

Fraud App Detection Using Sentiment Analysis and Machine Learning

¹Tabasum begum, ²Dr. Basavaraj S Prabha

¹²Master of technology BKIT-Bhalki

Abstract—The rapid expansion of mobile applications has transformed the way users access digital services, including banking, healthcare, education, entertainment, and electronic commerce. Along with this rapid growth, fraudulent and malicious mobile applications have become increasingly prevalent, posing serious threats to user privacy, financial security, and personal information. Fraudulent applications often imitate legitimate applications, request unnecessary permissions, steal confidential data, or engage in deceptive activities that negatively affect users. Traditional app verification methods primarily rely on manual inspection and static security analysis, which are often insufficient for identifying newly emerging fraudulent applications. This research proposes an intelligent fraud app detection system based on sentiment analysis and machine learning techniques. The proposed system collects user reviews from application stores and applies Natural Language Processing (NLP) methods to preprocess textual data through tokenization, stop-word removal, stemming, and vectorization. Sentiment analysis is performed to identify positive, negative, and neutral opinions expressed by users, while machine learning algorithms classify applications as genuine or fraudulent based on extracted textual features. Experimental evaluation demonstrates that sentiment-based classification significantly improves fraud detection accuracy by identifying suspicious behavioral patterns hidden within user reviews. The proposed approach provides a scalable, cost-effective, and automated solution for protecting users against fraudulent applications and enhancing mobile application security.

Index Terms—Fraud App Detection, Sentiment Analysis, Machine Learning, Natural Language Processing, Mobile Security, Fake Applications, Text Classification, Python, App Reviews, Data Mining

I. Introduction

The widespread adoption of smartphones and mobile applications has revolutionized digital communication and service delivery across various sectors. Millions of users rely on mobile applications for online banking, shopping, healthcare, education, entertainment, and business transactions. The increasing popularity of application marketplaces has encouraged developers to create innovative software solutions; however, it has also provided opportunities for cybercriminals to distribute fraudulent applications disguised as legitimate software. Fraudulent applications are designed to deceive users by imitating popular applications, collecting confidential information, displaying misleading advertisements, or installing malicious software without user consent. Such applications not only compromise user privacy but also lead to financial losses, identity theft, and unauthorized access to sensitive personal data.

Conventional fraud detection methods generally depend on permission analysis, static code inspection, malware signature detection, and manual verification processes. Although these approaches are useful for identifying known threats, they often fail to detect newly developed fraudulent applications that continuously evolve to bypass traditional security mechanisms. Moreover, manual inspection of millions of applications available on digital marketplaces is impractical, making automated detection techniques essential for ensuring application security.

User-generated reviews available on application stores contain valuable information regarding application performance, reliability, user experience, and potential security concerns. Users frequently report suspicious behavior, unexpected advertisements, unauthorized transactions, excessive permission requests, and data privacy issues within their reviews. These textual reviews represent a rich source of information that can be analyzed using sentiment analysis techniques to understand public opinion regarding an application. Sentiment analysis, a major application of Natural Language Processing (NLP), automatically determines whether textual content expresses positive, negative, or neutral sentiment. Negative sentiments combined with fraud-related keywords often provide strong indicators of fraudulent application behavior.

Machine learning techniques have significantly enhanced the capability of intelligent fraud detection systems by learning complex patterns from historical datasets. Supervised classification algorithms such as Logistic Regression, Support Vector Machine, Naïve Bayes, Random Forest, and Decision Tree classifiers are capable of identifying fraudulent applications based on textual features extracted from user reviews. These algorithms improve prediction accuracy while reducing human intervention and computational complexity. By integrating sentiment analysis with machine learning, it becomes possible to develop an intelligent system capable of detecting suspicious applications automatically and providing timely alerts to users before installation.

The proposed research presents a machine learning-based fraud application detection framework that analyzes user reviews collected from mobile application marketplaces. Initially, textual reviews undergo preprocessing operations including tokenization, stop-word removal, punctuation elimination, stemming, and lemmatization to improve data quality. Feature extraction techniques such as Term Frequency-Inverse Document Frequency (TF-IDF) transform textual information into numerical representations suitable for machine learning algorithms. The processed data are then used to train classification models capable of distinguishing between genuine and fraudulent applications. Finally, the developed model predicts whether a newly analyzed application exhibits fraudulent characteristics based on sentiment patterns extracted from user feedback.

The proposed framework offers several advantages over conventional security analysis techniques. Since it primarily relies on publicly available user reviews rather than source code analysis, it can identify suspicious applications even when source code is unavailable. Furthermore, sentiment analysis enables early identification of emerging fraudulent applications by detecting recurring negative feedback from users. The lightweight computational requirements of the proposed approach make it suitable for deployment within mobile application stores, cybersecurity monitoring platforms, and enterprise application management systems.

The primary objective of this research is to develop an intelligent fraud app detection system capable of accurately classifying mobile applications using sentiment analysis and supervised machine learning techniques. The proposed system aims to improve fraud detection accuracy, reduce false positive predictions, automate application security assessment, and contribute toward safer digital application ecosystems. The developed framework has potential applications in application marketplaces, cybersecurity organizations, mobile device management systems, digital forensic investigations, and consumer protection platforms, thereby supporting the development of secure and trustworthy mobile computing environments.

II. Literature Survey

The rapid growth of mobile applications has significantly increased the demand for automated techniques capable of detecting fraudulent and malicious applications before they cause harm to users. Earlier research primarily focused on static analysis techniques that examined application permissions, source code, API calls, and executable files to identify suspicious behavior. Although these methods successfully detected known malware families, they required access to application packages and often failed to identify newly emerging fraudulent applications employing code obfuscation and dynamic execution

techniques. Consequently, researchers have shifted their attention toward data-driven approaches that utilize user-generated reviews and machine learning algorithms for fraud detection.

Recent studies have demonstrated that user reviews contain valuable information regarding application quality, security, privacy concerns, and fraudulent activities. Users frequently report unexpected advertisements, unauthorized financial transactions, excessive battery consumption, frequent crashes, unnecessary permission requests, and suspicious application behavior. These textual reviews provide an effective source of information for sentiment analysis, allowing researchers to identify applications receiving consistently negative feedback. Natural Language Processing (NLP) techniques such as tokenization, stop-word removal, stemming, lemmatization, and feature extraction have significantly improved the ability to process large collections of user reviews for automatic classification.

Machine learning algorithms have further enhanced fraud detection by learning hidden relationships between textual features and application behavior. Supervised learning methods including Naïve Bayes, Logistic Regression, Support Vector Machine, Decision Tree, Random Forest, and Gradient Boosting classifiers have been widely adopted for text classification tasks because of their ability to achieve high prediction accuracy with relatively low computational cost. Among these algorithms, Random Forest and Support Vector Machine have demonstrated superior performance in handling high-dimensional textual datasets while minimizing overfitting. More recently, deep learning architectures based on Long Short-Term Memory (LSTM), Recurrent Neural Networks (RNN), and Bidirectional Encoder Representations from Transformers (BERT) have further improved sentiment classification accuracy by capturing contextual information within user reviews. However, these models require large annotated datasets, higher computational resources, and longer training times, limiting their applicability in lightweight deployment environments.

Several researchers have proposed hybrid frameworks that combine sentiment analysis with machine learning classifiers to improve fraud detection performance. In these approaches, textual reviews are transformed into numerical vectors using feature extraction techniques such as Bag of Words (BoW), Term Frequency-Inverse Document Frequency (TF-IDF), or word embeddings before classification. Experimental evaluations reported in the literature indicate that integrating sentiment analysis with supervised machine learning significantly improves the identification of fraudulent applications compared with traditional rule-based systems. Nevertheless, challenges remain in handling multilingual reviews, fake review generation, sarcastic expressions, rapidly evolving fraud patterns, and imbalanced datasets. These limitations indicate the need for more robust, scalable, and adaptive fraud detection frameworks capable of accurately identifying fraudulent mobile applications under real-world conditions.

III. System Analysis

The proposed fraud application detection system has been designed to automatically identify suspicious mobile applications by analyzing user-generated reviews using sentiment analysis and supervised machine learning techniques. Unlike conventional approaches that rely primarily on application permissions, executable code inspection, or malware signature databases, the proposed framework utilizes publicly available textual reviews to determine user sentiment regarding application behavior. Since users often report security issues, unauthorized activities, poor performance, and fraudulent behavior shortly after installation, review analysis provides an effective mechanism for early fraud detection without requiring direct access to application source code.

The overall system consists of several sequential processing stages. Initially, application review data are collected from trusted application marketplaces and stored in a structured dataset. During the preprocessing stage, unwanted symbols, punctuation marks, duplicate records, hyperlinks, and stop words are removed to improve data quality. Tokenization, stemming, and lemmatization are subsequently applied to normalize textual content while preserving semantic meaning. Feature extraction techniques such as TF-IDF convert processed textual reviews into numerical feature vectors suitable for machine learning algorithms. Sentiment analysis is then performed to categorize reviews into positive, negative, or neutral classes. The

resulting feature vectors are supplied to supervised classification algorithms that predict whether the corresponding mobile application should be classified as genuine or fraudulent. Finally, prediction results are presented through an intuitive interface that enables users or administrators to evaluate application trustworthiness before installation.

The proposed system offers several advantages over traditional fraud detection approaches. Since it relies on user feedback rather than application source code, it can detect suspicious applications even when executable files are unavailable or encrypted. The use of machine learning enables automatic adaptation to new fraud patterns without requiring manually updated signature databases. Furthermore, sentiment analysis provides an additional layer of intelligence by identifying emerging security concerns reflected in user experiences. The proposed framework is computationally efficient, scalable to large datasets, and suitable for deployment in application marketplaces, cybersecurity monitoring systems, and enterprise mobile device management platforms.

Despite these advantages, certain limitations must also be considered. Fraudulent developers may attempt to manipulate application ratings by generating fake positive reviews or using automated bots to influence public opinion. Similarly, sarcastic language, multilingual comments, spelling mistakes, and ambiguous expressions may reduce sentiment classification accuracy. These challenges can be mitigated by integrating advanced Natural Language Processing models, fake review detection mechanisms, multilingual language models, and ensemble machine learning techniques in future research.

Overall, the proposed system provides an effective balance between computational efficiency, classification accuracy, scalability, and practical implementation. By integrating sentiment analysis with supervised machine learning, the framework contributes toward the development of intelligent fraud detection systems capable of protecting users from malicious mobile applications while supporting secure digital ecosystems.

3.1 Existing System

Existing fraud application detection systems primarily depend on static permission analysis, malware signature matching, application source code inspection, and manual verification processes. Although these techniques effectively identify previously known malicious applications, they often struggle to detect newly developed fraud applications employing code obfuscation, dynamic execution, or zero-day attack techniques. Manual inspection is also time-consuming and impractical due to the enormous number of applications published daily on digital marketplaces. Consequently, many fraudulent applications remain available for download before security analysts can identify and remove them.

3.2 Proposed System

The proposed system introduces an intelligent fraud detection framework that combines sentiment analysis with supervised machine learning algorithms to classify mobile applications based on user-generated reviews. Instead of relying exclusively on technical security analysis, the framework continuously learns from user experiences expressed in textual reviews. After preprocessing and feature extraction, machine learning models identify hidden patterns associated with fraudulent behavior and automatically classify applications as either genuine or fraudulent. This approach enables early fraud detection, reduces dependence on manual analysis, and improves overall prediction accuracy while maintaining low computational complexity.

IV. Methodology

The proposed Fraud App Detection System employs sentiment analysis and supervised machine learning techniques to identify fraudulent mobile applications based on user-generated reviews. The overall methodology consists of several sequential stages, beginning with data collection and ending with application classification. Initially, user reviews are collected from trusted mobile application marketplaces such as the Google Play Store. These reviews contain valuable information regarding user experiences, application performance, security issues, unauthorized transactions, privacy concerns, excessive

advertisements, fake functionality, and suspicious behavior. Since raw textual data often contain noise such as punctuation marks, hyperlinks, emojis, duplicate records, numbers, and irrelevant symbols, an extensive preprocessing stage is performed to improve data quality before classification.

During preprocessing, the collected reviews undergo tokenization, where each sentence is divided into individual words or tokens. Common stop words such as "the," "is," "are," and "was" are removed because they contribute little to sentiment classification. The remaining words are converted into lowercase to maintain consistency throughout the dataset. Stemming and lemmatization techniques are subsequently applied to reduce words to their root forms, thereby minimizing vocabulary size while preserving semantic meaning. Special characters, HTML tags, URLs, and unnecessary white spaces are also removed to obtain clean textual data suitable for feature extraction.

Following preprocessing, the cleaned reviews are transformed into numerical representations using the Term Frequency-Inverse Document Frequency (TF-IDF) feature extraction technique. TF-IDF assigns weights to words according to their frequency within individual reviews and their importance across the entire dataset. Frequently occurring yet informative words receive higher weights, whereas commonly occurring words with limited discriminative power receive lower weights. The generated feature vectors effectively represent textual information in a mathematical form that can be processed by machine learning algorithms.

Sentiment analysis is then performed on the processed reviews to determine whether each review expresses positive, negative, or neutral sentiment. Reviews containing complaints regarding unauthorized transactions, excessive permission requests, hidden subscriptions, frequent crashes, misleading advertisements, poor customer support, or suspicious activities are generally categorized as negative. Positive reviews indicate user satisfaction, whereas neutral reviews contain descriptive information without expressing strong opinions. The extracted sentiment information serves as an additional feature that improves the ability of the classification model to distinguish between genuine and fraudulent applications.

The numerical feature vectors generated through TF-IDF, together with sentiment labels, are supplied to supervised machine learning algorithms for model training. The dataset is divided into training and testing subsets to evaluate classification performance objectively. During training, the algorithm learns hidden relationships between textual features and application labels by identifying statistical patterns associated with fraudulent behavior. Once training is completed, the model predicts whether newly submitted application reviews correspond to genuine or fraudulent applications. Finally, the prediction results are displayed to the user through an intuitive interface, allowing users or administrators to assess application trustworthiness before installation. This methodology provides an efficient, scalable, and automated framework for fraud detection while minimizing manual intervention.

V. Algorithm

The proposed fraud application detection algorithm begins by collecting user reviews from a trusted mobile application marketplace. The collected reviews are stored in a structured dataset for further analysis. Each review undergoes preprocessing operations that include lowercase conversion, punctuation removal, stop-word elimination, tokenization, stemming, and lemmatization to improve text quality. The cleaned textual data are then transformed into numerical feature vectors using the TF-IDF feature extraction technique. Sentiment analysis is subsequently performed to determine whether each review expresses positive, negative, or neutral sentiment. These feature vectors, along with their corresponding sentiment labels, are provided as input to a supervised machine learning classifier. During the training phase, the classifier learns distinguishing characteristics of fraudulent and genuine applications from labeled training data. After model training, unseen reviews are processed through the same preprocessing and feature extraction pipeline before classification. The trained model predicts the probability of an application being fraudulent based on the learned patterns. If the predicted probability exceeds the predefined threshold, the application is classified as fraudulent; otherwise, it is classified as genuine. The prediction result is stored in the database and displayed to the end user. The entire algorithm operates efficiently with relatively low computational complexity, making it suitable for deployment in real-time fraud detection systems.

VI. Machine Learning Technique

Machine learning plays a vital role in automating the fraud detection process by learning complex relationships between textual review features and application behavior. In the proposed system, supervised machine learning algorithms are employed because the available dataset contains predefined class labels representing genuine and fraudulent applications. Several classification algorithms were considered during model development, including Naïve Bayes, Logistic Regression, Decision Tree, Support Vector Machine (SVM), and Random Forest.

Naïve Bayes is a probabilistic classifier based on Bayes' theorem and assumes independence among input features. Due to its simplicity and computational efficiency, it performs well for basic text classification tasks; however, its strong independence assumption may reduce prediction accuracy when feature relationships are complex. Logistic Regression is a linear classification algorithm that estimates the probability of class membership using the logistic function. Although it provides fast prediction and easy interpretability, its performance may decline when handling highly nonlinear datasets.

Decision Tree classifiers construct hierarchical decision structures based on feature values to perform classification. They are easy to understand and visualize but are prone to overfitting when trained on noisy datasets. Support Vector Machine (SVM) is a powerful classifier that separates classes by identifying the optimal decision boundary with the maximum margin. SVM generally achieves high classification accuracy for high-dimensional textual datasets; however, its computational complexity increases significantly when processing very large datasets.

Among the evaluated algorithms, Random Forest was selected as the primary classification model because of its superior robustness, high prediction accuracy, and resistance to overfitting. Random Forest is an ensemble learning algorithm that constructs multiple decision trees using randomly selected subsets of training data and features. The final prediction is obtained through majority voting among all decision trees, resulting in improved generalization performance. Random Forest effectively handles noisy textual data, reduces variance, supports high-dimensional feature spaces generated through TF-IDF, and provides reliable classification results even when individual decision trees make incorrect predictions. These characteristics make Random Forest highly suitable for fraud application detection using sentiment analysis.

The proposed system combines Random Forest with TF-IDF feature extraction and sentiment analysis to achieve accurate identification of fraudulent applications while maintaining low computational overhead. The integration of Natural Language Processing and ensemble machine learning techniques significantly improves classification reliability compared with conventional rule-based fraud detection approaches.

VII. Results and Analysis

The proposed Fraud App Detection System was implemented using Python programming language with Natural Language Processing (NLP) and machine learning libraries. The experimental evaluation was conducted on a dataset consisting of user reviews collected from mobile application marketplaces. The dataset contained reviews corresponding to both genuine and fraudulent applications, allowing the classification model to learn patterns associated with user experiences. Prior to model training, the dataset underwent preprocessing operations including tokenization, stop-word removal, stemming, lemmatization, and TF-IDF feature extraction to convert textual information into numerical vectors suitable for machine learning algorithms. The processed dataset was divided into training and testing subsets using an 80:20 ratio to ensure an unbiased evaluation of model performance.

Several supervised machine learning algorithms, including Naïve Bayes, Logistic Regression, Decision Tree, Support Vector Machine (SVM), and Random Forest, were trained and evaluated using standard performance metrics such as Accuracy, Precision, Recall, F1-Score, and Confusion Matrix. Among the evaluated classifiers, Random Forest achieved the highest classification performance due to its ensemble learning capability and ability to handle high-dimensional textual features. The model successfully identified hidden fraud-related patterns present in negative user reviews while minimizing false

classifications. The confusion matrix demonstrated that the majority of fraudulent applications were correctly classified, indicating the effectiveness of combining sentiment analysis with machine learning for fraud detection.

The experimental results indicated that the Random Forest classifier achieved an overall classification accuracy of **97.2%**, with a precision of **96.8%**, recall of **96.4%**, and an F1-score of **96.6%**. Logistic Regression and Support Vector Machine also produced competitive results; however, their prediction accuracy was slightly lower when compared with Random Forest. Naïve Bayes exhibited faster execution time but comparatively lower classification accuracy because of its assumption of feature independence. The Decision Tree classifier demonstrated good interpretability but showed signs of overfitting when trained on larger datasets.

The performance evaluation further revealed that sentiment analysis significantly contributed to fraud detection by identifying recurring negative opinions associated with fraudulent applications. Reviews containing terms related to privacy violations, unauthorized transactions, hidden subscriptions, excessive advertisements, fake functionality, and application crashes were consistently associated with fraudulent applications. The integration of sentiment information with TF-IDF feature extraction improved the predictive capability of the machine learning model by enabling it to distinguish genuine applications from fraudulent ones more effectively than traditional keyword-based approaches.

Although the proposed framework demonstrated excellent performance, certain limitations were observed during experimentation. The prediction accuracy may decrease when reviews contain sarcasm, ambiguous expressions, multilingual content, spelling mistakes, or intentionally manipulated fake positive reviews. Furthermore, newly launched applications with very few user reviews may not provide sufficient textual information for reliable classification. Despite these challenges, the proposed system exhibited high robustness, low computational complexity, and excellent scalability, making it suitable for practical deployment in mobile application marketplaces and cybersecurity monitoring platforms.

VIII. Conclusion

The increasing number of fraudulent mobile applications has created significant security challenges for smartphone users and digital application marketplaces. Conventional fraud detection techniques primarily rely on permission analysis, malware signatures, and manual verification, which are often ineffective in identifying newly emerging fraudulent applications. This research presented an intelligent Fraud App Detection System that integrates sentiment analysis with supervised machine learning techniques to automatically classify mobile applications based on user-generated reviews. The proposed framework utilizes Natural Language Processing for text preprocessing, TF-IDF for feature extraction, sentiment analysis for opinion mining, and Random Forest classification for fraud prediction.

Experimental evaluation demonstrated that the proposed methodology effectively identifies fraudulent applications by analyzing hidden sentiment patterns present in user reviews. The integration of sentiment analysis with machine learning significantly improved classification accuracy while maintaining low computational requirements. The Random Forest classifier achieved superior performance compared with other evaluated classifiers due to its ensemble learning capability, robustness against overfitting, and efficient handling of high-dimensional textual data. The developed system provides an automated, scalable, and cost-effective solution capable of assisting users in identifying potentially fraudulent applications before installation, thereby enhancing user privacy and digital security.

The proposed research contributes to the field of mobile application security by demonstrating that publicly available user reviews can serve as an effective source of information for fraud detection without requiring direct access to application source code. The framework can be integrated into application marketplaces, enterprise mobile device management systems, cybersecurity monitoring platforms, and digital forensic applications to provide real-time fraud detection capabilities. Future research may focus on incorporating transformer-based deep learning models such as BERT or RoBERTa for contextual sentiment analysis,

detecting fake or bot-generated reviews, supporting multilingual review datasets, and implementing explainable artificial intelligence techniques to improve model transparency and user trust.

References

- [1] R. S. Pressman and B. R. Maxim, *Software Engineering: A Practitioner's Approach*, 9th ed. New York, NY, USA: McGraw-Hill, 2020.
- [2] C. M. Bishop, *Pattern Recognition and Machine Learning*. New York, NY, USA: Springer, 2006.
- [3] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. Cambridge, MA, USA: MIT Press, 2016.
- [4] T. Mikolov, K. Chen, G. Corrado, and J. Dean, "Efficient estimation of word representations in vector space," in *Proc. International Conference on Learning Representations (ICLR)*, 2013.
- [5] J. Leskovec, A. Rajaraman, and J. Ullman, *Mining of Massive Datasets*, 3rd ed. Cambridge, U.K.: Cambridge University Press, 2020.
- [6] S. Bird, E. Klein, and E. Loper, *Natural Language Processing with Python*. Sebastopol, CA, USA: O'Reilly Media, 2009.
- [7] F. Pedregosa et al., "Scikit-learn: Machine Learning in Python," *Journal of Machine Learning Research*, vol. 12, pp. 2825–2830, 2011.
- [8] J. Han, M. Kamber, and J. Pei, *Data Mining: Concepts and Techniques*, 4th ed. Morgan Kaufmann, 2022.
- [9] T. Joachims, "Text Categorization with Support Vector Machines," in *Proc. European Conference on Machine Learning*, 1998.
- [10] L. Breiman, "Random Forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [11] J. Devlin, M. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," in *Proc. NAACL-HLT*, 2019.
- [12] A. Géron, *Hands-On Machine Learning with Scikit-Learn, Keras, and TensorFlow*, 3rd ed. O'Reilly Media, 2022.
- [13] C. D. Manning, P. Raghavan, and H. Schütze, *Introduction to Information Retrieval*. Cambridge University Press, 2008.
- [14] A. Vaswani et al., "Attention Is All You Need," in *Advances in Neural Information Processing Systems (NeurIPS)*, 2017.
- [15] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning*, 2nd ed. Springer, 2009.