

JOURNAL OF INFORMATION
SYSTEM AND TECHNOLOGY
MANAGEMENT (JISTM)www.jistm.comPENGECAMAN AKSARA TULISAN TANGAN MENGGUNAKAN
PEMBELAJARAN MENDALAM

HAND WRITING CHARACTER RECOGNITION WITH DEEP LEARNING

Nur Fauzuna Hannan Ahmad¹, Afzan Adam^{1*}, Mohammad Faizul Nasrudin¹¹ Center for Artificial Intelligence Technology, Fakulti Teknologi dan Sains Maklumat, Universiti Kebangsaan Malaysia, Malaysia

Email: a180091@siswa.ukm.edu.my; afzan@ukm.edu.my; mfn@ukm.edu.my

* Corresponding Author

Article Info:

Article history:

Received date: 10.09.2023

Revised date: 28.09.2023

Accepted date: 23.10.2023

Published date: 05.12.2023

To cite this document:

Ahmad, N. F. H., Adam, A., & Nasrudin, M. F. (2023). Pengecaman Aksara Tulisan Tangan Menggunakan Pembelajaran Mendalam. *Journal of Information System and Technology Management*, 8 (33), 31-41.

DOI: 10.35631/JISTM.833003

This work is licensed under [CC BY 4.0](https://creativecommons.org/licenses/by/4.0/)

Abstract:

Pengecaman aksara, merupakan komponen utama Pengecaman Aksara Optik (OCR) yang melibatkan kenalpastian dan penukaran teks daripada imej kepada format yang boleh dibaca komputer. Walau bagaimana pun, sukar untuk mengenal pasti aksara daripada rentetan tulisan tangan pada dokumen kertas, terutamanya dokumen lama, selain imej dokumen tidak cukup jelas untuk menunjukkan tulisan itu. Oleh itu, kajian ini dibuat untuk melatih algoritma pembelajaran mendalam bagi mengecam aksara rentetan tulisan tangan menggunakan set data sendiri. Set data Tugasan ini terdiri dari tentetan tulisan tangan yang diambil dari penghantaran tugasan pelajar-pelajar tahun tiga di Universiti Kebangsaan Malaysia. Selain itu, model pengecaman akan dilatih dengan set data awam: EMNIST dan First Name. Kedua-dua set data ini mempunyai aksara-aksara tidak bersambung, Cuma EMNIST mengandungi aksara individu manakala First Name mengandungi rentetan aksara tidak bersambung. Model pengecaman aksara akan dibangun dan dilatih menggunakan algoritma Rangkaian Neural Konvolusi (CNN). Eksperimen pertama adalah untuk membuktikan hipotesis mengenai prestasi model pengecaman aksara dengan latihan dan ujian menggunakan aksara tidak bersambung. Ini telah dijalankan di mana model dilatih dan diuji dengan set data EMNIST dan First Name Model yang terhasil memberikan ketepatan pengecaman yang sempurna. Eksperimen kedua pula adalah untuk menguji prestasi model yang sama, bagi mengecam aksara-aksara bersambung yang ada pada set data Tugasan. Ketepatan asal model ini pada aksara bersambung adalah sangat rendah, tetapi hasil latihan semula yang lebih kukuh dan suntikan teknik pengoptimuman, ketepatan model pengecaman ini meningkat dari 4% kepada 43%.

Keywords:

Pengecaman Aksara; Rangkaian Neural Convolusi; Aksara Bersambung; SDG4

Abstract:

Optical character recognition (OCR) is primarily used to identify text in photographs and convert it into a readable format for computers. However, recognizing handwritten characters on paper documents, especially old ones, can be challenging. This study aimed to develop deep learning algorithms using our dataset to recognize characters in handwritten text sequences. The dataset consisted of text sequences from the assignments of third-year students at the National University of Malaysia. To train the recognition model, we utilized the EMNIST and First Name common datasets, which contain both disconnected and individual characters. We employed Convolutional Neural Network (CNN) algorithms to build and train the character recognition model. In the first experiment, we tested the performance of the model by training and testing it on disconnected characters from the EMNIST and First Name datasets. The model achieved perfect accuracy in recognizing these characters. The second experiment aimed to assess the model's ability to recognize connected characters from the assignment dataset. Initially, the model had low accuracy in recognizing connected characters. However, through intensive retraining and optimization techniques, the accuracy of the character recognition model improved from 4% to 43%.

Keywords:

Character Recognition, Convolution Neural Network, Connected Characters, SDG4

Pengenalan

Pembelajaran mendalam merupakan sejenis pembelajaran mesin berdasarkan rangkaian neural buatan di mana mempunyai pelbagai lapisan pemprosesan yang digunakan untuk mengekstrak ciri yang lebih terperinci secara progresif (Alaslani, M. G., 2018). Algoritma pembelajaran mendalam dapat mempelajari dan membuat keputusan berdasarkan corak atau tren dalam data, dan dapat menambah baik prestasi pembelajarannya dari masa ke semasa apabila terdedah kepada data yang lebih banyak dan lasak (Yap, M. I., et. al., 2021). Senibina rangkaian pembelajaran mendalam boleh menggabungkan beberapa model rangkaian. Walaupun tiada satu rangkaian dianggap sempurna, sesetengah algoritma lebih sesuai untuk melaksanakan tugas tertentu. Antaranya seperti AlexNet, ResNet dan EfficientNetB4.

Kepentingan pengecaman aksara boleh dilihat dalam beberapa penggunaan seperti keselamatan dan servis pengecaman plat kenderaan (Arshad, H., Abidin, R. Z., & Obeidy, W. K., 2017). Selain itu, model pengecaman ini juga telah cuba dibangunkan untuk tulisan jawi (Nayef, B. H. et. al., 2022), manakala Yanto, W. et. al. (2023) pula mengimplementasi pengecaman aksara untuk menulis tulisan tersebut. Kajian ini memfokuskan kepada pengecaman aksara tulisan tangan seperti huruf dan nombor. Ini kerana, terdapat beberapa dokumen yang masih diisi secara manual untuk tujuan keselamatan. Borang yang diisi dengan tangan dan mempunyai gaya tulisan yang berbeza-beza (Hamdan, Y. B., & Sathesh, A., 2021) dan boleh digunakan sebagai identiti seseorang. Dokumen tulisan tangan seperti borang, laporan, dan cek bil akan dikumpul untuk kerja pejabat. Tulisan tangan juga merupakan salah satu cara untuk menyiapkan tugas atau kerja selain daripada menaip melalui komputer.

Sesetengah laporan atau borang perlu ditulis menggunakan tulisan tangan seperti borang perjanjian untuk pinjaman pembelajaran. Oleh itu, terdapat pertambahan kerja dari segi menukar dokumen yang ditulis secara manual kepada digital bagi tujuan simpanan data sedangkan penyimpanan Salinan atau imej dokumen memakan ruang storan yang besar. Jadi, penukaran automatik sangat diperlukan untuk menukarkan rentetan tulisan tangan manual kepada digital.

Dewan Bahasa dan Pustaka Malaysia menakrifkan; rentetan tulisan tangan merujuk kepada jujukan atau kumpulan aksara yang disusun dalam satu baris panjang atau bersambungan. Proses menenalpasti aksara daripada rentetan tulisan tangan lebih sukar kerana setiap abjad, nombor ataupun simbol seperti noktah dan koma, adakalanya bersambung dan berubah bentuk. Selain itu, ada antara dokumen lama sukar dibaca dan difahami kerana kebanyakan laporan sudah lama dan dakwat pen pada tulisan semakin pudar selain imej laporan yang kurang jelas (Qureshi, R et. al., 2019). Ini akan menimbulkan kekeliruan dalam mentafsirkan maksud tulisan apabila membacanya dan mengurangkan keberkesanan carian teks atau isi kandungan sesebuah dokumen seperti borang, teks atau laporan. Oleh itu, tulisan tangan yang diperolehi perlu nyata dan algoritma yang digunakan boleh mengecam aksara tulisan tangan di samping mengenal pasti aksara yang terlibat dalam tulisan tangan dan memberikan ketepatan yang lebih bagus.

Oleh itu, kajian ini akan dilakukan untuk mencari model pengecaman rentetan aksara tulisan tangan yang terbaik. Algoritma pembelajaran mendalam berasaskan Rangkaian Neural Konvolusi akan ditala dan dilatih bagi memahami kelainan parameter dan prestasi antara pengecaman aksara individu, dan rentetan aksara. Satu lagi eksperimen tambahan untuk melihat prestasi model ini mengecam rentetan aksara bersambung turut dibuat.

Metodologi

Terdapat 3 topik utama yang akan dibentangkan dalam fasa metodologi; iaitu penyediaan data latihan dan ujian, pembangunan model pengecaman aksara individu, dan pembangunan model pengecaman rentetan aksara.

Penyediaan Data

Set data awam dalam talian yang diperolehi dari Kaggle, seperti EMNIST di (Cohen G et al, 2017) dan First Name digunakan bagi kajian ini. Set data EMNIST terdiri dari alphabet A sehingga Z, nombor (0 sehingga 9) dan aksara lain (@, #, \$, dan &); telah digunakan untuk kajian ini dan mempunyai 39 kategori untuk setiap aksara yang melibatkan abjad, nombor dan aksara istimewa. Set data First Name pula merupakan nama-nama dari permohonan kad keselamatan sosial Amerika Syarikat (<https://data.world/len/us-first-names-database>), mempunyai nama rentetan aksara tidak bersambung terdiri dari nama keluarga, dan nama pertama. Kedua-dua set data merangkumi pelbagai gaya tulisan tangan, variasi dan potensi cabaran yang akan dihadapi oleh sistem dalam scenario dunia sebenar. Jumlah keseluruhan dataset terdiri daripada 155209 aksara yang telah dimuat turun ke dalam fail projek. Sebagai tambahan, satu set data Tugas yang dikumpul sendiri dari tugas pelajar-pelajar tahun tiga Fakulti Teknologi dan Sains Maklumat, Universiti Kebangsaan Malaysia turut disediakan. Sebanyak 140 dokumen telah diimbas, menghasilkan sebanyak 16905 patah perkataan.

Kaedah analisis data merupakan salah satu cara untuk mengetahui atau mengenal bentuk sesebuah data. Data dianalisis dengan kaedah seperti penerokaan dan pembahagian data. Penerokaan data merupakan penerapan kaedah yang jelas dalam memahami ciri-cirinya seperti

bilangan jumlah sampel set data untuk setiap kelas dan kemungkinan ketidakseimbangan sampel. Pada bahagian prapemprosesan data imej, kod telah diuji untuk mendapatkan bilangan sampel yang seimbang bagi setiap kelas iaitu lebih daripada 4000 imej setiap kelas. Keseimbangan data penting bagi melatih model dengan berkesan. Ini kerana, set data yang tidak seimbang boleh membawa kepada hasil yang berat sebelah dan prestasi yang kurang baik atau teruk pada kelas yang kurang diwakili oleh set data.

Prapemprosesan imej ialah proses menyediakan imej untuk digunakan sebagai input kepada model pembelajaran mendalam. Tujuan prapemprosesan adalah untuk meningkatkan kualiti imej supaya boleh dianalisis dengan lebih berkesan. Hal ini kerana, model pembelajaran mendalam yang akan digunakan dalam projek memerlukan ke semua imej input berada dalam tata susunan dengan saiz yang sama. Selepas menukar kepada skala kelabu, imej perkataan akan ditukar kepada vektor binary. Kelas *LabelBinarizer* daripada modul *sklearn.preprocessing* digunakan untuk menukar label kategori kepada matriks vektor binari. *train_Y* dan *val_Y* diubah menggunakan *LabelBinarizer* untuk mendapatkan perwakilan label *binary*. Setiap baris dalam matriks yang sepadan dengan label dan lajur mewakili kategori yang berbeza.

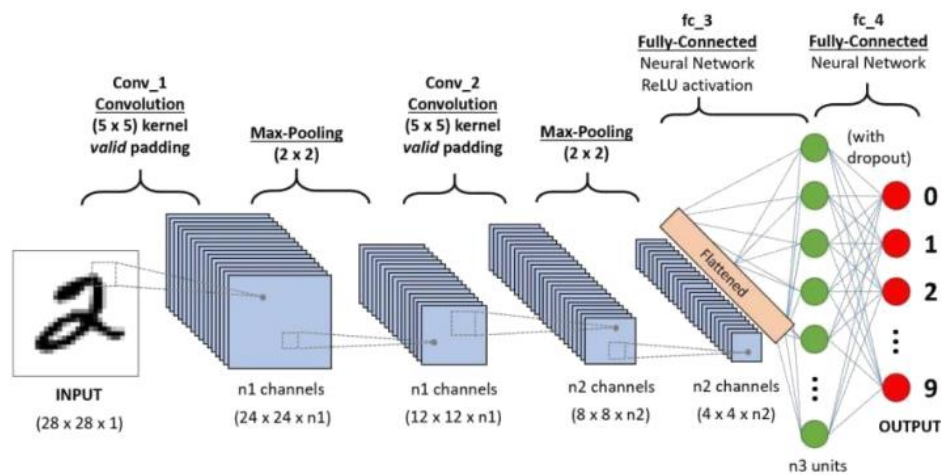
Kesemua data adalah data bersih yang telah dibahagikan dalam fail iaitu latihan, pengesahan dan ujian. Data bersih bagi kategori data latihan dan data ujian menggunakan pustaka *openCV*. *OpenCV* ialah pustaka yang sangat baik untuk pemprosesan imej dan boleh digunakan untuk memuatnaik, memproses dan memanipulasi imej. Walau bagaimanapun, set data itu sendiri tidak akan berada dalam 'format *OpenCV*.' Sebaliknya, imej dan label akan disediakan dalam salah satu format standard yang digunakan untuk menyimpan imej.

Seterusnya set data EMNIST dan First Name telah dibahagikan kepada 3 bahagian iaitu set data latihan, data pengesahan dan data ujian dengan nisbah 7:2:1. Set data Tugas digunakan sepenuhnya sebagai data ujian bagi eksperimen kedua.

Pembangunan Model Pengesanan Aksara Individu

Rangkaian Neural Konvolusi atau Convolutional Neural Network (CNN) merupakan reka bentuk rangkaian bagi pembelajaran mendalam dan belajar terus daripada data yang diperolehi. Di dalam CNN bahagian konvolusi berfungsi untuk mengenal pasti tren set data. Lapisan (layer) ini direka bentuk untuk mempelajari set filter secara automatik dan boleh melakukan lingkaran data input untuk mengestrak fitur yang akan digunakan dalam proses seterusnya. Radu, V., et al. (2018) hanya menggunakan model asas CNN yang sebenar untuk melakukan pengesanan imej fitur dengan cara yang pelbagai. Untuk lapisan hidden yang pertama ialah lapisan konvolusi yang mempunyai 30 peta fitur, setiap satu kernel bersaiz 55 piksel dan menggunakan fungsi pengaktifan ReLu. Khandokar et al. (2021) pula menggunakan CNN pada bahagian konvolusi di mana memproses satu nombor matriks yang dikenali sebagai kernel kemudian submatriks imej akan memilih nilai piksel untuk didarabkan dengan nilai piksel yang sepadan daripada kernel dan kemudiannya menghasilkan jumlah yang perlu dirumuskan bagi satu nilai piksel untuk imej filter. Khandokar et al. (2021) menggunakan CNN bagi mengenal pasti aksara daripada set data imej dan mendapatkan ketepatan untuk pengesanan menggunakan data ujian dan latihan. Selain mereka, Guha R, et. al (2020); Bohra, N., & Bhatnagar, V. (2021), and Kaladgi, Z et. al (2022) turut menggunakan CNN pada dataset EMNIST.

Selepas bahagian konvolusi selesai, bahagian pooling digunakan untuk *down-sample* peta fitur yang dihasilkan oleh lapisan konvolusi. Objektif utama pooling adalah untuk mengurangkan dimensi ruang dalam peta fitur disamping mengekalkan maklumat penting. Terdapat beberapa jenis lapisan pooling yang berbeza, tetapi yang paling sering digunakan ialah *max-pooling* dan *average pooling*. *Max-pooling* memilih nilai maksimum daripada setiap tetingkap pooling, manakala *average pooling* mengambil nilai purata. Untuk mengetahui keberkesanan CNN pada pengecaman imej menggunakan tulisan tangan Inggeris, Jinfeng Bai et al. (2021) menggunakan CNN model di mana struktur telah ditetapkan dengan lapisan *max-pooling* berada di atas kawasan pertindanan dengan saiz 3x3 piksel, dan pemberat lapisan konvolusi yang dikongsi dengan 64 peta dengan tapisan 9x9 piksel.



Rajah 1. Senibina Rangkaian CNN Yang Digunakan Bagi Pengecaman Aksara. Pelbagai Fitur Akan Disari Dari Imej Input Dalam 2 Set Lapisan Konvolusi Dan Max Pooling Sebelum Dimasukkan Kedalam Rangkaian FCN Bagi Proses Pengecaman Aksara.

Akhir sekali ialah, bahagian lapisan fully-connected (FCN), merujuk kepada rangkaian neural di mana setiap neuron menggunakan transformasi linear kepada vektor input melalui matriks pemberat. Kesannya, kemungkinan sambungan antara lapisan muncul dan ini bermakna setiap input dalam input vector mempengaruhi setiap output dalam vektor output. Modhej, N., et al. 2020 menggunakan CNN tapi CNN tersebut menggunakan AlexNet dalam rangkaian neural dalam berbilang lapisan lanjutan dengan mengumpul fitur nilai data mentah (raw data) dengan menetapkan 3 terakhir adalah lapisan rangkaian penuh. Rajah 1 menunjukkan senibina piawai CNN yang digunakan bagi pengecaman aksara ini.

Set data aksara tulisan tangan perlu menjalani latihan untuk membolehkan model pembelajaran mesin mengenali dan memahami aksara tulisan tangan dengan berkesan. Proses latihan melibatkan memasukkan set data ke dalam algoritma pembelajaran mesin, seperti rangkaian neural, dan melaraskan parameter model secara berulang untuk mempelajari dan menyamaratakan pola daripada data.

Secara keseluruhan, reka bentuk model berfungsi sebagai tulang belakang kod pengecaman aksara, bertanggungjawab untuk mengekstrak ciri bermaklumat, corak pembelajaran, dan membuat ramalan yang tepat untuk aksara tulisan tangan. Pilihan reka bentuk yang berbeza, seperti bilangan dan jenis lapisan, corak ketersambungan atau pengubahsuaian reka bentuk,

boleh memberi kesan ketara kepada prestasi model dalam tugas pengecaman aksara tulisan tangan. Dalam projek ini, akan menggunakan pengestrak fitur. Reka bentuk model direka untuk mengekstrak ciri yang boleh digunakan daripada input seperti aksara tulisan tangan. Fitur ini mungkin termasuk corak *stroke*, kelengkungan, orientasi garisan atau fitur tersendiri lain yang boleh membantu membezakan antara aksara yang berbeza. Reka bentuk biasanya terdiri daripada berbilang lapisan neuron, seperti lapisan konvolusi, lapisan *pooling* dan lapisan bersambung sepenuhnya, yang mengekstrak perwakilan hierarki data input. Selain itu, pengesanan kotak sempadan, pengecaman dan pengelasan corak turut diaplikasikan dalam kod. Contoh kod pengesanan kotak sempadan aksara secara automatik boleh dilihat dalam Rajah 2.

```
def sort_contours(cnts, method="left-to-right"):
    reverse = False
    i = 0
    if method == "right-to-left" or method == "bottom-to-top":
        reverse = True
    if method == "top-to-bottom" or method == "bottom-to-top":
        i = 1
    boundingBoxes = [cv2.boundingRect(c) for c in cnts]
    (cnts, boundingBoxes) = zip(*sorted(zip(cnts, boundingBoxes),
    key=lambda b:b[1][i], reverse=reverse))
    # return the list of sorted contours and bounding boxes
    return (cnts, boundingBoxes)
```

Rajah 2. Contoh Kod Aturcara Bagi Pengesanan Kotak Sempadan Aksara.

Pembangunan Model Pengecaman Rentetan Aksara

Model pengecaman aksara yang telah dilatih pada eksperimen pertama, digunakan pula pada data ujian First Name. Ini dilakukan untuk melihat jika model yang dilatih dengan aksara-aksara individu, boleh mengecam aksara sama tetapi dalam bentuk rentetan. Seterusnya, model juga diuji dengan data dari set Tugas bagi melihat prestasi pengecaman.

Penalaan Parameter Dan Fungsi Pengoptimuman Adam

Dalam proses penilaian, ujian pengesanan akan dilakukan untuk mengenal pasti *overfitting*, pemilihan model dan hiperparameter. Hiperparameter atau parameter merupakan parameter model pembelajaran mesin yang tidak dipelajari daripada data tetapi ditetapkan untuk pembangunan model. Dalam projek ini, parameter iaitu "*img_size*" boleh dianggap sebagai hiperparameter kerana ia menentukan saiz imej yang diubah saiznya. Model pembelajaran mendalam diketahui terdedah kepada *overfitting*, yang berlaku apabila model berprestasi baik pada data latihan tetapi kurang pada data baharu yang tidak kelihatan dan tidak digunakan dalam data ujian. Dengan menilai prestasi model pada set data pengesanan, dapat mengenal pasti sama ada model mempunyai *overfitting* atau pun tidak dan mengambil langkah untuk menangani isu tersebut. Penilaian model merupakan penilaian dimana model yang digunakan sesuai atau tidak dengan set data yang ada dalam projek.

Pengelas dilatih pada set data aksara berlabel dalam data latihan untuk mempelajari hubungan antara ciri *contour* dan kelas aksara yang sepadan. Dengan itu, setiap satu perkataan dapat diestrak dan dicam bermula dari kiri sehingga kanan aksara. Fungsi *get_letters* bertanggungjawab bagi memperoleh semula senarai abjad manakala fungsi *get_word* bertanggungjawab dalam mendapatkan dan mengenalpasti setiap satu perkataan yang berbeza-beza.

Selain itu, algoritma pengoptimuman Adam turut disuntik kedalam model pengecaman ini. Algoritma Adam merupakan salah satu teknik pengomtimuman adaptif yang semakin terkenal kebelakangan ini. Ia boleh diuji dengan mudah kerana tidak perlu menala nilai kadar pembelajaran atau juga dikenali sebagai *learning rate* sesuatu model. Ia mengemaskini nilai pemberat rangkaian berulang kali berdasarkan data latihan.

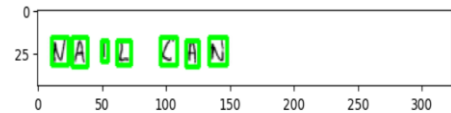
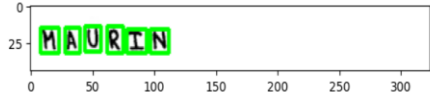
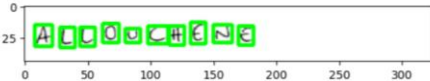
Keputusan dan Perbincangan

Penilaian yang digunakan bagi kajian ini adalah seperti hasil ketepatan dan kehilangan model semasa fasa latihan, pengesanan dan ujian. Hasil dapatan latihan model boleh dijangkakan melalui sejauh mana keberkesanan model capai. Dalam kajian ini, ketepatan model adalah metrik penilaian kritikal. Keberkesanan model boleh dijangkakan melalui keputusan ketepatan dan kehilangan untuk set data latihan dan data pengesanan. Ketepatan merupakan metric yang mengukur prestasi model pada set data tertentu. Ketepatan latihan yang tinggi menunjukkan bahawa model CNN membuat ramalan yang betul dan tepat dengan menggunakan data latihan.

Hasil Eksperimen Pembangunan Model Pengecaman Aksara Individu

Pada fasa ujian, set data ujian telah digunakan untuk mengetahui keberkesanan model dan kaedah yang digunakan. Jadual dibawah menunjukkan hasil projek melalui fasa ujian. Model berhasil mengecam tulisan yang mempunyai jarak diantara aksara dan output akan mengeluarkan hasil yang sama seperti dalam imej input .

Jadual 1
Contoh Hasil Pengecaman Aksara Pada Dataset Ujian First Name

Aksara tulisan tangan	Hasil keputusan	Aksara yang sepadan dengan tulisan tangan	Ketepatan (%)
NAIL CAN	NAILCAN Out[39]: <matplotlib.image.AxesImage at 0x245ffb43e20>  Output: NAILCAN	Semua aksara sepadan	100
MAURIN	AAURIH Out[24]: <matplotlib.image.AxesImage at 0x245fef5b1c0>  Output: AAURIH	3 daripada 7 adalah sepadan iaitu U, R, dan I.	57
ALLOUC#ENE	ALLOICAENE Out[25]: <matplotlib.image.AxesImage at 0x245fef975b0>  Output: ALL01CAENE	7 daripada 10 adalah sepadan iaitu A, L, L, C, E, N, dan E	70

	FERAAHDEZ Out[26]: <matplotlib.image.AxesImage at 0x245ff289bd0> Output: FERAAHDEZ	6 daripada 9 adalah sepadan iaitu F, E, R, D, E, dan Z	66.67
--	---	---	--------------

Hasil ketepatan hanya untuk set data EMNIST dan First Name sahaja dengan menggunakan *epochs* 30 adalah 0.9406. Nilai *epochs* lain telah diuji dan tidak memberi kesan ketara kepada prestasi model. Oleh itu, nilai *epochs* lebih rendah dipilih bagi mengurangkan *overhead*. Jadual 1 menunjukkan hasil akhir pengecaman huruf pada set latihan First Name.

Hasil Eksperimen Pembangunan Model Pengecaman Aksara Bersambung

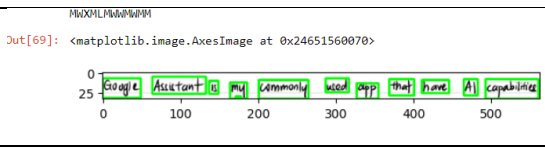
Prestasi model pengecaman yang baik tadi, tidak Berjaya dikekalkan apabila model pengecaman diuji dengan set data ujian Tugasan. Ini kerana, set data Tugasan mengandungi rentetan aksara yang bersambung-sambung dan ada kalanya berubah bentuk dari aksara individu yang digunakan dalam latihan model.

Kelemahan utama yang dikesan adalah pengesanan kotak sempadan yang lemah kerana ia hanya dibuat apabila ada jarak antara tulisan. Ini boleh dilihat dari Jadual 2 dimana bagi rentetan aksara bersambung, keseluruhan perkataan dikotakkan dan digunakan untuk pengecaman. Akibatnya, aksara-aksara yang ada dalam perkataan tersebut tidak berjaya dicam kerana corak yang amat berbeza dengan aksara yang digunakan semasa latihan model.

Jadual 2 menunjukkan hasil kualitatif dan kuantitatif setiap aksara yang telah diuji. Nilai ketepatan diambil dalam dua tempat titik perpuluhan. Nilai ketepatan menunjukkan bahawa aksara sepadan lebih tinggi terhadap aksara yang mempunyai jarak yang diantara satu sama lain iaitu dengan 100% nilai ketepatan, manakala untuk aksara yang bersambung nilai ketepatan adalah sangat rendah iaitu 3.92%.

Jadual 2
Contoh Hasil Pengecaman Aksara Pada Set Data Ujian Tugasan

Aksara tulisan tangan	Hasil keputusan	Aksara yang sepadan dengan tulisan tangan	Ketepatan (%)
	Out[41]: <matplotlib.image.AxesImage at 0x245ffc29030>	Tiada yang sepadan.	0
	Out[29]: <matplotlib.image.AxesImage at 0x245fec0fa0>	2 daripada 51 adalah sepadan iaitu T dan T.	3.92
	Out[64]: <matplotlib.image.AxesImage at 0x2465029c400>	Tiada yang sepadan.	0

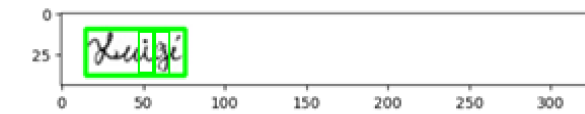
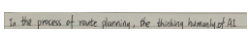
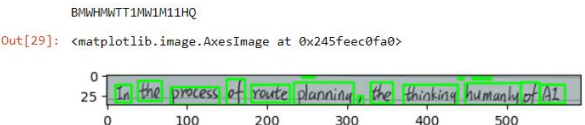
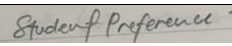
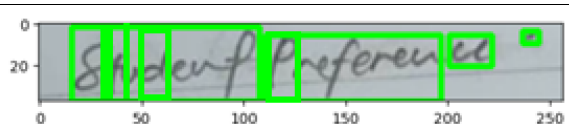
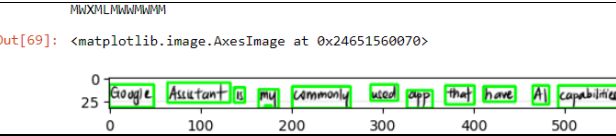
Google Assistant is my commonly used app that have AI capabilities		Tiada yang sepadan.	0
--	---	---------------------	---

Hasil Penalaan Parameter Dan Adam Booster

Hasil dari suntikan algoritma pengoptimuman Adam, dimana fungsi *Stochastic Gradient Descent* bagi pencarian kadar pembelajaran rangkaian, dilanjutkan dengan nilai yang berubah-ubah. Ini menjadikan model pengecaman ini boleh berfungsi jauh lebih baik dari model asal. Seperti yang dapat dilihat dari Jadual 3, pengesanan kotak sempadan bagi beberapa aksara bersambung boleh dilakukan, dan seterusnya proses pengecaman bagi aksara tersebut, berjaya dilakukan dengan baik.

Jadual 3

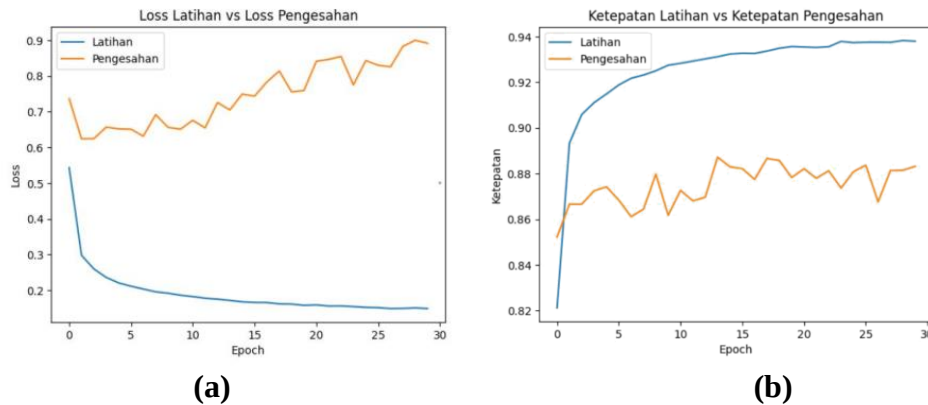
Contoh Hasil Pengecaman Aksara Selepas Pengoptimuman Adam Pada Set Data Ujian Tugasan

Aksara tulisan tangan	Hasil keputusan	Aksara yang sepadan dengan tulisan tangan	Ketepatan (%)
		i, g, i sahaja yang sepadan.	60%
		Tiada yang sepadan.	0
		S, t, d, P sahaja sepadan.	30%
Google Assistant is my commonly used app that have AI capabilities		Tiada yang sepadan.	0

Seterusnya, graf perbandingan diantara ketepatan latihan dan ketepatan pengujian ditunjukkan pada Rajah 3. Garis berwarna biru menunjukkan ketepatan latihan yang memfokuskan kepada tahap pencapaian model dalam data latihan manakala garis berwarna oren menunjukkan ketepatan pengesanan yang merujuk kepada keberkesanan pencapaian model terhadap data yang berlainan dengan data latihan. Paksi X memaparkan bilangan *epoch* iaitu bilangan masa yang diambil oleh model untuk melakukan keseluruhan set data.

Rajah 3(a) menunjukkan pola *overfitting* yang berlaku semasa latihan kerana pola kehilangan yang sangat berbeza antara data latihan dan ujian. Manakala Rajah 3(b) menunjukkan peningkatan pembelajaran model setelah algoritma pengoptimuman Adam digunakan. Pola

pembelajaran model pada data latihan dan ujian adalah sama walaupun pada kadar yang berbeza.



Rajah 3. Graf Kehilangan Dan Ketepatan Model Pengesanan Aksara. (A) Kadar Kehilangan Bagi Data Latihan EMNIST Dan Data Ujian Tugasan, Dimana Kehilangan Amat Rendah Pada Data Latihan Tetapi Semakin Meningkat Pada Data Ujian. (B) Kadar Ketepatan Pengesanan Model Setelah Suntikan Algoritma Pengoptimuman Adam Pada Data Latihan EMNIST Yang Sangat Baik, Berbanding Pada Data Ujian Tugasan Yang Lebih Rendah.

Kesimpulan

Keputusan hasil akhir mendapati bahawa model pengesanan aksara ini berpotensi untuk digunakan untuk mengesan tulisan tangan bersambung. Ini menunjukkan, CNN piawai dan kaedah contour hanya boleh berkesan dalam mengenali aksara yang terpisah dengan baik dan berbeza diantara satu sama lain. Namun begitu, CNN dan kaedah berasaskan *contour* berpotensi untuk mempelajari kerumitan dan gaya tulisan tangan dengan pengoptimuman dan penalaan parameter yang baik.

Pelbagai aplikasi penukaran aksara tulisan tangan kepada bentuk digital telah berjaya dibangunkan dan ini dimulakan dengan model pengesanan seperti ini. Penukaran aksara bersambung yang ditulis dengan tangan di atas kertas, kepada bentuk digital membolehkan ia dikongsi secara lebih meluas secara atas talian dengan penjimatan ruang storan, selain membenarkan manipulasi digital dilakukan terhadap tulisan tersebut.

Penghargaan

Penulis merakamkan setinggi-tinggi penghargaan kepada Kementerian Pengajian Tinggi Malaysia (KPT) atas sokongan mereka kepada projek ini melalui geran penyelidikan FRGS/1/2022/ICT02/UKM/02/5.

References

- Alaslani, M. G. (2018). Convolutional neural network based feature extraction for iris recognition. *International Journal of Computer Science & Information Technology (IJCSIT)* Vol, 10.
- Arshad, H., Abidin, R. Z., & Obeidy, W. K. (2017). Identification of vehicle plate number using optical character recognition: a mobile application. *Pertanika Journal of Science and Technology*, 25, 173-180.

- Balci, B., Saadati, D., & Shiferaw, D. (2017). Handwritten text recognition using deep learning. CS231n: Convolutional Neural Networks for Visual Recognition, Stanford University, Course Project Report, Spring, 752-759.
- Bohra, N., & Bhatnagar, V. (2021). Tuning of CNN Architecture by CSA for EMNIST Data. In *Advances in Information Communication Technology and Computing: Proceedings of AICTC 2019* (pp. 45-55). Springer Singapore.
- Cohen, G., Afshar, S., Tapson, J., & Van Schaik, A. (2017, May). EMNIST: Extending MNIST to handwritten letters. In *2017 international joint conference on neural networks (IJCNN)*
- Guha, R., Das, N., Kundu, M., Nasipuri, M., & Santosh, K. C. (2020). Devnet: an efficient cnn architecture for handwritten devanagari character recognition. *International Journal of Pattern Recognition and Artificial Intelligence*, 34(12), 2052009.
- Hamdan, Y. B., & Sathesh, A. (2021). Construction of statistical SVM based recognition model for handwritten character recognition. *Journal of Information Technology and Digital World*, 3(2), 92-107.
- Kaladgi, Z., Gupta, S., Jesawada, M., & Kandale, V. (2022) Handwritten Character Recognition Using CNN with Extended MNIST Dataset. *International Research Journal of Modernization in Engineering Technology and Science*, 4(5): 5577-5580.
- Khandokar, I., Hasan, M., Ernawan, F., Islam, S., & Kabir, M. N. (2021, June). Handwritten character recognition using convolutional neural network. In *Journal of Physics: Conference Series* (Vol. 1918, No. 4, p. 042152). IOP Publishing.
- Nayef, B. H., Abdullah, S. N. H. S., Sulaiman, R., & Alyasseri, Z. A. A. (2022). Optimized leaky ReLU for handwritten Arabic character recognition using convolution neural networks. *Multimedia Tools and Applications*, 1-30.
- Modhej, N., Bastanfard, A., Teshnehlab, M., & Raiesdana, S. (2020). Pattern separation network based on the hippocampus activity for handwritten recognition. *IEEE Access*, 8, 212803-212817.
- Qureshi, R., Uzair, M., Khurshid, K., & Yan, H. (2019). Hyperspectral document image processing: Applications, challenges and future prospects. *Pattern Recognition*, 90, 12-22.
- Radu, V., Tong, C., Bhattacharya, S., Lane, N. D., Mascolo, C., Marina, M. K., & Kawsar, F. (2018). Multimodal deep learning for activity and context recognition. *Proceedings of the ACM on interactive, mobile, wearable and ubiquitous technologies*, 1(4), 1-27.
- Yanto, W., Nayan, N. M. N. M., & Sulaiman, R. S. S. Optical Character Recognition (OCR) Based on Text Image Using Raspberry Pi with Arduino to Write Text. (2023) *Available at SSRN 4341207*.
- Yap, M. I., Moorthy, K., Daud, K. M., & Ernawan, F. (2021, August). Optical Character Recognition using Backpropagation Neural Network for Handwritten Digit Characters. In *2021 International Conference on Software Engineering & Computer Systems and 4th International Conference on Computational Science and Information Management (ICSECS-ICOCSIM)* (pp. 167-171). IEEE.