

Cybershield : Email Spoofing Detection System

Sharad S. Bolde, Prof. Rahul P. Bembade, Shri P. Ingale, Sahil M. Patil

MIT Art ,Design and Technology University, Pune, India

Abstract: While most communication today has shifted to the digital world, email remains the backbone of both personal and organizational communication. Also, notwithstanding the coming of more advanced communication channels, email remains a preferred medium for information interchange-and unfortunately, for cybercriminals too. Of these, one of the most enduring threats in this domain has been email spoofing: a cunning technique allowing attackers to impersonate trusted entities by manipulating email headers, sender identities, and message metadata. Traditional security measures, such as SPF, DKIM, and DMARC, were designed to protect the authenticity of emails; however, recent studies show that those mechanisms can still be bypassed through delegation loopholes, inconsistencies in forwarding, and improper configuration across domains. Therefore, spoofing remains a valid and prevalent kind of cyber threat in 2025. This research proposes Cybershield, a hybrid intelligent framework for the real-time detection of email spoofing at the server level, to address these evolving risks. Cybershield takes inspiration from the reliable spoofing detection using artificial intelligence by Mane et al. (2025), extending beyond static header verification to machine learning-based anomaly detection and adaptive trust scoring. The system will focus on in-depth analysis of email header fields such as "Received," "Return-Path," "From," and "Replyto," and will then apply classification algorithms like Random Forest, Support Vector Machines, and ensemble-based predictors that identify discrepancies pointing to spoofing attempts. Each incoming email is analyzed for syntactic validation, combined with its behavioral pattern and historical sender behavior, thus making detections proactive rather than reactive. The system is implemented on a Python-based backend with Flask and Node modules for easy integration with existing mail servers, ensuring scalability and minimum latency in processing. What makes the difference is that Cybershield's context-aware learning model evolves with new spoofing patterns. Unlike protocol-based validation, which fails when attackers manipulate delegation mechanisms, Cybershield constantly improves its detection accuracy through learning from false positives, user feedback, and cross-domain anomalies. Furthermore, the proposed framework provides an easy-to-use graphical interface-a "Spoof Guardian"-that allows both technical administrators and ordinary users to derive proper interpretations of the detection results and perform corrective actions. This bridges the gap between forensic-level spoofing analyses and practical, deployable defense mechanisms. Preliminary testing of Cybershield on a mixed dataset of legitimate and spoofed emails yielded a detection accuracy exceeding 96%, outperforming traditional SPF/DKIM-based validation. It had successfully detected spoofed emails that were misclassified as legitimate by traditional filters, particularly those based on forwarding and sender-inconsistency vulnerabilities. By integrating machine learning with cybersecurity principles and email protocol forensics, Cybershield takes a robust and reliable approach toward one of the most stubborn problems in cybersecurity: trust verification in digital communication. Beyond detection, this work relates to the larger discourse on the security of the email ecosystem, reinforcing the importance of intelligent automation in maintaining both the integrity of communication and the trust of users..

Keywords: Cybershield's

I. INTRODUCTION

E-mail has become an essential part of our digital life: it is the basis of international communication both for private persons, businesses, and governments. It is fast, reliable, and universally accessible, which makes it ideal for official correspondence, information exchange, and authentication processes. However, the very openness and ubiquity that



make email so useful also make it a major target for cyberattacks. Among the wide range of email-based threats, email spoofing continues to be one of the most deceptive and damaging. In email spoofing, attackers forge the sender's address or manipulate header fields to make their messages appear to originate from trusted sources. This simple yet powerful deception often leads to severe consequences such as financial fraud, phishing, ransomware distribution, and large-scale data breaches.

Email Spoofing: The problem lies deep inside the Simple Mail Transfer Protocol — the basic standard that governs the sending of emails. SMTP was designed in the early days of the internet when network trust was assumed and authentication was not considered important.

Because SMTP does not verify the sender's true identity, malicious actors can manipulate the From, Reply-To, and Return-Path fields of an email header to deceive recipients and even get through many spam filters. Over time, in order to overcome such vulnerabilities, several authentication protocols were brought into being, namely SPF, DKIM, and DMARC, each addressing specific gaps in sender validation. SPF checks if the sending server is authorized by the domain owner, DKIM uses cryptographic signatures to verify integrity, and DMARC aligns both mechanisms to define policies regarding suspicious emails.

While these technologies represent major advancements, research consistently shows that they are not foolproof. Configuration errors, lack of global adoption, and design limitations continue to allow spoofed messages to slip through. Studies by Shen et al. (2021) and Ma et al. (2024) revealed that even with all three protocols correctly implemented, spoofing attacks can still succeed due to delegation loopholes, forwarding inconsistencies, and weak user interface protections

Their large-scale analyses of more than 30 e-mail services showed that popular platforms such as Gmail and Outlook are still vulnerable to several "imperceptible" spoofing methods. Similarly, Chauhan and Shah (2023) were able to show that an attacker could easily manipulate the header fields in such a way as to trick a recipient, especially when receiving mail servers fail to perform proper domain alignment checks.

These findings bring to the fore a critical truth: even robust authentication standards cannot guarantee security when human error, inconsistent configurations, and evolving attack techniques are involved.

Consequently, cybersecurity researchers have started adopting intelligent detection mechanisms that transcend static rule-based validation. Rather than mere verification as to whether an email passes SPF or DKIM, the focus now lies in subtly identifying patterns and behavioral anomalies within the metadata of an email. This is where AI and ML have emerged as game-changing tools. By evaluating a vast feature space-encompassing attributes like IP reputation, header inconsistencies, and linguistic tone, among others-ML algorithms can use all of these features to detect spoofed or phishing emails with great precision, even if the attackers craft messages that technically pass authentication checks.

Building on this vision, our research introduces Cybershield, an AI-driven, server-level spoofing detection system that incorporates forensic email header analysis with machine learning classification. Guided by the "Spoof Guardian" framework proposed by Mane et al. (2025),

Cybershield augments, rather than replaces, existing email security infrastructure through automated intelligence. It systematically analyzes every incoming email for abnormalities in key header attributes such as sender domain alignment, relay path, timestamp consistency, and content-encoding irregularities-and flags suspicious messages before they reach the user's inbox. Unlike many previous models that rely purely on static rules or manual inspection,

Accordingly, Cybershield uses a dynamic learning mechanism that can adapt to new spoofing strategies over time. The increasing rate of spoofing-related cyber incidents emphasizes the need for such a solution. According to various reports by cybersecurity firms and academic studies, more than 90% of data breaches originate from either a spoofed or phishing email that targets unsuspecting employees or customers. These not only attack technical gaps but also exploit human psychology by utilizing urgency, authority, or familiarity to deceive users into performing harmful actions. This makes the challenge not only technical but also behavioral: designing systems that are precise in their detection and intuitive in communicating risks with end users. The Cybershield project aims to bridge this gap. By integrating AI-based detection with an easy-to-understand user interface, it transforms spoofing prevention from a purely back-end process into an interactive security experience. Users receive transparent notifications when an email is flagged, while administrators gain detailed analytics to investigate spoofing sources and trends. In essence, Cybershield is not just a



filter — it is a trust enforcer within the digital communication chain. This paper explores the theoretical foundation, architecture, and implementation of Cybershield, evaluates its performance, and situates it within the broader evolution of email security.

II. LITERATURE SURVEY

During the past two decades, the problem of email spoofing has been much discussed from multiple perspectives: technical vulnerabilities, forensic investigation, protocol design, and machine learning-based prevention. Despite the major improvements achieved so far in the field of email authentication, spoofing attacks remain a persistent and challenging threat. This section reviews the literature in chronological order, highlighting that early studies were focused on the understanding and definition of spoofing, while recent studies are targeted to assess its persistence in contemporary infrastructures and propose intelligent detection mechanisms.

2.1 Early Understanding and Foundational Research

The first comprehensive study on email spoofing was presented in the International Journal of Computer Applications, 2010, by Pandove, Jindal, and Kumar.

Their paper provided the building blocks necessary for defining spoofing as the intentional modification of sender identity fields within a message to make the recipient believe that the email originates from a trusted source. The authors explained how the attackers took advantage of weaknesses in the Simple Mail Transfer Protocol-SMTP-a protocol that was never designed with built-in authentication. They also explained common misuse scenarios, including phishing, spam propagation, and malware distribution.

The study presented some basic countermeasures of authenticity verification of messages through digital signatures, encryption through PGP, and protocols like SSL/TLS. They estimated limitations to those solutions as well because of high implementation costs and a lack of universal compatibility between mail servers. Though pioneering for the time, this work emphasized user-side caution and administrative awareness rather than server-level automation, making it an early but necessary step in recognizing the spoofing problem.

2.2 Forensic Investigation of Email Header Manipulation

A more technical and investigative approach emerged with Pilli, Mishra & Joshi, 2014 in their paper entitled: Forensic Analysis of Email Address Spoofing.

This study examined how digital forensics can trace spoof emails. The authors identified discrepancies in the email header metadata, such as the Received, Return-Path, Message-ID, and Date fields, which would indicate tampering or falsification. They were able to show that, by reconstructing the transmission path from the raw header information, investigators can trace the source of a spoofed email to determine at which point in the routing process the message was tampered with.

While this work provided important insights into the forensic reconstruction of spoofing incidents, it was not automated and lacked scalability. Expert knowledge and manual inspection were required for this kind of analysis, which cannot be feasible in a large-scale enterprise environment that processes millions of emails on a daily basis. This study was necessary to establish the idea that there are verifiable traces of authenticity in the headers of emails, further influencing automatic header-based spoofing detection systems such as Cybershield.

2.3 Protocol-Level Countermeasures: SPF, DKIM, and DMARC

Next came the formulation and analysis of authentication protocols to perform validation checks on the sender identity. Chauhan and Shah (2023) from the Journal of Harbin Engineering University gave an in-depth overview of those anti-spoofing mechanisms, in their study Email Spoofing: In Today's Era.

Their work systematically investigated the three core standards: SPF (Sender Policy Framework), DKIM (DomainKeys Identified Mail), and DMARC (Domain-based Message Authentication, Reporting & Conformance). They did empirical tests on ten major email service providers, including Gmail, Outlook, Yahoo, and ProtonMail, by sending them controlled spoofed messages and analyzing the outcomes. Their results showed that even when these protocols



were set up correctly, it was possible for spoofed emails to reach the user's inbox under specific conditions, such as domain misalignment or forwarding through third-party servers.

The study concluded that, although SPF, DKIM, and DMARC have improved email integrity, their effectiveness is greatly dependent on correct configuration and domain-level policy enforcement, both of which are frequently overlooked. Their work truly emphasized how essential it is to have intelligent, adaptive detection systems that can analyze behavioral anomalies rather than depending wholly on inflexible authentication rules.

2.4 Large-Scale Vulnerability Studies

More recently, researchers have shifted to large-scale empirical studies to uncover systemic weaknesses in email ecosystems.

In one of the most comprehensive analyses, Weak Links in Authentication Chains: A Large-Scale Analysis of Email Sender Spoofing Attacks, Shen et al. (2021) at Tsinghua University present an analysis that includes

. Their study revealed that the modern email infrastructure behaves like a fragile “chain of trust,” where a single weak link — whether in sending, receiving, forwarding, or user-interface rendering — can compromise the entire system. They discovered 14 new attack types that can bypass SPF, DKIM, and DMARC validations through inconsistent protocol implementations, forwarding services, and visual deception in mail clients. Testing on 30 popular email services and 23 clients demonstrated that all were vulnerable to at least one spoofing method.

The authors also introduced a “cocktail attack,” combining several weaknesses to create realistic spoofed emails that bypassed both Gmail's and Outlook's filters. Their findings underscored the urgent need for holistic, cross-layer security solutions beyond mere standard protocol checks.

In this light, Ma et al. (2024) further extended the line of research with FakeBehalf: Imperceptible Email Spoofing Attacks Against the Delegation Mechanism in Email Systems at the USENIX Security Symposium.

They found that the Sender field of the email delegation mechanism can be arbitrarily forged to masquerade as a legitimate delegate and can evade even high-end authentication systems. They did controlled experiments with 16 service providers and showed that almost half of them are vulnerable to “FakeBehalf” attacks. The study also performed a user survey with 50 participants, where 50% were misled by spoofed delegate information. Their work highlighted the human factor in spoofing detection, reinforcing that even secure systems may fail if users misjudge trust indicators.

2.5 AI-Driven Spoofing Detection and Hybrid Approaches

The latest generation of research involves the application of artificial intelligence and machine learning for dynamic detection of spoofing. Mane et al. proposed such a system in their paper Reliable Email Spoofing Detection using Enhanced Cybersecurity Approaches, 2025.

Their framework contained header analysis, anomaly detection, and machine learning models to identify mismatched or changed email attributes. The system was implemented using a web-based platform built with Flask and Node modules, thus providing easy deployment at the server level.

Unlike static authentication mechanisms, their model used a training dataset of legitimate and spoofed emails to classify threats in real time. The “Spoof Guardian” application demonstrated an accuracy rate above 95%, showing the potential of AI-driven cybersecurity tools. However, their approach primarily focused on detection within a controlled dataset and lacked adaptive learning capabilities to handle new spoofing techniques emerging over time.

This research serves as the foundation for Cybershield, which extends these principles with continuous learning and user feedback integration. Cybershield builds upon Mane et al.'s design but introduces context-aware classification, leveraging behavioral and temporal features in addition to header-based indicators. By combining forensic insights from Pilli et al. (2014) with protocol intelligence from Chauhan and Shah (2023), and awareness of delegation vulnerabilities from Ma et al. (2024), Cybershield represents a synthesis of past research into a unified, adaptive defense mechanism.



2.6 Summary of Literature Gaps

Following the above review, certain gaps can be noticed: Limitations of the protocol: SPF, DKIM, and DMARC can't completely stop spoofing because of misconfigurations or issues within email forwarding. 2. Lack of adaptive intelligence: The majority of the detection mechanisms rely on rules and are static; hence, they cannot keep up with the evolving spoofing tactics. 3. Forensic complexity: Manual header analysis is accurate but not feasible for real-time detection. 4. User unawareness: When systems flag suspicious emails, users often fail to interpret the warnings correctly. 5. Challenges with integration: Not many of the systems integrate AI-based detection with intuitive, deployable interfaces for either users or administrators. The proposed framework of Cybershield will address these limitations by integrating real-time machine learning.

III. PROBLEM DEFINITION AND SYSTEM OVERVIEW

A. Problem Definition

In today's digital communication landscape, email spoofing has emerged as one of the most common and dangerous forms of cyber deception. Attackers exploit weaknesses in the Simple Mail Transfer Protocol (SMTP) to forge sender identities and make emails appear to originate from trusted sources. These spoofed messages are frequently used in phishing campaigns, financial fraud, and credential theft, posing serious risks to both individuals and organizations.

Although mechanisms such as Sender Policy Framework (SPF), **DomainKeys Identified Mail (DKIM)**, and **Domain-based Message Authentication, Reporting and Conformance (DMARC)** were introduced to authenticate email origins, they have several limitations. These include **inconsistent configuration**, **lack of universal adoption**, **forwarding loopholes**, and **vulnerabilities in email delegation**. As a result, even properly configured mail servers remain susceptible to sophisticated spoofing attempts that bypass these checks.

The primary challenge lies in the **dynamic and evolving nature of spoofing techniques**, which cannot be fully addressed by static rule-based systems. Hence, there is a pressing need for an **intelligent, adaptive, and real-time detection mechanism** that goes beyond protocol verification to identify spoofing patterns from email behavior itself.

Therefore, the problem can be formally stated as:

"To design and implement a machine learning-based intelligent system capable of detecting email spoofing attacks through comprehensive header analysis and anomaly detection, thereby improving accuracy and reliability over traditional authentication mechanisms."

B. System Overview

To address this issue, the proposed system — **Cybershield** — integrates traditional email authentication with **machine learning-driven anomaly detection** to create a robust and adaptive defense mechanism. The system is designed to function as an intermediate security layer within a mail server or as a standalone detection application for email analysis.

The system workflow can be summarized as follows:

Email Header Extraction: When a new email is received, Cybershield extracts key metadata fields such as *From*, *Return-Path*, *Received*, *Reply-To*, and *Message-ID*.

Feature Analysis: These header attributes are compared against known behavioral and structural norms. Domain alignment, timestamp consistency, and IP mapping are evaluated for authenticity.

Machine Learning Classification: Extracted features are passed to a trained machine learning model — primarily **Random Forest** and **Support Vector Machine (SVM)** — to classify the email as *legitimate* or *spoofed*.

Result Generation: The classification result, along with a confidence score, is displayed to the user or mail administrator through a simple interface.

Adaptive Feedback: The model continuously learns from new samples and false classifications to improve accuracy over time.

The **system architecture** consists of five primary modules:

Input Module – accepts raw emails or header files.

Header Extractor – parses and structures header data.

Feature Analyzer – identifies and encodes relevant header-level features.

Copyright to IJAR SCT

www.ijarsct.co.in

DOI: 10.48175/IJAR SCT-29736



Machine Learning Engine – processes inputs using trained algorithms.

Reporting Module – generates detection reports and visual feedback.

This modular design ensures scalability and ease of integration with existing email infrastructures. The combination of rule-based checks and intelligent classification enables Cybershield to detect spoofing attempts that traditional methods overlook.

IV. METHODOLOGY

The proposed methodology for **Cybershield** focuses on the accurate and intelligent detection of spoofed emails using a hybrid combination of email header analysis and machine learning techniques. The system operates in multiple stages, from data collection to classification and reporting, ensuring both accuracy and adaptability.

A. Data Collection and Preprocessing

A dataset comprising **5,000 emails** (2,500 legitimate and 2,500 spoofed) was prepared for experimentation. Legitimate samples were collected from verified institutional and corporate mail servers, while spoofed samples were generated using controlled phishing simulations and public repositories such as *SpamAssassin* and *PhishTank*.

Each email's **header** was extracted using Python's email library. Key attributes such as *From*, *Received*, *Return-Path*, *Reply-To*, and *Message-ID* were isolated.

To ensure uniformity, irrelevant metadata and duplicate fields were removed. The data was then structured into a tabular format suitable for machine learning models.

B. Feature Extraction

In this stage, relevant header-level features were extracted to distinguish spoofed and legitimate emails.

Features included:

Domain alignment between *From* and *Return-Path*

Sender IP–domain consistency

Time stamp mismatch in *Received* fields

Unusual header sequence lengths

Missing or forged authentication fields

These features were encoded numerically using **label encoding** and **binary transformation** for compatibility with machine learning algorithms.

C. Model Training

The preprocessed dataset was split into **80% training** and **20% testing** subsets.

Two machine learning algorithms were evaluated:

Support Vector Machine (SVM) – for boundary-based classification.

Random Forest Classifier – for feature-level decision aggregation.

After testing, Random Forest achieved higher accuracy and stability due to its ensemble decision-making process. The model was trained using the **Scikit-learn** library in Python

D. Detection Workflow

When a new email is received, its header is parsed by the **Header Extractor Module**. The **Feature Analyzer** then compares the attributes against learned patterns from the training dataset.

The **ML Engine** computes a confidence score between 0 and 1, where a higher score indicates higher spoofing probability.

Finally, the **Reporting Module** labels the message as *Legitimate* or *Spoofed* and provides a confidence percentage to the administrator or user.



V. EXPERIMENTAL SETUP AND RESULTS

The implementation of **Cybershield: Email Spoofing Detection** was carried out with the objective of creating a reliable, adaptive, and lightweight system capable of identifying spoofed emails in real time. This section elaborates on the technologies used, data preparation, algorithmic configuration, and the testing environment that were utilized for the development and evaluation of the system.

A. System Environment

Cybershield was implemented using **Python 3.10** as the core programming language due to its rich ecosystem of machine learning and cybersecurity libraries. The web-based interface was developed using the **Flask framework**, which allowed the model to interact with a user-friendly dashboard for uploading and analyzing email samples. The experiments were conducted on a local machine with the following configuration:

Processor: Intel Core i5 (11th Gen)

RAM: 16 GB DDR4

Operating System: Windows 11

Storage: 512 GB SSD

This environment ensured fast execution, low latency, and sufficient memory for training models and processing large email datasets.

B. Tools and Libraries Used

The system was developed using the following major tools and libraries:

Scikit-learn: For implementing and training machine learning models such as Random Forest and Support Vector Machine (SVM).

Pandas and NumPy: For data preprocessing, feature extraction, and dataset manipulation.

Matplotlib: For visualizing accuracy, confusion matrices, and comparison graphs.

SQLite: Used as a lightweight database to store user-uploaded email headers and detection results.

Flask: For creating the REST-based web interface for real-time spoofing detection.

C. Dataset Preparation

The dataset used for training and testing Cybershield consisted of **5,000 email samples** — 2,500 legitimate and 2,500 spoofed. Legitimate emails were collected from verified institutional and corporate mail servers, while spoofed samples were obtained from phishing simulation frameworks and open-source datasets such as *SpamAssassin* and *PhishTank*. Each email was parsed using Python's built-in email library to extract header fields like *From*, *Return-Path*, *Reply-To*, and *Received*. Non-essential metadata was removed, and missing values were handled using mean and mode imputation.

The dataset was then divided into **80% training** and **20% testing** subsets for model validation.

D. Machine Learning Implementation

The processed dataset was passed through two classification algorithms — **Support Vector Machine (SVM)** and **Random Forest (RF)** — using the Scikit-learn library.

The **SVM model** used a radial basis function (RBF) kernel with grid search optimization for parameter tuning.

The **Random Forest classifier**, which eventually became the final model, was configured with 100 decision trees ($n_estimators = 100$) and used Gini impurity as the splitting criterion.

The Random Forest model demonstrated superior accuracy and resilience to overfitting, making it ideal for practical deployment.

E. Integration and User Interface

Cybershield integrates all components through a web-based dashboard built on Flask. The dashboard allows users to upload email header files (.eml format) or paste raw headers directly. Upon submission, the system performs real-time analysis through the ML engine and displays:

Copyright to IJARSCT

www.ijarsct.co.in



DOI: 10.48175/IJARSCT-29736



273

Spoofing detection result (*Legitimate* or *Spoofed*)

Confidence score (0–100%)

Highlighted mismatched fields for manual inspection

The interface is lightweight and responsive, designed for both desktop and server-side deployment.

F. Experimental Setup and Testing

To validate the system's accuracy and reliability, several test runs were conducted using unseen email samples from both corporate and phishing domains. Performance metrics — **accuracy, precision, recall, and F1-score** — were calculated to evaluate the classification results.

A confusion matrix was generated to visualize correct and incorrect predictions. The model achieved the following results:

Accuracy: 96.4%

Precision: 0.96

Recall: 0.96

F1-score: 0.96

The experimental results confirmed that Cybershield performs consistently across varied datasets and is capable of detecting spoofed messages that bypass SPF, DKIM, and DMARC validation.

VI. RESULTS AND DISCUSSIONS

The performance of the proposed **Cybershield** system was evaluated using multiple metrics including accuracy, precision, recall, and F1-score. The model was trained and tested using the dataset described in the implementation section, which included an equal number of legitimate and spoofed emails. The results were compared with existing standard authentication-based systems such as **SPF (Sender Policy Framework)** and **DKIM (DomainKeys Identified Mail)** to demonstrate the improvement achieved by integrating machine learning techniques.

A. Performance Metrics

To assess the detection capability of Cybershield, four evaluation metrics were considered:

Accuracy: Measures the overall correctness of the classification.

Precision: Indicates how many predicted spoofed emails were actually spoofed.

Recall: Measures how effectively the system identifies all spoofed messages.

F1-Score: Represents the harmonic mean of precision and recall, providing a balance between the two

The following table summarizes the comparative results:

Approach	Accuracy (%)	Precision	Recall	F1-Score
SPF / DKIM Verification	85.2	0.83	0.82	0.82
Support Vector Machine (SVM)	93.7	0.94	0.92	0.93
Cybershield (Random Forest)	96.4	0.96	0.96	0.96

The table indicates that Cybershield significantly improves upon traditional authentication mechanisms by leveraging header-level learning and anomaly detection.

B. Discussion of Results

The results clearly show that Cybershield achieved a high level of accuracy and reliability. Unlike SPF and DKIM, which rely solely on static DNS and cryptographic verification, Cybershield analyzes behavioral and structural inconsistencies within email headers. This allows it to detect spoofed emails that traditional methods fail to identify, especially in cases of **forwarding, delegation, and domain misalignment**.

The **Random Forest algorithm** provided robustness due to its ensemble learning approach, minimizing the chances of overfitting and improving prediction stability. The system's **adaptive learning** capability enables it to update its model as new spoofing techniques emerge, ensuring long-term reliability.



The confusion matrix generated during evaluation revealed that false positives (legitimate emails marked as spoofed) were minimal — below **3.5%** — which is acceptable for operational environments. This confirms the system's practicality for integration in enterprise mail servers.

C. Practical Implications

From a practical standpoint, the Cybershield system can be integrated into existing **email gateways** and **mail servers** with minimal overhead. Since it uses lightweight header analysis and classification rather than full content scanning, it offers faster execution and reduced computational load.

Organizations implementing Cybershield can benefit from an **additional defense layer**, capable of flagging spoofed emails that pass conventional validation checks. Moreover, with periodic retraining, the system remains effective against evolving spoofing methods and newly registered malicious domains.

VII. CONCLUSION AND FUTURE SCOPE

This research presented **Cybershield**, a machine learning-driven system for the detection of email spoofing through comprehensive header analysis. The system effectively identifies forged sender addresses and mismatched metadata that often go undetected by traditional authentication protocols such as SPF, DKIM, and DMARC.

By integrating intelligent header evaluation with machine learning algorithms—specifically the **Random Forest Classifier**—Cybershield demonstrated an overall **accuracy of 96.4%**, significantly outperforming static rule-based verification methods. The system's design emphasizes scalability, adaptability, and real-time detection, ensuring that it can function efficiently across diverse email infrastructures.

One of Cybershield's major advantages is its **adaptive learning capability**. Unlike static filtering systems, it retrains periodically using new datasets and user feedback. This allows the model to evolve as spoofing techniques advance, maintaining high accuracy over time. Additionally, the use of lightweight header features reduces computational cost, making the system suitable for both enterprise-level mail servers and smaller organizational setups.

In conclusion, Cybershield proves that combining **forensic header analysis with artificial intelligence** offers a powerful solution to modern email security challenges. The results validate its potential as a practical cybersecurity enhancement tool capable of mitigating spoofing-based social engineering and phishing attacks.

Future Scope

While the current version of Cybershield performs exceptionally well in header-based spoofing detection, several areas remain open for improvement:

Integration with Content-Based Detection: Incorporating **Natural Language Processing (NLP)** to analyze email body text and detect phishing or social engineering cues.

Cloud-Based Deployment: Developing a cloud API version for seamless integration with popular platforms like Gmail, Outlook, and enterprise mail servers.

User Feedback Mechanism: Implementing a continuous feedback and retraining system to enhance accuracy and adapt to emerging spoofing tactics.

Hybrid Security Layer: Combining Cybershield's ML engine with traditional SPF/DKIM/DMARC checks for multi-layered security.

With these future enhancements, Cybershield has the potential to become a **comprehensive and intelligent anti-spoofing framework**, capable of defending against evolving email threats in real-world communication systems.

REFERENCES

- [1] K. Pandove, A. Jindal, and R. Kumar, "Email Spoofing," *International Journal of Computer Applications*, vol. 5, no. 1, pp. 27–31, 2010. <https://www.ijcaonline.org/archives/volume5/number1/974-1320>
- [2] S. Pilli, D. Mishra, and R. Joshi, "Forensic Analysis of Email Address Spoofing," *IEEE International Conference on Digital Security and Forensics*, 2014. <https://ieeexplore.ieee.org/document/6918812>
- [3] R. Chauhan and P. Shah, "Email Spoofing: In Today's Era," *Journal of Harbin Engineering University*, 2023. <https://harbinengineeringjournal.com/article/view/3625>



- [4] J. Shen, J. Jiang, and B. Zhao, "Weak Links in Authentication Chains: A Large-Scale Analysis of Email Spoofing Attacks," *USENIX Security Symposium*, 2021. <https://www.usenix.org/conference/usenixsecurity21/presentation/shen>
- [5] J. Ma, R. Jinrui, and H. Jin, "FakeBehalf: Imperceptible Email Spoofing Attacks and Defenses," *USENIX Security Symposium*, 2024. <https://www.usenix.org/conference/usenixsecurity24/presentation/ma-jinrui>
- [6] R. Mane, A. Joshi, and A. Patil, "Reliable Email Spoofing Detection Using Enhanced Cybersecurity Approaches," *International Journal of Computer Applications (IJCA)*, 2025

