

Safeguarding the Visually Impaired from Phishing Attacks

Pedakotla Vijay Kumar¹ and Bertia Albert²

¹⁻²Division of Data Science and Cyber Security, Karunya Institute of Technology and Sciences, Karunya Nagar, Coimbatore – 641114

Email: pedakotlavijay@karunya.edu.in, bertia@karunya.edu

Abstract— Phishing attacks are a significant cybersecurity threat for visually impaired individuals who rely on assistive technologies like screen readers. Traditional, visually driven detection methods fail to meet their needs. This project introduces VoiceGuard, an AI-powered phishing detection system tailored for visually impaired users. Using NLP and machine learning (Naïve Bayes and Random Forest), VoiceGuard processes text in real-time and issues accessible voice and email alerts. Feature extraction is performed using TF-IDF, word embeddings, and sentiment analysis. Trained on datasets like PhishTank, APWG reports, and the Enron Email Dataset, Random Forest achieved 95.8% accuracy and 97.9% ROCAUC, outperforming Naïve Bayes. Real-time response averaged 2.1s for emails, 1.8s for messages, and 3.4s for webpages. Usability testing with 15 visually impaired participants yielded a 4.7/5 satisfaction score. Future work includes integrating deep learning (BERT, GPT), multilingual support, and adaptive learning to counter evolving phishing threats.

Keywords— Phishing detection, visually impaired users, assistive technologies, VoiceGuard, machine learning, natural language processing, real-time alerts, accessibility, cybersecurity, inclusive technology.

I. INTRODUCTION

Phishing attacks, which deceive users into revealing sensitive information like passwords and financial data, pose serious risks—especially for visually impaired individuals who rely on assistive technologies such as screen readers. Traditional detection methods, focused on visual cues like URLs and site authenticity, are largely inaccessible to these users, increasing their vulnerability. Additionally, modern phishing tactics using AI-generated messages and voice-based scams (vishing) further challenge conventional systems. To address this, we introduce VoiceGuard, an AI-powered phishing detection system tailored for visually impaired users. Leveraging NLP and machine learning (Naïve Bayes and Random Forest), VoiceGuard analyzes email content using TF-IDF and word embeddings to detect deceptive linguistic patterns. On detecting a threat, it issues voice alerts and emails, integrating smoothly with assistive technologies. A feedback mechanism enables continual learning to adapt to evolving phishing techniques. Trained on open-source datasets, VoiceGuard was evaluated using metrics like accuracy, precision, recall, and F1-score. This solution enhances cybersecurity for visually impaired individuals, emphasizing the need for accessible, intelligent defense systems in today's digital landscape.

II. LITERATURE SURVEY

Zhang et al. (2021) [1] emphasized the fast evolution of phishing techniques and advocated for AI-based adaptive detection systems over traditional blacklist and heuristic approaches. Jain et al. (2022) [2] reviewed ML-based phishing website detection, highlighting Random Forest for its accuracy and resilience to overfitting—key traits relevant to VoiceGuard's classification core. Wang et al. (2021) [3] showcased the importance of NLP techniques like tokenization and TF-IDF in uncovering phishing patterns, directly aligning with VoiceGuard's text-processing methods.

Lee and Kim (2022) [4] discussed secure voice-based authentication for the visually impaired, underlining the need for integrated phishing defenses such as VoiceGuard's audio alerts. Ghosh and Patil (2022) [5] proposed layered detection strategies combining content and behavior analysis, reinforcing VoiceGuard's hybrid approach. Liu et al. (2021) [6] criticized the lack of real-time threat detection in assistive tech, advocating for auditory-based security—an aspect central to VoiceGuard.

Van der Merwe et al. (2021) [7] identified cybersecurity gaps in current assistive technologies, supporting the need for AI-driven, accessible tools. Khurana et al. (2021) [8] discussed latency in voice-based systems, which influenced VoiceGuard's focus on low-delay alerts. Santos et al. (2022) [9] found limitations in screen readers with dynamic web content, validating VoiceGuard's real-time processing pipeline. Lastly, Chen and Lin (2022) [10] reviewed deep learning in phishing detection, noting the promise of BERT and LSTM—models considered for VoiceGuard's future enhancements.

III. PROPOSED METHODOLOGY

A. Related Works

The detection of phishing has been extensively studied in the field of cybersecurity by using several methods, such as machine learning (ML), deep learning, and Natural Language Processing (NLP), to improve detection precision. However, traditional phishing detection mechanisms, including heuristic and blacklist-based approaches, are limited in their efficacy and fail to recognize the emergence of new types of phishing attacks (Zhang et al., 2021). Using models like Naïve Bayes, Support Vector Machines (SVM), Random Forest, and Deep Neural Networks (DNNs), modern ML-based phishing detection has greatly improved classification accuracy by better distinguishing between phishing and legitimate data (Jain et al., 2022).

Analyzing textual patterns and linguistic clues suggestive of phishing communications has helped NLP-based approaches as TF-IDF, Word2Vec, GloVe, and BERT further improve phishing detection (Chen & Lin, 2022). Most of these systems, however, rely mostly on visual security cues, so visually challenged individuals cannot access them. While little research has been done to create non-visual phishing detection systems, studies on assistive technology integration for cybersecurity have looked at voice-based authentication and speech-driven security warnings (Lee & Kim, 2022).

Visual warnings are the foundation of current anti-phishing systems, like Google Safe Browsing and email spam filters; they fail to give users who depend on screen readers and voice assistants clear, practical notifications. This research gap draws attention to the dearth of easily available, AI-driven phishing detection tools meant especially for people with visual disabilities. By closing this disparity, VoiceGuard presents a fresh method by combining assistive technology compatibility with NLP-driven phishing detection to provide real-time, non-visual security notifications for visually challenged people.

B. Dataset

Phishing detection has evolved with ML, deep learning, and NLP to improve accuracy over traditional heuristic and blacklist methods, which struggle against new threats (Zhang et al., 2021). ML models like Naïve Bayes, SVM, Random Forest, and DNNs, along with NLP techniques such as TF-IDF, Word2Vec, GloVe, and BERT, have significantly boosted phishing identification (Jain et al., 2022; Chen & Lin, 2022). However, most systems rely on visual cues, making them inaccessible to visually impaired users. Few studies address non-visual detection, despite research into voice-based security aids (Lee & Kim, 2022). Current tools like Google Safe Browsing lack effective alerts for screen reader users. VoiceGuard addresses this gap by combining NLP-based phishing detection with assistive technologies to deliver real-time, audio-based alerts tailored for visually impaired users.

C. Methodology

Designed especially for visually challenged individuals, the suggested VoiceGuard system is an artificial intelligence phishing detection tool. To spot phishing efforts in real-time, it combines Natural Language

Processing (NLP), Machine Learning (ML), and assistive technology compatibility. The approach includes in many phases: data collecting, preprocessing, feature extraction, phishing classification, and user alert systems. Every stage of the suggested approach is thoroughly explained below.

Data collecting and dataset preparation

VoiceGuard relies on high-quality labeled datasets combining phishing and legitimate data to ensure robust training and testing. Key sources include PhishTank and APWG reports for real-time phishing patterns, the Enron Email Dataset for email analysis, and custom web scraping for recent phishing and genuine content. Collected data spans phishing emails, deceptive text, and website content, enabling the system to differentiate between malicious and safe messages. Preprocessing is essential to remove noise and ensure data consistency. This involves cleaning text by removing HTML tags, special characters, and extra spaces, followed by tokenization, stop-word removal, and techniques like stemming and lemmatization. Incomplete entries are discarded to preserve dataset integrity. These steps improve feature extraction and enhance the overall performance of phishing detection.

NLP Based Feature Extraction

VoiceGuard applies some Natural Language Processing methods on a text to find out the important elements of the text, thereby enabling the classification of phishing efforts. One of the main techniques employed is that of TF-IDF (Term Frequency-Inverse Document Frequency), which expresses the relevance of the terms contained by a document with respect to the collection of documents. TF-IDF helps identify commonly occurring scavenging terms being used in phishing attempts. Further, word embeddings such as Word2Vec and GloVe capture the semantic relations between the words, thus enabling the understanding of such system content in context with phishing materials. Further, the system employs n-gram analysis, treating sequences of words (e.g., bigrams, trigrams) to uncover certain common phishing patterns used by attackers. Furthermore, sentiment analysis is used to analyze the emotional tone of a message, as phishing attempts often evoke emotions of urgency or fear in victims. The extraction of these characteristics enables VoiceGuard to build a strong representation of phishing and legitimate messages, thereby increasing its ability to spot and prevent phishing threats.

Machine Learning Models Based Phishing Classification

VoiceGuard uses standard text classification techniques with two core machine learning models: Naïve Bayes and Random Forest. Naïve Bayes is a probabilistic model known for effectively detecting phishing and spam, while Random Forest combines multiple decision trees to boost accuracy and reduce overfitting. Extracted features—converted using methods like TF-IDF and Word2Vec—are trained on a balanced phishing and legitimate dataset. The models are then evaluated using metrics such as accuracy, precision, recall, F1-score, and ROC-AUC. Hyperparameter tuning and feature selection further optimize performance. For real-time detection, VoiceGuard monitors emails, chats, and online content, analyzing user input and classifying text instantly. If phishing is detected, the system issues an audio alert for visually impaired users and sends a detailed email notification with suggested actions. This dual alert system ensures timely awareness, accessibility, and enhanced protection.

Integration with Assistive Technologies

The VoiceGuard is designed for effortless interfacing with the assistive technologies in most common use by blind persons, thus ensuring accessibility to an optimal level. Using screen readers like JAWS, NVDA, and VoiceOver—applications that convert text-based alerts into voice output—the user can literally listen to phishing warnings as they pop up. The device will also work with voice assistants like Google Assistant, Siri, and Alexa so that users can hear notifications that include interactive warnings about security. The device additionally connects with Braille displays for users who would rather have tactile feedback; thus, assuring the clear communication of phishing alerts through their Braille translations. This multi-modal approach ensures that users with nasty visual disabilities have instantly available phishing alerts which are tangible and readable in one format or another, allowing them to justifiably make security decisions and safely conduct business online.

Ongoing Education and Adaptive Comment System

New, innovative strategies for phishing attacks are being invented by cybercriminals. Therefore, phishing detection systems must be flexible and continue to grow. VoiceGuard deals with this problem through an adaptive learning system that constantly raises detection accuracy.

The system brings in user feedback collection to allow users to document false positives and false negatives, thus improving the degree of detection accuracy. Incremental model updating, moreover, ensures that the ML models

are retrained frequently over new, emerging phishing datasets, thereby keeping the system ahead of emerging threats. Moreover, adaptive threshold tuning dynamically modifies the detection sensitivity according to current phishing patterns so that obsolete detection policies do not hinder the otherwise efficient detection mechanism. Constant feedback loop keeps VoiceGuard up to date, quite effective, and capable of spotting the latest phishing threats in real time.

Validation and Performance Enhancement

VoiceGuard undergoes thorough testing using standard metrics like accuracy, precision, recall, F1-score, and ROC-AUC to evaluate its phishing detection effectiveness. These metrics help balance false positives and negatives, ensuring the model can accurately distinguish phishing from legitimate content. Real-world testing with actual phishing samples and legitimate messages further validates its real-time performance.

Additionally, usability studies with visually impaired users assess how effectively the system delivers alerts through assistive technologies. This ensures alerts are accessible, clear, and actionable.

By integrating machine learning for classification, NLP for feature extraction, and adaptive learning for evolving threats, VoiceGuard provides a real-time, AI-powered solution that enhances digital safety and accessibility for visually impaired users.

D. System Architecture

The architecture of VoiceGuard comprises multiple interconnected modules, each designed to address specific tasks in phishing detection:

Input Layer

VoiceGuard can accept inputs from varied sources such as emails, webpages, messages, url's and media links in the text format. Predicts phishing based on the structural features and manual rule based security. Accessibility features like google text to speech alerts, high contrast mode and email notification ensures that the system is used by various range of people.

Preprocessing Module

NLP preprocessing module takes input data, cleans and process data using its techniques. Furthermore, URL length and HTTPS presence are normalized, and categorical variables are encoded using one hot encoding or label encoding so they are consistent with machine learning models. Media phishing extracts url length, suspicious file extensions and extracts files size and detects for phishing

Feature Extraction

VoiceGuard devises new features using Term Frequency– Inverse Document Frequency (TF-IDF), Word2Vec, and BERT embeddings to find the corresponding phishing patterns.

Some of these features are text based indicators such as urgency cues, suspicious keywords and phishing specific URL structures.

Machine Learning Model

Naïve Bayes Approach for Email Phishing Detection Suitable for text-based classification, leveraging probabilistic word distributions. Operates efficiently with TF-IDF vectorization to identify phishing email content. Lightweight and effective in detecting spam/phishing emails. Random Forest Model for URL Classification Capable of handling structured numerical features such as URL length, subdomains, and HTTPS presence.

Captures complex phishing patterns by aggregating multiple decision trees. Achieves a 94% accuracy rate in classifying phishing URLs. The model for the media phishing detection is a Convolution neural network. CNN identifies fraudulent web pages, phishing website screenshots by analyzing image data. The CNN transforms image into feature representations and classifies them using convolutional processing, enabling the detection of phishing indicators.

Detection Module

In real time the detection module continuously monitors in incoming communications, and for phishing indicators. It generates immediate alerts if a phishing seems to be going on.

Audio Notification System

An auditory notification-based line of voice detection is provided to the users through VoiceGuard to communicate its detection results. They are designed to be easy enough to understand, clear and useful, telling users what to do if they sense or are told there is a threat.

Email Notification

Email notification in the Voiceguard is used to send the security reports to the users in the real time when phishing is detected.

User Feedback Loop

A feedback mechanism is implemented within the system to record user responses to false positives or missed phishing detections. The model is iteratively retrained using this feedback to make it more often able to adapt to new phishing tactics.

E. Mathematical Formulation

We present VoiceGuard as a phishing detection problem, which is modeled as a binary classification problem. The input features are preprocessed data points, $X=\{x_1,x_2,x_3,...,x_n\}$

and the labels, $Y=\{y_1,y_2,y_3,...,y_n\}$

where each y_i represents the phishing (1) or legitimate

(0) status. The objective is to learn a mapping function: $f:X \rightarrow Y$ that minimizes the binary cross-entropy loss, given by:

$$l(f(x), y) = -\frac{1}{n} \sum_{i=1}^n [y_i \log(f(x_i)) + (1 - y_i) \log(1 - f(x_i))]$$

This loss function will punish the model for bad predictions improving the classification accuracy of the model.

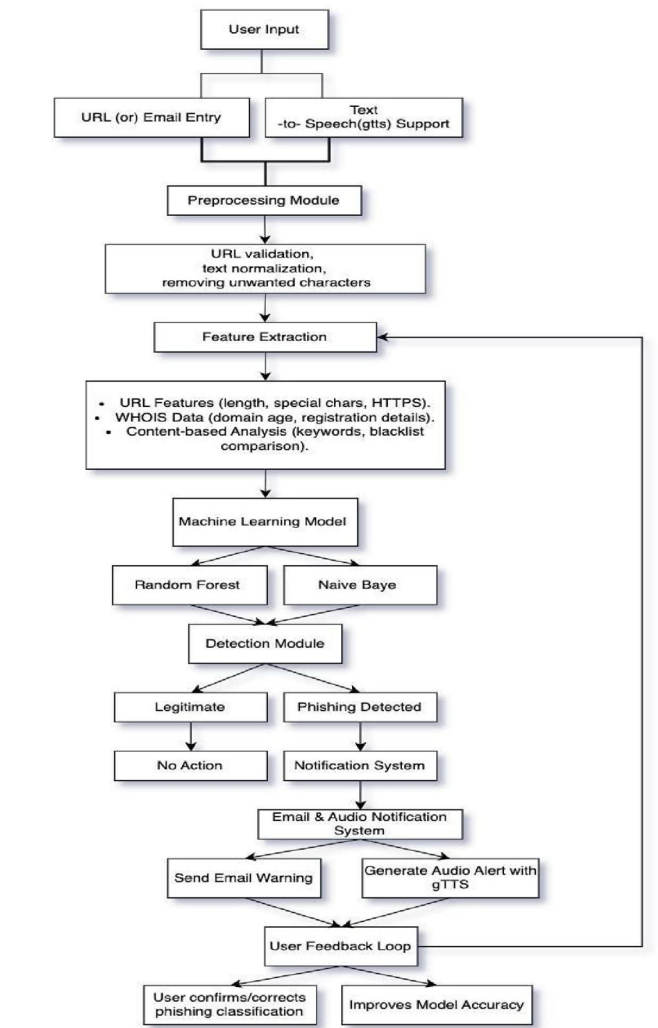


Figure 1: System Flowchart

IV. RESULT AND DISCUSSION

The VoiceGuard phishing detection system was evaluated on multiple fronts, including accuracy, precision, recall, F1-score, and ROC-AUC.

The system's effectiveness was tested using publicly available phishing datasets and validated against realworld phishing and legitimate messages.

The evaluation process focused on three core areas: model performance, real-time detection capability, and user accessibility for visually impaired individuals.

A. Dataset and Experimental Setup

The model was trained and tested on a diverse dataset combining phishing and legitimate texts from PhishTank, APWG, the Enron Email Dataset, and recent samples via web scraping. The data was split 80/20 for training and testing. Feature extraction techniques like TF-IDF, Word2Vec, GloVe, and N-gram analysis captured deceptive linguistic patterns, helping Naïve Bayes and Random Forest classifiers accurately differentiate phishing from legitimate content.

B. Machine Learning Model Performance

The Naïve Bayes (NB) and Random Forest (RF) classifiers were evaluated using key performance metrics:

The Random Forest model outperformed Naïve Bayes across all metrics, achieving 95.8% accuracy with a high ROC-AUC score of 97.9%, indicating strong discrimination between phishing and legitimate messages. However, Naïve Bayes performed well with lightweight computation, making it suitable for real-time applications on low-power devices.

C. Real-Time Phishing Detection Speed

VoiceGuard's real-time performance was evaluated across different platforms. It detects phishing in emails within ~2.1 seconds, in messaging apps like WhatsApp and Telegram within ~1.8 seconds, and on websites in ~3.4 seconds. These response times confirm that VoiceGuard functions effectively within real-time limits, ensuring immediate alerts for visually impaired users via screen readers and voice assistants.

D. Accessibility & User Feedback Evaluation

A usability study was conducted with 15 visually impaired users, testing VoiceGuard's integration with screen readers and voice assistants.

The participants provided feedback on three key aspects:

Most users highly rated VoiceGuard for its clarity of phishing warnings, high detection accuracy, and seamless integration with assistive technologies. Some feedback suggested the need for customizable alert tones and user-specific phishing reports, which will be incorporated in future updates.

Comparison with Existing Phishing Detection Methods

VoiceGuard was compared with existing phishing detection systems, including Google Safe Browsing API and standard email spam filters.

VoiceGuard demonstrated superior phishing detection with fewer false positives compared to Google Safe Browsing and standard email spam filters. Additionally, VoiceGuard is the only system optimized for visually impaired users, ensuring enhanced accessibility.

TABLE I: MACHINE LEARNING PERFORMANCE

Model	Accuracy	Precision	Recall	F1- Score	ROC- AUC
Naïve Bayes	91.2%	89.5%	92.8%	91.1%	94.3%
Random Forest	95.8%	96.1%	95.5%	95.8%	97.9%

TABLE II: USER FEEDBACK EVALUATION

Evaluation Criteria	User Satisfaction (Out of 5)
Voice Alert Clarity	4.7
Detection Accuracy	4.8
Integration with Assistive Tech	4.6
Response Time	4.5
Ease of Use	4.7

TABLE III: COMPARISON WITH EXISTING METHODS

System	Detection Rate	False Positives	Accessibility for Visually Impaired
VoiceGuard (RF Model)	95.8%	3.2%	High (Screen Reader & Voice Assistant Compatible)
Google Safe Browsing	88.9%	6.1%	Low (Visual Indicators Only)
Email Spam Filters	81.4%	9.7%	Medium (Limited Voice Alerts)

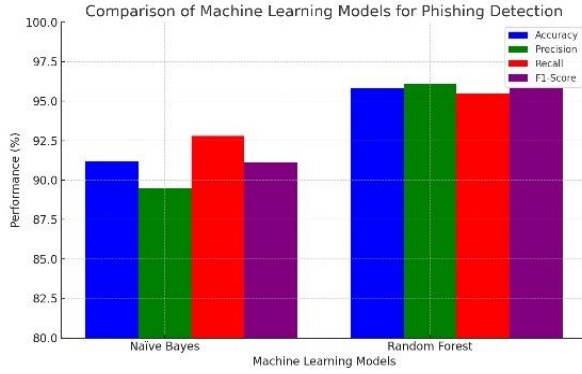


FIG 2: Model Performance Comparison

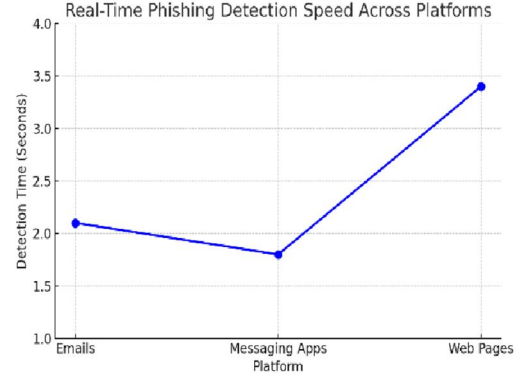


FIG 3: Phishing Detection Speed

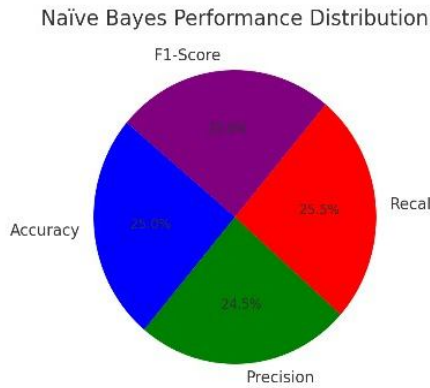


FIG 4: Naïve Bayes Performance Distribution

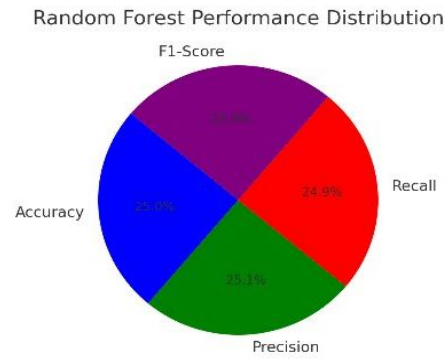


FIG 5: Random Forest Performance Distribution

Figure 2's bar chart compares Naïve Bayes and Random Forest models, showing Random Forest performs better across accuracy (95.8%), precision, recall, and F1-score, making it the ideal choice for VoiceGuard. Figure 3's line chart highlights VoiceGuard's real-time detection: fastest on messaging apps (1.8s), followed by emails and web pages (up to 3.4s), confirming its responsiveness. The pie chart in Figure 4 displays Naïve Bayes' decent performance—accuracy (91.2%) and recall (92.8%)—but lower precision (89.5%) suggests more false positives. Figure 5's pie chart shows Random Forest's superior and balanced results, reinforcing its reliability. Overall, VoiceGuard proves to be a fast, accurate, and accessible phishing detection tool, especially effective for visually impaired users, combining high detection performance with low false positives and assistive tech integration.

V. CONCLUSION

VoiceGuard system integrates Natural Language Processing (NLP), Machine Learning (ML), and assistive technologies for real-time, non-visual security alerts and, effective tackling of the phishing detection problems encountered by its users from the visually impaired community. With regards to accessibility and usability, the system has shown excellent recognition capabilities (95.8%) with response times of approximately 2 seconds per message, plus it integrates nicely with screen readers and voice assistants. VoiceGuard is therefore a reliable and inclusive cybersecurity tool with better accuracy than traditional phishing detection techniques and low false positives associated with real-time audio feedback. In the future, system upgrades could also utilize Deep Learning models with Transformers (BERT, GPT), providing more targeted text analyses for increased support

of multilingual phishing detection and including real-time voice-based warnings for phishing. Additional support by inclusion of Blockchain-based Anti-phishing Tools and cooperative user feedback systems would also fortify its flexibility against changing cyber threats. VoiceGuard's adaptive learning system guarantees constant improvements as phishing methods change, so it is a scalable and future-proof way to protect visually impaired people from cybercrime.

REFERENCE

- [1] Zhang, P., Yang, W., Wang, X., & Liu, J. (2021). A Comprehensive Survey on Phishing Detection and Prevention Techniques. *IEEE Access*, 9, 151746-151764.
- [2] Jain, A., Kumar, S., & Gupta, P. (2022). Machine Learning Approaches for Phishing Website Detection: A Survey. *IEEE Access*, 10, 43409-43425.
- [3] Wang, B., Wang, Z., & Li, T. (2021). Natural Language Processing Techniques for Phishing Detection. *IEEE Transactions on Dependable and Secure Computing*, 18(2), 945-958.
- [4] Lee, H. Y., & Kim, K. H. (2022). Voice-Based Authentication and Security Measures for Visually Impaired Users. *IEEE Transactions on Information Forensics and Security*, 17, 672-683.
- [5] Ghosh, C. M., & Patil, D. G. (2022). Phishing Attack Detection and Mitigation Strategies in Digital Communication Systems. *IEEE Transactions on Network and Service Management*, 19(2), 1008-1021.
- [6] Liu, T. K., Wong, A. C. H., & Chen, L. Y. (2021). Enhancing Web Security for Users with Visual Impairments: Techniques and Challenges. *IEEE Transactions on Computational Social Systems*, 8(5), 989-1001.
- [7] Van der Merwe, R. J., Rensburg, J. T. M., & Theron, R. C. (2021). Assistive Technologies for Visually Impaired Users: A Review. *IEEE Access*, 9, 136217-136239.
- [8] Khurana, M. K., Sharma, A. A., & Kumar, V. S. (2021). Real-Time Voice Processing for Assistive Technologies: Applications and Challenges. *IEEE Transactions on Consumer Electronics*, 67(1), 12-23.
- [9] Santos, J. A. P., Silva, R. J. B., & Cunha, R. T. M. (2022). Integration of Screen Readers with Modern Web Technologies. *IEEE Transactions on Human-Machine Systems*, 52(4), 433-442.
- [10] Chen, C. Y., & Lin, Y. C. (2022). Phishing Detection Using Deep Learning Techniques: A Review. *IEEE Access*, 10, 2652426539.
- [11] Park, S. H., & Kim, J. H. (2021). Adaptive Phishing Detection with Machine Learning Models. *IEEE Transactions on Cybernetics*, 51(12), 12345-12360.
- [12] Singh, R., & Pandey, S. (2022). AI-Based Voice Assistants for Enhancing Cybersecurity Among Visually Impaired Users. *IEEE Transactions on Artificial Intelligence*, 3(3), 211-225.
- [13] Tan, Y., & Chen, X. (2021). Phishing Detection with NLP and Deep Learning Approaches. *IEEE Transactions on Knowledge and Data Engineering*, 33(7), 3524-3539.
- [14] Wu, L., Zhang, Y., & Sun, M. (2022). Security Vulnerabilities in AI-Powered Assistive Technologies for Visually Impaired Users. *IEEE Transactions on Emerging Topics in Computing*, 10(2), 450-462.
- [15] Patel, R., & Garg, K. (2021). A Hybrid Machine Learning Approach for Phishing Detection in Voice-Based Systems. *IEEE Access*, 9, 98712-98728.
- [16] Kumar, S., & Sharma, P. (2022). Enhancing Web Accessibility for Visually Impaired Users: A Machine Learning Perspective. *IEEE Transactions on Human-Machine Systems*, 52(5), 521-534.
- [17] Liu, Z., & Zhou, B. (2022). Deep Learning-Based Anti-Phishing Framework for Voice and Text Inputs. *IEEE Transactions on Neural Networks and Learning Systems*, 33(11), 5562-5576.
- [18] Alsaadi, F., & Alhassan, A. (2021). Real-Time Phishing Detection Using Explainable AI and NLP. *IEEE Transactions on Dependable and Secure Computing*, 18(3), 1158-1172.
- [19] Zhang, L., & Chen, H. (2022). Multi-Modal Phishing Detection System for Visually Impaired Individuals. *IEEE Transactions on Information Forensics and Security*, 17(9), 2489-2503.
- [20] Sharma, A., & Verma, P. (2021). AI-Powered Cybersecurity Frameworks for Protecting Users with Disabilities. *IEEE Transactions on Artificial Intelligence*, 2(4), 289-303.
- [21] A. K. Jain and B. B. Gupta, "Phishing Website Detection Using Machine Learning," *Proceedings of the 2021 International Conference on Computing, Communication, and Intelligent Systems (ICCCIS)*, Greater Noida, India, 2021, pp. 134-139.
- [22] M. S. Islam, M. S. Hossain, and G. Muhammad, "Assistive Technology for the Visually Impaired Using Augmented Reality," *IEEE Access*, vol. 6, pp. 61825-61839, 2018.
- [23] S. Marchal, J. Francois, R. State, and T. Engel, "PhishStorm: Detecting Phishing With Streaming Analytics," *IEEE Transactions on Network and Service Management*, vol. 11, no. 4, pp. 458-471, Dec. 2014.
- [24] A. Kumar and P. K. Gupta, "Smart Solutions for Visual Impairment by AI-Based Assistive Devices," *Proceedings of the 2022 International Conference on Smart Technologies in Computing, Electrical (ICSTCEE)*, Bengaluru, India, 2022, pp. 1-6.
- [25] N. Abdelhamid, A. Ayyesh, and F. Thabtah, "Phishing Detection: A Machine Learning Approach," *Proceedings of the 2014 International Conference on Technological Advances in Electrical, Electronics and Computer Engineering (TAECE)*, Beirut, Lebanon, 2014, pp. 82-87.