



Sistem Identifikasi Citra Huruf Aksara Minangkabau Berbasis Convolutional Neural Network

Riyan Saputra¹, Agung Ramadhanu², Rini Sovia³

¹²³Program Magister Teknologi Informasi, Ilmu Komputer, Universitas Putra Indonesia “YPTK” Padang

¹alamat.emailnya.riyan@gmail.com, ²agung_ramadhanu@upiypk.ac.id, ³rini_sovia@upiypk.ac.id

Abstract

Preserving regional scripts is crucial for safeguarding the nation's cultural heritage. The Minangkabau script, as one of Indonesia's cultural assets, remains under-researched and lacks an adequate digitization system. This study represents an initial exploration of Minangkabau script recognition using a Convolutional Neural Network (CNN) approach as an effort to document the script and assess its digitization potential. CNNs are a type of deep learning model designed to process structured grid data such as images and prior studies have shown that they perform very well in handwriting recognition. The character images used in this research were obtained from museum sources and handwritten samples from 31 volunteers. The dataset comprises 4,650 character images across 75 classes with various combinations of diacritics on the five vowels then processed through grayscale conversion, contrast enhancement, segmentation, and augmentation, yielding a total of 8,537 images. The designed CNN model classifies characters into 75 classes. The evaluation indicates that the model can recognize characters very well. Testing showed 99% accuracy in a limited evaluation scenario on 500 test samples. These findings provide an initial foundation for subsequent academic studies and broader cultural discussions regarding the Minangkabau script.

Keywords: *Minangkabau script, convolutional neural network (CNN), image processing, artificial intelligence, cultural documentation.*

Abstrak

Pelestarian aksara daerah penting untuk menjaga warisan budaya bangsa. Aksara Minangkabau, sebagai salah satu kekayaan budaya Indonesia, masih minim penelitian dan belum memiliki sistem digitalisasi memadai. Penelitian ini merupakan tahap awal eksplorasi pengenalan aksara Minangkabau menggunakan pendekatan *Convolutional Neural Network* (CNN) sebagai upaya mendokumentasikan dan menguji potensi digitalisasi aksara tersebut. CNN merupakan salah satu model *deep learning* yang dirancang untuk memproses *data grid* terstruktur seperti citra. Penelitian sebelumnya menunjukkan kinerja CNN sangat baik dalam pengenalan tulisan tangan. Citra aksara yang digunakan dalam penelitian ini diperoleh dari sumber museum dan tulisan tangan dari 31 sukarelawan. Dataset terdiri dari 4.650 citra karakter dari 75 kelas dengan berbagai kombinasi tanda baca pada lima huruf vokal, yang kemudian diproses melalui konversi *grayscale*, peningkatan kontras, segmentasi, dan augmentasi hingga menghasilkan total 8.537 citra. Model CNN yang dirancang mengklasifikasikan karakter ke dalam 75 kelas. Hasil pengujian mengindikasikan bahwa model dapat mengenali karakter dengan sangat baik. Pengujian menunjukkan akurasi 99% dalam skenario pengujian terbatas pada 500 data uji. Temuan ini memberikan landasan awal untuk digunakan dalam kajian akademis lanjutan maupun diskusi kultural yang lebih luas terkait keberadaan aksara Minangkabau.

Kata kunci: aksara Minangkabau, convolutional neural network (CNN), pemrosesan citra, kecerdasan buatan, dokumentasi budaya.

1. Pendahuluan

Indonesia merupakan negara *multicultural* [1], terdiri dari sekitar 300 kelompok etnis dan menggunakan hampir 200 bahasa [2]. Salah satunya adalah masyarakat Minangkabau yang sebagian besar berasal dari wilayah Sumatera Barat, Indonesia [3]. Bahasa Minang adalah bahasa daerah yang sebagian besar digunakan oleh orang-orang Minangkabau [4]. Bahasa Minang termasuk dalam kelompok bahasa Austronesia [5]. Bahasa Minang terdiri dari 5 vokal (a, i, u, e, o) dan 20 konsonan, bersama dengan 6 diftong (i, u, aw, ay, uy, e). Struktur fonemik ini sangat penting untuk memahami pengucapan dan tata bahasa [4]. Bahasa ini mencerminkan identitas budaya Minangkabau dan merupakan bagian integral dari interaksi sosial mereka [6].

Sistem penulisan (aksara) untuk bahasa adalah bentuk material yang mewakili makna dan suara bahasa tertentu [7]. Aksara Minangkabau masih perlu pembuktian secara ilmiah agar diakui keberadaannya. Di luar literatur arkeologis arus utama, ada klaim lokal bahwa Minangkabau memiliki aksara sendiri, dengan dua varian temuan, yaitu Pariangan dan Ruweh Buku Silek Aia. Varian Pariangan berasal dari Padang Panjang, disebut memiliki 15 grafem dasar. bunyi vokal ditandai titik (atas=i, bawah=u), dan ada klaim arah tulis kanan ke kiri. Bukti rujukan utama berupa artikel komunitas/media lokal dan testimoni penggiat [8]. Sedangkan varian Ruweh Buku Silek Aia dari Solok, dilaporkan alfabetis (~21 huruf), ada tanda baca modern, dan konon penataan kata ke bawah pada penulisan tertentu. Sumber juga didominasi tulisan populer dan wawancara tokoh lokal [9].

Penelitian sebelumnya menjelaskan bahwa pentingnya digitalisasi terhadap aksara Suriah Timur karena terancam punah [10]. Penelitian yang sama menunjukkan bahwa mendigitalkan aksara Indus dapat menghasilkan wawasan tentang aspek budaya dan membangun hubungan dengan keluarga linguistik yang dikenal [11]. Penelitian lainnya menjelaskan digitalisasi aksara Bali dapat memicu ketertarikan generasi muda untuk mempelajari aksara Bali [12].

Implementasi dalam penyelesaian masalah identifikasi aksara dapat mengadopsi konsep *Neural Network*. *Neural Network* merupakan sebuah model komputasi yang terinspirasi dari struktur dan mekanisme kerja jaringan saraf biologis pada otak [13], yang terdiri dari kumpulan unit pemroses sederhana (*neuron* buatan) yang saling terhubung melalui bobot (*weights*) dan bekerja secara paralel [14]. Jaringan ini mempelajari pola atau representasi dari data melalui proses pelatihan berbasis pengoptimalan [15]. Konsep tersebut erat kaitannya dengan perkembangan kecerdasan buatan, yang dalam penerapannya terbukti meningkatkan efisiensi sistem informasi, misalnya pada otomasi perpustakaan [16]. Selain itu, kecerdasan buatan juga berkontribusi positif dalam konteks pendidikan, terutama dalam meningkatkan efektivitas pembelajaran mahasiswa [17].

Salah satu metode yang terdapat pada konsep *Neural Network* yaitu *Convolutional Neural Network* (CNN). CNN merupakan model *Deep Learning* yang dirancang untuk memproses *data grid* terstruktur, seperti gambar [18]. CNN menggunakan lapisan konvolusi untuk secara otomatis mengekstrak fitur dari gambar, mengurangi kebutuhan akan rekayasa fitur manual [19]. Arsitektur dari CNN terdiri dari lapisan seperti konvolusi, aktivasi, dan *down-up-sampling* [20].

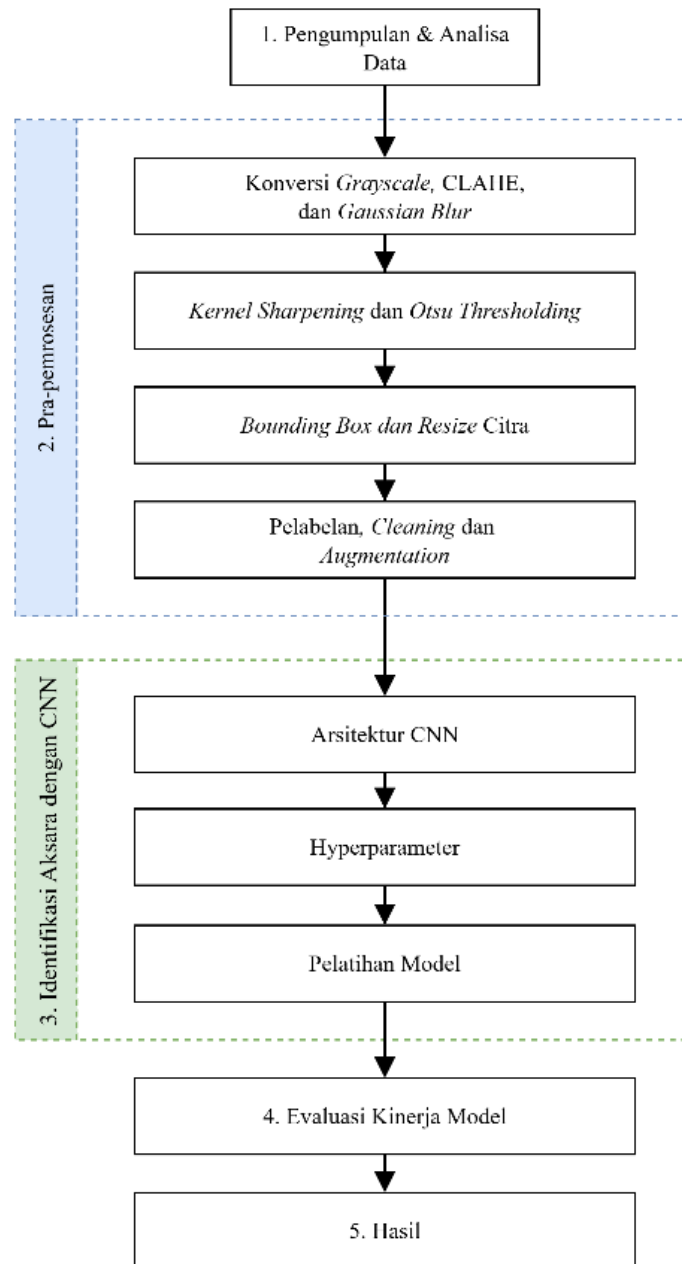
Berdasarkan kajian terhadap penelitian-penelitian sebelumnya mengenai performa metode CNN dalam penyelesaian permasalahan identifikasi aksara, diperoleh temuan bahwa model CNN mampu mencapai tingkat akurasi hingga 97,34% [21]. Penelitian lain menunjukkan bahwa penerapan model CNN dalam tugas identifikasi aksara Sasak mampu menghasilkan akurasi sebesar 92% [22]. Penelitian terhadap identifikasi aksara Lampung dengan CNN menunjukkan akurasi 100% pada perhitungan *confusion matrix* [23].

Berdasarkan uraian sebelumnya, penelitian ini bertujuan untuk melakukan identifikasi aksara Minangkabau secara otomatis dengan memanfaatkan metode CNN. Penerapan CNN dalam penelitian ini diharapkan mampu menghasilkan model identifikasi yang memiliki tingkat ketepatan dan akurasi tinggi. Model yang dikembangkan menawarkan kebaruan dalam konteks pengenalan aksara Minangkabau, baik dari segi metodologi maupun efektivitas proses identifikasi. Selain itu, penelitian ini diharapkan dapat menghasilkan sistem identifikasi yang tepat dan akurat, sehingga dapat menjadi landasan awal bagi pengembangan kajian akademis lanjutan maupun memperkaya diskusi kultural yang lebih luas terkait pelestarian dan pemanfaatan aksara Minangkabau.

2. Metode Penelitian

Pendekatan penelitian yang digunakan adalah pendekatan kuantitatif. Seluruh proses penelitian mulai dari pengumpulan dan persiapan data, pelatihan model, hingga evaluasi kinerja sistem dilaksanakan secara terukur, terstruktur, dan berbasis pada data numerik. Pendekatan ini memungkinkan peneliti memperoleh hasil yang objektif melalui analisis statistik terhadap metrik evaluasi, sehingga dapat divalidasi secara ilmiah.

Proses penelitian dilakukan secara bertahap. Tahapannya dimulai dari persiapan dataset, pra-pemrosesan gambar, perancangan arsitektur CNN, pelatihan model, hingga evaluasi kinerja menggunakan metrik tertentu. Tahapan penelitian ini divisualisasikan pada Gambar 1.

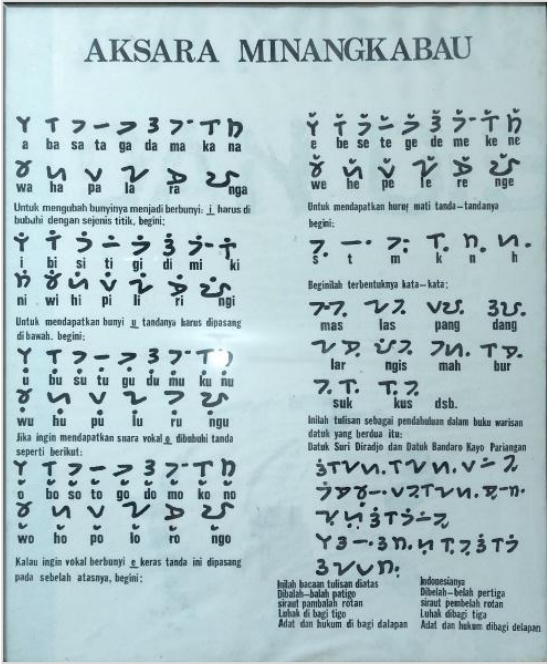


Gambar 1. Kerangka Kerja Penelitian

Kerangka kerja penelitian pada Gambar 1 terdiri atas lima tahap utama yang dimulai dari proses pengumpulan dan analisis data, kemudian dilanjutkan dengan tahap pra-pemrosesan berupa konversi citra, penajaman, segmentasi, dan augmentasi data. Tahap berikutnya adalah identifikasi aksara menggunakan CNN melalui perancangan arsitektur, penentuan *hyperparameter*, serta pelatihan model, yang kemudian menghasilkan model klasifikasi. Setelah model terbentuk, dilakukan evaluasi kinerja untuk menilai performa sistem hingga akhirnya diperoleh hasil yang menjadi output penelitian ini.

2.1. Pengumpulan dan Analisis Data

Pengumpulan data dalam penelitian ini dilakukan melalui dua sumber utama, yaitu dokumentasi langsung artefak aksara dan kontribusi sukarelawan. Aksara Minangkabau varian Pariangan yang terpajang di ruang penyimpanan Museum Adityawarman didokumentasikan menggunakan kamera smartphone Poco X6. Dokumentasi ini bertujuan merekam bentuk visual aksara asli sebagai acuan dalam proses pengembangan dataset awal. Hasil dokumentasi ini dapat dilihat pada Gambar 2



Gambar 2. Dokumentasi Aksara Minangkabau

Pada Gambar 2 diketahui bahwa setiap huruf dari aksara Minangkabau memiliki variasi bentuk dasar dan tanda diakritik. Terdiri dari 75 huruf yang merepresentasikan kombinasi dari 5 huruf vokal dasar dengan 15 kemungkinan kombinasi tanda baca (gabungan dari titik dan garis di atas/bawah). Adapun daftar dari aksara tersebut disajikan pada Tabel 1 berikut.

Tabel 1. Aksara Minangkabau

	b	d	g	h	k	l	m	n	p	r	s	t	ng	W	
a	a	b	d	g	h	k	l	m	n	p	r	s	t	v	w
e	a	b	d	g	h	k	l	m	n	p	r	s	t	v	w
	e	e	e	e	e	e	e	e	e	e	e	e	e	e	e
i	ai	bi	di	gi	hi	ki	li	mi	ni	pi	ri	si	ti	vi	wi
o	a	b	d	g	h	k	l	m	n	p	r	s	t	v	w
	o	o	o	o	o	o	o	o	o	o	o	o	o	o	o
u	a	b	d	g	h	k	l	m	n	p	r	s	t	v	w
	u	u	u	u	u	u	u	u	u	u	u	u	u	u	u

Tabel 1 menyajikan huruf aksara Minangkabau yang memiliki aturan baca bergantung pada tanda yang menyertainya. Apabila huruf b ditulis tanpa tanda baca, maka ia dibaca “ba”. Jika diberi garis atas seperti be, maka bacaan berubah menjadi “be”, sedangkan penambahan titik di atas huruf seperti bi menjadikannya dibaca “bi”. Selanjutnya, ketika digaris di bagian bawah bo, bacaannya menjadi “bo”, dan apabila ditambah titik di bawah seperti bu, maka huruf tersebut dibaca “bu”.

Mengingat keterbatasan jumlah sampel, dilakukan pengumpulan data tambahan melalui partisipasi sukarelawan. Formulir isian yang telah dirancang memuat 75 jenis karakter aksara, masing-masing ditulis sebanyak dua kali,

sehingga menghasilkan total 150 sampel karakter tulisan tangan per partisipan, Desain dari fomulir tersebut dapat dilihat pada Gambar 3.

Gambar 3. Fomulir Isian Aksara Minangkabau

Gambar 3 merupakan formulir isian aksara Minangkabau yang dicetak pada kertas berukuran A4. Hasil isian dipindai menggunakan pemindai Epson L360 dengan resolusi 100 dpi dan disimpan dalam format JPEG. Total partisipasi sebanyak 31 orang menghasilkan 4.650 sampel karakter tulisan tangan, yang kemudian digunakan untuk membentuk dataset penelitian. Data citra aksara yang terkumpul mengandung noise, latar belakang yang kompleks, serta ukuran karakter yang tidak seragam. Kondisi ini menyulitkan proses pelabelan dan ekstraksi fitur yang akurat, sehingga diperlukan tahapan pra-pemrosesan yang cermat untuk memastikan kualitas data sebelum masuk ke tahap identifikasi atau pelatihan model.

2.2. Pra-pemrosesan

Pra-pemrosesan merupakan aspek fundamental dari pemrosesan citra yang melibatkan penerapan berbagai teknik untuk menyiapkan citra mentah untuk analisis. Proses ini mencakup langkah-langkah seperti pengurangan noise, peningkatan citra, dan transformasi geometris, yang meningkatkan kualitas citra dan memfasilitasi ekstraksi informasi berharga [24]. Tahapan ini dilakukan untuk meningkatkan performa CNN dalam identifikasi aksara Minangkabau. Tahapan pra-pemrosesan yang dilakukan dalam penelitian ini meliputi:

- Konversi citra menjadi *grayscale*, penerapan *Contrast Limited Adaptive Histogram Equalization* (CLAHE), dan *Gaussian Blur* untuk perbaikan kontras dan pengurangan noise.
- Kernel sharpening* dan *Otsu thresholding* untuk mempertajam tepi serta memisahkan objek dari latar belakang.
- Bounding box* dan *resize* citra untuk menyeragamkan ukuran input.
- Pelabelan data, *cleaning* untuk menghapus data yang tidak valid, serta *augmentation* untuk memperbanyak variasi data pelatihan.

2.3. Identifikasi Aksara dengan CNN

Perancangan arsitektur CNN dalam penelitian ini dilakukan secara menyeluruh, mencakup pemilihan jumlah lapisan, jenis fungsi aktivasi, serta struktur filter yang disesuaikan dengan karakteristik data dan tingkat kompleksitas klasifikasi [25]. Penetapan *hiperparameter* seperti *learning rate*, ukuran *batch*, dan jumlah *epoch* dilakukan secara hati-hati, karena parameter-parameter ini memiliki pengaruh signifikan terhadap stabilitas dan konvergensi model selama proses pelatihan [26]. Dataset yang digunakan telah melalui tahapan pra-pemrosesan, termasuk normalisasi dan augmentasi data, untuk memastikan konsistensi dan kualitas input yang optimal. Pelatihan model dilakukan secara sistematis berdasarkan arsitektur dan parameter yang telah ditentukan, dengan

tujuan menghasilkan model CNN yang andal, stabil, dan memiliki kemampuan generalisasi yang baik terhadap data baru [27].

2.3. Evaluasi Kinerja Model

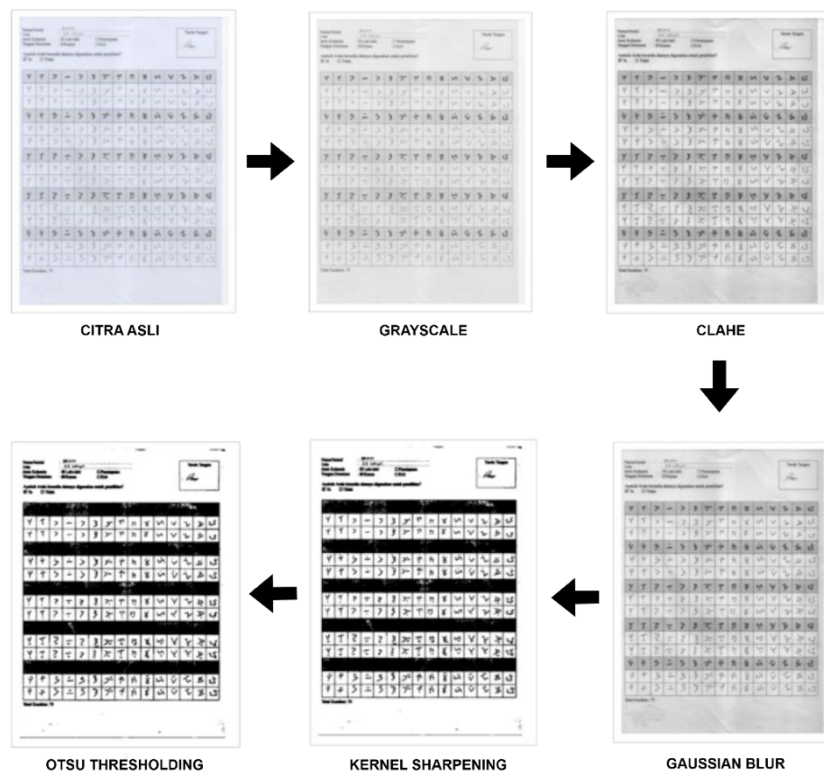
Model yang telah selesai dilatih kemudian dievaluasi secara menyeluruh menggunakan dataset pengujian (*testing set*), dengan tujuan mengukur kinerja prediktif secara objektif. Proses evaluasi ini mencakup penyusunan *confusion matrix*, yang merupakan alat penting dalam klasifikasi yang memungkinkan visualisasi distribusi prediksi benar (*true positives*, *true negatives*) dan kesalahan (*false positives*, *false negatives*) untuk setiap kelas [28]. Selain itu, perhitungan *classification report* dilakukan untuk mendapatkan metrik kunci seperti akurasi, presisi (*precision*), daya tangkap (*recall*), dan *F1-score*. Metrik-metrik tersebut memberikan gambaran menyeluruh mengenai efektivitas model dalam memprediksi setiap kelas secara seimbang [29]. Evaluasi berbasis metrik ini tidak hanya mencerminkan performa agregat, tetapi juga membantu mengidentifikasi kekuatan dan kelemahan model pada klasifikasi kelas tertentu, di mana perbedaan presisi dan recall antar kelas dianalisis secara mendalam melalui *confusion matrix* dan metrik *F1-score* [30]. Dengan demikian, pendekatan evaluasi ini memastikan bahwa model CNN yang dikembangkan tidak hanya akurat secara keseluruhan tetapi juga andal dan stabil dalam menangani variasi kelas secara individual.

3. Hasil dan Pembahasan

Bab ini menyajikan hasil penelitian sekaligus pembahasannya. Hasil ditampilkan sesuai dengan fokus penelitian, kemudian dianalisis dan dikaitkan dengan teori serta temuan terdahulu guna melihat kesesuaian maupun perbedaannya. Pembahasan bertujuan memberi interpretasi terhadap temuan serta implikasinya terhadap kajian yang diteliti.

3.1. Pra-pemrosesan

Pra-pemrosesan citra memainkan peran penting dalam menyiapkan data input yang berkualitas tinggi untuk pelatihan model CNN, terutama dalam kasus pengenalan karakter aksara dari citra hasil pindai kertas fomulir isian. Tahapan ini tidak hanya membantu meningkatkan kualitas visual citra, tetapi juga meminimalkan noise yang dapat mengganggu proses ekstraksi fitur. Adapun hasil dari tahapan tersebut dapat dilihat pada Gambar 4.



Gambar 4. Hasil Implementasi Pra-pemrosesan

Gambar 4 menjelaskan bahwa proses pra-pemrosesan dimulai dengan konversi citra hasil pemindaian ke format *grayscale*. Konversi ini menyederhanakan data dengan menghilangkan informasi warna, sehingga pemrosesan dapat difokuskan pada intensitas piksel dan pola tekstur yang lebih relevan untuk mengenali bentuk karakter.

Setelah itu, dilakukan peningkatan kontras menggunakan metode CLAHE (*Contrast Limited Adaptive Histogram Equalization*) dengan *clip limit* = 2.0 dan *tileGridSize* = (8, 8). Teknik ini memperjelas detail karakter di dalam kotak isian, khususnya pada area dengan pencahayaan tidak merata. Untuk mengurangi noise yang mungkin timbul dari proses scan, digunakan *Gaussian blur* dengan *kernel* 3×3 dan *sigmaX* = 0, yang membantu meratakan nilai piksel tanpa mengaburkan struktur penting. Dilanjutkan dengan penerapan *kernel sharpening* menggunakan nilai [[0, -0.5, 0], [-0.5, 5, -0.5], [0, -0.5, 0]] untuk menegaskan tepi dan garis karakter agar fitur-fitur penting dapat lebih mudah dikenali oleh model. Selanjutnya, dilakukan segmentasi menggunakan *Otsu thresholding*, yang secara otomatis menentukan nilai ambang optimal berdasarkan histogram intensitas untuk memisahkan karakter dari latar belakang.

Setelah segmentasi, diterapkan proses pembungkaihan (*bounding box*) pada masing-masing kotak isian dalam form untuk mengekstrak setiap karakter secara individual. Proses ini sangat penting untuk membentuk dataset per karakter, yang kemudian akan digunakan untuk pelatihan model klasifikasi 75 kelas aksara. Setiap karakter hasil ekstraksi kemudian diubah ukurannya menjadi 64×64 piksel, guna menyeragamkan dimensi input sesuai dengan kebutuhan arsitektur CNN.

Sebagai langkah augmentasi data, dilakukan *elastic distortion* dengan parameter *probability*=1, *grid_width*=4, *grid_height*=4, dan *magnitude*=2. Teknik ini mendistorsi bentuk karakter secara elastis, menciptakan variasi tambahan pada dataset untuk meningkatkan robustnes model terhadap deformasi, rotasi, dan perubahan posisi karakter. Urutan pra-pemrosesan ini secara langsung berkontribusi pada peningkatan akurasi dan kemampuan generalisasi model CNN dalam mengenali pola karakter aksara yang kompleks.

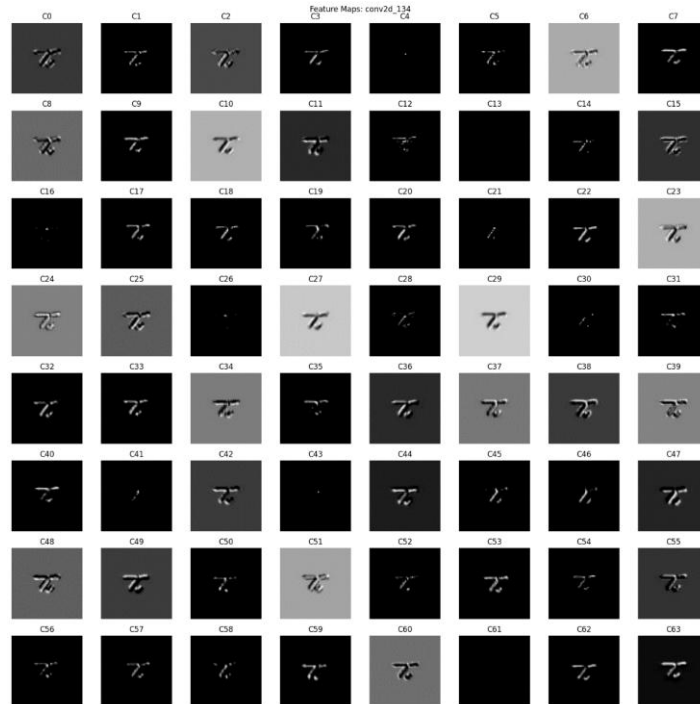
3.2. Identifikasi Aksara dengan CNN

Dalam menyelesaikan permasalahan identifikasi citra aksara sebanyak 75 kelas, digunakan arsitektur CNN dengan beberapa blok konvolusi dan lapisan *dense* sebagai bagian dari jaringan saraf tiruan. Hasil proses identifikasi aksara dengan CNN disajikan pada Tabel 2.

Tabel 2. Arsitektur CNN Usulan

Lapisan	Output Shape	Jumlah Parameter
Conv2D (64 filter, 3x3)	(62, 62, 64)	640
BatchNormalization	(62, 62, 64)	256
Conv2D (64 filter, 3x3)	(60, 60, 64)	36.928
BatchNormalization	(60, 60, 64)	256
MaxPooling2D (2x2)	(30, 30, 64)	0
Conv2D (128 filter, 3x3)	(28, 28, 128)	73.856
BatchNormalization	(28, 28, 128)	512
Conv2D (128 filter, 3x3)	(26, 26, 128)	147.584
BatchNormalization	(26, 26, 128)	512
MaxPooling2D (2x2)	(13, 13, 128)	0
Conv2D (256 filter, 3x3)	(11, 11, 256)	295.168
BatchNormalization	(11, 11, 256)	1.024
Conv2D (256 filter, 3x3)	(9, 9, 256)	590.080
BatchNormalization	(9, 9, 256)	1.024
MaxPooling2D (2x2)	(4, 4, 256)	0
Conv2D (512 filter, 3x3)	(2, 2, 512)	1.180.160
BatchNormalization	(2, 2, 512)	2.048
MaxPooling2D (2x2)	(1, 1, 512)	0
Flatten	(512)	0
Dense (1024 neuron)	(1024)	525.312
Dense (512 neuron)	(512)	524.800
Dropout	(512)	0
Dense (75 neuron)	(75)	38.475
Total Parameter	-	3.419.167

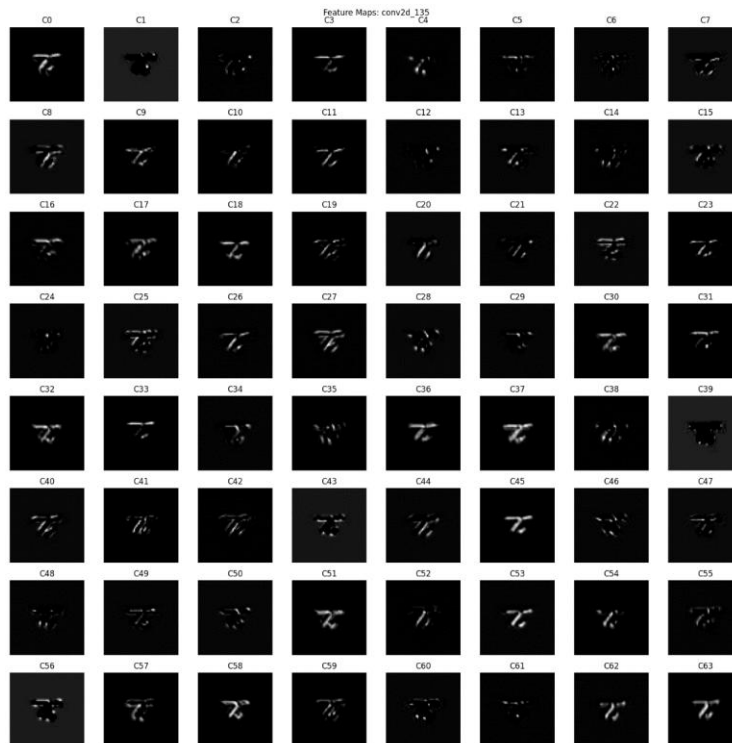
Lapisan pertama dari arsitektur CNN pada Tabel 2 adalah Conv2D dengan 64 *filter* berukuran 3x3, *stride* 1, dan *padding* 'valid', yang menghasilkan output berukuran 62x62x64. Lapisan ini diikuti oleh BatchNormalization untuk menstabilkan distribusi aktivasi dan mempercepat konvergensi. Hal ini dibuktikan dengan *feature maps* yang divisualisasikan pada Gambar 5.



Gambar 5. *Feature Maps* dari Conv2D Lapisan Pertama

Pada Gambar 5 terlihat bahwa hampir seluruh kanal menampilkan bentuk huruf secara utuh dengan garis tepi yang jelas dan kontras tinggi di sekitar batang huruf maupun tanda diakritik. Aktivasi ini menunjukkan bahwa lapisan awal berperan dalam menangkap fitur-fitur dasar (*low-level features*) seperti tepi horizontal, vertikal, lengkungan, dan titik kecil. Pada tahap ini, detail spasial masih terjaga secara optimal untuk digunakan pada tahap ekstraksi berikutnya.

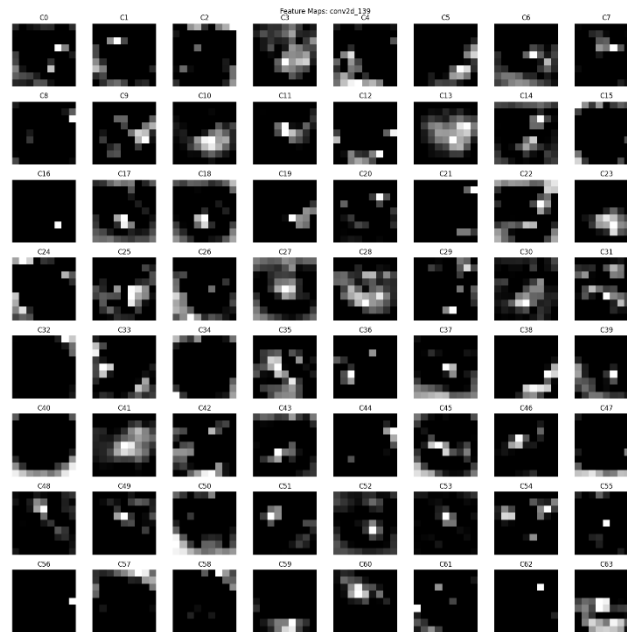
Lapisan Conv2D kedua mempertahankan jumlah *filter* (64) namun mengurangi dimensi spasial menjadi 60x60, dan kembali diikuti oleh BatchNormalization. Gambar 6 merupakan *feature maps* dari Conv2D lapisan kedua.



Gambar 6. *Feature Maps* dari Conv2D Lapisan Kedua

Pada Gambar 6 memperlihatkan bahwa aktivasi mulai terfokus pada area tertentu dari huruf, khususnya pada bagian yang mengandung tanda diakritik, sementara latar belakang menjadi lebih redup. Hal ini menandakan bahwa lapisan ini memperkuat fitur lokal yang relevan dan mengurangi informasi yang tidak diperlukan, sehingga jaringan mulai sensitif terhadap perbedaan posisi diakritik

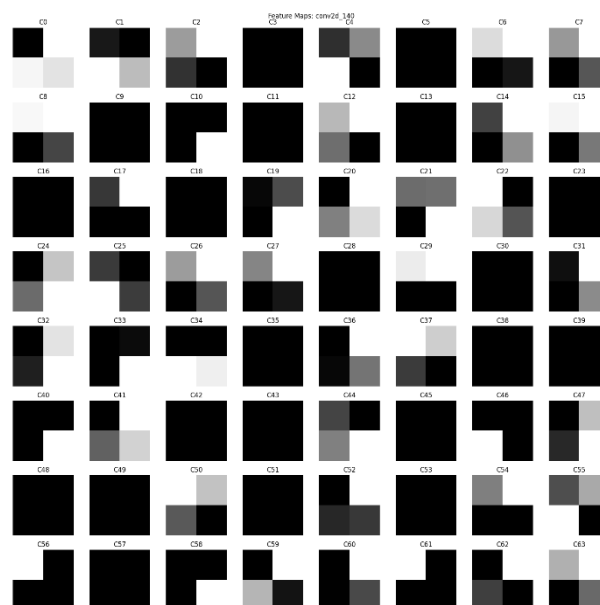
Setelah dua lapisan konvolusi awal, digunakan MaxPooling2D dengan ukuran 2x2 untuk mereduksi dimensi spasial menjadi 30x30 sambil mempertahankan kedalaman (jumlah *filter*). Pola ini diulang pada blok berikutnya dengan jumlah *filter* yang ditingkatkan menjadi 128 dan kemudian 256, memungkinkan model menangkap fitur yang lebih kompleks dan abstrak. Fitur yang ekstrak oleh CNN dapat dilihat pada Gambar 7.



Gambar 7. *Feature Maps* dari Conv2D Lapisan Keenam

Visualisasi pada Gambar 7 menunjukkan bahwa aktivasi menjadi semakin terfokus dan latar belakang hampir sepenuhnya gelap. Fitur yang tersisa pada tahap ini adalah fitur tingkat tinggi (*high-level features*) yang mewakili identitas kelas aksara secara unik, tanpa dipengaruhi oleh informasi yang tidak relevan.

Di bagian akhir *feature extraction*, digunakan lapisan Conv2D dengan 512 filter dan ukuran kernel 3x3, diikuti oleh MaxPooling2D terakhir yang mereduksi dimensi menjadi 1x1x512. Fitur yang ditangkap pada lapisan ini divisualisasikan pada Gambar 8.



Gambar 8. *Feature Maps* dari Conv2D Lapisan Keenam

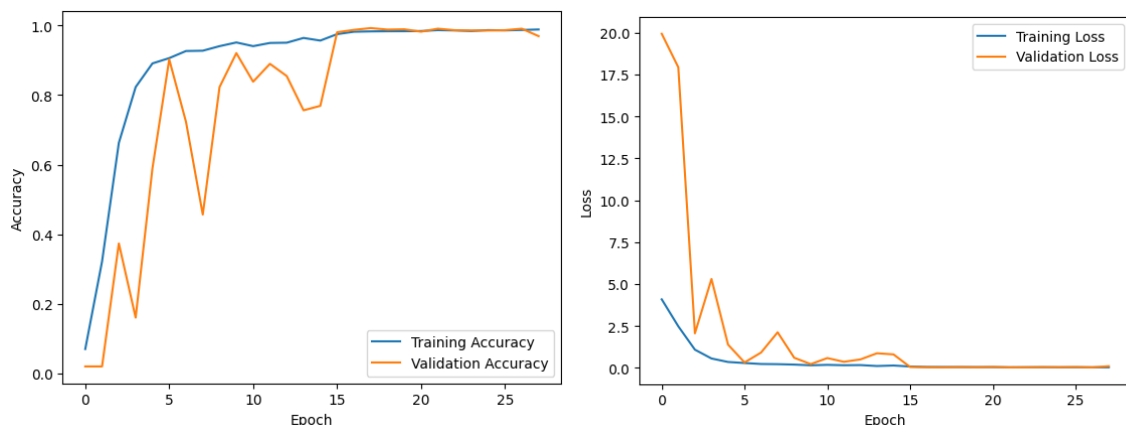
Pada Gambar 8 menampilkan ukuran *feature map* yang sangat kecil (2×2 piksel) dengan aktivasi yang sangat terpadatkan pada area tertentu. Representasi ini mengandung informasi yang sangat ringkas namun kaya akan ciri-ciri diskriminatif. Hasil dari lapisan ini menjadi masukan utama bagi lapisan Dense untuk menentukan kelas akhir pada proses klasifikasi.

Total parameter yang digunakan sebesar $\pm 3,4$ juta, arsitektur ini mampu menyeimbangkan antara kedalaman (*depth*) dan efisiensi komputasi, mengingat ukuran dataset yang relatif terbatas. Pemilihan jumlah filter dan penggunaan BatchNormalization serta Dropout menjadi bagian dari strategi untuk meningkatkan generalisasi model dan menghindari *overfitting*.

Proses pelatihan model ini dilakukan menggunakan pendekatan *supervised learning* pada data citra berlabel untuk menyelesaikan tugas klasifikasi multi-kelas. Pelatihan dijalankan dengan algoritma optimasi Adam menggunakan *learning rate* awal sebesar 0.0005 serta *loss function categorical crossentropy*, yang sesuai untuk permasalahan klasifikasi multi-kelas. Model dilatih dengan *batch size* 32 hingga maksimal 500 *epoch*, namun dibekali dengan mekanisme regulasi untuk menjaga stabilitas pelatihan dan mencegah terjadinya *overfitting*.

Regulasi dilakukan melalui dua *callback* utama, yaitu *EarlyStopping* dan *ReduceLROnPlateau*. *Callback EarlyStopping* secara aktif memantau *validation accuracy* dan akan menghentikan proses pelatihan lebih awal apabila tidak terdapat peningkatan akurasi validasi selama 10 *epoch* berturut-turut, sekaligus mengembalikan bobot terbaik yang diperoleh selama pelatihan. Sementara itu, *ReduceLROnPlateau* bertugas menurunkan *learning rate* sebesar 20% (faktor 0.2) apabila *validation loss* tidak menunjukkan perbaikan selama 5 *epoch*, dengan batas minimum *learning rate* sebesar 0.0001.

Model identifikasi aksara Minangkabau menggunakan CNN akan diukur tingkat akurasi dengan mengadopsi hasil *metrik accuracy*, baik pada data pelatihan maupun validasi, untuk memastikan proses pembelajaran berjalan secara optimal. Kombinasi strategi regulasi ini bertujuan meningkatkan efisiensi pelatihan, mempercepat konvergensi, dan menghasilkan model yang lebih generalisasi terhadap data baru. Adapun hasil grafik *Accuracy and Loss Pelatihan Model* tersaji pada Gambar 9.



Gambar 9. Accuracy and Loss Pelatihan Model

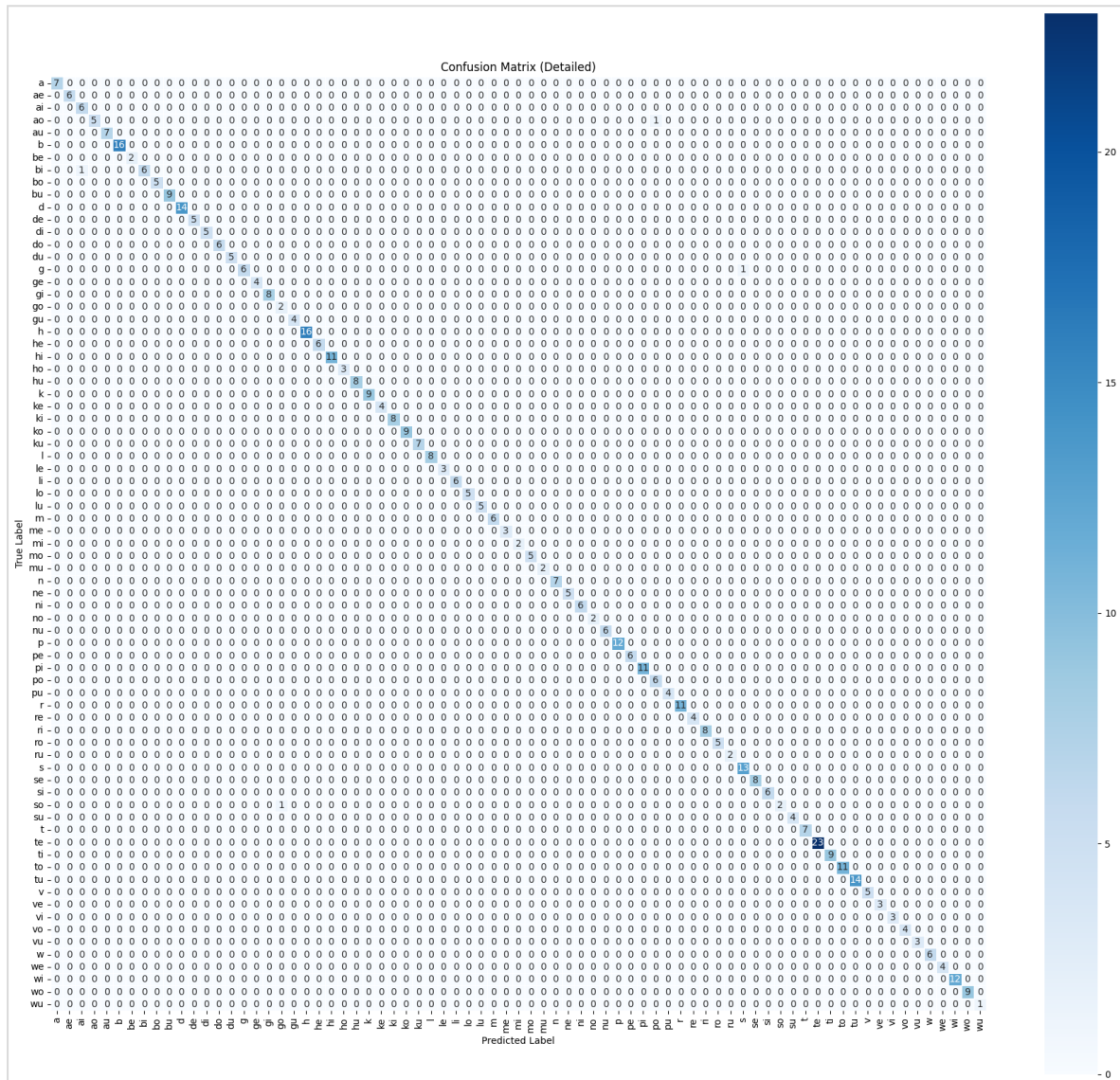
Gambar 9 merupakan hasil grafik *Accuracy and Loss* pelatihan model. Hasil akurasi pelatihan model berhenti pada *epoch* ke-28 karena *EarlyStopping*, ketika akurasi validasi tidak membaik selama 10 *epoch* berturut-turut. Pada grafik *Accuracy and Loss*, terlihat performa model meningkat tajam pada awal pelatihan, lalu stabil. Akurasi pelatihan terus naik mendekati 100%, sementara akurasi validasi sempat berfluktuasi namun stabil di atas 90% setelah *epoch* ke-15, menunjukkan kemampuan generalisasi model yang baik. Kurva *loss* menunjukkan penurunan konsisten pada data pelatihan dan validasi hingga mendekati nol, menandakan tidak terjadi *overfitting* yang berarti. Strategi *EarlyStopping* dan *ReduceLROnPlateau* terbukti efektif menjaga kualitas pelatihan, menghasilkan model klasifikasi yang stabil dan berkinerja tinggi.

3.2. Evaluasi Kinerja Model

Evaluasi performa model dilakukan terhadap data uji yang sebelumnya tidak digunakan selama proses pelatihan. Evaluasi ini bertujuan untuk menilai kemampuan generalisasi model dalam mengklasifikasikan data baru secara objektif dan menyeluruh. Teknik yang digunakan yaitu *Confusion Matrix* dan *Classification Report* (*Precision*, *Recall*, *F1-Score*).

Confusion Matrix memberikan gambaran mengenai distribusi hasil prediksi model pada masing-masing kelas, sekaligus memungkinkan identifikasi terhadap kelas-kelas yang masih sering mengalami kesalahan klasifikasi.

Visualisasi *Confusion Matrix* digunakan untuk merepresentasikan performa model dalam mengklasifikasikan data uji. Hasil dari perhitungan *Confusion Matrix* terhadap model CNN identifikasi aksara Minangkabau divisualisasikan pada Gambar 10.



Gambar 10. Accuracy and Loss Pelatihan Model

Visualisasi *Confusion Matrix* pada Gambar 10 menunjukkan performa klasifikasi yang sangat baik pada data uji, ditandai oleh dominasi nilai pada diagonal utama matriks yang mencerminkan prediksi sesuai label sebenarnya pada mayoritas sampel. Sebagian besar dari 75 kelas memiliki akurasi prediksi mendekati sempurna, meskipun masih terdapat beberapa sel di luar diagonal yang bernilai lebih dari nol sebagai indikasi *misclassification*. Kesalahan ini kemungkinan disebabkan oleh kemiripan visual antar kelas, terutama kelas dengan bentuk atau tanda baca yang hampir serupa, maupun keterbatasan variasi sampel pada subset pengujian. Tidak terlihat adanya pola kesalahan sistematis, karena *misclassifications* tidak terkonsentrasi pada kelas tertentu. Hal ini menunjukkan bahwa model memiliki kemampuan generalisasi yang baik dan merata terhadap seluruh kelas.

Selain memeriksa performa model melalui *confusion matrix*, evaluasi dapat diperluas menggunakan *classification report* guna melihat metrik seperti *precision*, *recall*, dan *f1-score*. *Classification report* memberikan gambaran menyeluruh mengenai performa model dalam mendeteksi masing-masing kelas. Proses *Classification report* mampu membantu dalam analisa performa model CNN dalam identifikasi. Metrik evaluasi model secara lebih rinci, meliputi nilai *precision*, *recall*, *f1-score*, dan *support* untuk setiap kelas tersaji pada Tabel 3.

Tabel 3. Classification Report

Class	Precision	Recall	F1-Score	Support
a	1,00	1,00	1,00	7
ae	1,00	1,00	1,00	6
ai	0,86	1,00	0,92	6
ao	1,00	0,83	0,91	6
au	1,00	1,00	1,00	7
b	1,00	1,00	1,00	16
be	1,00	1,00	1,00	2
bi	1,00	0,86	0,92	7
bo	1,00	1,00	1,00	5
bu	1,00	1,00	1,00	9
...	...	...	...	...
wi	1,00	1,00	1,00	12
wo	1,00	1,00	1,00	9
wu	1,00	1,00	1,00	1
accuracy			0,99	500
macro avg	0,99	0,99	0,99	500
weighted avg	0,99	0,99	0,99	500

Tabel 3 menunjukkan bahwa model memiliki performa klasifikasi yang sangat tinggi dengan akurasi keseluruhan 0,99, serta nilai *precision*, *recall*, dan *F1-score* rata-rata (baik *macro* maupun *weighted*) juga sebesar 0,99. Sebagian besar kelas mencapai nilai sempurna (1.00), menandakan model mampu mengenali pola dengan sangat baik. Namun, terdapat beberapa kelas dengan performa lebih rendah, seperti so (*F1-score*: 0.80), go (0.80), dan ao (0.91), yang umumnya disebabkan oleh kemiripan visual antar kelas dan jumlah sampel yang relatif kecil. Secara umum, model sangat efektif, namun beberapa kelas masih dapat ditingkatkan dengan pelatihan tambahan atau data augmentation.

4. Kesimpulan

Berdasarkan hasil eksperimen dan analisis yang dilakukan dapat disimpulkan bahwa sistem identifikasi citra huruf aksara Minangkabau dapat mengidentifikasi aksara dengan baik. Arsitektur yang diusulkan menunjukkan performa pelatihan yang stabil dan mampu menggeneralisasi data dengan baik. Evaluasi terhadap kinerja model menggunakan *Confusion Matrix* dan *Classification Report* (precision, recall, f1-score) menunjukkan hasil akurasi keseluruhan mencapai 99%. Hasil evaluasi menunjukkan bahwa sistem identifikasi aksara Minangkabau memiliki potensi untuk menjadi solusi awal dalam upaya digitalisasi aksara Minangkabau. Penelitian selanjutnya diharapkan mengembangkan model CNN yang lebih kompleks guna menghasilkan identifikasi lebih baik diwaktu yang akan datang.

Daftar Rujukan

- [1] T. Fannia, K. Johana, F. Awliya, F. Rismawaty, M. Ameri, and S. Lianto Lau, "Intercultural Communication Interaction of Multicultural Society in West Kalimantan Province: Ethnographic Studies," *KnE Social Sciences*, vol. 2023, pp. 632–644, 2023, doi: 10.18502/kss.v8i12.13711.
- [2] M. Mahmudah and T. Noor, "Internalisasi Nilai Multikultural pada Pola Asuh Anak Urang Banjar (Studi Etnografi Di Kabupaten Tanah Bumbu Dan Batola)," *Al-Madrasah: Jurnal Pendidikan Madrasah Ibtidaiyah*, vol. 7, no. 1, p. 443, 2023, doi: 10.35931/am.v7i1.1927.
- [3] D. A. Darwis and N. Muslim, "Minangkabau Cultural Identity: History And Development," *International Journal of Religion*, vol. 5, no. 10, pp. 794–805, Jun. 2024, doi: 10.61707/fbvrnv21.
- [4] R. Sovia, S. Defit, and Yuhandri, "Development of the Minangkabau Local Language Translation Machine Based on Stemming," *Proceeding - 2022 International Symposium on Information Technology and Digital Innovation: Technology Innovation During Pandemic, ISITDI 2022*, pp. 195–198, 2022, doi: 10.1109/ISITDI55734.2022.9944457.
- [5] Rini Andriani, S. F. Rahmah, and D. A. K. Putra, "Sistem Fonem Vokal dalam Bahasa Minang," *Wacana : Jurnal Bahasa, Seni, dan Pengajaran*, vol. 7, no. 2, pp. 112–120, Oct. 2023, doi: 10.29407/jbsp.v7i2.20430.

-
- [6] R. Rahim, “Etnomatematika Melalui Simbol dan Komunikasi Nonverbal dalam Tradisi Marosok di Minangkabau,” *Jurnal Ilmiah Universitas Batanghari Jambi*, vol. 23, no. 2, p. 1571, 2023, doi: 10.33087/jiubj.v23i2.3179.
- [7] K. A. Overmann, “Writing system transmission and change: A neurofunctional perspective,” Jun. 04, 2023. doi: 10.31235/osf.io/esmkw.
- [8] A. Hermilah, “Urang Minang harus tahu! Aksara Minangkabau dikenal dengan aksara Nusantara yang jarang diketahui,” *Harian Haluan*, 2023.
- [9] H. Tanjung, “Aksara Minangkabau: Minangkabau punya aksara?,” *TopSumbar.co.id*, 2023.
- [10] A. Majeed and H. Hassani, “Ancient but Digitized: Developing Handwritten Optical Character Recognition for East Syriac Script Through Creating KHAMIS Dataset,” pp. 1–17, 2024.
- [11] G. Kothari, B. Jatav, O. Bhimani, and Prof. S. Bangal, “Exploring OCR for Historical Document Preservation (Indus Script),” *INTERANTIONAL JOURNAL OF SCIENTIFIC RESEARCH IN ENGINEERING AND MANAGEMENT*, vol. 07, no. 09, pp. 1–49, Sep. 2023, doi: 10.55041/IJSREM25807.
- [12] N. Made and L. Parwati, “Penguatan Eksistensi Aksara Bali Melalui Digitalisasi Aksara Bali Bagi Generasi Muda,” pp. 65–72, 2024.
- [13] S. Schmidgall, R. Ziaei, J. Achterberg, L. Kirsch, S. P. Hajiseyedrazi, and J. Eshraghian, “Brain-inspired learning in artificial neural networks: A review,” *APL Machine Learning*, vol. 2, no. 2, pp. 1–13, 2024, doi: 10.1063/5.0186054.
- [14] M. Thorat, S. Pandit, and S. Balote, “Artificial Neural Network: A brief study,” *Asian Journal of Convergence in Technology*, vol. 8, no. 3, pp. 12–16, 2022, doi: 10.33130/ajct.2022v08i03.003.
- [15] D. Bisant, “An Introduction to Neural Networks,” *Proceedings of the International Florida Artificial Intelligence Research Society Conference, FLAIRS*, vol. 36, p. 2009, 2023, doi: 10.32473/flairs.36.134019.
- [16] F. Ramadhany Budyman and E. Rahmah, “Implementasi Stock Opname Berbasis Teknologi di UPA Perpustakaan ISI Padang Panjang: Potensi Penerapan AI dalam Sistem Otomasi Perpustakaan,” *Jurnal Pustaka AI (Pusat Akses Kajian Teknologi Artificial Intelligence)*, vol. 5, no. 1, pp. 113–118, Apr. 2025, doi: 10.55382/jurnalpustakaai.v5i1.1113.
- [17] E. Ramadani and D. Desriyeni, “Dampak Artificial Intelligence (AI) terhadap Proses Pembelajaran Mahasiswa Program Studi Perpustakaan dan Ilmu Informasi Universitas Negeri Padang,” *Jurnal Pustaka AI (Pusat Akses Kajian Teknologi Artificial Intelligence)*, vol. 5, no. 1, pp. 89–93, May 2025, doi: 10.55382/jurnalpustakaai.v5i1.950.
- [18] E. S. Agung, A. P. Rifai, and T. Wijayanto, “Image-based facial emotion recognition using convolutional neural network on emognition dataset,” 2024. doi: 10.1038/s41598-024-65276-x.
- [19] W. H. Lim, M. B. Bonab, and K. H. Chua, “An Aggressively Pruned CNN Model With Visual Attention for Near Real-Time Wood Defects Detection on Embedded Processors,” *IEEE Access*, vol. 11, pp. 36834–36848, 2023, doi: 10.1109/ACCESS.2023.3266737.
- [20] L. A. Zavala-Mondragón, P. H. N. de With, and F. van der Sommen, “A signal processing interpretation of noise-reduction convolutional neural networks,” vol. 14, no. 8, 2023.
- [21] A. A. Handoko, M. A. Rosid, and U. Indahyanti, “Implementasi Convolutional Neural Network (CNN) Untuk Pengenalan Tulisan Tangan Aksara Bima,” *Smatika Jurnal*, vol. 14, no. 01, pp. 96–110, 2024, doi: 10.32664/smatika.v14i01.1196.
- [22] E. Alfariza, D. A. C. Tue, A. S. Anas, M. Tajuddin, and A. Adil, “Klasifikasi Aksara Sasak Menggunakan Convolutional Neural Networks (CNN),” *JTIM : Jurnal Teknologi Informasi dan Multimedia*, vol. 6, no. 3, pp. 346–353, 2024, doi: 10.35746/jtim.v6i3.623.
-

- [23] N. Sazqiah *et al.*, “Pengenalan Aksara Lampung Menggunakan Metode CNN (Convolutional Neural Network),” *Seminar Nasional Insinyur Profesional (SNIP)*, vol. 2, no. 1, pp. 1–5, 2022, doi: 10.23960/snip.v2i1.165.
- [24] M. Zhang, W. Deng, and X. Li, “A Unified Image Preprocessing Framework For Image Compression,” 2022.
- [25] A. Ramadhanu, H. Hendri, Mardison, L. Navia Rani, S. Enggari, and M. Reza Putra, “THREE LAYER MEDIAN FILTER METHOD FOR IDENTIFYING CONCRETE STRENGTH LEVELS BASED ON CONCRETE IMAGES,” *International Journal of Software Engineering and Computer Systems*, vol. 10, no. 2, pp. 159–172, Feb. 2025, doi: 10.15282/ijsecs.10.2.2024.13.0131.
- [26] G. Latif, J. Alghazo, M. A. Khan, G. Ben Brahim, K. Fawagreh, and N. Mohammad, “Deep convolutional neural network (CNN) model optimization techniques—Review for medical imaging,” *AIMS Mathematics*, vol. 9, no. 8, pp. 20539–20571, 2024, doi: 10.3934/math.2024998.
- [27] M. Ghayoumi, “Enhancing Efficiency and Regularization in Convolutional Neural Networks: Strategies for Optimized Dropout,” *AI (Switzerland)*, vol. 6, no. 6, 2025, doi: 10.3390/ai6060111.
- [28] O. Rainio, J. Teuho, and R. Klén, “Evaluation metrics and statistical tests for machine learning,” *Scientific Reports*, vol. 14, no. 1, pp. 1–14, 2024, doi: 10.1038/s41598-024-56706-x.
- [29] A. K. Saha, M. Rabbani, A. S. I. Sum, M. F. Mridha, and M. M. Kabir, “An enhanced deep learning model for accurate classification of ovarian cancer from histopathological images,” *Scientific Reports*, vol. 15, no. 1, pp. 1–16, 2025, doi: 10.1038/s41598-025-07903-9.
- [30] M. Jawarneh, A. Marwanto, D. Syamsuar, and M. Kusnandar, “A Comparative Study of Convolutional Neural Networks and Vision Transformers for Fruit Classification,” *International Journal of Advances in Artificial Intelligence and Machine Learning*, vol. 2, no. 2, pp. 104–112, 2025, doi: 10.58723/ijaauml.v2i2.435.
