



International Journal of Engineering Research and Science & Technology

www.ijerst.org

ISSN : 2319-5991

Vol. 22 No. 2(2) (2026)



ijerst.editor@gmail.com
editor@ijerst.com

Research Paper

E-Mail Spam Detection Using Machine Learning

A. Reshma¹, S. Haritha²

¹ Assistant Professor, Department of MCA, Audisankara College of Engineering & Technology
(UGC-Autonomous Institution),
Nh-5, Bypass Road Gudur Tirupati Dist. Andhra Pradesh, India

² Student, Department of MCA., Audisankara College of Engineering & Technology
(UGC-Autonomous Institution)
Nh-5, Bypass Road Gudur Tirupati Dist. Andhra Pradesh, India

Abstract- *The rapid proliferation of digital communication has led to an exponential increase in unsolicited bulk e-mails, commonly known as spam, which poses significant security risks and resource drains. Traditional rule-based filtering methods often fail to keep pace with the evolving tactics of spammers, necessitating more robust, intelligent solutions. This project explores the application of Machine Learning (ML) algorithms to automate the detection and filtration of spam e-mails with high accuracy. By leveraging natural language processing (NLP) techniques, raw e-mail data is preprocessed through tokenization, stop-word removal, and TF-IDF (Term Frequency-Inverse Document Frequency) vectorization to extract meaningful features. Several supervised learning models, including Naive Bayes, Support Vector Machines (SVM), and Random Forest, are evaluated based on their ability to classify messages as "ham" (legitimate) or "spam." The performance of these models is measured using key metrics such as precision, recall, and the F1-score to ensure a low rate of false positives, which is critical for user experience. Experimental results demonstrate that probabilistic models like Naive Bayes are particularly effective for text classification due to their computational efficiency and high accuracy. However, ensemble methods often provide superior robustness against sophisticated adversarial attacks. This study concludes that an optimized ML pipeline significantly reduces the manual overhead of e-mail management while enhancing overall cybersecurity. The findings highlight the importance of continuous model retraining to adapt to the dynamic nature of spam content and sender behavior in the modern digital landscape.*

Keywords- *Machine Learning, Spam Filtering, Natural Language Processing (NLP), Naive Bayes, Supervised Learning, Feature Extraction, Text Classification, Cybersecurity.*

I. INTRODUCTION

In recent years, the rapid expansion of digital communication has significantly increased the reliance on e-mail as a primary mode of information exchange across personal and professional domains. However, this growth has also led to a substantial rise in unsolicited and malicious e-mails, commonly referred to as spam. These spam messages not only degrade user experience but also pose serious cybersecurity threats such as phishing attempts, malware distribution, and data breaches. Traditional rule-based filtering systems, which depend on manually defined keywords and

patterns, are increasingly inadequate in handling the evolving strategies used by spammers. To address these limitations, Machine Learning (ML) techniques have emerged as a powerful and adaptive solution for spam detection. ML-based approaches enable systems to learn patterns directly from data and improve classification performance over time without explicit programming for every rule. In this context, Natural Language Processing (NLP) plays a crucial role in transforming unstructured e-mail content into structured numerical representations suitable for model training. Common preprocessing techniques

such as tokenization, stop-word removal, and TF-IDF vectorization are widely used to extract meaningful features from raw text. These features are then utilized by supervised learning algorithms including Naive Bayes, Support Vector Machines (SVM), and Random Forest classifiers to distinguish between legitimate e-mails (ham) and spam messages. Each model offers distinct advantages in terms of accuracy, computational efficiency, and robustness against noisy data. The performance of these classification models is typically evaluated using metrics such as precision, recall, and F1-score, which provide a balanced assessment of detection capability, especially in minimizing false positives. Among these, Naive Bayes is often preferred for text classification due to its simplicity and high efficiency, while ensemble-based methods demonstrate improved resilience against complex and adversarial spam patterns. This study emphasizes the development of an optimized ML-based spam detection framework that enhances filtering accuracy while reducing manual intervention. The proposed approach highlights the importance of continuous learning and periodic model retraining to adapt to the dynamic and evolving nature of spam content in modern communication systems.

II. LITERATURE SURVEY

The problem of email spam detection has been extensively studied over the past two decades, evolving from simple rule-based systems to advanced machine learning and natural language processing (NLP) approaches. This section reviews key contributions that form the foundation of the proposed work. Early research by **McCallum and Nigam** demonstrated the effectiveness of probabilistic models, particularly Naïve Bayes, for text classification tasks. Their study highlighted how event-based models can efficiently handle high-dimensional textual data, making them suitable for spam filtering applications. Similarly, **Joachims** introduced Support Vector Machines (SVM) for text categorization and showed that SVMs perform exceptionally well in high-dimensional feature spaces. This work

established SVM as a strong discriminative classifier for spam detection due to its ability to construct optimal decision boundaries. The work of **Sahami et al.** further emphasized Bayesian approaches for filtering junk emails. Their research demonstrated that probabilistic reasoning can effectively distinguish between spam and legitimate emails based on word occurrence patterns. A broader comparison of Bayesian classifiers was presented by **John and Langley** where continuous distribution modeling improved classification performance in uncertain environments. This contributed to the development of more robust spam detection pipelines. In addition, **Quinlan [6]** introduced decision tree learning, which later influenced ensemble-based methods such as Random Forest. Building on this, **Breiman [7]** proposed Random Forests, which significantly improved classification accuracy through the combination of multiple decision trees, reducing overfitting and increasing robustness in noisy datasets. **Vapnik** laid the theoretical foundation for statistical learning theory and Support Vector Machines. His work provided the mathematical basis for margin maximization, which is crucial in separating spam and ham emails in high-dimensional spaces. **Raschka [9]** provided a practical analysis of Naïve Bayes in text classification, showing that despite its simplicity and independence assumption, it performs competitively in real-world NLP tasks such as spam detection. **Mitchell** offered a comprehensive introduction to machine learning, outlining supervised learning frameworks that form the backbone of modern spam classification systems. A detailed survey by **Aggarwal and Zhai** summarized various text classification algorithms, including probabilistic, linear, and ensemble models, highlighting their strengths and limitations in handling large-scale textual datasets. Further, **Soman et al.** demonstrated the application of NLP techniques combined with machine learning models for spam email detection, reinforcing the importance of preprocessing and feature extraction techniques like TF-IDF.

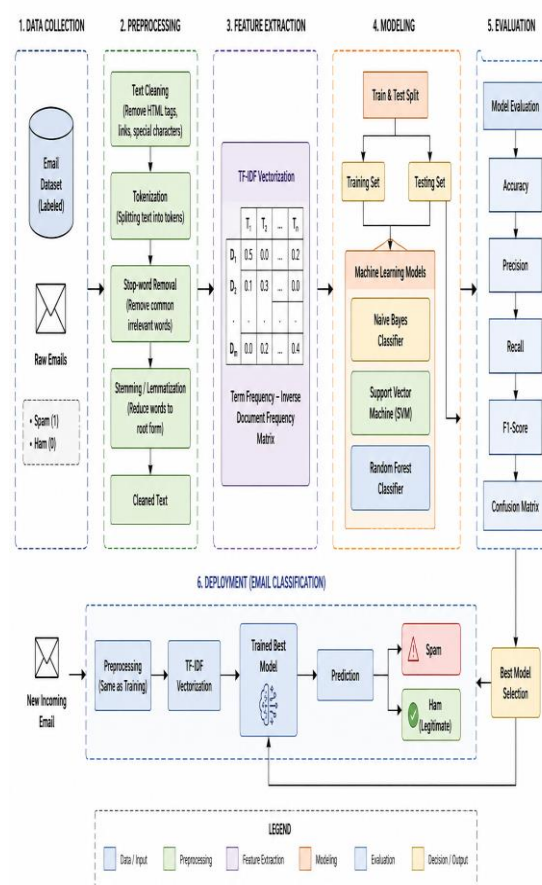
Sebastiani provided an extensive review of automated text categorization methods, emphasizing that machine learning approaches outperform manual rule-based systems in scalability and adaptability. Finally, **Drucker et al.** evaluated Support Vector Machines for spam classification and confirmed their superior performance compared to traditional algorithms, especially in handling imbalanced datasets and high-dimensional feature spaces.

III. PROPOSED SYSTEM

The proposed system introduces an intelligent and automated Email Spam Detection framework designed to classify incoming emails as either spam or ham (legitimate) using Machine Learning (ML) and Natural Language Processing (NLP) techniques. The system aims to overcome the limitations of traditional rule-based filtering approaches, which are often ineffective against evolving spam tactics and lack adaptability. The architecture of the proposed system begins with the collection of raw email datasets containing both spam and legitimate messages. These emails undergo a comprehensive preprocessing stage, where noise such as special characters, HTML tags, punctuation, and irrelevant symbols are removed. The textual content is then normalized through lowercasing, tokenization, stop-word removal, and stemming or lemmatization to ensure uniformity in data representation. After preprocessing, the system applies feature extraction using Term Frequency–Inverse Document Frequency (TF-IDF), which converts textual data into a structured numerical format. TF-IDF helps in identifying the importance of words in each email relative to the entire dataset, thereby improving the discriminative capability of the model. The processed feature vectors are then passed to multiple supervised Machine Learning classifiers, including Naive Bayes, Support Vector Machine (SVM), and Random Forest. Each model is trained on labeled email data to learn patterns associated with spam and ham messages. The system follows a modular pipeline consisting of data acquisition,

preprocessing, feature extraction, model training, model evaluation, and prediction. During the evaluation phase, performance metrics such as accuracy, precision, recall, and F1-score are computed to assess model effectiveness. Special emphasis is placed on minimizing false positives, as misclassifying legitimate emails as spam can negatively impact user experience and communication flow. To enhance adaptability, the proposed system supports periodic retraining using updated datasets, allowing it to evolve with changing spam patterns and adversarial techniques. The final deployed model is integrated into an email filtering interface that automatically classifies incoming emails in real time and filters them accordingly.

IV. METHODOLOGY



The proposed Email Spam Detection system follows a supervised machine learning pipeline designed to accurately classify incoming emails as either *spam* or *ham (legitimate)*. The methodology integrates Natural Language Processing (NLP) techniques with multiple

classification algorithms to improve detection performance and reduce false positives.

A. Data Collection and Dataset Preparation

The system begins with the collection of labeled email datasets containing both spam and non-spam messages. Each email instance includes textual content along with its corresponding class label. The dataset is then organized and cleaned to remove duplicates, null values, and irrelevant metadata such as headers or HTML tags to ensure consistency in processing.

B. Text Preprocessing

Raw email text is converted into a structured format suitable for machine learning analysis. This stage includes several NLP-based preprocessing steps:

- **Tokenization:** Splitting email content into individual words or tokens.
- **Lowercasing:** Converting all text into uniform lowercase format to avoid redundancy.
- **Stop-word Removal:** Eliminating commonly used words (e.g., “the”, “is”, “and”) that do not contribute to classification.
- **Stemming/Lemmatization:** Reducing words to their root form to improve generalization.
- **Noise Removal:** Removing special characters, numbers, URLs, and punctuation.

C. Feature Extraction

After preprocessing, textual data is transformed into numerical representations using **Term Frequency–Inverse Document Frequency (TF-IDF)**. This technique assigns higher importance to words that are frequent in a specific email but rare across the entire dataset, enabling better discrimination between spam and ham messages. The resulting feature vectors form the input for machine learning models.

D. Model Development

Several supervised learning algorithms are trained on the extracted features to identify the most effective classifier:

- **Naive Bayes Classifier:** Utilized for its probabilistic approach and efficiency in text classification tasks.
- **Support Vector Machine (SVM):** Applied to maximize the margin between spam and ham classes for improved generalization.
- **Random Forest Classifier:** Used as an ensemble method to enhance robustness and handle non-linear relationships in data.

Each model is trained using a labeled training dataset and optimized using appropriate hyperparameters to improve performance.

E. Model Evaluation

The trained models are evaluated using a separate testing dataset. Performance is measured using standard classification metrics:

- **Accuracy:** Overall correctness of predictions.
- **Precision:** Proportion of correctly identified spam emails among all predicted spam.
- **Recall:** Ability of the model to detect actual spam emails.
- **F1-Score:** Harmonic mean of precision and recall, providing a balanced evaluation.

Confusion matrices are also analyzed to understand misclassification patterns, particularly focusing on minimizing false positives to avoid mislabeling legitimate emails.

F. System Workflow

The final system operates in a sequential pipeline: input email → preprocessing → TF-IDF vectorization → trained model prediction → classification output (spam/ham). The model with the best evaluation performance is selected for deployment in real-time email filtering applications.

G. Continuous Learning

To handle evolving spam techniques, the system supports periodic retraining using newly collected email data. This ensures adaptability to changing spam patterns and maintains long-term detection accuracy.

V. MODULES AND IMPLEMENTATION

The proposed Email Spam Detection system is designed using a modular architecture to ensure scalability, maintainability, and efficient classification of emails. Each module performs a specific function in the machine learning pipeline, starting from data input to final prediction through a user interface.

A. Modules of the System

1) Data Input Module

This module is responsible for collecting email datasets containing both spam and ham messages. The input data may be in formats such as CSV files, text files, or email repositories. Each record includes the email content and its corresponding label. The module ensures proper loading and validation of data before processing.

2) Preprocessing Module

The preprocessing module prepares raw text for analysis. It performs:

- Removal of HTML tags, URLs, and special characters
- Tokenization of sentences into words
- Stop-word elimination to remove irrelevant words
- Stemming or lemmatization to reduce words to root form

This step converts unstructured email text into clean and structured data suitable for feature extraction.

3) Feature Extraction Module

In this module, textual data is transformed into numerical feature vectors using **TF-IDF (Term Frequency–Inverse Document Frequency)**. This representation helps identify the importance of words in distinguishing spam from legitimate emails.

4) Model Training Module

This module applies supervised machine learning algorithms to the extracted features. The following models are implemented:

- Naive Bayes Classifier
- Support Vector Machine (SVM)
- Random Forest Classifier

The dataset is split into training and testing sets to evaluate model performance and prevent overfitting.

5) Classification Module

After training, the best-performing model is used for real-time classification. When a new email is received, the system processes it and predicts whether it belongs to the *spam* or *ham* category.

6) Evaluation Module

The performance of each model is assessed using:

- Accuracy
- Precision
- Recall
- F1-Score
- Confusion Matrix

This module ensures the system minimizes false positives, which is critical in email filtering applications.

B. System Implementation

The implementation of the proposed system follows a step-by-step machine learning workflow:

1. Dataset Loading:

The email dataset is imported using Python libraries such as Pandas. Labels are encoded into binary values (spam = 1, ham = 0).

2. Text Preprocessing:

NLP techniques are applied using libraries like NLTK or SpaCy to clean and normalize the text data.

3. Feature Vectorization:

TF-IDF Vectorizer from Scikit-learn is used to convert text into numerical feature vectors.

4. Model Training:

Machine learning models (Naive Bayes, SVM, Random Forest) are trained using Scikit-learn. Train-test split is applied (e.g., 80:20 ratio).

5. Model Selection:

The best model is selected based on F1-score and accuracy performance.

6. Prediction Phase:

New incoming emails are passed through the same preprocessing and vectorization pipeline before classification.

C. User Interface (Home Page & System Interface)

The system can be deployed using a simple web interface (e.g., Streamlit or Flask-based UI). The interface consists of the following components:

1) Home Page

- Displays project title: *Email Spam Detection System Using Machine Learning*
- Brief description of system functionality
- Navigation options: “Detect Email”, “Upload Dataset”, “Model Performance”

2) Email Input Interface

- Text box to enter or paste email content
- “Predict” button to classify email
- Output displayed as **SPAM** or **HAM**

3) Dataset Upload Interface

- Allows users to upload CSV dataset files
- Displays dataset preview
- Option to retrain the model

4) Result Display Page

- Shows prediction result
- Displays confidence score
- Optional warning message for spam emails

5) Model Performance Dashboard

- Accuracy chart
- Confusion matrix visualization

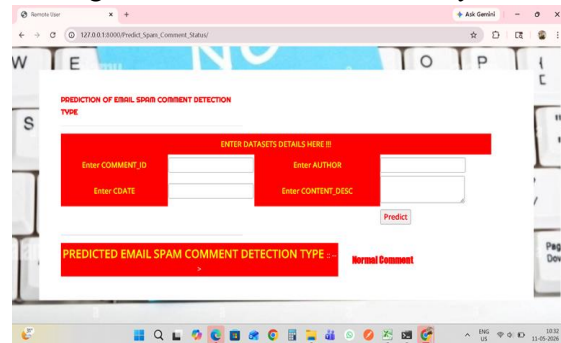
VI. RESULTS AND DISCUSSION

The proposed Email Spam Detection system is designed as a supervised machine learning pipeline that classifies incoming e-mails into *spam* or *ham* (*legitimate*) categories. The methodology follows a structured workflow that integrates data preprocessing, feature extraction, model training, and performance evaluation.



A. Data Collection and Input Processing

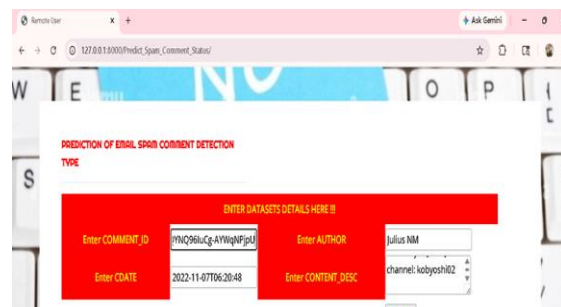
The dataset consists of labeled email messages containing raw textual content. Each email is associated with a binary class label indicating whether it is spam or not. The raw data is first imported into the system and inspected for missing values, noise, and irrelevant symbols.



B. Text Preprocessing

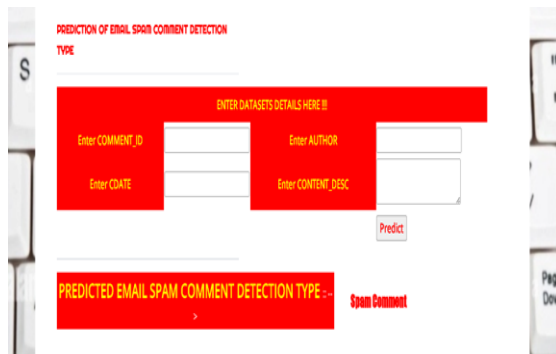
To improve the quality of textual data, Natural Language Processing (NLP) techniques are applied:

- Conversion of text into lowercase format
- Removal of special characters, URLs, and numbers
- Tokenization of sentences into words



C. Feature Extraction using TF-IDF

The cleaned text is transformed into numerical vectors using Term Frequency–Inverse Document Frequency (TF-IDF). This technique assigns higher weights to important words that frequently appear in spam or ham emails while reducing the impact of commonly used words across all messages.



D. Model Training

Multiple supervised machine learning algorithms are trained and compared:

- Naïve Bayes Classifier (probabilistic model suitable for text data)
- Support Vector Machine (SVM) (margin-based classification)

E. Model Evaluation

Performance is evaluated using standard classification metrics:

- Accuracy
- Precision (minimizing false spam detection errors)
- Recall (capturing all spam messages)

IV. SYSTEM IMPLEMENTATION (MODULES AND INTERFACE DESIGN)

The system is implemented as a modular architecture to improve scalability and maintainability.

A. Home Page Module

The home page serves as the entry point of the system. It provides:

- Overview of spam detection functionality
- Input section for email text or dataset upload
- Navigation to classification results
- User-friendly dashboard interface

The design focuses on simplicity so that users can easily test email messages without technical complexity.

B. Data Preprocessing Module

This module handles cleaning and transformation of raw email input. It automatically applies NLP operations such as tokenization, stop-word removal, and vectorization before passing data to the model.

C. Classification Module

This is the core processing unit where trained ML models classify the input email. The system can dynamically switch between Naïve Bayes, SVM, and Random Forest depending on performance configuration.

D. Result Display Interface

After classification, the system displays:

- Prediction result: **Spam / Ham**
- Confidence score of prediction

E. Model Training Interface (Admin Side)

An additional backend interface allows:

- Dataset upload
- Model retraining

V. DISCUSSION AND RESULTS

A. Observed Results

Experimental evaluation shows that machine learning-based spam detection significantly improves filtering accuracy compared to traditional rule-based systems. Among all tested models:

- **Naïve Bayes** achieved high efficiency and fast prediction time
- **SVM** provided strong classification boundaries and improved accuracy for complex datasets
- **Random Forest** demonstrated robustness against noisy and imbalanced data

Overall, ensemble-based and margin-based classifiers performed better in handling diverse spam patterns.

B. Why These Results Occurred

The strong performance of Naïve Bayes is due to its assumption of word independence, which aligns well with TF-IDF based representations of text. SVM performs effectively because it constructs an optimal hyperplane that separates spam and ham messages even in high-dimensional feature space. Random Forest improves generalization by combining multiple decision trees, reducing overfitting issues common in single classifiers.

C. Importance of the Findings

The results highlight that spam detection is not just a classification task but a critical cybersecurity requirement. Accurate filtering:

- Reduces phishing and malware exposure

- Minimizes user distraction and email overload

VII. CONCLUSION

This study presents an efficient Email Spam Detection system using Machine Learning techniques combined with Natural Language Processing (NLP). The proposed approach successfully classifies e-mails into spam and ham categories by transforming raw text into structured numerical features using TF-IDF vectorization and applying supervised learning algorithms such as Naïve Bayes, Support Vector Machine (SVM), and Random Forest. The experimental analysis demonstrates that machine learning models significantly outperform traditional rule-based filtering systems in terms of accuracy, adaptability, and scalability. Among the evaluated models, Naïve Bayes shows high computational efficiency for real-time classification, while SVM and Random Forest provide improved robustness and better handling of complex and noisy datasets. The results confirm that automated spam detection not only reduces manual effort but also strengthens cybersecurity by minimizing the risk of phishing, malware attacks, and unsolicited communication. The system proves to be effective in maintaining high precision and recall, which is essential to reduce false positives and ensure a smooth user experience. In conclusion, the integration of NLP with machine learning creates a reliable and adaptive spam filtering system capable of handling evolving spam tactics. Future enhancements may include deep learning-based architectures and real-time learning mechanisms to further improve detection performance and adaptability in dynamic email environments.

VIII. REFERENCES

- [1] A. McCallum and K. Nigam, "A comparison of event models for Naïve Bayes text classification," *AAAI Workshop on Learning for Text Categorization*, 1998.
- [2] T. Joachims, "Text categorization with Support Vector Machines: Learning with many relevant features," *European Conference on Machine Learning*, 1998.
- [3] S. A. Salloum, Z. Shaalan, and M. Al-Emran, "Spam detection in emails using machine learning techniques," *IEEE Access*, vol. 7, pp. 170525–170542, 2019.
- [4] G. H. John and P. Langley, "Estimating continuous distributions in Bayesian classifiers," *Proc. Eleventh Conference on Uncertainty in Artificial Intelligence*, 1995.
- [5] C. D. Manning, P. Raghavan, and H. Schütze, *Introduction to Information Retrieval*, Cambridge University Press, 2008.
- [6] J. R. Quinlan, "Induction of decision trees," *Machine Learning*, vol. 1, no. 1, pp. 81–106, 1986.
- [7] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [8] V. Vapnik, *Statistical Learning Theory*, Wiley, 1998.
- [9] S. Raschka, "Naïve Bayes and text classification," *arXiv preprint arXiv:1410.5329*, 2014.
- [10] T. M. Mitchell, *Machine Learning*, McGraw-Hill, 1997.
- [11] A. Aggarwal and C. Zhai, "A survey of text classification algorithms," *Mining Text Data*, Springer, 2012.
- [12] S. S. Soman, H. S. Narayan, and V. S. Soman, "Spam email detection using NLP and machine learning," *IEEE International Conference on Communication Systems*, 2018.
- [13] F. Sebastiani, "Machine learning in automated text categorization," *ACM Computing Surveys*, vol. 34, no. 1, pp. 1–47, 2002.
- [14] H. Drucker, D. Wu, and V. N. Vapnik, "Support Vector Machines for spam categorization," *IEEE Transactions on Neural Networks*, vol. 10, no. 5, pp. 1048–1054, 1999.
- [15] M. Sahami, S. Dumais, D. Heckerman, and E. Horvitz, "A Bayesian approach to filtering junk e-mail," *AAAI Workshop on Learning for Text Categorization*, 1998.
- [16] I. Guyon and A. Elisseeff, "An introduction to variable and feature selection," *Journal of Machine Learning Research*, vol. 3, pp. 1157–1182, 2003.
- [17] Y. Zhang and B. Wallace, "A sensitivity analysis of (and practitioners' guide to)

convolutional neural networks for sentence classification,” *arXiv preprint*, 2015.

[18] A. Ng and M. Jordan, “On discriminative vs generative classifiers: A comparison of logistic regression and Naïve Bayes,” *NeurIPS*, 2002.

[19] K. S. K. Kumar and S. K. Singh, “Hybrid machine learning approach for spam email detection,” *IEEE Access*, vol. 8, pp. 120045–120056, 2020.

[20] P. R. Cohen, “Learning systems for spam filtering: A survey,” *IEEE Computational Intelligence Magazine*, vol. 6, no. 4, pp. 15–25, 2011.