



INTERDISCIPLINARY MAJOR PROJECT

FINAL REPORT

B.E.W.A.R.E: Battling Employment Woes and Recruitment Exploit

Submitted on Dec 10, 2025

In Partial Fulfilment of the Requirements for the UG Program 2022-2026

FLAME University, Pune

CERTIFICATE

This is to certify that the work incorporated in this Interdisciplinary Major Project report titled “B.E.W.A.R.E: Battling Employment Woes and Recruitment Exploit” submitted by the undersigned project team was carried out under my mentorship. All such material (like journal articles, book chapters, videos, photographs, websites, etc) that was referenced from secondary sources has been duly cited and acknowledged.

S.No	Name of the Student	Designation	Signature
1.	Aaryan Mehta	Group Representative	
2.	Ameesha Talreja	Group Member	
3.	Dhriti Chheda	Group Member	
4.	Moksha Shah	Group Member	
5.	Neha Salvi	Group Member	
6.	Rachna Govindayapalli	Group Member	
7.	Shanaya Bajaj	Group Member	

Name of the Faculty Advisor: Prof. Snehal Shekatkar

Signature of the Faculty Advisor:



Name of the Subject Matter Expert (SME): Prof. Arnab Chakrabarti

Signature of the Subject Matter Expert (SME):



Date: 10th December, 2025

Abstract

Fake job postings have become a serious problem for students and recent graduates. These scams appear on job portals, WhatsApp groups, and social media, leading people to lose money, waste time, and lose confidence in legitimate hiring platforms. Despite how widespread this issue is, there aren't many tools that both detect fake postings and help job seekers avoid them.

Our project tackles this problem by combining computer science, economics, marketing, and data science. We're building a machine learning model using a dataset of real and fake job postings to flag suspicious listings. Additionally, we will create awareness materials that actually connect with students: a strategically designed social media campaign that helps identify fraudulent job listings in real-time. We're also developing an interactive dashboard that visualises scam patterns across different industries and locations, giving users a clear picture of how these scams operate.

The final deliverable is a Data-Driven Awareness Platform with three parts: a scam-detection model, a student-friendly educational campaign, and a visual dashboard.

With the help of this project, job seekers will be better equipped to handle the complex job market of today with greater awareness and confidence.

Table of Contents

Introduction.....	4
Context of the Study.....	4
Rationale for the Problem Chosen.....	5
Research into developing Fraudulent Job Listing Detection Models.....	6
The Economic Impact of Fake Job Postings.....	8
Identification of Gaps in Literature.....	9
Objectives.....	11
Data Analysis.....	12
Quantitative Findings.....	12
Qualitative Patterns observed:.....	13
Solution Concept and Design.....	14
Model.....	14
Social Media Campaign.....	18
Results and Evidence.....	21
For the Campaign.....	26
Interdisciplinary Integration.....	27
Solution Contribution to SDGs.....	29
Reflection and Limitations.....	30
References.....	32
Appendix.....	33

Introduction

Context of the Study

In the current job market, digital platforms like LinkedIn, Naukri, Internshala, WhatsApp groups, and other social media platforms have become popular channels for job-seeking. However, fake job postings have become one of the most common threats. Victims often end up paying application or training fees, disclosing sensitive personal data, or simply wasting time on processes that never lead to employment.

This problem is quite significant, and according to reports by the Indian Computer Emergency Response Team (CERT-In), job frauds have become the fastest-growing type of cybercrime in India. According to a study done by the Indian Ministry of Labour in 2023, 1 out of 5 fresh graduates encounter fraudulent postings in their first year of job searching. This issue is quite widespread. In 2022, the Federal Trade Commission (FTC) in the United States reported that over \$86 million had been lost as a result of employment and job scams. These frauds not only take advantage of people but also damage the reputation of digital hiring platforms, which are otherwise essential for fair and effective labour markets.

Rationale for the Problem Chosen

In recent times, while awareness about email scams and financial scams has grown, the area of fake job postings and job scams is under-researched and not adequately dealt with by institutions and policymakers. The systems that do exist, like reporting options on job portals, tend to work

only after someone has already fallen victim. Moreover, the majority of current solutions only focus on detecting technical fraud without consideration for the psychological, economic, or social vulnerabilities that may put job seekers at even greater risk.

Fresh graduates, many of whom are navigating employment processes for the first time, are particularly at risk. For first-generation workers, the stakes are even higher: scams take advantage of their limited financial resources, reducing their confidence and they fall prey. Its effects extend beyond individuals, affecting families and weakening the job recruitment credibility.

Literature Review

Research into developing Fraudulent Job Listing Detection Models

The detection of fraudulent job listings can be accomplished by the use of intelligent Machine Learning models, trained with existing job listing datasets, so as to reliably detect and flag these listings online. This problem is closely related to online recruitment fraud (ORF).

Various existing studies explore the use of machine learning and text analysis in developing reliable detection models for job listing fraud, and reviewing such studies helps understand which approaches are most suitable for the development of our project.

In a study by Marcel Naude et. al, an array of classification models is used in combination with empirical rule set features. Approaching a similar ORF problem, they identified Gradient Boosting to be the best performing model when combined with the empirical rule set features, parts-of-speech tags, and bag of words vectors. The evaluation metrics returned an F1 score of ~0.88. This study highlights the benefits of the empirical rule set features in building a reliable model, showcasing the importance of using handcrafted features.

K. Swetha et al. use a similar approach, applying a range of common classifiers in order to ascertain the best-performing model. This study, instead of using text data, has used the categorical characteristics of the EMSCAD dataset - the authors observed that this method efficiently lowered the number of trainable attributes, hence reducing the complexity of the model.

In this approach, they found that the Deep Neural Network Classifier outperformed other models, giving a prediction accuracy of approximately 98%. This sheds light upon a different approach to our project, which uses the categorical attributes of the dataset, that efficiently and quickly lowers the model's complexity.

The dataset used in these papers, EMSCAD, is sourced from the initial study Vidros et al. (2017), which created the EMSCAD dataset and primarily used the Random Forest Classifier. This set the foundation for ORF studies that emerged later, providing a stable and reliable dataset and background for the problem.

Moreover, these papers highlight the importance of specific practices and how they impact the model's reliability and performance. Looking over the various classifiers used in these studies, we can narrow down certain classifiers that performed better than others - a study by Fahad Alharby showcases the successful performance of Random Forest classifiers, while the study by K. Swetha et al. highlighted their inclination towards the Deep Neural Network Classifier. From these papers, the Random Forest Classifier, Deep Neural Network Classifier, and Gradient Boosting (in combination with empirical rule set features) performed the best when applied to these ORF problems. Comprehensive analysis of these studies helps us narrow down the best methodology to apply to our project, in order to take forward the formulation of reliable job fraud detection models for a safer, more secure recruitment experience.

The Economic Impact of Fake Job Postings

The economic consequences of fake job postings and employment scams represent a rapidly escalating crisis that affects vulnerable jobseekers, particularly fresh graduates and first-time applicants. According to the Federal Trade Commission (FTC), reported losses to job scams tripled between 2020 and 2023, surpassing 220 million USD in just the first six months of 2024, with task scams accounting for nearly 40% of these losses (FTC, 2024). Beyond the direct financial losses incurred by victims who are manipulated into investing their own money through various schemes, the broad economic impact includes significant opportunity costs, including

wasted time spent on fraudulent applications, delayed entry into legitimate employment and diminished human capital development during vital early career stages.

A qualitative study by Alexandra J Ravenelle et al. examines how job seekers during the COVID-19 pandemic encountered and responded to fraudulent job postings on digital platforms. The study revealed that major job platforms like LinkedIn estimate approximately one per cent of their listings are fraudulent (roughly 60,000 postings at any given time), and that the burden of scam detection has been shifted from platforms and law enforcement officials to individual job seekers themselves. The research documents how exposure to repeated scams leads to profound psychological and behavioural impacts: workers either adopt a volume-based approach by applying indiscriminately to more positions or experience discouragement and withdraw from job searching altogether, some rejecting legitimate offers themselves. Despite the study involving the COVID-19 pandemic, the concerns are still legitimate as jobs are still sought online. This project provides empirical evidence of the real-world opportunity cost and behavioural consequences of job scams and extends beyond direct financial losses. Furthermore, the study's documentation of how job seekers develop 'detection strategies' yet still fall victim to sophisticated scams underscores the need for technological detection tools.

An empirical study by Sharoom Ansar et al. employs a cross-sectional online survey (n=284) to examine how exposure to online shop scams affects users' trust in and reliance on digital job search platforms. A key finding indicates that victims' willingness to use a job board has reduced by approximately 40% following scam exposure, representing a substantial behavioural shift with economic implications. The study emphasises that the economic burden is disproportionately concentrated among recent graduates, low-income individuals, and those with limited digital literacy and references aggregate losses of hundreds of millions of dollars

annually to job scams. This paper is highly relevant to our study as it provides empirical evidence of two distinct economic channels through which job scams impose costs: direct monetary losses and opportunity costs arising from behavioural responses to scam exposure. These findings support our approach of quantifying direct financial losses and opportunity costs to demonstrate the regressive nature of jobs, scam burdens, and strengthen our case for targeted interventions and awareness campaigns.

Identification of Gaps in Literature

1. Dataset Limitations: The chosen dataset, while relevant to the model's training, contains data mainly from the years 2014 to 2016. A majority of the literature surrounding ORF studies uses this dataset to train their models, published in the years 2019 - 2023, and so the models may not be able to capture language patterns used in current job listings as efficiently as they would have if it had been trained on a more updated dataset.
2. Geographic and Cultural Contextualisation: The existing literature on job scams and economic impact is predominantly centred on Western countries, particularly the United States, creating a significant knowledge gap regarding how fraudulent job postings manifest in emerging markets such as India. The vulnerability profiles of job seekers combined by factors such as education systems, family employment histories and digital infrastructure access differ substantially across geographic contexts, yet existing studies do not adequately capture these variations.
3. Low Focus on User-Centred Design: Existing literature focuses almost exclusively on model performance without addressing how these tools would be accessed and understood by actual job seekers. This substantial gap shows that even highly accurate

models may fail to protect those most at risk if insights cannot be communicated effectively.

Objectives

- To build a supervised machine learning model that is capable of identifying fake job postings and preventing job seekers from such scams.
- To estimate the opportunity and financial costs of fake job postings and emphasise the impact on fresh graduates and first-generation job seekers.
- Creating accessible awareness-based content on social media to educate people on how to identify fraudulent postings.
- To highlight scam patterns across regions and jobs with the help of an interactive dashboard.
- To ultimately provide insights along with a tool to help candidates make a more informed decision.

Data Analysis

Quantitative Findings

The dataset used for the model consists of 27,880 postings, which merges the Fake job posting dataset (10,000 records) and the EMSCAD dataset (17,880 records). Post cleaning, 10,866 postings (39%) were labelled as fraudulent and 17,014 (61%) legitimate. This is significantly higher than the scams present in the real world (usually under 5%). This balance promotes better learning outcomes by giving the model enough examples of both types.

Statistical Insights:

- Frauds cannot be accurately predicted by a single feature. For example, payment requests appear in only 68% of frauds, fake partnerships in 29%, and dodgy contact details in 34%. For the model to detect effectively, a combination of multiple features are required.
- Not all frauds have the same patterns; 43% were quite obvious, 38% professionally written but suspicious, and 19% were indistinguishable from real postings. This creates the need for probability-based decisions rather than binary classification.
- False positives are worse than false negatives, especially when real opportunities are fewer, impacting threshold choice to balance recall with user trust.

Qualitative Patterns observed:

- Geographical context plays a role. In India, fraud patterns caught local terms such as UPI references, walk-in interviews, and hiring through WhatsApp. On a global scale, fraud detection rules miss these hints.

- With time, scam patterns evolve. After comparing both datasets, we saw changes in language, tone, fewer mentions of fees, and increased use of direct messaging platforms, requiring future training of the model.

Solution Concept and Design

Model

It comprises three different components:

- i) a hybrid machine-learning model for detecting fraudulent job postings.
- ii) the BEWARE web application for simple user interaction.
- ii) an analytical dashboard that visualises broader fraud patterns

Different components were used to separate computational logic from the presentation, allowing an effective model execution and clear user interface. The model generates fraud probabilities and explanations, the web application communicates these insights to end users, and the dashboard gives descriptive analytics that inspect fraud activity across sectors and regions.

Machine Learning Framework

Feature Engineering

Each job posting is represented using 427 features made up of

- 43 engineered fraud indicator features
- 384 semantic embedding features generated using the *all-MiniLM-L6-v2* SentenceTransformer.

Fraud Indicator Features (43 features)

These features are identifiable characteristics of fraud mainly in the Indian context.

1. Weighted detection of terms such as registration fee, security deposit, UPI-based payment requests, and digital wallet mentions. Combined occurrences increase the fraud signal.

These features constantly rank among the strongest predictors.

2. It detects fraudulent claims of connection with reputable organisations through phrases such as authorised partner, tie-up with or collaboration with Flipkart and Amazon, as such. This differentiates legitimate from untrue endorsements.
3. It identifies common scam lingo like data entry, form filling, copy and paste work, and typing jobs, with additional weighting when multiple terms appear simultaneously.
4. Regex-based extraction detects exaggerated or guaranteed income statements, for eg, statements such as earn ₹5000/day, assured payout, guaranteed income.
5. Captures communication through WhatsApp only, absence of a professional email, ten-digit mobile numbers, shortened URLs, and informal messaging prompts. These features represent high-risk features.
6. Includes exclamation mark frequency, uppercase ratio, repeated punctuation, readability scores, extremely short or excessively long posts, and sentence-length structure. These stylistic indicators can be used to separate informal scam messages from legitimate communication.

Model Selection

LightGBM outperformed alternative models such as LSTM, BERT, Random Forest, and XGBoost by speed, memory efficiency, and explainability. It trains substantially faster than deep learning models and integrates directly with SHAP, enabling transparent prediction analysis. This made it suitable for high-feature and mixed-signal fraud detection.

BEWARE Application

The BEWARE interface transforms model outputs into understandable information for the users

Input Methods

- **Text input:** Direct pasting from job portals, emails, or messages
- **Image input:** OCR-based extraction from flyers, screenshots, or photos, or job posting PDFs

Output Components

1. **Fraud Probability Gauge** showing four risk zones (low, moderate, high, critical).
2. **Binary Verdict** was fraudulent or legitimate
3. **Confidence Level** derived from prediction margin
4. **Red Flag System** highlighting critical, high-risk risk or moderate indicators.
5. **Actionable Guidance** aligned with risk severity.
6. **SHAP Explanations** identifying top contributing features per prediction.

These outputs allow users to evaluate postings even without the technical knowledge.

Analytical Dashboard

The dashboard, constructed using insights from the EMSCAD dataset, provides descriptive analytics including:

- Fraud distribution across industries.
- Prevalence across job types (full-time, remote, contract-based).
- Geographic concentration of fraudulent postings.
- Comparison of description length and completeness.
- Missing information ratios.

These visualisations support user understanding and validate patterns learned by the model.

Design Choices

i) **LightGBM** was used for its ability to handle a high-dimensional mixed feature space capably, its strong performance on tabular data and its compatibility with SHAP for interpretability. It offered accuracy comparable to BERT-based models while reducing training time a lot, as well making it preferable for repeated testing.

ii) **Hybrid feature** strategy, which is engineered fraud indicators combined with 384-dimensional MiniLM embeddings, was used to capture both explicit “scam signals” and semantic patterns not detectable through rule-based features alone. This mix improved generalisation across posting types and enabled the model to detect subtly hidden fraud patterns.

iii) **Multiple datasets** approach broadened the distribution of fraud behaviours available during training. Adding the dataset source indicators reduced the area shift by allowing the model to adjust for structural and linguistic differences across both our datasets, hence giving us more stable performance across both.

iv) **Threshold** selection was treated as a model parameter rather than just a fixed constant. The final threshold (0.45) was based on maximising fraud recall while keeping the false positive rates extremely low. Additional thresholds were retained to support different operational sensitivity needs.

v) **Optuna** was used to optimise significant LightGBM parameters due to its efficient guided search and ability to balance AUC and fraud recall-oriented objectives. This approach avoided a long, time-consuming grid search process while consistently producing higher-performing configurations.

Social Media Campaign

Scams usually target volume and attention, which spreads quickly and preys on gaps in awareness. A campus centric, social media campaign leverages existing social networks and high share ability to surface scam patterns, build awareness and tunnel students towards using our ML tool for verification. Our primary audience consists of Flame students who are undergraduates and fresh graduates, campus job seekers, and internship applicants; Focusing on Flame first creates a tight testing ground for the campaign, before scaling outwards.

Our core idea consists of using bite-size, audience-focused storytelling on Instagram. Each post, reel and story is designed to be instantly useful as it informs students of our model being accessible and easy to use, showcasing red flags, including a real micro example. Another core aspect of our campaign was the on-campus activation, a physical interactive booth manned by our team, that allowed us to connect with the student population of FLAME directly. This way, we were able to have a stronger impact on our target audience, as well as directly assess the effectiveness of our social media campaign thus far.

Upon meeting with Mr. Padmanava Das From Flame's Career Service Office, we received useful insights that not only helped us shape our campaign directed at Flame students but also receive in depth knowledge ourselves about fraudulent job postings in the market. He also shared fraudulent job listing data that CSO had come across this year to aid our project.

With the help of the campaign, we wish to drive increased usage of the verification tool from within Flame and higher rate of scam reports to campus authorities. Reach and engagement are the next success signals that would indicate active learning.

Implementation

Before we began the campaign execution, we outlined a tentative content calendar - a timeline of when each content creative would be posted, or organized. This helped us understand which content pieces needed to be ready by when, in order to execute the campaign in a timely manner. We chose our on-campus activation date strategically, a few weeks into the social media campaign; this allowed us to first establish a strong base presence, so that we could attract more students to our booth based on recognition and interest.

We strategised our timeline in a way that each of our posts would be able to garner broad reach and engagement. Shares and comments, especially, are key drivers to building reach, so we focused on creating shareable content in order to expand our reach and page visibility. We decided to stretch the campaign across one month, with 3 to 4 key posts assigned per week - to ensure regularity, while avoiding overload of content.

Our posts were also designed in a uniform manner, the design choices reflecting the brand image, tone and personality. We chose the colours neon yellow, charcoal black, and white as our brand palette, to make our posts visually impactful. These colours attract audiences' interest, while also maintaining the serious undertone of our posts. We kept our language friendly and energetic in order to resonate with our target group, so that our content would be fully understood and appreciated; we kept the content language simple and screen reader friendly.

By starting hyperlocal at FLAME, we were able to build trust and iterate content based on campus feedback. As messaging and formats are validated, we can further target other universities and communities, preserving the peer led tone while adjusting language and platforms for new audiences.

Results and Evidence

Model Performance

The final LightGBM model trained on the combined dataset achieved consistently high performance . The extended test set (n = 6,970), the model showed the following metrics:

Metric	Value	Interpretation
AUC-ROC	0.9980	Almost perfect distinction between fraudulent and legitimate postings.
F1 Score	0.9811	Under the set threshold, very good balance between precision and recall.
Accuracy	98.6%	Correctly classifies 6,886 out of 6,970 cases.

Fraud-specific performance

Metric	Value	Interpretation
--------	-------	----------------

Recall (Frauds detected)	96.4%	2,618 out of 2,716 fraudulent postings were correctly identified.
False Alarms	0.1%	Only 3 legitimate postings were misclassified as fraud.
Precision	99.9%	When the model predicts fraud, the majority of the time this prediction is correct.

These performance metrics show a greater improvement compared to the one-class approaches used earlier, which detected only 5.3% of fraud. The confusion matrix below illustrates the classification behaviour:

Predicted	Legitimate	Predicted Fraud
Actually Legitimate	4251	3
Actually Fraud	98	2618

This tells us that the model has low false positives (3 cases), and a high fraud detection rate (2,618 cases), remaining errors are focused towards false negatives (98 cases), which represent newer or unusual fraud patterns.

Insights:

To quantify features, we used SHAP analysis. The top predictors include both features from the dataset and engineered.

Rank	Feature	Importance	Category
1	source_fake_postings	64,528	Dataset metadata
2	source_recruitment	6,168	Dataset metadata
3	low_barrier_count	5,152	Fraud heuristic
4	word_count	4,254	Structural indicator
5	has_logo	1,711	Metadata
6	wfh_keywords_count	-	Fraud heuristic
7	avg_sentence_length	-	Stylometry
8	avg_word_length	-	Stylometry
9	embed_157	-	Semantic embedding
10	text_length	-	Structural indicator

- Manually engineered fraud features are positioned higher, reassuring the effectiveness of designing domain-based features.
- Complex indicators (low_barrier_count) are more predictive than individual keywords.
- Embedding features individually stand weak, but are meaningful, unique and unusual frauds.
- Metadata features (has_logo) hold a significant predictive value since a lot of fraudulent postings lack this feature.

Comparison to Alternative Approaches

Method	Frauds Caught	False Alarms	Training Time
Keyword Matching	43.2%	8.3%	Instant
Rule Based	58.7%	12.1%	N/A
Random Forest	91.4%	2.1%	8 mins
XG Boost	93.2%	1.2%	5 mins
LightGBM (ours)	96.4%	0.1%	3 mins
BERT	94.1%	0.8%	45 mins

The proposed LightGBM model gives us the best balance between accuracy, computational efficiency, and understanding and training time.

Output Variability

When analysing images, the system first uses OCR and then extracted text can vary slightly depending on lighting, angle or image quality. Because of this fraud probabilities from image inputs may fluctuate ($\pm 3-7\%$) which is expected. This variability does not affect typed or pasted text, where predictions remain stable and consistent.

Real-World Impact

If we have 1,000 monthly users and a real-world fraud rate of 4%:

- Approximately 40 individuals would encounter fraudulent postings.
- With a 96.4% detection rate, the system would prevent around 38-39 fraudulent crossings per month.
- Averaging victim losses of ₹15,000 to 50,000, this corresponds to ₹5.7–19.2 lakh in prevented losses monthly.

For the Campaign

The two main objectives the social media campaign addressed were: creating accessible awareness content and helping candidates make more informed choices. To meet these expectations, we approached them with both qualitative and quantitative methods. We created ten posts ranging from informational carousels to reels and on-ground activation. Instagram insight since late November shows that we have received around 22,000 views and nearly 800 interactions, such as likes, comments, shares and around 50 followers expressing interest in under a month. A content explicitly explains detection signals across the posts, and the traffic shows initial interest. High engagement reels, user-generated content, such as crossword

participation and the on-site stall conversations, indicate active attention and pure discussion. On the other hand, testimonials and reaction videos demonstrate increased awareness and intent to verify listings.

We faced a couple of limitations, such as a short timeframe, the target audience being campus-skewed, a small qualitative sample and self-selection bias in testimonials and the lack of and guarantee that engagement equals change in behaviour among our peers.

Interdisciplinary Integration

Detecting fake job postings requires more than just technology, as it demands understanding economics, human behaviour, and effective communication. Our project integrates multiple disciplines to create a comprehensive solution.

Computer Science, Data Science, and Business Analytics form the technical core. We have used machine learning algorithms and natural language processing to automatically identify patterns in fraudulent postings. Business analytics helps us structure the data and focus on features that matter in real-world applications, not just theoretical models.

As for the project campaign design, it combines insights from marketing, advertising, branding, economics and visual design. Marketing perspectives frame the problem of employment scams as one requiring awareness, repeated exposure and relatable messaging. Advertising and branding shaped our campaign strategy, choosing Instagram as the core distribution channel and guided the development of content format to maximise reach, engagement and simplification of complex information. Economics added depth to the approach by explaining why awareness generation is itself a meaningful intervention. Visual design translated these conceptual inputs into clean and accessible formats that Gen-Z users can digest quickly.

A key integration challenge was supplying economics without quantitative data sets. We resolve this by shifting from econometrics to conceptual economics and using theory to justify campaign mechanics rather than trying to run numerical analysis.

Solution Contribution to SDGs

Our solution contributes to SDG 8 by reducing employment fraud and supporting fair access to genuine job opportunities, especially for vulnerable young applicants. It also supports SDG 4 by improving digital literacy through scam awareness education, empowering students to make informed and safer employment decisions.

Reflection and Limitations

The social media campaign demonstrates that simplified behavioural insights can make fraud awareness more accessible and actionable. Peer-driven formats, humour and UGC boosted engagement, helping us make the topic less intimidating. However, our biggest limitation was attribution, as we cannot yet confirm whether the event is translated into safer job decisions. Engagement metrics are proxies and thus not behavioural proof. Time constraints, campus-centred sampling and the absence of economic data limited deeper analysis. A future iteration should add measurable CTAs, longer tracking and a structured feedback service to assess actual behavioural change rather than online interaction.

As for the model, here are several additional factors that can affect how the it works:

- Fraud patterns shift and evolve, particularly in the Indian job-fraud market. Regular retraining and frequent updating are required.
- The model can only understand English text. Postings in Hindi or any other Indian regional languages might not register.
- Frauds that are visually unstructured (unprofessional layout, manipulated logos) are not modelled unless the OCR module is used; even then, OCR comes with noise.
- The 98 missed frauds could be underrepresented patterns or newly introduced scamming techniques.
- Dataset Imbalance Relative to Real-World Rates

- The combined dataset has a very high percentage (39) for fraud cases, which is way above what is actually present in the job market. The desired model performance will require some fine-tuning.

References

- Alharby, F. (2019). An Intelligent Model for Online Recruitment Fraud Detection. *Journal of Information Security*. <https://doi.org/10.4236/JIS.2019.103009>
- Naudé, M., Adebayo, K.J. & Nanda, R. (2023). A machine learning approach to detecting fraudulent job types. *AI & Soc*, 38 (3), 1013–1024. <https://doi.org/10.1007/s00146-022-01469-0>
- Swetha, K., Reddy, M. T., Sravani, K., & Subramanyam, B. (n.d.). FAKE JOB DETECTION USING MACHINE LEARNING APPROACH. *Journal of Engineering Sciences*, 14(02). <https://jespublication.mlsoft.in/upload/2023-V14I209.pdf>
- Ansar, S., & Hussain, S. (2025). *Effect of Online Job Scams on Trust and Dependence on Digital Platform*. ResearchGate. https://www.researchgate.net/publication/395914252_Effect_of_Online_Job_Scams_on_Trust_and_Dependence_on_Digital_Platform <https://doi.org/10.63468/jpsa.3.3.86>
- Ravenelle, A., Janko, E., & Kowalski, K. C. (2022). *Good jobs, scam jobs: Detecting, normalizing, and internalizing online job scams during the COVID-19 pandemic*. University of Carolina Digital Repository. <https://cdr.lib.unc.edu/concern/articles/2f75rj88m?locale=en> <https://doi.org/10.17615/91b3-hg17>

Appendix

BEWARE: Spotting Fake Job Postings with Data

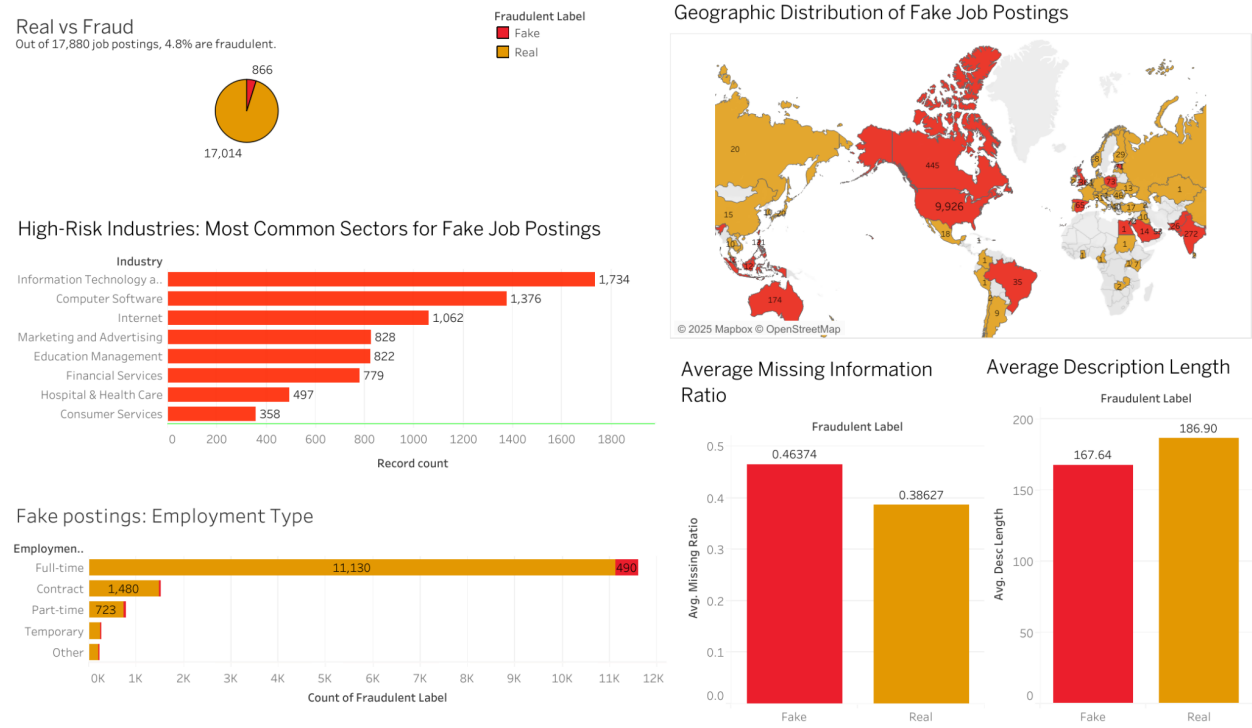


Fig 1:Interactive Dashboard

The Key Metrics of our Social Media Campaign (figures and images)

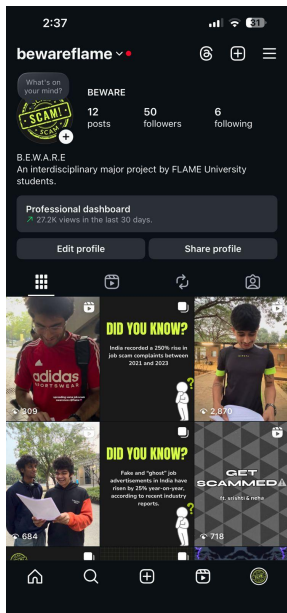


Fig 2: Our social media page

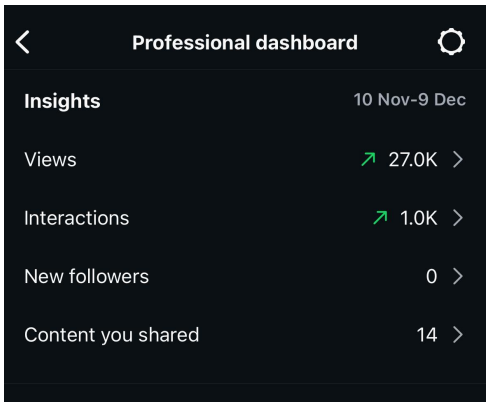


Fig 3: Our professional dashboard for the last month



Fig 4: No. of views received by our page

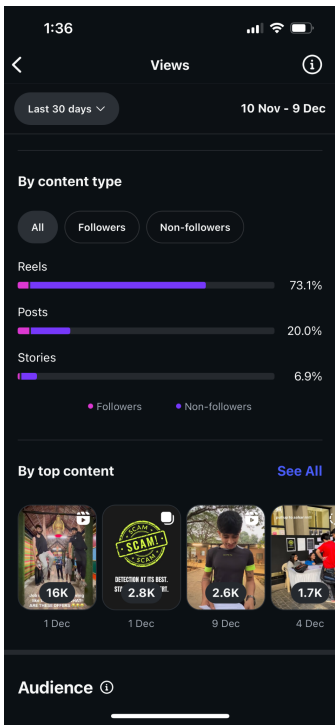


Fig 5: Our views based on the content we posted

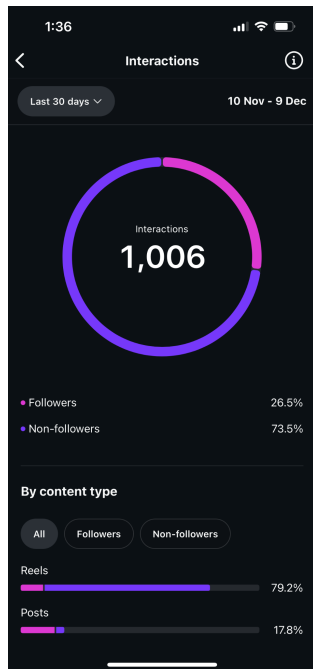


Fig 6: Overall Interactions (likes, shares, comments, reposts)

AI Disclosure Statement

AI tools were used in this project only for technical challenges in line with IMP 2025 guidelines. They were not used to generate any part of the report, design decisions, model interpretations, or analysis. All conceptual work, modelling choices, the written content, and system architecture were produced entirely by the student team.

The AI assistance provided was limited to help understanding and resolve specific implementation issues that occurred during the development , mainly backend errors, deployment configuration problems, frontend and backend connectivity. At no point did we use AI to produce explanations, summaries, or the content in this report .

Examples of AI Use

1. Backend Debugging

Prompt used:

“Uvicorn fails to start my FastAPI application and returns ‘AttributeError: module has no attribute app’. My main file is structured with an app = FastAPI() object. What changes should I make to ensure the server recognises the correct entry point?”

AI guidance received:

The AI identified that the issue came from pointing Uvicorn to the wrong module path. It recommended updating the server command to `uvicorn main:app --reload`, which resolved the startup error.

2. Frontend-Backend Integration

Prompt used:

“The frontend generated on Lovable cannot reach my model endpoint after deployment. The interface loads correctly, but the prediction request fails. How should I modify my fetch call to point to an externally deployed FastAPI endpoint instead of localhost?”

AI guidance received:

The AI explained that a deployed frontend cannot call a localhost URL. It suggested using an environment variable (e.g., `NEXT_PUBLIC_API_URL`) and provided an example of a fetch request pointing to the production endpoint. After updating the call, the integration worked correctly.

3. Deployment Diagnostics

Prompt used:

“My Railway deployment logs show a module import error for scikit-learn even though it installs fine locally. How do I ensure the platform installs all required Python dependencies?”

AI guidance received:

The AI noted that Railway builds its environment strictly from the `requirements.txt` file. It advised regenerating the file with explicit versions for all dependencies. After updating it, the deployment succeeded without further errors.

Verification and Human Oversight

All suggestions provided by AI were reviewed, adapted, and manually validated by the team before implementation. AI did not generate any system design, documentation text, analysis, or model related insights. It functioned solely as a technical assistant to resolve specific coding and deployment issues when required.