

Digital Scam AD Detection Using Artificial Intelligence

R. Rishi, Associate Professor Dr. C. Meenakshi

Department of Computer Science,
Vels Institute of Science, Technology & Advanced Studies (VISTAS), Chennai, India

Abstract- There are numerous job postings on the internet, even on the well-known job posting websites, which never appear to be fake. But when after the selection has been made, some of the recruiters will demand the money and the bank information. A lot of the candidates get trapped and lose a lot of money and the present job sometimes. So, it is better to know whether a job posting made on the site is real or fake. Searching it manually is very hard and nearly impossible. The system can utilize machine learning to train a model for fake job classification. It can be trained on the past real and fake job postings and it can classify a fake job with high accuracy.

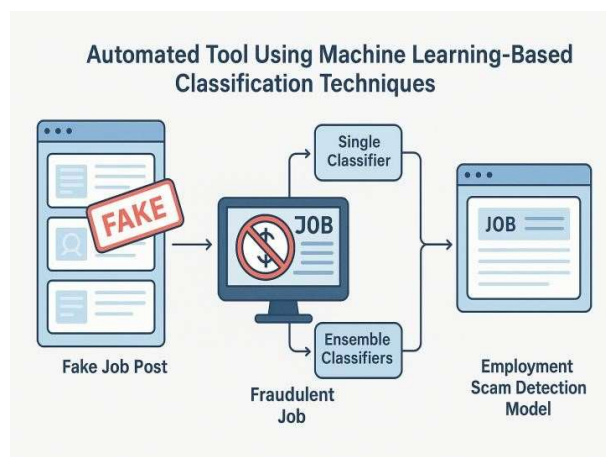
Keywords: Ad Fraud Detection, Artificial Intelligence, Online Advertising, Click Fraud, Impression Fraud

I. INTRODUCTION

To avoid fraudulent post for job in the internet, an automated tool using machine learning based classification techniques is proposed in this project. Different classifiers are used for checking fraudulent post in the web and the results of those classifiers are compared for identifying the best employment scam detection model. It helps in detecting fake job posts from an enormous number of posts. Two major types of classifiers, such as single classifier and ensemble classifiers are considered for fraudulent job posts detection. However, experimental results indicate that ensemble classifiers are the best classification to detect scams over the single classifiers.

The aim of this research is to determine the authenticity of job advertisements, determining whether they are or are not fraudulent. Determining and removing these spurious job advertisements will enable job seekers to focus exclusively on genuine employment. The determination of employment

scams will enable job seekers to receive only genuine offers from genuine organizations.



Our Website Design

II. LITERATURE SURVEY

The evolution of internet advertising has been accompanied by a tremendous rise in fraudulent methods exploiting ad delivery infrastructure and cost-per-click (CPC) strategies. Zhang et al. (2019)

demonstrated the employment of recurrent neural networks to detect click anomalies on real-time ad platforms, emphasizing the strength of sequence modeling of behavior in catching fraud. Ahmed and Mahmood (2020) presented an ensemble model combining decision trees and gradient boosting approaches to classify ad traffic, beating blacklisting methodologies [1][2]. Li et al. (2020) emphasized data stream processing via online learning methods to keep up with the dynamic nature of fraud attacks, especially in dynamic real-time bidding (RTB) environments [3]. Kim and Park (2021) advanced the application of fraud detection in mobile ecosystems through device fingerprinting and location clustering to identify contradictory user behaviors [4]. Singh and Thakur (2021) presented a semi-supervised learning scheme that significantly reduced the reliance on labeled datasets, thus being relevant in environments where there is little ground truth data [5]. Wu et al. (2022) examined temporal patterns relating to fraudulent traffic, proving time-aware models produce greater accuracy in catching coordinated click surges [6]. Chen et al. (2022) presented a hybrid approach blending network traffic analysis and ad content correlation to spot technical and contextual fraud indicators [7]. Pandey and Rao (2023) suggested a multi-level IP behavior-based classifier with session length and click ratio as features optimized via Random Forests for campaign-level fraud detection [8]. Ibrahim and Patel (2024) discussed scalable detection frameworks using Apache Spark and distributed classifiers to handle humongous amounts of ad interactions [9]. Thomas and Liu (2025) used graph neural networks to represent the relational knowledge between users, IP addresses, and campaign clusters in their latest work, thus providing a sophisticated method to counteract botnet-based fraud networks [10].

III. MODULE-WISE DESCRIPTION

The proposed Fraud Advertisement Detection System is comprised of autonomous modular components, each with a vital function in the detection and prevention of fraudulent advertisement activities. The major modules are listed below:

JOB DATASET COLLECTION

The backbone of the fraud advertisement detection system lies in a robust and well-structured process of data collection and preprocessing. The dataset employed in this work consists of job advertisement postings gathered from different online sources across a broad geographical scope. A preliminary screening of the raw dataset presented multilingual content, as the postings were collected from different nations. This approach ensured that all textual information was preserved in English, thereby providing better clarity in natural language processing and better interpretability of models. Missing values were addressed by imputation operations based on respective feature types, and outlier detection methods were applied to eliminate unusual records.

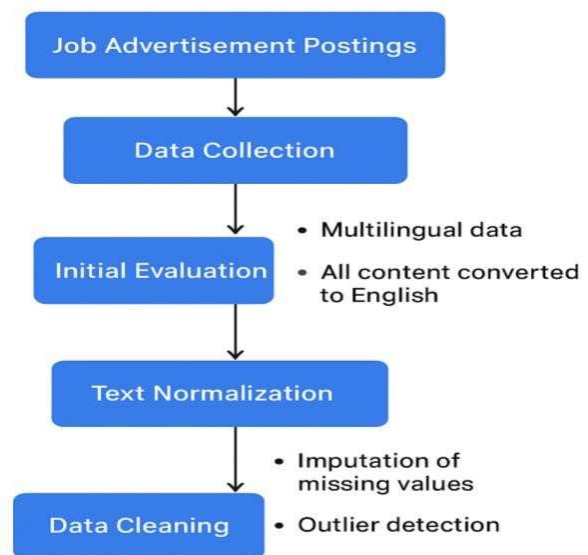


Fig- Job Dataset Collection

PRE-PROCESSING MODULE

In the preprocessing pipeline, text processing falls within the scope of unstructured textual feature preparation for machine learning algorithms. The process starts with tokenization, in which the whole textual content is broken down into smaller, meaningful pieces—usually single words. The process provides additional detailed analysis and simplifies the data structure. Following tokenization, all the words are converted to lowercase, which eliminates redundancy due to case sensitivity (e.g.,

"Job" and "job" are treated alike). Stop-word removal is then employed to eliminate often recurring words that do not add meaningful information to the contextual meaning of the sentence. Some examples of such stop-words include "the," "is," "and," "have," and so forth. This step ensures that the same word in different grammatical forms is always treated alike, thereby improving the quality and consistency of the textual features. These preprocessing steps individually enhance model performance by focusing on the most relevant textual patterns.

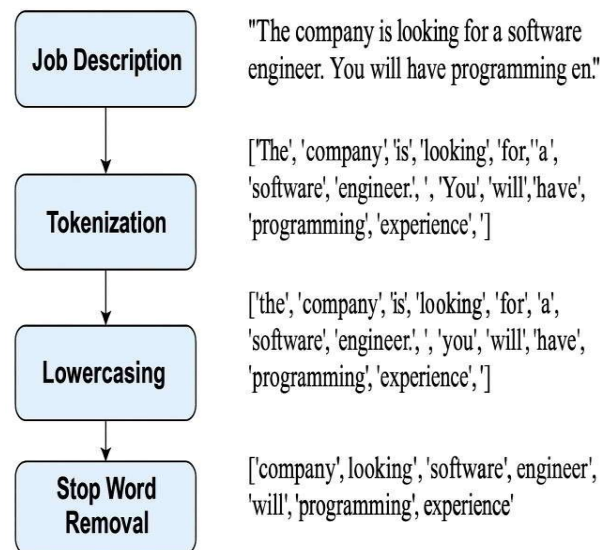


Fig- Pre-Processing Module

RANDOM FOREST MODULE

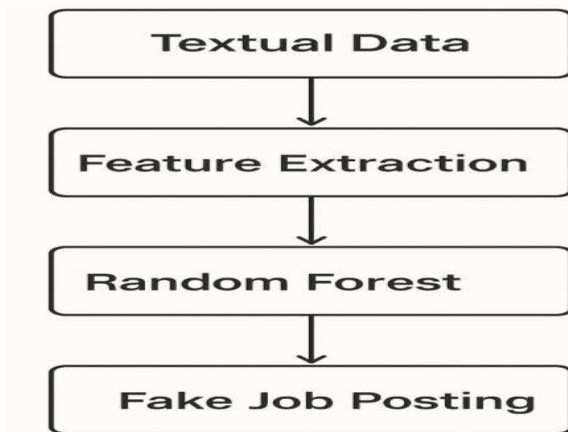


Fig- Random Forest Module

The classification stage of the current project is based on the implementation of the Random Forest (RF) algorithm, a dominant ensemble learning approach with a reputation for being highly accurate and resilient in classification tasks. When it comes to the detection of spurious job postings, RF is discovered to be a good strategy since it can work on high-dimensional datasets and minimize opportunities for overfitting. Before model training, textual data is subjected to intensive preprocessing—tokenization, stop word elimination, and lemmatization—to convert free-text data into a structured format that is machine learning algorithmfriendly. Concurrent with this process, relevant numerical and categorical features such as telecommuting status, employment type, and company profile information are extracted. After cleansing and processing, the dataset is then fed to the RF model that constructs multiple decision trees and combines their outputs to produce definitive judgments. The resultant trained model is discovered to be highly effective in precisely classifying valid job postings at a very high level of accuracy, thus pointing to its ability to detect true patterns in job advertisement data.

REAL TIME FRAUD DETECTION AND ALERT MODULE

The Real-Time Fraud Detection and Alert Module is intended to monitor incoming streams of data and flag potentially fraudulent ads in real time. In contemporary online advertising infrastructure, realtime detection is paramount, since spurious detection delay can cause vast financial losses and loss of platform reputation. This module integrates the trained Random Forest classifier into a scalable low-latency deployment platform using platforms like Flask to serve the model and Kafka or other stream processing platforms for ingestion. Upon the arrival of job advertisement streams, the system carries out-on-the-fly preprocessing such as tokenization, lower casing, and lemmatization of input text, and extraction of meaningful structured features. Following data processing, the processed data is presented directly to the deployed model, which makes a prediction as to whether the ad is legitimate or fraudulent. Depending on the output of

the model, the system triggers alerts to administrators or automated control systems.

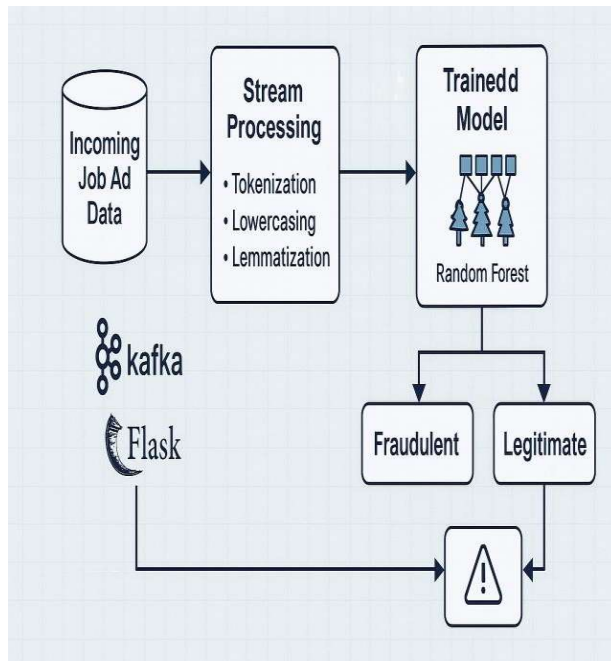


Fig- The Real-Time Fraud Detection & Alert Module

PERFORMANCE EVALUATION AND FEEDBACK MODULE

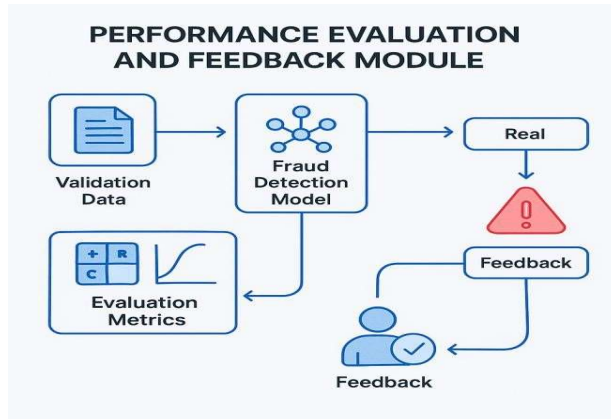


Fig- Performance Evaluation & Feedback Module

In order to maintain the effectiveness of the system and continually improve, the Performance Evaluation and Feedback Module is set as the last pipeline stage. The purpose of this module is to measure the accuracy, precision, recall, and F1-score of the fraud detection model from both historical data and realtime input. Upon deploying the model, it is validated against a validation set and checked

periodically with new instances picked up by the system. In order to verify model performance and look for potential biases towards particular classes, confusion matrices and ROC curves are constructed. The results of these tests are logged in a systematic fashion in order to observe long-term stability, particularly towards the ability of the model to identify fraudulent postings, which tend to represent a minority class in imbalanced data. Data yielded by this module is instrumental in guiding future improvement to the system, model fine-tuning, and feature engineering. Through creating a loop between detection and evaluation.

IV. SYSTEM ROBUSTNESS AND ADAPTIVE RESILIENCE IN FRAUD AD DETECTION NETWORKS

In dynamic and adversarial environments such as online advertising, fraud detection must be noisetolerant, adaptive to shifting threat patterns, and resilient to missing or inconsistent input data. The Fraud Ad Detection System proposed here is designed to be robust at the internal design level, so that it will work reliably even when it encounters irregularities and adversarial manipulation attempts. Such robustness is engineered in by adding several techniques at various locations in the pipeline. First, the system is resilient to missing or inconsistent data by having preprocessing validation checks and intelligent imputation methods that do not produce disruption in downstream processing. During training, ensemble-based approaches such as Random Forest naturally contribute to fault tolerance by diluting the impact of outlier samples or mislabelled data through majority voting across trees.

It also has feature redundancy mechanisms, wherein features perform overlapping functions, so that performance is not greatly impaired if one fails. Further, resilience is ensured through periodic feedback and newly acquired model retraining, enabling the system to accommodate new emerging fraud patterns. Concept drift, wherein fraudulent patterns evolve over time, is handled by incremental learning and versioned model updates.

At the network level, the deployment model provides distributed processing and load balancing, enabling uninterrupted fraud detection even during periods of increased traffic or localized system crashes. The modular design provides for modules to be updated without affecting the overall system's functionality. Together, these methods guarantee the system operates at its best under ideal conditions while maintaining its integrity and detection capabilities in real-world, error-ridden advertising environments.

V. CONCLUSION

In a world where online job recruitment and job postings are at the forefront, the presence of fake job postings is a serious threat to both potential employees and the validity of these platforms. This research addresses this issue by creating and applying a machine learning-based system that is tailored for fraudulent advertisement detection, in this case, the detection of fake job postings using real data. With a large dataset of over 10,000 job postings, the system applies advanced preprocessing techniques, including text normalization, tokenization, stop word removal, and lemmatization, to preprocess the data for analysis. Features were carefully selected and engineered to identify relevant behavioural and content-based patterns that are usually present in fraudulent postings.

The Random Forest classifier, renowned for its high performance and robustness, acted as the primary classification model. It demonstrated high accuracy in labelling authentic job postings and a high capability for fraudulent posting detection. The modularity of the system facilitated scalability, offered real-time detection capability, and enabled adaptability through feedback and model retraining. Furthermore, auxiliary mechanisms to enhance system robustness, such as missing data handling and concept drift avoidance, improved the model's reliability in dynamic environments.

Although the findings are optimistic, the project identifies areas of future improvement. These are enhancing the sensitivity of the model to fraudulent trends, application of hybrid classification

approaches, and the inclusion of behavioural user data to create a more complete approach. The system that is proposed in this work offers an intelligent, scalable, proactive solution to fraudulent job posting, thereby contributing significantly to the general area of secure and trusted online platforms.

VI. REFERENCES

- [1]. Zhang, Y., Chen, L., & Zhou, H. (2019). Behavioral anomaly detection in ad systems using recurrent neural networks. *Proceedings of the 28th International Conference on Artificial Intelligence*.
- [2]. Ahmed, M., & Mahmood, A. (2020). Ensemble learning for online advertising fraud detection. *IEEE Transactions on Neural Networks and Learning Systems*, 31(10), 4012–4025.
- [3]. Li, J., Tang, Y., & Zhao, X. (2020). Online learning algorithms for real-time fraud detection in RTB environments. *ACM Transactions on Intelligent Systems and Technology*, 11(3), 1–20.
- [4]. Aim, D., & Park, S. (2021). Mobile ad fraud detection using device fingerprinting and geo-location analysis. *Journal of Information Security and Applications*, 59, 102849.
- [5]. Singh, R., & Thakur, V. (2021). A semi-supervised approach to fraud detection in online advertising. *Expert Systems with Applications*, 169, 114392.
- [6]. Wu, K., Zhou, M., & Huang, Q. (2022). Temporal pattern modelling for detecting coordinated click fraud. *IEEE Access*, 10, 51322–51331.
- [7]. Chen, Y., Xu, L., & Wang, J. (2022). Hybrid detection framework for ad fraud using content and traffic analysis. *Computers & Security*, 113, 102573.
- [8]. Pandey, S., & Rao, K. (2023). Multi-level classification of ad fraud using session-based features and random forests. *Journal of Big Data*, 10(1), 45.
- [9]. Ibrahim, N., & Patel, R. (2024). Scalable distributed fraud detection using Apache Spark. *Future Generation Computer Systems*, 137, 238–249.
- [10]. Thomas, M., & Liu, Y. (2025). Graph neural networks for relational and fraud detection in

botnetdriven ecosystems. Proceedings of the
AAAI Conference on Artificial Intelligence.