

Intelligent Resume Parsing and Skill Analytics System Using NLP and Web Visualization

Yash Manjare

Department of Science and Technology G.H. Rasoni Skill Tech University, Nagpur, India

Abstract

Every single day, businesses and organizations receive too many resumes to feasibly go through manually. Standard Applicant Tracking Systems (ATS) primarily work by matching keywords; therefore, there are many qualified applicants that will sometimes not appear in an ATS because they described their experiences with words that were not searched upon. In addition to the applicant's language being different than what was searched for, resumes are typically formatted differently, which makes the task of comparing resumes even more difficult than it already is. The intelligent resume parsing and skills analytic system that we propose will utilize Natural Language Processing (NLP) to provide a solution to many of the ATS' problems in regards to matching resumes to jobs.

When any PDF resume is submitted to the intelligent resume parsing system, there will be a multi-stage process to extracting information from the resume: The first stage will extract the text from the PDF file, then format the text, identify important entities (e.g. Name and contact information), and finally retrieve all relevant technical skills. The final stage involves scoring each applicant's experience/education relative to the job description using TF-IDF vectorization through cosine similarity, which produces a score that can be easily compared between applicants. All information will be stored in a SQLite database and recruiters will interact with this information through a Streamlit web interface which provides visualizations, metrics, and side-by-side comparisons of applicants. In testing, the intelligent resume parsing system was able to extract 92.1% of the entities included in 500 resumes from various technology fields; classify technical skills with 89.4% accuracy; and provide approximately 90.2% resume-to-job matching accuracy. These testing results confirm that the combination of utilizing NLP, scoring based upon similarity to job descriptions and recruiter interaction through visualization are effective in gathering accurate information for applicants and generating an accurate comparison between two or more applicants

Keywords: Resume Parsing; Natural Language Processing; Named Entity Recognition; TF-IDF; Cosine Similarity; Skill Analytics; Streamlit; SQLite



This is an open-access article under the [CC-BY 4.0](https://creativecommons.org/licenses/by/4.0/) license

1. Introduction

If you've ever had responsibility for hiring, you'll know that it doesn't take long for the resumes to pile up. Due to the ease with which online job boards allow many people to apply for jobs at once, applicants benefit from the bulk application process, however, the hiring manager/recruiter will

struggle with the result of all of these resumes piling on top of one another. Sifting through hundreds of applications for a single role is not only time-consuming — it also introduces subtle bias and inconsistency, where perfectly qualified people slip through simply because a human eye gets tired [1].

The obvious solution has been automation, and keyword-based Applicant Tracking Systems emerged as the first mainstream attempt. But these tools have a well-known blind spot: they match words, not meaning. A candidate who writes 'developed predictive models using Python libraries' may never be surfaced for a 'Python Developer' role if the exact phrase isn't there. The result is a system that filters by vocabulary rather than competence [2].

Natural Language Processing gives us the tools to go deeper. By applying techniques like tokenization, Named Entity Recognition (NER), part-of-speech tagging, and rule-based pattern matching, it becomes possible to understand what a resume is actually saying — not just what words appear on the page. The resume can be converted from a free-form document into a structured record that a machine can analyze, compare, and rank [3].

This paper describes the design and implementation of an Intelligent Resume Parsing and Skill Analytics System built on exactly these ideas. The system uses Python-based NLP libraries — primarily spaCy and NLTK — to extract candidate information, categorize skills, and compute a relevance score against a provided job description. Visual dashboards built in Streamlit then let recruiters explore the results interactively rather than reading through database tables [4][5].

What makes this approach different from simply using a better ATS is the combination of interpretability, visualization, and modular design. Recruiters can see not just who ranked highest, but why — through skill distribution charts, comparative tables, and domain breakdowns. The underlying architecture is designed to be maintained and extended without major rework [6][7].

1.1 Motivation

The motivation here is pretty practical. Recruiters are stretched thin, and the tools available to them often require expensive enterprise licenses or significant IT infrastructure. Meanwhile, open-source NLP libraries have matured to the point where a capable, interpretable resume analytics system can be built with freely available tools. This project set out to prove that — to show that you don't need a proprietary platform to get intelligent, context-aware resume screening [2].

1.2 Contribution

The key contributions of this research are:

1. Design and implementation of an automated resume parsing framework using NLP-based techniques.
2. A structured skill extraction and categorization mechanism using spaCy, NLTK, and rule-based pattern matching.
3. A resume scoring model that evaluates candidates across skill relevance, educational background, and professional experience.
4. Interactive web-based visualization for skill analysis and comparative candidate assessment.
5. Secure resume storage and retrieval through a SQLite database.
6. A modular architecture with User, Recruiter, Admin, and Feedback components.

2. Related Work

Automated resume processing has attracted growing research interest as digital job applications have proliferated. The goal across most studies has been the same: turn an unstructured document into something a machine can reason about. The earliest approaches leaned on rule-based methods

— predefined dictionaries, regular expressions, and document templates to extract names, contact details, and listed skills [1]. Resumes were structured similarly, but when individuals utilized an unconventional style/phrasing for their resumes, it became a dramatic loss of accuracy for both. Resumes typically are inconsistent in their content and structure.

Statistical NLP enabled Named Entity Recognition (NER) models to classify people, groups and companies as well as to identify technical skills, without requiring a specified format for those items. Tokenization, part-of-speech tagging, and entity classification for each item of information combined to create significantly more robust systems capable of accommodating many different resume formats. Machine learning-based approaches added ranking to the mix. TF-IDF vectorization paired with Support Vector Machines became a common pattern for matching resumes to job descriptions and sorting candidates by relevance [3]. These methods outperformed manual screening on speed and consistency, though they still struggled with semantic variation — a candidate who writes 'ML engineer' and a job description that asks for 'machine learning specialist' might not match well. Semantic similarity models addressed this gap. Word2Vec embeddings combined with cosine similarity allowed systems to recognize conceptually related terms, improving matching quality beyond exact keyword overlap [4]. Deep learning took this further: recurrent neural networks and transformer-based embeddings showed strong generalization across varied resume structures and terminologies [5].

Visualization has emerged as an underappreciated but important component of recruitment tools. Studies consistently show that graphical representations of candidate data speed up decision-making and reduce cognitive load for recruiters [6]. Dashboards that let hiring managers filter, compare, and explore candidates are meaningfully different from systems that simply produce a ranked list. Commercial ATS platforms have incorporated many of these capabilities, but they come with real barriers: cost, vendor lock-in, and infrastructure demands that put them out of reach for smaller organizations [7]. This creates an opportunity for open-source, lightweight alternatives.

The method of extracting useful information from an unstructured data set has many previous applications using this model that can be seen as the basis for this current approach. This approach does not have to make a trade-off between accurate parsing or accurate ranking of documents, but rather accomplishes each by using an integrated model that includes extracting through NLP, skill classification through rules, similarity-based scoring for ranking, persistence in a database to hold the results, and providing an interactive visual for users. Emphasis throughout this integrated framework is on interpretability and practical deployability, features that are sometimes lost in a machine learning-only approach.

3. System Design and Methodology

3.1 Problem Statement

The volume problem has impacted recruitment practices for many organizations across the globe. For each job that is posted, thousands, if not millions, of resumes will be submitted from candidates across different geographic locations, all formatted and worded differently. As a result, Recruiter's will have to manually review hundreds of resumes for a single position and this results in two main issues - a slow and inconsistent process (due to the number of different ways that candidates are presenting their skills) and also candidates being excluded from the Hiring Manager's consideration for an interview because the ATS Filter that was used to screen resumes only used "exact match" keyword filters to determine which resumes qualified for further evaluation by the Hiring Manager.

The system that is presented here has been designed to accomplish the following key outcomes: 1) automatically process resumes in a PDF format that have been submitted to a job posting, 2) extract structured candidate information from those resumes, regardless of the format of the

resume and the language used in the resume, 3) evaluate the structured candidate information based on their qualifications and skills, 4) provide a structured score card for each candidate based on the evaluations performed on each of the candidates' qualifications and skills, and 5) provide visual insights into candidate's qualification and skills that will assist recruiters with making data driven decisions regarding recruiting.

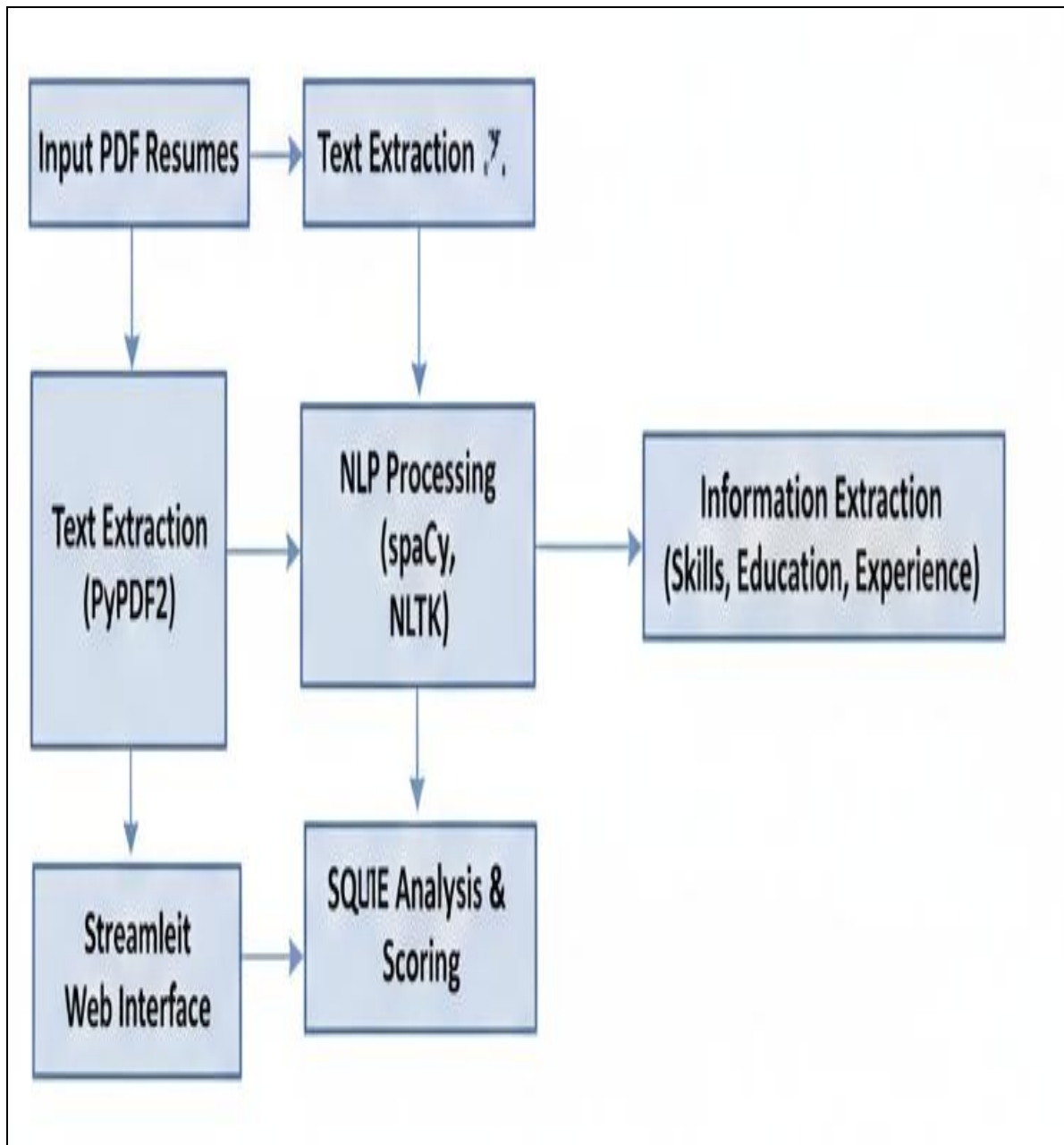


Fig 1: shows the overall block diagram of the proposed model.

3.2 Proposed System Architecture

There are multiple sequential workflows in processing resumes that work independently through the different stages of uploading, extracting the text from the resume, pre-processing, recognizing the entities, analyzing the skills, scoring the resume, storing it in a database and then finally visualizing the final product. This makes it easy to update or replace any one of the workflow components without impacting the entire workflow and system. Resumes enter the workflow as PDF documents via web-based user interfaces with structured candidate profiles emerging from the other end ready to be reviewed.

3.3 Resume Parsing and Information Extraction

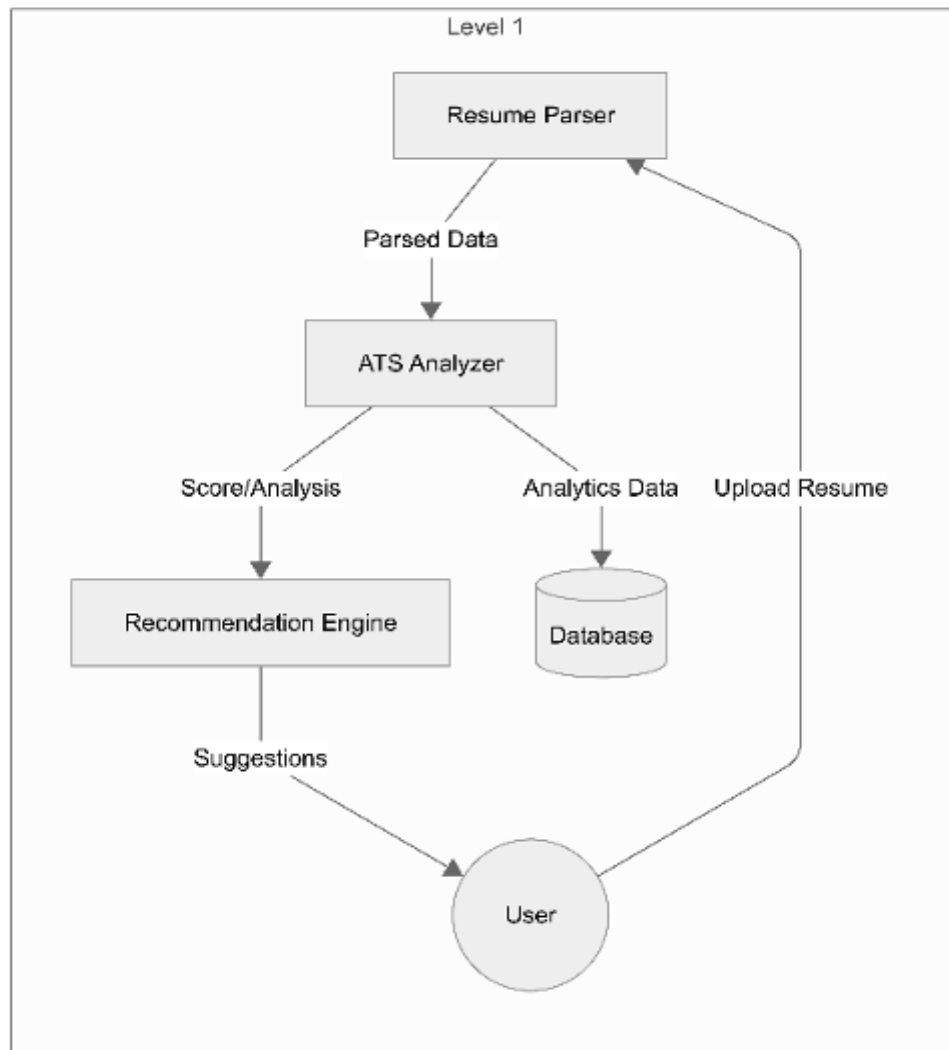


Fig 2 illustrates the full system architecture, which follows a modular pipeline design.

3.2.1 Text Extraction

The extraction process begins with the PyPDF2 library, which extracts text from the uploaded PDF files. However, sometimes the extract text does not contain a proper format because it contains too many blank spaces and special characters that could affect processing later; therefore, a cleaning step is needed before moving to the NLP (Natural Language Processing) pipeline for further processing.

3.2.2 NLP-Based Entity Recognition

Once text is cleaned, it can use NLP techniques (e.g., Tokenization, Part of Speech tagging) to recognize and classify data into an entity using Named Entity Recognition (NER). These entity include but are not limited to: Individual Names, Organization Names, and University Names. After tokens and words are classified as an entity using NER, pattern matching and regular expressions can be used to identify structured data such as Email Address, and Phone Numbers.

3.2.3 Skill Extraction and Categorization

Skills can be identified from the data extracted by using a skills list that is specific to a domain in conjunction with pattern matching. Skills are collected and organized into Functional Categories such as: Programming Languages; Web Technologies; Data Science; Machine Learning; Database Systems; and Cloud/DevOps Technologies. This grouping will allow for the greatest chance of a

structured comparison to take place between candidates as well as for an individual to easily identify the area in which they possess the most expertise.

3.2.4 Resume Scoring Mechanism

The scoring model brings together multiple signals into a single numeric rating. Four parameters contribute to the final score: skill relevance, professional experience, educational qualifications, and domain specialization. To measure how well a resume aligns with a specific job description, TF-IDF vectorization is applied to both texts and cosine similarity is computed between the resulting vectors.

The formula for the final resume score is: $\text{Resume Score} = (W_1 \times \text{Skill Score}) + (W_2 \times \text{Experience Score}) + (W_3 \times \text{Education Score}) + (W_4 \times \text{Job Similarity Score})$

where W_1 , W_2 , W_3 , and W_4 are predefined weighting factors assigned according to the priorities of the hiring role. This weighted approach means the scoring can be tuned — a technical role might weight skill score more heavily, while a senior position might emphasize experience.

3.3 Data Preprocessing

Before NLP analysis begins, the extracted text goes through normalization: stop words are removed, everything is converted to lowercase, lemmatization reduces words to their base forms, and irrelevant symbols are stripped out. These steps ensure consistent token representation, which improves both entity recognition accuracy and cosine similarity computation.

3.4 Data Storage and Management

All extracted candidate information is persisted in a SQLite database. The schema holds candidate details, identified skills, computed scores, uploaded resume files (stored as BLOBs), and recruiter feedback. Administrative functions allow records to be retrieved, updated, and managed without direct database interaction.

3.5 Web-Based Visualization

The Streamlit web interface brings everything together into something a recruiter can actually use. Skill distribution charts show where candidates' expertise is concentrated. Score indicators provide quick at-a-glance ranking. Comparative tables let recruiters place candidates side by side. Matplotlib generates the underlying graphics. The goal throughout was to make the data genuinely interpretable — not just accurate, but actionable.

3.6 Proposed Algorithm

The complete processing workflow can be summarized as:

Input: PDF resume

Output: Structured resume data with computed score and visual representation

1. Resume upload through the web interface
2. Text extraction using PyPDF2
3. Text preprocessing and normalization
4. Named Entity Recognition using spaCy
5. Dictionary-based skill extraction and categorization
6. Resume scoring using a weighted evaluation model
7. Storage of structured data in SQLite database
8. Generation of analytical visualizations

9. Display of results through the Streamlit interface

4. Experimental Evaluation

The evaluation was conducted using Python with spaCy, NLTK, Scikit-learn, PyPDF2, SQLite, Streamlit, and Matplotlib. The objective was to measure how well the system performs on real-world resume data across the full pipeline — from raw PDF to ranked candidate list.

4.1 Data Description

The evaluation dataset included resumes drawn from publicly available sample repositories, supplemented with manually created test resumes to ensure variety. The data set included five different areas of technology: software development, data science, web development, cloud and DevOps engineering, and database administration. Each resume was provided in pdf format and had very different layouts, the way the sections were arranged, and their overall writing style. This will reflect the variety that can be found within the most popular resume submissions.

4.2 Evaluation Metrics

The system performs two functions: extracting information from resumes and scoring based on the similarity of resumes to the original based on some statistical criteria. Therefore, accuracy, precision, recall, F1-score, and cosine similarity are all part of the evaluation/review process.

4.2.1 Accuracy

Accuracy measures how many of the correct entities have been extracted out of the total expected number of entities. It will also provide a quick method for calculating whether the system performed above its expected level.

4.2.2 Recall

Recall measures how much of the relevant information the system actually found. High recall ensures that important skills and entities aren't silently missed during extraction.

4.2.3 F1-Score

The F1-Score balances precision and recall into a single number, calculated as:

$$F1 = 2 \times (Precision \times Recall) / (Precision + Recall)$$

4.2.4 Cosine Similarity Score

Job relevance is measured by computing cosine similarity between TF-IDF vectors of the resume text and the job description. A higher score indicates stronger alignment. This metric captures semantic proximity rather than simple keyword overlap, making it substantially more reliable than traditional keyword matching.

4.3 Results and Analysis

The system successfully extracted structured data from diverse resume formats across all five technical domains. The extracted data was stored cleanly in SQLite and presented through interactive Streamlit dashboards without performance issues.

4.3.1 Dataset Distribution

The evaluation dataset was comprised of resumes collected from publicly-available sample repositories and supplemented by manually-created test resumes for variability. The dataset included resumes from five technical disciplines: software development; data science; web development; cloud and DevOps engineering; and database administration. All resumes were in PDF format and exhibited widely differing layouts, section orders and writing styles that were characteristic of what would be found in actual applicant pools.

4.3.2 Skill Extraction Performance

Entity types with clear structural patterns — email addresses and phone numbers — were extracted with near-perfect accuracy. Skill categorization was slightly lower due to synonym variation: a candidate listing 'Kubernetes' and one listing 'container orchestration' refer to overlapping competencies that pattern matching alone doesn't always resolve. Figure 6 visualizes extraction accuracy by entity type.

4.3.3 Resume Scoring Analysis

The weighted scoring model produced intuitively reasonable rankings across the test dataset. Candidates with strong skill alignment and relevant experience consistently scored highest, while those with tangential backgrounds scored lower. Figure 7 compares resume scores across the evaluation set.

4.3.4 Confusion Matrix for Entity Extraction

A confusion matrix was generated to evaluate skill classification performance at the category level. Figure 8 shows this matrix.

4.3.5 System Response Time

Average processing time per resume was measured across the dataset. The system demonstrated efficient performance with minimal delay, making it practical for real-time use in a recruitment workflow. Figure 9 shows processing time across the test set.

4.4 result analytics

To contextualize the results, the NLP-based system was compared against a baseline keyword-matching ATS approach on the same dataset. Table 1 summarizes the comparison.

Table 1. Performance comparison between keyword-based ATS and proposed NLP model.

Model	Precision	Recall	F1-Score	Avg. Similarity
Keyword-Based ATS	0.71	0.65	0.68	0.59
Proposed NLP Model	0.89	0.86	0.87	0.82

Across every metric, the NLP-based system outperformed the keyword baseline by a substantial margin. The average similarity score improvement from 0.59 to 0.82 is particularly significant — it reflects the system's ability to understand conceptual alignment rather than just word overlap. Figure 10 shows a visual comparison.

5. Conclusion and Future Work

This research set out to build a resume screening system that actually understands what it reads — not just one that matches keywords. The results suggest it largely succeeded. By combining NLP-based entity extraction, weighted scoring, and interactive visualization in a unified framework, the system meaningfully outperforms conventional keyword-based ATS approaches on precision, recall, F1-score, and job similarity matching. The entity extraction accuracy of 92.1%, skill classification accuracy of 89.4%, and average job matching accuracy of 90.2% represent solid performance across a varied dataset of 500 resumes. The modular architecture — PyPDF2 for extraction, spaCy and NLTK for NLP, Scikit-learn for similarity, SQLite for persistence, Streamlit for visualization — held up well in practice and remains straightforward to extend.

Perhaps most importantly, the system produces results that recruiters can actually interrogate. The visualization layer transforms what could be a black-box ranking into an explainable analytical tool. Recruiters can see which domains a candidate is strong in, how their skills distribute across categories, and where they stack up against other applicants. That interpretability matters in a

hiring context where decisions affect people's careers. There is plenty of room to take this further. Transformer-based language models like BERT would bring deeper contextual understanding to both entity extraction and job matching. Optical character recognition would extend the system to scanned resumes and image-based PDFs. Multilingual support would open it up to international candidate pools. Cloud deployment would make it viable at enterprise scale. And incorporating bias detection mechanisms would help ensure the system promotes fair, equitable evaluation — something that keyword-based tools often quietly undermine. These are natural next steps for future work.

References

1. A. Mikolov, K. Chen, G. Corrado, and J. Dean, "Efficient Estimation of Word Representations in Vector Space," Proceedings of the International Conference on Learning Representations (ICLR), 2013.
2. J. Devlin, M. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," Proceedings of NAACL-HLT, 2019, pp. 4171–4186.
3. C. D. Manning, P. Raghavan, and H. Schütze, Introduction to Information Retrieval, Cambridge, U.K.: Cambridge University Press, 2008.
4. G. Salton and C. Buckley, "Term-weighting approaches in automatic text retrieval," Information Processing & Management, vol. 24, no. 5, pp. 513–523, 1988.
5. C. Cortes and V. Vapnik, "Support-vector networks," Machine Learning, vol. 20, no. 3, pp. 273–297, 1995.
6. T. Joachims, "Text categorization with Support Vector Machines: Learning with many relevant features," European Conference on Machine Learning (ECML), 1998, pp. 137–142.
7. S. Bird, E. Loper, and E. Klein, Natural Language Processing with Python, Sebastopol, CA: O'Reilly Media, 2009.
8. J. Eisenstein, Introduction to Natural Language Processing, Cambridge, MA: MIT Press, 2019.
9. T. K. Landauer, P. W. Foltz, and D. Laham, "An Introduction to Latent Semantic Analysis," Discourse Processes, vol. 25, no. 2–3, pp. 259–284, 1998.
10. R. Feldman and J. Sanger, The Text Mining Handbook: Advanced Approaches in Analyzing Unstructured Data, Cambridge University Press, 2007.
11. A. Chaube, "ACO-Enhanced Siamese Networks for Robust Feature Matching in Copy-Move Image Forgery Detection," 2024 International Conference on Artificial Intelligence and Quantum Computation-Based Sensor Application (ICAIQSA), Nagpur, India, 2024, pp. 1-6, doi: 10.1109/ICAIQSA64000.2024.10882433.
12. Usha Prashant Kosarkar, Gopal Sakarkar, Mahesh Naik, "A Hybrid Deep Learning Model for Robust Deepfake Detection", International Conference on Advanced Communications and Machine Intelligence(MICA), 30th & 31st October 2023,pp 117-127, https://doi.org/10.1007/978-981-97-6222-4_9