



University of Tehran Press

Interdisciplinary Journal of Management Studies (IJMS)

Online ISSN: 2981-0795

Home Page: <https://ijms.ut.ac.ir>

Sharable Device-Aware Phishing Attack Detection in Cloud Environments Using VAE with Clustering Mechanism and SVM

Venkat Garikipati^{1*} | Charles Ubagaram² | Narsing Rao Dyavani³ | Bhagath Singh Jayaprakasam⁴ | Rohith Reddy Mandala⁵ | Gabriel Ayodeji Ogunmola⁶

1. Corresponding Author, Innosoft, Windsor Mill, Maryland, USA. Email: venkatgarikipati@ieee.org
2. Tata Consultancy Services, Milford, Ohio, USA. Email: charlesubagaram@ieee.org
3. Uber Technologies Inc, San Francisco, California, USA. Email: narsingraodyavani@ieee.org
4. Cognizant Technology Solution, College Station, Texas, USA. Email: bhagathsinghjayaprakasam@ieee.org
5. Tekzone Systems Inc, Rancho Cordova, California, USA. Email: rohithreddymandala@ieee.org
6. Associate Professor, Department of Economic Theory, Faculty of Economics, Tashkent State University of Economics, Uzbekistan. Email: drGabrielao@yahoo.com

ARTICLE INFO

Article type:
Research Article

Article History:
Received: 28 March 2025
Revised: 04 November 2025
Accepted: 10 November 2025
Published Online: 03 September 2026

Keywords:
Phishing attack detection,
Sharable devices,
Device-aware detection,
Variational autoencoders (VAE),
Clustering mechanism.

ABSTRACT

Phishing attacks pose a significant risk to cybersecurity and are especially troublesome in cloud-based settings as the risk is heightened due to shared and multi-device access. To counteract this, the current paper presents a Sharable Device-Aware Phishing Attack Detection system that integrates Variational Autoencoder (VAE) with a Clustering Mechanism and Support Vector Machine (SVM) to make the detection of phishing attacks more effective. The VAE is employed for feature extraction and to facilitate the unsupervised learning of phishing behavior. This approach enables the clustering mechanism to group threats effectively, after which the SVM model is applied to accurately classify the phishing cases. The presented model is evaluated on a dataset of phishing behaviors collected from cloud-based IoT environments, demonstrating strong performance in terms of detection accuracy, recall, and F1-score. Results show 99.5% accuracy, 99% precision, 98.5% recall, and 99.95% F1-score, outperforming the current phishing detection algorithms. The combination of device-aware learning with more sophisticated machine learning concepts offers an effective, scalable, and flexible phishing detection algorithm for cloud security.

Cite this article: Garikipati, V.; Ubagaram, Ch.; Dyavani, N.; Singh Jayaprakasam, B.; Mandala, R. & Ayodeji Ogunmola, G. (2026). Sharable Device-Aware Phishing Attack Detection in Cloud Environments Using VAE with Clustering Mechanism and SVM. *Interdisciplinary Journal of Management Studies (IJMS)*, 19 (4), 995-1015. <http://doi.org/10.22059/ijms.2025.392673.677516>



© Authors retain the copyright and full publishing rights.
DOI: <http://doi.org/10.22059/ijms.2025.392673.677516>

Publisher: University of Tehran Press.

1. Introduction

Cloud computing has revolutionized data storage and access, offering businesses and individuals scalable and flexible computing power (Vichare et al., 2024). Security has consequently become a primary concern for organizations that rely on cloud environments as the foundation of their operations (Mohammed et al., 2024). The proliferation of IoT devices has further complicated cloud security, as diverse connected devices introduce new vulnerabilities that malicious actors can exploit (Torres et al., 2021). Among the various cyber threats confronting cloud environments, phishing attacks remain among the most widespread and dangerous, targeting both users and organizations by deceiving them into disclosing login credentials, financial information, and personal data (Wilson et al., 2024).

Phishing attacks involve misleading e-mails or websites that pretend to be legitimate and that try to fraudulently extract information from victims (James & Rabbi, 2023). Phishing is a major threat in cloud environments since attack operatives can take advantage of users' trust and resources provided by the cloud for more advanced, targeted operations (Hoefer et al., 2024). Even though cloud infrastructure is often secured by the cloud providers, end-users—especially those accessing services on shared or vulnerable devices—remain the weakest link in the security chain (Igwenagu et al., 2024).

Several causes account for this increasingly high phishing threat level in cloud environments. Services such as cloud computing have extended their widespread nature, allowing the integration of several user devices, many of which are sharable and may not be designed to give adequate security assurance (Akinade et al., 2024). Moreover, the ever-evolving and wide-ranging environment of cloud apps helps phishing campaigns to be more accurate in targeting users, employing social-engineering techniques, and eventually, getting past the conventional security measures (Odeleye et al., 2023). The wide range of cloud communication channels, including email, combined with the evolving tactics of attackers, presents an additional challenge. This makes phishing attacks exceedingly difficult to detect through human analysis alone, thereby necessitating sophisticated automated detection systems (Kasri et al., 2025).

There are more sophisticated attacks techniques that increase phishing attacks. Among them, spear-phishing targets individuals according to their information collected from social media or any other compromised source, making these messages appear legitimate to the user (Schmitt & Flechais, 2024). The anonymity of the cloud environment enables hackers to conceal activities and attacks, thereby posing significant detection and prevention challenges worldwide (Singh et al., 2024). This typically has the advantage of traditional methods used for detecting phishing through the use of a signature-based approach that uses an email's content and metadata for detection (Kyaw et al., 2024). A limitation, however, in signature-based approaches is their inability to detect new, advanced, or unique campaigns which have not yet been cataloged (King & Huang, 2023). Additionally, such approaches create false positives and bring high noise ratios, decreasing the trust of a security system for its users. It usually overlooks contextual issues, such as the type of attack that user behavior may indicate, or the devices involved in accessing the cloud system (Sicari et al., 2022). Consequently, even within environments comprising shared and unmanaged devices, certain attacks may evade detection entirely (Niwa, 2024). Furthermore, old systems cannot make rapid adjustments regarding changing phishing styles since they totally depend on historical data or prespecified rules, which creates considerable risk in clouds because phishing changes continuously (Chen et al., 2025). Now, with sophistication regarding phishing types, traditional methodology loses the feasibility of effectively eliminating and detecting sophisticated phishing attacks (Owa & Adewole, 2025).

To overcome these limitations, this paper proposes a novel scheme for the detection of phishing attacks in cloud environments that makes use of the VAE combined with clustering. The VAE model facilitates the unsupervised learning of data patterns, enabling the detection of previously unseen phishing attacks. By using clustering techniques, the model combines similar attack patterns, thereby enhancing its ability to identify new and mutating threats. This is dynamic enough to update based on evolving phishing methods. Accordingly, the proposed method bases its detection on the characteristics of users and their devices, enabling more contextual and comprehensive phishing detection. By learning to recognize even subtle differences in data patterns, the method effectively reduces false positives and filters out only genuinely benign activity. VAE with clustering together

ensures the best defense mechanism to detect phishing attacks even in complex and large-scale cloud environments. Scalable and adaptive, it becomes an ideal solution for cloud-based applications and systems requiring real-time threat detection. The primary benefit of the proposed method is high accuracy in phishing attack detection and the reduction of false positives, even in large-scale cloud environments. The system would be able to detect novel phishing attacks that signature-based systems could not detect otherwise. Moreover, the VAE-based model is adaptive to changes in the threat landscape, thereby continuously improving its phishing detection performance as new patterns of attacks arise. Such improved flexibility and precision are vital in protecting sensitive data and ensuring the security of cloud services, particularly from sharable devices that could be compromised. Employing this process, the method will produce a more effective, scalable, and context-aware solution to phishing attack detection that assists in the improvement of cloud security.

1-1. Problem Statement

This paper (Mali et al., 2025) focuses on the increasing sophistication of cyber threats—such as phishing attacks and encrypted network traffic—which enables them to evade detection, and addresses this challenge using machine learning and deep neural networks (DNNs). Traditional methods often fail to identify adversarial phishing attacks with minimal changes or cannot classify encrypted traffic when the approaches used rely on the analysis of unencrypted data or where effective feature selection is not available (Sudar et al., 2024). The rise of advanced phishing attacks, coupled with existing datasets and suboptimal feature selection, hinders precise detection. A further challenge in applying DNNs to web phishing detection lies in the high-dimensional nature of the data and the presence of irrelevant features (Champa et al., 2024). These papers propose exploring intelligent techniques, curated datasets, feature analysis, and employing methods like ANOVA for enhanced accuracy, efficiency, and robustness of phishing and encrypted traffic detection systems beyond the drawbacks of traditional systems (Lestari & Ulina, 2024).

1-2. Objectives

- Gather data from IoT devices related to phishing
- Apply pre-processing techniques, including noise removal and data normalization
- Use VAE with clustering and SVM to detect phishing attacks effectively

2. Literature Review

Daengsi et al. (2022) proposed a mixed-approach training methodology that would enhance cybersecurity awareness and decrease phishing attacks. Certainly, this method depends on human intervention, which cannot be adequate in the rapidly changing world of automated environments such as cloud systems. On the other hand, the current study fills that gap leveraging machine learning techniques, such as Variational Autoencoders (VAE), to detect phishing patterns in real time autonomously without manual input, thereby increasing scalability in cloud environments.

Færøy et al. (2023) investigated the weaknesses in Internet of Things (IoT) devices and made use of automated penetration testing for cybersecurity. Nonetheless, their technique is highly reliant on scenarios that are pre-configured, which might not be able to identify emerging attack patterns. The current research takes advantage of their work by introducing unsupervised learning techniques, such as VAE, to recognize novel, invisible phishing attacks, and to apply clustering for real-time adaptation to new threat situations.

Kondoyanni et al. (2022) concentrated on security vulnerabilities in robotic systems and offered recommendations for their protection through the use of cryptographic algorithms and multi-factor authentication, thus adding security to these systems. One of the drawbacks of this approach is the need for more computational resources which might pose a problem in less powerful environments. On the other hand, this research emphasizes the use of machine learning techniques, particularly SVM and VAE, which provide more scalable, real-time solutions for cloud environments with limited resources.

In Yaacoub et al.'s (2022) study, the authors suggested that they analyzed security vulnerabilities in robotics and proposed measures to prevent such attacks by using and recommending techniques like multi-factor authentication and cryptographic algorithms. In contrast, the current study takes the

opposite approach and applies machine learning algorithms, such as SVM together with VAE, for efficient phishing detection without the burden of implementing multiple security layers. Achaal et al. (2024) highlighted the cybersecurity issues surrounding smart grid technologies and placed special emphasis on the communication networks. They introduced a classification of cyber-attacks; however, their model focused merely on a few types of attacks, thereby, not detecting some new threats. The present research extends prior work by proposing an agile machine learning system that integrates a VAE and clustering, enabling it to adapt to a broader range of evolving phishing threats in smart grid systems.

Paul et al. (2024) studied the security issues of smart grids and developed a machine learning-based method to detect abnormal activities in communication networks. However, their method had a significant drawback—a high rate of false positives—resulting in the generation of unnecessary alerts. This study proposes a novel approach by integrating clustering mechanisms to minimize false positives while simultaneously enhancing the accuracy of real-time phishing detection in changing environments.

In Alissa et al.'s (2022) study, the main objective was to detect botnet attacks in IoT networks via machine learning algorithms. Their methodology relied on specific datasets, which limits its applicability to other types of IoT networks. The present study, however, overcomes this limitation by incorporating a VAE into the unsupervised learning process, enabling the model to identify phishing attacks across diverse IoT environments without being confined to predetermined datasets.

Kumar and Sinha (2021) suggested a zero-day cyber-attack detection model that employed heavy-hitter techniques along with graph-based methods for detection. Nevertheless, their method has to deal with scalability problems in vast IoT settings due to its computational complexity. This work challenges that assumption by applying SVM and VAE, which provide more efficient and scalable solutions for real-time phishing detection in large, dynamic environments.

Azeem et al. (2024) conducted thorough research on the malware detection in Internet of Things gadget using machine learning models, such as Random Forest and K-Nearest Neighbours. The major drawback of their method is the difficulty in dealing with imbalanced datasets, which in turn can lead to a decline in detection accuracy. The current research manages to eliminate this issue by using clustering methods together with SVM to boost detection accuracy and to take care of class imbalances in phishing detection more effectively.

Thapa et al. (2023) investigated phishing email detection using federated learning techniques, which offer the advantage of privacy preservation. However, their method struggles to achieve high accuracy when working with limited data or in scenarios where few devices participate in the federated network. A more powerful technique combining VAE with SVM is proposed in this research, which allows effective phishing detection that is not dependent on the drawbacks of federated learning frameworks.

Several works in the literature (e.g., Madhavaram et al., 2025) have proposed AI-based techniques for detecting malicious email motifs, a form of big data analysis commonly employed for large-scale transactions. A primary disadvantage of these approaches is their high computational cost. Consequently, such models may not be well suited to ultra-low-resource environments, such as IoT devices operating in the cloud. In this regard, VAE and SVM offer a more efficient solution due to their improved balance between computational cost and detection performance.

High detection accuracy outputs are both maintained and elapsed into a correct direction of lower computational resources in cloud-constrained surroundings.

Stoleriu et al. (2023) employed deep neural networks (DNNs) and ensemble methods for phishing detection in large-scale datasets. However, their method is vulnerable to adversarial attacks, which can reduce detection accuracy. This study mitigates this risk by leveraging VAE for robust feature extraction, making the system more resistant to adversarial manipulation and better suited to evolving phishing techniques.

According to Liubchenko and Volkov (2024), malware and phishing threats in cloud computing were certainly the main concerns, and the introduction of a hybrid security framework was an accompanying suggestion. However, the framework's reliance on large-scale data processing may hinder real-time detection in resource-constrained environments. The present study addresses this

limitation by employing SVM and VAE, which provide a lightweight and scalable solution for phishing attack detection without requiring extensive data processing.

In Lokesh et al. (2023), a policy-driven framework was proposed for the detection of phishing in cloud storage services. Nevertheless, the limitations of the system include its slow adaptation to the changing techniques of phishing and the possibility of false positives. The present research increases the capability of phishing detection by the use of VAE and SVM for the purposes of dynamic detection and classification, which results in the substantial reduction of false positives and the improvement of the system's adaptability.

3. Methodology

Figure 1 is a representation of the phishing attacks which takes place using IoT devices. The whole thing begins with gathering the data; then, pre-processing comes in, which is a highly important part of the process, as it consists of cleaning the data, reducing noise, and normalizing the data. The process of feature extraction is subsequently conducted, where it applies behavioral pattern analysis in identifying the phishing patterns.

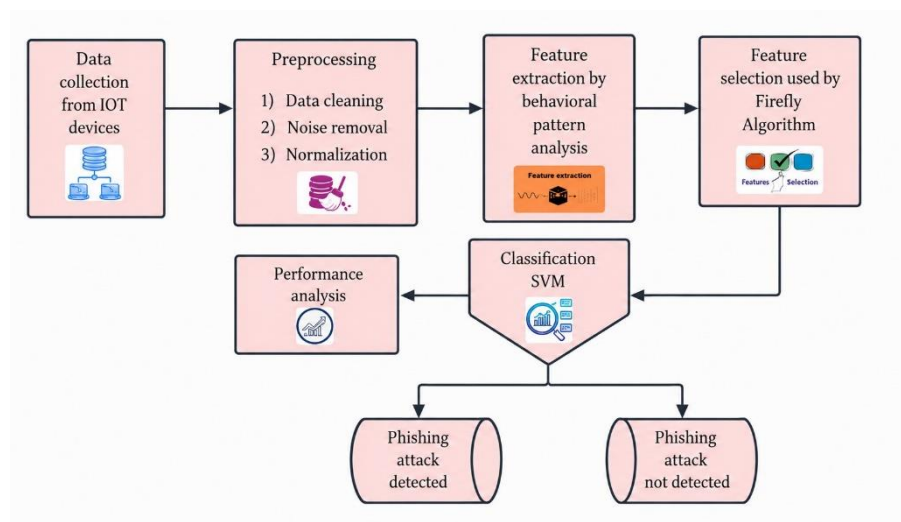


Fig. 1. Architectural Diagram

Next, feature selection is carried out using the Firefly Algorithm to enhance the accuracy of detection by choosing the most appropriate features. Later, classification with a Support Vector Machine (SVM) model is used to classify whether the attack is phishing or not. The output at the final stage indicates whether a phishing attack has been detected. Furthermore, the performance of the system is studied to assess the suitability of the detection method. Noise reduction is an important step in data pre-processing to remove transient disturbances and irrelevant variations from the input. This ensured clear behavioral patterns whose extraction would be further subjected to a feat. Thus, as the VAE and SVM models unreservedly focused on the relevant data, it became easier to ascertain phishing attacks with the increase of true positives and with the lowering of false positives; therefore, the overall classification process is made more reliable.

3-1. Data Collection

Data acquisition from IoT Devices involves collecting real-time information from sensors, smart home appliances, and other connected devices. These devices generate a continuous stream of data over time, capturing user activities, device-to-device interactions, and how individual devices interact with their surrounding environment. Actually, the actual data collected here includes temperature reading, device status, network activities, or direct user commands. Such raw data forms the basis for pattern analysis and anomaly detection as well as for decision-making in various applications whether in security, automation, or predictive maintenance. Data, therefore, must be of high quality and accuracy for effective analysis and the appropriate action that may then follow.

Behavioral and interaction data are collected from network-connected devices, such as smart TVs, routers, and wearables. The streams of information include logs of the device status, access time stamps, IP connection records, user commands for unit control, and flow measurements. Such detailed, time-stamped data provide behavioral insights to detect allocable patterns of phishing, particularly when devices experience unusual patterns of use or access that can be mapped back to compromised endpoints. At a high level, data collection via an IoT setup must be robust to preserve data integrity and to allow for real-time threat analysis. IoT environments create dynamic data streams in constant interaction with users, applications, and other networks. Recording these streams live, whether from smart cameras, sensors, or smart home hubs, allows for real-time behavior profiling.

3-2. Pre-processing

Data preparation involves eliminating every type of errors or inconsistencies within the dataset, including duplicate entries or outliers and missing values. Quality and accuracy will therefore ensure the correct type of data used for analyses or training to build models.

Noise removal is such an important step in reducing the random variations in data that interfere with learning processes. Techniques such as Gaussian smoothing or median filtering are typically applied to remove noise or noisy data points. Scaling of the numerical features normalizes them into a comparable range so that none of the features would dominate other ones in terms of their magnitudes. It may help improve the model's performance by keeping the training process stable and thus predictable.

3-2-1. Data Cleaning

Data cleaning involves the identification of errors, inconsistencies, and inaccuracies in a dataset and its rectification. It involves eliminating missing values and duplicate entries along with correcting any erroneous data point. Irrelevant or incomplete data is removed, and the remaining data is transformed into an appropriate format suitable for subsequent analysis or modeling. Data cleaning is essential for improving data quality, thereby enabling machine learning models to produce accurate and meaningful results.

$$f'(a,b) = \sum_{i=-k}^k \sum_{j=-k}^k f(c,d) \cdot G(c,d) \quad (1)$$

where $f'(a,b)$ denotes smoothed value, $f(c,d)$ is original value at position, and $G(c,d)$ represents Gaussian kernel.

Missing values and incomplete records in phishing-related data can significantly degrade the prediction performance by causing bias, reducing training efficiency, and distorting behavioral patterns. Data gaps could result from technician errors, transmission delays, or simply plain apathy on the part of the users. In cases where these issues are not resolved, the outcome can be a series of predictions that are not very reliable, particularly when the features that are influential are the ones affected, such as, login time, packet size, or device status. The removal of missing values and duplicate records is a fundamental yet essential step in ensuring data quality and integrity before training any machine-learning model. If not properly addressed, missing values can introduce noise during model training, distort feature distributions, and potentially lead algorithms to perform inaccurate computations. Similarly, duplicate records can distort the underlying data patterns, potentially resulting in overfitting and artificially inflating the importance of certain features beyond their actual predictive contribution.

3-2-2. Noise Removal

Noise removal refers to the technique of eliminating nonessential or random changes/distortions that might hinder further processing or hurt model accuracy. It lowers the level of noise that is unwanted or random in the data, hence making it easier for patterns or signals to become visible. Filtering, smoothing, and various algorithms for removing irrelevant data and outliers are among the commonly employed data-pre-processing techniques. The process of noise removal mainly increases the correctness and is also widely adopted since the important data in the machine is the one that is being focused.

$$f'(g,h) = \text{median}\{f(g+i, h+j) | -k \leq i, j \leq k\} \quad (2)$$

where f represents original value and $f(g + i, h + j)$ is the neighbouring values in a square window centered at (g, h) , k and determines the window size.

Gaussian smoothing and median filtering are crucial for removing noise while retaining the meaningful signal patterns present in cybersecurity data. Gaussian smoothing employs a weighted averaging process to reduce high-frequency fluctuations and suppress background noise. Median filtering, in contrast, removes outliers by replacing the central value with the median of the values within a specified neighborhood, making it particularly effective against impulsive and random noise. Together, these techniques enhance the quality of the input data, enabling the VAE to detect subtle anomalies and the SVM to classify phishing behavior with greater precision.

3-2-3. Normalization

Normalization refers to scaling the features or variables to values in the range of $[0, 1]$. Other transformations of data can be done into a consistent scale or distribution by altering the range or the standard deviation of numerical data. This ensures that each input feature equally contributes to the training of the model. There will not be biasing toward features having large values. Thus, normalization enables the model to understand every feature equally. Therefore, convergence speed is higher in training, and the model's performance improves. This expression is presented in equation (3).

$$y_{\text{norm}} = \frac{y}{y_{\text{max}} - y_{\text{min}}} \quad (3)$$

where y is original value, y_{min} represented the minimum value of the feature, and y_{max} is the maximum value of feature.

3-3. Feature Extraction

The behavioral feature extraction of phishing attack detection focuses on discovering and analysing interaction patterns between the devices and the networks, the users, or services in identifying unusual or malicious activities through abnormal behavior. This includes features that are device interaction patterns, including the frequency at which the interactions occur, signalling abnormal behavior in the form of bot activity or phishing link distribution. Patterns of usage, such as attempts to log in during unusual hours, may indicate that compromised systems are being exploited to conduct phishing operations. Other features indicative of malicious network traffic include packet length and traffic volume, as unusually large data transfers or sudden surges in traffic may signal data exfiltration or ongoing phishing campaigns. Furthermore, abnormal access patterns involving known phishing sites, as well as unusually frequent or atypical HTTP requests, may provide additional indicators of phishing activity. The distribution of IP addresses can also contribute to identifying such anomalous patterns and potential phishing attempts.

In contrast, extremely long or short session durations could indicate the execution of automation scripts or some compromised access tokens being exploited. Real-time analysis of these patterns enables the model to promptly identify and generate alerts for potential threats that may escalate into full-scale phishing attacks. By incorporating such time-oriented signals, the detection system provides a proactive response mechanism, thereby improving detection accuracy and reducing the response time for compromised user devices or accounts under attack.

Other time-related features would consist of access times and session lengths where uncharacteristic access times or unusually long connection could highlight possible attacks. A login pattern, where repetitive attempts or a sudden increase of failed logins are visible, indicates brute-force or phishing. Usage of devices statistics, like IP geolocation and device fingerprinting, can define suspect activities. In case the device unexpectedly connects from another geographic location or has a much different fingerprint, it could mean that the attackers have taken the identity of legitimate devices and are performing hijacking and phishing activities. These features allow detection of phishing attacks with behavioral analysis of the devices in a holistic way. Combination of expedient usage profiles with behavioral analysis makes for a more comprehensive scheme for fraud detection. Expedient statistics, including the number of logins, how often the device is switched out, and the locations of access, provide almost ad hoc and quantifiable signs of abnormal behavior. These are then

merged with considerations at the level of session durations, sequences of commands, and flows of network interactions, to yield a layered picture of user intention.

Working with a device fingerprinting system and IP geolocation system provides clues to help detect phishing-related anomalies. A device fingerprint is a unique set of characteristics that partly includes operating system, browser type, installed fonts, and hardware characteristics. A system distinguishes between familiar and unfamiliar devices. When a known user account is accessed using a fingerprint that is significantly different from or unrecognized by the system, this may indicate that the account credentials have been compromised or misused. IP geolocation maps IP addresses to the geographic locations associated with access requests. A sudden login from a distant or unusual location that is inconsistent with the user's previous access patterns may therefore be considered suspicious. When combined with login time, access frequency, and network traffic characteristics, these features can help the model identify anomalous behavioral patterns that may indicate phishing attempts or account takeover.

3-3-1. Traffic Volume

This is the amount of data transferred over a network within any given time. It is one of the most important network activity metrics and can be utilized to identify an anomaly or security issue. This term is presented in equation (4).

$$V = \sum_{i=1}^N \text{DataSize}(i) \quad (4)$$

where N denotes the number of data packets and $\text{DataSize}(i)$ is the size of each packet.

By monitoring such anomalies, the system can identify suspicious activities that might otherwise evade content-based detection mechanisms. This approach is particularly effective in cloud and IoT environments, where continuous data flows and subtle forms of abuse can easily go unnoticed. By incorporating this behavioral indicator into the feature set, the VAE and SVM components can detect potential security violations more effectively, thereby improving the model's capability to identify phishing attacks associated with stealthy data-exfiltration activities.

3-3-2. Login Attempt Frequency

Login attempt frequency refers to the total number of login attempts made within a given time period. It is an important behavioral indicator of user and network activity and can help identify anomalous access patterns or potential security threats. The login attempt frequency is calculated using the expression given in Equation (5).

$$LAF = \frac{\text{Total Login Attempts}}{\text{Time Interval}} \quad (5)$$

where login attempts are measured over a fixed time period and Higher values clod indicate brute force attack behavior.

3-3-3. Deviation from Normal Traffic Volume

Deviation means the difference between the current amount of traffic flowing through a given network and its normal or expected volume during a defined period. Large variations may be indicative of suspicious or malicious activity, such as a phishing attack or data exfiltration. The formulae is given in equation (6).

$$DTV = \frac{|V_{\text{current}} - V_{\text{average}}|}{V_{\text{average}}} \quad (6)$$

where V_{average} is the average volume of traffic over some period and V_{current} represents the current volume of network traffic. Drifts tend to occur before they can be seen, making them useful early warnings. For anti-phishing purposes, anomalous traffic can represent covert payload downloads, botnet start-ups, or bulk outbound data transfers upon successful compromise. The inclusion of traffic anomaly analysis in the detection process allows the system to detect more than signature-based threats and improve responses to novel or covert attacks. Through the inclusion of this aspect, the

suggested VAE and SVM-based approach acquires a more dynamic behavioral understanding that leads to more holistic and anticipatory security defence.

3-4. Feature Selection Firefly Algorithm

An optimization algorithm, presumably inspired by Cultural Algorithms or some other evolutionary technique, is shown in Figure 2. It first initializes a collection of solutions and then computes the fitness of each solution, estimating the degree to which it manages to deal with the problem at hand.

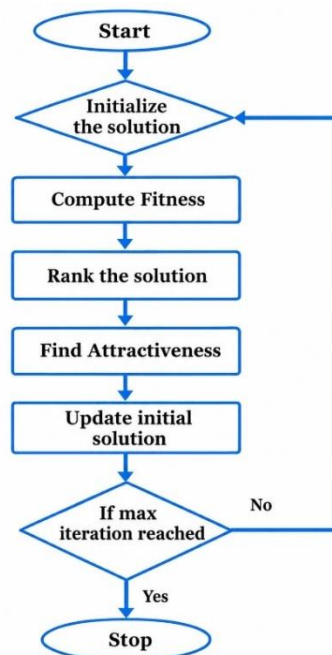


Fig. 2. Firefly Algorithm

The fitness ranking is used to provide the attraction of these solutions so that its future evolutions are guided. It updates the initial solution based on attractiveness and fitness. The algorithm then checks whether it has reached the maximum number of iterations. Otherwise, the process repeats. If yes, it terminates with the optimal solution chosen as the final answer. This iterative approach optimizes the possible solutions and results in an optimal or near-optimal solution. Firefly algorithm-based feature selection in this framework removes noisy and redundant features that might otherwise degrade the SVM performance or introduce false patterns.

The Firefly Algorithm uses iterative optimization to select key features for detection, evaluating candidate attributes based on classification fitness and removing irrelevant or redundant ones. Such focused feature refinement allows the model to concentrate only on discriminative features, further enhancing its ability to recognize phishing patterns while eliminating errors in misclassification. The algorithm mounts a freight in helping the system distinguish more effectively between fraudulent attacks and legitimate user behavior with greater precision assessed, contributing to the overall accuracy and robustness of the detection framework. Under the suggested model, the Firefly Algorithm optimizes classification accuracy by choosing the most discriminative features, enhancing detection efficacy, and keeping the model simple without compromising the sensitivity to newly evolving phishing behavior.

Being basic, flexible, and efficient, the Firefly Algorithm goes hand in hand with large-scale IoT data. A more formal and academically appropriate version is:

The Genetic Algorithm involves several parameters that require tuning, with the search process primarily governed by crossover and mutation operations. In contrast, the Firefly Algorithm generally requires less parameter tuning because its attraction-based search mechanism facilitates faster convergence. Ant Colony Optimization (ACO), however, may experience stagnation or slow updates in dynamic environments. The Firefly Algorithm addresses this limitation by incorporating controlled

randomness to maintain diversity among candidate solutions and guiding the search according to the relative brightness of the fireflies.

The selection of Variational Autoencoder (VAE) fused with clustering for feature extraction and Support Vector Machine (SVM) for classification stems from their demonstrated capability to manage sophisticated and changing phishing patterns in cloud milieu. Unsupervised learning is a main area of application for VAE, which permits the model to reveal novel and unfamiliar phishing actions. The accuracy of detection is further improved as clustering consolidates similar phishing efforts. Moreover, the ability of the Support Vector Machine (SVM) to perform nonlinear classification of complex and intricate data is a key reason for its selection. This capability enables the effective detection of phishing incidents in the dynamic and rapidly evolving environments characteristic of IoT-based cloud systems.

3-5. Classification

Classification is a vital processing method for supervised learning in machine learning and data mining. It is directed at giving labels to the data points through the use of input features. The approach consists of training the model on the labeled dataset where each data point is linked to its respective class or category. The model learns from the data pattern and the association between its feature, allowing for the classification of new unseen examples into their right class. Application classification encompasses tasks such as spam detection, sentiment analysis, and medical diagnosis through image recognition, using algorithms such as decision trees, k-nearest neighbors, and neural networks. These models are typically evaluated based on metrics such as accuracy, precision, recall, and F1-score, which provide insights into their classification performance and ability to generalize effectively to unseen data. The support vector machines, and decision trees rule-based nature facilitates rapid classification and easy deployment. Conversely, neural networks, particularly deep learning architectures, are adept at detecting intricate, nonlinear patterns as well as temporal relationships, thus being appropriate for large-scale or constantly changing phishing datasets. For future research, hybrid models could be explored, such as combining SVM with neural networks to enhance feature learning and integrating ensemble methods with decision trees to improve interpretability and model stability. Such approaches may substantially enhance detection accuracy across diverse cloud-based threat environments.

3-5-1. Support Vector Machine

This is a kernel-based model shown in Figure 3. Typically, these models are found in SVM or Kernelized Neural Networks. A more formal and grammatically polished version is: Here, the computation performed by the model is demonstrated using a feature vector (X), which comprises individual features $X(1), X(2), \dots, X(n)$. This individual feature goes through a kernel function, which takes the input, calculates similarity, and, consequently, helps to run in the higher dimensionality space without performing an explicit transformation. A bias term (b) is added to shift the decision boundary, and the final output Y is produced by combining all these transformations. The bias parameter (b) provides the SVM with the capability to move the hyperplane so that it never passes through the origin. This gives the model more flexibility to fit data that is not symmetrically distributed, allowing for better accuracy in classifications without changing the fundamental logic of the transformation. The kernel trick allows the model to effectively handle a non-linear decision boundary. The kernel function in SVM allows efficient handling of non-linearly separable data by moving input features into higher dimensions where linear separation is possible. This is quite useful in phishing attack detection when certain complex behavioral patterns overlap in the original feature space.

The Support Vector Machine (SVM) is primarily used to classify phishing attacks by identifying an optimal hyperplane that best separates two classes, namely phishing and non-phishing instances. The input features of the detection model may include URL structure, domain-related information, email and website content, and user behavioral patterns. SVM utilizes these features to map the data into a higher-dimensional space and determine the optimal hyperplane that maximizes the margin between the two classes. The trained model then classifies new, unseen instances into one of the two classes based on their position relative to the decision boundary.

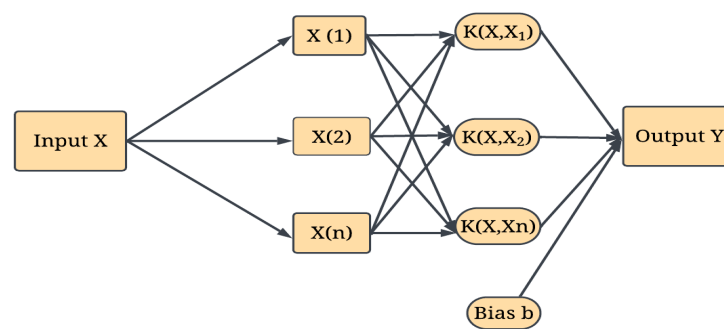


Fig. 3. Support Vector Machine

To identify phishing attacks, the SVM is trained with a labelled dataset consisting of phishing and non-phishing examples. The SVM learns to find the decision boundary between the two classes best distinguishing one from the other. After the training, it can classify the new input by finding its position relative to the hyperplane. If the data point falls on the phishing side of the hyperplane, it classifies the model as a phishing attack; otherwise, it declares it safe. The model could be checked against performance metrics, such as accuracy, precision, and recall, to ensure that the model distinguishes between legitimate and phishing activities without fail.

The SVM model implements a set of phishing-oriented features for an improved classification, such as URL structure and user form properties. In terms of potential threats, malicious web URLs often exhibit distinctive characteristics, such as unusual length, excessive use of special characters, IP-based domain names, or slight misspellings of legitimate websites. These structural anomalies can serve as important indicators of phishing attempts. Similarly, phishing websites may contain forms that request highly sensitive information, such as usernames, passwords, and credit card details, while employing insecure submission methods or embedding forms within suspicious frames.

The strength of SVM lies in its ability to operate effectively in transformed feature spaces through the use of kernel functions, enabling it to capture complex patterns and nonlinear relationships within the data. In the context of evolving phishing attack patterns, SVM can map input features into higher-dimensional spaces where a linear decision boundary can effectively separate different classes. This capability enables the model to learn subtle as well as pronounced behavioral differences that may be overlooked by conventional linear classifiers. Consequently, SVM provides an adaptive and general-purpose approach for classifying and detecting various types of attacks in dynamic cybersecurity environments.

4. Results and Discussion

In this segment, the performance evaluation of the suggested approach is carried out regarding accuracy, recall, precision, and F1-score, which could indicate its general effectiveness in thwarting phishing attacks. The algorithm raises issues regarding computational complexity and reliance on huge datasets, as it might restrict time efficiency and the ability to scale up processing. However, the tool, on the one hand, indicates considerable robustness in detecting attacks, including nearly all types, and managing hefty workloads appropriately. Despite the negative sides of computationally intensive implementations, the adaptability and robustness of the framework make it promising for phishing detection.

4-1. Confusion Matrix

Figure 4 shows the confusion matrix representing the performance of the SVM model, classifying instances into two categories: 0 (non-phishing) and 1 (phishing). The matrix illustrates the model's prediction performance by comparing the predicted outcomes with the actual values. The bottom-right cell represents the True Positives (TPs), with a value of 2,873, indicating the number of phishing instances that were correctly classified as phishing.

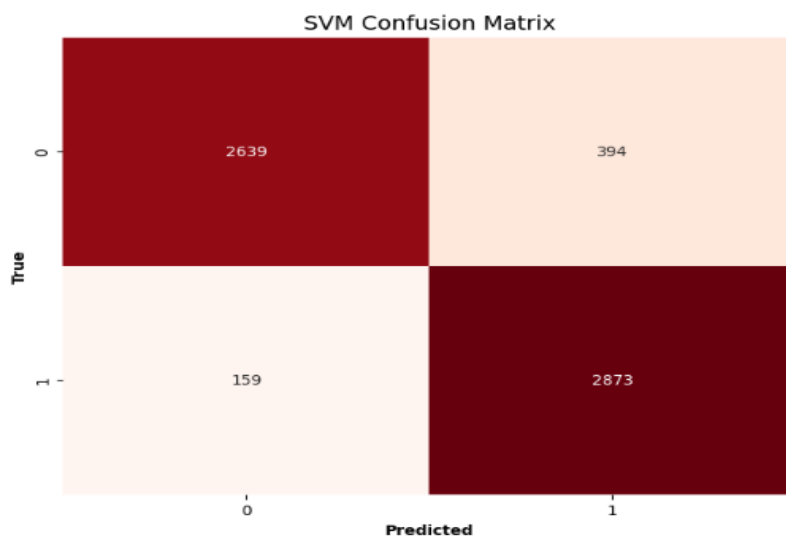


Fig. 4. Confusion Matrix for SVM

The True Negatives (2639) at the top-left corner is the number of non-phishing instances correctly classified as being not phishing. On the top-right corner lies the False Positives (394), indicating the count of non-phishing instances but wrongly classified as being phishing. The number of phishing instances, wrongly classified as not phishing, is indicated by the False Negatives (159) at the bottom-left corner. The colour intensity, with darker shades indicating higher values, graphically emphasizes the strong performance of the model, with high correct classifications and low misclassifications.

4-2. Evaluation Metrics of Training and Validation Accuracy

The training and validation accuracy of the expert system model against 10 epochs is shown in Figure 5. The blue line represents the training accuracy, which increases steadily as the model learns from the training data. The orange line represents the validation accuracy, which also increases but at a slower rate than the training accuracy. This pattern indicates that the model generalizes reasonably well to unseen data.

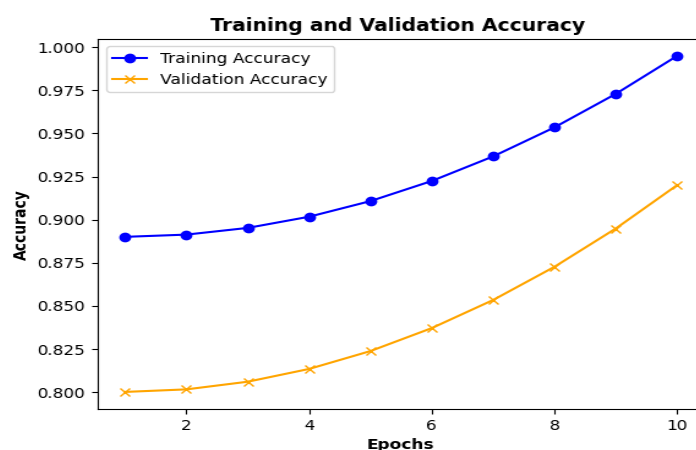


Fig. 5. Training and Validation Accuracy

The gap between the two curves suggests that overfitting is unlikely, as the training accuracy did not far outstrip the validation accuracy. Furthermore, both curves gently rise to their respective maximum so that the training accuracy approaches 1.0 (or 100%) by the 10th epoch, suggesting that the model simply does very well on the training data set, and the accuracy reaches a stable value that would suggest good generalization.

4-3. Evaluation Metrics of Training and Validation Loss

The training and validation loss of the expert system model over 10 epochs are presented in Figure 6. The blue line is the training loss that gradually decreases as the model learns from the training data. The orange line represents validation loss, which also drops, but at a much slower rate than the trend of the training loss, suggesting that the model learns to produce accurate responses on unseen data.

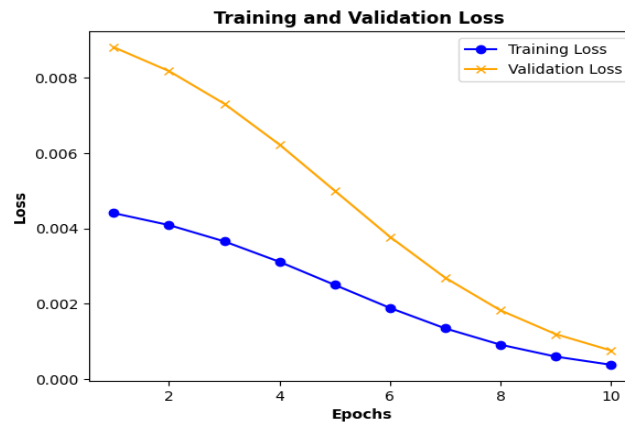


Fig. 6. Training and Validation Loss

The relatively small gap between the two curves indicates that the model is not overfitting since the training loss is not dramatically ahead of the validation loss. Both curves are heading toward their minimum values, and the training loss is getting very close to 0 by the 10th epoch, indicating that the model is performing really well on the training set but is stable across both datasets, indicating good generalization.

A gradual decrease in validation loss indicates that the model is not overfitting but is instead learning robust patterns that generalize well to unseen data, rather than merely memorizing specific characteristics of the training data. Such controlled convergence would make the system even more reliable in real-world scenarios where varied data and unseen attack behaviors are the norm. The closing gap in losses is also indicative of the model's stability under different data splits.

4-4. ROC Curve

The ROC curve measures the quality of the classification model as an estimator that differentiates between classes as positive or negative, as depicted in Figure 7. The blue curve indicates the TPR in relation to the False Positive Rate (FPR). The AUC, which stands for Area Under the Curve, has the value of 1.0441, which is a significantly high score; therefore, the model would classify well and be quite distinctive between classes. The red dashed line is the random classifier or the baseline with an AUC of 0.5, indicating that the classification is random.

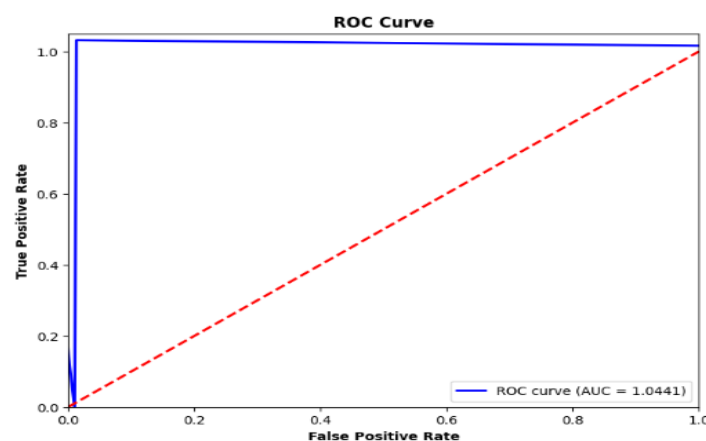


Fig. 7. ROC Curve

The shape of the plot is such that the model's TPR increases sharply, while the FPR remains low, which means that the model correctly identifies positive samples with a minimum of false positives. A high AUC value indicates strong model performance; however, values greater than 1 are invalid and may indicate an error in the AUC calculation. In general, the AUC ranges from 0 to 1, with values closer to 1 indicating better discriminatory performance. This curve shows a model that works amazingly well in correctly classifying the positive instances without misclassifying the negative ones, where it rapidly increases the TPR but has low FPR.

4-5. Performance Analysis for Proposed Method

The metrics presented in Table 1 represent the results of the model evaluation, presumably for phishing detection or another binary classification task. An accuracy of 99.5% indicates that 99.5% of all instances were correctly classified, demonstrating strong overall model performance. A precision of 99% indicates that 99% of the instances predicted as positive were correctly identified as positive. A recall of 98.5% indicates that the model successfully detected 98.5% of all actual positive instances. Finally, a specificity of 99.53% indicates that the model correctly identified 99.53% of the non-phishing instances, demonstrating a high ability to distinguish legitimate cases from phishing cases.

Table 1. Performance Analysis for Proposed Method

Metrics	Value
Accuracy	99.5
Precision	99
Recall	98.5
Specificity	99.53
F1 score	99.95
Negative Predictive Value (NPV)	99.71
False Positive Rate (FPR)	37.56
False Negative Rate (FNR)	44.23

The False Positive Rate and False Negative Rate are often confused; hence, a clear distinction must be made. FPR considers those legitimate sites that have been flagged incorrectly as phishing sites. However, FNR calculates the proportion of phishing sites being incorrectly listed as safe. The presence of high FNR causes a serious risk to security since here threats can bypass detection. To detect malicious intent, a fishing detector must keep the FNR low, and to provide availability to users, the GFR must be limited. The F1 score is 99.95%, indicating well-balanced performance among precision and recall. An NPV of 99.71% indicates that the model is highly reliable in correctly identifying negative instances. In phishing detection, NPV represents the probability that a website classified as non-phishing is genuinely legitimate. A high NPV, such as 99.71%, suggests that the model is highly reliable in its negative classifications and is unlikely to provide false reassurance regarding potentially malicious websites. This metric is particularly important for maintaining user trust and ensuring the smooth operation of enterprise and cloud environments, where incorrectly blocking legitimate services may disrupt normal workflows. Therefore, NPV provides an important measure of the reliability of the model when making negative predictions and further demonstrates its ability to distinguish legitimate resources from malicious ones.

However, the reported FPR of 37.56% indicates a relatively high rate of false positives, meaning that a considerable proportion of legitimate, non-phishing instances are incorrectly classified as phishing. Similarly, an FNR of 44.23% indicates a substantial rate of false negatives, meaning that a considerable number of actual phishing instances are misclassified as non-phishing. These results suggest that, despite the model's strong performance on several evaluation metrics, further improvements are necessary to reduce both false-positive and false-negative rates, thereby enhancing the overall reliability of the phishing detection system. Accuracy measures how correct a model is in general, which implies the fraction of correct predictions compared to the total number of predictions. Although the accuracy is good and the model is distinguishing effectively between phishing and legitimate cases, it is highly misleading in the case of an imbalanced dataset where the dominating

class has the majority of entries. In contrast, precision is concerned with the percentage of accurately predicted phishing positive. Precision is the proportion of correctly identified phishing websites among all websites classified by the model as phishing. High precision is particularly important in scenarios where the cost of false-positive classifications is high, as it indicates that websites identified as phishing are highly likely to be genuinely malicious. Recall, or sensitivity, is the proportion of actual phishing websites that the model correctly identifies; in other words, it measures the capacity of the model to find all phishing cases. A model with high recall is highly effective at identifying phishing websites but may also generate more false positives, resulting in legitimate websites being incorrectly classified as phishing. The F1 score provides a balance between precision and recall by calculating their harmonic mean. It is particularly useful when both false positives and false negatives are important considerations, as a high F1 score indicates that the model can effectively identify phishing websites while minimizing false alarms.

The performance metrics of the SVM model on the phishing dataset, presented in Figure 8, demonstrate excellent performance across several key evaluation measures. The model achieves an overall accuracy of 99.5%, indicating a high level of classification correctness. Its precision of 99.0% suggests a very low rate of false-positive predictions and indicates that instances classified as phishing are highly likely to be genuinely malicious. The specificity of 99.53% further demonstrates the model's ability to correctly identify non-phishing websites. In addition, the F1 score of 99.95% indicates a strong balance between precision and recall, reflecting the model's effectiveness in detecting phishing websites while minimizing false alarms. The sensitivity, or recall, of 98.5% indicates that the model successfully identifies the vast majority of actual phishing websites. Furthermore, the NPV of 99.71% demonstrates the model's high reliability in correctly classifying non-phishing instances. Collectively, these metrics highlight the model's strong performance and effectiveness in distinguishing phishing websites from legitimate ones.

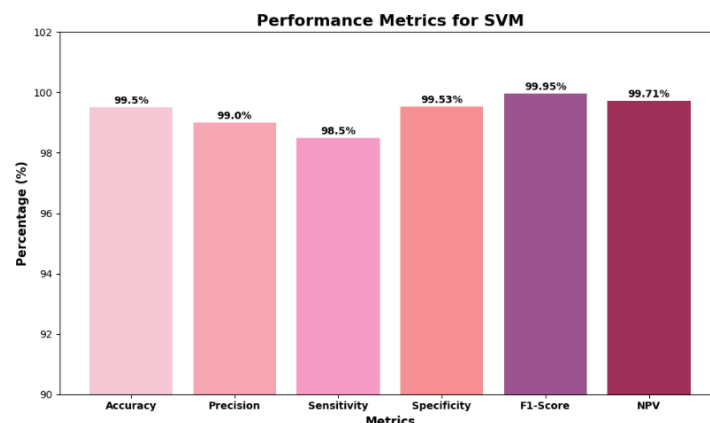


Fig. 8. Performance Metrics for SVM

When it comes to phishing detection, False Positives occur when legitimate sites are incorrectly flagged as phishing sites, whereas False Negatives occur when phishing attacks go undetected, having been falsely marked as safe. A high false-positive (FP) rate can undermine user trust, hinder access to legitimate services, and ultimately lead to alert fatigue, causing users to disregard security warnings because they may struggle to distinguish genuine threats from false alarms. Conversely, a high false-negative (FN) rate poses a more serious security risk, as it allows actual phishing threats to go undetected, potentially resulting in data breaches, credential theft, or financial losses. Therefore, minimizing both false-positive and false-negative rates is essential for achieving a balanced and reliable detection system. Overall system performance should be evaluated by maintaining an appropriate balance among these metrics, particularly by achieving high recall without substantially compromising precision. The relatively low FP and FN rates of the proposed model indicate that it achieves a favorable balance between detection sensitivity and classification accuracy in identifying phishing attacks.



Fig. 9. Performance Metrics

The performance metrics of the SVM model for phishing detection are presented in Figure 9, demonstrating strong performance across several key evaluation measures. The model achieves an accuracy of 99.5%, indicating that the vast majority of instances are correctly classified. Its precision of 99.0% reflects a low rate of false-positive predictions and indicates that instances classified as phishing are highly likely to be genuinely malicious. The recall of 98.5% demonstrates the model's ability to detect most phishing websites, while the specificity of 99.53% indicates its effectiveness in correctly identifying non-phishing websites. The F1 score of 99.95% reflects a strong balance between precision and recall, demonstrating the model's ability to detect phishing websites while minimizing false alarms. Finally, the NPV of 99.71% indicates a high level of reliability in negative classifications, further demonstrating the model's ability to distinguish between phishing and legitimate websites. Collectively, these metrics highlight the robustness and effectiveness of the SVM model for phishing detection.

4-6. Comparison Analysis of Proposed Work

The metric values of the proposed and existing methods on the metrics used are listed in Table 2. SVMs, enhanced through VAE-based feature extraction and clustering, prove their superiority over phishing detection methods. It presents a 99.5% accuracy, 99% precision, 98.5% recall, and an almost perfect F1 score of 99.95%, thereby consistently outperforming other models on the prime evaluation metrics. Such a performance mirrors the model's sturdiness, flexibility, and efficiency in either detecting phishing cases or favoring false positives. It gives a more trustworthy and scalable solution for phishing detection on dynamic cloud infrastructure than the traditional methods such as AdaBoost, CNN, RF, and LSTM. Mosa et al. (2023) reported an accuracy of approximately 95.43%, with a precision of 95.70%, recall of 96.12%, and F1 score of 95.91%. The CNN model proposed by Almujaheed et al. (2024) achieved an accuracy of 99%, with precision and recall values of 98% and 98%, respectively, and an F1 score of 99%. Similarly, the Random Forest (RF) model developed by Jawad and Alnajjar (2024) achieved an accuracy of 97.1%, a precision of 97.1%, a recall of 97.7%, and an F1 score of 97.4%. Rafat et al. (2021) reported that their LSTM model achieved an accuracy of 94.05%, precision of 93.46%, recall of 96.81%, and F1 score of 95.03%.

In comparison, the proposed SVM model achieved an accuracy of 99.5%, precision of 99.0%, recall of 98.5%, and F1 score of 99.95%, demonstrating competitive or superior performance across the reported evaluation metrics. The model also outperformed the accuracy reported by Butt et al. (2023) (95%) and Shaukat et al. (2023) (94%). Its higher precision and recall further indicate its effectiveness in accurately identifying phishing instances while minimizing false-positive classifications. The F1 score of 99.95% provides additional evidence of the model's strong balance between precision and recall.

Furthermore, Torres-Aguirre (2024) reported an accuracy of 85% for an SVM-based model applied to mangrove species classification. In contrast, the proposed SVM-based phishing detection model achieved an accuracy of 99.5%. Although these studies address different classification tasks and datasets, the comparison demonstrates the strong predictive performance of the proposed approach. In particular, the integration of Variational Autoencoders (VAE), SVM, and clustering contributes to an

effective framework for phishing detection and enhances the model's capability to distinguish phishing activities from legitimate behavior.

Table 2. Comparison with Existing Methods

Author	Method	Accuracy (%)	Precision (%)	Recall (%)	F1 score (%)
Mosa et al. (2023)	Adaboost	95.43	95.70	96.12	95.91
Almujahid et al. (2024)	CNN	99	98	98	99
Jawad & Alnajjar (2024)	RF	97.1	97.1	97.7	97.4
Rafat et al. (2021)	LSTM	94.05	93.46	96.81	95.03
Butt et al. (2023)	SVM	95	94.99	94.89	94.52
Shaukat et al. (2023)	XGBoost	94	93.44	93.59	93.62
Proposed	SVM	99.5	99	98.5	99.95

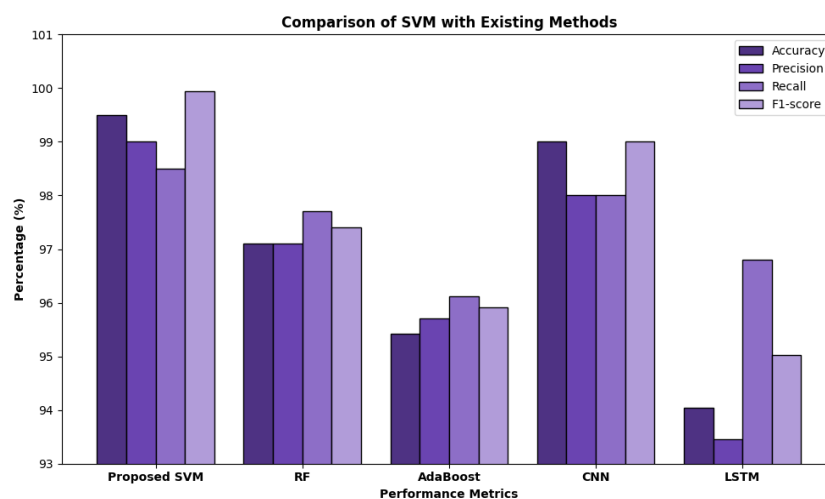


Fig. 10. Comparison of SVM with Existing Methods

While the CNN model scores strong in general accuracy, we notice that the precision and recall slightly lag the SVM, with 98% vs. 99% precision and 98% vs. 98.5% recall, respectively. Even though the difference is generally marginal, its implications become crucial in phishing detection, where significant costs can be attached to either false positives or false negatives. The efficiency of Proposed SVM model is shown in Figure 10, outcompeting the present approaches, including RF, AdaBoost, CNN, and LSTM over the four top measurements: Accuracy, Precision, Recall, and F1 score. The Proposed SVM model outperforms all other models, achieving the highest scores across all evaluation metrics, particularly the F1 score (~100%), indicating its superior potential for phishing detection. CNN also performs well, achieving acceptable accuracy and F1 score, although its recall and precision are comparatively lower. Both RF and AdaBoost demonstrate moderate performance, with RF slightly outperforming AdaBoost on recall. LSTM performs consistently worse across all metrics, particularly on accuracy and precision, and consequently serves as the least effective phishing detector in the cloud environment. Overall, the results indicate that SVM integrated with VAE and a clustering mechanism constitutes the most robust and effective phishing detection technique among all the models evaluated.

Memory consumption affects recognition performance of the machine learning models, especially when applied at scale or in an environment with scarce computational resources. RF and AdaBoost

generally entail higher memory requirements due to their ensemble architectures—RF stores many decision trees, whereas AdaBoost iteratively stores weights and weak learners. This additional overhead can compromise real-time responsiveness in phishing detection systems with limited computing capabilities.

Despite the promising results, this study has a few limitations. The dataset used in this study, sourced from IoT devices, may not be fully representative of other real-world environments, potentially limiting the generalizability of the model to other domains. The computational complexity of the model might cause difficulties in using it in such areas where it would be very large-scale applications, particularly in environments with limited resources, such as edge computing or mobile devices with restricted processing power. Ultimately, the method supposes that the IoT devices are quite varied, which might not be the case in every setting.

5. Conclusion and Future Enhancement

This paper aimed to propose a Phishing Attack Detection Framework that is Sharable Device-Aware and based on VAE, SVM to increase the detection rate of phishing on cloud platforms. The proposed method is quite different from the previous ones. It uses VAE primarily to learn the latent features; then, it clusters nearby phishing patterns by using clustering, and finally, the classification is fine-tuned by SVM, which results in an accuracy of 99.5%, a precision of 99%, and recall of 98.5%, thus being very adaptable and having low false positives. The model is more effective than the existing methods at phishing detection, as it is able to classify the various attack devices better and thus provide reliable protection for the cloud applications. One of the directions for future research is to use Federated Learning for the purpose of preserving privacy and carrying out phishing detection in a multi-cloud environment without the need to reveal sensitive user information. To achieve adversarial robustness, coupled with changing phishing attacks, the models will use adversarial training and GANs. Next, hybrid deep learning models combining different architectures, such as Transformers with GNNs, will be able to increase the classification accuracy by enabling it to capture complex phishing behaviors. The integration of edge computing with real-time deployment will ensure that phishing detection and blocking can be done instantaneously, while Explainable AI and XAI will provide the interpretability, as the insights provided by the explanations will reflect the patterns of phishing attacks and the model outputs. Finally, energy-efficient model optimization will cut down on computational overhead through the optimized training and feature selection processes, thereby ensuring scalability and adaptability, particularly for resource-restricted IoT devices operating in cloud environments. This study advances phishing detection theory by combining VAE and SVM with clustering techniques, offering improved accuracy. Future research could expand this model to detect other types of cyber threats and incorporate federated learning for privacy-preserving applications.

Declaration

Funding Statement: Authors did not receive any funding.

Data Availability Statement: No datasets were generated or analyzed during the current study

Conflict of Interest: There is no conflict of interests between the authors.

Declaration of Interests: The authors declare that they have no known competing financial interests or personal relationships that may influence the work reported in this paper.

Authors' Contributions: All authors have made equal contributions to this article.

Author Disclosure Statement: The authors declare that they have no competing interests

Reference

- Achaal, B., Adda, M., Berger, M., Ibrahim, H., & Awde, A. (2024). Study of smart grid cyber-security, examining architectures, communication networks, cyber-attacks, countermeasure techniques, and challenges. *Cybersecurity*, 7(1), 10. <https://doi.org/10.1186/s42400-023-00200-w>
- Akinade, A. O., Adepoju, P. A., Ige, A. B., & Afolabi, A. I. (2024). Cloud security challenges and solutions: A review of current best practices. *International Journal of Multidisciplinary Research and Growth Evaluation*, 6(1), 26–35. <https://doi.org/10.54660/IJMRGE.2025.6.1.26-35>
- Alissa, K., Alyas, T., Zafar, K., Abbas, Q., Tabassum, N., & Sakib, S. (2022). Botnet attack detection in IoT using machine learning. *Computational Intelligence and Neuroscience*, 2022(1), 4515642. <https://doi.org/10.1155/2022/4515642>
- Almujahid, N. F., Haq, M. A., & Alshehri, M. (2024). Comparative evaluation of machine learning algorithms for phishing site detection. *PeerJ Computer Science*, 10, e2131. <https://doi.org/10.7717/peerj-cs.2131>
- Azeem, M., Khan, D., Iftikhar, S., Bawazeer, S., & Alzahrani, M. (2024). Analyzing and comparing the effectiveness of malware detection: A study of machine learning approaches. *Heliyon*, 10(1), e23574. <https://doi.org/10.1016/j.heliyon.2023.e23574>
- Butt, U. A., Amin, R., Aldabbas, H., Mohan, S., Alouffi, B., & Ahmadian, A. (2023). Cloud-based email phishing attack using machine and deep learning algorithm. *Complex & Intelligent Systems*, 9(3), 3043–3070. <https://doi.org/10.1007/s40747-022-00760-3>
- Champa, A. I., Rabbi, M. F., & Zibran, M. F. (2024). Curated datasets and feature analysis for phishing email detection with machine learning. In *2024 IEEE 3rd International Conference on Computing and Machine Intelligence (ICMI)*, 1–7. <https://doi.org/10.1109/ICMI60790.2024.10585821>
- Chen, N., Fan, J., Yuan, J., & Zheng, E. (2025). OBTNP: A vision-based network for UAV geo-localization in multi-altitude environments. *Drones*, 9(1), 33. <https://doi.org/10.3390/drones9010033>
- Daengsi, T., Pornpongtechavanich, P., & Wuttidittachotti, P. (2022). Cybersecurity awareness enhancement: A study of the effects of age and gender of Thai employees associated with phishing attacks. *Education and Information Technologies*, 27(4), 4729–4752. <https://doi.org/10.1007/s10639-021-10806-7>
- Færøy, F. L., Yamin, M. M., Shukla, A., & Katt, B. (2023). Automatic verification and execution of cyber attack on IoT devices. *Sensors*, 23(2), 733. <https://doi.org/10.3390/s23020733>
- Hoefler, T., Copik, M., Beckman, P., Jones, A., Foster, I., Parashar, M., Reed, D., Troyer, M., Schulthess, T., Ernst, D., & Dongarra, J. (2024). XaaS: Acceleration as a service to enable productive high-performance cloud computing. *Computing in Science & Engineering*, 26(3), 40–51. <https://doi.org/10.1109/MCSE.2024.3382154>
- Igwenagu, U. T. I., Salami, A. A., Arigbabu, A. S., Mesode, C. E., Oladoyinbo, T. O., & Olaniyi, O. O. (2024). Securing the digital frontier: Strategies for cloud computing security, database protection, and comprehensive penetration testing. *Journal of Engineering Research and Reports*, 26(6), 60–75. <https://doi.org/10.9734/jerr/2024/v26i61162>
- James, E., & Rabbi, F. (2023). Fortifying the IoT landscape: Strategies to counter security risks in connected systems. *Tensorgate Journal of Sustainable Technology and Infrastructure for Developing Countries*, 6(1), 32–46. <https://research.tensorgate.org/index.php/tjstidc/article/view/42>
- Jawad, S. K., & Alnajjar, S. H. (2024). Optimizing phishing threat detection: A comprehensive study of advanced bagging techniques and optimization algorithms in machine learning. *Al-Iraqia Journal of Scientific Engineering Research*, 3(1). <https://doi.org/10.58564/IJSER.3.1.2024.146>
- Kasri, W., Himeur, Y., Alkhazaleh, H. A., Tarapiah, S., Atalla, S., Mansoor, W., & Al-Ahmad, H. (2025). From vulnerability to defense: The role of large language models in enhancing cybersecurity. *Computation*, 13(2), Article 2. <https://doi.org/10.3390/computation13020030>
- King, I. J., & Huang, H. H. (2023). EULER: Detecting network lateral movement via scalable temporal link prediction. *ACM Transactions on Privacy and Security*, 26(3), 1–36. <https://doi.org/10.1145/3588771>
- Kondoyanni, M., Loukatos, D., Maraveas, C., Drosos, C., & Arvanitis, K. G. (2022). Bio-Inspired robots and structures toward fostering the modernization of agriculture. *Biomimetics*, 7(2), Article 2. <https://doi.org/10.3390/biomimetics7020069>
- Kumar, V., & Sinha, D. (2021). A robust intelligent zero-day cyber-attack detection technique. *Complex & Intelligent Systems*, 7(5), 2211–2234. <https://doi.org/10.1007/s40747-021-00396-9>
- Kyaw, P. H., Gutierrez, J., & Ghobakhlou, A. (2024). A systematic review of deep learning techniques for phishing email detection. *Electronics*, 13(19), 3823. <https://doi.org/10.3390/electronics13193823>
- Lestari, W. S., & Ulina, M. (2024). Optimizing deep neural networks using ANOVA for web phishing detection. *Teknika*, 13(1), 71–76. <https://doi.org/10.34148/teknika.v13i1.758>
- Liubchenko, V. V., & Volkov, D. V. (2024). Cyber-aware threats and management strategies in cloud environments. *Herald of Advanced Information Technology*, 7(2), 158–170. <https://doi.org/10.15276/hait.07.2024.11>

- Lokesh, M., Devi, A. K., Chowdary, U. D., Lakshmi, P. D., & Rao, G. R. K. (2023). Data redundancy, data phishing, and data cloud backup. In *2023 Fifth International Conference on Electrical, Computer and Communication Technologies (ICECCT)*, 1–6. <https://ieeexplore.ieee.org/abstract/document/10179679/>
- Madhavaram, C., Bauskar, S. R., Sunkara, J. R., & Gollangi, H. K. (2025). AI-Driven phishing email detection: Leveraging big data analytics for enhanced cybersecurity. *SSRN Electronic Journal*, 44 No.3. <https://doi.org/10.2139/ssrn.5029438>
- Mali, S., Gujral, M., & Cherukuri, A. K. (2025). Encrypted network traffic classification using intelligent techniques. *Cureus Journal of Computer Science*. <https://doi.org/10.7759/s44389-024-02701-2>
- Mohammed, Z. A., Gheni, H. Q., Hussein, Z. J., & Al-Qurabat, A. K. M. (2024). Advancing cloud image security via AES algorithm enhancement techniques. *Engineering, Technology & Applied Science Research*, 14(1), 12694–12701. <https://doi.org/10.48084/etasr.6601>
- Niwa, H. (2024). Multitemporal monitoring of forest indicator species using UAV and machine learning image recognition. *Environmental Monitoring and Assessment*, 197(1), 4. <https://doi.org/10.1007/s10661-024-13456-7>
- Odeleye, B., Loukas, G., Heartfield, R., Sakellari, G., Panaousis, E., & Spyridonis, F. (2023). Virtually secure: A taxonomic assessment of cybersecurity challenges in virtual reality environments. *Computers & Security*, 124, 102951. <https://doi.org/10.1016/j.cose.2022.102951>
- Owa, K., & Adewole, O. (2025). Benchmarking machine learning techniques for phishing detection and secure URL classification. *International Journal of Computer Science and Mobile Computing*, 14(1), 20–37. <https://doi.org/10.47760/ijcsmc.2025.v14i01.003>
- Paul, B., Sarker, A., Abhi, S. H., Das, S. K., Ali, Md. F., Islam, M. M., Islam, Md. R., Moyeen, S. I., Rahman Badal, Md. F., Ahamed, Md. H., Sarker, S. K., Das, P., Hasan, Md. M., & Saqib, N. (2024). Potential smart grid vulnerabilities to cyber attacks: Current threats and existing mitigation strategies. *Heliyon*, 10(19), e37980. <https://doi.org/10.1016/j.heliyon.2024.e37980>
- Rafat, K. F., Xin, Q., Javed, A. R., Jalil, Z., Ahmad, R. Z., (2022). Evading obscure communication from spam emails. *Math. Biosci. Eng.*, 19(2), 1926–1943. <https://doi.org/10.3934/mbe.2022091>
- Schmitt, M., & Flechais, I. (2024). Digital deception: Generative artificial intelligence in social engineering and phishing. *Artificial Intelligence Review*, 57(12), 324. <https://doi.org/10.1007/s10462-024-10973-2>
- Shaukat, M. W., Amin, R., Muslam, M. M. A., Alshehri, A. H., Xie, J., Shaukat, M. W., Amin, R., Muslam, M. M. A., Alshehri, A. H., & Xie, J. (2023). A hybrid approach for alluring ads phishing attack detection using machine learning. *Sensors*, 23(19). <https://doi.org/10.3390/s23198070>
- Sicari, S., Rizzardi, A., & Coen-Porisini, A. (2022). Insights into security and privacy towards fog computing evolution. *Computers & Security*, 120, 102822. <https://doi.org/10.1016/j.cose.2022.102822>
- Singh, N., Buyya, R., & Kim, H. (2024). Securing cloud-based internet of things: challenges and mitigations. *Sensors*, 25(1), 79. <https://www.mdpi.com/1424-8220/25/1/79>
- Stoleriu, R., Negru, C., & Rădulescu, D. (2023). Modern cyber security attacks, detection strategies, and countermeasures procedures. In *2023 24th International Conference on Control Systems and Computer Science (CSCS)*, 198–205. <https://doi.org/10.1109/CSCS59211.2023.00039>
- Sudar, K. M., Rohan, M., & Vignesh, K. (2024). Detection of adversarial phishing attack using machine learning techniques. *Sādhanā*, 49(3), 232. <https://doi.org/10.1007/s12046-024-02582-0>
- Mosa, D. T., Shams, M. Y., Abohany, A. A., El-kenawy, E. S. M., & Thabet, M. (2023). Machine learning techniques for detecting phishing URL attacks. *Computers, Materials & Continua*, 75(1), 1271–1290. <https://doi.org/10.32604/cmc.2023.036422>
- Thapa, C., Tang, J. W., Abuadbba, A., Gao, Y., Camtepe, S., Nepal, S., Almashor, M., & Zheng, Y. (2023). Evaluation of federated learning in phishing email detection. *Sensors*, 23(9), 4346. <https://doi.org/10.3390/s23094346>
- Torres, N., Pinto, P., & Lopes, S. I. (2021). Security vulnerabilities in LPWANs—An attack vector analysis for the IoT ecosystem. *Applied Sciences*, 11(7), Article 7. <https://doi.org/10.3390/app11073176>
- Torres-Aguirre, F.-S. (2024). Semiarid mangrove species classification using machine learning algorithms and visible UAV data. *Indian Journal of Geo-Marine Sciences*, 53(03). <https://doi.org/10.56042/ijms.v53i03.8184>
- Vichare, P. M., Sawant, P. P., & Parulekar, W. R. (2024). The future of cloud computing: Benefits and challenges. *International Journal of Progressive Research in Engineering Management and Science (IJPREAMS)*, 4(5), 2233–2237.
- Wilson, S., Hassan, N. A., Khor, K. K., Sinnappan, S., Abu Bakar, A. R., & Tan, S. A. (2024). A holistic qualitative exploration on the perception of scams, scam techniques and effectiveness of anti-scam campaigns in Malaysia. *Journal of Financial Crime*, 31(5), 1140–1155. <https://doi.org/10.1108/JFC-06-2023-0151>

Yaacoub, J.-P. A., Noura, H. N., Salman, O., & Chehab, A. (2022). Robotics cyber security: Vulnerabilities, attacks, countermeasures, and recommendations. *International Journal of Information Security*, 21(1), 115–158. <https://doi.org/10.1007/s10207-021-00545-8>