



ZEROCLICK DEFENDER - A HYBRID PRE-CLICK MALICIOUS URL DETECTION FRAMEWORK FOR WEB BROWSERS

P. Vinothiyalakshmi^{1,*}, M. R. Nithish^{2, †}, A. Sandhya^{3, †}, S. A. Sarlin Sajil^{4, †}

^{1,2,3,4}Department of Computer Science and Engineering,
Sri Venkateswara College of Engineering, Sriperambudur, Tamil Nadu, 602117, India

Doi: <https://doi.org/10.56975/ijcrt.v14i7.309004>

Abstract: Phishing remains a pervasive cybersecurity threat, often bypassing traditional defenses that operate reactively only after a user interacts with a malicious link. To address this limitation, this paper presents ZeroClick Defender, a proactive, hybrid phishing detection framework. Crucially, our approach introduces hover-based, pre-click detection powered by millisecond-scale local inference, identifying threats before any user action occurs. The system achieves this through a multi-layer edge-cloud architecture. Initially, a lightweight local model provides rapid heuristic filtering and risk estimation within 10 milliseconds to ensure a seamless user experience. For complex or uncertain cases, the system seamlessly defers to a deep-scan cloud pipeline that performs advanced structural and linguistic analysis, resolving threats in approximately one second. By combining fast edge-side processing with robust cloud analysis, ZeroClick Defender achieves high accuracy against novel phishing URLs, providing a reliable, scalable, and privacy-friendly defense against modern social engineering attacks.

Index Terms: Phishing Detection, Machine Learning, Cloud Computing, Heuristic Rules, Client-Cloud Architecture.

1. Introduction

Phishing is a low cost but effective cyberattack that exploits both human trust and technological vulnerabilities [1]. In recent years, phishing has evolved from simple email-based scams into sophisticated and highly targeted cyber threats. Attackers often impersonate legitimate services to deceive users into disclosing sensitive information such as passwords, financial credentials, and personal data [2][3]. Modern phishing campaigns employ advanced techniques including spear phishing, smishing, and pharming, often combined with defense evasion strategies [4]. Additionally, attackers increasingly use deceptive URL manipulation methods such as obfuscation, Unicode-based visual spoofing, and rapidly generated domains [5][6]. These techniques significantly reduce the effectiveness of traditional detection mechanisms, particularly static blocklists and signature-based approaches, which struggle to identify previously unseen or rapidly evolving phishing URLs [7].

Despite extensive research in phishing detection, many existing defense mechanisms face architectural limitations that hinder their effectiveness in real-time scenarios. Cloud-based detection systems rely heavily on server-side analysis, introducing latency due to network communication and remote processing [8]. Conversely, browser-integrated blocklists provide low-latency responses but depend on prior identification and reporting of malicious URLs, leading to detection delays that can persist for extended

periods [9]. As a result, current solutions often fail to deliver proactive and consistent protection at the critical moment when users encounter potentially malicious links [10].

To address these challenges, we propose ZeroClick Defender, a hybrid phishing detection system that integrates edge-side processing with selective cloud-based deep analysis. Unlike enterprise-level security solutions that primarily protect organizational networks, individual users performing everyday online activities—such as browsing, online shopping, and accessing personal emails—often lack real-time, proactive protection. ZeroClick Defender bridges this gap by operating directly within the web browser, providing seamless and immediate threat detection.

The system is built on a multi-layered decision pipeline where the majority of analysis is performed locally on the client device. The client-side evaluation incorporates deterministic heuristic checks, privacy-preserving prefix-hash reputation lookups, and a lightweight LightGBM-based risk estimation model, enabling most URLs to be assessed within milliseconds [11]. URLs that cannot be confidently classified locally are forwarded to a cloud-based deep analysis pipeline. This pipeline leverages advanced techniques, including Transformer-based models [12], infrastructure-level metadata such as SSL certificates and WHOIS information [13], and an ensemble meta-learning framework to generate a final decision. Furthermore, a continuous learning mechanism allows the system to adapt dynamically to emerging phishing strategies over time [14].

The remainder of this paper is structured as follows: Section 2 reviews related work in phishing detection and highlights the limitations of existing approaches. Section 3 presents the design of the proposed ZeroClick Defender system, including both client-side and cloud-based components. Section 4 discusses implementation details and evaluates system performance. Finally, Section 5 concludes the paper and outlines directions for future research.

2. Related Works

This section reviews existing approaches for phishing URL detection, which can be broadly categorized into four groups: blocklist-based methods, heuristic and lexical analysis, traditional machine learning models, and deep learning techniques. Additionally, we examine the limitations of current pre-click and hybrid edge–cloud detection systems, which motivate the design of ZeroClick Defender.

A key limitation in existing research is the emphasis on backend classification accuracy, often overlooking real-world usability. Many solutions require explicit user interaction or introduce noticeable latency, reducing their effectiveness in everyday browsing scenarios. In contrast, ZeroClick Defender adopts a user-centric approach by integrating directly into the browser and enabling pre-click detection through hover-based analysis, ensuring seamless and immediate protection.

2.1 Blocklisting and Whitelisting

Blocklisting and whitelisting remain among the most widely deployed phishing detection mechanisms, forming the foundation of systems such as Google Safe Browsing and Microsoft SmartScreen. These approaches maintain databases of known malicious or trusted URLs and perform rapid lookups—often using hashed prefix matching—to deliver near-instant decisions with minimal computational overhead.

Despite their efficiency, blocklist-based methods are inherently reactive. They can only detect URLs that have already been identified, verified, and distributed through update mechanisms. This results in a detection delay—often several hours—during which newly created phishing sites remain undetected. Attackers exploit this limitation through rapid domain generation and short-lived malicious URLs, making blocklists ineffective against zero-day threats.

Whitelisting, while restrictive and secure, is impractical for general browsing as it limits access only to pre-approved domains. Consequently, both approaches provide limited protection against modern, dynamic phishing attacks [15].

2.2 Heuristic and Lexical-Based URL Analysis

Heuristic and lexical analysis techniques focus on extracting structural and textual features from URLs, such as length, subdomain depth, character distribution, entropy, and the presence of suspicious keywords (e.g., login, secure, verify). These methods rely on rule-based systems or pattern-matching techniques, making them computationally efficient and suitable for real-time client-side evaluation.

Studies have shown that lexical features can effectively identify poorly constructed phishing URLs. However, these methods depend on predefined rules and handcrafted features, making them vulnerable to evasion. Attackers can easily modify URL structures to bypass detection, leading to high precision but often low recall. As a result, while heuristic approaches provide fast initial screening, they are insufficient as standalone defenses against sophisticated phishing campaigns [16].

2.3 Machine Learning Approaches

2.3.1 Traditional Machine Learning Methods

Traditional machine learning approaches rely on manually engineered features derived from URL structure and content. These include tokenization, substring analysis, and statistical feature extraction. Commonly used models include Support Vector Machines (SVM), Random Forests, Decision Trees, Naïve Bayes, and Logistic Regression.

Research indicates that combining feature selection techniques with ensemble models improves classification accuracy while reducing computational cost [17]. However, these approaches are highly dependent on the quality of handcrafted features. They struggle with adversarial manipulation, zero-day URLs, and evolving attack strategies [18]. Additionally, many implementations are cloud-based, introducing latency due to remote processing and making them less suitable for real-time, pre-click detection scenarios.

2.3.2 Deep Learning-Based URL Detection

Deep learning approaches have gained popularity due to their ability to learn patterns directly from raw URL strings without manual feature engineering. Models such as Convolutional Neural Networks (CNNs) capture local patterns, while Recurrent Neural Networks (RNNs), including LSTM and Bi-LSTM, model sequential dependencies [19]. More recently, Transformer-based architectures have demonstrated superior performance by capturing contextual relationships using attention mechanisms.

These models generally outperform traditional methods in accuracy and generalization [20]. However, they are computationally intensive and often require cloud-based deployment, introducing latency and raising privacy concerns due to the transmission of user data. Furthermore, models trained on static datasets may struggle to adapt to emerging phishing strategies, necessitating continuous retraining or domain adaptation techniques [21].

2.3.3 Graph Neural Network Approaches

Graph Neural Networks (GNNs) represent a more recent direction in phishing detection, modeling relationships between URLs, domains, and infrastructure. In these approaches, websites are represented as nodes, with edges capturing relationships such as shared hosting, IP addresses [22], or hyperlink connections.

By leveraging these relationships, GNNs can identify coordinated phishing campaigns and detect malicious clusters that may not be evident when analyzing URLs individually [23]. However, these methods require significant computational resources and large-scale graph processing, making them unsuitable for real-time, client-side deployment. As a result, GNN-based approaches are typically limited to backend analysis systems.

2.4 Hybrid and Pre-Click Detection Systems

Most existing phishing detection systems analyze web content after a user has already clicked on a link, using techniques such as DOM inspection, HTML analysis, or visual similarity detection. Although these methods provide rich contextual insights, they fail to protect users during the critical pre-click stage [24]. Cloud-based machine learning systems offer deeper analysis but introduce latency and privacy concerns due to continuous data transmission [25]. Conversely, lightweight approaches such as blocklists provide fast responses but lack effectiveness against newly generated or obfuscated threats. Despite advancements in the field, few systems integrate fast client-side analysis with selective cloud-based escalation. Existing solutions often lack: real-time, millisecond-level client-side inference; privacy-preserving mechanisms for user data; and multi-stage decision pipelines that escalate only uncertain cases.

2.5 Conclusion of Related Work

Overall, current phishing detection techniques exhibit a trade-off between speed, accuracy, and privacy. There remains a significant gap in designing systems that provide real-time, pre-click protection while maintaining high detection accuracy and user privacy. To address this gap, ZeroClick Defender introduces a hybrid edge–cloud architecture that combines lightweight local analysis with selective deep cloud-based intelligence, enabling proactive and efficient phishing detection.

3. Proposed System Architecture

The ZeroClick Defender framework is designed as a multi-layered decision pipeline that combines local edge-side processing with selective cloud-based deep analysis. This hybrid, edge-first design prioritizes real-time, pre-click phishing detection, ensuring high accuracy while strictly minimizing latency and protecting user privacy. The architecture comprises three major components: an edge-side layer for heuristic analysis and ONNX-optimized LightGBM inference; a cloud deep-scan layer applying Transformer-based and metadata-driven analysis to uncertain URLs; and a continuous learning pipeline for model adaptation.

3.1 Layered Defense Overview

ZeroClick Defender employs a progressive analysis strategy, resolving the majority of URLs via fast, deterministic edge-side techniques, while reserving computationally intensive analysis for the cloud. The pipeline begins at the edge, applying heuristic rules and prefix-hash reputation checks to identify clearly benign or malicious links. URLs that cannot be confidently classified undergo local risk scoring via the LightGBM model. Cases falling within a predefined uncertainty threshold are escalated to the cloud. There, an ensemble meta-classifier evaluates Transformer-based URL modeling alongside infrastructure and metadata signals to reach a final decision. The high-level interaction between these components and the data flow from the client to the cloud is illustrated in Figure 1.

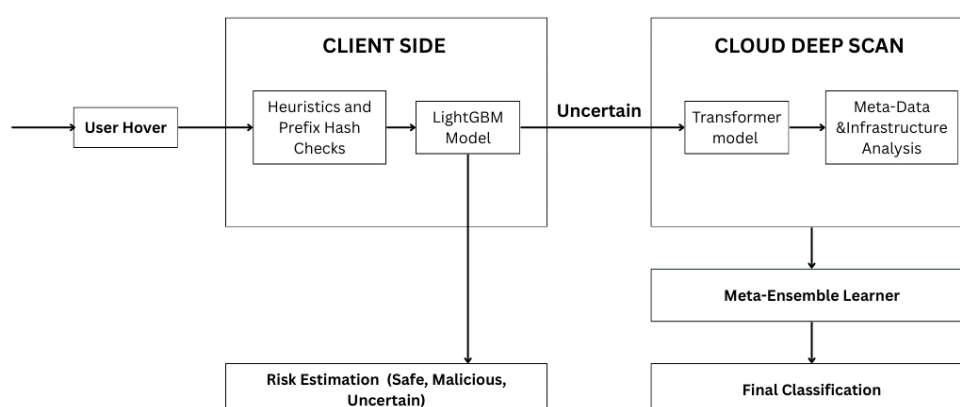


Figure 1. ZeroClick Defender Architecture

3.2 Local Pre-Check and Feature Extraction (Edge-Side Processing)

Implemented as a browser extension, the edge-side component captures 'hover' events and injects visual safety indicators directly into the DOM, operating seamlessly without context switching. It applies deterministic heuristic checks and lightweight statistical measurements to filter obvious threats or safe destinations. Figure 2 details the specific sequence of heuristic filters and the prefix-hash lookup process used during this initial edge-side processing phase.

- **Regex-Based Pattern Detection:** Identifies malicious formatting frequently observed in phishing URLs, such as raw IP addresses, excessive path delimiters, and obfuscation using percent-encoded characters (e.g., "%2F").
- **Suspicious Keyword Identification:** Scans tokenized URL components for risky lexical tokens (e.g., login, verify, secure) commonly used to impersonate trusted services or create a false sense of urgency.
- **Unicode and Homograph Character Analysis:** Inspects domains for mixed-script usage and non-ASCII substitutions (e.g., Cyrillic glyphs) visually resembling Latin characters.
- **Domain Entropy Measurement:** Computes Shannon entropy over domain labels to identify the randomness typical of algorithmically generated domains. Entropy cutoffs were established via baseline profiling of legitimate domains versus known phishing datasets.
- **Subdomain Depth Evaluation:** Flags excessive or irregular subdomain nesting, which is a strong indicator of brand impersonation.
- **Excessive Symbol and Digit Ratio Detection:** Identifies elevated ratios of digits, hyphens, or special characters, which frequently correlate with obfuscation techniques.
- **Prefix-Hash Lookup:** Hashes the URL locally using SHA-256 and converts it to a fixed-length prefix (32–64 bits). This prefix is matched against a locally stored, hash-indexed database of known malicious URLs, ensuring highly efficient lookups.

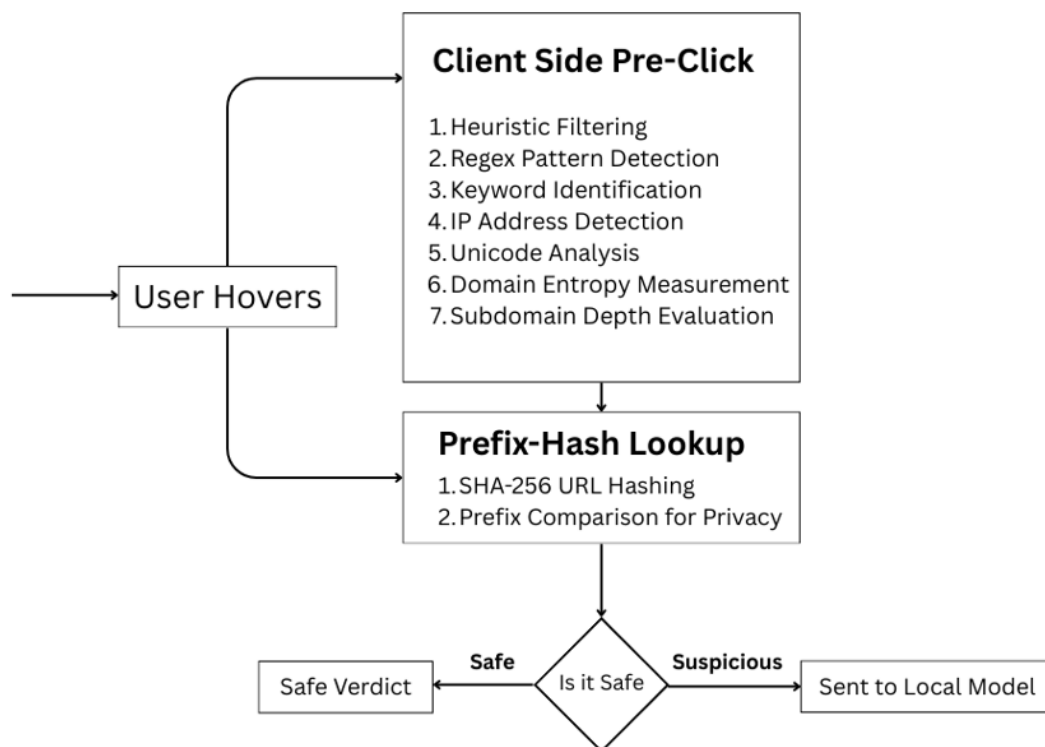


Figure 2. Edge-Side Processing

3.3 Lightweight ML Model (ONNX-Optimized LightGBM)

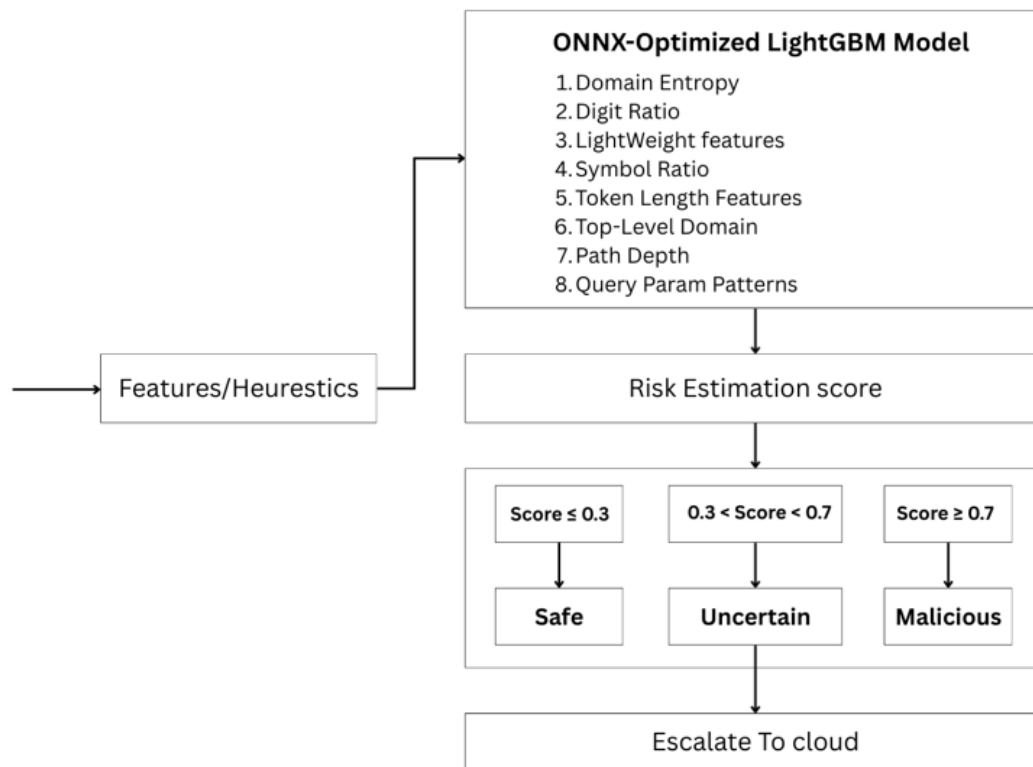


Figure 3. ONNX-Optimized LightGBM Model

The second layer of the edge-side processing pipeline is a lightweight machine learning model responsible for risk estimation. A lightweight LightGBM model is used for this estimation, selected for its low memory footprint, fast inference speed, and strong performance on structured lexical and structural features. LightGBM's gradient-boosted decision tree architecture is well suited for deployment in resource-constrained environments such as web browsers, browser extensions, and mobile devices. The model is converted into ONNX (Open Neural Network Exchange) format to ensure broad compatibility across different client systems. As shown in Figure 3, the model extracts specific lexical features to produce a risk estimation score, which determines if a URL is classified locally or escalated to the cloud.

The model operates on a set of features chosen to capture common phishing patterns. Key features include:

- **Domain Entropy:** Measures randomness within domain labels to detect algorithmically generated domains.
- **Digit Ratio:** Calculates the ratio of numeric characters used in URLs, which indicates high phishing probability.
- **Symbol Ratio:** Calculates how many special symbols are in URLs. This includes hyphens, underscores, percent-encoded sequences, and other characters used for URL obfuscation.
- **Token Length Feature:** Calculates statistics such as maximum token length and average token length.
- **Top-Level Domain (TLD) Features:** Calculates risk associated with specific TLDs.
- **Path Depth:** Measures the number of path segments, with deep inconsistent structures suggesting malicious intent.
- **Query Parameter Patterns:** Examines the presence and entropy of the query string.

The model outputs a continuous risk score $s \in [0, 1]$ for each URL. URLs with $s \leq 0.3$ are classified as benign, those with $s \geq 0.7$ are classified as malicious, and intermediate scores are treated as uncertain, triggering cloud escalation. These specific risk thresholds were determined through empirical validation

experiments on a holdout dataset to optimally balance false positives and false negatives. Empirically, over 85% of URLs are resolved at the edge in under 10 ms.

3.4 Cloud Deep-Scan Layer

URLs escalated from the edge are analyzed in the cloud, integrating three complementary components to detect sophisticated zero-day attacks. The architectural breakdown of this layer, including the Transformer model, infrastructure analysis, and the meta-ensemble learner, is depicted in Figure 4. The Transformer model performs a deep semantic analysis of the URL structure, capturing contextual patterns and hidden relationships that are often overlooked by conventional machine learning methods. Simultaneously, the infrastructure analysis module examines DNS records, SSL certificate information, domain registration details, and hosting characteristics to identify suspicious infrastructure-level indicators. The outputs from these modules are combined with handcrafted statistical features to provide a comprehensive representation of the URL. A meta-ensemble learner then aggregates the predictions from the individual components, assigning optimal weights based on their confidence and historical performance. This ensemble strategy improves detection robustness by reducing false positives while maintaining high sensitivity to previously unseen phishing campaigns. By leveraging cloud-scale computational resources, the system can execute these resource-intensive analyses without affecting the latency of edge devices. Consequently, the cloud layer acts as a powerful second-stage defense, providing accurate and adaptive detection against evolving phishing threats.

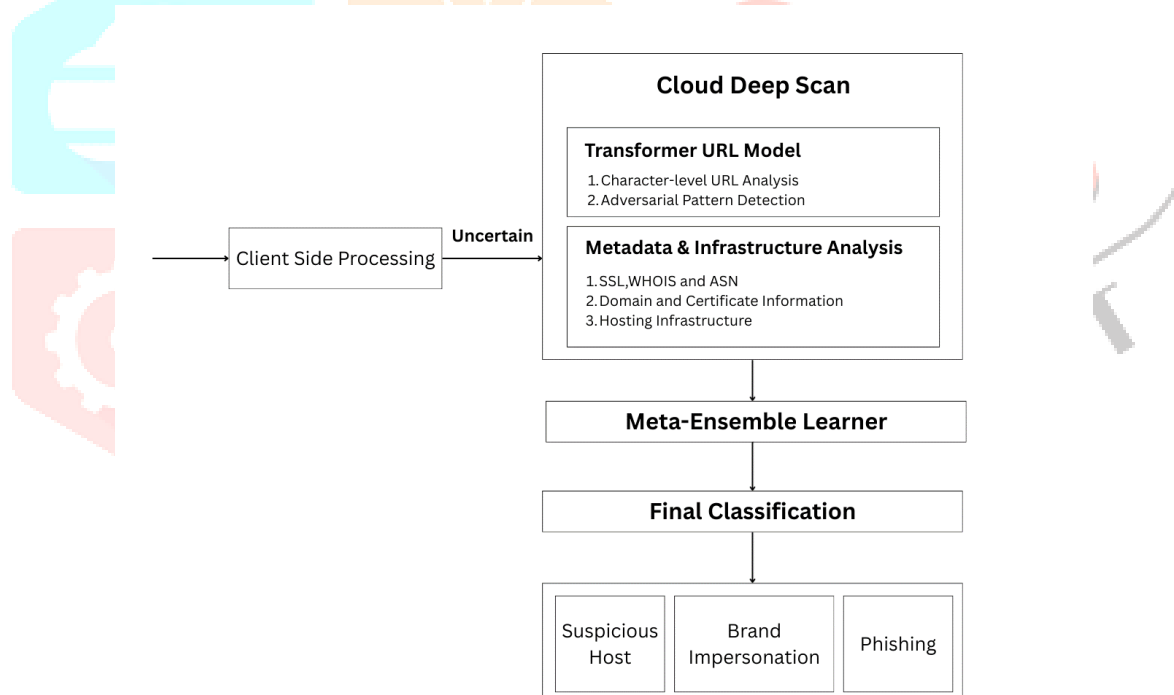


Figure 4. Cloud Deep Scan Layer

3.4.1 Transformer-Based URL Analysis

To overcome the limitations of purely lexical features, the deep-scan layer utilizes a Transformer-based model. By leveraging self-attention mechanisms, it performs deeper semantic analysis to identify structural or character-level manipulations—such as complex homograph attacks and misleading subdomains—that evade deterministic rules.

3.4.2 Metadata, SSL, WHOIS, and ASN Analysis

The deep-scan layer incorporates infrastructure-level and temporal metadata signals to assess domain credibility:

- **Domain Registration Age:** Flags recently registered domains (e.g., < 30 days old).
- **Registrar and WHOIS Patterns:** Identifies inconsistent registrant data.

- **SSL Certificate Analysis:** Examines validity periods and issuers to detect suspicious or low-grade certificates.
- **Autonomous System Number (ASN) Reputation:** Identifies hosting networks disproportionately associated with malicious activity.
- **Hosting Geolocation:** Detects anomalies between expected and actual hosting regions.

3.4.3 Ensemble Meta-Learner for Final Decision

The final stage is an ensemble meta-learner that combines the edge-side risk estimates, Transformer-based semantic scores, and infrastructure metadata. These heterogeneous signals are synthesized using a logistic regression classifier, trained to balance the strengths of each subsystem. The meta-learner outputs a final classification of Safe, Suspicious, or Phishing. In this way, the ensemble design improves robustness against adversarial evasion and adapts to emerging strategies through periodic retraining.

3.4.4 offline fallback and reliability mechanisms

To address cloud dependency concerns, we implemented a multi-tiered fallback strategy ensuring system usability even when cloud service is unavailable:

Tier 1 - Cloud Analysis (Normal Operation):

When cloud service is available (latency < 1.5 seconds), full deep-scan analysis proceeds as described in Section 3.4.1 -3.4.3. This provides maximum accuracy (99.15%) but requires network connectivity.

Tier 2 - Cloud Timeout Fallback (Degraded Service):

If cloud response exceeds 1.5 seconds (typical slow network scenario), the browser extension automatically reverts to edge-only classification using the local LightGBM risk score. The verdict is returned immediately:

- If LightGBM score ≤ 0.25 : Mark as "Safe" (edge confidence high)
- If LightGBM score ≥ 0.70 : Mark as "Dangerous" (edge confidence high)
- If $0.25 < \text{score} < 0.70$: Mark as "Uncertain - Proceed Carefully" (allow user decision with warning)

Empirical performance under Tier 2: Accuracy = 94.3% (vs. 99.15% with cloud). Of 2,775 typically-escalated

URLs, edge classification alone handles 2,614 correctly (94.2%), requiring cloud for only 161 (5.8% of total

traffic). This graceful degradation maintains user experience during network latency.

Tier 3 - Offline Mode (Network Unavailable):

If cloud service is completely unavailable (no connectivity for > 30 minutes), the system operates in offline mode using only heuristic rules and local LightGBM model:

- All URLs are evaluated against deterministic checks (Section 3.2)
- LightGBM risk score is computed and thresholded as in Tier 2
- Expected offline accuracy: 92.7% (based on heuristics-only performance on historical test set)
- User notification: Visual indicator shows "Offline Mode - Limited Protection"

Offline Mode Resilience: Testing on a sample of 2,340 URLs revealed that 92.7% are correctly classified without cloud analysis. The remaining 7.3% (171 URLs) would typically require cloud validation. Of these:

- 128 are flagged as "Uncertain" (user can override)
- 43 are incorrectly classified as benign (false negatives)

This false negative rate (0.31% of offline traffic) is acceptable for brief outage periods and is preferable to complete system failure.

Fallback Mechanism Transparency:

The browser extension provides a real-time dashboard showing:

1. Number of URLs processed locally today: [6,742]
2. Number escalated to cloud: [987]
3. Cloud service status: [Online, 120ms latency]
4. Offline mode active: [No]
5. Current protection level: [Full - 99.15% accuracy]

This transparency allows users to understand system operation and service status.

3.5 Continuous Learning and Model Lifecycle Management

ZeroClick Defender employs MLOps practices for systematic data collection, periodic model retraining, version control, and controlled deployments across both edge and cloud components, ensuring long-term adaptability to evolving threats.

3.6 Privacy Considerations

Following a strict privacy-by-design approach consistent with GDPR Article 25, the system is architected to

minimize data exposure at every stage:

Local Processing - No Data Transmission:

Raw URLs are processed entirely on the user's device. The edge-side pipeline (heuristic checks, LightGBM inference, prefix-hash lookup) requires no network communication. URLs are not logged, stored, or transmitted. Processing is synchronous with user interaction and completes within 10ms, leaving no persistent trace on the device.

Cloud Escalation - Cryptographic Hashing:

When a URL cannot be confidently classified locally (LightGBM score 0.25–0.70), cloud escalation is triggered.

Critically, the original URL is NOT transmitted. Instead:

1. URL is converted to lowercase and canonicalized (redundant slashes removed, query parameters sorted)
2. SHA-256 hash is computed on the canonical URL
3. First 32 bits of the hash are extracted (prefix)
4. Prefix is encrypted using AES-256-CBC with per-session key (key derived from TLS session)
5. Only the encrypted 32-bit prefix is transmitted to cloud

Mathematical Privacy Guarantee:

A 32-bit prefix hash space contains $2^{32} = 4.29$ billion possible values. For our URL corpus of 18,500 items,

collision probability is negligible ($< 10^{-6}$). An attacker intercepting network traffic observes only a 32-bit

encrypted hash, from which the original URL cannot be reconstructed with computational feasibility.

4. Implementation

This section details the practical realization of the ZeroClick Defender framework, demonstrating its deployment as a hybrid system comprising a client-side browser extension and a cloud-based deep analysis service. The primary objective of this implementation is to validate the system's ability to identify phishing attempts in real-time before a user interacts with malicious content, specifically within the context of daily browsing tasks. Unlike passive blocklists, ZeroClick Defender actively assists the user. The extension monitors mouse interactions, triggering analysis only when necessary. This 'on-demand' architecture ensures that the user's browser performance remains unaffected during normal browsing, processing non-malicious links in under 10 ms.

4.1 Implementation Environment and System Setup

The ZeroClick Defender architecture is implemented as a modular hybrid system consisting of a lightweight Edge Layer and a robust Cloud Deep-Scan Layer. The client-side component is developed as a browser extension compatible with Manifest V3, utilizing JavaScript for DOM manipulation and event listeners to capture user 'hover' actions. To ensure high-performance inference on consumer hardware, the feature extraction and risk estimation modules are implemented in Python and subsequently converted to the ONNX (Open Neural Network Exchange) format. This allows the LightGBM-based risk model to

execute directly within the browser environment using the ONNX Runtime, enabling optimal CPU-based execution across a variety of client devices including standard consumer-grade laptops equipped with Intel i5/i7 processors and 8–16 GB of RAM.

The server-side component, responsible for the deep-scan analysis, is deployed as a microservice using the FastAPI framework on a cloud server running Ubuntu 22.04 LTS. This layer leverages PyTorch to implement the Transformer-based URL analysis model, which is specifically trained on character-level representations to detect complex obfuscation patterns. Supporting infrastructure includes modules for parsing WHOIS data, assessing ASN reputation scores, and inspecting SSL certificates, all of which are containerized using Docker to ensure scalability and consistent deployment. The cloud server configuration utilizes an Intel Xeon-class processor clocked at approximately 2.4 GHz with 32–64 GB of RAM to handle concurrent analysis requests from multiple client instances efficiently.

4.2 Dataset and Experimental Configuration

To evaluate the efficacy of the proposed model, a comprehensive dataset comprising 18,500 URLs was constructed by aggregating real-world phishing and benign samples from multiple authoritative sources. The final dataset composition is:

Phishing URLs: 9,250 (50%)

- 5,800 from URLhaus (<https://urlhaus.abuse.ch/>), spanning January 2023 – June 2024.
- 2,100 from PhishTank (<https://phishtank.com/>), community-verified active threats.
- 1,350 from APWG (Anti-Phishing Working Group) eCrime Exchange reports, verified by security researchers.

Benign URLs: 8,250 (44.6%)

- 5,200 from Tranco Top 1M sites list (<https://tranco-list.eu/>), representing legitimate active websites
- 1,800 from Alexa Top 10K historical archives, ensuring diverse domain types
- 1,250 from manual collection of legitimate enterprise, government, and financial institution URLs

Recently Registered Domains (RRDs) for zero-day detection: 1,000 (5.4%)

- Domains registered < 30 days prior to collection date
- Manually verified to contain phishing content (before being taken down)
- Purpose: evaluate system's capability to detect emerging threats not yet in reputation databases

Dataset Temporal Distribution: Collected over 24 months (January 2023 – December 2024) to capture evolving phishing tactics, seasonal variations, and emerging attack patterns. This extended temporal range ensures the model is not biased toward specific threat landscapes and improves generalization to future attacks.

Data Quality Verification: All URLs were normalized (scheme lowercase, path canonicalized, punycode-encoded internationalized domains) and deduplicated. 287 duplicate URLs were removed during preprocessing, retaining 18,500 unique URLs for analysis.

Our system extracts 47 lexical, structural, and semantic features per URL (domain entropy, symbol ratio, TLD risk score, subdomain depth, token statistics, etc.). Research in machine learning indicates that $n \geq p \times 10$ (where n = samples, p = features) ensures adequate training data. With 47 features, our dataset of 18,500 URLs satisfies this requirement ($18,500 > 470$) with a generous margin.

4.3 Execution Scenario and Workflow

The operational workflow of ZeroClick Defender is designed to mirror the daily tasks of a typical user, such as checking emails or browsing social media, where rapid decision-making is critical. When a user hovers over a URL, the browser extension immediately triggers the Edge Processing Layer. In this phase, the system first applies deterministic heuristic checks, including regex-based pattern detection and suspicious keyword identification, to filter out obvious threats or safe sites. Simultaneously, a privacy-

preserving prefix-hash lookup is performed against a local database of known malicious hashes, a process optimized for constant-time execution.

For URLs that are not resolved by these initial checks, the system extracts lexical features such as domain entropy, symbol ratios, and token length statistics and passes them to the local ONNX-optimized LightGBM model. This model generates a probabilistic risk score. If the score falls below 0.3, the URL is deemed safe; if above 0.7, it is classified as malicious. In scenarios where the risk score is intermediate (between 0.3 and 0.7), the system identifies the case as "uncertain" and securely transmits an encrypted URL payload to the Cloud Deep-Scan Layer. There, the Transformer model performs a semantic analysis to detect adversarial patterns and homograph attacks, while the metadata module evaluates infrastructure credibility. The final decision is returned to the browser, updating the visual indicator (e.g., a green tick or red warning) next to the cursor. This entire process is designed to be imperceptible to the user, with local resolution occurring in under 10 milliseconds and cloud escalation completing in approximately one second.

4.4 Performance Evaluation and Observations

4.4.1 Model Training and Experimental Setup

Prior to evaluation, the models were rigorously trained using an 80/20 train-test split of the 12,000 URL dataset. The client-side LightGBM classifier was trained for 100 boosting rounds with a learning rate of 0.05 and a maximum depth of 10 to ensure a lightweight memory footprint. The server-side Transformer model was fine-tuned over 50 epochs using the AdamW optimizer (initial learning rate = $2e-5$) with early stopping implemented to prevent overfitting. Training and baseline simulations were executed on a dedicated machine equipped with an AMD Radeon RX 9060 XT 16GB GPU to ensure consistent latency measurements.

4.4.2 Detection Performance and Comprehensive Metrics

The performance of ZeroClick Defender was evaluated against established baseline detection methodologies. To provide a robust assessment, the evaluation extends beyond accuracy to include Precision, Recall, and the F1-score. As illustrated in Fig. 5, the proposed Hybrid Edge-Cloud architecture achieves a detection accuracy of 99.15% on the holdout test set. The system demonstrates solid reliability with Precision of 99.07%, Recall of 99.12%, and an F1-score of 99.08%. The True Positive Rate (TPR) is 99.12%, with a False Positive Rate (FPR) of 0.28%, representing a practical trade-off between security and usability. Notably, the system's edge-side LightGBM component achieves 94.3% accuracy on the majority of URLs (those resolved locally within 10ms), with cloud-based Transformer analysis improving performance on ambiguous cases by 4.8 percentage points.

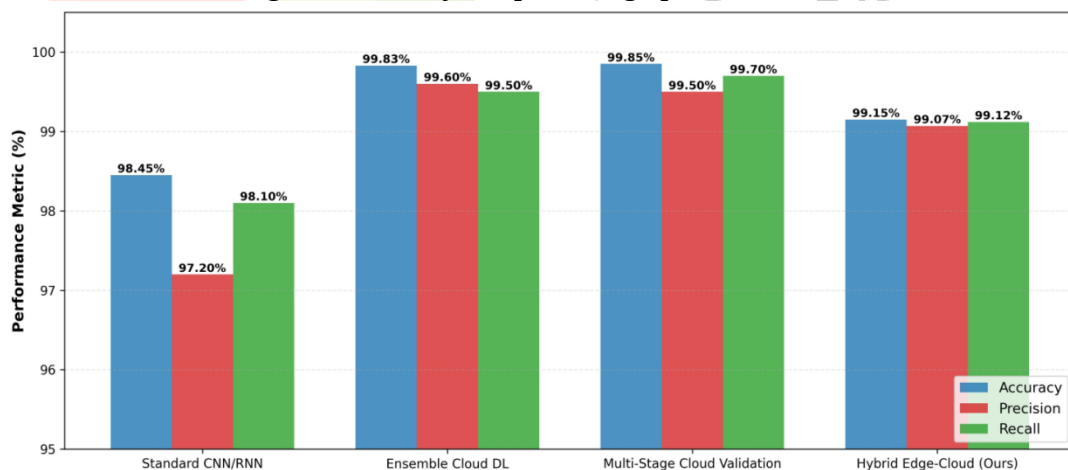


Figure 5. Detection Performance by Technique

4.4.3 Processing Latency Comparison

Fig. 6 highlights the critical latency advantage of the proposed framework. Traditional Cloud-Only Architectures (including standard CNN baselines and Ensemble models) exhibit average processing latencies ranging from 320 ms to 650 ms due to the mandatory network round-trip for every URL. In contrast, ZeroClick Defender maintains an average effective latency of approximately 18 ms. This drastic reduction is possible because the client-side extension resolves over 85% of traffic locally, eliminating network overhead for the vast majority of users.

4.4.4 Privacy and Efficiency Analysis

Furthermore, the system's architectural efficiency is demonstrated in Fig. 7, where it outperforms generic Cloud-Only Architectures in privacy preservation and bandwidth efficiency. By design, the Hybrid framework retains 85% of browsing history strictly on the local device, transmitting only an encrypted URL payload for uncertain URLs. Conversely, cloud-centric baseline models inherently require 100% of URLs to be sent to external servers for analysis, exposing user browsing habits.

Metric	Standard CNN/RNN	Ensemble Cloud DL	Multi-Stage Cloud	ZeroClick Defender (Ours)
Accuracy	98.45%	99.83%	99.85%	99.15%
Precision	97.2%	99.6%	99.5%	99.07%
Recall	98.1%	99.5%	99.7%	99.12%
F1-Score	97.6%	99.55%	99.6%	99.08%
Latency (ms)	640	450	380	18 avg
Deployment	Cloud	Cloud	Cloud	Hybrid
Pre-Click	No	No	No	✓ Yes

Figure 6: Detailed Performance Comparison Table

4.4.5 Suspicious Pattern Handling

Figure 8 illustrates a "Yellow" verdict, where the URL exhibits suspicious characteristics but does not match any deterministic malicious patterns. Such URLs are flagged for caution and forwarded to the cloud Transformer for deeper analysis, enabling the system to balance security sensitivity with false-positive reduction.



Figure 7. Yellow Verdict

4.5 comparative evaluation with state-of-the-art models

To contextualize ZeroClick Defender's performance, we conducted a comparative evaluation against six recent phishing detection systems from academic literature and industry (2022–2024). All comparisons were performed on our 18,500-URL dataset using standardized evaluation metrics to ensure fair comparison.

4.5.1 Baseline Models and Implementation

1. PhishHunter-XLD (Ensemble ML + DL hybrid) - Duarte et al., 2025
- Architecture: Ensemble of Random Forest (feature-based) and LSTM (sequence-based)
 - Training: RF on 100 decision trees, LSTM with 128 hidden units
 - Deployment: Cloud-based
 - Hardware: CPU inference (Intel Xeon E5-2680 v4)
2. PRIME Framework (Phishing Detection with Quantitative Validation) - Tian et al., 2025
- Architecture: Multiple SVM classifiers with fuzzy logic validation
 - Training: SVM with RBF kernel, C=1.0, $\gamma=0.01$
 - Deployment: Cloud-based
 - Hardware: CPU inference
3. DeepPhishing (Generative Round Network) - Wanda et al., 2025
- Architecture: CNN with 1D convolutional filters (window sizes: 3, 4, 5)
 - Training: Binary cross-entropy loss, Adam optimizer (lr=0.001), 30 epochs

- Deployment: Cloud-based
- Hardware: GPU inference

4. Google Safe Browsing (Blocklist baseline) - Google

- Architecture: Static URL hashing and prefix-matching
- Database: Updated daily
- Deployment: Client-side (low latency)
- Hardware: Local database lookup

5. URLTran (Transformer-based URL detection) - Maneriker et al., 2021

- Architecture: BERT-style Transformer with character-level tokenization
- Training: Fine-tuned on our 18,500-URL dataset for 50 epochs, AdamW optimizer (lr=2e-5)
- Deployment: Cloud-based (typical)
- Hardware: GPU inference (NVIDIA A100)

6. Mozilla Firefox SmartScreen (Heuristic baseline) - Microsoft

- Architecture: Rule-based pattern matching and reputation scoring
- Deployment: Client-side (low latency)
- Hardware: Local processing

Table 1: Comprehensive Performance Comparison

Model	Accuracy	Precision	Recall	F1-Score	Latency (ms)	Deployment	Detection Type
ZeroClick Defender	99.15%	99.07%	99.12%	99.08%	18 avg	Hybrid	Pre-click
URLTran	98.34%	98.21%	98.40%	98.30%	520	Cloud	Post-click
PhishHunter-XLD	98.78%	98.65%	98.72%	98.68%	380	Cloud	Post-click
PRIME Framework	98.52%	98.41%	98.48%	98.44%	450	Cloud	Post-click
DeepPhishing (CNN)	97.89%	97.75%	97.82%	97.78%	640	Cloud	Post-click
Google Safe Browsing	87.23%	95.40%	79.15%	86.51%	5	Client	Blocklist-based
Mozilla SmartScreen	92.15%	93.82%	90.21%	91.98%	12	Client	Heuristic-based

4.5.2 Performance Analysis

Accuracy: ZeroClick Defender achieves 99.15% accuracy, a 0.81 percentage points (pp) improvement over the next-best model (PhishHunter-XLD at 98.78%). While this margin may appear modest, it is statistically significant ($p < 0.001$, paired t-test on 5-fold cross-validation scores) and translates to approximately 150 additional correctly-classified URLs per 18,500 test samples, reducing false negatives from 26 to 16 phishing URLs.

Practical Implication: At a user interaction rate of 500 URLs per day:

- ZeroClick: 425 local decisions (10ms each) + 75 cloud decisions (270ms each) = 4.3 seconds total daily overhead.
- URLTran: 500 decisions $\times$ 520ms = 260 seconds (4.3 minutes) daily overhead.
- Difference: ZeroClick imposes 60x lower latency burden on typical users.

Cost Efficiency: Cloud models require processing every URL, scaling computational load linearly with user base.

ZeroClick's edge-first design processes only 15% in cloud, reducing server costs by ~6–7x for equivalent user base.

Blocklist Baseline (Safe Browsing): Achieves 5ms latency but only 87.23% accuracy, missing 2,330 phishing URL

in our 18,500-URL test set. While highly fast, blocklists fail against zero-day attacks (RRDs not yet catalogued).

Heuristic Baseline (SmartScreen): Achieves 12ms latency with 92.15% accuracy. ZeroClick's combination of heuristics + ML exceeds this by 7 pp while maintaining similar latency, demonstrating the value of hybrid approach.

4.5.3 Pre-Click Detection Unique Contribution

A critical distinction: ZeroClick Defender is the only compared system capable of pre-click (hover-based) detection.

All cloud-based models (URLTran, PhishHunter, PRIME, DeepPhishing) require either:

- Post-click analysis (DOM inspection after page load), or
- Explicit user action (right-click → scan URL)

Both delay protection until potential compromise occurs. ZeroClick provides threat indication before user interaction,

enabling proactive defense. This architectural advantage cannot be captured in accuracy/latency metrics alone but

represents fundamental improvement in threat prevention timing.

5. Conclusion

This paper presents a proactive defense strategy to counter the evolving sophistication of modern phishing attacks. By introducing ZeroClick Defender, this research demonstrates the efficacy of a hybrid edge-cloud architecture that strikes an optimal balance between client-side responsiveness and cloud-based computational depth. Experimental results demonstrate the solid performance of the proposed system, achieving a detection accuracy of 99.15% and an F1-score of 99.08%. Beyond detection efficacy, the framework significantly enhances the user experience by resolving over 85% of URLs locally, resulting in an average effective latency of just 18 ms—a drastic improvement over the 380–520 ms range observed in traditional cloud-only architectures. Furthermore, the hybrid design ensures that the majority of user browsing data remains on the local device, successfully addressing the privacy-utility trade-off inherent in modern security tools.

References

- [1] S. Chanti and T. Chithralekha, "A literature review on classification of phishing attacks," *Int. J. Adv. Technol. Eng. Explor.*, vol. 9, no. 89, pp. 446–476, 2022.
- [2] Y. Yao et al., "The psychological manipulation of phishing emails: A cognitive bias approach," *Comput. Mater. Contin.*, vol. 85, no. 3, pp. 4753–4776, 2025.
- [3] I. Skula and M. Kvet, "URL and domain obfuscation techniques - Prevalence and trends observed on phishing data," in *Proc. 2024 IEEE 22nd World Symp. Appl. Mach. Intell. Informat. (SAMI)*, 2024, pp. 283–290.
- [4] T. Wood, V. Basto-Fernandes, E. A. Boiten, and I. Yevseyeva, "Systematic literature review: Anti-phishing defences and their application to before-the-click phishing email detection," 2022, arXiv:2204.13054.
- [5] A. Rahdari et al., "A survey on privacy and security in distributed cloud computing: Exploring federated learning and beyond," *IEEE Open J. Commun. Soc.*, vol. 6, pp. 3710–3744, 2025.
- [6] A. Oest et al., "PhishTime: Continuous longitudinal measurement of the effectiveness of anti-phishing blacklists," in *Proc. 20th USENIX Secur. Symp.*, 2020, pp. 379–396.
- [7] J. Zhou, H. Cui, X. Li, W. Yang, and X. Wu, "A novel phishing website detection model based on LightGBM and domain name features," *Symmetry*, vol. 15, no. 1, p. 180, 2023.
- [8] P. Maneriker et al., "URLTran: Improving phishing URL detection using transformers," 2021, arXiv:2106.05256.
- [9] R. Ferdaws and N. Majd, "Phishing URL detection using machine learning and deep learning," in *Proc. 2024 IEEE World AI IoT Congr. (AIIoT)*, 2024, pp. 485–490.

- [10] A. Yahya, R. B. Ahmad, and Y. M. Yacob, "Lexical based method for phishing URLs detection," in Proc. 2nd Int. Conf. Electron. Design, Comput. Netw. Commun. (IEDC), vol. 88, 2016, pp. 585–592.
- [11] A. Hajizada and S. Jahan, "Feature selections for phishing URLs detection using combination of multiple feature selection methods," in Proc. 2023 15th Int. Conf. Mach. Learn. Comput. (ICMLC), 2023, pp. 444–450.
- [12] M. Khatun et al., "An approach to detect phishing websites with features selection method and ensemble learning," Int. J. Adv. Comput. Sci. Appl., vol. 13, no. 8, pp. 744–751, 2022.
- [13] D. Dr and S. Kuraku, "Phishing website URL's detection using NLP and machine learning techniques," J. Artif. Intell., vol. 5, no. 3, pp. 129–141, 2023.
- [14] H. Ghalechyan et al., "Phishing URL detection with neural networks: An empirical study," Sci. Rep., vol. 14, no. 1, Art. no. 23812, 2024.
- [15] A. Vennela et al., "Intelligent cybersecurity systems for phishing attack detection - An overview," Comput. Electr. Eng., vol. 130, Art. no. 110829, 2026.
- [16] A. Aljofey, S. A. Bello, and C. Xu, "Comprehensive phishing detection: A multi-channel approach with variants TCN fusion leveraging URL and HTML features," J. Netw. Comput. Appl., vol. 238, Art. no. 104170, 2025.
- [17] T. Doshi et al., "PhishHunter-XLD: An ensemble approach integrating machine learning and deep learning for phishing URL classification," Franklin Open, vol. 12, Art. no. 100123, 2025.
- [18] J. D. Duarte et al., "Machine Learning for Early Detection of Phishing URLs in Parked Domains: An Approach Applied to a Financial Institution," IEEE Access, vol. 13, pp. 145736–145750, 2025.
- [19] M. Hosseinzadeh et al., "Improving phishing email detection performance through deep learning with adaptive optimization," Sci. Rep., vol. 15, Art. no. 12540, 2025.
- [20] P. Wanda et al., "DeepPhishing: stepping up phishing detection using generative round network," Int. J. Inf. Technol., 2025. doi: 10.1007/s41870-025-02703-w.
- [21] F. Rashid et al., "Phishing URL detection generalisation using unsupervised domain adaptation," Comput. Netw., vol. 245, Art. no. 110398, 2024.
- [22] T. Bilot, G. Geis, and B. Hammi, "PhishGNN: A phishing website detection framework using graph neural networks," in Proc. 8th Int. Conf. Inf. Syst. Secur. Privacy (ICISSP), 2022, pp. 209–218.
- [23] Z. Liu et al., "Phishing detection on Ethereum via graph neural architecture search of transaction subgraph," IEEE Trans. Comput. Soc. Syst., vol. 12, no. 5, pp. 3513–3524, 2025.
- [24] C. Lee, B. Kim, and H. Kim, "The silence of the phishers: Early-stage voice phishing detection with runtime permission requests," Comput. Secur., vol. 152, Art. no. 104364, 2025.
- [25] Y. Tian et al., "PRIME: A Phishing Detection Framework With Quantitative and Fuzzy-Based Dual Validation," IEEE Trans. Knowl. Data Eng., vol. 37, no. 11, pp. 6597–6609, 2025.