



AI POWERED RECRUITMENT SCAM DETECTION

Machine Learning

Dr. E. S. Shamila	Professor & Head, Department of Artificial Intelligence and Data Science, Jansons Institute of Technology, Karumathampatti, Coimbatore
Aswin R	UG Student, Jansons Institute of Technology.
Madala Siva Naga Poojitha	UG Student, Jansons Institute of Technology.
Mani Maran P	UG Student, Jansons Institute of Technology.
Potti Likithasri	UG Student, Jansons Institute of Technology

ABSTRACT

The rapid growth of online recruitment platforms has led to a significant rise in fake job postings and internship scams, exposing job seekers to financial losses and identity theft. Existing detection systems rely on single-source analysis such as keyword filtering or basic spam detection, which fails to address the multi-dimensional nature of modern recruitment fraud. This paper proposes RecruitGuard AI, a multi-source scam detection system integrating Machine Learning (ML), Large Language Models (LLMs), and Retrieval-Augmented Generation (RAG) for intelligent, explainable fraud detection. The system analyzes job descriptions, recruiter messages, job URLs, and uploaded offer letters using Optical Character Recognition (OCR) for document processing. Supervised ML models detect suspicious patterns while LLM-assisted RAG enhances contextual reasoning by retrieving verified external knowledge. An explainable trust score classifies job opportunities as legitimate or fraudulent. Experimental results demonstrate an accuracy of 91%, precision of 89%, recall of 87%, and F1-score of 88%, outperforming all baseline models.

KEYWORDS

Fake Job Detection, LLM, Retrieval-Augmented Generation, OCR, Recruitment Scam, Trust Score, Explainable AI

INTRODUCTION

The rapid expansion of online recruitment platforms such as LinkedIn, Naukri, and Indeed has broadened employment access globally. However, this growth has simultaneously enabled a sharp rise in recruitment fraud. Fraudsters exploit job seekers through unrealistic salary offers, fake recruiter identities, malicious URLs, and forged offer letters, targeting particularly students and early-career professionals [1].

The Microsoft Cyber Signals Report (2023) states that online job scams increased by over 300% globally. The U.S. Federal Trade Commission reported losses exceeding \$2 billion between 2020 and 2023, with similar trends observed in India through NCRP reports. Traditional rule-based and single-source machine learning approaches are increasingly insufficient, as fraudsters closely mimic legitimate recruitment communications. Most existing systems also lack explainability, offering only binary outputs without reasoning, reducing user trust [2].

To address these challenges, this paper proposes RecruitGuard AI, a multi-source recruitment scam detection system that integrates ML, LLMs, and RAG for accurate, explainable fraud detection across diverse recruitment inputs.

Related Works

Early research on recruitment fraud focused on supervised ML techniques. Reddy et al. [1] proposed fake job detection using Logistic Regression, Random Forest, and SVM on labeled job advertisement datasets. While effective on structured data, these models struggle with contextually sophisticated fraud patterns. Kumar et al. [13] extended this work using IEEE Access datasets, achieving improved detection accuracy.

Deep learning approaches by Zhang et al. [11] and Pillai [6] introduced Bi-LSTM models that capture sequential dependencies in recruitment text, outperforming shallow classifiers. NLP-based techniques from Sowmya et al. [9] and Brown and Smith [12] employed TF-IDF and keyword analysis for linguistic fraud detection, though they lack semantic depth for advanced scam identification.

Singh et al. [2] demonstrated that LLMs significantly outperform traditional classifiers in detecting contextually deceptive patterns. However, standalone LLMs risk hallucination. Lewis et al. [14] introduced the foundational RAG framework (NeurIPS, 2020), showing that grounding LLM outputs in retrieved external facts

substantially reduces false predictions. Al Barwani et al. [3] applied RAG directly to phishing and scam detection with strong results. Despite these advances, existing systems analyze only single input sources and lack user-facing support features, gaps that RecruitGuard AI directly addresses.

Proposed System

Recruit Guard AI is a modular, web-based application that processes multiple recruitment inputs through an interconnected detection pipeline. The system accepts four input types: job descriptions and recruiter messages, job posting URLs, and uploaded offer letters processed via Tesseract OCR.

The Data Preprocessing Module cleans, tokenizes, and normalizes all text inputs, removing stop words and applying stemming and lemmatization. URL inputs undergo domain metadata extraction and link structure parsing. The Machine Learning Classification Module applies Logistic Regression, Random Forest, and SVM trained on 6,000 samples (3,200 legitimate, 2,800 fraudulent) using a 70/30 train-test split with 5-fold cross-validation.

The LLM Reasoning Engine performs deep contextual and semantic analysis of recruitment content, identifying misleading language and deceptive communication patterns beyond the reach of traditional classifiers. The RAG Module retrieves verified information from external knowledge sources including company databases, known scam pattern repositories, domain reputation systems, and cybersecurity intelligence feeds, grounding LLM predictions in factual context.

The Trust Score Computation Module combines outputs from all analysis layers using the weighted formula:

$$T = w_1 \cdot \text{SML} + w_2 \cdot \text{SURL} + w_3 \cdot \text{SOCR} + w_4 \cdot \text{SLLM}$$

where $w_1 + w_2 + w_3 + w_4 = 1$. A job is classified as Legitimate when $T \geq 0.5$ and Fraudulent when $T < 0.5$. The weight configuration and score components are detailed in Table 2. In addition to detection, the platform provides AI-based interview preparation, calendar reminders, performance tracking, and AI-assisted email generation for job seekers.

METHODOLOGY

The system follows a layered pipeline architecture. User inputs are ingested through the web interface and passed to the preprocessing module for text normalization. For uploaded documents, Tesseract

OCR extracts textual content enabling fraud analysis of offer letters. Feature engineering then extracts text-based features such as suspicious keywords and salary anomaly indicators, network-context features including email domain credibility and URL metadata, and OCR-derived document features covering logo consistency and signatory verification.

The ML ensemble generates an initial prediction score. Concurrently, the LLM Reasoning Engine analyzes full text for semantic fraud indicators. The RAG module retrieves corroborating or contradicting evidence from external knowledge bases. All module outputs are fused in the Trust Score Module, producing a final explainable classification with flagged suspicious indicators highlighted for the user.

RESULTS

The system was evaluated on a dataset of 6,000 samples collected from public job portals, phishing datasets, and scam report repositories. Table 1 presents the performance comparison between the proposed integrated system and individual baseline models. RecruitGuard AI achieves 91% accuracy, outperforming standalone Logistic Regression (82%), SVM (85%), and Random Forest (86%) explanations. The HR module has improved the process of shortlisting candidates by using the insights gained.

Table – 1 Performance Metrics Comparison

Model	Acc (%)	Prec (%)	Rec (%)	F1 (%)
Logistic Regression	82	80	79	79
Random Forest	86	84	83	83
Support Vector Machine	85	83	82	82

The integration of LLM reasoning and RAG-based knowledge retrieval provides significant improvement over standalone models. Real-world validation confirmed correct detection across all major fraud scenarios: salary anomaly patterns, payment request keywords, phishing URLs, and OCR-based offer letter inconsistency detection. Legitimate postings with normal indicators were also correctly classified, confirming low false-positive rates.

Table – 2 Trust Score Weight Configuration

Component	Description	Weight	Threshold
SML	ML Classifier Score	w1 = 0.30	T ≥ 0.5 → Legit
SURL	URL & Domain Score	w2 = 0.25	T < 0.5 → Fraud
SOCR	Offer Letter Score	w3 = 0.15	—
SLLM	LLM-RAG Reasoning Score	w4 = 0.30	—

DISCUSSION

The proposed multi-layer architecture demonstrates that combining ML, LLM, and RAG techniques yields superior performance compared to any single-method approach. The ML module efficiently identifies structured fraud patterns such as suspicious keywords and salary anomalies. The LLM Reasoning Engine captures contextual and semantic deception that evades rule-based systems. The RAG module reduces hallucination risk by grounding predictions in verified external knowledge, a key limitation of standalone generative AI models [14].

The explainable trust score addresses a major limitation of existing systems by providing detailed reasoning for each classification decision, improving user trust and transparency. The integrated user-support features transform RecruitGuard AI into a comprehensive career assistance platform beyond mere detection, serving the safety and preparedness needs of job seekers simultaneously. Current limitations include dependence on training data quality, challenges with highly sophisticated scams that precisely mimic legitimate postings, and requirements for stable internet connectivity for API-based LLM services.

CONCLUSIONS

In this paper, RecruitGuard AI was proposed, a multi-source recruitment scam detection system integrating Machine Learning, Large Language Models, and Retrieval-Augmented Generation. The system analyzes job descriptions, recruiter messages, URLs, and offer letters to generate explainable trust scores achieving 91% accuracy, 89% precision, 87% recall, and 88% F1-score, outperforming all baseline approaches.

Future enhancements include integration with real-time company verification APIs and official corporate registries, incorporation of

BERT and GPT models fine-tuned for recruitment fraud, mobile application development for real-time push alerts, multilingual and voice input support, direct integration with major job portals such as LinkedIn and Naukri, and federated learning for privacy-preserving distributed model training.

REFERENCES

[1] V. B. Reddy et al., "Fake Job Recruitment Detection Using Machine Learning Approach," International Journal of Engineering Trends, 2025.
[2] G. Singh et al., "Advanced Real-Time Fraud Detection Using RAG-Based LLMs," arXiv:2501.15290, 2025.
[3] A. H. Al Barwani et al., "User-Centric Phishing Detection: A RAG and LLM-Based Approach," arXiv, 2026.
[4] C. W. Chang et al., "Scam Shield: Multi-Model Voting and Fine-Tuned LLMs Against Adversarial Attacks," arXiv, 2025.
[5] P. Kumar and R. Sharma, "Fake Job Detection Using Machine Learning and Deep Learning," International Journal of AI Research, 2024.
[6] A. S. Pillai, "Detecting Fake Job Postings Using Bidirectional LSTM," arXiv preprint, 2023.
[7] A. H. M. Hanif et al., "Machine Learning Approach in Predicting Fraudulent Job Advertisement," International Journal of Academic Research, 2024.
[8] L. Zhang et al., "Deep Learning-Based Detection of Online Recruitment Fraud," Elsevier Journal of Information Security, 2024.
[9] A. Sowmya et al., "Detection of Online Scam Messages Using NLP Techniques," International Journal of Cybersecurity, 2023.
[10] T. Nguyen and P. Tran, "Phishing Email Detection Using Machine Learning and NLP," IEEE International Conference on Cyber Security, 2021.
[11] J. Brown and M. Smith, "Fake Job Posting Detection Using Natural Language Processing," Proc. ACM Conference on Data Science, 2023.
[12] A. Kumar et al., "Detection of Fraudulent Job Advertisements Using Machine Learning," IEEE Access, 2022.
[13] P. Lewis et al., "Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks," NeurIPS, 2020.
[14] J. Devlin et al., "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," NAACL-HLT, 2019.