

BAB I

PENDAHULUAN

1.1. Latar Belakang

Proses rekrutmen merupakan salah satu tahapan krusial bagi perusahaan dalam menyaring dan memperoleh talenta terbaik sesuai dengan kebutuhan organisasi. Di era digital saat ini, perusahaan dihadapkan pada tantangan dalam menyaring ribuan dokumen resume yang masuk untuk setiap lowongan pekerjaan. Alur kerja rekrutmen tradisional sering kali kewalahan, bergantung pada evaluasi manual yang memakan waktu lama, atau bergantung pada sistem otomatis seperti *Applicant Tracking System* (ATS). Namun, sistem ATS konvensional yang mengandalkan pencocokan kata kunci (*keyword matching*) terbukti memunculkan masalah baru, yaitu tingginya tingkat ketidaksesuaian (*mismatch*) antara kandidat yang lolos seleksi dengan deskripsi pekerjaan yang sebenarnya. ATS sering gagal memahami sinonim, konteks pengalaman, dan keahlian tersirat dari pelamar. Hal ini menyebabkan ketidakakuratan yang fatal dalam penilaian; di mana kandidat potensial justru tereliminasi hanya karena perbedaan pemilihan kata (*False Negative*), sementara kandidat yang kurang relevan dapat lolos seleksi awal. Akibatnya, perekrut kesulitan mendapatkan hasil seleksi yang benar-benar akurat dan representatif terhadap kebutuhan perusahaan [1], [2].

Seiring berkembangnya bidang *Natural Language Processing* (NLP), berbagai pendekatan otomatisasi dalam merangkum teks telah muncul. Metode *extractive summarization* memungkinkan sistem memilih kalimat paling relevan dari teks asli, sementara *abstractive summarization* menghasilkan ringkasan baru dengan memahami konteks dan menyusun kalimat baru yang lebih padat [3]. Dengan memanfaatkan model *pre-trained* tersebut, resume dapat diproses secara otomatis, menghasilkan ringkasan yang ringkas namun tetap mempertahankan esensi informasi penting mengenai kandidat. Kemajuan teknologi LLM juga menghadirkan peluang baru untuk menyederhanakan dan mengotomatiskan alur kerja rekrutmen. Dalam analisis dokumen resume, LLM dapat diintegrasikan ke dalam kerangka kerja yang terdiri dari empat komponen utama, yaitu ekstraktor informasi dari resume, evaluator kesesuaian terhadap deskripsi pekerjaan, peringkasan isi dokumen, dan pemformat skor untuk meningkatkan kesesuaian kontekstual penilaian kandidat [4]. Masing-masing komponen berperan dalam mengolah data secara sistematis untuk menghasilkan penilaian yang lebih objektif dan efisien. Meskipun demikian, implementasi LLM dalam sistem analisis dokumen juga menghadirkan tantangan

yang perlu diperhatikan seperti kebutuhan komputasi yang tinggi, keterbatasan interpretabilitas hasil, serta potensi bias dalam data pelatihan yang dapat memengaruhi objektivitas sistem [5]. Selain itu, penelitian yang secara spesifik mengkaji efektivitas penerapan LLM dalam konteks pencocokan resume terhadap deskripsi pekerjaan di pasar kerja Indonesia masih sangat terbatas, padahal karakteristik bahasa dan struktur dokumen pada pasar kerja lokal memiliki perbedaan signifikan dibandingkan konteks global.

Model LLaMA merupakan salah satu keluarga LLM *open source* yang terdiri dari berbagai varian dengan jumlah parameter mulai dari 7B hingga 65B. Model ini diakui sebagai salah satu model bahasa besar (LLM) *open source* paling kuat dan tulang punggung LLM yang populer dari model bahasa multi-modal (MLLM) [6]. Hasil penelitian menunjukkan dengan menyempurnakan LLM mampu memberikan peningkatan performa yang signifikan, dengan skor F1 mencapai 87,73% pada tahap klasifikasi kalimat dalam dokumen resume. Selain itu, pada tahap peringkasan dan evaluasi resume, model yang telah disempurnakan tersebut terbukti melampaui kinerja model dasar GPT-3.5, menunjukkan potensi besar LLM dalam meningkatkan akurasi dan efisiensi proses analisis dokumen rekrutmen [7]. Secara khusus, LLaMA-3.1-8B dipilih dalam penelitian ini karena memiliki *context window* hingga 128 K token, yang memungkinkan model memahami dokumen panjang seperti resume tanpa kehilangan konteks informasi penting [8]. Kemampuan ini memberikan keunggulan signifikan dibandingkan model lain yang memiliki keterbatasan pada panjang konteks. Berdasarkan evaluasi pada *LongBench Benchmark*, LLaMA-3.1-8B-*Instruct* menunjukkan performa unggul dalam tugas-tugas *summarization*, menjadikannya efektif dalam mengekstraksi dan menilai informasi utama dari dokumen teks panjang seperti resume [9]. Selain itu, Llama-3.1-8B memiliki fleksibilitas yang tinggi untuk adaptasi domain melalui metode *Low-Rank Adaptation* (LoRA), yang memungkinkan proses *fine-tuning* dengan data lokal tanpa memerlukan pelatihan ulang berskala besar [10]. Dengan keunggulan-keunggulan tersebut, LLaMA-3.1-8B menjadi pilihan yang sangat efisien, adaptif dan tepat untuk implementasi sistem penyaringan dan analisis resume berbasis LLM.

Berdasarkan permasalahan ketidakefisienan dan keterbatasan akurasi pada sistem penyaringan resume yang ada saat ini, penelitian ini mengusulkan implementasi *Large Language Model* (LLM) untuk membangun sistem analisis kesesuaian dokumen resume terhadap deskripsi pekerjaan. Penelitian ini bertujuan untuk merancang sebuah model yang mampu memberikan skor kesesuaian secara akurat antara resume kandidat dengan kualifikasi yang dibutuhkan.

1.2. Rumusan Masalah

Berdasarkan latar belakang yang telah diuraikan, maka rumusan masalah dalam penelitian ini adalah sebagai berikut:

1. Bagaimana mengatasi masalah ketidakakuratan dan ketidaksesuaian (*mismatch*) dalam proses pencocokan dokumen resume pelamar dengan deskripsi pekerjaan yang selama ini terjadi pada sistem ATS konvensional?
2. Bagaimana merancang dan membangun sistem seleksi kandidat berbasis *Large Language Model* (LLM) yang mampu melakukan penilaian kesesuaian secara kontekstual dan objektif?

1.3. Tujuan

Tujuan dari tugas akhir ini adalah untuk merancang, membangun, dan menguji sebuah sistem berbasis *Large Language Model* (LLM) dengan model LLaMA-3.1-8B yang mampu memberikan analisis kesesuaian dokumen resume dengan deskripsi pekerjaan secara kontekstual, akurat, serta objektif untuk menghasilkan proses identifikasi, kesesuaian dan kualifikasi yang lebih efektif dan efisien.

1.4. Manfaat

Manfaat yang diperoleh dari tugas akhir ini ialah:

1. Membantu meningkatkan objektivitas dalam penilaian kandidat melalui pendekatan analisis berbasis konteks yang meminimalisir ketergantungan pada pencocokan kata kunci semata.
2. Memberikan studi mengenai efektivitas penerapan teknik *prompt engineering* model LLaMA-3.1-8B untuk analisis dokumen rekrutmen dalam konteks bahasa Indonesia dan bahasa Inggris.
3. Menyediakan data evaluasi dan kerangka kerja yang dapat dijadikan dasar bagi peneliti selanjutnya dalam mengembangkan sistem seleksi berbasis *Large Language Model* (LLM) dengan model LLaMA-3.1-8B.

1.5. Ruang Lingkup

Ruang lingkup untuk implementasi dan penelitian sistem analisis kesesuaian dokumen resume terhadap deskripsi pekerjaan ini adalah sebagai berikut:

1. Menggunakan model dasar Meta yaitu LLaMA-3.1-8B sebagai fondasi *Large Language Model*.

2. Metode pengembangan sistem yang digunakan adalah metode *Waterfall* yang bersifat linear dan berurutan (*sequential*), dimulai dari tahap analisis kebutuhan, perancangan, implementasi, pengujian.
3. Penelitian ini memanfaatkan data yang bersumber dari perusahaan Sinar Jasa Powerindo. Pada setiap data telah diberi label "Sesuai" atau "Tidak sesuai" oleh pihak penilai yaitu HR perusahaan Sinar Jasa Powerindo, HR Staff Perusahaan Jaya Baru Pertama dan juga HR Staff Perusahaan Laris Jaya Kendari. Untuk mengukur tingkat kesepakatan antar-penilai, dilakukan proses menghitung skor menggunakan metode *Fleiss' Kappa*. Berdasarkan hasil penilaian tersebut, sistem menetapkan label kesimpulan akhir (*Ground Truth*) melalui metode *majority vote* (voting mayoritas), di mana kandidat akan dikategorikan "sesuai" jika setidaknya dua dari tiga HR memberikan persetujuan. Dataset yang digunakan dalam penelitian ini dapat diakses melalui tautan berikut ini "https://mikroskilacid-my.sharepoint.com/:x/g/personal/221111978_students_mikroskil_ac_id/IQBRW3k8yQ6qSqaWprsubfLJAbUx6aN7zvKAvho77DmM18k?e=vOnxcI"
4. Pembangunan sistem dan proses eksekusi pemodelan dilakukan menggunakan bahasa pemrograman Python versi 3.10, dengan dukungan pustaka spesifik yaitu *PyTorch*, *Transformers*, *BitsAndBytes* dengan kuantisasi 8-bit, serta *PEFT/LoRA*.
5. File *input* untuk dokumen yang dapat diolah oleh sistem adalah dalam format *Portable Document Format* (PDF) dan format *Microsoft Word* (.docx) untuk fleksibilitas pengguna dengan kapasitas maksimal 30 MB.
6. Pengujian sistem dan model akan dilakukan untuk memastikan fungsionalitas dan performa berjalan sesuai dengan tujuan, yang meliputi:
 - a. Pengujian fungsional terhadap sistem dilakukan dengan metode *Black Box Testing* guna memvalidasi keseluruhan alur kerja aplikasi. Pengujian ini secara khusus menilai keandalan proses pengolahan dokumen, mulai dari tahap memasukkan data *input* hingga keluaran hasil analisis oleh model *Large Language Model* (LLaMA).
 - b. Pengujian performa model dilakukan secara kuantitatif menggunakan metrik *Precision*, *Recall*, dan *F1-Score* untuk mengukur akurasi klasifikasi. Agar evaluasi klasifikasi ini dapat dilakukan, *output* akhir dari sistem ditetapkan berupa label ("Sesuai" atau "Tidak Sesuai"), di mana pelabelan ini ditentukan berdasarkan nilai ambang batas (*threshold*) optimal.