

BAB 5. PENUTUP

5.1. Kesimpulan

Berdasarkan hasil penelitian, pengembangan, dan pengujian yang telah dilakukan untuk menjawab rumusan masalah mengenai rancang bangun website deteksi lowongan kerja palsu, dapat ditarik kesimpulan sebagai berikut:

1. Pipeline deteksi menggunakan model Support Vector Machine (SVM) berhasil dikembangkan pada dua jalur. Pada jalur teks (mengintegrasikan OCR, NLP, dan TF-IDF), model mencapai akurasi sebesar 92,75%. Hasil ini terbukti lebih unggul dibandingkan model SVM Visual berbasis handcrafted features yang hanya mencapai akurasi 82,28%. Hal ini menunjukkan bahwa indikator linguistik pada poster penipuan lebih mudah dan efektif untuk diidentifikasi secara algoritmik dibandingkan dengan sekadar mengevaluasi statistik piksel atau warna dasar citra.
2. Model Convolutional Neural Network (CNN) dengan arsitektur VGG16 berhasil dibangun dan mencapai tingkat akurasi yang jauh lebih tinggi, yakni 94,94%, mengungguli kedua model SVM tunggal. Keunggulan ini dicapai karena arsitektur deep learning mampu mengekstrak fitur spasial hierarkis secara otomatis, sehingga model dapat mengenali konteks tata letak visual, inkonsistensi tipografi, dan komposisi desain yang tidak bisa ditangkap oleh ekstraksi fitur statistik manual.
3. Pendekatan Ensemble yang menggabungkan probabilitas dari SVM Teks, SVM Visual, dan CNN sukses memberikan performa paling optimal dengan tingkat akurasi 90,88% dan skor AUC 0.9632. Skor tertinggi ini diraih karena pendekatan ensemble memiliki sifat komplementer; kelemahan klasifikasi visual pada satu model dapat ditutupi dengan sempurna oleh ketajaman analisis teks dari model lainnya.
4. Website Scanix berhasil diimplementasikan secara fungsional. Hasil pengujian backend API menggunakan Postman menunjukkan bahwa seluruh test assertion berhasil dijalankan dengan hasil 176 passed dan 0 failed. Selain itu, pengujian End-to-End Testing yang mencakup 42 skenario (melibatkan alur guest, user terdaftar, dan administrator)

membuktikan bahwa seluruh fitur utama sistem seperti pemindaian, riwayat, laporan anomali, hingga dashboard berhasil berjalan sesuai rancangan pada lingkungan pengujian.

5.2. Saran

Berdasarkan hasil penelitian yang telah diselesaikan, peneliti memberikan sejumlah saran yang diharapkan dapat menjadi acuan bagi peneliti selanjutnya dalam mengembangkan sistem deteksi seperti yang dibuat oleh peneliti, antara lain:

1. Kinerja model ekstraksi fitur visual (*SVM Visual*) masih dapat ditingkatkan. Disarankan untuk mengeksplorasi penambahan metode ekstraksi fitur citra lainnya guna meningkatkan presisi sistem dalam membedakan tekstur gambar asli dan palsu.
2. Untuk menjaga stabilitas server saat terjadi lonjakan permintaan pemindaian dari banyak pengguna secara bersamaan, disarankan agar pemrosesan model kecerdasan buatan dikembangkan lebih lanjut menggunakan sistem antrean tugas secara asinkron.
3. Modul ekstraksi teks (*OCR*) saat ini masih rentan mengalami kesalahan pembacaan pada poster beresolusi rendah atau berdesain rumit. Pengembangan selanjutnya dapat memanfaatkan pustaka *OCR* berbasis *Deep Learning* agar hasil ekstraksi karakter menjadi lebih akurat.
4. Data laporan keluhan yang dikumpulkan dari pengguna dapat diintegrasikan ke dalam sebuah alur pelatihan ulang model secara berkala. Hal ini bertujuan agar sistem dapat terus beradaptasi dengan modus penipuan lowongan kerja yang baru.
5. Mengingat batasan pada analisis kontekstual tautan (*URL*) saat ini, disarankan untuk mengembangkan fitur pelacakan *URL* yang lebih mendalam secara forensik. Pengembangan selanjutnya diharapkan dapat mendeteksi manipulasi tautan yang menyaru sebagai domain resmi (seperti *.com* atau *.co.id*) serta mengevaluasi keabsahan layanan formulir pihak ketiga (seperti *Google Forms* atau *Microsoft Forms*) yang sering disalahgunakan.