

An Integrated Pipeline with Hierarchical Segmentation and CNN for Automated KTP-el Data Extraction on the e-Magang Platform

Nuansa Syafrie Rahardian^{*1}, Eddy Maryanto², Devi Astri Nawangnugraeni³

^{1,2,3}Informatics, Universitas Jenderal Soedirman, Indonesia

Email: ¹nuansa.rahardian@mhs.unsoed.ac.id

Received : Aug 18, 2025; Revised : Sep 1, 2025; Accepted : Sep 2, 2025; Published : Oct 16, 2025

Abstract

In alignment with Indonesia's digital transformation agenda, this research addresses the inefficiencies and error-prone nature of manual data entry on the Foreign Policy Strategy Agency's (BSKLN) e-magang platform. This study introduces a comprehensive, end-to-end Optical Character Recognition (OCR) pipeline, specifically designed for structured identity documents and real-world government platform integration. The proposed methodology features a robust workflow, including image preprocessing with histogram matching, hierarchical segmentation using vertical projection, and intelligent postprocessing to structure the output. To overcome the limitations of a small dataset, three specialized Convolutional Neural Network (CNN) models were rigorously trained and validated using a stratified 5-fold cross-validation technique. The final system was successfully integrated, connecting a Flask-based model engine with the existing Laravel and React platform. End-to-end testing demonstrated strong performance, achieving an average character-reading accuracy of 93.31% with a mean processing time of 14.48 seconds per image. The primary contribution of this research to the field of informatics is the development of a complete and deployable system architecture that ensures data interoperability and reliability, providing a practical blueprint for integrating intelligent automation into digital public services.

Keywords: character segmentation, CNN, flask API, K-Fold Cross Validation, KTP-el, OCR

This work is an open access article and licensed under a Creative Commons Attribution-Non Commercial 4.0 International License



1. PENDAHULUAN

Kartu Tanda Penduduk Elektronik (KTP-el) merupakan dokumen identitas resmi bagi penduduk Indonesia yang diterbitkan oleh instansi berwenang dalam administrasi kependudukan [1] [2]. Di balik fungsinya sebagai dokumen legal, KTP-el menyimpan informasi kunci seperti NIK, nama, alamat, dan tanggal lahir yang menopang berbagai proses administratif lintas sektor serta menegaskan perannya sebagai *source of truth* identitas pada layanan digital [3]. Sejalan dengan itu, hadirnya kerangka hukum Pelindungan Data Pribadi (UU No. 27/2022) memperkuat kebutuhan akan pengelolaan data yang sah, minim risiko, dan aman. Regulasi ini menuntut agar proses verifikasi identitas tidak hanya akurat dan efisien, tetapi juga sesuai prinsip legalitas pemrosesan serta keamanan informasi [4]. Dengan demikian, urgensi standarisasi dan integrasi data muncul sejak tahap awal pelayanan publik, sehingga KTP-el dapat berfungsi optimal bukan sekadar sebagai alat validasi identitas, melainkan juga sebagai pilar konsistensi data lintas sistem [5].

Dalam konteks organisasi, e-magang merupakan *platform* pendaftaran daring untuk program magang reguler di Badan Strategi Kebijakan Luar Negeri (BSKLN) Kementerian Luar Negeri yang dirancang dengan pendekatan *first-come-first-served* demi mengurangi beban administratif dan mempercepat seleksi [6]. Namun, pada implementasi saat ini KTP-el belum menjadi dokumen wajib unggah. Akibatnya, data diri (nama, jenis kelamin, alamat) masih di-*input* manual memakan waktu, rawan salah ketik, dan menimbulkan ketidakseragaman format antarpelamar serta belum selaras dengan agenda transformasi layanan publik yang menekankan integrasi dan standarisasi data. Penambahan

KTP-el sebagai dokumen wajib unggah karenanya menjadi strategis, selain menjadi identitas resmi pendaftar, data dari KTP-el dapat menjadi sumber ekstraksi data untuk mengotomatisasi pengisian data diri di modul isi data diri platform e-magang, tetapi perlu diimbangi validasi otomatis agar berkas yang diunggah benar-benar KTP (bukan dokumen lain) sehingga integritas sistem dan proses seleksi tetap terjaga.

Secara teknis, KTP-el memiliki tata letak visual yang relatif tetap (bidang, tipografi, dan posisi label) yang sangat sesuai untuk diproses dengan *Optical Character Recognition* (OCR), yaitu teknologi untuk mengenali dan mengekstrak teks dari citra atau dokumen digital sehingga dapat diolah dan disimpan dalam bentuk *file* digital [7] [8]. Proses OCR pada KTP-el memiliki tantangan unik [9]. Meskipun formatnya tetap, kualitas gambar yang diunggah pengguna sangat bervariasi, sering kali mengandung *noise*, pencahayaan tidak merata, dan sedikit kemiringan [10]. Selain itu, beberapa karakter memiliki bentuk visual yang sangat mirip (misalnya, 'O' dan '0', 'l' dan 'I'), yang dapat menyebabkan kesalahan klasifikasi jika tidak ditangani dengan baik [11]. Sejumlah penelitian terdahulu telah mengkaji ekstraksi informasi KTP-el dengan dua arus utama. Pertama, M. Haris et al melakukan pendekatan OCR dengan algoritma berbasis *template matching* dan operasi morfologis klasik, yang relatif sederhana dan cepat, namun sensitif terhadap variasi tata letak, orientasi, kualitas pencahayaan, sehingga kinerjanya mudah menurun pada kondisi dunia nyata menghasilkan akurasi akhir 79% [12]. Persentase tersebut menunjukkan bahwa tingkat pendeteksian pada teks di e-KTP masih memerlukan perbaikan agar dapat mendeteksi teks pada KTP-el dengan lebih baik, dan implementasi OCR hanya dilakukan pada aplikasi berbasis desktop (MATLAB) [12]. Kedua, pendekatan OCR berbasis *deep learning* dengan *Convolutional Neural Network* (CNN) dilakukan P. Fatih et al melaporkan sistem deteksi KTP dan OCR dengan kinerja tinggi (akurasinya sekitar 92% dengan presisi 100% dan F1-score 92%) [13]. Sedangkan, Sugiarta et al menunjukkan OCR dengan empat lapis jaringan yang mengekstraksi informasi e-KTP dengan waktu rata-rata 30 detik dan *error rate* $\pm 5\%$, dan menekankan pentingnya kualitas citra [14]. Meski cukup menjanjikan, masih ada ruang pengoptimalan untuk waktu ekstraksi, akurasi, serta integrasi *end-to-end* ke *platform* operasional.

Celah penelitian yang ada terletak pada kurangnya solusi terintegrasi yang menggabungkan seluruh alur kerja. Meskipun penelitian sebelumnya telah berhasil mencapai akurasi yang tinggi, implementasinya sering kali bersifat parsial dan belum menunjukkan sebuah arsitektur *end-to-end* yang terintegrasi penuh ke dalam platform operasional. Menjawab celah penelitian pada solusi OCR yang cenderung parsial, penelitian ini memperkenalkan kebaruan berupa arsitektur *pipeline* OCR holistik yang bekerja secara *end-to-end*. Arsitektur ini dirancang untuk mengatasi tantangan secara menyeluruh melalui tiga tahapan kunci: validasi dokumen di hulu, segmentasi karakter yang tangguh, dan *postprocessing* terstruktur di hilir untuk memastikan data siap diintegrasikan ke sistem operasional. Untuk inti dari *pipeline* ini, yakni ekstraksi karakter, CNN dipilih karena kemampuannya mempelajari fitur visual secara bertingkat sehingga lebih tahan terhadap variasi kondisi citra dibandingkan *template matching* [15] [16]. Selain itu, pendekatan klasifikasi per karakter pada CNN ini menyederhanakan proses pembuatan *dataset*, karena tidak memerlukan anotasi dalam bentuk kata atau kalimat utuh yang lebih sulit diperoleh dalam jumlah besar, tidak seperti model sekuensial (misalnya LSTM atau CRNN) [17] [18]. Model ini bekerja *end-to-end* dan otomatis menemukan fitur penting tanpa *feature engineering* manual, sehingga cocok untuk dokumen berformat tetap seperti KTP-el [19], [20]. Melalui rangkaian operasi konvolusi pada tiap lapisan, CNN menangkap keterkaitan spasial dan pola yang lebih kompleks pada citra, yang pada praktiknya meningkatkan akurasi saat mengenali pola visual, termasuk karakter pada dokumen identitas [21] [22]. Selanjutnya, untuk memastikan model yang dihasilkan *robust* dan dapat diandalkan meskipun dilatih pada *dataset* yang terbatas sekaligus mengukur performanya secara objektif, penelitian ini menerapkan skema evaluasi *stratified k-fold cross-validation*

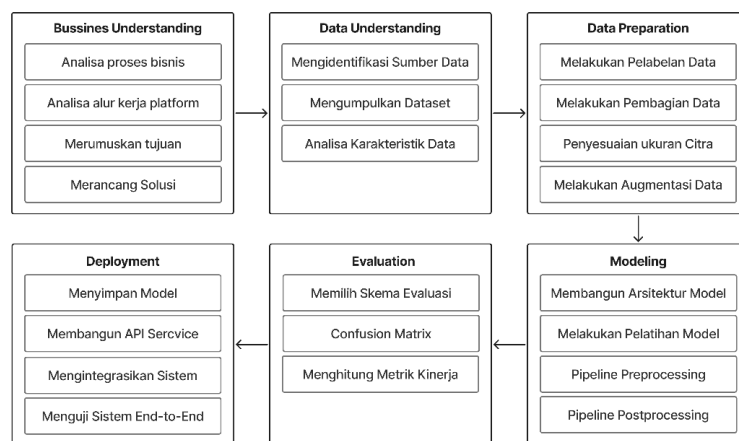
[23][24], sehingga distribusi kelas tetap seimbang di setiap *fold* dan hasil yang diperoleh lebih representatif untuk implementasi di platform e-magang [25].

Berdasarkan latar belakang tersebut, tujuan dari penelitian ini adalah:

1. Merancang dan membangun sebuah *pipeline* OCR *end-to-end* yang mencakup validasi dokumen, *preprocessing*, ekstraksi karakter menggunakan CNN, dan *postprocessing*.
2. Mengintegrasikan sistem OCR yang dikembangkan ke dalam platform e-magang BSKLN untuk mengotomatisasi pengisian formulir data diri.
3. Mengevaluasi performa sistem secara keseluruhan, baik dari segi akurasi maupun efisiensi waktu, dalam skenario penggunaan nyata.

2. METODE PENELITIAN

Penelitian ini mengadopsi kerangka *Cross Industry Standard Process for Data Mining* (CRISP-DM) yang terbukti efektif dalam pengembangan sistem berbasis pembelajaran mesin karena sifatnya yang iteratif dan fleksibel [26]. Kerangka kerja ini memandu penelitian melalui enam tahapan utama yang saling berhubungan, mulai dari *business understanding* hingga *deployment* sistem. Alur dan detail aktivitas pada setiap tahapan penelitian ini dirangkum secara visual pada Gambar 1. [27].



Gambar 1. Tahap Penelitian dengan Kerangka CRISP-DM

2.1. Business Understanding

Business Understanding merupakan tahap awal dan krusial pada kerangka CRISP-DM, yang bertujuan memahami konteks bisnis serta merumuskan tujuan penelitian [27]. Pada penelitian ini, analisis dilakukan terhadap proses bisnis dan alur kerja modul pengisian data diri pada platform e-magang BSKLN Kemlu untuk mengidentifikasi titik integrasi teknologi OCR berbasis *deep learning*. Fokus utamanya adalah merancang metode integrasi yang optimal sehingga sistem dapat melakukan validasi dokumen KTP dan ekstraksi data otomatis, guna meningkatkan efisiensi serta meminimalkan kesalahan *input*.

2.2. Data Understanding

Data Understanding bertujuan untuk memahami karakteristik dan kualitas data yang digunakan dalam pembangunan sistem OCR berbasis CNN untuk ekstraksi data KTP-el. *Dataset* yang digunakan terdiri atas citra KTP-el utuh, citra karakter individu KTP (huruf, angka, dan tanda baca), serta citra non-KTP yang mencakup gambar dokumen selain KTP maupun gambar acak lainnya. Format citra yang diterima adalah PNG, JPG, atau JPEG. Citra KTP-el utuh berperan sebagai data uji alur sistem, sumber pembuatan algoritma *preprocessing*, dan sebagai kelas positif pada model validasi KTP, sedangkan kelas negatif diperoleh dari citra non-KTP. *Dataset* karakter diperoleh dari pemotongan manual setiap

karakter individu pada citra KTP-el, yang kemudian dibagi menjadi dua *subset*, yaitu *dataset* untuk model NIK dan model non-NIK. Pemisahan ini dilakukan karena perbedaan *font* yang signifikan, sehingga diharapkan dapat mengoptimalkan performa model pengenalan karakter [28].

2.3. Data Preparation

Tahap *data preparation* bertujuan menyiapkan *dataset* untuk pelatihan dan evaluasi model melalui langkah *data labeling*, *data splitting*, *resize*, dan *data augmentation* [29]. Setiap citra diberi label sesuai kategori (karakter NIK, karakter non-NIK, KTP-el utuh, dan non-KTP). *Dataset* kemudian dibagi menjadi data pelatihan (80%) dan validasi (20%) menggunakan metode *stratified 5-fold cross-validation*, yaitu teknik validasi silang yang membagi data menjadi lima lipatan (*fold*) dengan tetap menjaga keseimbangan proporsi tiap kelas pada setiap lipatan. Pada tiap iterasi, empat lipatan digunakan untuk melatih model dan satu lipatan sisanya untuk validasi, lalu hasil dievaluasi secara bergantian hingga seluruh data terpakai, sehingga metrik yang diperoleh lebih stabil dan representatif [30]. Pendekatan ini dipilih untuk memaksimalkan pemanfaatan data yang terbatas sekaligus memastikan evaluasi model lebih menyeluruh dan adil terhadap semua kelas [31]. Citra di-*resize* menjadi 64×64 piksel (*grayscale*) untuk model pengenalan karakter dan 224×224 piksel (RGB) untuk model validasi KTP, sebagai keseimbangan optimal antara retensi fitur bentuk dan efisiensi komputasi. *Data augmentation* diterapkan secara *real-time* hanya pada pelatihan, mencakup rotasi, translasi, *zoom*, penyesuaian kontras, dan *Gaussian noise* ringan, sedangkan data validasi tidak mengalami augmentasi untuk menjaga keaslian evaluasi [21]. *Splitting* dan augmentasi dilakukan *real-time* saat *training* di Google Colab.

2.4. Modeling

Tahap *modeling* pada penelitian ini, memiliki 3 tahapan utama, yang pertama adalah pengembangan dan pelatihan model, pengembangan *pipeline preprocessing* dan *postprocessing*. Tiga model CNN yang dikembangkan masing-masing memiliki fungsi spesifik, yaitu model validasi KTP yang berfungsi untuk klasifikasi biner untuk membedakan citra KTP vs non-KTP, model CNN karakter NIK untuk klasifikasi *multi-class* untuk karakter angka pada NIK, dan model CNN karakter non-NIK untuk klasifikasi *multi-class* untuk huruf, angka dan tanda baca. Seluruh model dilatih menggunakan skema *stratified 5-Fold cross-validation* untuk memastikan evaluasi yang *robust* pada *dataset* yang terbatas. Proses pelatihan dioptimalkan menggunakan *optimizer* Adam dengan *learning rate* awal 0.001, yang dipilih karena kemampuannya yang adaptif dan terbukti efektif. *Batch size* sebesar 16 digunakan sebagai keseimbangan antara efisiensi memori komputasi dan stabilitas gradien. Untuk mencegah *overfitting* dan menemukan durasi pelatihan yang optimal, pelatihan dibatasi hingga maksimum 80 *epoch* dan dikontrol oleh *callback EarlyStopping*. Selain itu, *callback ReduceLROnPlateau* diterapkan untuk menyesuaikan *learning rate* secara dinamis guna mencapai konvergensi yang lebih baik. Demi menjamin *reproducibility* atau keterulangan hasil penelitian, *random seed* ditetapkan pada nilai 42.

2.4.1. Arsitektur CNN Validasi KTP

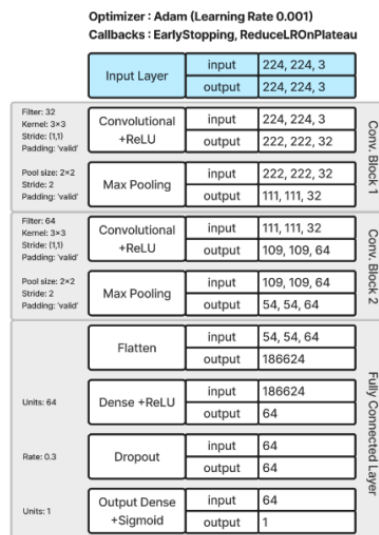
Model validasi KTP-el, yang berfungsi sebagai klasifikasi biner, menerima *input* citra berwarna (RGB) berukuran 224x224 piksel. Arsitektur terdiri dari dua blok konvolusi dengan jumlah filter yang meningkat secara bertahap (32 dan 64 filter), yang memungkinkan model untuk menangkap fitur dari pola sederhana hingga yang lebih kompleks. Perhitungan *output* pada operasi konvolusi ditunjukkan pada Persamaan (2). Setiap blok berisi lapisan Conv2D dengan *kernel* 3x3 dan *padding* "same" untuk menjaga dimensi spasial, diikuti oleh fungsi aktivasi ReLU untuk memperkenalkan non-linearitas, serta lapisan MaxPooling2D untuk mereduksi dimensi dan meningkatkan efisiensi komputasi. Setelah fitur diekstraksi, *feature map* dipipihkan (*Flatten*) dan diteruskan ke lapisan *Dense* dengan 128 neuron untuk

melakukan klasifikasi tingkat tinggi. Sebuah lapisan *Dropout* dengan *rate* 0.30 diterapkan sebagai teknik regularisasi untuk mencegah *overfitting*. Akhirnya, lapisan *output* menggunakan satu neuron dengan *aktivasi Sigmoid* untuk menghasilkan nilai probabilitas biner (0 hingga 1), yang secara matematis sesuai untuk tugas klasifikasi dua kelas. Lapisan *Batch Normalization* tidak diimplementasikan pada arsitektur ini karena sifat model yang lebih sederhana dan untuk menghindari potensi ketidakstabilan statistik yang dapat timbul saat digunakan pada *dataset* dengan ukuran *batch* yang kecil.

$$f(x) = \max(0, x) \quad (1)$$

$$Y[i]_{j,k} = \left(\sum_m \sum_n N_{[j-m],[k-n]} F_{[m,n]} \right) + b \quad (2)$$

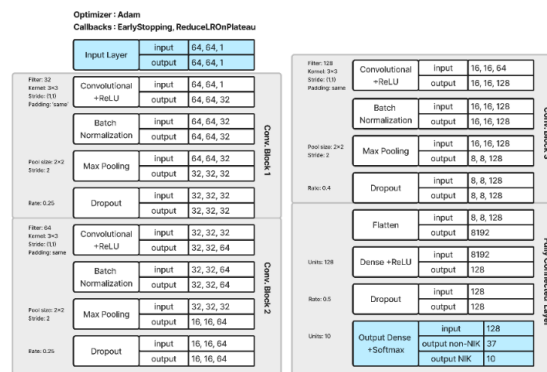
Dalam operasi konvolusi, $Y[i]$ adalah *feature map* (output), N matriks *input*, dan F *kernel* dengan bobot yang dipelajari. BBB berperan sebagai bias, sedangkan indeks j,k menunjukkan posisi piksel *input* dan m,n posisi piksel pada kernel. Arsitektur CNN model validasi KTP ditunjukkan oleh Gambar 2.



Gambar 2. Arsitektur CNN Model Validasi KTP

2.4.2. Arsitektur CNN OCR NIK dan non NIK

Model OCR untuk karakter NIK dan non-NIK dibangun dengan arsitektur CNN yang menerima masukan berupa citra *grayscale* berukuran 64x64 piksel. Kedua model berbagi struktur dasar, dengan perbedaan hanya pada jumlah neuron di lapisan *output* 10 untuk model NIK dan 37 untuk model non-NIK yang diakhiri dengan aktivasi *Softmax* untuk klasifikasi multi-kelas. Arsitektur terdiri dari tiga blok konvolusi bertingkat dengan filter (32, 64, 128) untuk ekstraksi fitur hierarkis. Setiap blok berisi lapisan Conv2D, aktivasi ReLU, dan lapisan *Batch Normalization* untuk menstabilkan serta mempercepat pelatihan. Setelah itu, diterapkan MaxPooling2D untuk mereduksi dimensi, diikuti oleh lapisan *Dropout*. Strategi regularisasi ini diimplementasikan secara bertingkat: *rate* 0.25 diterapkan setelah dua blok konvolusi pertama, meningkat menjadi 0.4 setelah blok ketiga, dan mencapai 0.5 setelah lapisan *Dense* dengan 128 neuron. Pendekatan ini secara efektif mencegah *overfitting* pada lapisan dengan parameter terbanyak sambil tetap mempertahankan fitur fundamental di lapisan awal, sehingga meningkatkan kemampuan generalisasi model secara keseluruhan. Arsitektur model ditunjukkan oleh Gambar 3.



Gambar 3. Arsitektur CNN Model OCR NIK dan non-NIK

2.4.3. Pipeline Preprocessing

Setiap gambar yang diunggah ke *platform* e-magang akan melalui sebuah *pipeline preprocessing* otomatis yang dirancang untuk mengisolasi setiap karakter secara individual sebelum diklasifikasikan oleh model CNN. Langkah pertama dalam *pipeline* adalah validasi citra. Citra KTP-el yang menjadi inputan harus utuh, tidak memiliki latar belakang, tidak miring dan tidak terpotong. Gambar *input* diubah ukurannya menjadi 224x224 piksel dan dimasukkan ke dalam model validasi KTP-el. Jika citra terklasifikasi sebagai KTP yang valid, proses dilanjutkan; jika tidak, proses akan dihentikan. Citra yang lolos validasi kemudian di standardisasi dengan mengubah ukurannya menjadi 1720x906 piksel, sebuah resolusi yang ditemukan secara empiris optimal untuk menjaga kejelasan teks sebelum segmentasi. Selanjutnya dilakukan penerapan *histogram matching* untuk menyeragamkan pencahayaan, dan pemotongan area teks relevan dengan menghilangkan 23% bagian kiri, 24% bagian kanan, dan 8% bagian bawah citra. Persentase ini ditetapkan untuk secara konsisten membuang area non-teks seperti pas foto, tanda tangan, dan margin kosong.

Selanjutnya, dilakukan segmentasi baris secara hierarkis. Citra yang telah distandardisasi diproses melalui serangkaian filter, termasuk *gaussian blur* (kernel 7x7) dan peningkatan kontras ($\alpha=1.6$, $\beta=55$), sebelum dikonversi menjadi citra biner menggunakan metode *otsu's thresholding*. Teknik kunci pada tahap ini adalah operasi morfologi *closing* dengan *kernel* horizontal berukuran 32x1 piksel yang berfungsi menyatukan semua karakter dalam satu baris menjadi blok solid. Bentuk *kernel* yang memanjang secara horizontal ini dipilih secara spesifik untuk menyatukan karakter-karakter dalam satu baris menjadi blok solid tanpa secara tidak sengaja menggabungkannya dengan baris di atas atau di bawahnya. Kontur dari setiap blok baris ini kemudian dideteksi untuk diekstraksi [8].

Setiap baris yang telah terisolasi kemudian melalui tahap segmentasi kata. Proses ini menggunakan metode proyeksi vertikal untuk mendeteksi spasi dan memisahkan setiap blok kata. Pada tahap ini, parameter peningkatan kontras disesuaikan secara adaptif tergantung pada baris yang diproses (misalnya, $\alpha=1.8$ untuk baris pertama KTP). Terakhir, segmentasi karakter dilakukan pada setiap kata. Setelah melalui serangkaian filter yang lebih agresif untuk penajaman, metode proyeksi vertikal kembali digunakan untuk memisahkan karakter individual, perhitungan pada proyeksi vertikal seperti yang ditunjukkan pada Persamaan (3).

$$P_v[i] = \sum_{j=1}^M S[i, j] \quad (3)$$

Sebuah logika berbasis lebar *bounding box* juga diimplementasikan untuk menangani karakter yang menyatu dengan memotongnya secara otomatis. Sebagai tahap final, setiap citra karakter yang berhasil diisolasi dari *pipeline* ini dinormalisasi ukurannya menjadi 64x64 piksel dan dikonversi ke format *grayscale*, agar seragam dan sesuai dengan format *input* yang dibutuhkan oleh model-model

OCR. Ukuran ini merupakan standar umum dalam tugas pengenalan karakter yang memberikan keseimbangan antara retensi detail fitur visual dan efisiensi komputasi saat pelatihan model OCR.

2.4.4. Pipeline Postprocessing

Setelah model CNN menghasilkan prediksi teks mentah, *pipeline postprocessing* dijalankan untuk mengoreksi kesalahan umum OCR dan menyusun data ke format JSON terstruktur [32]. Proses ini dimulai dengan menentukan struktur baris (10–12 baris), lalu melakukan koreksi berbasis konteks: pada kolom teks digit dikonversi ke huruf ('0' menjadi 'O', '5' menjadi 'S'), sedangkan pada kolom numerik huruf dikonversi ke angka ('I' menjadi '1'). Data kompleks juga ditangani, misalnya tempat dan tanggal lahir dipisah lalu distandardisasi ke format *dd-mm-yyyy*, jenis kelamin dinormalisasi menjadi "LAKI-LAKI" atau "PEREMPUAN", serta komponen alamat (jalan, RT/RW, kelurahan, kecamatan) digabungkan kembali. Hasil akhirnya berupa JSON bersih dan konsisten yang siap diintegrasikan otomatis ke formulir e-magang.

2.5. Evaluation

Evaluasi kinerja model CNN dilakukan dengan skema *stratified 5-fold cross-validation*. Pada tiap *fold*, model yang telah selesai dilatih (*pretrained fold model*) dimuat kembali tanpa proses *retraining*, sehingga evaluasi dilakukan murni menggunakan parameter akhir hasil pelatihan [33]. Setiap *fold* dievaluasi baik pada data latih maupun data validasi untuk membandingkan performa antara proses pembelajaran dan kemampuan generalisasi.

Prediksi model dikonversi ke label diskrit menggunakan ambang batas 0,5, kemudian dibandingkan dengan label sebenarnya untuk menghasilkan *confusion matrix* pada masing-masing *fold* [34]. Selanjutnya, *confusion matrix* dari seluruh *fold* dijumlahkan secara agregat sehingga diperoleh satu *confusion matrix* akhir yang merepresentasikan keseluruhan performa model pada skema *cross-validation*. Dari *confusion matrix* agregat ini kemudian dihitung empat metrik utama, yaitu *accuracy*, *precision*, *recall*, dan *F1-score* yang ditunjukkan oleh persamaan (4)-(7). Perhitungan metrik untuk *multi-class* dilakukan pada dua skema, yakni *macro average* untuk menilai performa rata-rata antar kelas secara seimbang, dan *weighted average* untuk memperhitungkan perbedaan jumlah sampel pada tiap kelas [24].

$$Prec_k = \frac{TP_k}{TP_k + FP_k}, Rec_k = \frac{TP_k}{TP_k + FN_k}, F1_k = \frac{2 Prec_k Rec_k}{Prec_k + Rec_k} \quad (4)$$

$$Prec_{macro} = \frac{1}{K} \sum_{k=1}^K Prec_k, Rec_{macro} = \frac{1}{K} \sum_{k=1}^K Rec_k, F1_{macro} = \frac{1}{K} \sum_{k=1}^K F1_k \quad (5)$$

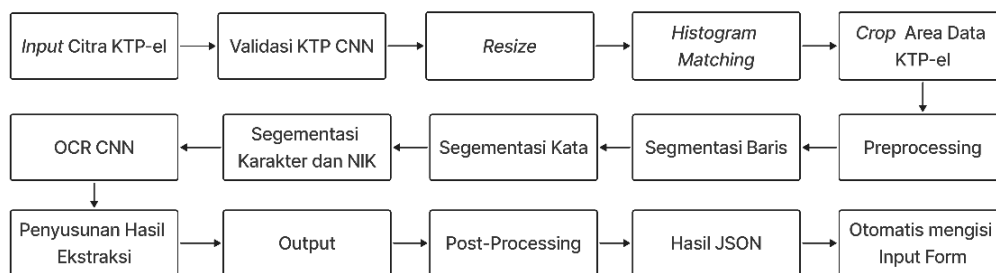
$$F1_{weighted} = \sum_{k=1}^K \left(\frac{n_k}{N} \right) F1_k \quad (6)$$

$$Accuracy = \frac{\sum_{k=1}^K C_{k,k}}{N} \quad (7)$$

2.6. Deployment

Tahap akhir penelitian adalah *deployment* model OCR berbasis CNN ke dalam platform e-magang. Model final yang diperoleh melalui *retraining* dari lima *fold cross-validation* berhasil disimpan dalam format .h5 dan dihubungkan dengan *backend Laravel* melalui *RESTful* API berbasis Flask yang dijalankan secara lokal. API ini menangani proses ekstraksi data KTP-el, mulai dari validasi citra,

preprocessing, segmentasi karakter, prediksi, hingga *postprocessing* dengan hasil dikirim kembali dalam format JSON. Pendekatan ini dipilih karena model hanya dimuat sekali saat server Flask dijalankan, sehingga lebih efisien dibanding memuat ulang pada setiap permintaan. Pada sisi *frontend*, antarmuka ReactJS memungkinkan pengguna mengunggah citra KTP-el, yang kemudian diproses otomatis dan hasilnya langsung mengisi formulir digital pada modul pengisian data diri, alur ekstraksi yang terjadi pada sistem secara *end-to-end*, ditampilkan pada diagram blok sistem yang ditunjukkan pada Gambar 4. Setelah integrasi, sistem dipantau untuk mengevaluasi akurasi serta waktu respons. Parameter evaluasi adalah akurasi penempatan label, akurasi pembacaan karakter dan waktu ekstraksi. Akurasi penempatan label dan akurasi pembacaan karakter dihitung dengan persamaan (8) dan (9).



Gambar 4. Diagram Blok Alur Sistem

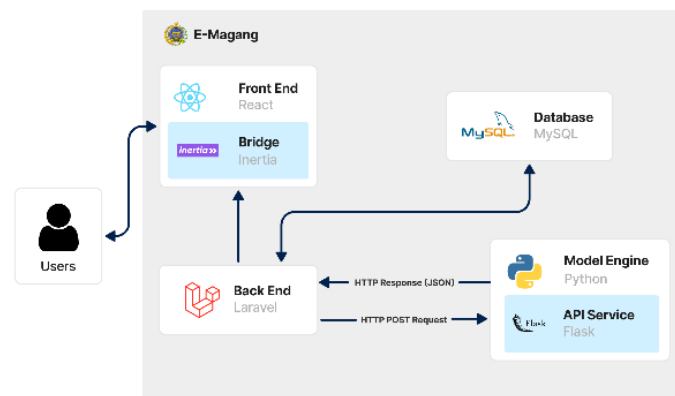
$$\text{Akurasi Penempatan Label}(\%) = \left(\frac{\text{Jumlah Label Data Benar}}{\text{Jumlah Kolom Input Form}} \right) \times 100 \quad (8)$$

$$\text{Akurasi Pembacaan Karakter}(\%) = \left(\frac{\text{Jumlah Karakter Benar}}{\text{Jumlah Ground Truth}} \right) \times 100 \quad (9)$$

3. HASIL DAN PEMBAHASAN

3.1. Bussines Understanding

Tujuan penelitian ini adalah mengembangkan sistem otomatisasi berbasis OCR untuk validasi dokumen KTP sekaligus mengekstraksi data utama yang dibutuhkan platform e-magang, yaitu NIK, nama, jenis kelamin, alamat (provinsi, kabupaten/kota, kecamatan, kelurahan, RT/RW, jalan), serta tempat dan tanggal lahir. Sistem ini diharapkan dapat meningkatkan efisiensi pengisian formulir digital, meminimalkan kesalahan *input*, dan mendukung transformasi layanan publik yang lebih praktis dan akurat.



Gambar 5. Arsitektur Sistem Platform e-magang

Platform e-magang yang dibangun dengan Laravel, Inertia.js, dan ReactJS dihubungkan dengan layanan Flask API yang menjalankan model CNN seperti yang ditunjukkan pada Gambar 4. Alur bisnis pengisian data diri dimodifikasi agar unggahan citra KTP dapat secara otomatis divalidasi, diproses, dan hasil ekstraksinya ditampilkan pada *form*, dengan opsi koreksi oleh pengguna sebelum penyimpanan. Integrasi ini memungkinkan proses pengisian data lebih cepat, akurat, dan sesuai kebutuhan digitalisasi layanan publik.

3.2. Data Understanding

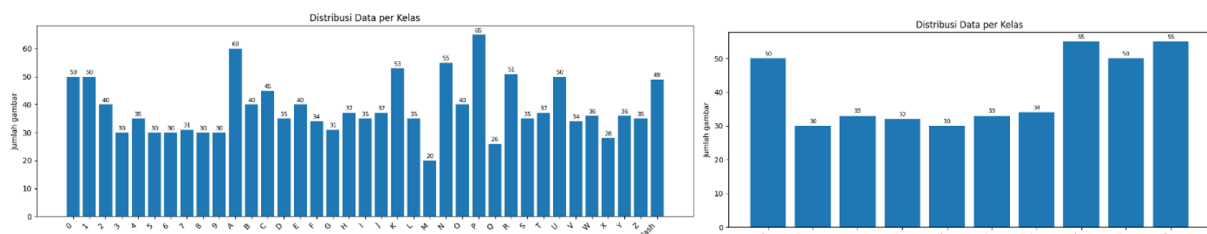
Data citra KTP-el yang didapat pada penelitian ini bersumber dari situs Roboflow dan melalui pengumpulan manual pada sukarelawan [35]. Data yang dikumpulkan sejumlah 100 data citra KTP-el dengan format yang sesuai, dan 100 citra non-KTP berupa gambar dokumen seperti SIM, KTM, kartu ATM, dan beberapa dokumen lain selain KTP. Dari total 100 citra KTP yang dikumpulkan, sebanyak 80 citra akan dipotong secara manual pada setiap karakter individunya untuk membentuk *dataset* pelatihan. Sementara itu, 20 citra lainnya dari KTP maupun non-KTP disisihkan sebagai data uji (*unseen data*) yang akan digunakan untuk menguji performa sistem setelah model selesai dilatih dan di-*deploy*.

3.3. Data Preparation

Pada tahap ini dimulai dengan data *labeling*, pada setiap *dataset* untuk 3 model yang akan dikembangkan. Untuk *dataset* model validasi KTP terdiri dari 2 kelas dengan label KTP dan non-KTP seperti yang ditunjukkan pada Gambar 5, sementara itu *dataset* model OCR NIK terdiri dari 10 kelas yang berisi citra berupa angka dari 0-9, *dataset* model OCR NIK terdiri dari 37 kelas yang berisi citra angka (0-9), huruf (A-Z) dan tanda baca slash (/). Distribusi data tiap kelas ditunjukkan oleh Gambar 6. Hasil dari proses *splitting* menggunakan *stratified 5-fold cross-validation* adalah tiap *fold* memiliki porsi data pelatihan sekitar 80% dan validasi 20% dengan augmentasi diterapkan hanya pada data pelatihan.



Gambar 6. Dataset KTP dan non-KTP



Gambar 7. Distribusi data tiap kelas untuk *dataset* OCR

3.4. Modeling

Tahap ini menyajikan hasil pelatihan tiga model CNN yang dikembangkan. Evaluasi dilakukan secara menyeluruh dengan distribusi kelas yang seimbang pada tiap *fold*. Tabel 1 menampilkan akurasi hasil pelatihan pada setiap *fold* untuk menggambarkan konsistensi performa model.

Tabel 1. Akurasi *Train* dan *Validation* tiap *Fold*

Model	Fold-1		Fold-2		Fold-3		Fold-4		Fold-5	
	<i>Train</i> (%)	<i>Val</i> (%)	<i>Train</i> (%)	<i>Val</i> (%)	<i>Train</i> (%)	<i>Val</i> (%)	<i>Train</i> (%)	<i>Val</i> (%)	<i>Train</i> (%)	<i>Val</i> (%)
Validasi KTP	100	100	82.8	86.67	100	100	92	93.7	100	100
NIK	99.4	100	98.1	100	98.44	98.78	98.14	100	97.8	97.4
non-NIK	98.95	96.6	99.7	99.3	100	97.6	99.2	97.9	98.96	97.5

Hasil evaluasi menunjukkan bahwa seluruh model CNN mencapai akurasi tinggi pada skema *stratified 5-fold cross-validation*. Model validasi KTP memperoleh akurasi validasi berkisar antara 86.67% hingga 100%, dengan beberapa *fold* mencapai performa sempurna. Rentang akurasi validasi yang konsisten tinggi di kelima *fold* bahkan mencapai 100% pada beberapa iterasi untuk model validasi KTP mengindikasikan bahwa arsitektur yang diusulkan bersifat stabil dan mampu melakukan generalisasi dengan baik pada partisi data yang berbeda. Perbedaan kecil antara akurasi pelatihan dan validasi menunjukkan bahwa ketiga model mampu melakukan generalisasi dengan baik tanpa indikasi *overfitting* yang signifikan.

3.4.1. Pipeline Preprocessing

Pipeline preprocessing yang dirancang pada penelitian ini berhasil menormalkan citra KTP sehingga lebih seragam untuk tahap pengenalan karakter. Pada tahap validasi awal, model validasi KTP mampu menyaring dokumen non-KTP dengan akurasi tinggi, sehingga hanya citra KTP yang relevan yang diteruskan ke tahap berikutnya. Proses standardisasi ukuran (1720×906 piksel) dan *histogram matching* terbukti efektif dalam mengurangi variasi pencahayaan maupun kualitas kamera, menghasilkan citra dengan kontras yang lebih konsisten, dengan hasil seperti yang dapat dilihat pada Gambar 7.

Gambar 8. Proses *Histogram Matching*

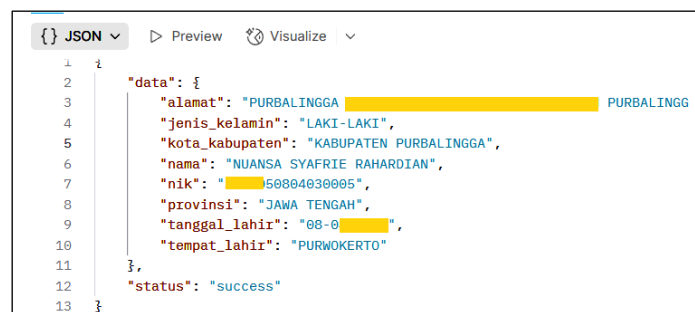
Tahap segmentasi hierarkis juga menunjukkan hasil yang memadai. Segmentasi baris berhasil memisahkan area teks sesuai dengan struktur dokumen, diikuti segmentasi kata yang dapat mendeteksi spasi dengan baik meskipun terdapat variasi jarak antarhuruf seperti yang ditunjukkan pada Gambar 8. Selanjutnya, segmentasi karakter individual menghasilkan citra huruf yang sesuai. Visualisasi hasil seperti yang ditunjukkan pada Gambar menunjukkan bahwa mayoritas karakter dapat diisolasi dengan baik dan sesuai format masukan model OCR.



Gambar 9. Hasil Segmentasi Karakter

3.4.2. Pipeline Postprocessing KTP

Pipeline postprocessing yang diusulkan terbukti efektif dalam meningkatkan kualitas hasil OCR mentah dari model CNN. *Pipeline* mampu menyusun keluaran teks ke format JSON yang rapi dan konsisten seperti yang ditunjukkan pada Gambar 9. sehingga dapat langsung digunakan modul integrasi e-magang. Kontribusi utamanya adalah kemampuan menangani variasi struktur baris, misalnya nama atau alamat yang terdeteksi lebih dari satu baris otomatis digabung tanpa kehilangan konteks. Aturan koreksi berbasis konteks juga berhasil mengurangi kesalahan karakter mirip, seperti digit '0' dan huruf 'O'. Selain itu, *pipeline* menormalisasi data kompleks: tempat dan tanggal lahir dikonversi ke format *dd-mm-yyyy*, jenis kelamin dinormalisasi menjadi kategori baku, dan komponen alamat digabungkan kembali menjadi *string* utuh. Secara keseluruhan, pasca-pemrosesan ini tidak hanya meningkatkan akurasi representasi teks, tetapi juga memastikan data memenuhi standar interoperabilitas dengan format JSON siap integrasi, sehingga meminimalkan koreksi manual.

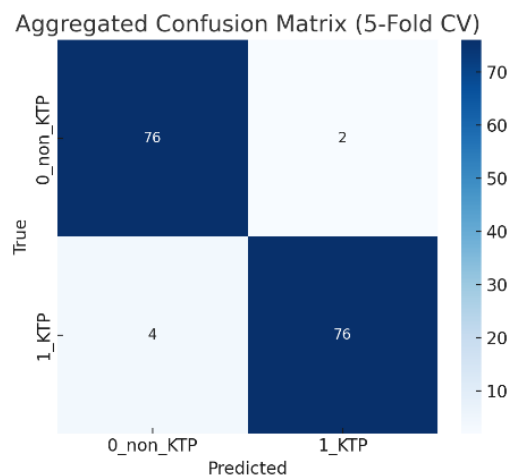


Gambar 10. Hasil JSON Ekstraksi KTP-el

3.5. Evaluation

3.5.1. Model Validasi KTP

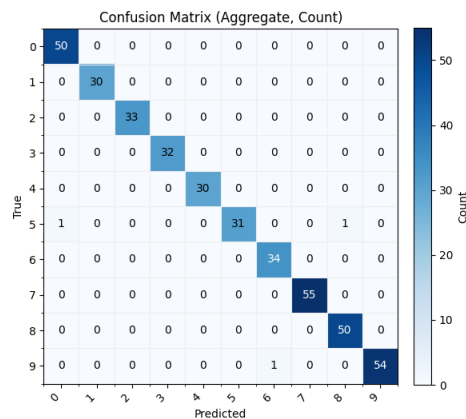
Model klasifikasi biner yang dirancang untuk membedakan citra KTP vs non-KTP menunjukkan performa yang sangat baik. Berdasarkan *confusion matrix* agregat yang ditunjukkan pada Gambar 10, model mencapai *accuracy* sebesar 0.9620, *precision* 0.9744, *recall* 0.9500, dan *F1-score* 0.9620. Keseimbangan metrik ini, ditambah dengan hanya 6 kesalahan klasifikasi dari total 160 data, membuktikan bahwa model sangat andal dalam meminimalkan *false positive* maupun *false negative*, sehingga efektif digunakan sebagai gerbang penyaringan awal pada *pipeline*.



Gambar 11. *Confusion Matrix Agregat* Model Validasi KTP

3.5.2. Model OCR NIK

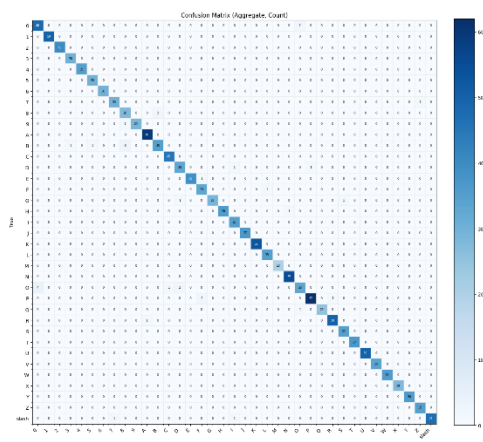
Model klasifikasi *multi-class* untuk karakter numerik pada NIK (0–9) menunjukkan hasil yang hampir sempurna. Ringkasan evaluasi menunjukkan *accuracy* 99.3%, *macro-precision* 0.993, *macro-recall* 0.992, serta *macro* dan *weighted F1-score* 0.992. *Confusion matrix* pada Gambar 11 memperlihatkan distribusi prediksi yang sangat konsisten, dengan kesalahan klasifikasi yang sangat minim antar digit. Analisis kesalahan pada *confusion matrix* menunjukkan bahwa error minor yang terjadi umumnya disebabkan oleh kemiripan bentuk visual antar digit, seperti antara '5' dan '8' serta '9' dan '6'.



Gambar 12. *Confusion Matrix Agregat Model OCR NIK*

3.5.3. Model OCR non-NIK

Model OCR untuk karakter huruf (A–Z), angka (0–9), dan tanda baca (/) juga menunjukkan performa yang kuat pada skema evaluasi. Berdasarkan *confusion matrix* agregat pada Gambar 12, model memperoleh *accuracy* 0.9777, *macro-precision* 0.9775, *macro-recall* 0.9778, *macro F1-score* 0.9774, serta *weighted F1-score* 0.9776. Hasil ini menegaskan kemampuan model dalam menjaga konsistensi prediksi meskipun dihadapkan pada tugas klasifikasi yang jauh lebih kompleks. Analisis kesalahan pada *confusion matrix* mengungkapkan bahwa error yang terjadi bersifat minor dan sistematis, umumnya disebabkan oleh ambiguitas visual antar karakter seperti '0' dan 'O' atau '1' dan 'l', terutama pada citra *input* berkualitas rendah. Temuan ini secara langsung memperkuat pentingnya peran *pipeline postprocessing*, yang dirancang khusus untuk melakukan koreksi berbasis konteks guna memitigasi pola kesalahan tersebut.



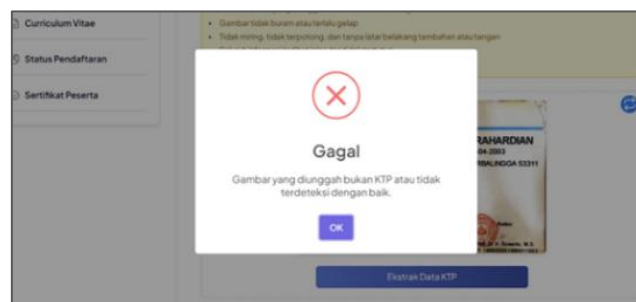
Gambar 13. *Confusion Matrix Agregat Model OCR non-NIK*

3.6. Deployment

Pengujian *end-to-end* dilakukan untuk mengevaluasi kinerja sistem setelah seluruh komponen diintegrasikan pada platform e-magang. Antarmuka pengguna pada modul pengisian data diri telah diperbarui untuk mendukung fitur unggah dan ekstraksi otomatis, seperti yang disajikan pada Gambar 13.

Gambar 14. Antarmuka pengguna *platform* e-magang

Tahap pengujian pertama adalah memverifikasi fungsi validasi dokumen. Sistem diuji dengan mengunggah 20 jenis gambar non-KTP, seperti SIM, kartu pelajar, dan KTM. Hasilnya, model validasi berhasil menolak seluruh gambar non-KTP tersebut dan menghentikan proses ekstraksi. Pengguna kemudian diberikan notifikasi kesalahan yang informatif, seperti yang divisualisasikan pada Gambar 14, yang membuktikan bahwa sistem dapat menyaring *input* yang tidak relevan secara efektif.



Gambar 15. Notifikasi Gambar tidak Valid

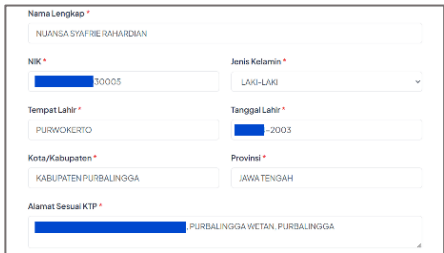
Setelah lolos validasi, pengujian dilanjutkan dengan citra KTP-el yang valid. Sistem mampu menjalankan seluruh *pipeline* dan secara otomatis mengisi kolom-kolom formulir. Contoh visual dari hasil ekstraksi yang berhasil disajikan pada Tabel 2, yang menampilkan tiga sampel citra KTP-el yang merupakan (*unseen data*) dan formulir data diri pada *platform* yang telah terisi otomatis. Secara kuantitatif, pengujian pada 20 sampel menunjukkan performa sistem yang cukup konsisten baik dari aspek akurasi maupun efisiensi seperti yang ditunjukkan pada Tabel 3. Rata-rata akurasi penempatan label mencapai 95.63%, dengan 13 dari 20 sampel (65%) berhasil memperoleh nilai sempurna (100%).

Adapun 7 sampel lainnya hanya mencapai 87.5%, terutama disebabkan oleh kesalahan pada kolom tanggal lahir. Kesalahan ini dipicu oleh segmentasi yang kurang tepat sehingga format hasil ekstraksi tidak sesuai standar *dd-mm-yyyy* yang menyebabkan data tidak lolos validasi *Date Picker*. Meskipun demikian, logika *postprocessing* terbukti mampu mengatasi variasi jumlah baris teks (10–12 baris), seperti pada KTP ke-1 yang memiliki nama dua baris atau KTP ke-11 dengan alamat hingga 12 baris, sehingga struktur data tetap tersusun rapi ke dalam format JSON.

Dari sisi akurasi pembacaan karakter, sistem mencapai rata-rata 93.31%, dengan performa terbaik (100%) ditunjukkan pada KTP ke-1 dan ke-10. Sebaliknya, akurasi terendah terjadi pada KTP ke-12 (80.48%) dan KTP ke-19 (80.8%) akibat kualitas citra yang buruk, misalnya karena adanya *noise*, *blur*, atau tinta yang *bleeding* sehingga mengganggu pemisahan dan pengenalan karakter.

Selain itu, sistem juga menunjukkan efisiensi waktu yang cukup baik dengan rata-rata waktu ekstraksi 14.48 detik, dihitung sejak tombol ekstraksi ditekan hingga seluruh hasil ditampilkan di antarmuka pengguna. Waktu ini menandai peningkatan signifikan dibandingkan proses entri data manual, sehingga mendukung tujuan utama penelitian untuk mempercepat dan mempermudah validasi serta ekstraksi data KTP secara otomatis.

Tabel 2. Hasil Pengujian Ekstraksi data KTP sampel pada platform e-magang

KTP ke-	Data Uji	Hasil
1		
2		
3		

Tabel 3. Hasil Uji Akurasi dan Waktu Ekstraksi pada Platform e-mangang

KTP ke-	Akurasi Penempatan Label(%)	Akurasi Pembacaan Karakter (%)	Waktu Ekstraksi (s)
1	100	100	16.03
2	100	99.186	13.98
3	100	97.71	15.35
4	100	88.46	14.50
5	100	96	14.33
6	100	98.34	14.73
7	100	95.65	13.23
8	100	95.45	14.95
9	100	99.15	13.37
10	100	100	14.14
11	100	96.57	16.70
12	87.5	80.48	13.56
13	87.5	89.42	14.92
14	100	95.72	13.40
15	87.5	87.59	14.44
16	87.5	91.36	14.72
17	100	96.24	14.27
18	87.5	90.90	15.64
19	87.5	80.8	13.38
20	87.5	87.2	14.01
Rata-rata	95.625	93.31	14.48

4. HASIL DISKUSI

Kinerja sistem yang diusulkan, dengan akurasi pembacaan karakter *end-to-end* sebesar 93.31%, menunjukkan keunggulan yang jelas dibandingkan pendekatan klasik berbasis *template matching* yang dilaporkan oleh M. Haris et al., yang hanya mencapai akurasi 79%. Performa ini sebanding dengan penelitian berbasis CNN lainnya oleh P. Fatih et al. (akurasinya sekitar 92%). Namun, *novelty* utama dari penelitian ini tidak hanya terletak pada akurasi, melainkan pada pengembangan arsitektur *pipeline* yang terintegrasi penuh, mulai dari validasi dokumen di hulu hingga penyajian data JSON yang siap pakai, sebuah aspek yang tidak dibahas secara mendalam pada penelitian-penelitian tersebut.

Secara eksplisit, kontribusi penelitian ini bagi bidang informatika adalah perancangan arsitektur sistem yang memastikan interoperabilitas data antara sistem pengenalan pola dan sistem informasi. Dengan mengubah data visual yang tidak terstruktur menjadi format JSON yang terstandarisasi, penelitian ini menyajikan sebuah cetak biru praktis untuk integrasi sistem cerdas ke dalam layanan publik digital, di mana konsistensi dan integritas data menjadi fondasi yang krusial.

Meskipun demikian, sistem ini memiliki keterbatasan yang signifikan, yaitu sensitivitasnya terhadap kualitas citra *input*. Gambar yang memiliki kemiringan, *noise*, atau *blur* terbukti dapat menurunkan akurasi, khususnya pada tahap segmentasi. Untuk mengatasi ini, pengembangan di masa depan harus mencakup implementasi modul *preprocessing* yang lebih adaptif, seperti deteksi dan koreksi kemiringan otomatis (*automatic skew detection and correction*) serta pemotongan gambar otomatis (*automatic cropping*) untuk membuang area latar belakang yang tidak relevan.

Pendekatan klasifikasi per karakter menggunakan CNN dipilih karena menyederhanakan proses pembuatan *dataset*. Namun, perlu diakui bahwa model sekuensial seperti CNN-LSTM atau CRNN

sering kali menawarkan keunggulan dalam memahami konteks antar karakter. Sebagai perbandingan, studi yang menggunakan CRNN untuk pengenalan teks pada plat nomor sering kali melaporkan akurasi yang lebih tinggi untuk rangkaian karakter yang panjang. Oleh karena itu, eksplorasi model sekuensial tersebut dapat menjadi langkah logis berikutnya untuk meningkatkan akurasi kontekstual pada pembacaan nama dan alamat.

5. KESIMPULAN

Studi ini telah berhasil merancang, membangun, dan mengintegrasikan sebuah *pipeline* OCR end-to-end yang fungsional sesuai dengan tujuan penelitian, yang terbukti mampu mengekstraksi data identitas dari citra KTP-el dengan akurasi pembacaan karakter rata-rata 93.31% dalam skenario penggunaan nyata. Kontribusi ilmiah utama dari penelitian ini adalah pengembangan arsitektur sistem terintegrasi yang memastikan interoperabilitas data, menyajikan model praktis tentang bagaimana teknologi *deep learning* dapat diterapkan untuk meningkatkan efisiensi dan integritas data dalam sistem informasi pemerintahan. Meskipun demikian, penelitian di masa depan disarankan untuk fokus pada tiga area utama: peningkatan *robustness* pada tahap *preprocessing* melalui deteksi dan koreksi kemiringan otomatis, eksplorasi model sekuensial seperti CNN-LSTM atau CRNN untuk meningkatkan akurasi kontekstual, serta perluasan keragaman *dataset* untuk meningkatkan kemampuan generalisasi sistem.

REFERENCES

- [1] Republik Indonesia, “Undang-Undang Republik Indonesia Nomor 24 Tahun 2013 Tentang Administrasi Kependudukan,” 2013
- [2] S. Salsabila, A. Zetra, dan R. E. Putera, “Penerapan E-Government Dalam Pelayanan KTP Pada Dinas Kependudukan dan Pencatatan Sipil Kota Padang,” *J. Ilmu Adm. Negara ASIAN (Asosiasi Ilmuwan Adm. Negara)*, vol. 9, no. 2, hal. 314–324, 2022, doi: 10.47828/jianaasian.v9i2.65.
- [3] M. Tampang, I. Sartika, dan F. Ruhana, “Kualitas Pelayanan Publik Dalam Pembuatan Kartu Tanda Penduduk Elektronik (E-Ktp) Di Suku Dinas Kependudukan Dan Catatan Sipil Kota Jakarta Selatan,” *J. Kaji. Pemerintah J. Gov. Soc. Polit.*, vol. 10, no. 1, hal. 73–85, 2024, doi: 10.25299/jkp.2024.vol10(1).16958.
- [4] A. Doramia Lumbanraja, “Urgensi Transformasi Pelayanan Publik melalui E-Government Pada New Normal dan Reformasi Regulasi Birokrasi,” *Adm. Law Gov. J.*, vol. 3, no. 2, hal. 220–231, 2020, doi: 10.14710/alj.v3i2.220-231.
- [5] M. Alfarizi, “Digitalisasi Kartu Tanda Penduduk dan Partisipasi Milenial-Gen Z: Investigasi Penerimaan Transformasi Digital dalam Kebijakan Kependudukan Indonesia,” *J. Stud. Kebijak. Publik*, vol. 2, no. 1 SE-, hal. 41–54, Mei 2023, doi: 10.21787/jskp.2.2023.41-54.
- [6] K. Luar Negeri, “Peraturan Menteri Luar Negeri Republik Indonesia Nomor 1 Tahun 2024 Tentang Standar Layanan Informasi Publik Di Kementerian Luar Negeri Dan Perwakilan Republik Indonesia,” 2016
- [7] M. R. Reyvansyah, “Penerapan Metode Optical Character Recognition (OCR) Untuk Mengambil Data Arsip,” *J. Tek. Elektro dan Komput. TRIAC*, vol. 10, no. 2, hal. 44–50, 2023, doi: 10.21107/triac.v10i2.20809.
- [8] J. Felisa, D. Setiawan, dan I. Khalisa, “Perancangan Perangkat Lunak Pengenalan Karakter Plat Nomor Kendaraan dengan Metode Convolutional Neural Network,” *Media Inform.*, vol. 21, no. 3, hal. 280–306, 2023, doi: 10.37595/mediainfo.v21i3.156.
- [9] I. Wijaya dan C. Lubis, “Pengimplementasian Ocr Menggunakan Cnn Untuk Ekstraksi Teks Pada Gambar,” *J. Ilmu Komput. dan Sist. Inf.*, vol. 10, no. 1, 2022, doi: 10.24912/jiksi.v10i1.17836.
- [10] F. Imran, M. A. Hossain, dan M. Al Mamun, “Identification and Recognition of Printed Distorted Characters Using Proposed DCR Method,” *2020 IEEE Reg. 10 Symp. TENSYP 2020*, no. March 2023, hal. 1478–1481, 2020, doi: 10.1109/TENSYP50017.2020.9230646.
- [11] M. S. Gumilang dan D. Avianto, “RECOGNITION OF REAL-TIME HANDWRITTEN CHARACTERS USING CONVOLUTIONAL NEURAL NETWORK ARCHITECTURE,” *J.*

- Tek. Inform.*, vol. 4, no. 5 SE-Articles, hal. 1143–1150, Okt 2023, doi: 10.52436/1.jutif.2023.4.5.993.
- [12] M. Haris, M. G. Suryanata, dan M. Yetri, “Implementasi OCR Menggunakan Algoritma Template Matching Correlation pada Pengarsipan e-KTP,” *J-SISKO TECH (Jurnal Teknol. Sist. Inf. dan Sist. Komput. TGD)*, vol. 6, no. 2, hal. 281, 2023, doi: 10.53513/jsk.v6i2.8134.
- [13] Fatih Gesang Panuntun dan Rr. Hajar Puji Sejati, “Sistem Otomatisasi Deteksi dan Ekstraksi Data KTP Berbasis Convolutional Neural Network dan Optical Character Recognition,” *JSIAI (Journal Sci. Appl. Informatics)*, vol. 7, no. 3, hal. 464–471, 2024, doi: 10.36085/jsai.v7i3.7269.
- [14] G. Sugiarta, D. P. Andini, dan S. Hidayatullah, “Ekstraksi Informasi/Data e-KTP Menggunakan Optical Character Recognition Convolutional Neural Network,” *JTERA (Jurnal Teknol. Rekayasa)*, vol. 6, no. 1, hal. 1, 2021, doi: 10.31544/jtera.v6.i1.2021.1-6.
- [15] A. R. Irawati, D. Kurniawan, Y. T. Utami, dan R. Taufik, “An Exploration of TensorFlow-Enabled Convolutional Neural Network Model Development for Facial Recognition: Advancements in Student Attendance System,” vol. 11, no. 2, hal. 413–428, 2024, doi: 10.15294/sji.v11i2.3585.
- [16] A. Kabir Rifai, M. Rafi Muttaqin, D. Irmayanti, S. Tinggi Teknologi Wastukencana, J. Cikopak No, dan J. Barat, “Pemanfaatan Algoritma Convolutional Neural Network Dengan Untuk Mendeteksi Penyakit Pada Tumbuhan Jagung,” *Sist. J. Ilm. Sist. Inf.*, vol. Vol.1, no. 1, hal. 18–26, 2024, [Daring]. Tersedia pada: <https://ejournal.rizaniamedia.com/index.php/sistematis>
- [17] C. A. Maharani, B. Warsito, dan R. Santoso, “Analisis Sentimen Vaksin Covid-19 Pada Twitter Menggunakan Recurrent Neural Network (Rnn) Dengan Algoritma Long Short-Term Memory (Lstm),” *J. Gaussian*, vol. 12, no. 3, hal. 403–413, 2024, doi: 10.14710/j.gauss.12.3.403-413.
- [18] F. G. Safinatunnajah, A. Prasetiadi, dan M. Wibowo, “CLASSIFICATION OF CAT SOUNDS USING CONVOLUTIONAL NEURAL NETWORK (CNN) AND LONG SHORT-TERM MEMORY (LSTM) METHODS,” *J. Tek. Inform.*, vol. 3, no. 5 SE-Articles, hal. 1349–1353, Okt 2022, doi: 10.20884/1.jutif.2022.3.5.373.
- [19] M. R. R. Allaam dan A. T. Wibowo, “Klasifikasi Genus Tanaman Anggrek Menggunakan Metode Convolutional Neural Network (CNN),” *eProceedings Eng.*, vol. 8, no. 2, hal. 1153, 2021, [Daring]. Tersedia pada: <https://openlibrarypublications.telkomuniversity.ac.id/index.php/engineering/article/view/14708>
- [20] P. F. Johari, N. Arifin, M. Muzaki, dan M. S. A. Utama, “Corn Leaf Diseases Classification Using CNN with GLCM, HSV, and L*a*b* Features,” *J. Tek. Inform.*, vol. 6, no. 2, hal. 709–722, 2025, doi: 10.52436/1.jutif.2025.6.2.4345.
- [21] K. Azmi, S. Defit, dan S. Sumijan, “Implementasi Convolutional Neural Network (CNN) Untuk Klasifikasi Batik Tanah Liat Sumatera Barat,” *J. Unitek*, vol. 16, no. 1, hal. 28–40, 2023, doi: 10.52072/unitek.v16i1.504.
- [22] S. H. Apandi, J. Sallim, dan R. Mohamed, “A Convolutional Neural Network (CNN) Classification Model for Web Page: A Tool for Improving Web Page Category Detection Accuracy,” *JITSI J. Ilm. Teknol. Sist. Inf.*, vol. 4, no. 3, hal. 110–121, 2023, doi: 10.30630/jitsi.4.3.181.
- [23] M. T R, V. K. V, D. K. V, O. Geman, M. Margala, dan M. Guduri, “The stratified K-folds cross-validation and class-balancing methods with high-performance ensemble classifiers for breast cancer classification,” *Healthc. Anal.*, vol. 4, no. July, hal. 100247, 2023, doi: 10.1016/j.health.2023.100247.
- [24] H. Ma’we, A. Y. Husodo, dan B. Irmawati, “Performance Comparison of Naive Bayes and Bidirectional Lstm Algorithms in Bsi Mobile Review Sentiment Analysis,” *J. Tek. Inform.*, vol. 6, no. 1, hal. 159–172, 2024, doi: 10.52436/1.jutif.2024.5.6.4178.
- [25] S. Widodo, H. Brawijaya, dan S. Samudi, “Stratified K-fold cross validation optimization on machine learning for prediction,” *Sinkron*, vol. 7, no. 4, hal. 2407–2414, 2022, doi: 10.33395/sinkron.v7i4.11792.
- [26] A. Rianti, N. W. A. Majid, dan A. Fauzi, “CRISP-DM: Metodologi Proyek Data Science,” *Pros. Semin. Nas. Teknol. ...*, hal. 107–114, 2023, [Daring]. Tersedia pada: <http://ojs.udb.ac.id/index.php/Senatib/article/view/3015>
- [27] S. Alden dan B. N. Sari, “Implementasi Algoritma CNN Untuk Pemilahan Jenis Sampah Berbasis

- Android Dengan Metode CRISP-DM,” *J. Inform.*, vol. 10, no. 1, hal. 62–71, 2023, doi: 10.31294/inf.v10i1.14985.
- [28] R. Khanam, M. Hussain, R. Hill, dan P. Allen, “A Comprehensive Review of Convolutional Neural Networks for Defect Detection in Industrial Applications,” *IEEE Access*, vol. 12, hal. 94250–94295, 2024, doi: 10.1109/ACCESS.2024.3425166.
- [29] D. B. Santosa, A. Wahana, dan W. Uriawan, “Implementation of Convolutional Neural Network Using Mobilenetv2 To Distinguish Human and Artificial Intelligence Painting,” *J. Tek. Inform.*, vol. 6, no. 1, hal. 441–452, 2025, doi: 10.52436/1.jutif.2025.6.1.3827.
- [30] W. Wijiyanto, A. I. Pradana, S. Sopingi, dan V. Atina, “Teknik K-Fold Cross Validation untuk Mengevaluasi Kinerja Mahasiswa,” *J. Algoritma*, vol. 21, no. 1, hal. 239–248, 2024, doi: 10.33364/algoritma/v.21-1.1618.
- [31] B. S. Abunasser, M. R. J. AL-Hiealy, I. S. Zaqout, dan S. S. Abu-Naser, “Convolution Neural Network for Breast Cancer Detection and Classification Using Deep Learning,” *Asian Pacific J. Cancer Prev.*, vol. 24, no. 2, hal. 531–544, 2023, doi: 10.31557/APJCP.2023.24.2.531.
- [32] I. S. Had, W. Maulana Baihaqi, dan D. Putriana Nuramanah Kinding, “Improving Tesseract OCR Accuracy Using SymSpell Algorithm on Passport Data,” *Sinkron*, vol. 9, no. 1, hal. 374–381, 2025, doi: 10.33395/sinkron.v9i1.14395.
- [33] Y. Widyaningsih, G. P. Arum, dan K. Prawira, “Aplikasi K-Fold Cross Validation Dalam Penentuan Model Regresi Binomial Negatif Terbaik,” *BAREKENG J. Ilmu Mat. dan Terap.*, vol. 15, no. 2, hal. 315–322, 2021, doi: 10.30598/barekengvol15iss2pp315-322.
- [34] J. A. Wuisan, A. Jacobus, dan S. R. U. A. Sompie, “Data Balancing Methods on Radiographic Image Classification on Unbalance Dataset (Perbandingan Metode Penyeimbangan Data pada Klasifikasi Citra Radiografi pada Dataset Tidak Seimbang),” *J. Tek. Elektro dan Komput.*, vol. 11, no. 1, hal. 1–8, 2022.
- [35] M. R. Hartono, C. A. Sari, dan R. R. Ali, “Football Player Tracking, Team Assignment, and Speed Estimation Using Yolov5 and Optical Flow,” *J. Tek. Inform.*, vol. 6, no. 1, hal. 51–62, 2025, doi: 10.52436/1.jutif.2025.6.1.4165.