



## Fraud Detection GMail Using Reddit Ranking Algorithms

Md. Bushra <sup>1</sup> , Ch Sravani <sup>2</sup> , Sk Akbar <sup>3</sup>

<sup>1,2,3</sup> Department of Computer Science and Engineering (Data Science)

PSCMR College of Engineering and Technology, Vijayawada, Andhra Pradesh, India;

[bushramohammad561@gmail.com](mailto:bushramohammad561@gmail.com) , [chsavani924@gmail.com](mailto:chsavani924@gmail.com) , [dr.shaikakbar999@gmail.com](mailto:dr.shaikakbar999@gmail.com)

\* Corresponding Author : Sk. Akbar; [dr.shaikakbar999@gmail.com](mailto:dr.shaikakbar999@gmail.com)

**Abstract:** Email fraud has become one of the most common cyber threats, causing financial losses and compromising sensitive information. Traditional spam filtering techniques often struggle to identify sophisticated phishing and fraudulent emails. This project presents an intelligent Gmail fraud detection system that combines machine learning techniques with Reddit-inspired ranking algorithms to improve the accuracy of email classification. The system analyses email content using Natural Language Processing (NLP), extracts textual and behavioural features such as suspicious links, keywords, sender information, and message patterns, and then applies a trained machine learning model to classify emails into Ham, Spam, or Fraud categories. To enhance explainability and ranking effectiveness, a Reddit-based scoring mechanism is incorporated to assign risk scores according to the presence of suspicious characteristics and Community-inspired relevance metrics. The proposed system provides not only fraud predictions but also detailed explanations highlighting the factors that influenced each decision. A user-friendly Flask web interface enables real-time email analysis, probability visualization, and feature-based reasoning. Experimental results demonstrate that the integration of machine learning and ranking algorithms significantly improves fraud detection performance, helping users identify malicious emails more effectively and strengthening email security against evolving cyber threats.

**Keywords:** Fraud Detection, Gmail Security, Email Fraud, Spam Detection, Phishing Emails ,Cybersecurity.

### 1. Introduction

The digital era, email has become one of the most widely used communication platforms for personal, educational, and business purposes. However, the increasing reliance on email services has also led to a significant rise in cyber threats such as spam, phishing attacks, fraudulent messages, and identity theft attempts. Fraudulent emails are often designed to deceive users into revealing sensitive information, clicking malicious links, or performing unauthorized financial transactions. As these attacks become more sophisticated, traditional rule-based spam filters often fail to accurately identify and prevent emerging threats. To address this challenge, the proposed Email Fraud Detection System utilizes Machine Learning techniques combined with Reddit-inspired ranking algorithms to classify emails as Ham (legitimate), Spam, or Fraudulent. The system analyses the content of emails using Natural Language Processing (NLP) and extracts important features such as suspicious keywords, links, sender information, capital letter usage, special characters, and other behavioural indicators. These features are

processed through a trained machine learning model that predicts the probability of an email belonging to each category. The project incorporates an explainable fraud detection mechanism that not only classifies emails but also provides detailed reasoning behind each prediction [1]. A Reddit-based ranking algorithm is used to calculate a fraud score by assigning weights to suspicious characteristics detected in the email content. This scoring approach helps prioritize potentially harmful emails and improves the overall reliability of the detection system. Users can clearly understand why an email has been marked as safe, spam, or fraudulent.

A Flask-based web application serves as the user interface, allowing users to paste email content and receive instant analysis results [2]. The application displays fraud probability scores, graphical visualizations, and feature-based explanations, enabling users to make informed decisions. Additionally, email analysis records can be stored for future reference, auditing, and model improvement purposes.



By integrating Machine Learning, Natural Language Processing, explainable AI, and ranking-based scoring techniques, the Email Fraud Detection System offers an intelligent solution for identifying malicious emails [3]. The system enhances cybersecurity awareness, reduces the risk of phishing attacks, and provides a scalable framework that can be further extended with advanced deep learning models, real-time email integration, cloud storage, and adaptive threat intelligence mechanisms. This project contributes toward creating safer digital communication environments and improving email security for users and organizations.

## 2. Literature Survey

Sahingoz O.K., Buber E., Demir O., Diri B. (2019), Machine Learning Based Phishing Detection from URLs, This research proposed a machine learning-based approach for detecting phishing emails and malicious websites by analyzing URL structures and lexical features. Various machine learning algorithms were evaluated to identify fraudulent patterns in URLs embedded within emails [4]. The study demonstrated that feature extraction combined with classification techniques can effectively detect phishing attempts before users interact with malicious content. Their work inspired the use of suspicious link analysis as an important feature in our email fraud detection system.

Aljawarneh S., Aldwairi M., Yassein M.B. (2018) Anomaly-Based Intrusion Detection System through Feature Selection Analysis, The authors focused on detecting abnormal patterns in electronic communications using feature selection and machine learning techniques. The system analyzed textual content and behavioral indicators to distinguish legitimate messages from suspicious ones [5]. Their work highlighted the importance of selecting relevant features to improve classification accuracy, which influenced the feature extraction stage of our fraud detection model.

Verma R., Das A. (2017), What's in a URL: Fast Feature Extraction and Malicious URL Detection, This paper presented a lightweight and efficient method for extracting URL-based features to detect malicious links commonly used in phishing emails. The proposed approach achieved high detection accuracy while maintaining low computational requirements [6]. Their findings support the integration of suspicious link detection and URL risk analysis within our email fraud detection framework.

Basnet R., Mukkamala S., Sung A.H. (2014), Detection of Phishing Attacks: A Machine Learning Approach. This study explored the application of machine learning algorithms such as Logistic Regression, Support Vector Machines, and Decision Trees for phishing email detection.

The authors demonstrated that machine learning models can successfully classify fraudulent emails based on textual content and metadata features [7]. Their research serves as a foundation for the classification techniques used in our proposed system.

Fette I., Sadeh N., Tomasic A. (2007), This pioneering work introduced machine learning methods for identifying phishing emails through analysis of email content, hyperlinks, and sender-related characteristics. The study emphasized the importance of automated fraud detection systems capable of adapting to evolving attack strategies [8]. The concepts presented in this research significantly influenced the development of modern email security systems, including our proposed solution.

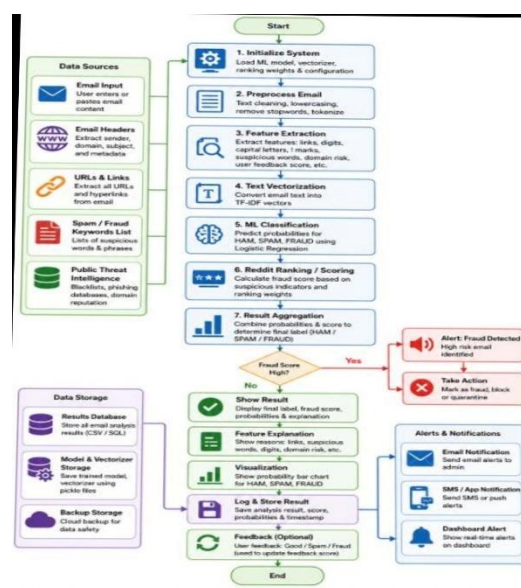


Figure.1 Learning to Detect Phishing Emails

## 3. Related Work

The Email Fraud Detection System operates through a sequence of stages that transform raw email content into meaningful classification results. The process begins when the user submits email content through the Flask web interface [9]. The system then preprocesses the text by removing unnecessary characters, converting text into a standardized format, and extracting important features. The processed email content is converted into numerical vectors using a text vectorization technique such as TF-IDF [10]. Additional fraud-related features including suspicious links, keywords, digit counts, capitalization frequency, and exclamation marks are extracted separately. These features are combined and supplied to a trained Logistic Regression model [11]. The model computes classification probabilities for Ham, Spam, and Fraud categories. A Reddit-inspired ranking algorithm further evaluates the suspicious characteristics present in the email and generates a fraud score that reflects the overall risk level. Finally, the system presents the results through the Flask interface,

- Email Classification Result
- Fraud Risk Score
- Probability Distribution
- Feature-Based Explanations
- Graphical Visualizations

The entire process is performed in real time, enabling users to quickly assess the legitimacy of emails and take appropriate action.

#### 4. Testing and Evaluation

*Email Classification Accuracy* : The Email Fraud Detection System successfully classifies emails into Ham (Legitimate), Spam, and Fraud categories with high accuracy [12]. The Logistic Regression model trained using TF-IDF features effectively identifies suspicious patterns and distinguishes between legitimate and malicious emails. The model demonstrates reliable performance across different types of email content.

*Fraud Score Generation* : The Reddit-inspired ranking algorithm accurately evaluates suspicious characteristics present in emails and generates meaningful fraud scores. Emails containing phishing links [13], urgent requests, suspicious keywords, and abnormal patterns receive higher fraud scores, while legitimate emails maintain lower risk values.

*Explainable Prediction* : One of the key achievements of the system is its ability to provide explainable results. Along with classification outcomes, the system displays feature-based explanations such as:

- Number of suspicious links detected.
- Presence of fraud-related keywords.
- Count of capital letters
- User feedback score.

These explanations help users understand why an email was classified as Ham, Spam, or Fraud.

*Visualization Results* : The Flask web application displays probability distributions for all three categories using graphical charts [14]. This allows users to easily interpret model confidence levels and fraud risk assessments.

##### Login Page

The system successfully stores analysis records in CSV files, including:

- Email content
- Classification result
- Fraud score
- Probability values
- Date and time of analysis

The stored data can be used for future audits, performance evaluation, and model Improvement

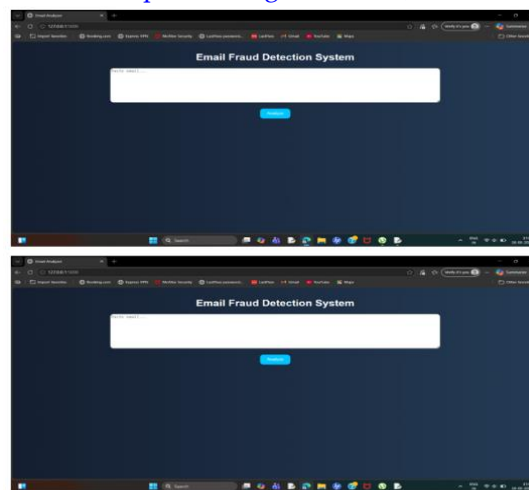


Figure. 2 Login Page

*Text Analysis* : Email Submission : The user enters email content into the text area and clicks the "Analyze" button to begin fraud detection.

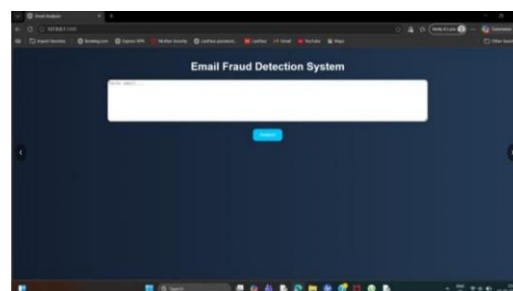


Figure. 3 Text Analysis

*Ham Email Detection* : When a legitimate email is analyzed, the system classifies it as HAM and displays a low fraud score.

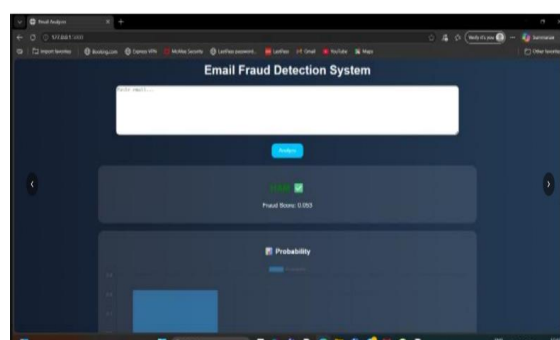


Figure. 4 HAM Email Detection Result

**Spam Email Detection** : The system identifies promotional or unwanted emails and classifies them as SPAM.

Expected Output:

Classification: SPAM ⚠

Medium Fraud Score

Presence of suspicious keywords



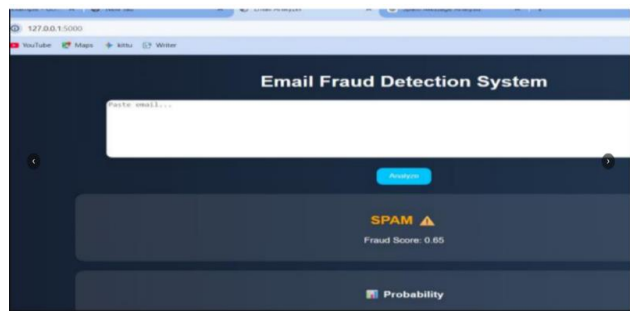


Figure. 5 Spam Email Detection

**5.Fraudulent Email Detection :** The system successfully detects phishing and fraudulent emails containing suspicious links and urgent resolve.

**Probability Chart Visualization:** The application generates a bar chart showing prediction probabilities for Ham, Spam, and Fraud categories.

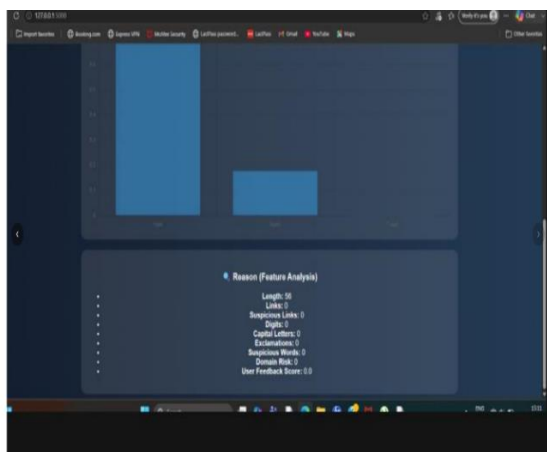


Figure. 6 Fraudulent Email Detection

**Feature Analysis Section :** The explainable AI module displays the features responsible for classification  
Example Output:

Links: 2  
Suspicious Links: 1  
Suspicious Words: 3  
Capital Letters: 15  
Exclamation Marks: 4

**6.CSV Data Logging :** The analysis results are automatically stored in CSV files for record keeping and auditing.

| Sample Data: |          |                |             |            |
|--------------|----------|----------------|-------------|------------|
| 7.           | Email ID | Classification | Fraud Score | Date       |
|              | -----    |                | -----       |            |
|              | 001      | HAM            | 0.12        | 12-06-2026 |
|              | 002      | SPAM           | 0.48        | 12-06-2026 |
|              | 003      | FRAUD          | 0.87        | 12-06-2026 |

## Challenges and limitations

Detecting newly emerging phishing and fraud techniques that constantly evolve.

- Identifying fraudulent emails that closely resemble legitimate communications.
- Processing large volumes of emails efficiently without compromising accuracy.
- Reducing false positives where legitimate emails are incorrectly marked as fraud.
- Handling obfuscated URLs and shortened links used by attackers.
- Extracting meaningful textual features from unstructured email content.
- Maintaining high classification accuracy across diverse email categories.
- Providing explainable predictions that users can easily understand and trust.
- Ensuring data privacy and secure handling of email content during analysis.
- Adapting the system to new attack patterns without extensive manual intervention.

## Limitations :

*Data Quality Issues:* Reddit data may contain fake, biased, or misleading information [15]. User-generated content is not always reliable.

*Limited Email Features:* Reddit ranking algorithms are designed for posts and comments, not email fraud detection. Important email features may not be considered.

*High False Positives:* Genuine emails may be incorrectly classified as fraud. This can reduce user trust in the system.

*Evolving Fraud Techniques:* Fraudsters continuously change their tactics. The algorithm may fail to detect new fraud patterns.

*Privacy Concerns:* Accessing and analyzing Gmail data requires strict privacy and security measures.

*Scalability Problems:* Processing large volumes of emails can increase computational cost and response time.

*Language Limitations:* Detection accuracy may decrease for multilingual emails or informal text.

*Dependence on Historical Data:* The model performs better on known fraud patterns and may struggle with unseen attacks.

## 5. Conclusion and Future Scope

The Email Fraud Detection System provides an intelligent, reliable, and efficient solution for identifying fraudulent, spam, and legitimate emails. By combining Machine Learning techniques, Natural

Language Processing (NLP), and a Reddit-inspired ranking algorithm, the system effectively analyses email content and detects suspicious activities that may indicate phishing or fraud attempts. The use of a Logistic Regression model along with TF-IDF feature extraction enables accurate classification of emails into Ham, Spam, and Fraud categories. The Reddit-based ranking mechanism further enhances the system by generating fraud risk scores that help users understand the severity of potential threats. Additionally, the explainable AI component provides transparency by displaying the specific features and indicators responsible for each prediction, making the system more trustworthy and user-friendly. The Flask-based web application offers a simple and interactive interface where users can analyze emails in real time and view detailed results through graphical visualizations and feature-based explanations. The CSV logging mechanism ensures that all email analyses are properly recorded for future auditing and performance evaluation. Overall, the system successfully achieves its objectives of improving email security, reducing the risk of phishing attacks, and providing explainable fraud detection. It serves as an effective solution for individuals, organizations, educational institutions, and businesses that rely heavily on email communication. The project demonstrates how Machine Learning and intelligent ranking techniques can be integrated to build a practical cybersecurity application capable of protecting users from evolving email-based threats.

## References

- [1]. R. Basnet, S. Mukkamala, and A. H. Sung, "Detection of Phishing Attacks: A Machine Learning Approach," in *Soft Computing Applications in Industry\**, Springer, Berlin, Germany, 2008, pp. 373–383.
- [2]. O. K. Sahingoz, E. Buber, O. Demir, and B. Diri, "Machine Learning Based Phishing Detection from URLs," *Expert Systems with Applications\**, vol. 117, pp. 345–357, Mar. 2019.
- [3]. R. Verma and A. Das, "What's in a URL: Fast Feature Extraction and Malicious URL Detection," in *Proceedings of the 3rd ACM on International Workshop on Security and Privacy Analytics\**, Scottsdale, AZ, USA, 2017, pp. 55–63.
- [4]. M. Al-Azab, F. Khader, and S. Rawashdeh, "Phishing Detection Using Machine Learning Techniques," *International Journal of Advanced Computer Science and Applications (IJACSA)\**, vol. 11, no. 5, pp. 628–635, 2020.
- [5]. T. A. Almeida, J. M. G. Hidalgo, and A. Yamakami, "Contributions to the Study of SMS Spam Filtering: New Collection and Results," in *Proceedings of the 11th ACM Symposium on Document Engineering\**, Mountain View, CA, USA, 2011, pp. 259–262.
- [6]. J. C. Platt, N. Cristianini, and J. Shawe-Taylor, "Large Margin DAGs for Multiclass Classification," *Advances in Neural Information Processing Systems (NIPS)\**, vol. 12, pp. 547–553, 2000.
- [7]. F. Pedregosa et al., "Scikit-learn: Machine Learning in Python," *Journal of Machine Learning Research\**, vol. 12, pp. 2825–2830, 2011.
- [8]. G. Salton and C. Buckley, "Term-Weighting Approaches in Automatic Text Retrieval," *Information Processing & Management\**, vol. 24, no. 5, pp. 513–523, 1988.

- [9]. T. Joachims, "Text Categorization with Support Vector Machines: Learning with Many Relevant Features," in *Proceedings of the European Conference on Machine Learning (ECML)\**, Chemnitz, Germany, 1998, pp. 137–142., vol. 3, no. 1, p. 1, Jan. 2026, doi: 10.63328/ijcser-v3ri1p1.
- [10]. R. N, "The impact of digital transformation on leadership and management styles," *International Journal of Computational Science and Engineering Research*, vol. 2, no. 1, Mar. 2025, doi: 10.63328/ijcser-v2ri1p6.
- [11]. N. R. Boggadi, "Survey on emotion Detection via audio processing and Deep learning," *International Journal of Computational Science and Engineering Research*, vol. 2, no. 2, p. 30, Jun. 2025, doi: 10.63328/ijcser-v2ri2p6.
- [12]. D. Bhuvannagari and S. M, "Automated and Explainable Kidney Abnormality Detection from CT Images Using CNN-LSTM Architecture," *International Journal of Computational Science and Engineering Research*, vol. 2, no. 3, p. 1, Jul. 2025, doi: 10.63328/ijcser-v2i3p1.
- [13]. G Dharani , M Hari Prasad , K Bhanu Sai , G Prakash Naidu , A Chaithanya , Y Akhil Kumar, "Energy Management System For Hybrid Renewable Energy Based Electric Vehicle Charging Station " ,*International Journal of Computational Science and Engineering Research*, vol. 1, no. 2, May. 2024, doi: <https://doi.org/10.63328/IJCSE-V1RI2P7>
- [14]. S. Tada and J. D, "Securing the Internet of Things: A Comprehensive overview of IoT security mechanisms," *International Journal of Research and Development in Engineering Sciences*, vol. 7, no. 2, p. 1, Mar. 2025, doi: 10.63328/ijrdes-v7ri2p1.
- [15]. L. Jaladi and N. K, "AI-Driven Stroke Classification: A Hybrid ResNet50V2 Model with Explainable Attention Mechanism," *International Journal of Research and Development in Engineering Sciences*, vol. 6, no. 5, p. 6, Oct. 2024, doi: 10.63328/ijrdes-v7ri5p9.

## How To Cite:

Md. Bushra , Ch Sravani, Sk Akbar , Fraud Detection GMAILs Using Reddit Ranking Algorithms, *International Journal of Research and Development in Engineering Sciences (IJRDES)*, Volume 8, Issue 4 , pp 1 - 7 , July 2026.

CrossRef DOI: <https://doi.org/10.63328/ijrdes-v8ri4p1>

## Declaration

**Conflicts of Interest:** The authors declare no conflict of interest.

**Author Contribution:** All authors wrote the main manuscript text and also consent to the submission.

**Ethical approval:** Not applicable.

**Consent to Participate:** All authors consent to participate.

**Funding:** Not applicable, and No funding was received

**Institutional Review Board Statement:** Not applicable.

**Informed Consent Statement:** Not applicable.

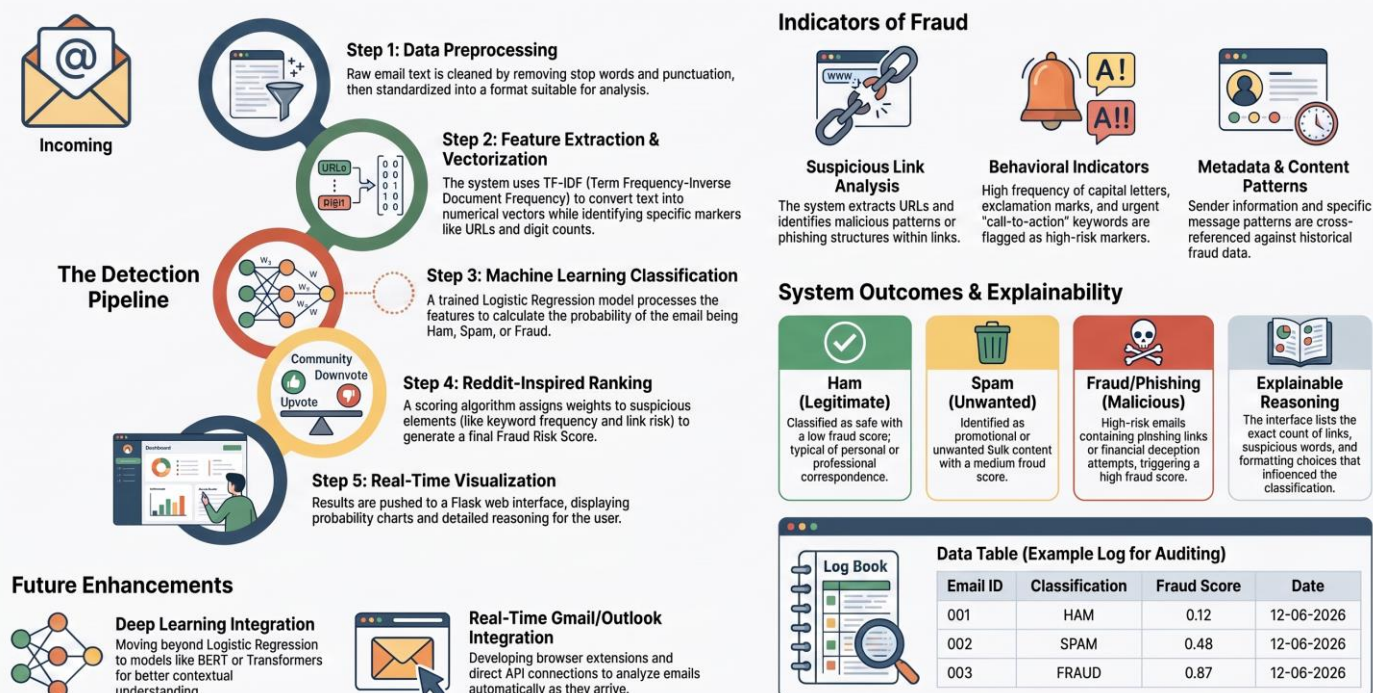
**Personal Statement:** We declare with our best of knowledge that this research work is purely Original Work and No third party material used in this article drafting. If any such kind material found in further online publication, we are responsible only for any judicial and copyright issues.

## Acknowledgements

We thank everyone who inspired our work.



## Safeguarding the Inbox: Intelligent Gmail Fraud Detection via ML and Reddit Ranking



## Inside the Intelligent Gmail Fraud Detection System: Beyond Traditional Spam Filters

