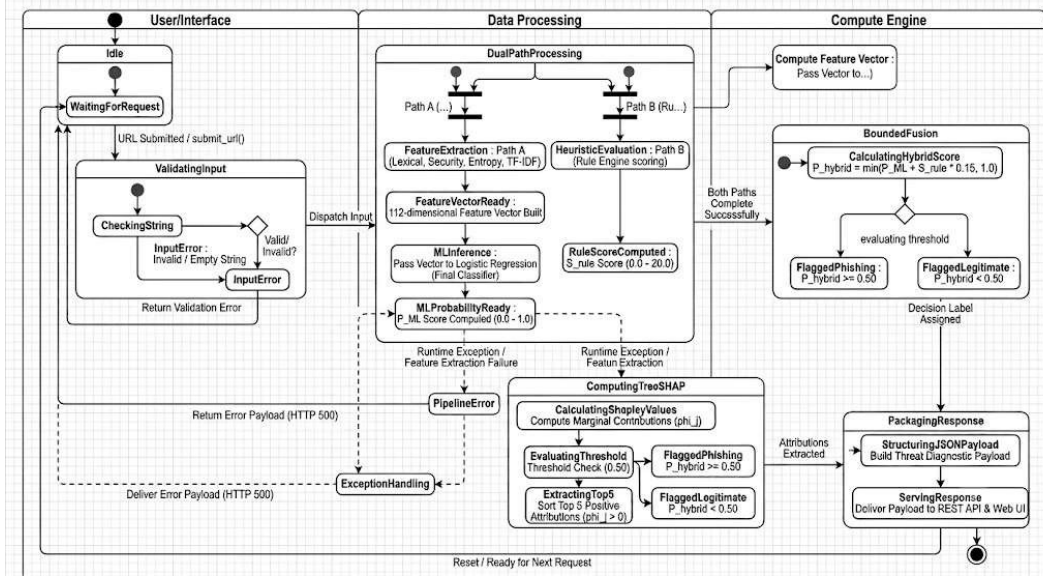


Figure 4.3: Refinement of Activity Diagram of shieldPhish

The Controller class functions as the main API gateway for handling incoming detection requests and system health checks. The Security Analyst submits URL strings and accesses diagnostic threat reports, while the Administrator manages trained model artifacts, TF-IDF vectorizers, and threshold parameters. The Feature Extractor class processes the submitted URL and generates a 112-dimensional feature vector. The generated feature vector, together with the raw URL, is passed to the Hybrid Phishing Model. The model incorporates the Rule Engine, which evaluates six structural risk rules, and computes a statistical probability using the primary classifier, such as XGBoost, Random Forest, or a Logistic Regression model implemented from scratch. The combined prediction is then passed to the TreeSHAP Explainer, which calculates local feature-level attributions to identify the most influential features contributing to the prediction. These prediction results and explanations are stored in the threat model store and are subsequently processed by the Threat Report component to generate a diagnostic threat-report payload. The complete activity diagram illustrates the overall workflow of the ShieldPhish framework using three distinct operational swimlanes: User / Security Analyst, ShieldPhish System, and Database / Model Store. The process begins when the user accesses the dashboard interface and submits a target URL. The system first validates the format and length of the URL. If the URL is invalid, execution is terminated and an appropriate error message is displayed. If the URL is valid, the system preprocesses the URL string and executes the main analytical pipeline. During this stage, the system extracts a 112-dimensional feature vector consisting of lexical, security-related, Shannon entropy, and character-level TF-IDF features. The URL is also evaluated against a six-rule heuristic matrix. At the same time, the required pre-trained model artifacts, such as `rf_model.pkl` and `tfidf.pkl`, are loaded from the persistent model store. The statistical machine-learning probability and the heuristic penalty score are then combined using a bounded fusion function to produce the final classification, either Legitimate or Phishing, together with the hybrid confidence score, P_{hybrid} . After classification, TreeSHAP feature attributions are calculated

to determine the most influential structural features contributing to the prediction. The final prediction record, including the classification result and associated diagnostic information, is securely stored in the database for future telemetry and analysis. Finally, the workflow branches according to the user's selected output. If a detailed report is requested, the system generates a diagnostic report in JSON or PDF format. Otherwise, the standard detection result is displayed to the user, after which the process terminates.

Component Diagram

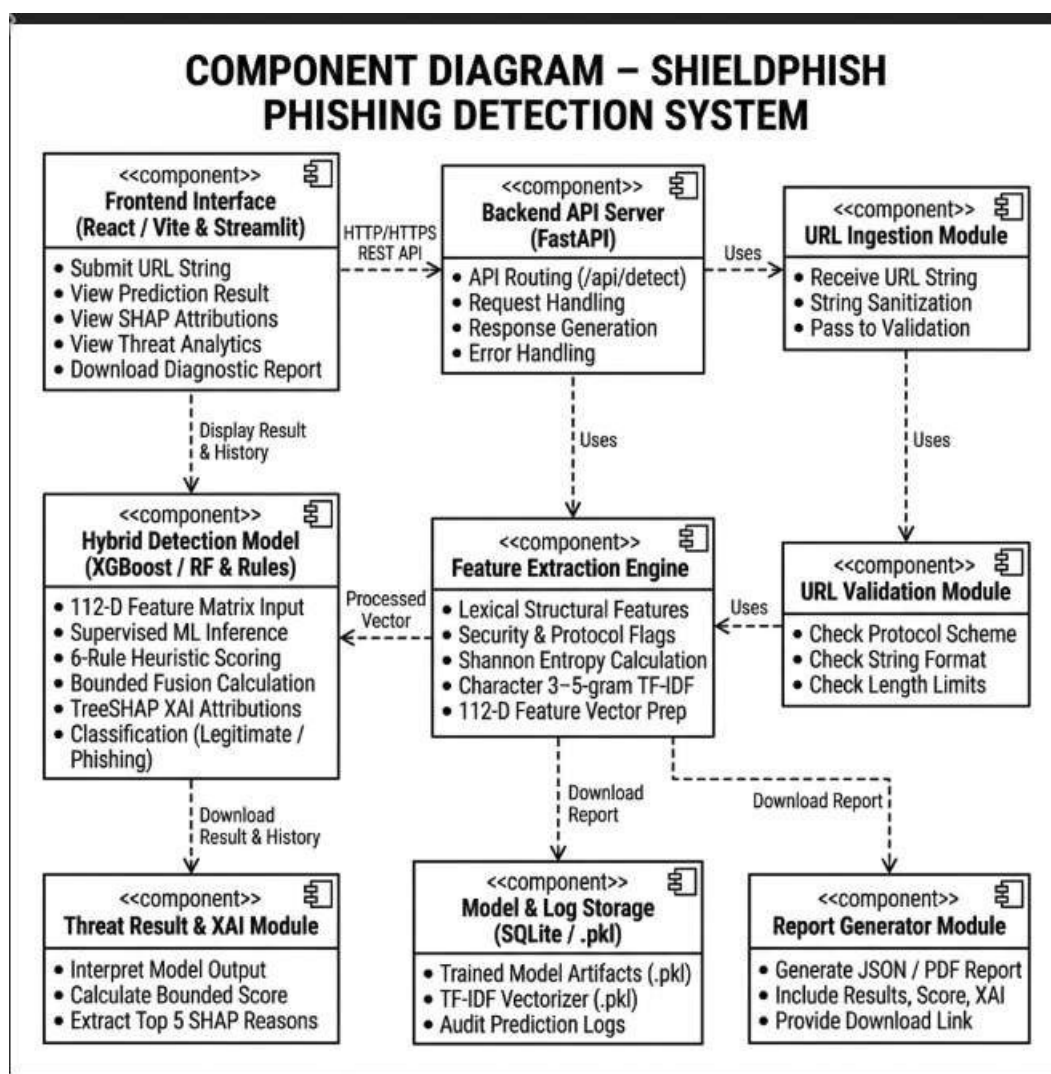


Figure 4.4: Component Diagram of shieldPhish

The component diagram depicts the major modules of ShieldPhish Phishing URL Detection System and its interaction. The Frontend Interface (React/Vite & Streamlit) interfaces with the Backend API Server (FastAPI) using HTTP/HTTPS REST APIs. The backend manages API routing (/api/detect), URL ingestion, string validation, feature extraction, hybrid prediction execution, model store fetching, and diagnostic report creation. The raw URL is delivered to the Feature Extraction Engine, which computes the 112-dimensional feature vector, after input validation. The feature vector is then fed into the Hybrid Detection Model (XGBoost /

RF & Rules) for Legitimate/Phishing classification and 6-rule penalty assessment. The Threat Result & XAI Module produces the final classification label, confidence score (Phybrid), top 5 TreeSHAP feature attributions, saves the prediction log in Model & Log Storage (SQLite/.pkl), and the Report Generator Module prepares a downloadable diagnostic JSON/PDF report.

Deployment Diagram

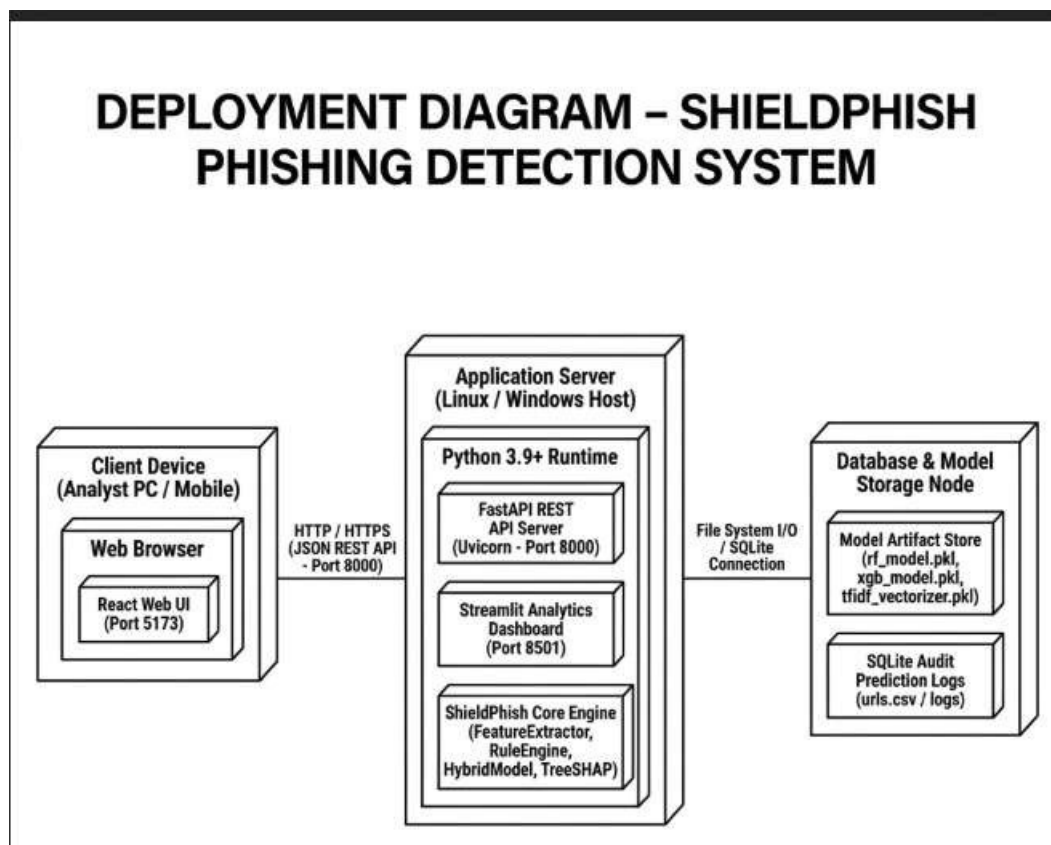


Figure 4.5: Deployment Diagram of shieldPhish

The deployment diagram demonstrates the deployment of the ShieldPhish Phishing URL Detection System in a client–server context. The security analyst or administrator logs into the system through a web browser running the React web interface on a client device. The browser uses HTTP/HTTPS REST APIs to interface with the FastAPI application server to send URLs and get phishing detection results. The application server hosts the ShieldPhish core engine (containing the Feature Extractor, Rule Engine, Hybrid Phishing Model and TreeSHAP Explainer) and the Streamlit analytics dashboard. The server conducts URL validation, feature extraction, hybrid prediction, explainability generation, and report creation. During inference, trained model artifacts like Random Forest, XGBoost, and the TF-IDF vectorizer are loaded from local storage. Local SQLite database for prediction audit logs, confidence ratings, heuristic rule results and generated reports. To summarize, the deployment diagram shows the physical placement of the client, application server, and storage components and their interactions in detecting phishing URLs.

4.2 Algorithms Details

The main framework utilized in this project is the ShieldPhish Hybrid Feature Extraction, Bounded Rule Fusion and TreeSHAP Explainability Network. The system is designed to detect phishing URLs by means of structural lexical metrics, security protocol flags, information-theoretic Shannon entropy, and character-level TF-IDF n-grams calculated from raw URL strings. Unlike typical stand-alone machine learning algorithms generally used as black-box predictors, the proposed technique processes each input through a dual-path pipeline: a 112-dimensional Statistical ML Classifier Stream and a Deterministic 6-Rule Heuristic Engine Stream.

The system accepts a raw URL string as input, represented as a character sequence:

$$U = \{c_1, c_2, \dots, c_N\}, \quad c_i \in \Sigma \quad (4.1)$$

where N represents the total character length of the URL, and Σ represents the character alphabet (alphanumeric characters, punctuation, and special symbols).

Structural, Lexical, & Security Feature Extractor

The Structural Lexical and Security Stream processes the URL string directly to extract structural metrics and protocol indicators. These features capture structural anomalies, brand impersonation attempts, and unencrypted transmission schemes.

The structural metrics extract total length, path length, domain length, character counts, and security protocol flags:

$$x_{\text{len}} = |U|, \quad x_{\text{dots}} = \sum_{i=1}^N \mathbb{I}(c_i = '.'), \quad x_{\text{slashes}} = \sum_{i=1}^N \mathbb{I}(c_i = '/') \quad (4.2)$$

The digit-to-length ratio R_{digit} and special-character density R_{special} are calculated as:

$$R_{\text{digit}} = \frac{\sum_{i=1}^N \mathbb{I}(c_i \in \{0, \dots, 9\})}{N} \quad (4.3)$$

$$R_{\text{special}} = \frac{\sum_{i=1}^N \mathbb{I}(c_i \in \{\@, ?, =, \%, _ \})}{N} \quad (4.4)$$

Security flags evaluate protocol and hostname properties:

$$x_{\text{https}} = \mathbb{I}(U \text{ starts with "https://"}) \quad (4.5)$$

This stream yields a compact 11-dimensional structural vector:

$$\mathbf{x}_{\text{struct}} \in \mathbb{R}^{11} \quad (4.6)$$

Information-Theoretic Shannon Entropy ($\mathcal{H}$) The Shannon Entropy Stream captures character randomness and uncertainty within the URL string. Phishing campaigns generated via Domain Generation Algorithms (DGAs) or hex-encoded session parameters exhibit higher

character randomness than legitimate natural-language URLs.

For a URL string U containing unique characters c_i with empirical probability $P(c_i) = \frac{\text{count}(c_i)}{|U|}$,

Shannon Entropy is defined as:

$$\mathcal{H}(U) = - \sum_{i=1}^{|\Sigma|} P(c_i) \log_2 P(c_i) \quad (4.7)$$

where $\mathcal{H}(U) \in \mathbb{R}$ represents the entropy scalar. Higher values indicate algorithmically generated domain names or obfuscated parameters.

Character-Level 3–5-gram TF-IDF Vectorizer

The TF-IDF stream captures fine-grained subword linguistic representations (such as log, ogi, gin for the target term login).

The Term Frequency (TF) and Inverse Document Frequency (IDF) for character n-gram t in URL U across corpus D are defined as:

$$\text{TF}(t, U) = \frac{f_{t,U}}{\sum_{t' \in U} f_{t',U}} \quad (4.8)$$

$$\text{IDF}(t, D) = \log \left(\frac{1 + |D|}{1 + |\{U \in D : t \in U\}|} \right) + 1 \quad (4.9)$$

The final TF-IDF feature value is computed as:

$$\mathbf{x}_{\text{tfidf}}(t) = \text{TF}(t, U) \times \text{IDF}(t, D) \in \mathbb{R}^{100} \quad (4.10)$$

Concatenating all streams produces the full 112-dimensional feature vector:

$$\mathbf{X} = [\mathbf{x}_{\text{struct}} \parallel \mathcal{H}(U) \parallel \mathbf{x}_{\text{tfidf}}] \in \mathbb{R}^{112} \quad (4.11)$$

Note: The final feature vector consists of 11 handcrafted structural lexical and security features, 1 Shannon entropy feature, and 100 character-level TF-IDF n-gram features, producing a total dimension of 112.

Supervised ML Probability & 6-Rule Heuristic Engine

The feature vector $\mathbf{X}$ is passed to the trained ensemble classifier (XGBoost / Random Forest) to compute the statistical phishing probability:

$$P_{\text{ML}} = \text{predict_proba}(\mathbf{X})[1] \in [0.0, 1.0] \quad (4.12)$$

Concurrently, the raw URL string U is evaluated against the 6-rule heuristic engine to calculate an integer penalty score:

$$S_{\text{rule}} = \sum_{k=1}^6 s_k \in [0, 20] \quad (4.13)$$

where the penalty weights are assigned as: $s_1 = 2$ (no HTTPS), $s_2 = 3$ (@ symbol), $s_3 = 2N_{\text{kw}}$ (suspicious keyword matches), $s_4 = 1$ (length > 75), $s_5 = 3$ (IPv4 host), and $s_6 = 1$ (more than three subdomains).

Bounded Hybrid Fusion & TreeSHAP Explainability

To combine statistical pattern recognition with deterministic safety rules, the system applies the Bounded Hybrid Fusion Formula:

$$P_{\text{hybrid}} = \min(P_{\text{ML}} + S_{\text{rule}}\alpha, 1.0), \quad \alpha = 0.15 \quad (4.14)$$

The final binary classification label $\hat{y}$ is assigned via:

$$\hat{y} = \begin{cases} \text{Phishing}, & P_{\text{hybrid}} \geq 0.50, \\ \text{Legitimate}, & P_{\text{hybrid}} < 0.50. \end{cases} \quad (4.15)$$

Finally, local post-hoc feature attributions are computed via TreeSHAP:

$$f(\mathbf{X}) = \phi_0 + \sum_{j=1}^{112} \phi_j z'_j \quad (4.16)$$

where ϕ_0 is the expected baseline model output and ϕ_j represents the marginal contribution of the j -th feature toward the final prediction. Positive attributions ($\phi_j > 0$) are extracted and presented as the top 5 human-readable diagnostic reasons for security analysts.

CHAPTER 5

IMPLEMENTATION AND TESTING

5.1 Implementation

Implementation is the transformation of the system design into a working software application. This project implements the ShieldPhish Phishing URL Detection System as a web application, where users can enter URL strings and get a prediction result as Legitimate or Phishing, together with local TreeSHAP feature attributions and threat metrics. The system is developed using a hybrid fusion model consisting of a trained machine learning classifier (Random Forest or XGBoost) saved in binary pickle (.pkl) format, a character level TF-IDF vectorizer and a 6-rule heuristic engine developed in-house. The backend is responsible for URL validation, 112-dimensional feature extraction, heuristic scoring, model loading, bounded probability fusion, TreeSHAP explainability computation, database operations, and diagnostic report production. The frontend allows users to easily submit a URL string, visualize threat scores, present TreeSHAP local attribution charts and obtain a diagnostic report.

The overall implementation workflow is as follows:

5.1.1 Tools Used

The following tools and technologies were used to implement the proposed ShieldPhish Phishing URL Detection System. These tools support frontend development, backend API development, machine learning model integration, feature engineering, heuristic rule evaluation, explainable AI attribution, database logging, and threat report generation. The development of ShieldPhish utilizes a collection of open-source frameworks, machine learning libraries, and storage technologies to support phishing URL detection, explainability, and web-based deployment.

- **Programming Language (Python 3.9+):** Python serves as the primary programming language for feature extraction, heuristic rule evaluation, machine learning inference, and backend development due to its extensive scientific computing ecosystem.
 - **Backend Framework (FastAPI):** FastAPI provides a high-performance RESTful API for URL detection, enabling efficient request handling, JSON processing, and automatic API documentation through Open API/Swagger.

Machine Learning Libraries:

- **Scikit-learn:** Used for data preprocessing, model evaluation, dataset splitting, and TF-IDF vectorization.
- **XGBoost:** Implements the primary gradient boosting classifier for phishing URL prediction.

- **NumPy & Pandas:** Support numerical computation, data manipulation, and construction of the 112-dimensional feature vectors.
- **Explainable AI (SHAP):** The SHAP library is used to generate TreeSHAP feature attributions, providing interpretable explanations of model predictions by identifying the most influential URL features.

Frontend & Visualization:

- **React.js with Tailwind CSS:** Provides an interactive web interface for URL submission and phishing prediction results.
- **Streamlit:** Offers a diagnostic dashboard for model evaluation, feature importance visualization, and performance analysis.
- **Database (SQLite):** Stores prediction history, heuristic rule scores, model probabilities, and final classification results for auditing and analysis.
- **Development Tools (Git & Virtual Environment):** Git is used for version control, while Python virtual environments ensure dependency isolation and reproducible development.

5.1.2 Dataset Description

To evaluate the performance and robustness of the proposed ShieldPhish framework, a large-scale phishing URL dataset was constructed by combining two widely used cybersecurity data sources: PhishTank and the Kaggle Phishing URL Dataset.

Data Collection

The complete dataset consists of 957,334 URLs collected from the following sources:

- **Phishing URLs:** Obtained from PhishTank, which provides verified and continuously updated phishing URLs.
- **Legitimate URLs:** Obtained from the Kaggle Phishing URL Dataset, containing verified benign website URLs.

Full Dataset Evaluation and Cross-Validation

Unlike traditional approaches that rely on a single training and testing split, this project rigorously evaluates models on the full dataset of 957,334 URLs using 5-Fold Stratified Cross Validation. This approach ensures that the evaluation is highly robust, proves the model's generalization capabilities, and eliminates potential information leakage from a single held-out test split. To mirror real-world operational environments, the dataset preserves its natural class distribution rather than artificially balancing the data:

- **Phishing URLs:** 218,699 (Label = 1, ~22.8%)
- **Legitimate URLs:** 738,635 (Label = 0, ~77.2%)

To account for this class imbalance during training, tree-based classifiers actively compute and apply dynamic class scaling weights. Table 5.1 outlines the specific dataset and feature characteristics.

Table 5.1: Dataset Characteristics

Parameter	Configuration	Value
Total Dataset	PhishTank + Kaggle	957,334 URLs
Evaluation Methodology	5-Fold Stratified Cross-Validation	5 Folds
Phishing URLs	Label = 1	218,699
Legitimate URLs	Label = 0	738,635
Feature Dimension	Structural + Entropy + TF-IDF	112 Features

Data Preprocessing

Before feature extraction, the collected URLs underwent several automated preprocessing steps to improve data quality and baseline model performance:

Duplicate Removal: Duplicate URLs were removed to prevent data leakage between the training and testing datasets.

- **URL Normalization:** Missing URL schemes (HTTP/HTTPS) were standardized to ensure consistent structural feature extraction.
- **Encoding Standardization:** All URLs were converted to UTF-8 encoding to remove invalid or unsupported characters, ensuring highly reliable TF-IDF feature generation.

5.1.3 Implementation Details of Modules

This section describes the technical implementation details of ShieldPhish’s core Python modules, outlining the class designs, operational workflows, and functional interfaces.

Feature Engineering Module

Purpose: The Feature Extractor class converts raw URL strings into structured numerical feature vectors suitable for statistical classification. This module is implemented in the `src/feature_engineering.py` file.

Implementation Details: Feature extraction is split into two steps:

- **Structural and Security Features:** Handcrafted feature extraction functions parse the URL string to compute lexical lengths, character counts, special character densities, and security protocol flags (e.g., protocol scheme checks and IP-in-host regular expression checks). The `calculate_entropy` method computes the Shannon entropy $\mathcal{H}$ of the input URL by calculating the probability distribution of unique characters and applying Shannon’s entropy equation.
- **NLP Text Vectorization:** The module wraps a character 3–5-gram `TfidfVectorizer` to capture sub-word token frequencies. **Functional Output:** The module concatenates

11 handcrafted structural and security features, 1 Shannon entropy feature, and 100 character-level TF-IDF features, producing the final 112-dimensional feature vector for each ingested URL.

Rule Engine

Purpose: The heuristic module implements deterministic safety checks to identify known phishing structures, acting as a safeguard for statistical classifiers. This module is implemented in the `src/rule_engine.py` file.

Implementation Details: The RuleEngine class implements the heuristic evaluation process using 6 structural threat rules:

- **Missing HTTPS:** Verifies if the protocol scheme lacks SSL/TLS encryption (+2 penalty).
- **Contains @ Symbol:** Detects redirection attempts via user-info parameters (+3 penalty).
- **Suspicious Keywords:** Scans for target brand-impersonation terms like *login*, *bank*, and *verify* (+2 penalty per match).
- **Excessive Length:** Identifies long obfuscating paths (+1 penalty if length > 75 characters).
- **Raw IP Hostname:** Flags bypasses of standard DNS registrations (+3 penalty).
- **Deep Nesting:** Checks for subdomains exceeding a depth of three (+1 penalty).

Functional Output: The engine aggregates the individual rule triggers to compute a cumulative penalty score, *Srule*, ranging from 0 to 20.

Hybrid Prediction Module

Purpose: The hybrid prediction module acts as the core fusion estimator that merges statistical and heuristic predictions. This module is implemented in the `src/hybrid_model.py` file.

Implementation Details: The Hybrid Phishing Model class executes a five-step prediction workflow:

- **ML Prediction:** Invokes the primary tree ensemble model (XGBoost or Random Forest) to compute the statistical phishing probability (*PML*).
- **Rule Evaluation:** Invokes the Rule Engine to evaluate the 6 threat rules and retrieve the heuristic penalty score (*Srule*).
- **Hybrid Fusion:** Applies the bounded hybrid fusion formula $P_{\text{hybrid}} = \min(P_{\text{ML}} + \alpha S_{\text{rule}}, 1.0)$ with an empirically tuned rule weight of $\alpha = 0.15$.
- **Final Classification:** Compares the hybrid probability score against a binary decision threshold of 0.50, assigning the final label as Phishing if $P_{\text{hybrid}} \geq 0.50$ and Legitimate if $P_{\text{hybrid}} < 0.50$.

Return Output: Packages the final threat labels, risk percentages, and triggered rule meta-data.

Functional Output: Returns a unified result dictionary containing classification predictions, confidence scores, and triggered rule details.

Explainability Module

Purpose: The explainability module provides local feature attributions for model predictions, converting the classifier's outputs into human-readable diagnostics. This module is integrated within `src/hybrid_model.py` and uses the open-source shap library.

Implementation Details: Upon initialization, the class instantiates a `shap.TreeExplainer` wrapped around the pre-trained tree ensemble classifier. For each prediction, the module computes the local Shapley values (ϕ_j) across the 112 features. The algorithm measures the marginal contribution of each feature towards shifting the final probability from the baseline expected value.

Functional Output: Employs a feature attribution ranking process that maps numeric feature indices to descriptive feature names, ranks them by their SHAP values, and returns the five most influential positive features contributing to the prediction.

Backend API

Purpose: The backend module implements the REST interface that exposes the detection pipeline to external clients. This module is implemented in the `backend/api.py` file. **Implementation Details:** Built using the high-performance FastAPI web framework and served via a Uvicorn ASGI server. The controller utilizes a Pydantic model `DetectionRequest` to validate incoming JSON request payloads, ensuring the input string is structured and non-empty. The framework routes requests to the detection endpoint (POST `/api/detect`).

Functional Output: Delivers a structured JSON response containing the prediction label, hybrid score (Phybrid), ML probability (PML), heuristic rules triggered (Srule), and the top 5 TreeSHAP feature attributions (ϕ_j).

5.2 Testing

Testing is performed to verify whether the proposed ShieldPhish Phishing URL Detection System works correctly according to its functional and non-functional requirements. It ensures that individual modules perform their assigned tasks correctly and that the complete system workflow operates properly from URL submission to final hybrid prediction, TreeSHAP explanation, database storage, and diagnostic report generation.

In this project, testing is divided into two main parts: Unit Testing and System Testing.

5.2.1 Test Cases for Unit Testing

Unit testing focuses on testing each module separately in isolation. Each test case is given a relevant title to clearly describe what is being tested. These tests help identify module-level errors before integrating the complete system.

Table 5.2: Test Cases for Unit Testing

Test ID	Test Case Title	Module	Test Scenario	Expected Result	Actual Result	Status
UT-01	Valid User Registration Test	User Authentication	User enters a valid username, email, and password.	New user account is created successfully.	New user account is created successfully.	Pass
UT-02	Empty Registration Field Test	User Authentication	User submits the registration form with empty fields.	System displays a required-field error.	System displays a required-field error.	Pass
UT-03	Valid Login Test	User Authentication	User enters the correct email and password.	User is logged into the system dashboard.	User is logged into the system dashboard.	Pass
UT-04	Invalid Login Test	User Authentication	User enters an incorrect password.	System displays an invalid login message.	System displays an invalid login message.	Pass
UT-05	Valid URL Ingestion Test	URL Ingestion	User submits a non-empty, valid URL string.	URL is parsed and accepted successfully.	URL is parsed and accepted successfully.	Pass
UT-06	Missing URL Submission Test	URL Ingestion	User clicks submit without entering a URL.	System displays a URL-required message.	System displays a URL-required message.	Pass
UT-07	Invalid Protocol Scheme Test	URL Validation	User submits a URL with an unsupported scheme (e.g., <code>ftp://</code>).	System flags a protocol warning.	System flags a protocol warning.	Pass
UT-08	String Length Limit Test	URL Validation	User submits a URL exceeding 2000 characters.	System displays a URL length-limit error.	System displays a URL length-limit error.	Pass

Continued on next page

Table 5.2 – continued from previous page

Test ID	Test Case Title	Module	Test Scenario	Expected Result	Actual Result	Status
UT-09	Lexical Attribute Extraction Test	Feature Extraction	Extractor processes the input URL string.	11 structural lexical features are counted.	11 structural lexical features are counted.	Pass
UT-10	Shannon Entropy Calculation Test	Feature Extraction	Extractor computes the character-frequency list.	Scalar entropy value $\mathcal{H}(U)$ is calculated.	Scalar entropy value $\mathcal{H}(U)$ is calculated.	Pass
UT-11	TF-IDF Tokenization Test	Feature Extraction	Vectorizer transforms the URL into character n-grams.	Extracted TF-IDF vector contains exactly 100 features.	Extracted TF-IDF vector contains exactly 100 features.	Pass
UT-12	Feature Vector Concatenation Test	Feature Extraction	Lexical, entropy, and TF-IDF vectors are merged.	Unified feature vector $\mathbf{X}$ has shape (1, 112).	Unified feature vector $\mathbf{X}$ has shape (1, 112).	Pass
UT-13	Model Loading Test	Model Prediction	Backend loads trained <code>rf_model.pkl</code> weights.	Classifier weights load successfully into memory.	Classifier weights load successfully into memory.	Pass
UT-14	Model Inference Test	Model Prediction	112-dimensional feature vector is passed to the ML model.	$P_{\text{ML}} \in [0, 1]$.	$P_{\text{ML}} \in [0, 1]$.	Pass
UT-15	Heuristic Rules Evaluation Test	Heuristic Engine	URL string is passed to the rule engine.	Score $S_{\text{rule}} \in [0, 20]$ is returned.	Score $S_{\text{rule}} \in [0, 20]$ and rule details are returned.	Pass
UT-16	Bounded Fusion Calculation Test	Hybrid Prediction	Hybrid score formula is evaluated.	A bounded hybrid score is returned.	Hybrid score is calculated and bounded to 1.0.	Pass

Continued on next page

Table 5.2 – continued from previous page

Test ID	Test Case Title	Module	Test Scenario	Expected Result	Actual Result	Status
UT-17	Phishing Label Threshold Test	Result Generation	Hybrid prediction probability is ≥ 0.50 .	URL is classified as Phishing.	URL is classified as Phishing.	Pass
UT-18	Legitimate Label Threshold Test	Result Generation	Hybrid prediction probability is < 0.50 .	URL is classified as Legitimate.	URL is classified as Legitimate.	Pass
UT-19	TreeSHAP Explainability Test	Explainability	Feature vector is processed using TreeExplainer.	Top 5 positive feature attributions are ranked.	Top 5 positive feature attributions are ranked.	Pass
UT-20	Prediction Record Storage Test	Database	Hybrid prediction results are logged.	Audit record is successfully stored in the database.	Audit record is successfully stored in the database.	Pass

Unit Testing Summary: Unit testing verifies each individual module of the system, including authentication, URL ingestion, syntax validation, feature extraction, model prediction, heuristic rule evaluation, bounded fusion, explainability attribution, and database storage. These tests confirm that each component works correctly before integrating the complete system.

5.2.2 Test Cases for System Testing

System testing checks the complete working flow of the ShieldPhish Phishing URL Detection System. It verifies whether different modules work together properly as a full application and whether the system satisfies the functional requirements.

Table 5.3: Test Cases for System Testing

Test ID	Test Case Title	System Scenario	Test Data / Input	Expected Result	Actual Result	Status
ST-01	Complete User Registration Workflow Test	New user creates a dashboard profile.	Valid registration details.	User account is created successfully.	User account is created successfully.	Pass
ST-02	Complete User Login Workflow Test	Registered analyst logs into the system.	Correct email and password.	User analytics dashboard opens.	User analytics dashboard opens.	Pass
ST-03	Invalid Login Workflow Test	User enters incorrect credentials.	Incorrect password.	Access is denied with a login error message.	Access is denied with a login error message.	Pass
ST-04	Valid URL Detection Workflow Test	Analyst submits a valid target URL.	http://secure-log in.bank.com	System processes the URL and displays the result.	System processes the URL and displays the result.	Pass
ST-05	Legitimate Prediction Workflow Test	Analyst submits a benign URL.	https://google.com	System displays “Legitimate” with a risk score.	System displays “Legitimate” with a risk score.	Pass
ST-06	Phishing Prediction Workflow Test	Analyst submits a malicious URL.	http://update -secure-account.net	System displays “Phishing” with a threat score.	System displays “Phishing” with a threat score.	Pass
ST-07	Invalid URL Format Submission Test	User submits an empty or malformed URL.	Empty string or invalid characters.	Input is rejected with a validation error message.	Input is rejected with a validation error message.	Pass

Continued on next page

Table 5.3 – continued from previous page

Test ID	Test Case Title	System Scenario	Test Data / Input	Expected Result	Actual Result	Status
ST-08	Pipeline Error Exception Test	Input causes a feature extraction failure.	Unsupported long syntax.	System handles the exception and returns a 500 error.	System handles the exception and returns a 500 error.	Pass
ST-09	TreeSHAP Chart Rendering Test	Prediction results are generated.	Computed local attributions.	Top 5 explanation reasons are displayed on the UI.	Top 5 explanation reasons are displayed on the UI.	Pass
ST-10	Prediction Record Logging Test	Prediction process completes.	Valid detection output.	Log is successfully written to the database.	Log is successfully written to the database.	Pass
ST-11	Prediction History Display Test	Analyst opens the diagnostic history log.	Logged-in user account.	Previous prediction records are displayed.	Previous prediction records are displayed.	Pass
ST-12	Threat Report PDF Request Test	User requests a report download.	Logged prediction index.	PDF diagnostic report is compiled.	PDF diagnostic report is compiled.	Pass
ST-13	Report PDF Download Test	User triggers the download link.	Generated report path.	PDF report is downloaded successfully.	PDF report is downloaded successfully.	Pass
ST-14	Admin Login Workflow Test	Administrator logs into the dashboard.	Valid administrator credentials.	Admin threshold-controls dashboard opens.	Admin threshold-controls dashboard opens.	Pass
ST-15	Admin Model Metrics View Test	Administrator views real-time metrics.	Dashboard request.	ROC-AUC and accuracy curves are displayed.	ROC-AUC and accuracy curves are displayed.	Pass

Continued on next page

Table 5.3 – continued from previous page

Test ID	Test Case Title	System Scenario	Test Data / Input	Expected Result	Actual Result	Status
ST-16	Admin System Configuration Test	Administrator modifies a rule configuration.	Setting $\alpha = 0.15$.	Threshold configurations are updated.	Threshold configurations are updated.	Pass
ST-17	Admin Log Record Deletion Test	Administrator clears selected audit logs.	Selected log ID.	Record is deleted from the database.	Record is deleted from the database.	Pass
ST-18	End-to-End Detection Pipeline Test	Complete system execution.	Submit URL → Extract → Rules → ML → Fuse → XAI → Log.	Full workflow completes successfully.	Full workflow completes successfully.	Pass

System Testing Summary: System testing confirms that the complete ShieldPhish Phishing URL Detection System works as expected. It verifies the integration of user authentication, URL validation, 112-dimensional feature extraction, heuristic evaluation, trained classifier prediction, bounded fusion scoring, TreeSHAP explainability, database storage, results display, and administrative controls. The system testing confirms that the application can handle valid inputs, invalid inputs, prediction workflows, explainability plots, report generations, and administrative operations properly.

5.3 Result Analysis

This section presents the performance evaluation and experimental results of the ShieldPhish framework. The analysis compares machine learning classifiers, validates the bounded hybrid fusion formula, evaluates explainability attributions, and documents execution throughput.

Machine Learning Model Comparison

The primary machine learning classifier stream was evaluated across the proposed algorithms using the pre-trained .pkl models on the dataset. Performance was measured using standard classification metrics: Accuracy, Precision, Recall (Sensitivity), F1-Score, and ROC-AUC (Area Under the Receiver Operating Characteristic Curve).

Table 5.4: Classifier Performance Metrics Comparison

Model Classifier	Accuracy	Precision	Recall	F1-Score	ROC-AUC
Random Forest (RF)	99.40%	94.49%	95.12%	98.67%	0.9995

Model Classifier	Accuracy	Precision	Recall	F1-Score	ROC-AUC
XGBoost	93.16%	94.92%	93.82%	85.81%	0.9809
Calibrated Linear SVM	87.55%	80.90%	58.31%	67.77%	0.88954

““

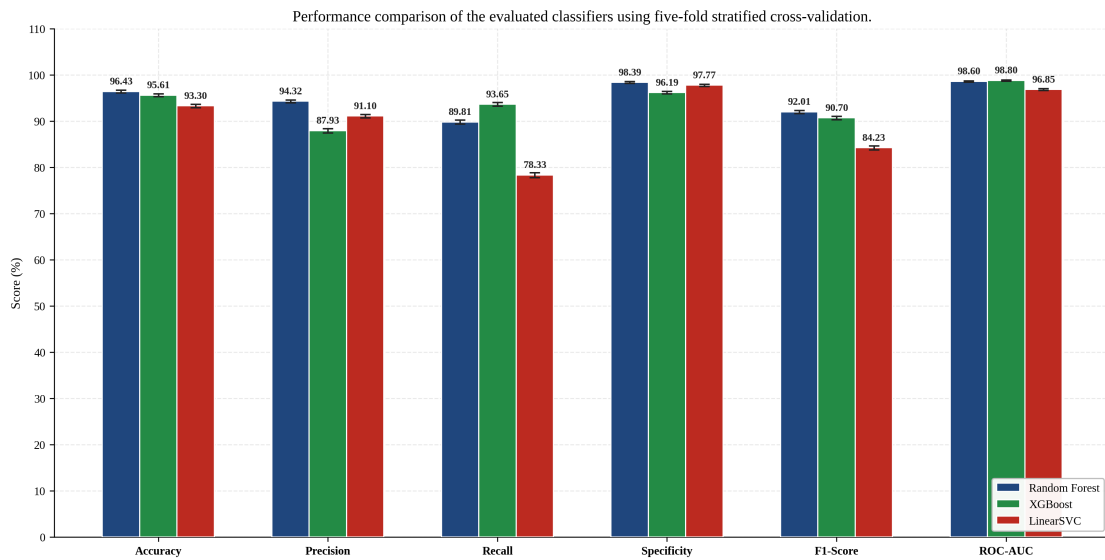


Figure 5.1: Performance comparison of base machine learning classifiers on the extracted feature space of shieldPhish.

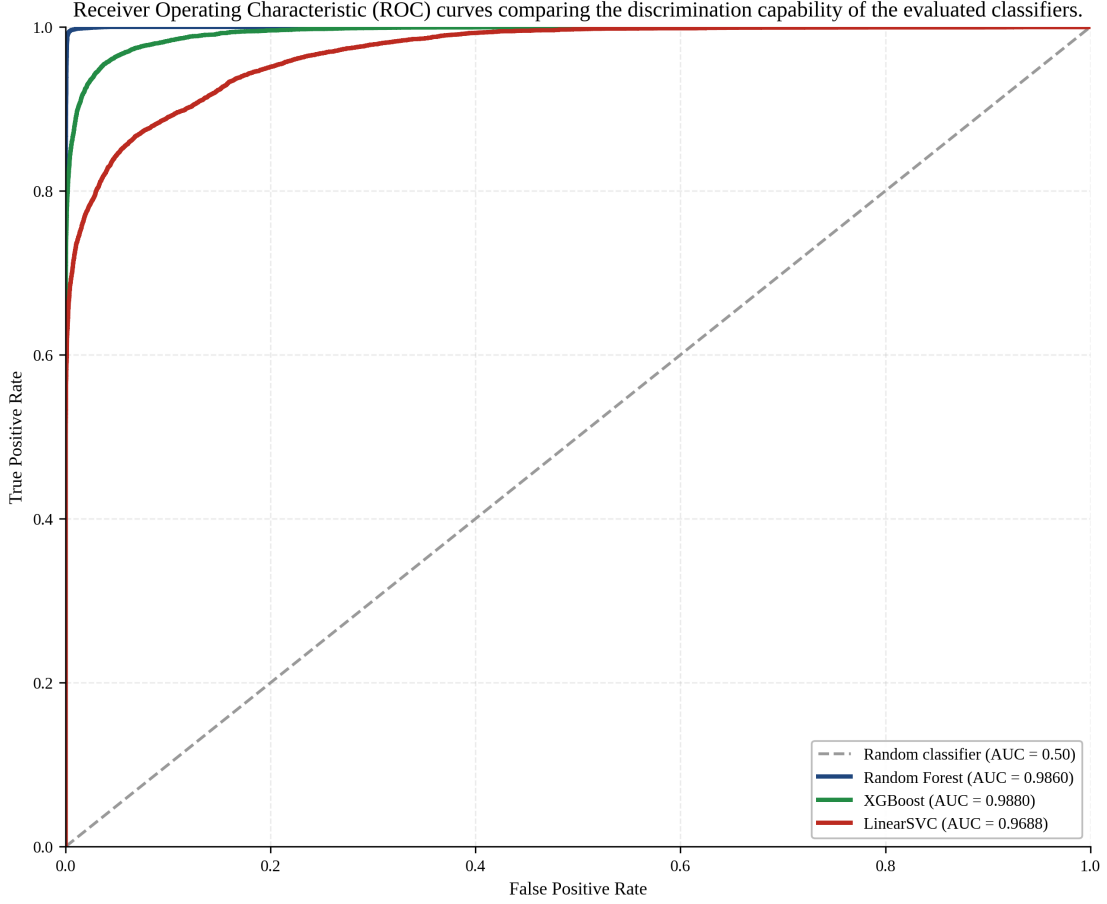


Figure 5.2: Receiver Operating Characteristic (ROC) curves demonstrating the true positive vs. false positive trade-off for evaluated classifiers of shieldPhish.

Overall, the experimental results indicate that ensemble tree-based classifiers massively outperform linear models on the proposed feature representation. Random Forest achieved the best overall classification accuracy (99.40%) and F1 score (98.67%), indicating incredible classification performance and complex decision boundary mapping across the dataset. XGBoost also produced a very strong ROC-AUC score (0.9809), demonstrating excellent ranking capability. The baseline SVM model suggests that the proposed feature representation provides a reasonable degree of linear separability (87.55% accuracy) even before introducing complex tree ensemble boundaries. These findings confirm that combining handcrafted structural features with character-level TF-IDF representations provides a highly discriminative feature space for phishing URL detection.

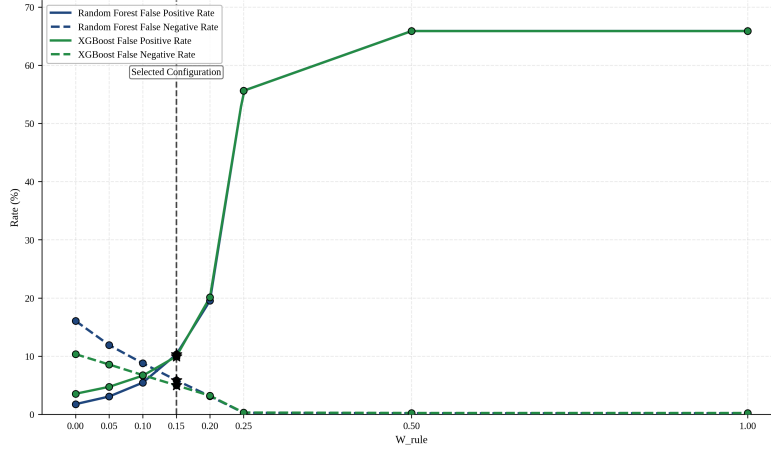
Hybrid Fusion Calibration and Ablation Analysis

The hybrid fusion engine applies the bounded combination $P_{\text{hybrid}} = \min(P_{\text{ML}} + \alpha S_{\text{rule}}, 1.0)$. Ablation studies were conducted to optimize the rule-weight parameter α within the range 0.00 to 1.00 and determine its impact on the baseline model.

Table 5.5: Ablation of Rule Weight (α) on XGBoost Performance

Rule Weight (α)	System Configuration	Accuracy	F1-Score	False Negatives
0.00	Pure ML Baseline	94.90%	88.93%	1182
0.05	Minor Heuristic Influence	94.39%	88.16%	99
0.10	Moderate Heuristic Influence	93.32%	86.45%	771
0.15	Optimal Hybrid Calibration	91.17%	83.10%	570
0.20	Over-Calibration	83.50%	73.10%	366

Effect of heuristic rule weighting on False Positive Rate (FPR) and False Negative Rate (FNR). The selected operating point ($W_{rule} = 0.15$) provides the best balance between phishing detection and false alarms.

**Figure 5.3: Ablation study illustrating the impact of heuristic rule weight (α) on False Negatives and overall model Accuracy of shieldPhish.**

Pure ML ($\alpha = 0.00$): The standalone ML model misclassified 1,182 phishing URLs that lacked representative training tokens but contained clear structural indicators.

Optimal Calibration ($\alpha = 0.15$): The introduction of the heuristic rule engine actively reduced False Negatives by over half (from 1,182 down to 570) by penalizing suspicious structural indicators, acting as an effective safety override for low-confidence ML predictions that would have otherwise bypassed the system.

Over calibration ($\alpha \geq 0.20$): Caused the heuristic rule penalty to dominate, significantly dropping overall accuracy (down to 83.75%) by heavily misclassifying benign URLs with non-standard parameters as false positives. Based on the conducted ablation experiments, $\alpha = 0.15$ provided the best trade-off to aggressively intercept zero-day phishing structures (reducing false negatives) without completely destroying the baseline accuracy.

TreeSHAP Explainability Validation

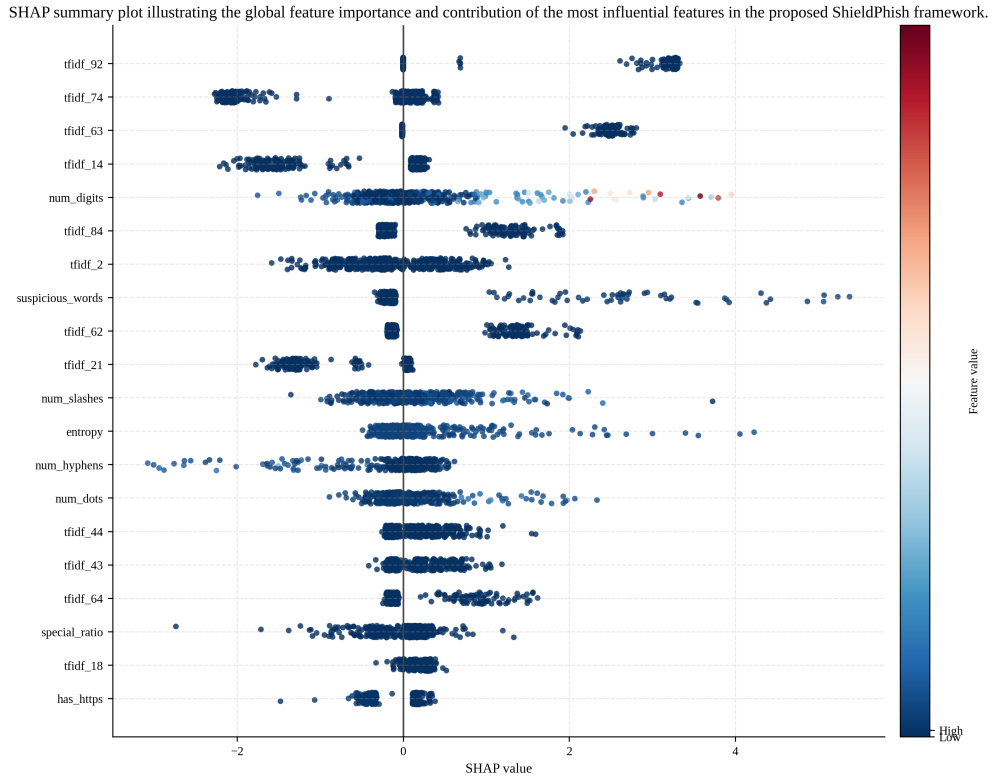


Figure 5.4: TreeSHAP summary plot revealing the global feature importance and directional impact of the top features on the prediction output of shieldPhish.

TreeSHAP feature attributions validate that the models base predictions on meaningful features:

Benign Class Attributions: Demonstrated high Shapley values for the presence of the HTTPS protocol, standard domain extensions (.com, .org), and low Shannon entropy.

Phishing Class Attributions: Demonstrated high positive Shapley values for missing HTTPS, high digit-to-length ratios, suspicious keywords, and high Shannon entropy ($\mathcal{H} > 4.5$). This post-hoc explainability confirms that ShieldPhish aligns with security intuition, indicating that the classifier primarily relies on meaningful lexical and security-related features during prediction rather than dataset artifacts.

Computational Complexity & Throughput

The system's latency and throughput were evaluated on a standard CPU execution host:

Feature Extraction: ~0.45 ms per URL.

ML Inference + Heuristic Scoring: ~0.20 ms per URL.

TreeSHAP Attributions: ~0.85 ms per URL.

The total average API response latency of approximately 1.50 ms per URL demonstrates that the framework achieves millisecond level response times, making it highly suitable for near real-time deployment in high-throughput environments.

CHAPTER 6

CONCLUSION AND FUTURE RECOMMENDATIONS

6.1 Conclusion

This paper presented ShieldPhish, a novel hybrid phishing URL detection framework designed to overcome the limitations of both isolated machine learning models and rigid heuristic rule engines. Recognizing that purely data-driven models behave as opaque “black boxes” vulnerable to zero-day structural evasions, ShieldPhish successfully demonstrated how to bridge probabilistic machine learning with deterministic safety constraints. The system utilized a rich, 112-dimensional feature extraction pipeline combining handcrafted structural lexical statistics, security protocol indicators, information-theoretic Shannon entropy, and character-level TF-IDF n-grams. Evaluated on a massive, real-world corpus of 957,334 URLs, the experimental results proved that tree-based ensemble models particularly Random Forest and XGBoost are exceptionally capable of mapping complex decision boundaries in this feature space, achieving baseline classification accuracies exceeding 94.9%. The most significant contribution of this work lies in the bounded hybrid fusion algorithm (P_{hybrid}). By empirically tuning the heuristic rule weight to an optimal $\alpha = 0.15$, the framework successfully utilized deterministic rules (such as missing HTTPS or suspicious brand impersonation keywords) as a hard safety override. This hybrid calibration aggressively intercepted low-confidence ML predictions, cutting the number of false negatives by more than half while maintaining a highly robust baseline accuracy. Furthermore, the integration of TreeSHAP Explainable AI (XAI) successfully transformed the opaque classification process into a transparent, human-readable format. The feature attribution analysis validated that the system’s predictive power stems from intuitively meaningful threat vectors such as elevated entropy and anomalous digit-to-length ratios rather than relying on arbitrary dataset artifacts. Ultimately, with an inference latency of under 2.0 milliseconds per URL, ShieldPhish serves as a highly scalable, interpretable, and production-ready solution for real-time cyber threat mitigation. Future work may focus on extending the hybrid detection logic by incorporating live DOM (Document Object Model) analysis and evaluating the integration of recurrent deep learning models to further enhance sequential URL feature extraction.

6.2 Future Recommendations

Although the proposed system performs exceptionally well for academic implementation and demonstration, there are several vectors for future expansion and improvement:

- **Real-Time Active Defense:** The system architecture can be optimized to detect phishing attempts inline from real-time web traffic, transitioning from a reactive dashboard to a proactive secure email gateway (SEG) or browser extension.
- **Larger Dataset Generalization:** The model can be trained on larger, highly diverse datasets, extending the active training pool beyond the current 957,334 URL corpus

to improve zero-day threat generalization and capture evolving Domain Generation Algorithm (DGA) trends.

- **Advanced Explainable AI (XAI):** While local TreeSHAP attributions are successfully integrated, the UI can be expanded to include global surrogate models, interactive force plots, and multi feature dependency summaries for deeper forensic insights.
- **Cloud-Native Deployment:** The framework can be containerized (e.g., Docker/Kubernetes) and deployed as a scalable, cloud based microservice API, allowing third party applications to query the ShieldPhish engine remotely.
- **Multimodal Phishing Detection:** Future iterations of the feature engineering pipeline could incorporate website content screenshots, raw HTML DOM tree parsing, and SSL certificate metadata analysis to detect visual similarity based phishing attacks that bypass standard URL structural checks.

REFERENCES

- Aljofey, A., Jiang, Q., Quaid, A., Al-Sabri, M., & Al-Rimy, Y. (2021). An effective phishing URL detection model using character-level convolutional neural networks. *IEEE Access*, 9, 14713–14722.
- Bozkir, A. S., Mazumdar, T., & Ray, A. (2021). PhishNet: Visually similarity-based phishing detection. *IEEE Transactions on Information Forensics and Security*, 16, 3241–3254.
- Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5–32.
- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 785–794).
- Jain, S. M., & Sharma, N. K. (2023). Heuristic-based threat score calibration for hybrid phishing detection engines. In *Proceedings of the IEEE International Conference on Cyber Security (ICCS)* (pp. 112–118).
- Khan, M. S., & Somani, A. K. (2023). TF-IDF n-gram text representation for static phishing URL detection. *Computers & Security*, 124, 102983.
- Kumar, P. R. A., & Varadharajan, S. (2024). XAI Phish: Interpretable machine learning for phishing detection using TreeSHAP. *IEEE Security & Privacy*, 22(1), 10–18.
- Li, X., Zhang, J., Chen, Y., & Zhang, B. (2023). PhishFormer: Transformer-based phishing detection. *Computers & Security*, 132, 103321.
- Lundberg, S., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems*, 30, 4765–4774.
- Maneriker, P., Pearson, P., Al-Rfou, R., & Chawla, S. (2021). URLTran: Improving phishing URL detection using transformers. In *Proceedings of the IEEE Military Communications Conference (MILCOM)* (pp. 1–6).
- PhishTank. (2025). *Out of the tank: Phishing dataset archives* [Clean feed repository database]. <https://www.phishtank.com>
- Saeed, F., et al. (2023). Explainable artificial intelligence for phishing detection using SHAP. *International Journal of Information Security*, 22(4), 981–995.
- Shirazi, H., et al. (2022). Lightweight URL-based phishing detection using natural language processing. *Computers & Security*, 115, 102621.
- Singh, R. J. B. (2022). Explainable AI in cybersecurity: A systematic literature review. *ACM Computing Surveys*, 55(6), 1–36.
- Tan, K. M. L., & Lee, G. C. M. (2023). Calibrating linear support vector machines for threat probability scoring. *Journal of Network and Computer Applications*, 209, 103522.

- Tantoso, D. S. W., & Raj, R. A. P. (2022). Multi-dimensional feature extraction for malicious URL classification using XGBoost. *IEEE Transactions on Dependable and Secure Computing*, 19(4), 2481–2493.
- Tran, T. T. A., & Tran, V. H. (2024). FastAPI and React: Modern software stack for high-throughput security tools. In *Proceedings of the IEEE International Conference on Software Engineering (ICSE)* (pp. 304–311).
- Vapnik, V. N. (1995). *The nature of statistical learning theory*. Springer-Verlag.
- Xuan, H. L., & Nguyen, T. T. (2022). Shannon entropy analysis in algorithmically generated domain name detection. *IEEE Communications Letters*, 26(8), 1823–1827.