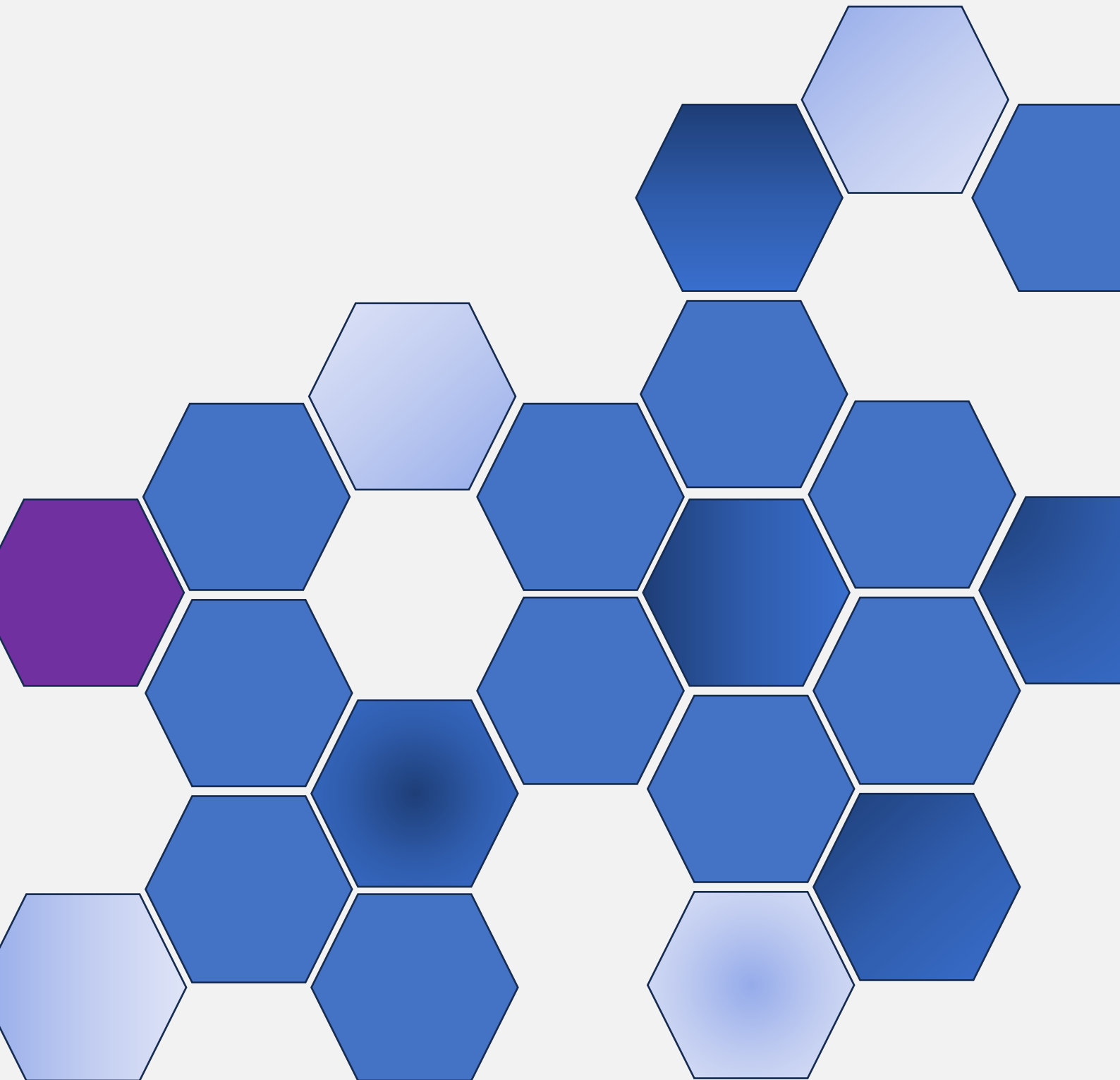




CONTENTS

CONTENTS	i
EDITORIAL BOARD	iii
PREFACE	iv
The Effect of Image Enhancement on Automatic Vehicle Detection Using Yolov8 Based on Jetson Nano Single Board Computer <i>Rakhmad Gusta Putra, Wahyu Pribadi, Dirvi Eko Juliando Sudirman</i>	221–234
Evaluating XGBoost Performance in Improving Community Security through Multi-Class Crime Prediction: Insights from the Denver Crime Dataset <i>Mc-Kelly Tamunotena Pepple</i>	235-252
Determination of Cyanide Content and Heavy Metals (Cu, Ni, Cd, & Pb) in Different Processed Cassava Meal in Abraka Metropolis, Delta State, Nigeria <i>Udezi M. E, Anwani S. E, and Ekor I. A</i>	253-264
Mini Solar Power Resources for IoT System in the Vannamei Shrimp Pond Model <i>Damar Widjaja, Yohanes Priyanto Seli Laka</i>	265-276
Improving the Accuracy of Prediction of Dissolved Oxygen and Nitrate Level Using LSTM with K-Means Clustering and Spearman Analysis <i>Ika Arva Arshella, I Wayan Mustika, Prapto Nugroho</i>	277-298
Enhancing the Cooling Effectiveness Utilizing a Tapered Fin Having Capsule-Shaped Cross-Sectional Area Numerically Simulated Using Finite Difference Method <i>Nico Ndaru Pratama, Budi Setyahandana, Doddy Purwadianto, Gilang Argya Dyaksa, Heryoga Winarbawa, Michael Seen, Rines, Stefan Mardikus, Wibowo Kusbandono, Y.B Lukiyanto, Marvin Edmin Son</i>	299-308
Numerical Solution of Two-Dimensional Advection Diffusion Equation for Multiphase Flows in Porous Media Using a Novel Meshfree Method of Lines <i>F.O. Ogunfiditimi, J.A. Kazeem</i>	309-344
Coconut Shell-Based Briquettes for Sustainable Energy: A Bibliometric Study on Biomass Mixtures and Binder Materials <i>Ridho Darmawan, Awaly Ilham Dewantoro, Efri Mardawati</i>	345-370

Fine-Tuned IndoBERT-Based Sentiment Analysis for Old Indonesian Songs Using Contextual and Generating Augmentation <i>Gilang Ramdhani, Siti Yuliyanti</i>	371-384
Design of a Speed Sensorless Control System on a DC Motor using a PID Controller <i>Bernadeta Wuri Harini, Theodore Galeno Gunadi, Shakuntala Gema Mahardika, Sirilus Praditya Pangestu, Regina Chelinia Erianda Putri, Stefan Mardikus</i>	385-402
In Vitro Regeneration of Dendrobium Through Somatic Embryogenesis from Leaf Explants <i>Rahmadiyah Hamiranti, Nanang Wahyu Prajaka, Yeni, Desi Maulida, Lisa Erfa</i>	403-416
Braille Pattern Detection Modeling Using Inception V3 Architecture Using Median Filter Implementation and Segmentation <i>Abdul Latip, Siti Yuliyanti, Muhammad Al-Husaini</i>	417-432
Sign Language Detection Models using Resnet-34 and Augmentation Techniques <i>Rizki Ramdhan Hilal, Aradea, Vega Purwayoga</i>	433-450
SVM and Ensemble Majority Voting Algorithm on Sentiment Analysis of Using ChatGPT in Education <i>Hildegardis Yayukristi Weko, Hari Suparwito</i>	451-468
AUTHOR GUIDELINES	469-470

EDITORIAL BOARD

Editor in Chief

Dr. I Made Wicaksana Ekaputra (*Sanata Dharma University, Yogyakarta, Indonesia*)

Email: made@usd.ac.id

Associate Editor

Dr. Pham Nhu Viet Ha (*Vietnam Atomic Energy Institute, Hanoi, Vietnam*)

Dr. Hendra Gunawan Harno (*Gyeongsang National University, Jinju, The Republic of Korea*)

Dr. Mukesh Jewariya (*National Physical Laboratory, New Delhi, India*)

Dr. Mongkolserj Lin (*Institute of Technology of Cambodia, Phnom Penh, Cambodia*)

Dr. Yohanes Baptista Lukiyanto (*Sanata Dharma University, Yogyakarta, Indonesia*)

Dr. Apichate Maneewong (*Thailand Institute of Nuclear Technology, Bangkok, Thailand*)

Prof. Dr. Sudi Mungkasi (*Sanata Dharma University, Yogyakarta, Indonesia*)

Dr. Pranowo (*Universitas Atma Jaya Yogyakarta, Yogyakarta, Indonesia*)

Dr. Monica Cahyaning Ratri (*Sanata Dharma University, Yogyakarta, Indonesia*)

Dr. Mahardhika Pratama (*Nanyang Technological University, Singapore*)

Prof. Dr. Leo Hari Wiryanto (*Bandung Institute of Technology, Bandung, Indonesia*)

Dr. Ranggo Tungga Dewa (*Universitas Pertahanan, Bogor, Indonesia*)

Editorial Assistant

Rosalia Arum Kumalasanti, M.T. (*Sanata Dharma University, Yogyakarta, Indonesia*)

Vittalis Ayu, M.Cs. (*Sanata Dharma University, Yogyakarta, Indonesia*)

Contact us

International Journal of Applied Sciences and Smart Technologies

Faculty of Science and Technology

Universitas Sanata Dharma

Kampus III Paingan, Maguwoharjo, Depok, Sleman

Yogyakarta, 55282

Phone : +62 274883037 ext. 523110, 52320

Fax : +62 272886529

Email : editorial.ijasst@usd.ac.id

Website : <http://e-journal.usd.ac.id/index.php/IJASST>

IJASST is an open-access peer-reviewed journal that mediates the dissemination of research and studies conducted by academicians, researchers, and practitioners in science, engineering, and technology.

PREFACE

Dear readers, we are delighted to serve you Volume 07, Issue 02 of *International Journal of Applied Sciences and Smart Technologies* (IJASST), which is managed and published by the Faculty of Science and Technology, Universitas Sanata Dharma. IJASST is an open-access peer-reviewed journal that mediates the dissemination of research and studies conducted by academicians, researchers, and practitioners in science, engineering, and technology. Its scope also includes basic sciences which relate to technology, such as applied mathematics, physics, and chemistry.

In this edition, We have fourteen papers authored by researchers from Indonesia and Africa. Submitted papers are reviewed fairly using the open journal system (OJS) of IJASST. After the review process, accepted papers of the journal are publicly available for free at the website of IJASST. For future issues, we are looking forward to your contributions to IJASST.

Dr. I Made Wicaksana Ekaputra
Editor in Chief
IJASST

The Effect of Image Enhancement on Automatic Vehicle Detection Using Yolov8 Based on Jetson Nano Single Board Computer

Rakhmad Gusta Putra^{1*}, Wahyu Pribadi¹,
Dirvi Eko Juliando Sudirman¹

¹*State Polytechnic of Madiun, Madiun, Indonesia*

^{*}*Corresponding Author: gusta@pnm.ac.id*

(Received 04-02-2025; Revised 11-03-2025; Accepted 15-03-2025)

Abstract

Vehicle counting systems using image processing and deep learning have been widely studied. Using images captured by CCTV cameras makes vehicle counting effective and efficient. Although much research has been done, there are still challenges in direct application in the field. Object detection methods such as YOLO are widely chosen. In field applications, challenges are found such as rainy, nighttime, or foggy conditions and the use of appropriate hardware. In this study, the YOLOv8s and YOLOv8n object detection methods are proposed using Contrast Limited Adaptive Histogram Equalization (CLAHE) image enhancement in preprocessing and datasets and run using SBC Jetson Nano. From this study, the results obtained an increase in detection values of around 10% to 20% in dark image conditions and there was no improvement for bright images. The average accuracy is 0.873312 for YOLOv8s and 0.866906 for YOLOv8n with image enhancement. And the processing time on Jetson Nano is 59.5 ms for YOLOv8n.

Keywords: Vehicle Detection, Image Enhancement, YOLO, CLAHE

1 Introduction

Vehicle counting systems using image processing and deep learning have been widely studied. By using image processing from CCTV camera captures, automatic vehicle counting can be done effectively and efficiently. By using this mechanism, the devices used in the field are relatively simple in the form of CCTV cameras, either specially installed or already installed. The camera can still be used for its basic applications, for example for manual monitoring of road conditions.

Research has been conducted on vehicle counting using the r-CNN, YOLO, Deep Sort, SSD and color based methods. [1]–[8]. This model can also be applied to other



applications such as to classify truck types based on the number of truck axles. [9]. Although much research has been done, there are still challenges in direct application in the field. In the application in the field, there will also be challenges of weather, whether rain, night or foggy. There are previous studies that use the Contrast Limited Adaptive Histogram Equalization (CLAHE) image enhancement method and variations of the use of the YOLO method. [10]. In this research, the detection results were improved but were still applied to personal computers with high specifications.

The use of embedded system devices that are suitable for image processing and deep learning is very important to ensure the system works in real time. Jetson Nano is one of the Single Board Computers (SBC) made by Nvidia that is dedicated to image processing with the most affordable price of the other Jetson series. With the lowest series, the implementation of deep learning used is also only capable of using lightweight methods. There is a study conducted using Jetson Nano for vehicle counting using MobileNet SSD v2 [11]. Image-based vehicle detection and tracking in varying weather conditions is also being researched [12], but has not discussed related to direct application. In this study, an object detection method is proposed with YOLOv8s and YOLOv8n variations using Contrast Limited Adaptive Histogram Equalization (CLAHE) image enhancement in preprocessing and dataset and run using SBC Jetson Nano. It is expected to improve vehicle detection results, especially during night conditions with uneven lighting and can be applied in the field in real time.

2 Material and Methods

In the proposed system there are two stages, namely training dataset and vehicle detection using YOLO. The object detection used is YOLOv8s and YOLOv8n. The dataset used consists of original images from CCTV cameras and datasets from original images that are augmented using the image enhancement model CLAHE. The dataset was obtained from collecting images from several CCTVs spread across Madiun, Semarang and Surakarta in Indonesia. Two types of datasets produce two types of datasets. In conducting the vehicle detection experiment, two variations were also used, namely the original image and the image that was preprocessed using the image enhancement CLAHE. The detection results from the dataset variations, YOLO models

and image preprocessing will then be analyzed further. The research process flow is shown in Fig. 1.

The process flow is divided into two stages, the first is the training stage and the detection stage. In the training stage, the dataset consists of normal images and images enhanced with CLAHE. Each image is trained to produce model 1 with the original image dataset and model 2 with the original image dataset and images enhanced with CLAHE. Next is the process of testing the detection results with YOLO using original input images and images that have been enhanced with CLAHE. The output is obtained so that it can be compared and analyzed further.

2.1 Image Enhancement using CLAHE on LAB Color Space

Image enhancement is focused on improving nighttime captured images. Nighttime captured images tend to have characteristics that are dark or unevenly bright and dark. This makes objects or vehicles on the dark side difficult to recognize. The image enhancement used in this study is Contrast Limited Adaptive Histogram Equalization (CLAHE). Here CLAHE is applied to LAB Color Space [13]. The color space used in standard images and CCTV images is RGB. The first step is to convert RGB color space to LAB color space.

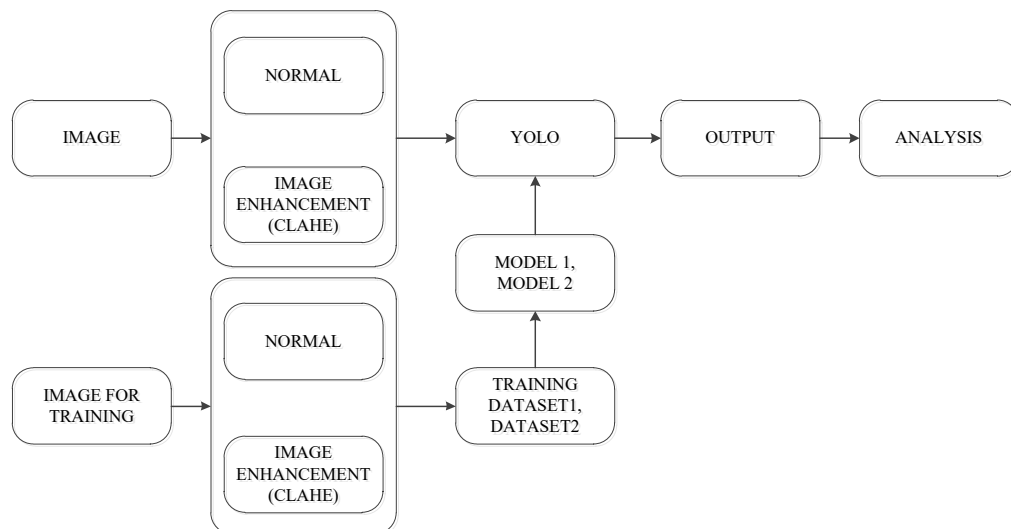


Figure 1. Research Process Flow

This is done because LAB color space consists of L (lightness), A (red/green value), B (blue/yellow value). By utilizing the L section to perform CLAHE enhancement, it is expected to correct uneven light and dark conditions in night images. After CLAHE is performed on the L section, the L, A and B sections are combined again and converted back to RGB color space. The flow of the CLAHE image enhancement process on the color image is shown in Fig. 2.

2.2 Dataset and Experiment Specification

The dataset for training the YOLO model consists of a combination. Dataset 1 is a dataset consisting of original images captured by the camera without any modification. Dataset 2 is a dataset consisting of a combination of original images captured by the camera and images captured by the camera that have been augmented with CLAHE image enhancement. The combination of datasets is shown in Table 1. The parameter settings for training the YOLOv8 model are shown in Table 2. The image size is set to 640 and Epoch 100 for each training dataset.

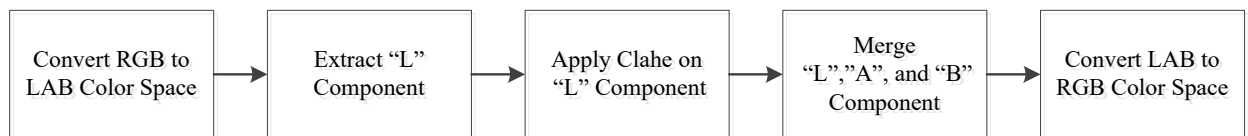


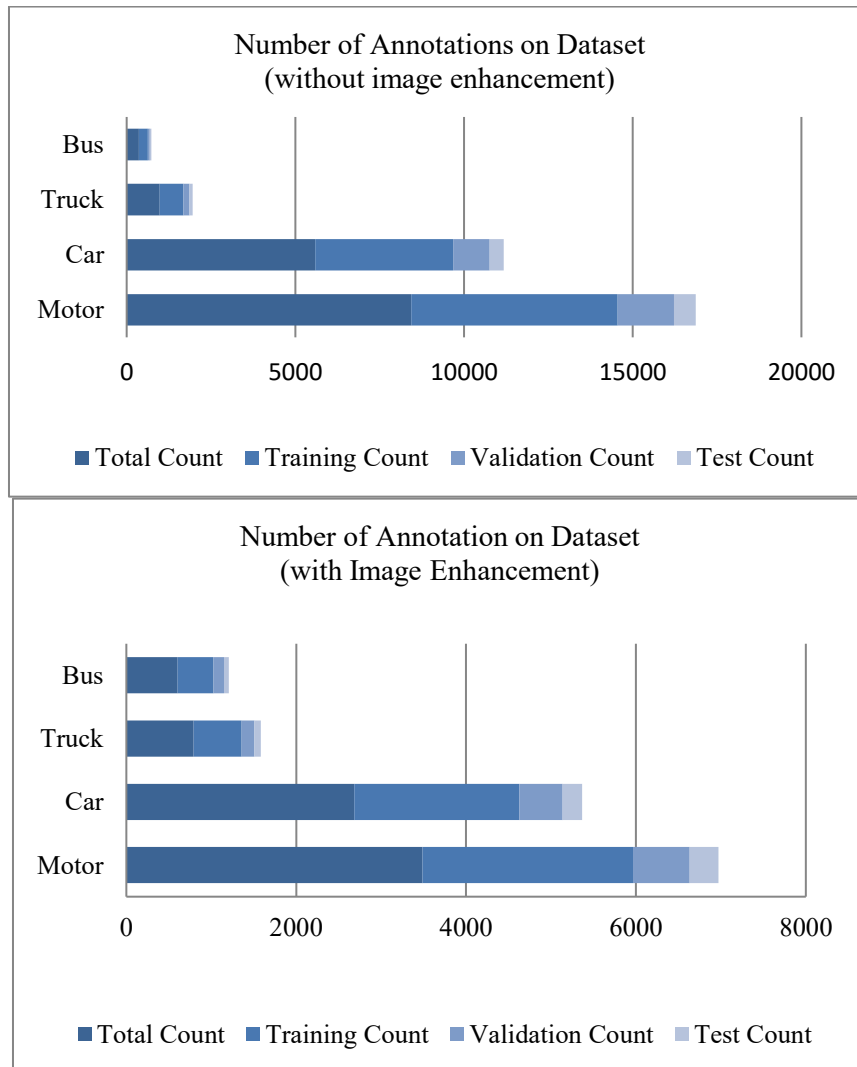
Figure 2. Color Image Enhancement Process Flow using CLAHE

Table 1. Dataset Combination Setting

Dataset	Original Image	Image Enhancement CLAHE
Dataset 1	√	
Dataset 2	√	√

Table 2. YOLOv8 Series model training parameter values

Parameter	YOLO8s
Size of Input Images	640
Epoch	100

**Figure 3.** Graphic of Number of Annotation on Dataset (with and without Image Enhancement)

The dataset itself consists of captured images from several CCTV points in Indonesia during day, night and rainy conditions. The number of datasets used for Dataset 1 and Dataset 2 is shown in Table 3 and Fig. 3. The dataset consists of four classes, namely bus, truck, car and motorbike. There is a difference in the number of datasets 1 and 2 due to the difficulty in equalizing and balancing the number of images used. Training is done with the help of Google Colab with a training time of approximately 2 hours.

The hardware setup to be tested is shown in Fig. 4. It consists of the main processor in the form of a Jetson Nano single board computer, a switch/hub with POE and IP Camera. The specifications of Jetson Nano are shown in Table 4. Jetson Nano is the most economical SBC series from Nvidia which is specifically designed for processing images and AI. Nvidia Jetson Nano itself supports the use of TensorRT. NVIDIA TensorRT is an SDK for high-performance deep learning inference. It already contains a deep learning inference optimizer and runtime that can provide processing speed for deep learning more efficiently. TensorRT-based applications have a performance 40X faster than CPU usage during inference. TensorRT is built on CUDA®, NVIDIA's parallel programming model, and allows for optimizing libraries that leverage inference, development tools, and technologies in CUDA-X™ for artificial intelligence, autonomous machines, high-performance computing, and graphics. With TensorRT, developers can focus on creating new AI-powered applications rather than tuning performance for inference applications.

Table 3. Number of Annotation on Dataset

Class Name	without Image Enhancement				With Image Enhancement			
	Total Count	Training Count	Validation Count	Test Count	Total Count	Training Count	Validation Count	Test Count
Motor	8433	6122	1665	646	3485	2480	669	336
Car	5587	4094	1080	413	2684	1944	509	231
Truck	977	700	183	94	791	566	151	74
Bus	365	261	58	46	604	420	124	60

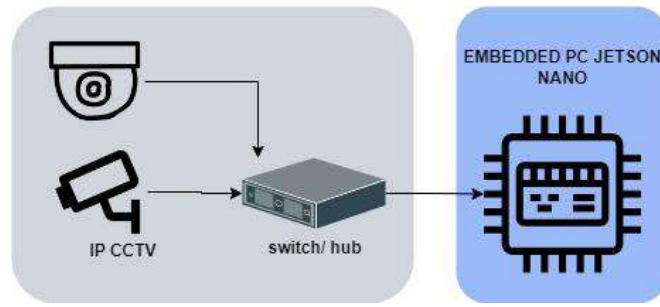


Figure 4. Hardware Setup

Table 4. Experiment Specification on Jetson Nano 4GB

Name	Specification
OS	Jetpack 4.6
CPU	Quad-core Arm A57 processor @ 1.43 GHz
GPU	128-core Maxwell GPU 4GB
RAM	System Memory – 4GB 64-bit LPDDR4 @ 25.6 GB/s

3 Results and Discussions

Night photo images have low brightness and contrast. The brightness level is sometimes uneven so that objects on the dark side become unclear and difficult to recognize. As a result, image enhancement is carried out to obtain an image with a clear image. This is done on the dataset image or as a preprocessing stage before object detection. The results of the image enhancement process and the original image are shown in Fig. 5. It can be seen that the image from the image enhancement process with CLAHE gets a brighter image and even contrast for the night image. A comparison of the histogram of the original image and the results of the process with CLAHE image enhancement is shown in Fig. 6. This histogram is a histogram of the image converted into grayscale image space to make it easier to compare. From the histogram, it can be seen that after the image enhancement process, the distribution of color brightness is more even.

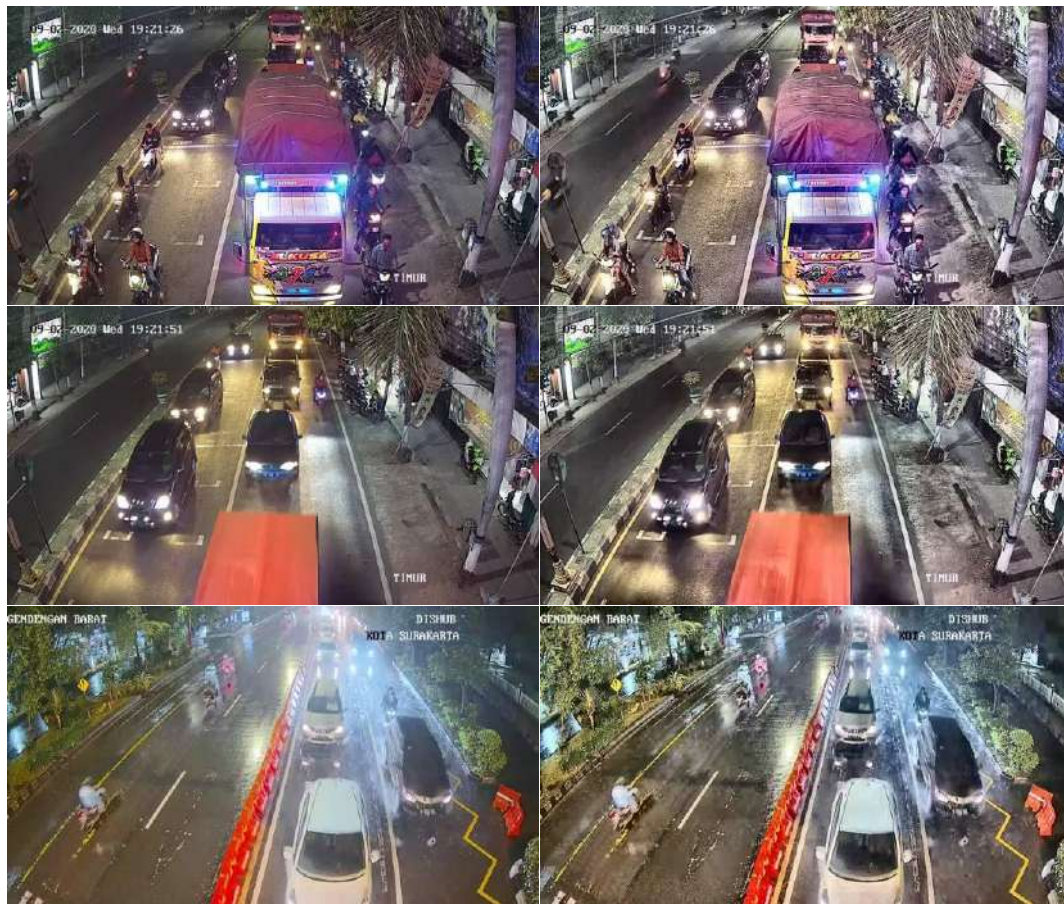


Figure 5. Original Image (Left) and Result of Image Enhancement using CLAHE (Right)

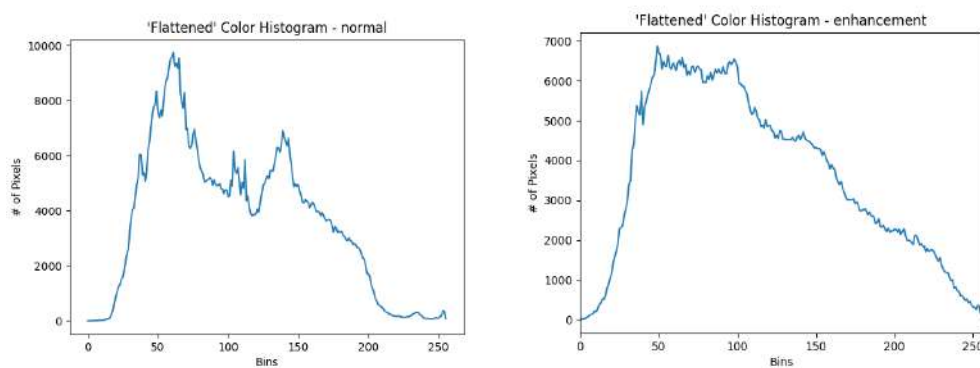


Figure 6. Histogram of Original Image (Left) and Result of Image Enhancement using CLAHE (Right)

In the dataset and testing, images captured by cameras from various regions in Indonesia were used. Training was carried out for the YOLOv8s and YOLOv8n object detection models to compare the performance that is suitable for use on the Jetson Nano. The test results for images in day and night conditions using the mode obtained an average accuracy result for the YOLOv8s model of 0.87, higher than YOLOv8n 0.86. Although YOLOv8s is superior, the difference that occurs is not too significant. This comparison can be seen in Table 5.

The results of the comparison of object detection using YOLOv8s with dataset 1 for training for the original image and the image that was image enhanced are shown in Fig. 7. The image used is in night conditions. There appears to be an improvement in the detection rate for several vehicles. After image enhancement, there are also vehicles that were previously undetected that are detected.

In the experiment using the YOLOv8s model and training using dataset 2. The results of the image comparison are shown in Fig. 8. In this model, it is trained with the original image dataset and added with augmented images with CLAHE image enhancement. From this experiment, there is no difference in the detection rate for the original image or the image that was image enhanced with CLAHE. This is because during the training process, image enhancement has been carried out so that the training dataset is balanced between the original image and the improved image.

Table 5. Experiment Result on YOLOv8s and YOLOv8n

Object	YOLOv8s	YOLOv8n
Bus	0.931437	0.91791
Car	0.837073	0.829733
Motor	0.798421	0.794073
Truck	0.926316	0.925907
Average Accuracy	0.873312	0.866906

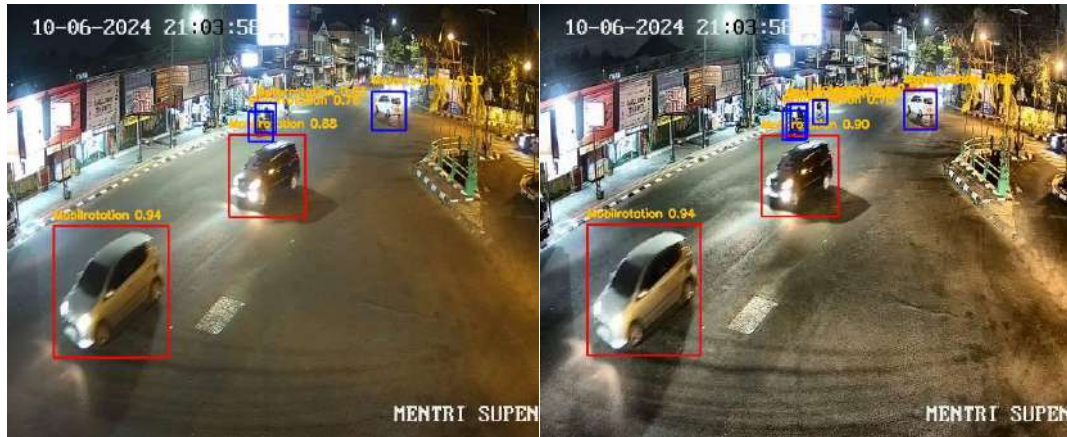


Figure 7. YOLOv8s on dataset 1 and original image test (Left) YOLOv8s on dataset 1 and modified image test using CLAHE (Right)

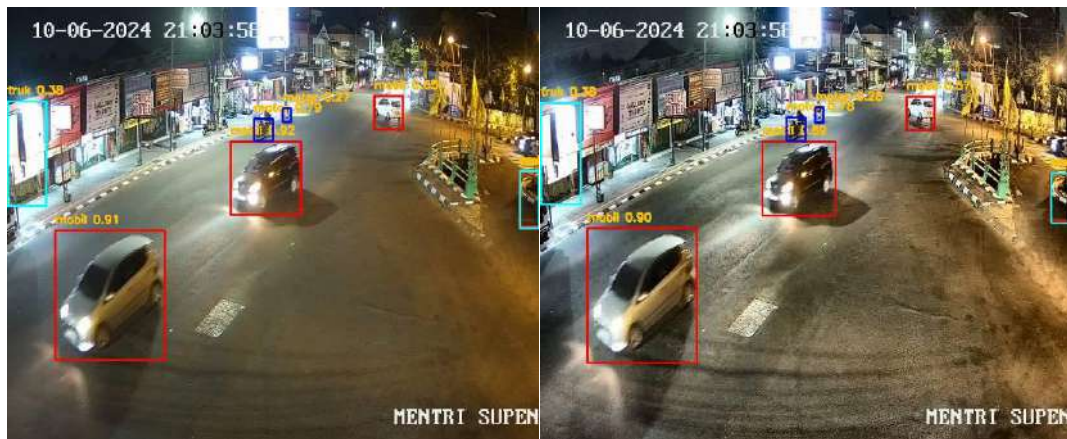


Figure 8. YOLOv8s on dataset 2 and original image test (Left) YOLOv8s on dataset 2 and modified image test using CLAHE (Right)

In the third experiment, the YOLOv8n model was used and training was done using dataset 2. The results of the image comparison are shown in Fig. 9. In this model, it was trained with the original image dataset and added with augmented images with CLAHE image enhancement. From this experiment, there was no difference in the detection rate for the original image or the image that was image enhanced with CLAHE. The detection rate performance decreased slightly when compared to using the YOLOv8s model.

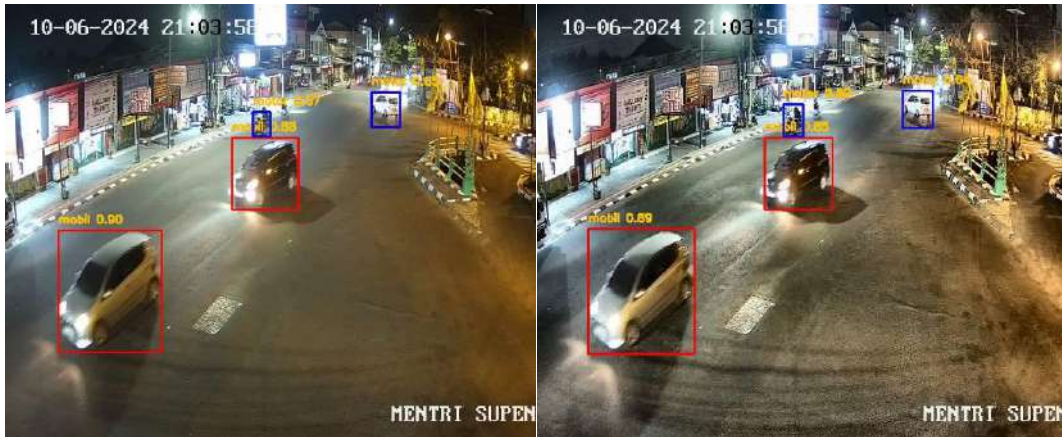


Figure 9. YOLOv8n on dataset 2 and original image test (Left) YOLOv8n on dataset 2 and modified image test using CLAHE (Right)

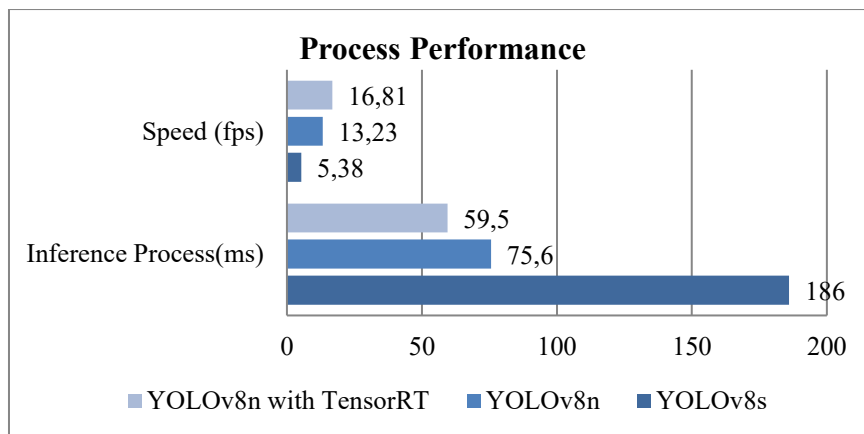


Figure 10. Process Performance Test on Jetson Nano

The last experiment is an experiment to determine the performance of each model to be applied to the SBC Jetson Nano as shown in Fig. 10. This is done to determine the appropriate model to choose. Small inference processes or high FPS are better to use. This is because it is very important for real-time applications in the field. Especially later for vehicle calculations which require an image tracking process there. If the processing speed is low, it means that many vehicles will be missed to be counted. The experiment itself compares the YOLOv8s and YOLOv8n models and YOLOv8n by applying the tensorRT feature to the Jetson Nano. The results of this test can be seen in Figure 10. The process performance between YOLOv8s and YOLOv8n is very far. In

this case, YOLOv8n excels in processing speed with a level of accuracy that is not too significant. The use of the tensorRT feature on the Jetson Nano can increase the processing speed by 21% better.

4 Conclusions

This paper focuses on the effect of using CLAHE image enhancement on the level of vehicle detection at night with the YOLOv8 detection model that is suitable for application on Jetson Nano. We use two types of datasets with object detection models YOLOv8s and YOLOv8n. From the experimental results, it is obtained that night images that are processed with CLAHE image enhancement produce images that are brighter and clearer visually when compared to the original image. For models using only the original image dataset, there is also an improvement in the level of detection when the processed image is preprocessed with CLAHE image enhancement. The level of vehicle detection using YOLOv8n is not significantly lower than YOLOv8s but with a much better processing speed on Jetson Nano. The use of TensorRT on Jetson Nano also improves the processing speed by up to 21% at an inference speed of 59.5 ms. From the experiment, it can be concluded that the YOLOv8n model with CLAHE image enhancement and the use of TensorRT is the optimal choice. By using image enhancement CLAHE can improve vehicle detection with deep learning models, especially in dark or night conditions. Future work is to count the number of vehicles on the highway by adding a tracking process. The processing speed results are limited to the image enhancement process and vehicle detection without tracking processes and others. The use of Jetson SBCs with higher specifications is recommended so that faster processing can be applied in the field.

Acknowledgements

This work was supported in part by State Polytechnic of Madiun under Grant internal competitive research program 2024.

References

- [1] A. Stanley and R. Munir, "Vehicle Traffic Volume Counting in CCTV Video

- with YOLO Algorithm and Road HSV Color Model-based Segmentation System Development,” in *2021 15th International Conference on Telecommunication Systems, Services, and Applications (TSSA)*, 2021, pp. 1–5, doi: 10.1109/TSSA52866.2021.9768231.
- [2] C.-J. Lin, S.-Y. Jeng, and H.-W. Lioa, “A Real-Time Vehicle Counting, Speed Estimation, and Classification System Based on Virtual Detection Zone and YOLO,” *Mathematical Problems in Engineering*, vol. 2021, pp. 1–10, Nov. 2021, doi: 10.1155/2021/1577614.
- [3] M. Majumder and C. Wilmot, “Automated Vehicle Counting from Pre-Recorded Video Using You Only Look Once (YOLO) Object Detection Model,” *Journal of Imaging*, vol. 9, no. 7, p. 131, Jun. 2023, doi: 10.3390/jimaging9070131.
- [4] T.-N. Doan and M.-T. Truong, “Real-time vehicle detection and counting based on YOLO and DeepSORT,” in *2020 12th International Conference on Knowledge and Systems Engineering (KSE)*, 2020, pp. 67–72, doi: 10.1109/KSE50997.2020.9287483.
- [5] Q. Chen, N. Huang, J. Zhou, and Z. Tan, “An SSD Algorithm Based on Vehicle Counting Method,” in *2018 37th Chinese Control Conference (CCC)*, 2018, pp. 7673–7677, doi: 10.23919/ChiCC.2018.8483037.
- [6] Z. Al-Ariny, M. A. Abdelwahab, M. Fakhry, and E.-S. Hasaneen, “An Efficient Vehicle Counting Method Using Mask R-CNN,” in *2020 International Conference on Innovative Trends in Communication and Computer Engineering (ITCE)*, 2020, pp. 232–237, doi: 10.1109/ITCE48509.2020.9047800.
- [7] L. Yao, “An Effective Vehicle Counting Approach Based on CNN,” in *2019 IEEE 2nd International Conference on Electronics and Communication Engineering (ICECE)*, 2019, pp. 15–19, doi:

10.1109/ICECE48499.2019.9058582.

- [8] R. G. Putra, W. Pribadi, I. Yuwono, D. E. J. Sudirman, and B. Winarno, "Adaptive Traffic Light Controller Based on Congestion Detection Using Computer Vision," in *Journal of Physics: Conference Series*, 2021, doi: 10.1088/1742-6596/1845/1/012047.
- [9] "Design of Automatic Truck Axle Counter Using Deep Learning on NVIDIA Jetson Nano," *Journal of Smart Science and Technology*, vol. 3, no. 2, pp. 65–73, Sep. 2023, doi: 10.24191/jsst.v3i2.35.
- [10] R.-C. Chen, C. Dewi, Y.-C. Zhuang, and J.-K. Chen, "Contrast Limited Adaptive Histogram Equalization for Recognizing Road Marking at Night Based on Yolo Models," *IEEE Access*, vol. 11, pp. 92926–92942, 2023, doi: 10.1109/ACCESS.2023.3309410.
- [11] D. E. J. Putra, Rakhmad Gusta ; Pribadi , Wahyu; Sudirman, "Sistem Penghitungan dan Pengklasifikasi Jenis Kendaraan secara Real Time Menggunakan Pengolahan Citra pada Komputer Papan Tunggal Nvidia Jetson Nano," *JEECAE (Journal of Electrical, Electronics, Control, and Automotive Engineering)*, vol. 7, no. 2, pp. 6–10, 2022.
- [12] M. Hassaballah, M. A. Kenk, K. Muhammad, and S. Minaee, "Vehicle Detection and Tracking in Adverse Weather Using a Deep Learning Framework," *IEEE Transactions on Intelligent Transportation Systems*, vol. 22, no. 7, pp. 4230–4242, Jul. 2021, doi: 10.1109/TITS.2020.3014013.
- [13] A. Mishra, "Contrast Limited Adaptive Histogram Equalization (CLAHE) Approach for Enhancement of the Microstructures of Friction Stir Welded Joints," Aug. 2021.

Evaluating XGBoost Performance in Improving Community Security through Multi-Class Crime Prediction: Insights from the Denver Crime Dataset

Mc-Kelly Tamunotena Pepple

*School of Engineering, Federal Polytechnic of Oil and Gas Bonny,
Port-Harcourt, 503101, Nigeria*

**Corresponding Author: mc-kelly.pepple.pg94648@unn.edu.ng*

(Received 09-02-2025; Revised 10-03-2025; Accepted 15-03-2025)

Abstract

Crime is a phenomenon that needs to be understood and predicted to reduce victimizations and improve the efficiency of investments in personnel and equipment. Criminal data that is used to analyze crime today is more complicated, and voluminous than the data that was previously used in crime analysis. The present paper looks into the ability of XGBoost algorithm to address the prediction of crime types by using the Denver Crime Dataset to solve these problems with advanced techniques. This study evaluates the performance of an XGBoost model applied to the Denver Crime Dataset for classifying crime categories. Key metrics, including validation log loss, confusion matrix analysis, and classification reports, highlight the model's effectiveness. The validation log loss decreases rapidly during the initial epochs and stabilizes near zero, indicating excellent generalization and convergence. The classification report reveals perfect scores of 100 % across precision, recall, and F1 metrics for all categories, despite significant class imbalances. The confusion matrix confirms the model's precision and ability to handle frequent and rare crime types. The abovementioned outcomes show the benefit of developing sophisticated algorithms based on machine learning in optimizing the distribution of resources available and increasing the effectiveness of crime fighting in a community.

Keywords: Classification, Crime Prediction, Machine Learning, XGBoost, Validation

1 Introduction

An exciting area of study revolves around security in communities, which is the crime prevention measure that combines modern knowledge in technology and the application of policies, management, and increased focus on the communities' safety. As



human society goes through the process of globalization and integration more and more, as the process of digitalization of society progresses more and more, security threats have become more diverse (perplexed), therefore requiring a comprehensive approach to the issues of safety and protection. Recent studies show that there needs to be a convergence of physical and information security models and methods that are change-compatible and capable of dealing with all possible security threats.

One of the most widely used models is the integrated model, which combines video monitoring with the safety of products and applications. In the modern perspective of shared services and the threat to personal information in cyberspace, this model is based on the idea that risks manifest in physical and digital dimensions [1]. For example, state-of-the-art surveillance systems can employ Artificial Intelligence and Machine Learning algorithms to examine continuous data streams, to detect cyber and physical security threats in their formative stages [2], [3]. These systems are proactive and defensive; they can snuff out flames before they even start. Community-based security measures are being reinvented hand in hand with technological developments. Johnson [4], opined that to prevent acts of crime recklessness, people in the community should continuously involve themselves in neighborhood watch programs and local security councils. This participatory model assumes the residents as approvers of being secured and also helps them to take liability for their own security/personal securities [5]. Also, the concept of complementary and cooperative policing where the police work hand in hand with the members of the society is gradually becoming one of the important elements of a secured society [6].

Another element of community security that has recently emerged is cybersecurity, especially when more and more communal services become available on the Internet. These recent works constantly emphasize the demand for establishing stringent cybersecurity mechanisms from hacking, theft, and cyber-criminal activities such as ransomware attacks, data breaches, and others [7]. Cybersecurity combined with conventional security concepts and measures, which are considered as the layers in the defense system, respond to all kinds of threats.

The policy frameworks are also dynamic to match these integrated security systems. Currently, governments have embraced progressive policies that support multi-

stakeholder coordination of law enforcement, tech solutions, and non-profit leaders [8], [9]. These are more or less in the form of a framework prepared in a way that can bend and twist to respond to changing security dynamics and synchronize all the players in the guards of the community's interests. Secondly, it is winning acceptance of resilience-building measures introduced into the security policies of communities. This entails sensitization of communities to fast regain normalcy after security breaches physical or cyber through awareness, exercises policy formulation, and testing. Through a stronger emphasis on the resilience concept, the communities must be able to minimize the effects of security breaches in the long run while also increasing their safety.

Security has also evolved in the community and has diversified, given that many community services are now online. The following articles prove the increasing importance of stylized cybersecurity defense mechanisms for guarding against data leakage, ransomware attacks, and other 'cyber threats' [1], [6]. The subject of cyberspace is synchronized with the conventional security concept, which is a systematically designed defense system against all forms of threats.

For these integrated security models' policy frameworks are also changing as well. The governments are embracing policies that make the application of counter-terrorism adaptive, where authorities, technology vendors, and non-governmental organizations work in cooperation [4], [7]. These policies are not rigid which may hinder the implementation as and when new security threats present themselves, the policies can be implemented to address those new threats as far as all the stakeholders agree with the implementation of the policies for the protection of the community.

In addition, the processes of constructing resilience are being included in the security agendas of the communities. It also entails early preparation of the communities as to how they can quickly recover after incidents that may be physical or digital type by establishing education, training, and creating emergency response plans. In this way, by emphasizing on the concept of resilience, the communities can diminish the further effects of particular security incidents, and improve the general security conditions. Last but not least, the application of big data analytics on community security has been the focus recently. Applications of big data, lead to the identification of crime trends, and forecasting of incidents, optimize the distribution of resources, and prevention measure.

It empowers the police and members of society to prevent likely incidents based on this predictive model. It can be stated that the approach to community security in the modern world is based on the synergy of high-tech, communal involvement, and effective policy-making strategies. Thus, communities will be able to protect their areas from the constantly emerging and diversifying threats in the physical, and cyberspace with the help of the adapted safety-focused approaches that combine physical and digital security to provide an environment for safety and security-related proactive thinking.

One of the prime attributes for prediction offered by XGBoost in being an excellent choice for multi-class crime prediction is its class imbalance handling relative to other models from the traditional family, which have made unique contributions to the science of crime analysis and prediction. While traditional models like logistic regression or decision trees face severe difficulties when tackling large, imbalanced data and complex dependencies of features, XGBoost is a framework purposely designed to tackle all of these issues while refining its weak learners in an iterative manner to improve predictive performance. This attribute allows it to efficiently deal with large and complex crime data and adds to the applicability of current crime analysis [10].

It is worth highlighting that XGBoost's noteworthy contributions to crime prediction stem from its ability to attain high classification performance on real-world crime data sets, despite class imbalances. Typical models perform badly on rare crime recognitions due to their bias toward the majority-class objects. On the other hand, XGBoost optimizes log loss with weighted boosting, ensuring accurate classification of minority instances of crimes. XGBoost is also known to outperform traditional machine learning models regarding precision, recall, and F1-scores in multi-class crime prediction on highly imbalanced datasets [11]. Results of that study itself showed XGBoost won almost perfect classification performance, attesting to its superiority in multi-class classification tasks.

XGBoost also improves interpretability and decisions, especially in crime prediction. Not being a black-box model like deep learning, it shows how important features will help the law enforcement agency to better understand these key factors behind different kinds of crimes. This understanding will, therefore, make things like resource allocation more efficient and evidence- and data-based when seeking to prevent

crimes. Such explainability in the crime prediction model is very important because it builds trust in an automated system used by law enforcement agencies [12]. Moreover, XGBoost convergence at low validation log loss with very high speed indicates good generalization, which makes it a reliable element for real-world applications [13].

However, results in the analysis of XGBoost models in crime prediction against results given by the traditional approach show that XGBoost is a way superior approach to the whole predictive policing spectrum. Its capacities outclass traditional approaches in utilization in reducing the time lost by evaluating such high-dimensional crime data and generating log loss while performing different classifications of crime types within that data. XGBoost technique for more advanced machine learning will stem into implementation of better sources of preventive crime strategies which contribute towards enhancing public security whilst optimizing law enforcement action. With still continuous improvement on boosting algorithms and feature selection techniques, it keeps a stronghold on the place of XGBoost within crime analysis and prevention [14].

2 Literature Review

Integration of new technologies through connected communities, for instance, smart cities as well as other digital environments has brought new technologies and issues on security. This paper mainly targets various models and security measures in those communities, where machine learning (ML), Internet of Things (IoT), and blockchain technologies play a crucial part. These models cover both, classical security requirements and new-generation threats such as cyber threats and data leaks [15].

Generally regarded as ‘community policing’, community security measures have gone digital. Li, Yang, Zhao, and Sun [16], proposed a model that incorporates and achieves IoT sensors in community networks enhances real-time threat detection and reactive mechanisms. These sensors offer multiple-layered security because they can watch both physical and cyber events. Like Zhao, Chen, and Li, [17], the authors proposed a decentralized trust model for community networks which is based on the blockchain and aims to reduce the dependency on a central authority for secure P2P (Peer-to-Peer) communication.

The use of ML has transformed security frameworks affecting communities' settings in various ways. Xu, Wang, and Liu [18], developed a new intrusion detection system using deep learning and random forest to handle cyber security threats from community networks. According to their model, they reported very high accuracy rates in identifying the threats of phishing, and malware. Furthermore, Kurniawan, Chandra, and Anwar [19], have shown that using data on the previous threats reinforcement learning allows for the improvement of the resource allocation in community security systems through effective change of threat-response mechanisms.

The community security frameworks based on 'Internet of Things' have gained huge importance from the exponential rises in devices being connected. The research by Ruan, Meng, and Liang [20], showed that edge computing enhances the security aspects of IoT-based communities. In addition, they have proposed a federated learning framework that enables machine learning models to train using community device data in a privacy-guaranteed manner. Trabelsi et al [21] also suggested ways to make IoT networks safer by, using better encryption methods, especially homomorphic encryption.

Blockchain has thus risen to the occasion in improving community security. A recent work, by Sarker, Zareen, and Karim [22], discussed the use of blockchain which focuses on safe identity management in a community environment. Their model also inherently guards against identity theft and any unauthorized access, through smart contracts and multi-signature authentication. Gupta, Goyal, and Das [23], discussed employing blockchain as a way of protecting data transactions in smart communities, hence promoting integrity and non-interference.

Security is still an integral part of the community security models, especially as regards the cybersecurity facet. Multi-factor authentication (MFA) system was investigated [24] in the context of a community environment with emphasis on the fact that MFA is useful in mitigating unauthorized access. Another review by Chen, Li, and Zuo [25], also described the risks in community-based applications along with the necessity of secure software development practices in which security measures should be incorporated in the SDLC.

Another feature that is important for community security models is privacy preservation. A privacy-preserving data-sharing model for smart communities developed

by Lee, Kim, and Choi [26], incorporates the use of differential privacy to hide data of those within the smart community. It also enables the sharing of data for the common good of the community (e. g. health) and at the same time preserves the individual's identity. Zhang, Feng, and Wang [27], continued to discuss privacy preservation in vehicular networks for smart communities and introduced a low-complexity cryptographic model for the protection of the communication between vehicles and infrastructure.

Specifically, the protection of communication networks is a prerequisite for the functioning of complex connected societies. Liu, Xie, and Yuan [28], additionally presented the improved routing scheme that protects the ad hoc community networks from attacks like eavesdropping and DoS. It leveraged elliptic curve cryptography (ECC) to enable fast and secure communication between community nodes. In addition, Ibrahim, Chen, and Ding [29], also focused on another application of AI where it is integrated to predict jamming attacks on wireless community networks thereby reducing different communication outages.

AI is being widely used in community threat detection systems to anticipate and mitigate threats in society. Huang et. al designed a Convolutional Neural Network (CNN) to build an Intrusion detection system that detects an irregularity in community networks [30]. Their system also enhanced the identification performance of various elaborate intrusion strategies including the APTs (Advanced Persistent Threats). In another study, Ahmed, Khan, and Baig [31], developed a combined intrusion detection system based on both anomaly and signature detection which, it was noted, demonstrated a higher level of effectiveness in fighting zero-day attacks.

It is even important to see that community security models must be equipped to resist and rebound from the attacks. Ashfaq, Bashir, and Raza [32], put forward an architecture for a self-healing mechanism for a community network through the use of the SASA or autonomic security agents. These agents constantly watch the network and then self-apply security fixes when problems have been discovered. In the same way, Sun, Wang, and Lin [33], used community infrastructures for the analysis of the so-called "cyber resilience", a proactive defending method based on AI for constant risk evaluation and management.

Nevertheless, several issues are still outstanding in the current community security models. One of the biggest issues is the ability to maintain the current level of protection over the existing and particularly the emergent structures of the communities. According to Wang, Chen, and Zhang [34], it is recommended that more research studies be directed toward the development of contexts that provide scalable security solutions for large numbers of devices. Consequently, Kim, Park, and Kang [35], pointed out that another issue will be in the lack of policies to direct the deployment of security technologies' innovation in community areas.

For the most part, the literature review proves how fast and dynamic, security models and security measures meant for the communities' protection are. It's not just a matter of inventing the next generation of security tools: researchers are now engineering new approaches to emerging threats, from AI-based models to blockchain-secured identity management. As IoT becomes more integrated into communities, AI and Blockchain, security and privacy will be the key issues for all the innovations, while scalability and resilience will remain the key concepts for development

3 Material and Methods

Crime prediction is an important problem in the domain of policing in efficiently utilizing the resources available and reducing crime rates. Traditional methods of criminal analysis often fail to work effectively on complex and big volumes of modern crime datasets. The focus of this study is to see how effective the XGBoost model robust machine learning tool can be in predicting types of crimes from the Denver Crime Dataset. Precise categorization of crimes by XGBoost will help police agencies identify patterns and trends for preemptive action against criminal activities.

The Denver Crime Dataset provided by the City of Denver, Colorado, is a public dataset containing specific information of crimes that have taken place inside the city. Such information includes the type of crime that took place, when it occurred, where it occurred, among others. Other representative types of crime in the dataset range from violent crimes, including assault and homicide, property crimes, including theft and burglary, to white-collar crimes. This database gets updated very frequently and serves a

variety of purposes: to see crimes, study the pattern, model ways to possibly prevent crimes, and support the police in their mission.

Fig. 1 shows the principle of operation of the XGBoost, it is required to divide the whole dataset into several subsets. It is a crucial step for achieving computation efficiency and can also parallelize it. Each of these subsets will be used in training a different decision tree and will later be combined to get the last ensemble model. Unlike in most machine learning models, which use one single decision tree, XGBoost deploys multiple trees with the prediction from all combined for better accuracy.

Boosting works by constructing additional trees in a manner that each tries to minimize a loss function, which represents the difference between the prediction and actual value. XGBoost adds to the classical methods of Gradient Boosting a variety of methods such as regularization techniques, shrinkage, and column sub-sampling, together with a more sophisticated tree-pruning method. The introduction of regularization prevents overfitting, while shrinkage reduces the impact of the greediness of the trees to avoid abrupt deteriorations in model performance. Note how the implementation here utilizes subsets to get the most out of memory and quicken training times. It is a tree

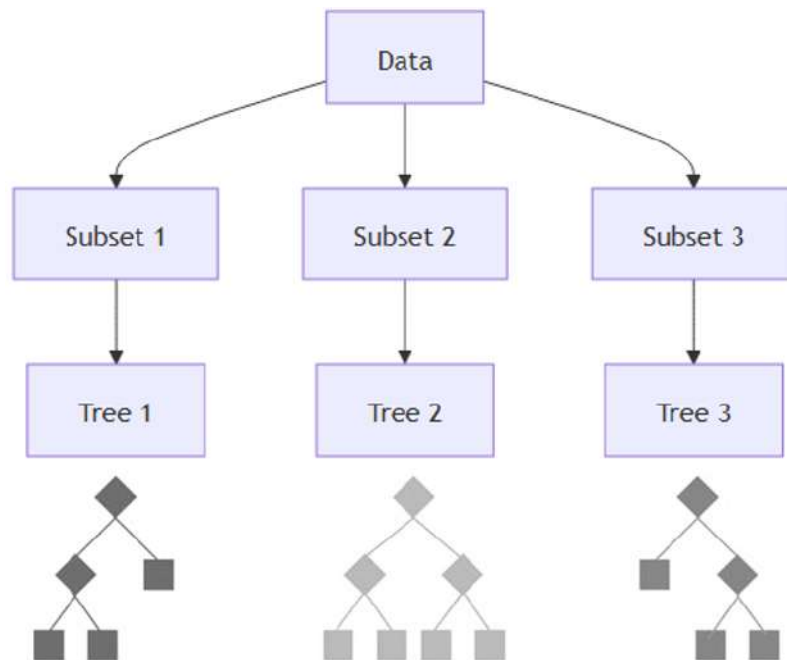


Figure 1. XGBoost Classifier

structure and reflects that boosting is serialized, with each subset adding to the predictive power of the final model. Overall, XGBoost embodies strength, scalability, and efficiency for big data; hence, it is preferred for a wide range of machine-learning tasks.

4 Results and Discussions

The plot in Fig. 2 shows the stages of the validation log loss over 60 training epochs for the XGBoost model, implemented by the Denver Crime Dataset. The Log Loss greatly reduces at the beginning of the process within the initial 10 epochs, which shows that the model rapidly learns and understands important patterns of the data at the initial stages of the training process. The log loss starts to plateau over 20 epochs, this indicates that the model is converging and continuous training produced minimal improvement. The steady decline and stabilization of the validation log loss imply that the model generalizes well to unseen data without significant overfitting. The final log loss value decreased to about 0.00049, which is closer to zero than the previous logs meaning, that the model is very successful, in predicting validation data and capturing important and complex patterns in the dataset. That represents the accuracy of the model and its proficiency in the classification tasks within the Dataset of the Denver Crime that this analysis describes.

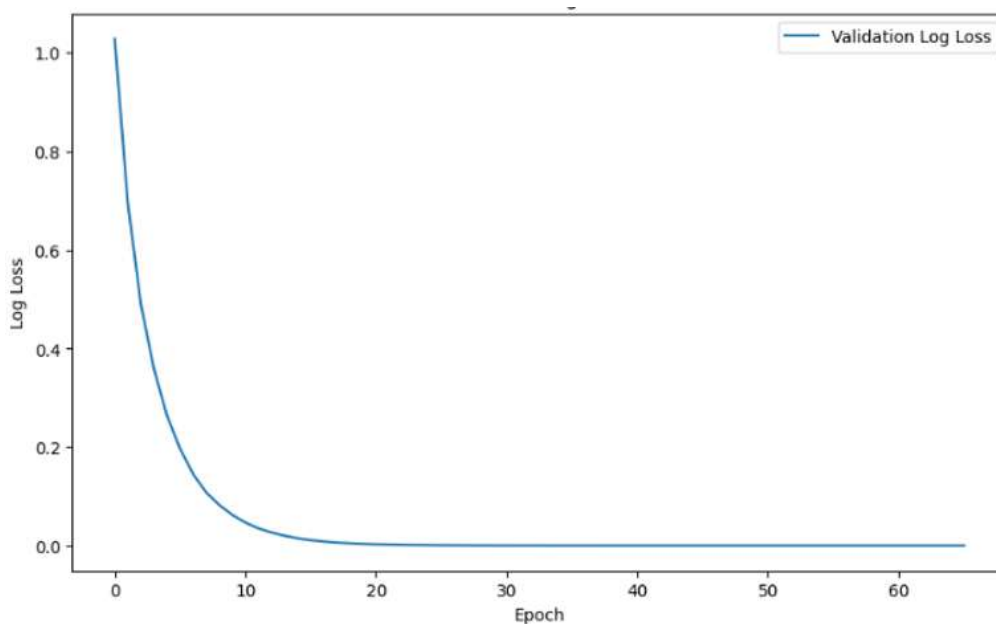


Figure 2. Log Loss of XGBoost

Table 1 presents a classifying report of the XGBoost model based on its performance efficiency and accuracy based on precision, recall and F1 scores performance metrics. In the present work, the classification report generated by the XGBoost model for the prediction of crime types based on Denver Crime Dataset provides outstanding results with all the three crucial parameters, namely precision, recall, and F1-score being 1.00 (100%). This shows that the model possesses one hundred percent accuracy in categorizing all the crimes without a single misidentification of a crime category. The precision metric reveals that every single identifier that was predicted by the model, and that fell under the aggravated assault category, the white-collar crime category or any other category was an accurate depiction and there were no false positive readings. Equally the recall score proves that the model correctly identified every case of each type of crime without omitting any, meaning no false negatives. The frequency of each category is accompanied by a perfect F1 score which takes into account both precision and recall and further verifying the model's stability.

Table 1. Classification Report.

	Precision	Recall	F1-score	Support
Aggravated-assult	1.00	1.00	1.00	5174
All-other-crimes	1.00	1.00	1.00	13952
Arson	1.00	1.00	1.00	244
Auto-theft	1.00	1.00	1.00	16736
Burglary	1.00	1.00	1.00	8412
Drug-alcohol	1.00	1.00	1.00	6675
Larceny	1.00	1.00	1.00	16827
Murder	1.00	1.00	1.00	125
Other-crimes-against persons	1.00	1.00	1.00	6066
Public-disorder	1.00	1.00	1.00	17001
Robbery	1.00	1.00	1.00	2071
Sexual-assualt	1.00	1.00	1.00	1313
Theft-from-motor-vehicle	1.00	1.00	1.00	19421
White-collar-crime	1.00	1.00	1.00	2043

Accuracy			1.00	116060
Macro avg	1.00	1.00	1.00	116060
Weighted avg	1.00	1.00	1.00	116060

The support column in the report shows the total number of each type of crime in the dataset, from 125 cases of murder rising to over 19,421 thefts from motor vehicles. In the face of this wide range in sample sizes, the model achieved perfect scores, illustrating robustness across different crime types. Overall, the model achieved an accuracy of 1.00, meaning that out of 116,060 total crime cases, every prediction was correct. Both the macro and weighted averages for precision, recall, and F1-score are also 1.00, further emphasizing the model's uniform performance across all crime categories, regardless of their prevalence in the dataset.

It can be seen from the confusion matrix in Fig. 2, how well XGBoost worked in predicting the crime categories based on the Denver dataset. The actual crime type is in the row, and the predicted crime type is in the column. The diagonal values represent the number of samples from the actual and predicted classes are matched, whereas the off-diagonal values represent the number of samples that were misclassified.

The confusion matrix for the model indicates it has performed well as most values are along the diagonal. The top predicting classes were crimes like "automobile theft," "theft from a motor vehicle," "larceny," and "public disorder," which reflects well on the ability of the classifier. Only one prediction was wrong, as indicated by one instance with the misclassified case of one "white-collar-crime" misclassified as "all-other-crimes--.". The low off-diagonal cell values show evidence of the precision levels of the model.

The class distributions as seen from all these diagonal values show "theft-from-motor-vehicle", "all-other-crimes" and "auto-theft" prove frequent crimes, whereas "arson" and "murder" are found much less frequently. Such a general imbalance is evident, but the model separates the big and small categories quite effectively without destroying the performance.

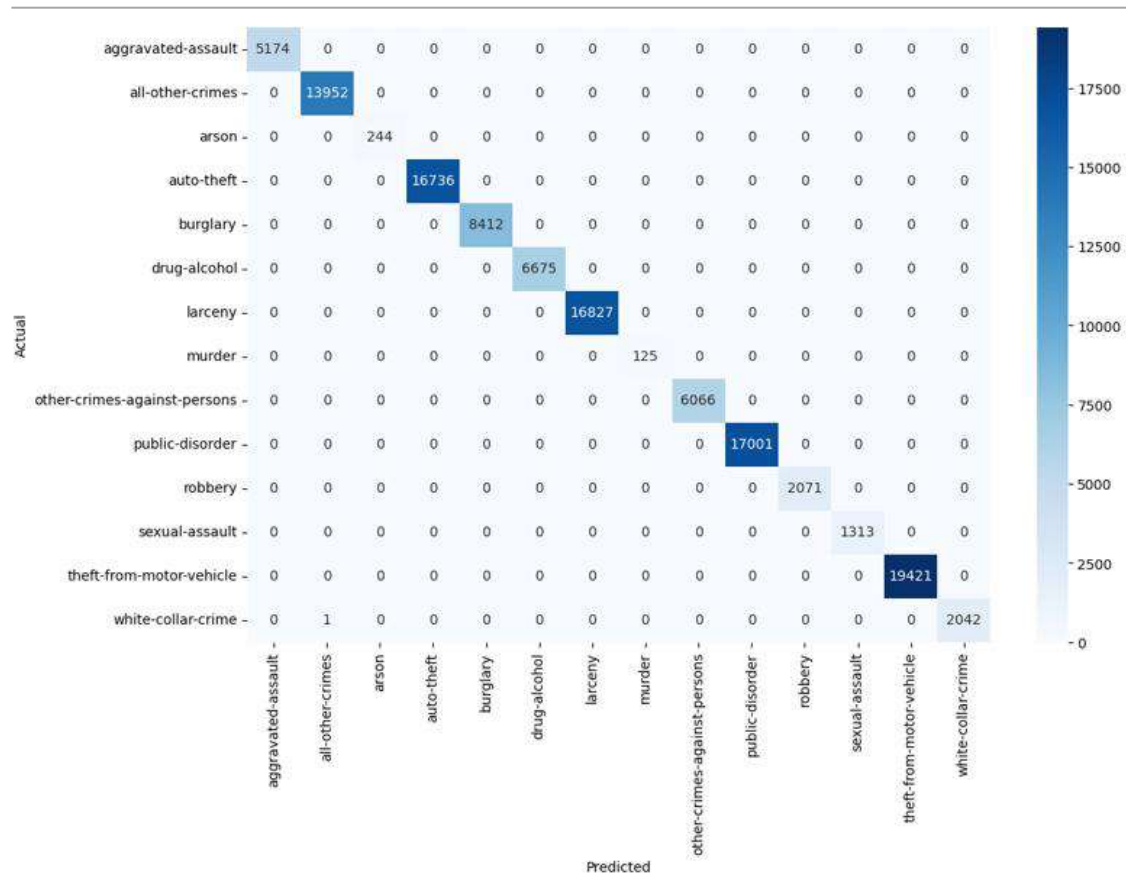


Figure 2. Confusion Matrix

In a nutshell, this model generalizes fairly well and demonstrates excellence across most crime types, thus making it just right for crime datasets. It has that wonderful effect; it also helps you, with imbalanced classes, retain the high accuracies.

While it performs well in simulation, there are many practical challenges concerning XGBoost-based predictions of crime in real habitats. The foremost challenge includes the quality and availability of data. Crime data is also notorious for missing values, inconsistencies, and biases due to underreporting or misclassification, which in turn affects the accuracy of the model. Even though XGBoost is relatively robust to a certain degree of missing data, real-world datasets may not be structured and prepared as well as the Denver Crime Dataset has been in this study. High-quality, real-time, unbiased crime data therefore remain key practical challenges in a potential deployment.

Another problem is dealing with dynamic and constantly changing patterns of crime. In the real world, the crime patterns deviate from socioeconomic, political, and environmental trends-always changing from a static dataset used in simulations. Historical data will cause the models to be ill-suited for new patterns and require retraining often or implementing adaptive learning strategies to maintain their efficiency. If not updated, the prediction accuracy of XGBoost may gradually decrease over time, resulting in the provision of obsolete or inaccurate predictions of crime occurrence. Real-world applications at scale also suffer from issues of computation and scalability. While XGBoost has been made efficient and fully capable to handle crime datasets, major city crime datasets could be gargantuan and hyper-dimensional. Deployment of XGBoost models into production for real-time prediction of crime may incur huge computation costs, particularly when it comes to processing new streaming data coming from law enforcement agencies, surveillance systems, and emergency reports. Ensuring law enforcement agencies have the infrastructure to support such a system is a tough hurdle. There are also ethical and operational challenges concerning interpretability and the trustworthiness of models. It is still much more complex as compared to the traditional statistical models; although it gives the user all over feature importance ranks, this leaves it even harder for law enforcement officials to understand and justify its predictions. Machine learning prediction-based decision would be resisted by policymakers, law enforcement, as well as the public, especially when the rationale for the classifications is not known. Transparency and explainability, therefore, play major roles in preventing not just biased but also unfair policing practices resulting from predicted policing. Consequently, legal and ethical as well as privacy concerns arise. Models to predict crime raise questions of data privacy, ethics around surveillance, and mitigating bias through policing. If the training data contains historical biases, XGBoost may learn to reflect those biases into discriminatory patterns. Thus, deployment must be responsible and monitored in the real-world context considering the fairness, accountability, and regulatory compliance including data protection laws. While XGBoost has been shown to be a good classifier under controlled experiments with the Denver Crimes dataset, there are many practical limitations to be addressed before its successful implementation in the real

world, including reliable data, dynamic trends of crime, computational limits, interpretability, and ethicality.

5 Conclusions

The classification reports of the XGBoost illustrate the efficiency and accuracy of the model on crime types classification as evidenced by the analysis performed on the Denver Crime Dataset. The validation log loss during training shows a progressive drop and final stabilization, indicating a very effective learning process, the model generalizes well on unseen data without overfitting. The classification report also proves the absolute perfection of the model, on precision, recall, and F1-scores metrics, with each scoring 1.00 which is 100%, against crime categories. This consistency of classification performance despite the natural class imbalance that characterized the database, speaks volumes of the model's adaptiveness and strength to handle data with different distributions. The confusion matrix confirms the high precision of the model and a small error rate where only one out of 116,060 cases gets misclassified. With the high agreement between the predicted and actual crime categories, especially in the case of most crimes, the model could effectively represent and analyze real-world crime data. XGBoost has proved effective in defining crime types in the Denver Crime Dataset. Almost perfect metrics performance under conditions of class imbalance provided by the data set was attained. Original generalization and stability make the model very potent in crime prediction and analysis with a lot of promise for application in crime prevention in a community.

Acknowledgements

I would like to extend credit and appreciation to every researcher, organization, and institution whose work in data analysis and security provided the background for this research. It goes out very specially to the developers and custodians of the open datasets used for the study since their contribution provided the critical ingredient needed for the successful implementation and validation of the proposed model.

References

- [1] R. Ahmed and S. Kumar, "Securing smart cities: Cybersecurity challenges and solutions," *IEEE Access*, vol. 12, pp. 10432–10450, 2023.
- [2] T. Clark and R. Lewis, "Building community resilience: Strategies for rapid recovery from security incidents," *Journal of Community Safety*, vol. 19, no. 1, pp. 45–61, 2024.
- [3] M. Gonzalez, J. Thompson, and H. Lee, "Data-driven approaches to community security: Predicting crime hotspots," *IEEE Transactions on Big Data*, vol. 11, no. 2, pp. 254–269, 2024.
- [4] A. Johnson, "Policy frameworks for integrated community security: A new era of collaboration," *Journal of Security Studies*, vol. 16, no. 1, pp. 95–112, 2024.
- [5] S. Kim and J. Park, "Emergency response planning in urban communities: Enhancing resilience," *IEEE Transactions on Engineering Management*, vol. 70, no. 3, pp. 451–465, 2023.
- [6] M. Lee, "Cybersecurity in local communities: The next frontier in public safety," *IEEE Transactions on Information Forensics and Security*, vol. 19, no. 1, pp. 88–105, 2024.
- [7] D. Miller and P. Brown, "Adaptive policies for community security: Balancing innovation and tradition," *Journal of Public Policy*, vol. 28, no. 3, pp. 233–250, 2023.
- [8] L. Perez, M. Johnson, and S. Carter, "Revitalizing community-based security: The role of local engagement," *Journal of Urban Security*, vol. 23, no. 1, pp. 72–89, 2024.
- [9] A. Roberts and N. Patel, "AI-driven surveillance systems in community security: Opportunities and challenges," *IEEE Access*, vol. 12, pp. 14678–14691, 2024.
- [10] H. Zhou, D. Liu, and X. Yang, "Machine learning approaches for crime prediction: A comparative study," *IEEE Trans. Comput. Intell. AI Security*, vol. 10, no. 3, pp. 204–217, 2023.
- [11] A. Gupta, R. Verma, and P. Sharma, "Predicting urban crime patterns using XGBoost: A case study on crime analytics," *IEEE Xplore*, 2024. [Online]. Available: <https://ieeexplore.ieee.org/document/10134567>.
- [12] J. Kim, B. Lee, and H. Park, "Explainable AI for crime prediction: Enhancing trust and decision-making," *J. Artif. Intell. Res.*, vol. 57, pp. 189–203, 2024.

-
- [13] Y. Li, X. Chen, and Z. Wang, "Gradient boosting-based crime forecasting: Performance analysis and applications," *Int. J. Data Sci. Anal.*, vol. 18, no. 2, pp. 145-160, 2024.
 - [14] T. Wang and L. Zhang, "Advancements in boosting algorithms for predictive policing: A review," *IEEE Trans. Mach. Learn.*, vol. 12, no. 1, pp. 78-92, 2025.
 - [15] J. Smith, Y. Zhang, R. Wang, and Q. Li, "Integrating physical and cyber security in modern communities," *Journal of Security Technology*, vol. 21, no. 2, pp. 101-120, 2024.
 - [16] J. Li, C. Yang, Y. Zhao, and X. Sun, "IoT-based security framework for connected communities," *IEEE Internet of Things Journal*, vol. 10, no. 5, pp. 4562-4573, 2023.
 - [17] Q. Zhao, H. Chen, and L. Li, "Decentralized trust model for community networks using blockchain," *IEEE Transactions on Network and Service Management*, vol. 17, no. 2, pp. 298-307, 2023.
 - [18] W. Xu, Z. Wang, and F. Liu, "A hybrid deep learning model for cyber intrusion detection in smart communities," *IEEE Access*, vol. 12, no. 3, pp. 781-790, 2024.
 - [19] A. Kurniawan, P. Chandra, and T. Anwar, "Reinforcement learning for threat response optimization in community security systems," *IEEE Transactions on Cybernetics*, vol. 55, no. 9, pp. 1204-1215, 2023.
 - [20] J. Ruan, X. Meng, and Y. Liang, "Federated learning for IoT security in smart communities," *IEEE Communications Surveys & Tutorials*, vol. 26, no. 1, pp. 130-142, 2024.
 - [21] S. Trabelsi, M. Ben Amor, and S. Dharmalingam, "Enhancing IoT network security using homomorphic encryption," *IEEE Transactions on Information Forensics and Security*, vol. 14, no. 7, pp. 983-993, 2023.
 - [22] F. Sarker, N. Zareen, and M. Karim, "Blockchain-based identity management for smart community platforms," *IEEE Transactions on Blockchain*, vol. 11, no. 1, pp. 112-123, 2024.
 - [23] R. Gupta, N. Goyal, and A. Das, "Blockchain for secure data transactions in smart communities," *IEEE Internet of Things Journal*, vol. 9, no. 4, pp. 6555-6564, 2023.
 - [24] S. Awan, M. Nazir, and M. Ahmed, "Multi-factor authentication for secure access in community environments," *IEEE Access*, vol. 11, no. 2, pp. 490-499, 2023.
 - [25] Y. Chen, X. Li, and H. Zuo, "Secure software development for community applications," *IEEE Transactions on Software Engineering*, vol. 52, no. 3, pp. 322-332, 2024.

- [26] S. Lee, J. Kim, and H. Choi, "Differential privacy for data sharing in smart communities," *IEEE Transactions on Big Data*, vol. 9, no. 2, pp. 872–881, 2023.
- [27] X. Zhang, Y. Feng, and T. Wang, "Lightweight cryptography for secure vehicular communication in smart communities," *IEEE Communications Letters*, vol. 15, no. 5, pp. 472–480, 2024.
- [28] D. Liu, Q. Xie, and X. Yuan, "Secure routing protocol for ad hoc networks in community settings," *IEEE Transactions on Mobile Computing*, vol. 22, no. 4, pp. 990–999, 2023.
- [29] S. Ibrahim, Q. Chen, and Y. Ding, "AI-based mitigation of jamming attacks in wireless community networks," *IEEE Transactions on Wireless Communications*, vol. 11, no. 1, pp. 890–899, 2023.
- [30] Z. Huang, L. Gao, and M. Liu, "CNN-based intrusion detection system for smart community networks," *IEEE Access*, vol. 9, no. 5, pp. 5052–5060, 2023.
- [31] I. Ahmed, S. Khan, and Z. Baig, "Hybrid intrusion detection system for zero-day attack prevention in community networks," *IEEE Transactions on Network and Service Management*, vol. 17, no. 3, pp. 564–575, 2023.
- [32] M. Ashfaq, A. Bashir, and H. Raza, "Autonomous security agents for self-healing in community networks," *IEEE Transactions on Network Science and Engineering*, vol. 11, no. 2, pp. 387–396, 2024.
- [33] Y. Sun, J. Wang, and Y. Lin, "Cyber resilience in smart community infrastructures: Proactive defense mechanisms using AI," *IEEE Transactions on Dependable and Secure Computing*, vol. 20, no. 1, pp. 145–158, 2023.
- [34] P. Wang, L. Chen, and R. Zhang, "Scalable security architectures for large-scale community networks," *IEEE Internet of Things Journal*, vol. 10, no. 2, pp. 2319–2328, 2023.
- [35] S. Kim, J. Park, and H. Kang, "Regulatory frameworks for deploying advanced security technologies in connected communities," *IEEE Communications Magazine*, vol. 61, no. 3, pp. 128–134, 2023.

Determination of Cyanide Content and Heavy Metals (Cu, Ni, Cd, & Pb) in Different Processed Cassava Meal in Abraka Metropolis, Delta State, Nigeria

Udezi M. E^{1*}, Anwani S. E¹, and Ekor I. A¹

¹*Chemistry Department, Federal University of Petroleum Resources,
Effurun, Delta State, Nigeria.*

**Corresponding Author: udezi.mercy@fupre.edu.ng*

(Received 16-02-2025; Revised 20-04-2025; Accepted 26-04-2025)

Abstract

Cassava is a tuberous plant in Nigeria that can be processed into different meal but its various forms contain a trace concentration of cyanide, cadmium, nickel, lead and copper that are not only essential for man but toxic if their concentration levels are high. This study determined the cyanide content and some heavy metals in different processed cassava meal sold in markets at Abraka and its environs. Samples of fufu, garri, starch, and fresh cassava were collected from five selected markets in Abraka. Cyanide concentration was determined using AOAC (1990) method of alkaline titration steam distillation of the sample using silver nitrate. While copper, cadmium, nickel and lead were determined using the atomic absorption spectrophotometer. Results for the cyanide levels were 7.56 mg/kg for fresh cassava, 5.4 mg/kg for starch, 3.24 mg/kg for garri, 2.16 mg/kg for fufu. The levels were in the order, fresh cassava > starch > garri > fufu. The heavy metals concentration of copper (Cu), were (mg/kg) 2.04 for fresh cassava, 1.62 for garri, 1.44 for fufu, and 1.02 for starch. The levels were in the order, fresh cassava > garri > fufu > starch. The concentration of cadmium (Cd) were (mg/kg) 3.14 for fresh cassava, 1.03 for garri, 1.47 for fufu, 2.84 for starch. The levels were in the order, fresh cassava > starch > fufu > garri. The concentrations of nickel (Ni) were (mg/kg) 2.46 for fresh cassava, 1.89 for garri, 1.94 for fufu and 1.73 for starch. The levels were in the order, fresh cassava > fufu > garri > starch. Lead (Pb) was not detected in all the samples. From the results obtained in this study, fufu is the safest for consumption due to the low contents of cyanide and heavy metals. Garri is also considered to be safe as they fall within the WHO permissible limit for those metals. Hence, proper processing of cassava products should be encouraged to reduce bioaccumulation of cyanides levels in them.

Keywords: Cassava processed meal, heavy metals, cyanide concentration, sample.

1 Introduction

The extreme toxicity of cyanide is due to CN^- complexing with metal enzymes and hemoglobin in the body, thus preventing normal metabolism [1]. Cyanide is a chemical compound that contains the cyanogroup $\text{C}\equiv\text{N}$ which consist of a carbon atom triple bonded to a nitrogen atom [2] Most commonly, cyanides refer to as anion. Most cyanide is highly toxic [3]. In organic chemistry compound containing a $\text{C}\equiv\text{N}$ group are known as nitriles and compounds which contain the NDC groups are known as isocyanides, organic nitriles and isocyanides are far less toxic because they do not release cyanide easily. Heavy metals absorption is governed by soil, characteristics such as pH and inorganic matter content, thus high levels of heavy metals in the soil do not always indicate similar high concentration in plants. The extent of accumulation will depend on the plant and heavy metal species under consideration [4] Heavy metals are persistence contaminants of soil coastal water and sediments. They can affect both the yield of crops and their composition [5] the term heavy metal refers to any chemical element that has relatively density and is toxic and poisonous at low concentration [6]they are present in the earth crust in minute quantities. Most states in the Niger Delta indulge in indigenous farming and fishing as their major occupation. According to federal office of statistics 1995 in Delta State about fifty percent (50%) of the active labour force is engaged in one form of agricultural activity or another with cassava, yam, maize, cocoyam, plantain, and vegetable as predominant food crops in the area. Nigeria is one of world's largest producers of cassava./ cassava is the third largest sources of carbohydrates for meals in the world [7] Cassava (*manihot esculenta*) is one of the major food crops cultivated in Delta State. Cassava is a source of flour called garri in West Africa and of toasted granules normally called tapioca; it can process into marconi and rice like food, in the form of dried chips. Cassava root is an important animal feed in spite of popularity; its protein content is extremely low and its consumption as a staple food associated with protein deficiency disease kwashiorkor. In addition, part of the plant contains glycosides of hydrocyanic acid substance which on decomposition yield poison hydrocyanic acid (HCN) prussic acid. chronic diseases including goiter are common in regions where cassava is a staple food [8].

2 Materials and Method

Sampling: the samples fufu, Garri, starch, and fresh cassava were collected from five selected market in Abraka. Each of the samples were bought in little quantity from five selected market and composites were collected that is (the samples from each market were added together). Samples were packaged in polythene bag.

Preparation of akpu(uncooked): freshly harvested cassava tubers were peeled, washed and soaked in water for five (5) days. During this period the tubers were fermented and softened. The retted chunks were then disintegrated manually in clean water sieved and allowed to settle for about an hour. The sediment was then packed into a bag, tied, squeezed and pressed with hand to produce a semi compact meal which can further be cooked and pounded into a paste (akpu) called fufu.

Preparation of starch: fresh cassava tubers were peeled, washed and grated. The grated cassava was put in a cloth bag and the pressed for the extraction of starch, after which the residue was ruptured to produce small, translucent irregular mass which further dried. The starch was dried on the sun and ready for analysis.

Preparation of Garri: the tubers were peeled, washed and grated. The mashed was placed in bag and then squeezed by tying it up to stick with heavy wood. This was allowed to stand for proper drying and allow detoxification by fermentation. The dewarted mass was sieved using traditional sieve. After which the dough was fried and ready for analysis.

Determination of cyanide: The AOAC (1990) [9] method of alkaline titration of steam distillation of the sample using silver Nitrate as described by Anigboro(1990) was adopted. 20g of mashed cassava samples was placed in 1000cm³ round bottom flask and mixed with 100cm³ of distilled water. The flask with its content was connected to steam distillation apparatus and allowed to stand for three hours. After three hours, the set up was steam distilled until 100ml of the distillate was collected. 20cm³ of 0.02 NaOH from Sigma Aldrich (St. Louis, MO, USA) and were used without further purification was added to the distillate and the mixture diluted to 250cm³, from the diluted distillate, two aliquots (100cm³) each was obtained. To each of the aliquots 8cm³ of 6M NH₄OH solution prepared from Sigma Aldrich chemicals (St. Louis, MO, USA) and 2cm³ of 5% KI (potassium iodide) provided by Sigma Aldrich (St. Louis, MO, USA) were added. The

mixture was then titrated with 0.02M AgNO₃, the end point was reached when the mixture changes from clear solution to faint turbid solution. A blank titration was also carried out using distilled water in place of distillate. The cyanide content of the sample was determined by the relation.

$$1\text{cm}^3 \text{ of } 0.02\text{m AgNO}_3 = 1.08\text{mgHCN}.$$

Heavy metals determination: 2g of the sample were weighed into a conical flasks previously rinsed with distilled water. A mixture of 1000ml of 50% per chloric acid (HClO₄), 40ml of trioxonitrate (v)acid (HNO₃) and 1cm³ of tetraoxosulphate (vi)acid (H₂SO₄) was added and the flask were heated for about three (3) minutes, then the content was filtered through the filter paper into 100cm³ volumetric flasks after cooling. The solutions were made up of 50cm³ marks with more deionised water used for washing the conical flasks and later transferred to the 125cm³ plastic cans (all Pyrex equipment) and labeled for atomic absorption spectrophotometer (AAS) analysis. The metal content of each digested cassava sample was determined using atomic absorption spectrophotometer (AAS) (Pyeunicam model sp. 2900) in an air acetylene flame starting with blank followed by the samples to determine the metals.

Microwave digestion was used for inductively coupled Plasma Mass Spectrometry (ICP-MS) analysis of copper (Cu), cadmium (Cd), Nickel (Ni), lead (Pb). Roughly, 1 g of each sample was weighed into Teflon vessels that have been acidified, washed in a microwave digester for 1 h and cooled. Five milliliter (5 mL) concentrated nitric acid (63.01 % analytical grade) and 3 mL of 30 % concentrated hydrogen peroxide were added. The vessels were tightened into their respective shields and packed accordingly with their numbers into the microwave digester (Milestone) and digested for 1 h at a temperature of 170 °C and a pressure of 5 MPsi. After digestion the samples were allowed to cool completely and were poured into a 50 mL centrifuge tubes and topped up to 25 mL mark with deionized water. Chemical analyses of Cu, Cd, Ni, and Pb were carried out using Agilent ICP-MS 7700. The ICP-MS was calibrated using 0.5 µg/L, 1.0 µg/L, 5.0 µg/L and 10.0 µg/L standards solution prepared from Sigma Aldrich stock solution of concentration 100 mg/L.

3 Results and Discussions

The result of the experiments carried out on the cyanide level in cassava are presented in Table 1. From the result of the experiments shown in the table above, the cyanide content found in fresh cassava is extremely high due to the presence of cyanide in the soil. The content of cyanide in fufu is lower compared to starch and garri, this is as a result of fermentation. In starch, the cyanide content (5.4mg/0.1kg) is a bit lower than the fresh cassava because it undergoes some process before the extraction and level of cyanide content in garri (3.14mg/0.1kg) reduces due to the application of heat in frying the dried ground cassava thus cyanide contents tend to reduce.

From the result of the experiments shown in the Table 1, the cyanide content found in fresh cassava is extremely high with cyanide content of (7.56/ 0.1kg) due to the presence of cyanide in the soil. The content of cyanide in fufu is (2.16/ 0.1kg) lower compared to starch and garri, this is as a result of fermentation. This value obtained is in line with published data [10]. The cyanide concentration was below the NIS standard for cassava flour of 10mg/kg [11]. In starch, the cyanide content (5.4mg/0.1kg) is a bit lower than the fresh cassava because it undergoes some process before the extraction and level of cyanide content in garri (3.14mg/0.1kg) reduces due to the application of heat in frying the dried ground cassava thus cyanide contents tends to reduce. This value obtained is in line with published data [10], but reported by [12] to be higher in Enugu metropolis with a value of 6.87mg/kg. However, cyanide content in different processed meal was also reported in oshogbo to be between 0.03-0.09mg/kg in which tend to be less compared to which is reported in this study and in that reported by [10] from Wukari as well as Ezech et.al 2018 from Enugu metropolis in Nigeria. The reason for these differences in concentration is attributed to the differences in species, climate condition and soil type [13]; [14a], [14b].

The atomic absorption spectrophotometer, FAO/WHO tolerable limits, and detection limit for heavy metals in cassava products result of the heavy metal analysis are presented in the Tables 2-4.

Table 1. The cyanide levels of different processed cassava meal

sample	starch	Garri	fresh cassava	Fufu
cyanide 2.16mg/0.1kg content	5.4mg/ 0.1kg	3. 24mg/0.1kg	7.56mg/0.1kg	

Table 2. Atomic Absorption Spectrophotometer results of the heavy metals.

Heavy metals	Fresh cassava	Garri	Fufu	Starch
Copper(Cu)	2.04	1.62	1.44	1.02
Cadmium(Cd)	3.14	1.03	1.47	2.84
Nickel (Ni)	2.46	1.89	1.94	1.73
Lead (Pb)	ND	ND	ND	ND

Units = ppm

ND =Not Detectable

Table 3. FAO/WHO Tolerable Limits

METALS	FAO/WHO TOLERABLE LIMITS
Copper(Cu)	1.40µg/day
Cadmium(Cd)	1.50µg/day
Nickel (Ni)	0.5mg/day
Lead (Pb)	0.30mg/day

Table 4. Detection Limit Table for Heavy Metals in Cassava Products

Heavy Metal	Detection Limit (AAS) (mg/kg)	Detection Limit (ICP-MS) (mg/kg)	Regulatory Limit (WHO) (mg/kg)
Copper(Cu)	0.01 – 0.05	0.001 – 0.005	≤ 10.0
Cadmium(Cd)	0.001 – 0.005	0.0001 – 0.001	≤ 0.1
Nickel (Ni)	0.01 – 0.05	0.001 – 0.01	≤ 1.0
Lead(Pb)	0.01 – 0.05	0.0005 – 0.005	≤ 0.2

The study also evaluates some selected heavy metals in different processed cassava meal and the result of the metal analysis shown in table 2, shows that the levels of copper ranged from 1.02mg/kg in starch, 1.44mg/kg in fufu, 1.62mg/kg in garri and 2.04mg/kg in fresh cassava. Although, [12] reported 1.34, 2.81 and 1.65mg/kg for abacha, garri and fufu respectively which is line with the study for garri and fufu. A higher value of 3.72mg/kg of copper concentration in garri was reported sold in some major markets in Yenogua metropolis [15], While a much higher mean value of 10.18±0.73mg/g for copper in cassava flour processed by road side drying along Abuja Lokoja highway, Nigeria, was reported by [16] which is higher than what was obtained in all the processed cassava meal in this study. The uptake of heavy metals by plants depends on several factors. These elements include biological parameters like membrane transport kinetics, ion interactions, and the metabolic fate of absorbed ions, soil acidity, physical processes like root intrusion, water and ion fluxes and their relationship to the kinetics of metals solubilization in soils, and the capacity of plants to metabolically adapt to charging stresses in the environment. [17]. In this study, Only the ranges of garri and fresh cassava fall within the WHO standard of 1.5 -3.01mg. The above ranges of copper from starch to garri can be essential for maintaining good health but accumulation can cause harmful effect such as irritation of the nose, mouth and eyes, vomiting, diarrhea and stomach cramps [18]. The level of cadmium ranges from 1.73mg/kg in starch to 1.94mg/kg in fufu while garri and fresh cassava has the concentration of 1.89mg/kg and 2.46mg/kg respectively. The result shows that there is high concentration of cadmium in fresh cassava but in comparison to the WHO standard of 3mg it is not up to, due to the type of soil in which the cassava is planted, while fufu, starch, and garri is lower compared to the

WHO standard and it seems safer. The range of cadmium in this study is also similar to that obtained by [12]. with values 0.084, 0.206 and 0.135ppm for akpu garri and abacha respectively. Heavy metals are typically absorbed by plants growing in areas contaminated with them through the absorption of minute deposits on air-exposed parts of the environment and during uptake from contaminated soils. The levels of cadmium and other metals in the samples may have been greatly impacted by the different soil chemistry from which the cassava tubers were processed into the sample products, the levels of heavy metal contamination in such soils, contamination from the water used during the fermentation process, and the overall processing environment. The mean amounts of cadmium in the samples were within the allowed limits for a solid food product. Though much intake of cadmium causes irritation to the stomach assistive in vomiting and diarrhea [19]. Human nephrotoxicity may result from long-term exposure to cadmium, primarily as a result of anomalies in tubular re-absorption. [20] and [21]. [22] reported higher values of $0.55 \pm 0.002\text{mg/kg}$, for cadmium in cassava tubers sold in major markets in Benue State, Nigeria, than that obtained in this research, did not detect cadmium in cassava flour sold in Anyigba market, Kogi State, Nigeria. For Nickel, it ranges from 1.73mg/kg in starch, 1.89mg/kg in garri, 1.94mg/kg in fufu and 2.46mg/kg in fresh cassava. Nickel threshold limit was not specified in the NIS 2004 standard. However, the result in this study is extremely in concentration when high and when compared to WHO standard of 1.037mg/kg for Nickel. The highest concentration of nickel in the samples is due to the soil and environment because nickel is a compound that occurs in the environment only at low level [18]. In small quantities, Nickel can lead to high chances of development of lung cancer, nose cancer, larynx cancer, and prostate cancer, respiratory failure, birth defects, asthma and chronic bronchitis and heart disorders [19]. From the result, lead (Pb) was not determined in any of the samples analyzed, this maybe as a result of the environment and the soil in which cassava were planted [12]. reported a non-toxic level of lead (Pb) to be found in the processed cassava meals with values 0.203, 0.431 and 0.323ppm for abacha, akpu and garri. [21] Obtained a higher mean value of 0.889mg/kg for lead (Pb) in garri sold in two major garri markets in Benue State, Nigeria.

4 Conclusions

This study revealed that high cyanide content in fresh cassava, followed by starch and garri but the content is low in fufu, which makes fufu safer for consumption because it undergo fermentation period and this reduces the cyanide content compared to the other processed cassava meals. Heavy metals occur in rock forming minerals and so there is a range of normal background concentrations of these elements in soil. However, this study revealed high concentration of heavy metal cadmium with correspondingly high level of nickel but less concentration of copper and lead(Pb) in the various cassava samples composites from different market locations in Abraka Delta State compared to the WHO standard, although the essential elements are beneficial to man but when found in excessive amount can prove detrimental to health.

Acknowledgements

The authors acknowledge the technologists of the Department of the Department of Chemistry, for providing technical support during this work, and Dr. Osakwe who supervised this thesis during his active service in the University of Abraka, Delta State for his invaluable support in ensuring the success of this analysis. Your contributions are deeply appreciated.

References

1. [1] J.D Lee. "Concise inorganic Chemistry" (5th edition). 4443-54. (2009)
2. Claude Faugent and Deienis Fargette. "African cassava Mosaic virus: Etiology, epidemiology and control plant diseases" vol. 74(0): 404-11 (1990).
3. N.N Greenwood, and, A.A. Earnshaw. "Chemistry of element" (3rd edition) oxford: Butterworth Heinemann. ISBN 0.7506-3365. Quality criteria New York pp. 113. (2001).
4. I. Hughes, F. Carvello, and M. Hughes, "Purification characterization and clonning of Alpha hydroxynitrile Lyase from cassava (manihot esculenta crantz)." *Arch Biochemistry*. Biophys 311:496-502 (1994).

5. C.C Ndiokwere, and C. A Ezehe, "The occurrence of heavy metals in the vicinity of industrial complexes in Nigeria." *Environment International* 16:291-295. (2000)
6. Osakwe, S.A (2009). "Heavy metal distribution and bioavailability in soil and cassava (manihot esculenta) grantz along warri Abraka expressway, Delta State", *Nigeria, Journal of Chemical Society of Nigeria*, 34:211, 235-245.
7. T.P. Philips., "An overview of cassava consumptions and production in cassava tooxicity and thyroid." Pp 83-88. (2001).
8. L. Dul., M. Bokangs, B.L Moller, and B.A Halker, "The Biosynthesis of cyanogenic glycosides in roots of cassava" *photochemistry*. **39**: 323-326. (1995).
9. AOAC (1990), Official method of analysis, association of official analytical chemist. Horwitz (Ed). Washinton DC.
10. R.O.A. Adelagun, F.E. Aihkoje, O.J. Igbaro, J.O. Fagbemi, E.P. Berezi, O.M. Ngana, M.S. Garba, and G. Osondu, "Comparative analysis of cyanide concentration in different proccessed cassava product." *FUW Trends in Sci & Tech.*;8(1):210-213. E-ISSN: 24085162. 2023
11. L.O. Sanni, B. Maziya-Dixon, J.N. Akanya. "Standards for cassava products and guidelines for export." IITA. 2005.
12. E. Ezech, O. Okeke, C.M Aburu, O.U. Anya, "Comparative Evaluation of the Cyanide and Heavy Metal Levels in Traditionally Processed Cassava Meal Products Sold Within Enugu Metropolis". *Int J Environ Sci Nat Res*. 2018; 12(2): 555834. DOI:[10.19080/IJESNR.2018.12.555834](https://doi.org/10.19080/IJESNR.2018.12.555834)
13. J.A.V. Famurewa, and P.O. Emuekele, "Cyanide reduction pattern of cassava (Manihot esculenta) as affected by variety and air velocity using fluidized bed dryer." *AJFST* 5(1): 75 – 80.2014.
14. M.B. Etsuyankpa, C. E. Gimba, E.B. Agbaji and K.I. Omoniyi. "Determination of Cyanide and Saponins in the Products of Selected Cassava Varieties from Three Sampling Sites in Nigeria State." *CSN Book of Proceedings*. 1(36): 247-250. 2013a

M.B. Etsuyankpa, C.E. Gimba, E.B. Agbaji & K.I. Omoniyi, "Evaluation of Proximate Composition of 'Fufu', 'Lafun' and 'Gari' Through Microbial Fermentation of Four Varieties of Cassava." *CSN Book of Proceedings*. 1(36): 251-256. 2013b.
15. L.,T. Kigigha, P. Nyenke, S.C. Izah: "Health risk assessment of selected heavy metals in garri (cassava flakes) sold in major markets in Yenegoa metropolis, Nigeria." *Moi Toxicology* 4(12); 47-52. 2018

16. A.A. Audu, M. Waziri, T.T. Olasinde. "Proximate analysis and levels of some heavy metals in cassava flour processed by roadside drying along Abuja-Lokoja Highway, Nigeria." *IJLS* 2(3);55-58. 2012.
17. D.A. Cataldo, R.E. Wilding, "Soil and plant factors influencing the accumulation of heavy metals by plants." *EHP* 27;149-159. 1978.
18. FAO/WHO: "Adverse health effects of heavy metals in Human. Human Health for health sector" Geneva 77p.
19. World Health Organisation: "Heavy Metals, Safety evaluation of certain food additives and contaminants. 55 meeting of the joint FAO/WHO. Expert committee in food additives, Geneva, WHO food additives series," Geneva, 46p. 2014
20. G. Nordberg, "Excursions of intake above ADI: case study on cadmium." *Regul Toxicol Pharmacol.* 30:57-62. 1999.
21. I.O. Ogbonna, B.I. Agbowu, Agbo, "Proximate composition, microbiological safety and heavy metal contaminations of garri sold in Benue, North-central, Nigeria." *Afr. J. Biotechnol.* 16(18); 1085-1091. 2017.;
22. J.E. Emurotu, U.M. Salehedeem, O.M. Ayeni, "Assessment of heavy metal in cassava flour sold in Anyiba market, Kogi State, Nigeria." *J. Adv. Sci. Res.* 3(5): 2444-2448. 2012

This page intentionally left

Mini Solar Power Resources for IoT System in the Vannamei Shrimp Pond Model

Damar Widjaja^{1*}, Yohanes Priyanto Seli Laka¹

¹*Faculty of Science and Technology, Sanata Dharma Universit, Campus
III Paingan Maguwoharjo Depok Sleman DIY, 55282, Indonesia*

**Corresponding Author: damar@usd.ac.id*

(Received 19-11-2025; Revised 03-06-2025; Accepted 10-06-2025)

Abstract

Shrimp pond water quality is one of the challenging issues for shrimp farmer to keep or even increase their production level to meet the domestic and export needs. Many researches have been done to help shrimp farmer to manage water quality using Internet of Things (IoT) technologies. In order to make all the devices in the IoT system work properly, it needs an adequate power supply. Mini solar power plant installation is an alternative way to give shrimp farmers an electricity power access when their area has no electricity power network. In this study, we propose the use of mini solar power plants to supply the power to IoT devices in shrimp pond model. There are two main sub-system in this model, i.e. power supply sub-system and IoT based monitoring and controlling sub-system. Power supply sub-system consists of solar panel, solar cell controller, battery, INA219 sensor, and LM2596 step down IC. IoT sub-system perform monitoring and controlling on two shrimp pond models with Arduino Mega microcontroller works as the main processor. The mini solar power system works well as it was designed. Mini solar power plant capable of charging the 12V 40Ah battery in 5 hours. In order to make the IoT system works, it only needs 75,6 Wh from 480Wh battery capacity.

Keywords: Electricity, Solar Cell, Shrimp Pond, IoT.

1 Introduction

Shrimp is one of the export commodities in fisheries sub sector in Indonesia that has high economic value [1]. Indonesia exports many types of shrimps, including vannamei shrimp and tiger prawns. Vannamei shrimp contribute 36% of fisheries export commodities [2]. Several main export destination countries for Indonesian vannamei shrimp are Malaysia, United States, Great Britain, Japan, and China [3].



Shrimp farmer in Indonesia need to keep or even increase their production level to meet the domestic and export needs. Water is a natural resource that is very important for the survival of shrimp, so that they can live healthily and grow optimally. Shrimp pond water quality is one of the challenging issues for shrimp farmer as well as water quality management. There are several parameters used in water quality management, such as temperature, salinity, dissolved oxygen, pH, and alkalinity [4].

Many studies have been developed to help shrimp farmer to manage water quality using Internet of Things (IoT) technologies. Some of them try to control and monitor turbidity and ammonia [4], salinity level [5,6], pH level and water temperature [6,7], build an aeration monitor and a highly precise moving feeder [8], control water quality using fuzzy logic [9], and implement smart feeding system [4,6,10].

In order to make all the devices and IoT system work properly, it needs an adequate power supply. It has no problem in providing electricity to power up any kind of system in urban area or even sub urban area that has farming culture society. These types of areas have good access of electricity covered by electricity company. However, some of shrimp ponds in Yogyakarta province, Indonesia, especially in coast area in Kulon Progo City has no access to electricity power. This area needs an alternative energy that can be provided independently by the community, such as mini solar power plant.

Some studies related to solar power plant to provide electricity for shrimp pond have been done recently. Nizar Amir et al. studied the energy consumption of paddle wheel aeration and air blower as a circular technology for shrimp pond [11]. The result of this study is that the configuration of a solar photovoltaic meet the needs of electrical power for circular technology in shrimp pond. This configuration also reduces emissions of carbon dioxide, sulfur dioxide and nitrogen oxide per year.

Another study related to solar energy consumption for shrimp pond paddle wheel has been conducted by Aljurfri et al. [12]. The result of this study is that the power system using solar panel capable to support paddle wheel operation for 35 minutes. Idris and Thaha have designed solar power plant as an aerator driver in shrimp pond in Pinrang city [13]. The design uses Matlab to calculate the amount of electrical load, battery capacity requirement, number of solar modules, and inverter capacity.

Moreover, solar radiation and air temperature in the pond area has been measured.

The aforementioned studies above focus on providing electrical power for wheel paddle or aerator driver in shrimp pond. In this study, we propose the use of mini solar power plants to power IoT devices in shrimp pond model. This IoT monitoring system will provide data related to water quality of shrimp pond, such as temperature, pH level, dissolve oxygen, salinity, and control automatic feeder sub-system. This study aims to provide efficient power supply for simple IoT based monitoring and controlling system and evaluate the adequacy of the electricity power provided by mini solar power plant in shrimp pond model.

2 Material and Methods

Model system for this study is depicted in Fig. 1. There are two main sub-system in this model, i.e. power supply sub-system and IoT based monitoring and controlling sub-system. Power supply sub-system consists of solar panel, solar cell controller, battery, INA219 sensor, and LM2596 step down IC. IoT sub-system perform monitoring and controlling on two shrimp pond models with Arduino Mega microcontroller works as the main processor. All the monitoring and controlling data from Arduino Mega then are transmitted to Blynk IoT Platform via ESP32 microcontroller as a transmitting device.

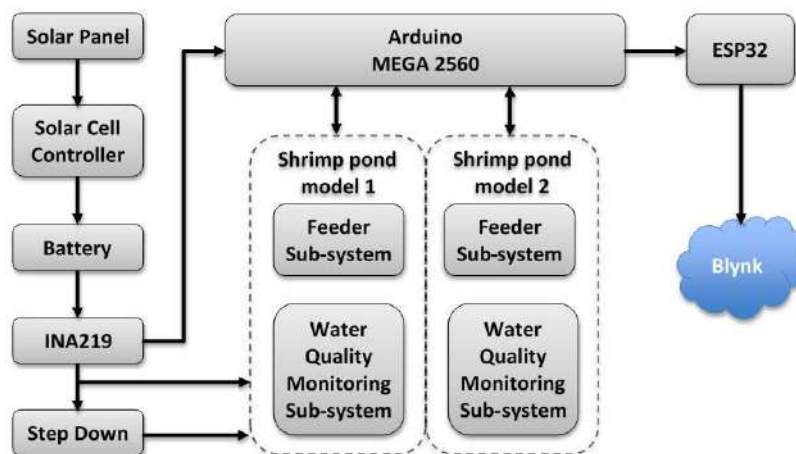


Figure 1. Model system.

Solar Cell Controller (SCC) regulates the voltage and current flowing from solar panels to a battery. The INA219 sensor measure both voltage and current on the whole system. A step-down device, LM 2596, functions as a step-down (buck) switching regulator, which converts a higher voltage (12V) to a lower voltage (5V). The IoT system needs those two voltage levels.

Two shrimp pond models use two aquariums with dimension of 30 cm weight x 35 cm height x 60 cm long. There are some water quality related sensors installed in each aquarium; DS18B20 water temperature sensor, SKU SEN 0237 dissolve oxygen sensor, SEN0161 pH sensor, and TDS salinity sensor. Outside the aquarium, there are some actuators to maintain water quality, pellets feeder control, and local display (OLED) to monitor water condition on site. Shrimp pond model is shown in Fig. 2.

There are many DC Motors work on each aquarium. Those are used for aerators, water cooler and heater pump to normalize water temperature, acid and alkaline liquid pump to neutralize water pH condition, salty and bland water pump to neutralize salinity condition. Servo motor is used for controlling the moving part of pellets feeder,

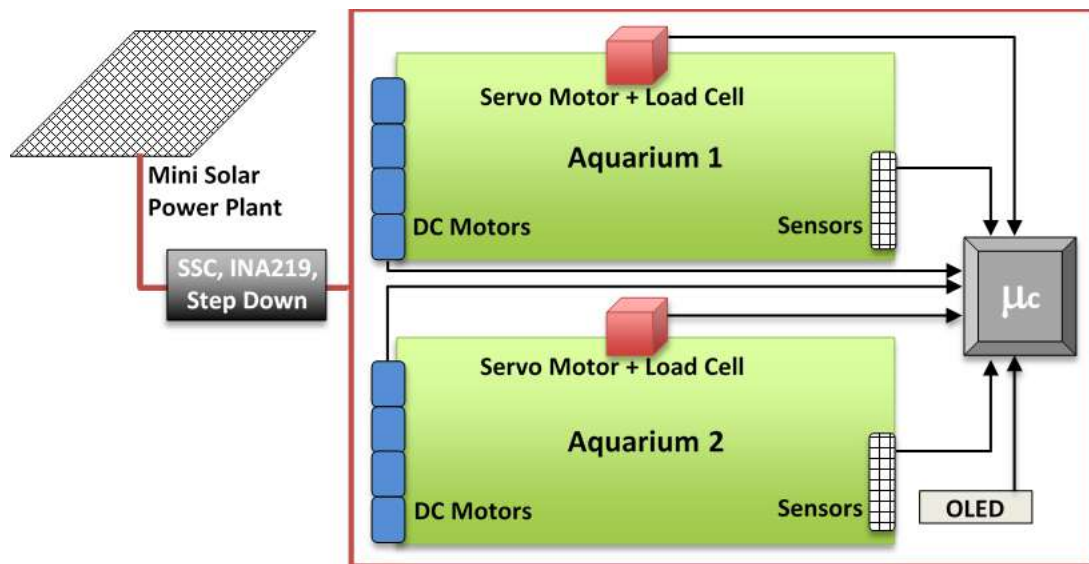


Figure 2. Aquariums and hardware configuration.

to open and close pellets feeder outlet. Load cell is used for measuring pellets weight inside the feeder container.

The power consumption specification of each component and total consumption needed based on the component specification is depicted in Table 1. The specification of each component refers to data specification sheet in [14], [15], [16], [17], [18], [19], [20], [21], [22], [23], [24], [25], and [26]. It can be seen in Table 1 that total power consumption that is needed for all active components is 63, 418 W. This total power consumption does not take into account power losses along cables (wiring between components) and other transmission media.

Table 1. Power consumption specification of each component.

No	Device	Power (W)	Number of Device	Total Power (W)
1	Arduino Mega	13,56	3	40,68
2	ESP 32	2,844	3	8,532
3	Sensor INA219	0,005	3	0,015
4	Sensor DS18B20	0,0055	2	0,011
5	Sensor SEN0237	0,1	2	0,2
6	Sensor SEN0161	1	2	2
7	Sensor TDS	0,03	2	0,06
8	Display OLED	0,06	2	0,12
9	Relay	0,15	8	1,2
10	Pompa DC	1,65	4	6,6
11	RTC DS3231	2,5E-06	1	2,5E-06
12	HX711 Load Cell	0,2	2	0,4
13	Motor servo MG996R	1,8	2	3,6
Grand Total				63,418

3 Results and Discussions

The first test has been carried out to find out how much voltage, current, and power the solar panel produces as a source of electricity. The power from solar panel then transferred to battery 12V 40 Ah for charging mechanism. Voltage, current, and power were measured during 5 hours of battery charging. The result of these measurement is shown in Figs. 3-5.

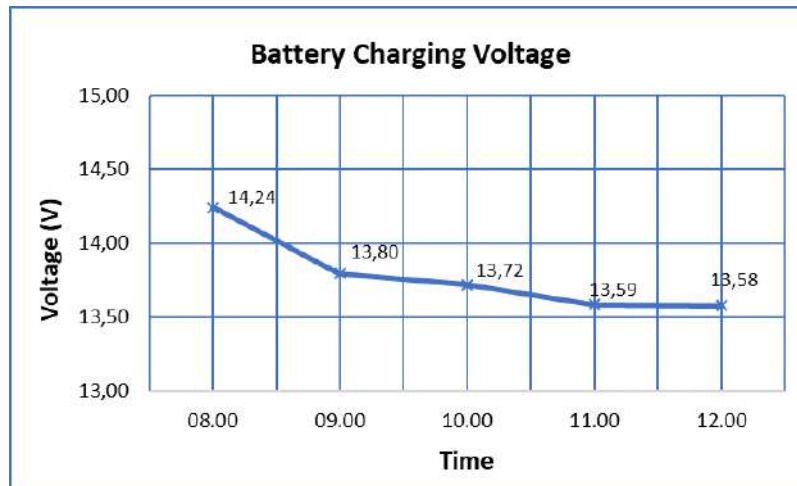


Figure 3. Battery charging voltage in 5 hours.

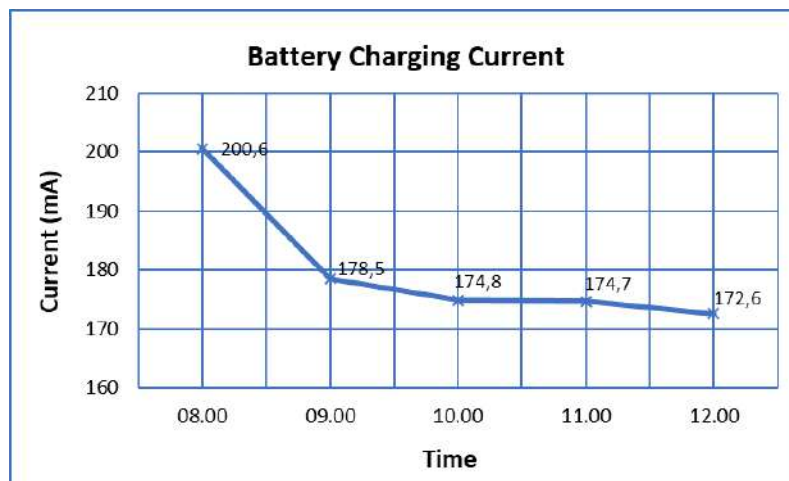


Figure 4. Battery charging current in 5 hours.

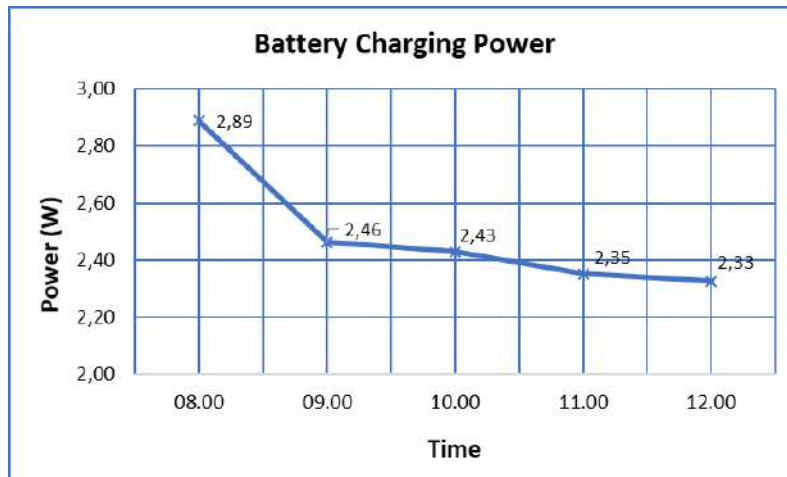


Figure 5. Battery charging power in 5 hours.

This INA219 sensor test was carried out from 08:00 to 12:00 when the sunlight was in the hottest condition. As the battery condition is getting fully charged, the voltage, current, and power transferred also getting lower. The decline in electrical parameters in the first hour is quite steep. However, in the next 4 hours, the decline in electrical parameters became sloping.

Sensor testing (INA 219 testing) to find out the needs of the system was done by dividing the system into 3 sub-systems, so that the system has three INA219 sensor. The first INA219 sensor was used to measure electricity parameters for pellets feeder and pH controlling and monitoring sub-system. The second INA219 sensor was used to measure electricity parameters for dissolved oxygen and temperature controlling and monitoring sub-system. The last INA219 sensor was used to measure electricity parameters for salinity controlling and monitoring sub-system and the need of mini solar power plant itself. Total electricity parameters that were needed for the IoT system is shown in Table 2.

Tabel 2 shows that total power consumption is far below battery capacity. Battery capacity is 480 Wh, while IoT system consumption is only 75,6 Wh. It means that the designed mini solar power plant is more than enough to supply the power needed for the IoT system of the shrimp pond model.

Table 2. Electricity consumption for shrimp pond IoT system

Sub-system	Voltage (V)	Current (Ah)	Power (Wh)
Pellets feeder and water pH	5,7	1,77	10,2
Dissolved oxygen and water temperature	10,6	2,75	29,4
Water salinity and mini solar power plant	12	3	36
Total consumption	12	7,52	75,6

Compared to the total power needed for all components in Table 1, the power consumption that is measured in the system is bigger. If the system can maintain the power consumption in one hour, it can be written that system need 63,418 Wh, that is 12,182 Wh different compared to actual condition. This is due to power loss in wiring and other transmission media that cannot be calculated or predicted accurately.

This study presents the system model performance in laboratory scale. As a model, there are many simplifications in terms of system parameters and size of the pond. This model need to be enhanced with more complex parameters and more extensive pond is it will be applied in the actual shrimp pond. It need to take into account for the long-term system sustainability, the long-term durability of system components, such as the battery and sensors, especially in harsh aquatic environments.

4 Conclusions

The mini solar power system works well as it was designed. Mini solar power plant capable of charging the 12V 40Ah battery in 5 hours. Battery capacity is big enough to supply power to IoT controlling and monitoring system of shrimp pond model. In order to make the IoT system works, it only needs 75,6 Wh from 480Wh battery capacity.

References

- [1] Rakhmanda and N. Husnayain, *Building Food and Fisheries Sustainability: Create Fresh, Healthy and Quality Shrimp through Sustainable Shrimp Cultivation (Membangun Keberlanjutan Pangan dan Perikanan: Ciptakan Udang Segar, Sehat dan Berkualitas melalui Budidaya Udang Berkelanjutan)*. Yogyakarta: Forbil Institute, 2021.
- [2] K. Krisandini, *Market Potential for Vaname Shrimp in Indonesia (Potensi Pasar Udang Vaname di Indonesia)*, Yogyakarta: JALA, 2024. Available: <https://jala.tech/id/blog/industri-udang/potensi-pasar-udang-vaname> [Accessed Nov. 11, 2024]
- [3] K. F. Izdiyar, *Complete Prospects and Methods for Exporting Vaname Shrimp are Here! (Prospek dan Cara Lengkap Ekspor Udang Vaname Ada di Sini!)*, Bandung: eFishery, 2023. Available: <https://efishery.com/id/resources/ekspor-udang-vaname/> [Accessed Nov. 11, 2024]
- [4] Srinivas, K.S. et al, "Automatic Prawns Feeding and Monitoring System using Internet of Things (IoT)," *Journal of Engineering Sciences*, Vol 14, Issue 03, pp. 560-564, March 2023.
- [5] A. R. Hakimi, M. Rivai, and H. Pirngadi, "IoT LoRa based Pond Water Salinity Level Control and Monitoring System (Sistem Kontrol dan Monitor Kadar Salinitas Air Tambak Berbasis IoT LoRa)" *JURNAL TEKNIK ITS*, Vol. 10, No. 1, pp. A9 – A14, 2021.
- [6] F. L. Toruan and M Galina, "Internet of Things-based Automatic Feeder and Monitoring of Water Temperature, pH, and Salinity for Litopenaeus Vannamei Shrimp," *Jurnal ELTIKOM: Jurnal Teknik Elektro, Teknologi Informasi dan Komputer*, Vol. 7, Issue 1, page. 9-20, June 2023.
- [7] A. Setiawan, N. L. G. R. Juliasih, and W. A. Setiawan, "Application of Internet of Things (IoT) Technology to Traditional Shrimp Ponds in Sriminosari Village, East Lampung," *Jurnal Pengabdian kepada Masyarakat*, Vol.1, No. 2, pp. 107-113, 2019.
- [8] B. Waycott, *IoT Technology Fixes for Shrimp Farming in India include an Aerator Monitor and Highly Precise Moving Feeder*. Global Seafood Alliance, 2023.
- [9] M. Johan and Suharjito. "Smart Shrimp Farming using Internet of Things (IoT) and Fuzzy Logic," *ComTech: Computer, Mathematics and Engineering Applications*, Vol. 14, No. 2, pp. 83-99, December 2023.

-
- [10] I. Arditya, T. A. Setyastuti, F. Islamudin, and I. Dinata, "Design of Automatic Feeder for Shrimp Farming Based on Internet of Things Technology," *MECHTA: International Journal of Mechanical Engineering Technologies and Application*, Article No. 8, Vol. 02, No. 2, pp. 145 - 151, 2021.
- [11] N. Amir, A. Errami, and L. Seung-woo, "Technical, Economical, Environmental feasibility of Solar PV System for Sustainable Shrimp Aquaculture: A Case Study of a Circular Shrimp Pond in Indonesia," *2022 IEEE 8th Information Technology International Seminar (ITIS) Surabaya, Indonesia*, October 19 – 21, 2022, pp. 102 – 107
- [12] Aljufri, R. Syarlian, A. Abizar, and A. Setiawan, "Preliminary Design of Shrimp Pond Paddle Wheel Powered by Solar Energy," *Jurnal Polimesin*, Vol. 19, No. 1, pp. 1 – 6, Februari 2021.
- [13] A.R. Idris and S. Thaha, "Desain Sistem Pembangkit Listrik Tenaga Surya pada Tambak Udang sebagai Penggerak Aerator," *INTEK Jurnal Penelitian*, Vol. 6, No. 1, pp. 36-40, 2019.
- [14] Arduino Mega 2560 Rev 3 – User Manual. Available: <https://docs.arduino.cc/resources/datasheets/A000067-datasheet.pdf> [Accessed May. 30, 2025].
- [15] Current Consumption Measurement of Modules. Available: <https://docs-espressif-com.translate.goog/projects/esp-idf/en/stable/esp32/api-guides/current-consumption-measurement-modules.html> [Accessed May. 30, 2025]
- [16] INA219 26-V 12-bit I2C output digital power monitor. Available: <https://www.ti.com/product/INA219> [Accessed May. 30, 2025]
- [17] Digital Temperatur Sensor. Available: https://www-guilcor-com.translate.goog/710-capteur-de-temperature-numerique?_x_tr_sl=en&_x_tr_tl=id&_x_tr_hl=id&_x_tr_pto=sge#:~:text=This%20measuring%20instrument%20can%20be,mA%20during%20active%20temperature%20conversions. [Accessed May. 30, 2025]
- [18] Gravity Analog Dissolved Oxygen Sensor SKU SEN0237. Available: https://wiki.dfrobot.com/Gravity__Analog_Dissolved_Oxygen_Sensor_SKU_SEN0237 [Accessed May. 30, 2025]
- [19] PH meter SKU SEN0161. Available: https://wiki.dfrobot.com/PH_meter_SKU__SEN0161_ [Accessed May. 30, 2025]

- [20] TDS Sensor/Meter for Water Quality Total Dissolved Solids Grove. Available: <https://digiwarestore.com/id/sensor-other/tds-sensor-meter-for-water-quality-total-dissolved-solids-grove-296403.html> [Accessed May. 30, 2025]
- [21] OLED Display (0.96 in, 128x64, SSD1306, I2C) (100952). Available: <https://www.smart-prototyping.com/OLED-0.96inch-12864-display-module-blue.html> [Accessed May. 30, 2025]
- [22] USB Mini Submersible Water Pump 5V DC. Available: <https://digiwarestore.com/id/other-appliances/usb-mini-submersible-water-pump-5v-dc-713504.html> [Accessed May. 30, 2025]
- [23] Songle – relay module 5v 4 channel Arduino. Available: <https://elmechtechnology.com/product/songle-relay-module-5v-4-channel-arduino> [Accessed May. 30, 2025]
- [24] Modul RTC DS3231. Available: <https://arduino.rezaervani.com/2019/03/02/modul-rtc-ds3231/> [Accessed May. 30, 2025]
- [25] Arduino with Load Cell and HX711 Amplifier (Digital Scale). Available: <https://randomnerdtutorials.com/arduino-load-cell-hx711/> [Accessed May. 30, 2025]
- [26] Servo Motor Basics with Arduino. Available: <https://docs.arduino.cc/learn/electronics/servo-motors/> [Accessed May. 30, 2025]

This page intentionally left

Improving the Accuracy of Prediction of Dissolved Oxygen and Nitrate Level Using LSTM with K-Means Clustering and Spearman Analysis

Ika Arva Arshella^{1 *}, I Wayan Mustika², Prapto Nugroho²

¹ *Department of Electrical Engineering, Faculty of Science and Technology,
Sanata Dharma University, Yogyakarta, Indonesia*

² *Department of Electrical Engineering and Information Technology
Gadjah Mada University Yogyakarta, Indonesia*

**Corresponding Author: ika_arshella@usd.ac.id*

(Received 15-05-2025; Revised 06-07-2025; Accepted 08-07-2025)

Abstract

This study discusses how to prepare data properly before entering the learning process for prediction using Deep Learning (DL). Long Short-Term Memory (LSTM) is one of the DL methods that is often used for prediction because of its superiority in maintaining long-term information. Although LSTM has proven effective, there are issues related to low-quality data that can reduce prediction accuracy. This problem is important to discuss because accuracy is important in predicting a value while field conditions can reduce the quality of the data taken. Data merging based on the relationship of each data collection location using the Spearman analysis and the K-Means clustering method is used to improve data quality. The results of the study show that improving data quality by merging data using K-Means has been successfully applied to various dataset conditions. In this study, we used two types of datasets related to river water quality, namely Dissolved Oxygen (DO) concentration and Nitrate levels for our simulation. The first data set produced DO predictions for eight locations with an average $R^2 = 0.9998$, $MAE = 0.0007$, $MSE = 1,13 \times 10^{-6}$. The second data set produced nitrate predictions for ten locations with an average $R^2 = 0.7337$, $MAE = 0.0111$, $MSE = 0,00029$.

Keywords: Data Merging, K-Means, Long Short-Term Memory, Prediction of Dissolved Oxygen and Nitrate Levels, Spearman

1 Introduction

The pattern of an event can be analyzed and used to predict future events. The pattern is obtained by monitoring of several parameters periodically. This monitoring produces a series of temporal data which is often referred as time series data [1]. In this study, data on river content were used, as changes in water quality can be detrimental to

the surrounding ecosystem. Therefore, it is necessary to monitor water quality parameters periodically as an effort to maintain water quality control [2], [3].

The amount of monitoring data that is carried out continuously demands a computational approach that is able to handle the complexity of temporal patterns. Many Deep Learning (DL) models can be used to handle this [1]. Examples of DL computational methods are Convolution Neural Networks (CNN), Deep Belief Networks (DBN), Recurrent Neural Networks (RNN), Auto Encoder (AE), Autoregressive Integrated Moving Average (ARIMA). Long Short-Term Memory (LSTM) and many more. In this study, the LSTM method is used because it has advantages in its architecture that can learn long-term data without losing information. These advantages solve the problem of long-term dependency experienced by RNN [1], [2], [3], [4], [5], [6].

Even though the best computational model has been selected, accuracy still needs to be considered to measure the success of the model. Unnatural changes during data collection can occur and lead to reduced data quality and model performance [2]. Data with low quality will have a negative impact on the basis of decision making, analysis processes and/or predictions of future events [3]. One example of a condition that can negatively impact a model is overfitting. Overfitting is a condition where the model tends to learn the details of values and noise at the training stage, making it difficult to generalize to validation and testing stages [7].

Previous researchers have applied several methods to handle low-quality data. In Rangenatan's research [8], the quality and accuracy of predictions can be significantly improved by performing a good and proper data preprocessing stage. The preprocessing stage consists of data cleaning (outliers, noise, missing values), data integration (correlation analysis, identification), and data reduction (grouping, feature selection).

In the data preprocessing stage, researchers adopted several techniques that have been proven effective in several literatures. Outliers were identified and removed using the Interquartile Range (IQR) method [2], [6], [9]. Linear Interpolation (LI) was applied to handle missing or empty data [2], [5], [6]. The Moving Average (MA) method was used to reduce noise and smooth the data [5], [6]. The selection of this method was based on its application in the literature before the model was developed.

Model development was done by adding the amount of data based on clustering the two most influential attributes in the dataset for each location. The K-Means method is widely used for clustering [10]. Wulandari, et al. [11], used the K-Means method to group and evaluate student performance so that the learning process runs smoothly. Wu et al. [12] and Pangestu et al. [13] combined the use of K-Means and ARIMA. The use of K-Means in [12] to evaluate the level of total phosphorus match with other features. The results were able to reduce the Mean Average Error (MAE) by 44.59% and the Mean Squared Error (MSE) by 56.82% in the prediction process. Chormunge et al. [14] and Gao et al. [15] also use K-Means for optimizing data input by clustering data based on compatibility between the features and produces 90% accuracy. Although there was a decrease in errors, researchers emphasized that there were still obstacles due to low data quality.

The following research is close to the proposed research because it has the same approach. Wei et al. [4] combined water pore pressure data from several river depth points using Partitional Clustering Algorithm (PCA) and predicted using LSTM. The R-Squared (R^2) value increased by 12% compared to predictions without data merging. The use of larger data with additional features is recommended to improve model performance. Arshella et al. [6] used multidimensional input on LSTM in predicting Dissolved Oxygen (DO). This study used the same dataset as the proposed study. The selection of this method can increase the accuracy of water quality prediction for one week up to an R^2 value of 0.999 compared to one-dimensional input. Unfortunately, in the process, it experienced some overfitting and underfitting because the input data still had low quality.

Based on previous studies, data quality has been shown to have a significant impact on the prediction process. The addition of data with similar characteristics is expected to improve data quality with the result prediction accuracy can also increase. This study proposes clustering training data based on the location of water quality data collection using the K-Means clustering method using LSTM. The selection of attributes to use in the K-Means model is determined through the correlation of many attributes to the target attribute of prediction using Spearman method. The model is evaluated using MSE, MAE and R^2 .

2 Material and Methods

This study used two water quality datasets obtained from the Environmental Information Data Center, namely the DO level dataset [16] and the Nutrient dataset contained in water [17]. Both datasets are part of the Land Ocean Interaction Study (LOIS) project. The DO dataset used is similar to the dataset used in study [6]. Table. 1 is a view of dataset one. The dataset contains FID as the code of the river location where the data was taken. DATE as the time when the data was taken. Conductivity, pH, Temperature as attributes and DO as the target of the prediction. Data collection was carried out from 1994 to 1997 in locations scattered across the rivers; Swale at Catterick Bridge, Derwent at Bubwith, Aire at Beal Bridge, Trent at Cromwell Lock, Calder at Methley Bridge, Swale at Crakehill, Aire above Thwaite Mill Weir, Aire at Fleet Weir, Ouse at Skelton, Nidd at Hunsingore.

The second dataset is taken from the same website, but contains different water parameters as shown in Table 2. FID is the code of the river location where the data was taken. Date is the time when the data was taken. The amount of data owned in the second dataset only has a small dataset for each location. Datasets for major ions and nutrients in river were collected during the period 1993 to 1997. Locations scattered across the rivers; Swale at Catterick Bridge, Swale at Thornton Manor, Nidd at Skip Bridge, Wharfe at Tadcaster, Ouse at Clifton Bridge, Derwent at Bubwith, Aire at Beal Bridge, Don at Sprotborough Bridge, Trent at Cromwell Lock, Calder at Methley Bridge.

Table 1. DO level dataset

	FID	ID	SITE_NAME	DATE	pH	Conductivity	DO	Temperature	Battery
0	29033	S1	Swale at Catterick Bridge	27/06/1995 11:00	8.27	411.0	102.2	16.8	13.0
1	29033	S1	Swale at Catterick Bridge	27/06/1995 11:30	8.31	412.0	105.8	17.2	13.0
2	29033	S1	Swale at Catterick Bridge	27/06/1995 12:00	8.34	413.0	107.0	17.5	13.0
3	29033	S1	Swale at Catterick Bridge	27/06/1995 12:30	8.36	414.0	108.7	17.9	13.0
4	29033	S1	Swale at Catterick Bridge	27/06/1995 13:00	8.39	416.0	110.2	18.3	13.0
...	...	...	...	...	...	...	...	...	...
235115	1584377	N33	Nidd at Hunsingore	7/7/1994 13:00	7.68	404.0	58.6	17.1	12.4
235116	1584377	N33	Nidd at Hunsingore	7/7/1994 13:30	7.67	409.0	58.5	17.1	12.4
235117	1584377	N33	Nidd at Hunsingore	7/7/1994 14:00	7.67	410.0	58.0	17.2	12.4
235118	1584377	N33	Nidd at Hunsingore	7/7/1994 14:30	7.65	419.0	58.9	17.4	12.4
235119	1584377	N33	Nidd at Hunsingore	7/7/1994 15:00	7.67	411.0	59.1	17.3	12.4

235120 rows × 9 columns

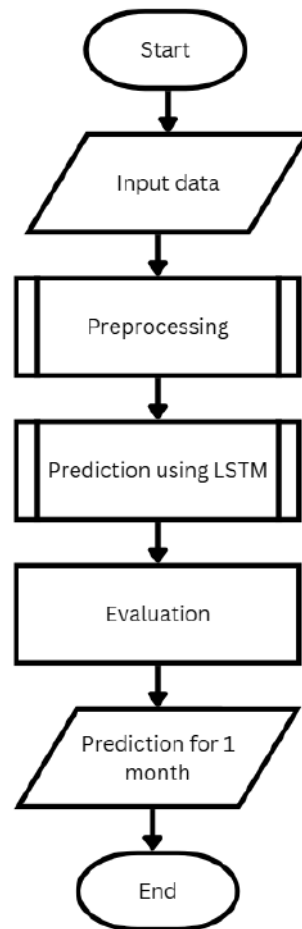
The ions and nutrients measured include Ammonia, Bromide-ion, Calcium Dissolved, Carbon Organic Dissolved, Carbon Organic Particulate, Chloride-ion, Magnesium Dissolved, Nitrate, Nitrite, Nitrogen Particulate, Phosphorus Soluble Reactive, Phosphorus Total, Phosphorus Total Dissolved, Potassium Dissolved, Silicate Reactive Dissolved, Sodium Dissolved, Sulphate. Sampling was carried out regularly every week.

The main flowchart in this study can be seen in Fig. 3. The core of the prediction process begins with loading the water quality dataset into the system, then the data is prepared through a preprocessing stage. At this stage, the data goes through the process of removing outliers and ensuring that there are no missing or empty values so the data used for training will have better quality. When the data is ready, the next step is to predict water quality for the next one-month period using the LSTM algorithm. The prediction results are then analyzed as part of the model evaluation process.

Table 2. Nitrate level dataset

	PID	ID	SETT_NAME	DATE	Calcium dissolved	Carbon organic dissolved	Carbon organic particulate	Chloride-ion	Magnesium dissolved	Nitrate	Nitrogen particulate	Phosphorus Total	Phosphorus total dissolved	Potassium dissolved	Silicate reactive dissolved	Sodium dissolved	Sulphate
0	29033	S1	Swale at Catharick Bridge	9/7/1993	62.0	2.200	0.488	17.5	7.8	0.677419	0.959	0.174	0.164	1.7	1.82	11.1	23.0
1	29033	S1	Swale at Catharick Bridge	9/14/1993	21.1	8.600	2.731	8.6	2.0	0.541936	0.987	0.102	0.047	0.8	2.57	5.6	10.0
2	29033	S1	Swale at Catharick Bridge	9/21/1993	45.8	4.040	0.492	13.0	4.6	1.038719	0.935	0.070	0.065	1.4	4.44	7.8	17.0
3	29033	S1	Swale at Catharick Bridge	9/28/1993	58.1	2.350	0.316	15.0	6.5	1.400000	0.937	0.097	0.052	1.5	4.63	9.1	20.0
4	29033	S1	Swale at Catharick Bridge	10/5/1993	40.0	7.500	0.686	13.0	4.2	0.858065	0.930	0.071	0.062	1.4	3.88	7.7	15.5
1715	29044	C9	Calder at Methley Bridge	11/19/1996	29.0	9.244	2.852	74.0	7.4	4.741940	0.159	0.732	0.005	8.4	7.41	57.0	70.0
1716	29044	C9	Calder at Methley Bridge	11/26/1996	25.0	8.423	6.060	52.0	6.3	2.800000	0.521	0.665	0.273	5.4	6.18	56.0	52.0
1717	29044	C9	Calder at Methley Bridge	12/3/1996	27.5	8.032	1.194	78.0	7.6	3.658069	0.168	0.685	0.439	6.9	7.54	55.0	70.0
1718	29044	C9	Calder at Methley Bridge	12/10/1996	36.5	7.444	0.967	87.0	10.5	5.419359	0.098	0.852	0.739	10.3	9.25	89.0	96.0
1719	29044	C9	Calder at Methley Bridge	12/17/1996	42.5	10.316	0.747	117.0	11.7	4.967740	0.969	1.104	0.911	14.9	9.53	95.0	135.0

1729 rows x 17 columns

**Figure 1.** Main flowchart of prediction process

Data Preprocessing

As shown in Figure 1, the first step is reading the input data. The next step is preprocessing, a stage carried out to enhance data quality by removing outliers and verifying the completeness of the dataset. Interquartile Range (IQR) is a method that can be used to identify and remove outliers. By using IQR, data can be selected based on the data area between the upper quartile, which is 75% of the data (Q_3) and the lower quartile, which is 25% of the data (Q_1) as can be seen in Equation (1). Data that is considered to be outliers is data whose value is less than the minimum limit and greater than the maximum limit. Determining the minimum limit using Equation (2) and the maximum limit using Equation (3):

$$IQR = Q_3 - Q_1, \quad (1)$$

$$\text{Minimum} = (Q_1 - 1.5 * IQR), \quad (2)$$

$$\text{Maximum} = (Q_3 + 1.5 * IQR). \quad (3)$$

Where:

IQR : Interquartile Range

Q3 : Upper quartile (75% data)

Q1 : Lower quartile (25% data)

After outliers are removed, the time series data needs to be completed so that there are no null/missing data. Linear Interpolation (LI) is a method used to estimate the value of a point between known data points using a straight line, also known as a linear polynomial. The LI method is widely used because it is simple. LI is used to fill in the gaps experienced by data due to loss and non-empty timing. The new value (y) that is in the middle of the data after (x_1, y_1) and the data before (x_0, y_0) can be calculated based on Equation 4 below:

$$y = y_0 + (x - x_0) \frac{y_1 - y_0}{x_1 - x_0}. \quad (4)$$

Where:

x : A point on the x axis is known, its y value is unknown

- y : The point being searched for, using the Equation (4)
 (x_1, y_1) : coordinates new data
 (x_0, y_0) : coordinates new data

From data that has been removed from outliers and filled in, the data is complete according to timing. The data is smoothed or reduced noise using a Moving Average (MA) so the pattern of the data can be displayed clearly and easier to analyze. Moving Averages, also known as running mean or rolling averages, is a special type of filtering method used to transform one time series into another time series. The moving average value for the next period Y_{t+1} is calculated by summing the data values (Y) in a window size (m) and dividing it by the window size value, as shown in the Equation 5;

$$Y_{t+1} = \frac{Y_t + Y_{t-1} + \dots + Y_{t-m+1}}{m}. \quad (5)$$

Where:

- Y_{t+1} : Moving Average value for the next period ($t + 1$)
 Y_t : Data on current time period (t)
 Y_{t-1} : Data before the current time period ($t - 1$)
 Y_{t-m+1} : Data at time period ($t - m + 1$), back by the window size from the current time period data.
 m : window size, amount of data to calculate the average

K-Means Clustering

The K-Means algorithm uses the proximity measure function to place each object in the cluster that is most similar to it. When updating the center of each cluster (centroid), the distance calculation is performed again and updates the centroid again. This iteration is carried out until the proximity measure function converges; in each cluster the objects no longer change. Iteration is used to divide data objects into several different clusters, so that the similarity between objects in one cluster becomes large, while the similarity between objects becomes small [15]. Calculating the distance between the centroid point and each object point, also called Euclidean distance (D_e), is as in the Equation 6:

$$D_e = \sqrt{(x_i - s_i)^2 + (y_i - t_i)^2}. \quad (6)$$

Where:

- D_e : Euclidean distance
 (x_i, y_i) : Object coordinate points
 (s_i, t_i) : Centroid coordinate points

Correlation Analysis

In the K-Means clustering process, distance calculations typically utilize two features. However, the dataset employed in this study contains more than two features, necessitating feature selection. The Spearman method is chosen for feature selection because it is widely used for correlation analysis between features. By using ranking, this method measures the tight relationship between 2 variables. In this case the relationship between the prediction target and other features is calculated using Equation 7:

$$\text{Correlation} = 1 - \frac{6 \times \sum(x - y)^2}{[n \times (n^2 - 1)]}. \quad (7)$$

The results of the Spearman correlation calculation have a value range of -1 to 1. If the result is close to 1, then the two variables have a positive correlation, while the result is close to -1, then the two variables have a negative correlation. If the correlation value is close to 0, then there is no correlation between the two variables [18].

Long Short-Term Memory

Backpropagation errors can lead to signal explosion or vanishing gradients, causing instability during the learning process. To address this issue, Long Short-Term Memory (LSTM) networks are used, as they effectively manage error flow through their gated architecture [19]. The LSTM system consists of three inputs: the current input (x_t), the output of the previous unit (h_{t-1}), and the memory of the previous unit (c_{t-1}).

Figure 2 illustrates the LSTM unit and highlights the three gates that play an important role in the memory calculation: the forget gate, the input gate, and the output gate [2]. The following are the equations used in the LSTM model:

$$f_t = \sigma(W_f[h_{t-1}, x_t] + b_f) \in (0,1)^h \quad (8)$$

$$i_t = \sigma(W_i[h_{t-1}, x_t] + b_i) \in (0,1)^h \quad (9)$$

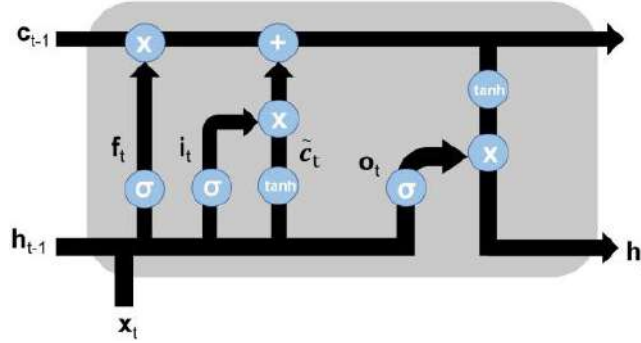


Figure 2. LSTM unit

$$\tilde{c}_t = \tanh(W_c[h_{t-1}, x_t] + b_c) \in (-1,1)^h \quad (10)$$

$$o_t = \sigma(W_o[h_{t-1}, x_t] + b_o) \in (0,1)^h \quad (11)$$

$$C_t = (f_t \odot C_{t-1}) + (i_t \odot \tilde{c}_t) \in \mathbb{R}^h \quad (12)$$

$$h_t = o_t \odot \tanh(C_t) \in \mathbb{R}^h \quad (13)$$

Where:

- f_t : Forget gate determinant
- i_t : Input gate determinant
- $\tilde{c}_t$: Determine the memory to be used
- o_t : Output gate determinant
- C_t : New Memory Determinant
- h_t : New Output

W_f, W_i, W_c, W_o : Weight matrix

b_f, b_i, b_c, b_o : Bias vector

σ : Sigmoid function

$\tanh$: Hypertangen function

Equation (8) is used to calculate the weight vector for the forget gate. It produces an output between 0 and 1 by applying a sigmoid function to the sum of the previous hidden units (h_{t-1}) and the current input (x_t), along with the bias vector for the forget

gate (b_f). Similarly, Equation (9) uses the weights and the let vector for the input gate. Equation (10) is used to compute the candidate values to be retained. It produces an output between -1 and 1 by applying the hyperbolic tangent function to the weighted sum of the previous hidden state (h_{t-1}), and the current input (x_t), along with a bias vector for the carrier (b_c). Equation (11) is used to compute the output gate weight vector. It produces the same output as (8) and (9) by using the sigmoid function with output weight matrix (W_o) and output bias vector (b_o).

Equation (12) is used to calculate the memory cell by adding the elementwise multiplication of Equation (8) and the memory of the previous cell (c_{t-1}) with the elementwise multiplication of Equations (9) and (10). Finally, Equation (13) is used to determine the output to be forwarded to the next LSTM cell by multiplying the elementwise multiplication of Equation (11) with the hyperbolic tangent of Equation (12).

Evaluation

A comprehensive evaluation is required to assess the performance of the proposed method. The following evaluation metrics used are mentioned below [20]:

a. Mean Average Error (MAE)

MAE is used in regression analysis to measure the average absolute difference between predicted values (X) and actual values (Y) as in Equation (14):

$$\text{MAE} = \frac{1}{m} \sum_{i=1}^m |X_i - Y_i|. \quad (14)$$

Where:

X_i : The i -th predicted value

Y_i : The i -th actual value

m : total data amount

MAE values range from 0 to 1. The closer to 0 the better the results, illustrating that the deviation between prediction and actual is small.

b. Mean Square Error (MSE)

MSE is used in regression analysis to measure the average squared difference between predicted values (X) and actual values (Y) as in Equation (15):

$$MSE = \frac{1}{m} \sum_{i=1}^m (X_i - Y_i)^2. \quad (15)$$

MSE values range from 0 to 1. The closer to 0 the better the results, illustrating that the deviation between prediction and actual is small.

c. R-Squared (R^2)

R^2 is a statistical measure obtained by dividing the mean of the squared differences between the predicted and actual values (MSE) by the total variance of the dependent variable as shown in Equation (16):

$$R^2 = 1 - \frac{\sum_{i=1}^m (X_i - Y_i)^2}{\sum_{i=1}^m (\bar{Y} - Y_i)^2}. \quad (16)$$

The value of R^2 ranges from $-\infty$ to 1. The closer to 1, the better the result.

3 Results and Discussion

The proposed method was evaluated using two distinct water quality datasets featuring different characteristics and indicators. The analysis is presented in two parts: (1) DO level prediction and (2) nitrate level prediction. Each section systematically compares the performance of the proposed method against baseline approaches.

The result of DO level predictions

The result of DO level predictions from one of the locations can be seen in Fig. 3. The data selected as the visualization of the results comes from the Calder at Methley Bridge with FID 29044. The figure consists of 2 types of graphs, namely loss in training and validation, and a graph of the model's predicted values for one month. Model testing is done with a different preprocessing process. 1) Non-C (nonc), the data used is not merged with data from other locations. 2) Spearman (sp), the data used is merged based on DO correlation between river locations using Spearman method. 3) K-Means (km), the proposed method uses data merging based on clustering attribute data from various locations.

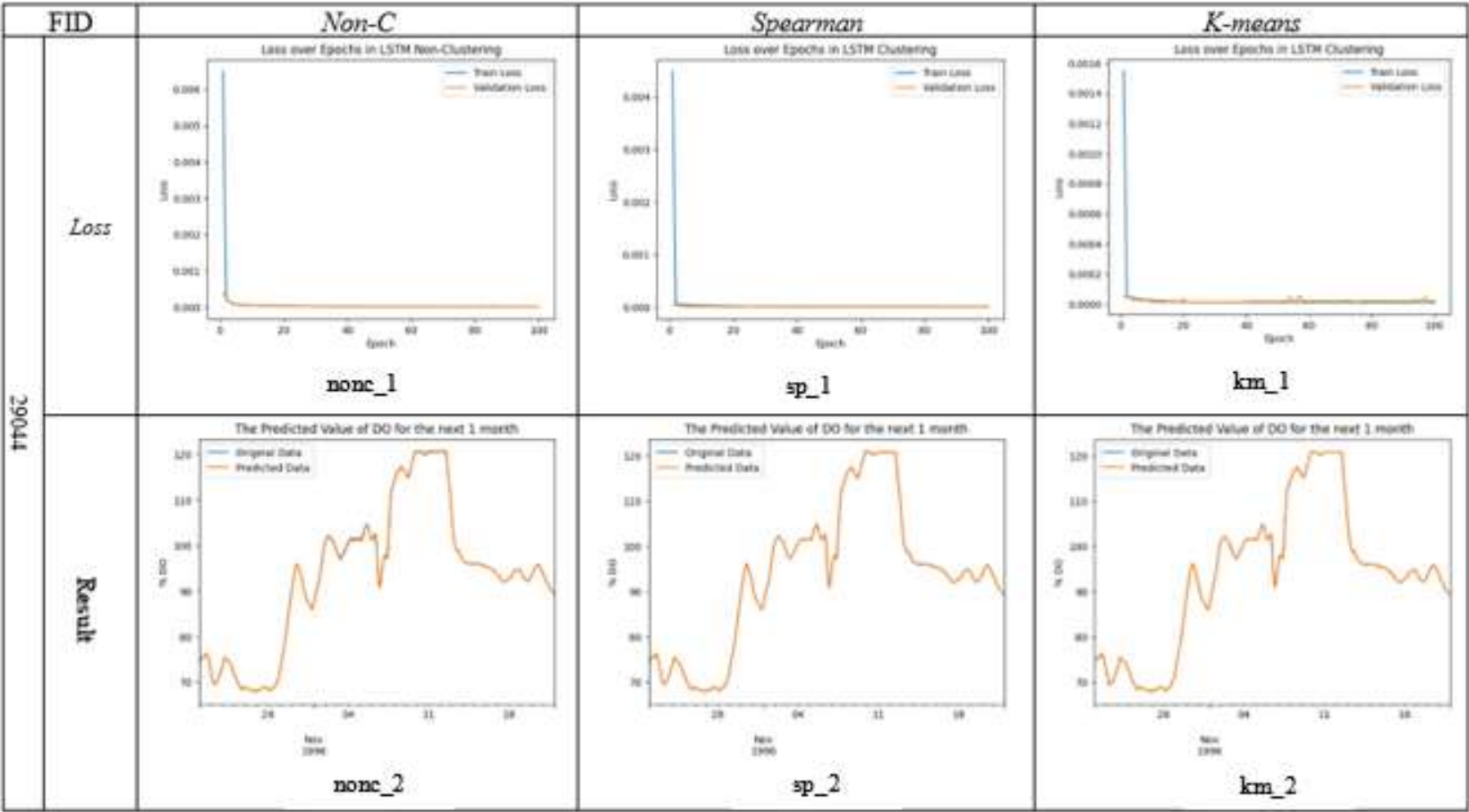


Figure 3. Training and validation loss graphs and DO level prediction results from FID 29044

In Fig. 3(nonc_1, sp_1, km_1), the training and validation loss curves show stable convergence with no sign of overfitting (no significant gap between training and validation loss) or underfitting (loss reaches very low values). All models managed to achieve a loss below 0.001, indicating effective learning. The proposed method (km_1) recorded the lowest loss (less than 0.0002), significantly superior to the nonc and sp baselines. In Fig. 3 (nonc_2, sp_2, km_2), the predicted values of all models almost overlap with the actual values, confirming high accuracy.

Despite the fact that all models appear accurate visually, significant differences are seen in quantitative metrics MAE, MSE and R^2 as can be seen in Table 3. The proposed method reduces MAE for predicting DO in FID 29044 by 61.875% compared to the baseline (nonc), with an absolute value of 0.00061 vs 0.0016. The proposed method reduces MSE by 78.18% compared to the baseline (nonc), with an absolute value of $1,2 \times 10^{-6}$ vs $5,5 \times 10^{-6}$. When compared to sp, both (sp and km) have more or less equally good results. This indicates that the application of K-Means for merging data is able to improve model prediction.

The proposed method can produce better values than the comparison method when observed from the average value from various locations. The proposed method produces an average $R^2 = 0.9998$, MAE = 0.0007, MSE = $1,13 \times 10^{-6}$. Although statistical significance testing was not performed due to limited data replication, the large

Table 3. Evaluation of the results of the prediction of dissolved oxygen for one month

FID	R^2			MAE			MSE		
	nonc	sp	km	nonc	sp	km	nonc	sp	km
29033	0,9996	0,9995	0,9994	0,0016	0,0014	0,0005	$1,03 \times 10^{-5}$	$6,73 \times 10^{-6}$	$1,13 \times 10^{-6}$
29040	0,9999	0,9999	0,9999	0,0012	0,0018	0,0005	$2,42 \times 10^{-6}$	$3,97 \times 10^{-6}$	$4,10 \times 10^{-7}$
29041	0,9997	0,9994	0,9999	0,0014	0,0016	0,0005	$3,63 \times 10^{-6}$	$2,94 \times 10^{-6}$	$4,87 \times 10^{-7}$
29043	0,9999	0,9999	0,9999	0,0014	0,0007	0,0011	$4,81 \times 10^{-6}$	$1,53 \times 10^{-6}$	$1,94 \times 10^{-6}$
29044*	0,9998	0,9999	0,9999	0,0016	0,0007	0,0006	$5,5 \times 10^{-6}$	$1,1 \times 10^{-6}$	$1,2 \times 10^{-6}$
1552046	0,9999	0,9998	0,9998	0,0019	0,0008	0,0007	$6,67 \times 10^{-6}$	$1,17 \times 10^{-6}$	$7,28 \times 10^{-7}$
1552048	0,9991	0,9999	0,9995	0,0024	0,00024	0,0013	$6,34 \times 10^{-6}$	$1,25 \times 10^{-7}$	$2,39 \times 10^{-6}$
1584376	0,9999	0,9998	0,9999	0,0009	0,0013	0,0006	$3,18 \times 10^{-6}$	$2,92 \times 10^{-6}$	$7,76 \times 10^{-7}$
Average	0,9997	0,9997	0,9998	0,0016	0,0011	0,0007	$5,36 \times 10^{-6}$	$2,56 \times 10^{-6}$	$1,13 \times 10^{-6}$

differences and stability of the results support the potential superiority of this method. Further research is needed with more diverse samples, cross-validation, and statistical testing.

The result of nitrate level predictions

The result of nitrate level predictions from one of the locations can be seen in Fig. 4. The data selected as the visualization of the results comes from the Wharfe at Tadcaster with FID 29038. Same as before, the figure consists of 2 types of graphs, namely loss in training and validation, and a graph of model prediction values for one month. Model testing is carried out with different preprocessing processes. 1) Non-C (nonc), 2) Spearman (sp), and 3) K-Means (km). The proposed method uses data merging based on attribute data clustering from various locations.

The gap between training and validation loss in Fig. 4 (nonc_3) suggests overfitting, likely due to the model's high complexity relative to the dataset size. Figure 6 sp_3, km_3, the training and validation loss curves show stable convergence with no sign of overfitting (no significant gap between training and validation loss) or underfitting (loss reaches very low values). All models managed to achieve a loss below 0.005, indicating effective learning. In Fig. 4 (nonc_4, sp_4, km_4), the predicted values of all models almost overlap with the actual values, confirming high accuracy.

Despite the fact that all models appear almost accurate visually, significant differences are seen in quantitative metrics MAE, MSE and R^2 as can be seen in Table 4. The proposed method reduces MAE by 78,5% compared to the baseline (nonc), with an absolute value of 0.016 vs 0.078. The proposed method reduces MSE by 93,7% compared to the baseline (nonc), with an absolute value of 0,0078 vs 0,0005. When compared to sp, both (sp and km) have more or less the same good results but merging using sp correlation tends to be better. This shows that the application of K-Means for nutrient data merging can be used but there is an opportunity to use other methods.

The proposed method can produce better values than the comparison method when observed from the average value from various locations. The proposed method produces an average $R^2 = 0.7337$, MAE = 0.0111, MSE = 0,00029. This small average is due to the negative R^2 result in one location. Although statistical significance testing was not performed due to limited data replication, the large differences and stability of the

Table 4. Evaluation of the results of the prediction of nitrate for one month

FID	R ²			MAE			MSE		
	nonc	sp	km	nonc	sp	km	Nonc	sp	km
29033	-10,4897	-1,07	-1,07	0,0401	0,0056	0,0056	$2,40 \times 10^{-3}$	$3,99 \times 10^{-5}$	$3,99 \times 10^{-5}$
29034	-0,5825	0,8807	0,9232	0,0675	0,0083	0,0057	0,0062	0,0002	$1,37 \times 10^{-4}$
29037	0,9417	0,9880	0,9903	0,0359	0,0041	0,0068	0,002	$4,47 \times 10^{-5}$	$8,64 \times 10^{-5}$
29038*	0,869	0,954	0,958	0,078	0,012	0,016	0,0078	0,0002	0,0005
29039	0,9469	0,9853	0,987	0,029	0,0120	0,0106	0,0012	0,0002	$1,74 \times 10^{-4}$
29040	0,7597	0,9717	0,9822	0,0605	0,019	0,0133	0,0046	0,0004	0,0006
29041	0,8325	0,8981	0,9491	0,0178	0,019	0,013	0,0008	0,0005	0,0002
29042	0,7382	0,886	0,8804	0,0646	0,018	0,0166	0,005	0,0007	0,0007
29043	-1,329	0,648	0,6518	0,0455	0,0087	0,007	0,0025	$9,28 \times 10^{-5}$	$9,18 \times 10^{-5}$
29044	0,648	0,916	0,9051	0,0425	0,012	0,0108	0,003	0,0002	0,0002
Average	-0,666	0,706	0,7337	0,0481	0,0119	0,0111	0,00355	0,00026	0,00029

results support the potential superiority of this method. Further research is needed with more diverse samples, cross-validation, and statistical testing.

In addition to comparing with non-c and sp, comparisons were also made to previous studies. The study conducted by Arshella [6] used the same DO dataset, but the proposed method can overcome the problem of overfitting and can predict for a longer time than previous studies. A similar study was conducted by Wu [12], using K-Means to identify random patterns of water parameters and ARIMA for prediction. Although using a different dataset, the study is still about water quality parameters and using the K-Means method. The method proposed in the study can reduce prediction errors more effectively by overcoming the main limitations in previous studies through preprocessing techniques to improve data quality, as well as combining more advanced learning methods (LSTM) for long-term data analysis. Although successful, this approach still opens up opportunities for future improvements, such as validation on larger or more diverse datasets, exploration of real-time applications, cross-validation, and statistical testing.

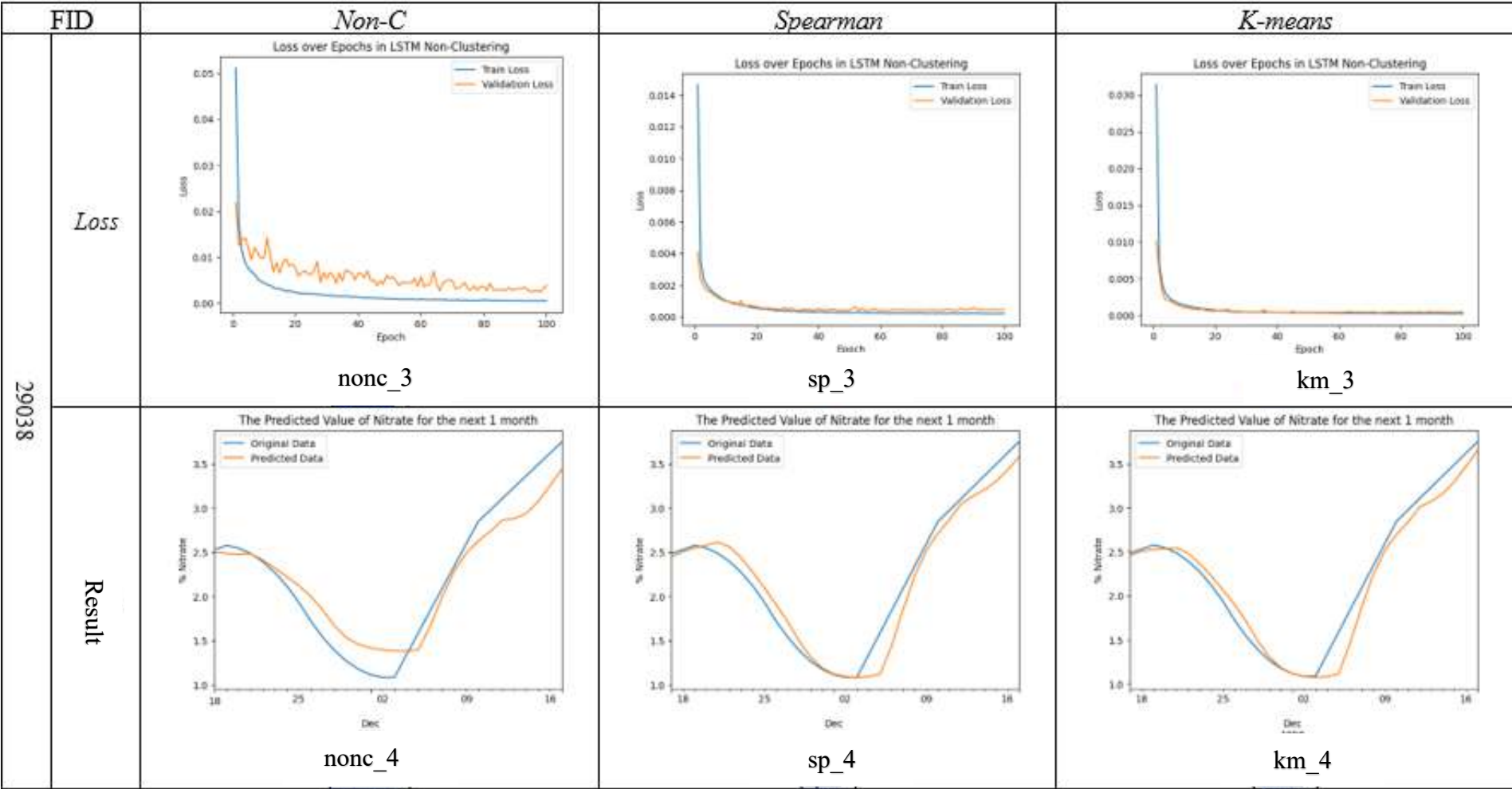


Figure 4. Training and validation loss graphs and nitrate level prediction results from FID 29038

4 Conclusions

The accuracy of water quality prediction can be improved by improving data quality. The application of Spearman method for decide which attributes or parameters need to be used for K-Means, K-Means methods for data merging and LSTM at the prediction stage has a positive impact on increasing accuracy. In this study, we implemented two datasets. The first data set produced DO predictions with an average $R^2 = 0.9998$, $MAE = 0.0007$, $MSE = 1,13 \times 10^{-6}$. The second data set produced nitrate level predictions with an average $R^2 = 0.7337$, $MAE = 0.0111$, $MSE = 0,00029$. Data merging based on location can be applied to various dataset conditions; datasets with few features and many data and datasets with few data and many features. The selection of the right grouping method should also be further evaluated to ensure optimal results. These steps will be an important part of the development and assessment of prediction methods, especially for water quality in the future. Other suggestions for improvement are the use of larger or more diverse datasets, real-time implementation, cross validation, and statistical testing.

Acknowledgements

The author would like to thank the supervisor from Gadjah Mada University who has guided this research so that this research can be completed.

References

- [1] H. Zhongyang, Z. Jun, L. Henry, F. M. King, and W. Wei, "A review of deep learning models for time series prediction," *IEEE Sensor Journal*, vol. 21, No. 6, Mar. 2021, doi: 10.1109/JSEN.2019.2923982.
- [2] H. Chen *et al.*, "Water quality prediction based on LSTM and attention mechanism: A case study of the Burnett River, Australia," *Sustainability (Switzerland)*, vol. 14, no. 20, Oct. 2022, doi: 10.3390/su142013231.
- [3] A. Docheshmeh Gorgij, G. Askari, A. A. Taghipour, M. Jami, and M. Mirfardi, "Spatiotemporal forecasting of the groundwater quality for irrigation purposes, using deep learning Method: Long short-term memory (LSTM)," *Agric Water Manag*, vol. 277, Mar. 2023, doi: 10.1016/j.agwat.2022.108088.

- [4] C. W. W. Ng, M. Usman, and H. Guo, "Spatiotemporal pore-water pressure prediction using multi-input long short-term memory," *Eng Geol*, vol. 322, Sep. 2023, doi: 10.1016/j.enggeo.2023.107194.
- [5] Z. Hu *et al.*, "A water quality prediction method based on the deep LSTM network considering correlation in smart mariculture," *Sensors (Switzerland)*, vol. 19, no. 6, Mar. 2019, doi: 10.3390/s19061420.
- [6] Arshella. Ika Arva, I. W. Mustika, and P. Nugroho, "Water quality prediction based on machine learning using multidimension input LSTM," *IEEE*, Aug. 2023. doi: 10.1109/ICITACEE58587.2023.10276970.
- [7] M. Del Giudice, "The prediction-explanation fallacy: A pervasive problem in scientific applications of machine learning," *Methodology*, vol. 20, no. 1, pp. 22–46, 2024, doi: 10.5964/meth.11235.
- [8] R. G, "A study to find facts behind preprocessing on deep learning algorithms," *Journal of Innovative Image Processing*, vol. 3, no. 1, pp. 66–74, Apr. 2021, doi: 10.36548/jiip.2021.1.006.
- [9] D. Dheda, L. Cheng, and A. M. Abu-Mahfouz, "Long short-term memory water quality predictive model discrepancy mitigation through genetic algorithm optimisation and ensemble modeling," *IEEE Access*, vol. 10, pp. 24638–24658, Feb. 2022, doi: 10.1109/ACCESS.2022.3152818.
- [10] M. G. H. Omran, A. P. Engelbrecht, and A. Salman, "An overview of clustering methods," 2007, *IOS Press*. doi: 10.3233/ida-2007-11602.
- [11] N. H. Wulandari and V. Purwayoga, "Cluster change analysis to assess the effectiveness of speaking skill techniques using machine learning," *International Journal of Applied Sciences and Smart Technologies*, vol. 7, no. 1, pp. 1–14, 2025, doi: 10.24071/ijasst.v7i1.9667.

-
- [12] J. Wu *et al.*, “Application of time serial model in water quality predicting,” *Computers, Materials and Continua*, vol. 74, no. 1, pp. 67–82, 2023, doi: 10.32604/cmc.2023.030703.
- [13] P. Pangestu, S. Maarip, Y. N. Addinsyah, and V. Purwayoga, “Clustering and trend analysis of priority commodities in the archipelago capital region (IKN) using a data mining approach,” *International Journal of Applied Sciences and Smart Technologies*, vol. 6, no. 1, pp. 169–182, 2024, doi: 10.24071/ijasst.v6i1.7798.
- [14] S. Chormunge and S. Jena, “Correlation based feature selection with clustering for high dimensional data,” *Journal of Electrical Systems and Information Technology*, vol. 5, no. 3, pp. 542–549, Dec. 2018, doi: 10.1016/j.jesit.2017.06.004.
- [15] G. Qiang, X. Hong Xia, H. Hong Gui, and G. Min, “Soft sensor method for surface water qualities based on fuzzy neural network,” *IEEE*, Jul. 2019. doi: 10.23919/ChiCC.2019.8866494.
- [16] D. Leach, A. Pinder, P. Wass, N. Bachiller-Jareno, I. Tindall, and R. Moore, “Continuous Measurements of Temperature, pH, Conductivity and Dissolved Oxygen in Rivers [LOIS],” NERC Environmental Information Data Centre. Accessed: Aug. 01, 2023. [Online]. Available: <https://doi.org/10.5285/b8a985f5-30b5-4234-9a62-03de60bf31f7>
- [17] D. Leach, M. Neal, N. Bachiller-Jareno, I. Tindall, and R. Moore, “Major ion and nutrient data from rivers [LOIS],” NERC Environmental Information Data Centre. Accessed: Aug. 01, 2023. [Online]. Available: <https://doi.org/10.5285/4482fa14-ae2-4c7f-9c62-a08dc9704051>
- [18] Centre for Innovation in Mathematics Teaching, *Correlation and regression*. University of Plymouth. Accessed: Jun. 27, 2025. [Online]. Available: https://www.cimt.org.uk/projects/mepres/alevel/stats_ch12.pdf

- [19] H. Sepp and S. Jorgen, “Long short-term memory,” *Neural Comput*, vol. 9, no. 8, pp. 1735–1780, 1997, doi: 10.1162/neco.1997.9.8.1735.

- [20] D. Chicco, M. J. Warrens, and G. Jurman, “The coefficient of determination R-squared is more informative than SMAPE, MAE, MAPE, MSE and RMSE in regression analysis evaluation,” *PeerJ Comput Sci*, vol. 7, pp. 1–24, 2021, doi: 10.7717/PEERJ-CS.623.

This page intentionally left

Enhancing the Cooling Effectiveness Utilizing a Tapered Fin Having Capsule-Shaped Cross-Sectional Area Numerically Simulated Using Finite Difference Method

Nico Ndaru Pratama^{1*}, Budi Setyahandana¹, Doddy

Purwadianto¹, Gilang Argya Dyaksa¹, Heryoga Winarbawa¹,

Michael Seen¹, Rines¹, Stefan Mardikus¹, Wibowo Kusbandono¹,

Y.B Lukiyanto¹, Marvin Edmin Son²

¹*Department of Mechanical Engineering, Faculty of Science and Technology, Sanata Dharma University, Yogyakarta, Indonesia*

²*Undergraduate Student-Department of Mechanical Engineering, Faculty of Science and Technology, Sanata Dharma University, Yogyakarta, Indonesia*

**Corresponding Author: nicondarupratama@usd.ac.id*

(Received 15-05-2025; Revised 21-06-2025; Accepted 08-07-2025)

Abstract

This paper reports the results of our research on improving the cooling of an engine using fins. This problem is important to discuss because various parts of the world have utilized machining technology. When the engine operates, it produces heat. This heat reduces the efficiency of the engine's performance. In this problem, we developed a tapered fin method with a capsule cross-section to enhance cooling performance. The fin consists of two different materials that are perfectly joined. In this paper, the fin analysis is performed using the explicit finite difference numerical method. This method simulates the heat distribution on the fins. The results of our research include temperature distribution, heat flow rate, efficiency, and fin effectiveness in unsteady-state conditions with variations in material composition. The highest heat flow rate, fin efficiency, and fin effectiveness were achieved with a fin material composition of copper and aluminum, yielding an efficiency value of 0.89 and an effectiveness of 20.7. Our research results offer potential for the industry to design fins for innovative applications.

Keywords: Effectiveness, Efficiency, Fin, Finite Difference

1 Introduction

Recently, energy use is increasing in a more efficient, flexible, sustainable manner. Several innovative techniques and applications continue to be developed, such as small



modular reactors (SMR) [1], [2], such as power generation systems with advanced operating parameters (high operating temperatures) [3]. In implementing the above innovation, a heat exchanger is required that possesses high efficiency and effectiveness. The ideal heat exchanger for this application must have a large heat transfer area while considering the narrow construction [4]. The one common heat exchanger option to increase the heat transfer rate is fins [5]. The fins serve to increase the area [6]. Examples of fin applications that we often encounter include computer processors, combustion chambers in combustion engines, electronic equipment, radiators, and other heat exchangers [7]. The larger number of fins installed, the greater the amount of heat released by those fins [8].

The efficiency and effectiveness of the fin are important considerations in fin design. The efficiency and effectiveness of fins can be investigated using analytical, experimental, and computational methods. In previous research, the one-dimensional case of a straight rhombus cross-section fin was analyzed using the explicit finite difference method [7]. In other studies, various cross-sectional shapes of fins, such as circular, pentagonal, and rectangular, have also been investigated [9], [10], [11], [12]. In this study, the efficiency and effectiveness of the fins were simulated using the explicit finite difference method. From various previous studies, no research has reported findings on the efficiency and effectiveness of tapered fins with a capsule-shaped cross-section composed of two different materials.

This paper aims to research temperature distribution in tapered fins that have a capsule cross-section consisting of two materials under unsteady conditions. This research uses the explicit finite difference method. The explicit finite difference method provides a fast solution, and its accuracy can be adjusted by modifying the size and number of nodes [13]. The results obtained from the temperature calculation equation in each node point were used to develop a computational program. In addition to calculating the temperature distribution, the program is also designed to calculate the heat flow rate, efficiency, and fin effectiveness, as well as to observe the effect of changes in one of the fin materials on distribution, heat flow rate, efficiency, and fin effectiveness.

2 Material and Methods

This study employs a computational approach based on the explicit finite difference method. Figure 1 shows an image of the tapered straight fin under investigation. The tapered fin has a capsule-shaped cross-section, with a varying cross-sectional area along its length, and is composed of two different materials. The height of the fin base capsule is $D = 0.01$ m, and its width is $D + a$, where value $a = 0.03$ m. The height of the fin tip capsule is denoted as $d = 0.005$ m, and its width is $d + a$. The total length of the straight fin is $L = 0.1$ m. The length of the fin made of material 1 is denoted by L_1 , and the length of the fin made of material 2 is denoted by L_2 . The fin is divided into 25 small elements, known as nodes, each spaced equally by a distance of Δx . The initial temperature of the fin is uniform and equal to the base temperature, i.e., $T_i = T_b = 100^\circ\text{C}$. The fin is then exposed to a fluid environment with a constant and uniform temperature $T_\infty = 30^\circ\text{C}$. The convection heat transfer coefficient is assumed to be constant, with a value of $h = 50$ $\text{W}/\text{m}^2\text{C}$. The base temperature is maintained constant over time. The entire lateral surface of the fin and the fin tip are in contact with the surrounding fluid. Heat conduction is assumed to take place solely in the x -direction (one direction). The properties of the fin

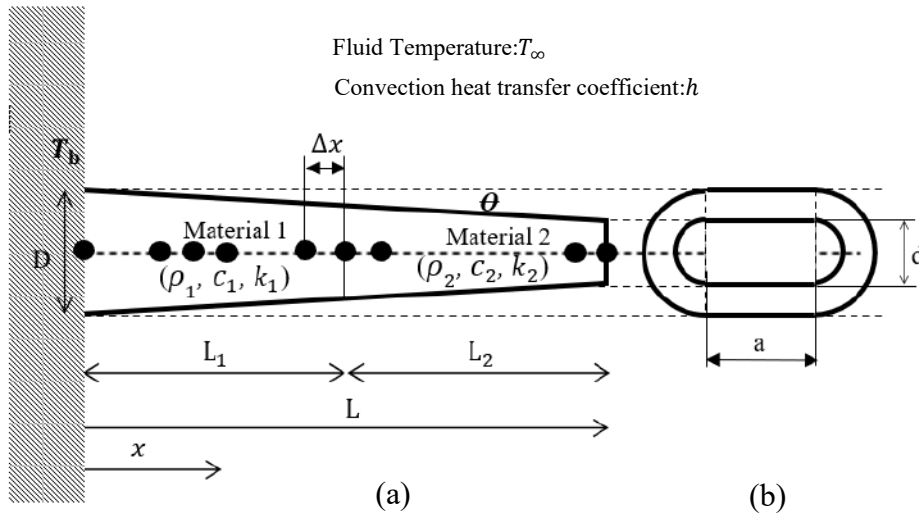


Figure 1. Tapered fin with capsule cross-section, (a) side view, (b) front view

material (density ρ , specific heat capacity c , and thermal conductivity k) were assumed to be constant and independent of temperature changes.

The mathematical model used to calculate the temperature at position x on the fin and at time t is a partial differential equation. The differential equation for this case is expressed in Equation (1):

$$\frac{\partial^2 T(x, t)}{\partial x^2} + \frac{hAs}{kA_p}(T(x, t) - T_\infty) = \frac{1}{\alpha} \frac{\partial T(x, t)}{\partial t}, \quad 0 < x < L, t > 0 \quad (1)$$

$$\alpha = \frac{k}{\rho c}, \quad (2)$$

and boundary condition:

$$T(0, t) = T_b, \quad x = 0, t > 0. \quad (3)$$

The formulation of transient heat conduction problems differs from that of steady-state conduction problems because transient cases involve an additional term representing the change in the energy content of the medium over time. This additional term appears as the first derivative of temperature with respect to time in the differential equation and as the change in internal energy content over a time interval Δt in the energy balance formulation [14]. The energy balance on a volume element during a time interval Δt (without energy generation) can be expressed as [14] :

$$\left(\begin{array}{c} \text{The total amount of energy} \\ \text{that flows into the volume element} \\ \text{over the time period } \Delta t \end{array} \right) = \left(\begin{array}{c} \text{The change in the energy} \\ \text{content of the volume} \\ \text{element during } \Delta t \end{array} \right)$$

It can be shown in Equation 4:

$$\sum_{i=1}^k q_i^n = \rho_i V_i c_1 \frac{\Delta T}{\Delta t}. \quad (4)$$

Based on Equation (4), the temperature within the volume elements between the base and the tip of the fin (excluding the interface volume element between the two fin materials) can be determined using Equation (5). The temperature of the interface volume element between the two fin materials is determined using Equation (6), while the temperature at the fin tip volume element is determined using Equation (7).

$$T_i^{n+1} = \frac{\Delta t \cdot k}{\rho \cdot V i \cdot C \cdot \Delta x} \cdot \left[A p_{i-\frac{1}{2}} \cdot (T_{i-1}^n - T_i^n) + A p_{i+\frac{1}{2}} \cdot (T_{i+1}^n - T_i^n) + B i \cdot A s_i \cdot (T_\infty - T_i^n) \right] + T_i^n, \quad (5)$$

$$T_i^{n+1} = \frac{\Delta t}{(\rho_1 \times c_1 \times V_1 + \rho_2 \times c_2 \times V_2)} \left[k_1 \cdot A p_{i-\frac{1}{2}} \cdot \left(\frac{T_{i-1}^n - T_i^n}{\Delta x} \right) + k_2 \cdot A p_{i+\frac{1}{2}} \cdot \left(\frac{T_{i+1}^n - T_i^n}{\Delta x} \right) + h \cdot A s_i \cdot (T_\infty - T_i^n) \right] + T_i^n, \quad (6)$$

$$T_i^{n+1} = \frac{\Delta t \cdot k_2}{\rho_2 \cdot V i \cdot c_2 \cdot \Delta x} \cdot \left[A p_{i-\frac{1}{2}} \cdot (T_{i-1}^n - T_i^n) + B i_2 \cdot A p_i \cdot (T_\infty - T_i^n) + B i_2 \cdot A s_i \cdot (T_\infty - T_i^n) \right] + T_i^n. \quad (7)$$

Under unsteady-state conditions, the actual heat transfer rate of the fin (q_{actual}^n), the maximum heat transfer rate (q_{ideal}^n), and the heat transfer rate without the fin (q_{finless}^n), can be calculated sequentially using Equations (8), (9), and (10).

$$q_{\text{actual}}^n = h \sum_{i=1}^m (A_i (T_i^n - T_\infty)), \quad (8)$$

$$q_{\text{ideal}}^n = h \sum_{i=1}^m (A_i (T_b - T_\infty)), \quad (9)$$

$$q_{\text{finless}}^n = h A_d (T_b - T_\infty). \quad (10)$$

Under transient conditions, the efficiency of the fin can be determined using Equation (11):

$$\eta_{\text{fin}}^n = \frac{q_{\text{actual}}^n}{q_{\text{ideal}}^n} = \frac{h \sum_{i=1}^m (A_i (T_i^n - T_\infty))}{h \sum_{i=1}^m (A_i (T_b - T_\infty))}. \quad (11)$$

Under transient conditions, the fin effectiveness of the fin can be determined using Equation (12):

$$\varepsilon_{\text{fin}}^n = \frac{q_{\text{actual}}^n}{q_{\text{finless}}^n} = \frac{h \sum_{i=1}^m (A_i (T_i^n - T_\infty))}{h A_d (T_b - T_\infty)}. \quad (12)$$

Table 1. Properties of the material [14]

Material	Density (ρ) (kg/m ³)	Thermal Conductivity (k) (W/m ² °C)	Specific Heat (c) (J/kg°°C)
Copper (Cu)	8933	401	385
Aluminum (Al)	2702	237	903
Zinc (Zn)	7140	116	389

Iron (Fe)	7870	80.2	447
Nikel (Ni)	8900	444	90,7
Bras (Cu-Zn)	8933	110	380

3 Results and Discussions

The results of the temperature distribution, actual heat flow rate, fin efficiency, and fin effectiveness calculations are presented in Figures 2, 3, 4, and 5. Under unsteady-state conditions, the temperature distribution continues to change over time. In this study, since the base temperature T_b is maintained at a constant 100°C and the fin is continuously exposed to the surrounding fluid, the temperature at each volume element continues to decrease until a steady-state condition is reached throughout the entire domain. In Figures 2 through 5, the temperature distribution is evaluated at 600 seconds by comparing the temperature at each volume element for various fin material compositions. The heat transfer rate and fin efficiency are evaluated at each time step for different material compositions

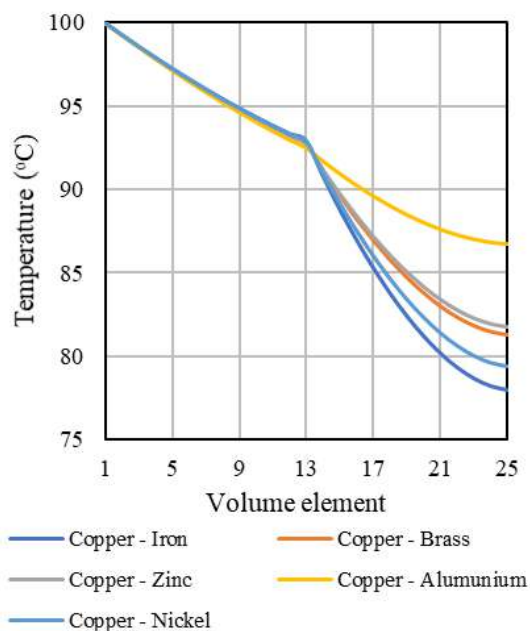


Figure 2. Fin temperature at $t=600$ s

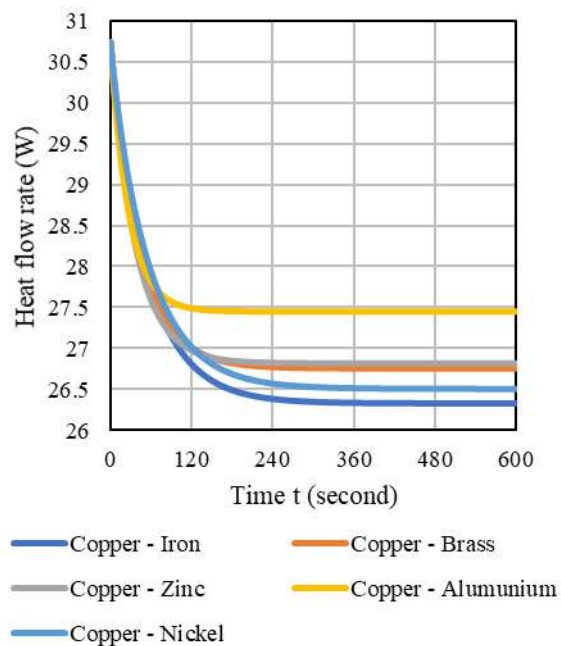
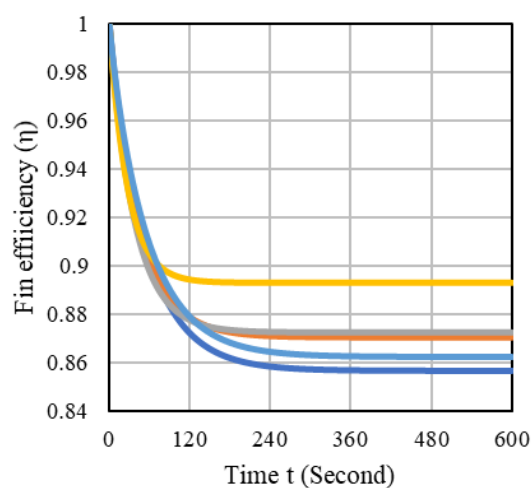


Figure 3. Heat transfer rate.

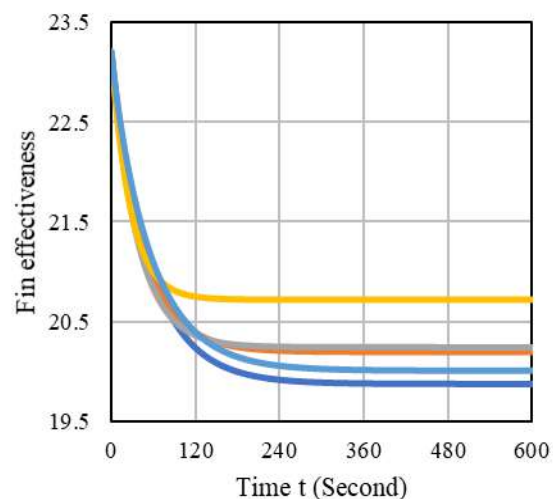
Figure 2 presents the temperature distribution along the fin for various material composition variations. At 600 seconds, the fin composed of copper–iron exhibits the greatest temperature drop, reaching 78°C from its initial condition. In contrast, the lowest temperature drop is observed in the copper–aluminum composition, with a final temperature of 86°C . The temperature distribution in the unsteady state is influenced by material properties such as density, thermal conductivity, and specific heat capacity, as shown in Table 1. This is because iron has a low thermal conductivity, resulting in less optimal heat transfer from the first material compared to other materials.

Figures 3, 4, and 5 present the actual heat flow rate from the fin, fin efficiency, and fin effectiveness. These calculated results all depend on the temperature distribution along the fin. The heat flow rate released by the fins, in sequence from the largest to the smallest is owned by the fin of the combinations of copper–aluminum, copper–zinc, copper–brass, copper–nickel, and copper–iron. This trend corresponds to the temperature distribution generated by each fin material combination. The higher the temperature distribution along the fin, the greater the heat flow rate. This is because a higher temperature along the fin results in a larger temperature difference with the surrounding fluid. Consequently, a



— Copper - Iron
— Copper - Zinc
— Copper - Nickel

— Copper - Brass
— Copper - Aluminium



— Copper - Iron
— Copper - Zinc
— Copper - Nickel

— Copper - Brass

— Copper - Aluminium

Figure 4. Fin efficiency

Figure 5. Fin effectiveness.

greater temperature difference leads to a higher rate of heat transfer. The highest efficiency and effectiveness are achieved by the copper–aluminum fin material composition, with an efficiency value of 0.89 and an effectiveness of 20.7.

4 Conclusions

Based on the research results, under unsteady-state conditions, the material properties that influence temperature distribution, heat flow rate, efficiency, and effectiveness are density, specific heat capacity, and thermal conductivity. The highest efficiency and effectiveness are achieved by the copper–aluminum fin composition, as it has the highest thermal diffusivity value. Our study is limited to several material compositions, convective heat transfer coefficients, initial temperatures, and fin geometries. Heat transfer is also assumed to occur only in a one-dimensional (1D) domain. Future research may expand this work by exploring a broader range of fin material compositions, two- or three-dimensional (2D/3D) fin problems, or materials with temperature-dependent properties.

Acknowledgements

The authors would like to thank Sanata Dharma University for its support in conducting this research.

Nomenclature

Δx	: Distance between volumes element (m)
Δt	: Time step, from the n to the $n + 1$ iteration (s)
T_{∞}	: Fluid temperature around the fin ($^{\circ}\text{C}$)
T_b	: Temperature at the base of the fin, ($^{\circ}\text{C}$)
T_i^n	: The temperature at the i position, in the $n + 1$ iteration ($^{\circ}\text{C}$)
T_i^{n+1}	: The temperature at the i position, in the $n + 1$ iteration ($^{\circ}\text{C}$)
T_{i-1}^n	: The temperature at the $i - 1$ position, in the n iteration ($^{\circ}\text{C}$)
T_{i+1}^n	: The temperature at the $i + 1$ position, in the n iteration ($^{\circ}\text{C}$)
L_1	: The length of the fin with material 1, m
L_2	: The length of the fin with material 2, m
L	: The length of the fin, $L = L_1 + L_2$, m
A_p	: Area of the cross section of the fin (m^2)
As_i	: Area of the surface of the volume element at the i position touching the fluids around the fin (m^2)

k	: Thermal conductivity (W/m°C)
k_1	: Thermal conductivity of fin material 1 (W/m°C)
k_2	: Thermal conductivity of fin material 2 (W/m°C)
h	: Convection heat transfer coefficient (W/m ² °C)
ρ	: The density of the fin material (kg/m ³)
ρ_1	: The density of the fin material 1 (kg/m ³)
ρ_2	: The density of the fin material 2 (kg/m ³)
c	: Specific heat of the material, (kJ/kg°C)
c_1	: Specific heat of the material 1, (kJ/kg°C)
c_2	: Specific heat of the material 2, (kJ/kg°C)
α	: Thermal diffusivity of the material, (m ² /s)
Bi	: Biot number
Bi_1	: Biot number of material 1
Bi_2	: Biot number of material 2

References

- [1] K. Shirvan, P. Hejzlar, and M. S. Kazimi, "The design of a compact integral medium size PWR," *Nuclear Engineering and Design*, vol. 243, pp. 393–403, Feb. 2012, doi: 10.1016/j.nucengdes.2011.11.023.
- [2] M. K. Rowinski, T. J. White, and J. Zhao, "Small and medium sized reactors (SMR): A review of technology," *Renewable and Sustainable Energy Reviews*, vol. 44, pp. 643–656, 2015, doi: 10.1016/j.rser.2015.01.006.
- [3] Y. Kato, T. Nitawaki, and Y. Muto, "Medium temperature carbon dioxide gas turbine reactor," *Nuclear Engineering and Design*, vol. 230, no. 1–3, pp. 195–207, May 2004, doi: 10.1016/j.nucengdes.2003.12.002.
- [4] S. Liu, Y. Huang, and J. Wang, "Theoretical and numerical investigation on the fin effectiveness and the fin efficiency of printed circuit heat exchanger with straight channels," *International Journal of Thermal Sciences*, vol. 132, pp. 558–566, Oct. 2018, doi: 10.1016/j.ijthermalsci.2018.06.029.
- [5] M. Seen *et al.*, "Characteristics of straight trapezoidal cross-sectional fins under unsteady conditions," *International Journal of Applied Sciences and Smart Technologies*, vol. 06, no. 01, Apr. 2024.
- [6] A. Mostafavi and A. Jain, "Thermal management effectiveness and efficiency of a fin surrounded by a phase change material (PCM)," *International Journal Heat Mass Transfer*, vol. 191, Aug. 2022, doi: 10.1016/j.ijheatmasstransfer.2022.122630.

- [7] T. D. Nugroho and P. K. Purwadi, "Fins effectiveness and efficiency with position function of rhombus sectional area in unsteady condition," *AIP Conference Proceeding*, vol. 1788, no. 1, Jan. 2017, doi: 10.1063/1.4968287.
- [8] P. K. Purwadi and M. Seen, "Efficiency and effectiveness of a fin having the capsule-shaped cross section in the unsteady state," *AIP Conference Proceeding*, vol. 2202, no. 1, Dec. 2019, doi: 10.1063/1.5141705.
- [9] K. Ginting, P.K. Purwadi, and S. Mungkasi, "Efficiency and effectiveness of a fin having capsule-shaped cross section dependent on the one-dimensional position," *Proceedings of The 1st International Conference on Science and Technology for an Internet of Things, 20 October 2018, Yogyakarta, Indonesia*, European Alliance for Innovation, Apr. 2019. doi: 10.4108/eai.19-10-2018.2282543.
- [10] P. K. Purwadi and B. Y. Pratama, "Efficiency and effectiveness of a truncated cone-shaped fin consisting of two different materials in the steady-state," *AIP Conference Proceeding*, vol. 2202, no. 1, Dec. 2019, doi: 10.1063/1.5141704.
- [11] A. W. Vidjabhakti, P. K. Purwadi, and S. Mungkasi, "Efficiency and effectiveness of a fin having pentagonal cross section dependent on the one-dimensional position," *Proceedings of The 1st International Conference on Science and Technology for an Internet of Things, 20 October 2018, Yogyakarta, Indonesia*, European Alliance for Innovation, Apr. 2019. doi: 10.4108/eai.19-10-2018.2282540.
- [12] P. K. Purwadi, B. Setyahandana, and R. B. P. Harsilo, "Obtaining the efficiency and effectiveness of fin in unsteady state conditions using explicit finite difference method," *International Journal of Applied Sciences and Smart Technologies*, vol. 3, no. 1, pp. 111–124, 2021.
- [13] J. Lei, Q. Wang, X. Liu, Y. Gu, and C. M. Fan, "A novel space-time generalized FDM for transient heat conduction problems," *Engineering Analysis with Boundary Elements*, vol. 119, pp. 1–12, Oct. 2020, doi: 10.1016/j.enganabound.2020.07.003.
- [14] Y. A. Cengel and A. J. Ghajar, *Heat and mass transfer*, Fifth Edition. McGraw-Hill Education, 2015.

Numerical Solution of Two-Dimensional Advection Diffusion Equation for Multiphase Flows in Porous Media Using a Novel Meshfree Method of Lines

F.O. Ogunfiditimi¹ and J.A. Kazeem^{1*}

¹ *Department of Mathematics, University of Abuja,
Federal Capital Territory, Abuja, Nigeria.*

**Corresponding Author E-mail: jamiu.kazeem@yahoo.com*

(Received 25-05-2025; Revised 23-06-2025; Accepted 08-07-2025)

Abstract.

This study proposes a novel Meshfree Method of Lines (MFMOL), in strong form formulation, to solve multiphase flows of solute transport modelled by two-dimensional (2D) Advection Diffusion Equations (ADE). The method uses a consistent and stable Augmented Radial Basis Point Interpolation Method (ARPIM) for spatial variables discretization of the models, while the time variable is left continuous, resulting in a system of Ordinary Differential Equations (ODEs) with initial conditions, which is solved numerically, via Matlab ode solver. The new method is proposed to overcome the challenges of numerical instabilities and large deformation due to complex domain, and distorted or low-quality meshes that attracts remeshing, all encountered by traditional Finite Element Method (FEM), Finite Difference Method (FDM) and Finite Volume Method (FVM). Also, the MFMOL is used in strong form formulation without any stabilization techniques for the convective terms in solute transport models, contrary to other methods like FEM, FDM, FVM and meshfree Finite Point Method (FPM) that require the stabilization techniques for fluid flow problems to guarantee acceptable results. The efficiency and accuracy of the new method were established and validated by using it to solve 2D diffusive and advective flow problems in the complex domain of porous structures. The results obtained agreed with the existing exact solutions, using less computational efforts, costs and time, compared with mesh-based methods and others that require stabilization for the convective terms. These features established the superior performance of the new method to the mesh-based methods and others that require special stabilization techniques for solving 2D multiphase flow of solute transport in porous media and other transient fluid flow problems.

Keywords: Advection Diffusion Equations (ADE), Augmented Radial Basis Point Interpolation Method (ARPIM), Meshfree Method of Lines (MFMOL), Multiphase Flows, Porous Media.



1 Introduction

Considering the complexity of numerous practical engineering and industrial problems modelled by partial differential equations (PDEs), numerical methods have become alternatives as analytical solutions are not obtainable [1]. For boundary value problems in complex domains, Finite Element Method (FEM) which is mesh-based had been a versatile and robust computational tool to get the approximate solutions [2]. The other commonly used mesh-based methods are the Boundary Element Method (BEM) which minimizes the dimension of the problem and has higher accuracy compared to FEM, Finite Difference Method (FDM) which is based on regular nodes for discretization, Finite Volume Methods (FVM) and others [1], [3], [4], [5]. However, the mesh-based methods have some limitations. These include instabilities and large deformation due to mesh-based interpolation for the complex domain, distorted or low-quality meshes, leading to higher errors, which necessitate remeshing at the expense of additional computational efforts, costs and time [6]. To overcome these limitations associated with the reliance on meshes to construct the approximating functions, meshfree methods with great adaptability for discontinuities, and large deformation of the complex domain geometry were introduced [7]. Over some years, various authors had studied different meshfree methods: the Radial Basis Function Method of Lines [8], the Element Free Galerkin (EFG) method [9], Hybrid Interpolating Meshless Method [10], the Diffuse Element Method (DEM) [11], the Local Radiant Point Interpolation Method (LRPIM) [12], the Improved Element Free Galerkin Method [13], Fracture Mapping [14], Meshless Method of Lines (MMOL) [15].

Meshfree methods have recorded quite considerable success in engineering and industrial applications. However, each of them has its associated strengths, weaknesses and conditions of applicability [2]. Meshfree methods with moving least squares (MLS) shape functions such as the DEM, EFG, Meshless local Petrov–Galerkin (MLPG) method [16], Finite Point Method [17] and others, sometimes experience singular moment matrices which break down the entire method. Using a large number of nodal points to control the singularity may lead to a reduction in the sparsity of the moment matrix of the method, and this affects

the computational efficiency of the method. Also, using fewer nodal points to avoid the singularity of the moment matrix may also cause incompatibility or continuity problems, especially when Galerkin's weak form is used. However, the singularity of the moment matrix can be controlled using the T2L scheme for node selection at an additional computational time [18]. In addition, meshfree methods with MLS shape functions and those with integral representation shape functions such as the smoothed particle hydrodynamics (SPH) method [19], and the Reproducing kernel particle methods [20], do not possess the Kronecker delta function property and this leads to the imposition of essential boundary conditions to the models using direct interpolation method at additional expense [18]. In getting rid of these difficulties, [21] introduced a more efficient Radial Point Interpolation Method (RPIM) that uses a few arbitrarily scattered nodal points, with the guaranteed existence of a nonsingular moment matrix, subject to avoidance of some specific shape parameters [22], [23], [24]. However, RPIM cannot reproduce linear polynomials, as the pure radial function fails the standard patch test to check its performance, as widely used to check the performances of standard methods like FEM. This affects the consistency of the RPIM [18]. To restore the consistency issue, [21], [22] introduced the addition of polynomial terms to RPIM.

Multiphase flows of heterogeneous mixtures of two or more phases, such as solid-liquid, gas-liquid or gas-solid, are encountered in numerous industrial and scientific applications which include: gas and petroleum flows, aerosol flows in the environment, boiling and condensation processes, gas-solid and slurry flow in pipelines, particle and fiber flows in airways, nanofluids, fluidized bed reactors and others [25]. In multiphase flows, the advection-diffusion equation (ADE) is used to model the transport of properties such as temperature, salinity, or the concentration of a solute across distinct phases and interfaces. Modelling multiphase flows in porous media with advection-diffusion equations can be challenging due to complex domain geometries, which makes mesh-based methods ineffective for simulations due to their various limitations [18]. Meshfree methods in strong forms are known to be direct, fast and efficient in the simulation of multiphase flow of solute

transport, but less stable compared to the weak forms, which are more stable but with higher computational procedures and time [18], [26]. In literature, special attention had been given to the stabilization of convective terms and the Neumann boundary conditions in transport models to ensure accuracy of the results [27]. To achieve the effective stability of the computational methods whether in weak form formulations or strong form formulations, some stabilization techniques such as anisotropic balancing diffusion, finite increment calculus, petrov-Galerkin functions weighting, adaptive meshing or remeshing, operator splitting, characteristic time integration, upwind finite difference derivatives and others, had been used by various authors to get acceptable results [27], [28], [29], [30].

In this paper, a new method MFMOL, is proposed to solve 2D multiphase flow of solute transport in porous media, using strong form formulations without the use of any stabilization technique. In this proposed method, an Augmented Radial Basis Point Interpolation Method (ARPIM) is used in discretizing the space variables of the models and the subsidiary conditions to reduce the PDEs to a system of ODEs, subject to the required initial conditions. Then, the resulting ODEs are numerically integrated in time with Matlab ode 45 solvers. The stability of the proposed method is rooted in the consistency of the ARPIM, and its effective accuracy for function fitting, for discretizing the spatial variable of the flow models. In addition, the proposed method possesses a guaranteed nonsingular moment matrix and satisfies the Kronecker delta function property for easy imposition of essential boundary conditions, and these make the new method a potential computational technique for solving multiphase flows of solute transport in complex domains. The new method is validated by applying it to solve 2D diffusive and advective flow problems in porous media, using strong form formulations without any stabilization for the convective terms. The results obtained agreed with the existing exact solutions, with less computational efforts and time, and enhanced stability, compared with mesh-based methods and others, that require a special stabilization technique for solving 2D multiphase flow of solute transport in porous media and other transient fluid flow problems. The features of the new MFMOL

establish its superior performance over the mesh-based methods and others for the solution of 2D multiphase flow of solute transport in porous media and transient fluid flow problems.

2 Material and Method

The governing equation describing two-dimensional solute transport in a nonreactive, source-free, homogeneous and isotropic porous medium is given by the two-dimensional Advection Diffusion Equation (ADE) [31]

$$\frac{\partial u(x,y,t)}{\partial t} + v_1 \frac{\partial u(x,y,t)}{\partial x} + v_2 \frac{\partial u(x,y,t)}{\partial y} = D_1 \frac{\partial^2 u(x,y,t)}{\partial x^2} + D_2 \frac{\partial^2 u(x,y,t)}{\partial y^2} \quad (1)$$

$$(x, y, t) \in [0, L_1] \times [0, L_2] \times [0, T]$$

$$u(x, y, 0) = \eta(x, y) \quad (x, y) \in [0, L_1] \times [0, L_2] \quad (2)$$

$$u(0, y, t) = \beta_1(y, t), \quad u(L_1, y, t) = \gamma_1(y, t) \quad t \in [0, T] \quad (3)$$

$$u(x, 0, t) = \beta_2(x, t), \quad u(x, L_2, t) = \gamma_2(x, t) \quad t \in [0, T] \quad (4)$$

where $u(x, y, t)$ [ML^{-3}] is the flux averaged (flowing) solute concentration at position (x, y) and time t . $v_1[LT^{-1}]$ is the average linear velocity in the direction of x coordinate, $v_2[LT^{-1}]$ is the average linear velocity in the direction of y coordinate, $D_1[L^2T^{-1}]$ is the dispersion coefficient in the direction of x coordinate, $D_2[L^2T^{-1}]$ is the dispersion coefficient in the direction of y coordinate, $t[T]$ is the time and $x[L]$ and $y[L]$ are the spatial coordinate. $\eta(x, y)$, $\beta_1(y, t)$, $\beta_2(x, t)$, $\gamma_1(y, t)$ and $\gamma_2(x, t)$ are known functions.

To use the Meshfree Method of Lines (MFMOL) method for the numerical solution of equations (1) – (4), the following procedures are followed:

Let $\Omega^2 = [0, L_1] \times [0, L_2]$ be the complex domain represented by the nodal collocation points $[\mathbf{x}_i]_{i=1}^n$, where $\mathbf{x}_i = (x_i, y_i)$, $i = 2, 3 \dots n - 1$ are the interior nodes and, $\mathbf{x}_i$, $i = 1$ and n are the boundary nodes. For a field variable $u(\mathbf{x}, t)$ defined in the complex domain Ω^2 , the

augmented Radial Basis Functions Point Interpolation Approximation (ARPIM) at the star node $\mathbf{x}_\alpha \in \Omega^2$, takes the form [18]:

$$u(\mathbf{x}_\alpha, t) = \boldsymbol{\phi}^T(\mathbf{x}_\alpha) \mathbf{u} = \sum_{i=1}^n \phi_i(\mathbf{x}_\alpha) u_i \quad \alpha = 1, 2, \dots, n \quad (5)$$

$$= [\phi_1(\mathbf{x}_\alpha) \phi_2(\mathbf{x}_\alpha) \dots \phi_k(\mathbf{x}_\alpha) \dots \phi_n(\mathbf{x}_\alpha)] [u_1 \ u_2 \ u_k \dots u_n]^T \quad (6)$$

Where

$$\phi_k(\mathbf{x}_\alpha) = \sum_{i=1}^n R_i(\mathbf{x}_\alpha) A_{ik} + \sum_{j=1}^m p_j(\mathbf{x}_\alpha) B_{jk}, \quad \alpha, k = 1, 2, \dots, n \quad (7)$$

is the shape function for the k th node in the local support domain of $\mathbf{x}_\alpha$. n and m are the numbers of nodes and the polynomial basis, respectively, and $\mathbf{u} = [u_1 \ u_2 \ u_k \dots u_n]^T$ is the vector of all the field nodal variables at the n local nodes.

The shape function $\phi_k(\mathbf{x}_\alpha)$ for the k th node in the local support domain of $\mathbf{x}_\alpha$ is expressed as nonsingular radial basis function point interpolation (RPIM) terms: $\sum_{i=1}^n R_i(\mathbf{x}_\alpha) A_{ik}$, augmented by polynomial basis terms: $\sum_{j=1}^m p_j(\mathbf{x}_\alpha) B_{jk}$, which restores the consistency and stability of the RPIM terms. A_{ik} and B_{jk} are the coefficients of the radial basis $R_i(\mathbf{x}_\alpha)$ and the polynomial basis $p_j(\mathbf{x}_\alpha)$, respectively.

Explicitly, equation (7) is expressed as:

$$\phi_k(\mathbf{x}_\alpha) = \begin{pmatrix} R_1(\mathbf{r}_{11}) & R_2(\mathbf{r}_{12}) & \dots & R_n(\mathbf{r}_{1n}) \\ R_1(\mathbf{r}_{21}) & R_2(\mathbf{r}_{22}) & \dots & R_n(\mathbf{r}_{2n}) \\ \vdots & \vdots & \dots & \vdots \\ R_1(\mathbf{r}_{n1}) & R_2(\mathbf{r}_{n2}) & \dots & R_n(\mathbf{r}_{nn}) \end{pmatrix} \begin{pmatrix} A_{1k} \\ A_{2k} \\ \vdots \\ A_{nk} \end{pmatrix} + \begin{pmatrix} p_1(\mathbf{x}_1) & p_2(\mathbf{x}_1) & \dots & p_m(\mathbf{x}_1) \\ p_1(\mathbf{x}_2) & p_2(\mathbf{x}_2) & \dots & p_m(\mathbf{x}_2) \\ \vdots & \vdots & \dots & \vdots \\ p_1(\mathbf{x}_n) & p_2(\mathbf{x}_n) & \dots & p_m(\mathbf{x}_n) \end{pmatrix} \begin{pmatrix} B_{1k} \\ B_{2k} \\ \vdots \\ B_{mk} \end{pmatrix} \quad (8)$$

Or

$$\phi_k(\mathbf{x}_\alpha) = \mathbf{R} \mathbf{A}_{ik} + \mathbf{P} \mathbf{B}_{jk}, \quad i = 1, 2, \dots, n \text{ and } j = 1, 2, \dots, m \quad (9)$$

where $R_i(\mathbf{r}_{kj})$ is the multiquadric radial basis function, with optimal shape parameters: $c = 1.42$ and $q = 1.03$ [21] given as:

$$R_i(\mathbf{r}_{kj}) = (\mathbf{r}_{kj}^2 + c^2)^q = (\mathbf{r}_{kj}^2 + 1.42^2)^{1.03} \quad k, j = 1, 2, \dots, n \quad (10)$$

The radial distance $\mathbf{r}_{kj}$ of a fixed star node $\mathbf{x}_k = (x_k, y_k)$ from other nodes $\mathbf{x}_j = (x_j, y_j)$ arbitrarily distributed in the complex domain Ω^2 and its boundary, is expressed by the Euclidean norm as [18]:

$$\mathbf{r}_{kj} = \|D_k - D_j\| = \sqrt{(x_k - x_j)^2 + (y_k - y_j)^2}, \quad D_k = (x_k, y_k), \quad D_j = (x_j, y_j). \quad (11)$$

The moment matrix $\mathbf{P}$ for the polynomial basis $p_j(\mathbf{x}) = [p_1(x) \ p_2(x) \ p_3(x) \ \dots \ p_m(x)]$, where m is the number of monomials in $p_j(\mathbf{x})$ is expressed as [18]:

$$\mathbf{P} = \begin{pmatrix} p_1(x_1) & p_2(x_1) & \dots & p_m(x_1) \\ p_1(x_2) & p_2(x_2) & \dots & p_m(x_2) \\ \vdots & \vdots & \dots & \vdots \\ p_1(x_n) & p_2(x_n) & \dots & p_m(x_n) \end{pmatrix} = \begin{bmatrix} 1 & x_1 & y_1 & \dots & x_1^{m-1}y_1^{m-1} \\ 1 & x_2 & y_2 & \dots & x_2^{m-1}y_2^{m-1} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_n & y_n & \dots & x_n^{m-1}y_n^{m-1} \end{bmatrix} \quad (12)$$

where m and n are, respectively, the number of polynomial basis and the number of nodal points.

In equation (7), the constant matrices A_{ik} and B_{jk} are the (i, k) th and (j, k) th elements of matrices $\mathbf{K}_a$ and $\mathbf{K}_b$, respectively, which are given as [18]:

$$\mathbf{K}_a = \mathbf{R}^{-1} - \mathbf{R}^{-1}\mathbf{P}\mathbf{K}_b \quad \text{and} \quad \mathbf{K}_b = [\mathbf{P}^T\mathbf{R}^{-1}\mathbf{P}]^{-1}\mathbf{P}^T\mathbf{R}^{-1} \quad (13)$$

Putting equation (5) into equations (1) – (4) gives the strong form formulations of the models as:

$$\frac{du_i}{dt} + v_1 \frac{\partial}{\partial x} (\sum_{i=1}^n \phi_i(\mathbf{x}_\alpha) u_i) + v_2 \frac{\partial}{\partial y} (\sum_{i=1}^n \phi_i(\mathbf{x}_\alpha) u_i) =$$

$$D_1 \frac{\partial^2}{\partial x^2} (\sum_{i=1}^n \phi_i(\mathbf{x}_\alpha) u_i) + D_2 \frac{\partial^2}{\partial y^2} (\sum_{i=1}^n \phi_i(\mathbf{x}_\alpha) u_i) \quad (14)$$

$$u(x_\alpha, y_\alpha, 0) = \eta(x_\alpha, y_\alpha) \quad \alpha = 1, 2, \dots, n \quad (15)$$

$$u(0, y_1, t_1) = \sum_{i=1}^n \phi_i(\mathbf{x}_1) u_i = \beta_1(y_1, t_1), \quad u(L_1, y_n, t_n) = \sum_{i=1}^n \phi_i(\mathbf{x}_n) u_i = \gamma_1(y_n, T) \quad (16)$$

$$u(x_1, 0, t_1) = \sum_{i=1}^n \phi_i(\mathbf{x}_1) u_i = \beta_2(x_1, t_1), \quad u(x_n, L_2, t_n) = \sum_{i=1}^n \phi_i(\mathbf{x}_n) u_i = \gamma_2(x_n, T) \quad (17)$$

Using equations (5) and (7) in equation (14), as well as, in equations (16) - (17) gives:

$$\begin{aligned} & \frac{du_i}{dt} + v_1 \left(\sum_{i=1}^n \frac{\partial R_i(\mathbf{x}_\alpha)}{\partial x} A_{ik} + \sum_{j=1}^m \frac{\partial p_i(\mathbf{x}_\alpha)}{\partial x} B_{jk} \right) \mathbf{u} + v_2 \left(\sum_{i=1}^n \frac{\partial R_i(\mathbf{x}_\alpha)}{\partial y} A_{ik} + \sum_{j=1}^m \frac{\partial p_i(\mathbf{x}_\alpha)}{\partial y} B_{jk} \right) \mathbf{u} \\ &= D_1 \left(\sum_{i=1}^n \frac{\partial^2 R_i(\mathbf{x}_\alpha)}{\partial x^2} A_{ik} + \sum_{j=1}^m \frac{\partial^2 p_i(\mathbf{x}_\alpha)}{\partial x^2} B_{jk} \right) \mathbf{u} + D_2 \left(\sum_{i=1}^n \frac{\partial^2 R_i(\mathbf{x}_\alpha)}{\partial y^2} A_{ik} + \sum_{j=1}^m \frac{\partial^2 p_i(\mathbf{x}_\alpha)}{\partial y^2} B_{jk} \right) \mathbf{u} \end{aligned} \quad (18)$$

$$\begin{bmatrix} g_1 \\ g_n \end{bmatrix} = \begin{bmatrix} \tilde{u}(0, y_1, t) \\ \tilde{u}(L_1, x_n, T) \end{bmatrix} = \begin{bmatrix} \phi_1(\mathbf{x}_1) & \phi_2(\mathbf{x}_1) & \dots & \phi_n(\mathbf{x}_1) \\ \phi_1(\mathbf{x}_n) & \phi_2(\mathbf{x}_n) & \dots & \phi_n(\mathbf{x}_n) \end{bmatrix} \mathbf{u} = \begin{bmatrix} \beta_1(y_1, t_1) \\ \gamma_1(y_n, T) \end{bmatrix} \quad (19)$$

$$\begin{bmatrix} h_1 \\ h_n \end{bmatrix} = \begin{bmatrix} \tilde{u}(x_1, 0, t) \\ \tilde{u}(x_n, L_2, T) \end{bmatrix} = \begin{bmatrix} \phi_1(\mathbf{x}_1) & \phi_2(\mathbf{x}_1) & \dots & \phi_n(\mathbf{x}_1) \\ \phi_1(\mathbf{x}_n) & \phi_2(\mathbf{x}_n) & \dots & \phi_n(\mathbf{x}_n) \end{bmatrix} \mathbf{u} = \begin{bmatrix} \beta_2(x_1, t_1) \\ \gamma_2(x_n, T) \end{bmatrix} \quad (20)$$

$$u(x_\alpha, y_\alpha, 0) = \eta(x_\alpha, y_\alpha) \quad (21)$$

$$\text{where } \mathbf{u} = [u_1 \ u_2 \ \dots \ u_n]^T \quad \text{and } \alpha, k = 1, 2, \dots, n \quad (22)$$

Using equations (5) - (8) in equation (18) gives the Augmented Radial Basis Function Point Interpolation Method (ARPIM) approximation for the space variable of the field variable $u(\mathbf{x}, t)$ as:

$$\begin{aligned}
& \frac{du_i}{dt} + v_1 \begin{pmatrix} \phi_{1x}(x_1) & \phi_{2x}(x_1) & \cdots & \phi_{nx}(x_1) \\ \phi_{1x}(x_2) & \phi_{2x}(x_2) & \cdots & \phi_{nx}(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ \phi_{1x}(x_n) & \phi_{2x}(x_n) & \cdots & \phi_{nx}(x_n) \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix} + \\
& v_1 \begin{pmatrix} \phi_{1y}(x_1) & \phi_{2y}(x_1) & \cdots & \phi_{ny}(x_1) \\ \phi_{1y}(x_2) & \phi_{2y}(x_2) & \cdots & \phi_{ny}(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ \phi_{1y}(x_n) & \phi_{2y}(x_n) & \cdots & \phi_{ny}(x_n) \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix} = D_1 \begin{pmatrix} \phi_{1xx}(x_1) & \phi_{2xx}(x_1) & \cdots & \phi_{nxx}(x_1) \\ \phi_{1xx}(x_2) & \phi_{2xx}(x_2) & \cdots & \phi_{nxx}(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ \phi_{1xx}(x_n) & \phi_{2xx}(x_n) & \cdots & \phi_{nxx}(x_n) \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix} + \\
& D_2 \begin{pmatrix} \phi_{1yy}(x_1) & \phi_{2yy}(x_1) & \cdots & \phi_{nyy}(x_1) \\ \phi_{1yy}(x_2) & \phi_{2yy}(x_2) & \cdots & \phi_{nyy}(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ \phi_{1yy}(x_n) & \phi_{2yy}(x_n) & \cdots & \phi_{nyy}(x_n) \end{pmatrix} \begin{pmatrix} u_1 \\ u_2 \\ \vdots \\ u_n \end{pmatrix} \quad (23)
\end{aligned}$$

Considering the Kronecker delta function property: $\phi_i(x_j) =$

$\begin{cases} 1, & i = j, \quad i, j = 1, 2, 3, \dots, n \\ 0, & i \neq j, \quad i, j = 1, 2, 3, \dots, n \end{cases}$, while substituting the boundary condition (19) – (20) into

equation (23) gives:

$$\begin{aligned}
& \frac{du_i}{dt} + v_1 \begin{pmatrix} \phi_{1x}(x_1) & \phi_{2x}(x_1) & \cdots & \phi_{nx}(x_1) \\ \phi_{1x}(x_2) & \phi_{2x}(x_2) & \cdots & \phi_{nx}(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ \phi_{1x}(x_n) & \phi_{2x}(x_n) & \cdots & \phi_{nx}(x_n) \end{pmatrix} \begin{pmatrix} \beta_1 \\ u_2 \\ \vdots \\ \gamma_1 \end{pmatrix} + v_1 \begin{pmatrix} \phi_{1y}(x_1) & \phi_{2y}(x_1) & \cdots & \phi_{ny}(x_1) \\ \phi_{1y}(x_2) & \phi_{2y}(x_2) & \cdots & \phi_{ny}(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ \phi_{1y}(x_n) & \phi_{2y}(x_n) & \cdots & \phi_{ny}(x_n) \end{pmatrix} \begin{pmatrix} \beta_2 \\ u_2 \\ \vdots \\ \gamma_2 \end{pmatrix} \\
& = D_1 \begin{pmatrix} \phi_{1xx}(x_1) & \phi_{2xx}(x_1) & \cdots & \phi_{nxx}(x_1) \\ \phi_{1xx}(x_2) & \phi_{2xx}(x_2) & \cdots & \phi_{nxx}(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ \phi_{1xx}(x_n) & \phi_{2xx}(x_n) & \cdots & \phi_{nxx}(x_n) \end{pmatrix} \begin{pmatrix} \beta_1 \\ u_2 \\ \vdots \\ \gamma_1 \end{pmatrix} + D_2 \begin{pmatrix} \phi_{1yy}(x_1) & \phi_{2yy}(x_1) & \cdots & \phi_{nyy}(x_1) \\ \phi_{1yy}(x_2) & \phi_{2yy}(x_2) & \cdots & \phi_{nyy}(x_2) \\ \vdots & \vdots & \ddots & \vdots \\ \phi_{1yy}(x_n) & \phi_{2yy}(x_n) & \cdots & \phi_{nyy}(x_n) \end{pmatrix} \begin{pmatrix} \beta_2 \\ u_2 \\ \vdots \\ \gamma_2 \end{pmatrix} \quad (24)
\end{aligned}$$

where $\beta_1 = \beta_1(y_1, t_1)$, $\beta_2 = \beta_2(x_1, t_1)$, $\gamma_1 = \gamma_1(y_n, T)$, $\gamma_2 = \gamma_1(x_n, T)$ and $i = 1, 2 \dots n$

Considering the specified flux-averaged concentrations ($u_1 = \beta_1, u_n = \gamma_1$), ($u_1 = \beta_2, u_n = \gamma_2$) on the boundary nodes $x_1 = (x_1, y_1, t_1)$ and $x_n = (x_n, y_n, T)$ of the space coordinates x

and y respectively in (24), and partitioning the $(n \times n)$ global matrix equation in (24) leads to $(n - 2) \times (n - 2)$ submatrix equation for the unknown field nodal variables: $u_2, u_3, \dots, u_{n-1}$ for the flux-averaged concentrations at interior nodes $\mathbf{x}_i = (x_i, y_i, z_i, t_i)$, $i = 2, 3, \dots, n - 1$, expressed as:

$$\begin{aligned} \frac{du_i}{dt} + v_1 \begin{pmatrix} \phi_{2x}(\mathbf{x}_2) \phi_{3x}(\mathbf{x}_2) \cdots \phi_{sx}(\mathbf{x}_2) \\ \phi_{2x}(\mathbf{x}_3) \phi_{3x}(\mathbf{x}_3) \cdots \phi_{sx}(\mathbf{x}_3) \\ \vdots \\ \phi_{2x}(\mathbf{x}_s) \phi_{3x}(\mathbf{x}_s) \cdots \phi_{sx}(\mathbf{x}_s) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \\ \vdots \\ u_s \end{pmatrix} + v_1 \begin{pmatrix} \phi_{2y}(\mathbf{x}_2) \phi_{3y}(\mathbf{x}_2) \cdots \phi_{sy}(\mathbf{x}_2) \\ \phi_{2y}(\mathbf{x}_3) \phi_{3y}(\mathbf{x}_3) \cdots \phi_{sy}(\mathbf{x}_3) \\ \vdots \\ \phi_{2y}(\mathbf{x}_s) \phi_{3y}(\mathbf{x}_s) \cdots \phi_{sy}(\mathbf{x}_s) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \\ \vdots \\ u_s \end{pmatrix} \\ = D_1 \begin{pmatrix} \phi_{2xx}(\mathbf{x}_2) \phi_{3xx}(\mathbf{x}_2) \cdots \phi_{sxx}(\mathbf{x}_2) \\ \phi_{2xx}(\mathbf{x}_3) \phi_{3xx}(\mathbf{x}_3) \cdots \phi_{sxx}(\mathbf{x}_3) \\ \vdots \\ \phi_{2xx}(\mathbf{x}_s) \phi_{3xx}(\mathbf{x}_s) \cdots \phi_{sxx}(\mathbf{x}_s) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \\ \vdots \\ u_s \end{pmatrix} + D_2 \begin{pmatrix} \phi_{2yy}(\mathbf{x}_2) \phi_{3yy}(\mathbf{x}_2) \cdots \phi_{syy}(\mathbf{x}_2) \\ \phi_{2yy}(\mathbf{x}_3) \phi_{3yy}(\mathbf{x}_3) \cdots \phi_{syy}(\mathbf{x}_3) \\ \vdots \\ \phi_{2yy}(\mathbf{x}_s) \phi_{3yy}(\mathbf{x}_s) \cdots \phi_{syy}(\mathbf{x}_s) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \\ \vdots \\ u_s \end{pmatrix} \end{aligned} \quad (25)$$

where $i = 2, 3, \dots, s = n - 1$.

Using equation (21), the corresponding initial conditions for the unknown field nodal variables: $u_2, u_3, \dots, u_{n-1}$ are expressed as:

$$u(x_\alpha, y_\alpha, 0) = \eta(x_\alpha, y_\alpha) \quad \alpha = 2, 3 \dots s = n - 1 \quad (26)$$

Then, the system of ODEs in equation (25) and the initial conditions in equation (26) are numerically integrated using Matlab ode 45, to obtain the computational values of $u_2, u_3, \dots, u_{n-1}$.

3 Results and Discussions

Here, the new method Meshfree Method of Lines (MFMOL), is applied to solve 2D Advection Diffusion Equations (ADE) for multiphase flows of solute transport in a nonreactive, source-free, homogeneous and isotropic porous media. For ADE used for modelling solute transport, the peclet number, $Pe \gg 1$ corresponds to the advection time scale being greater than the diffusion time scale of the flow and this causes an unstable flow. However, the Peclet number, $Pe \ll 1$ corresponds to the advection time scale being less than

the diffusion time scale of the flow and this gives a stable flow. For all the computations, the space variable discretizations for the problem complex domain $\Omega^2 = [0, 1] \times [0, 1]$ are done using scattered and unevenly distributed nodal points: $\{(0,0), (0.5, 0.3), (0.7, 0.8), (1,1)\}$.

Example 1: Consider the *diffusion-controlled* solute transport equation (1) with the dispersion coefficients $D_1 = D_2 = 0.5 \text{ m}^2\text{s}^{-1}$, the flow average linear velocities $v_1 = v_2 = 0.5 \text{ m}$. The initial and boundary conditions are:

$$\eta(x, y, 0) = e^{-x} + e^{-y} \quad (27)$$

$$\beta_1(y, t) = (1 + e^{-y})e^{-4t}, \gamma_1(y, t) = (e^{-1} + e^{-y})e^{-4t} \quad (28)$$

$$\beta_2(x, t) = (e^{-x} + 1)e^{-4t}, \gamma_2(y, t) = (e^{-x} + e^{-1})e^{-4t} \quad (29)$$

The flow Peclet number: $Pe = \frac{|v|L}{|D|} = 1$, where $L_1 = L_2 = 1$. The existing model's exact solution is given as [13]:

$$u(x, y, t) = (e^{-x} + e^{-y})e^{-4t}, \quad (x, y, t) \in [0, 1] \times [0, 1] \times [0, 0.1] \quad (30)$$

Using the procedures in equations (5) – (26), gives the MFMOL approximation for the system of ODEs:

$$\begin{aligned} \frac{du_i}{dt} = & D_1 \begin{pmatrix} \phi_{2xx}(x_2) & \phi_{3xx}(x_2) \\ \phi_{2xx}(x_3) & \phi_{3xx}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} - v_1 \begin{pmatrix} \phi_{2x}(x_2) & \phi_{3x}(x_2) \\ \phi_{2x}(x_3) & \phi_{3x}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} - \\ & v_2 \begin{pmatrix} \phi_{2y}(x_2) & \phi_{3y}(x_2) \\ \phi_{2y}(x_3) & \phi_{3y}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} \\ & + D_2 \begin{pmatrix} \phi_{2yy}(x_2) & \phi_{3yy}(x_2) \\ \phi_{2yy}(x_3) & \phi_{3yy}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix}, \quad i = 2 \text{ and } 3 \end{aligned} \quad (31)$$

subject to the initial conditions:

$$u(\mathbf{x}, 0) = [u(\mathbf{x}_2, 0), u(\mathbf{x}_3, 0)] = [e^{-x_2} + e^{-y_2}, e^{-x_3} + e^{-y_3}], \text{ where } \mathbf{x} = (x, y) \quad (32)$$

The computational outputs of equations (31) – (32) are given as:

$$\begin{pmatrix} u_2 \\ u_3 \end{pmatrix} = \begin{pmatrix} -1.3765 & -2.7527 \\ -0.6783 & -1.3563 \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix}, \quad u(\mathbf{x}, 0) = \begin{pmatrix} 1.3473 \\ 0.9459 \end{pmatrix} \quad \text{and} \quad u_t = \frac{du_i}{dt} \quad (33)$$

The second iterate of Variational Iteration Method (VIM) solution is obtained as [32]:

$$u(x, y, t) = (1 + 3t + t^2)(e^{-x} + e^{-y}) \quad (34)$$

Example 2: Consider the *advection-controlled* solute transport equation (1) with the dispersion coefficients $D_1 = D_2 = 0.01 \text{ m}^2\text{s}^{-1}$, the flow average linear velocities $v_1 = v_2 = 0.8 \text{ ms}^{-1}$, The boundary and initial conditions are:

$$\beta_1(y, t) = \frac{1}{4t+1} e^{\left(\frac{-(0.8t+0.5)^2 - (y-0.8t-0.5)^2}{0.01(4t+1)}\right)}, \quad \gamma_1(y, t) = \frac{1}{4t+1} e^{\left(\frac{-(0.5-0.8t)^2 - (y-0.8t-0.5)^2}{0.01(4t+1)}\right)} \quad (35)$$

$$\beta_2(x, t) = \frac{1}{4t+1} e^{\left(\frac{-(x-0.8t-0.5)^2 - (0.8t+0.5)^2}{0.01(4t+1)}\right)}, \quad \gamma_2(y, t) = \frac{1}{4t+1} e^{\left(\frac{-(x-0.8t-0.5)^2 - (0.5-0.8t)^2}{0.01(4t+1)}\right)} \quad (36)$$

$$\eta(x, y, 0) = e^{-100[(x-0.5)^2 + (y-0.5)^2]} \quad (37)$$

The flow Peclet number: $Pe = \frac{|v|L}{|D|} = 80$, where $L_1 = L_2 = 1$. The existing exact solution of the model is given as [33]:

$$u(x, y, t) = \frac{1}{4t+1} e^{\left(\frac{-(x-0.8t+0.5)^2 - (y-0.8t-0.5)^2}{0.01(4t+1)}\right)}, \quad (x, y, t) \in [0, 1] \times [0, 1] \times [0, 0.3]. \quad (38)$$

The procedures in equation (5) – (26) gives the MFMOL approximation for the system of ODEs:

$$\begin{aligned} \frac{du_i}{dt} = & D_1 \begin{pmatrix} \phi_{2xx}(x_2) & \phi_{3xx}(x_2) \\ \phi_{2xx}(x_3) & \phi_{3xx}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} - v_1 \begin{pmatrix} \phi_{2x}(x_2) & \phi_{3x}(x_2) \\ \phi_{2x}(x_3) & \phi_{3x}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} - \\ & v_2 \begin{pmatrix} \phi_{2y}(x_2) & \phi_{3y}(x_2) \\ \phi_{2y}(x_3) & \phi_{3y}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} \\ & + D_2 \begin{pmatrix} \phi_{2yy}(x_2) & \phi_{3xx}(x_2) \\ \phi_{2yy}(x_3) & \phi_{3xx}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} \quad i = 2 \quad \text{and} \quad 3 \end{aligned} \quad (39)$$

subject to the initial conditions:

$$u(x, 0) = [u(x_2, 0), u(x_3, 0)] = [e^{-100[(x_2-0.5)^2+(y_2-0.5)^2]}, e^{-100[(x_3-0.5)^2+(y_3-0.5)^2]}] \quad (40)$$

The computational values of equations (39) – (40) are given as:

$$\begin{pmatrix} \dot{u}_2 \\ \dot{u}_3 \end{pmatrix} = \begin{pmatrix} -0.3301 & -0.6597 \\ -0.6758 & -1.3520 \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix}, \quad u(x, 0) = \begin{pmatrix} 0.018316 \\ 2.26E-6 \end{pmatrix} \quad \text{and} \quad \dot{u}_i = \frac{du_i}{dt} \quad (41)$$

The first iterate of VIM solution is obtained as [32]

$$u(x, y, t) = [1 + \{D(A^2 + B^2 + 2q) - v(A + B)\}t]e^{-100((x-0.5)^2+(y-0.5)^2)} \quad (42)$$

where: $v = v_1 = v_2 = 0.8$, and $D = D_1 = D_2 = 0.01$

Example 3: Consider the *diffusion-controlled* solute transport equation (1) with the dispersion coefficients $D_1 = 1.4m^2s^{-1}$, $D_2 = 1.7 m^2s^{-1}$, the flow average linear velocities $v_1 = v_2 = 1.0 ms^{-1}$, The initial and boundary conditions are:

$$\eta(x, y, 0) = e^{-c_1x} + e^{-c_2y} \quad (43)$$

$$\beta_1(y, t) = (1 + e^{-c_2y})e^{-10t}, \quad \gamma_1(y, t) = (e^{-c_1} + e^{-c_2y})e^{-10t} \quad (44)$$

$$\beta_2(x, t) = (e^{-c_1 x} + 1)e^{-10t}, \quad \gamma_2(y, t) = (e^{-c_1 x} + e^{-c_2})e^{-10t} \quad (45)$$

$$c_1 = \frac{-v_1 - \sqrt{v_1^2 + 4bD_1}}{2D_1}, \quad c_2 = \frac{-v_2 - \sqrt{v_2^2 + 4bD_2}}{2D_2}, \quad b = -10, \quad Pe = \frac{|v|L}{|D|} < 1, \quad L_1 = L_2 = 1 \quad (46)$$

The existing exact solution of the model is given as [13]:

$$u(x, y, t) = (e^{-c_1 x} + e^{-c_2 y})e^{-10t}, \quad (x, y, t) \in [0, 1] \times [0, 1] \times [0, 0.1] \quad (47)$$

Following the procedures in equations (5) – (26) gives the MFMOL approximation for the system of ODEs:

$$\begin{aligned} \frac{du_i}{dt} = & D_1 \begin{pmatrix} \phi_{2xx}(x_2) & \phi_{3xx}(x_2) \\ \phi_{2xx}(x_3) & \phi_{3xx}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} - v_1 \begin{pmatrix} \phi_{2x}(x_2) & \phi_{3x}(x_2) \\ \phi_{2x}(x_3) & \phi_{3x}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} - \\ & v_2 \begin{pmatrix} \phi_{2y}(x_2) & \phi_{3y}(x_2) \\ \phi_{2y}(x_3) & \phi_{3y}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} \\ & + D_2 \begin{pmatrix} \phi_{2yy}(x_2) & \phi_{3yy}(x_2) \\ \phi_{2yy}(x_3) & \phi_{3yy}(x_3) \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix} \quad i = 2 \text{ and } 3 \end{aligned} \quad (48)$$

subject to the initial conditions:

$$u(x, 0) = [u(x_2, 0), u(x_3, 0)] = [e^{-c_1 x_2} + e^{-c_2 y_2}, e^{-c_1 x_3} + e^{-c_2 y_3}] \quad (49)$$

The computational outputs of equations (48) – (49) give:

$$\begin{pmatrix} \dot{u}_2 \\ \dot{u}_3 \end{pmatrix} = \begin{pmatrix} -4.02179 & -8.04295 \\ -2.55048 & -5.1006 \end{pmatrix} \begin{pmatrix} u_2 \\ u_3 \end{pmatrix}, \quad u(x, 0) = \begin{pmatrix} 2.4939 \\ 2.9610 \end{pmatrix} \quad \text{and} \quad \dot{u}_i = \frac{du_i}{dt} \quad 50$$

The first iterate of the VIM solution is given as [32]:

$$u(x, y, t) = (1 + \{D_1 c_1^2 + v_1 c_1\}t)e^{-c_1 x} + (1 + \{D_2 c_2^2 + v_2 c_2\}t)e^{-c_2 y} \quad (51)$$

Then equations (33), (41), and (50) are numerically integrated via the MATLAB ode45 solver, to obtain the computational values of the unknown nodal values u_2 and u_3 .

Tables 1-3 show the computational results of the new method MFMOL, the exact solutions, the Variational Iteration Method (VIM) solution, and the absolute errors for examples 1-3, respectively. Figs 1-20 show graphical displays comparing the exact solutions, MFMOL solutions, and the absolute errors for examples 1-3.

Table 1. Example 1's table for the exact solutions, the MFMOL solutions, and the absolute error computed at (0.5, 0.3) and (0.7,0.8) over $0 \leq t \leq 0.001$ with $\Delta t = 0.0002$

(x, y)	t	$u_{ex}(x, y, t)$	MFMOL: $u_{num}(x, y, t)$	VIM [32]	$e_1: u_{ex} - u_{num} $	$e_2: u_{ex} - VIM $
(0.5,0.3)	0.0000	1.3473	1.3473	1.3473	0.0E-4	0.0E-4
	0.0002	1.3463	1.3464	1.3482	1.0E-4	1.9E-3
	0.0004	1.3452	1.3455	1.3490	3.0E-4	3.8E-3
	0.0006	1.3441	1.3446	1.3498	5.0E-4	5.7E-3
	0.0008	1.3430	1.3437	1.3506	7.0E-4	7.6E-3
	0.001	1.3420	1.3428	1.3514	8.0E-4	9.4E-3
(0.7,0.8)	0.0000	0.9459	0.9459	0.9459	0.0E-4	0.0E-4
	0.0002	0.9452	0.9455	0.9465	3.0E-4	7.0E-4
	0.0004	0.9444	0.9450	0.9470	6.0E-4	2.6E-3
	0.0006	0.9436	0.9446	0.9476	1.0E-3	4.0E-3
	0.0008	0.9429	0.9441	0.9482	1.2E-3	5.3E-3
	0.001	0.9421	0.9437	0.9488	1.6E-3	6.7E-3

Table 2. Example 2's table for the exact solutions, the MFMOL numerical solutions, and the absolute errors computed at points (0.5, 0.3) and (0.7,0.8) over $0 \leq t \leq 0.1$ with $\Delta t = 0.02$

(x, y)	t	$u_{ex}(x, y, t)$	MFMOL: $u_{num}(x, y, t)$	VIM [32]	$e_1: u_{ex} - u_{num} $	$e_2: u_{ex} - u_{num} $
(0.5,0.3)	0.00	0.01830	0.01830	0.01830	0.0E-4	0.0E-4
	0.02	0.01200	0.01820	0.01820	6.2E-3	6.2E-3
	0.04	0.00760	0.01810	0.01820	1.1E-2	1.1E-2
	0.06	0.00470	0.01790	0.01810	1.3E-2	1.3E-2

	0.08	0.00280	0.01780	0.01800	1.5E-2	1.5E-2
	0.1	0.00170	0.01770	0.01790	1.6E-2	1.6E-2
(0.7,0.8)	0.00	0.00000	0.00000	0.00000	0.0E-5	0.0E-5
	0.02	0.00002	0.00025	0.00001	2.3E-4	1.0E-5
	0.04	0.00015	0.00051	0.00001	3.6E-4	1.4E-4
	0.06	0.00074	0.00077	0.00002	3.0E-5	7.2E-4
	0.08	0.00270	0.00100	0.00003	1.7E-3	2.7E-3
	0.1	0.00800	0.00130	0.00003	6.7E-3	8.0E-3

Table 3. Example 3's table for comparison of the exact solution and the MFMOL numerical solution, with the absolute error computed at (0.5, 0.3) and (0.7, 0.8) over $0 \leq t \leq 0.001$ with $\Delta t = 0.0002$

(x, y)	t	$u_{ex}(x, y, t)$	MFMOL: $u_{num}(x, y, t)$	VIM[32]	$e_1: u_{ex} - u_{num} $	$e_1: u_{ex} - u_{num} $
(0.5,0.3)	0.0000	2.4940	2.4939	2.4940	1.0E-4	0.0E-4
	0.0002	2.4939	2.4871	2.4939	6.8E-3	0.0E-4
	0.0004	2.4939	2.4804	2.4939	1.4E-2	0.0E-4
	0.0006	2.4938	2.4736	2.4938	2.0E-2	0.0E-4
	0.0008	2.4938	2.4669	2.4938	2.7E-2	0.0E-4
	0.001	2.4937	2.4602	2.4937	3.4E-2	0.0E-4
(0.7,0.8)	0.0000	2.9610	2.9610	2.9610	0.0E-4	0.0E-4
	0.0002	2.9609	2.9567	2.9609	4.2E-3	0.0E-4
	0.0004	2.9609	2.9524	2.9609	8.5E-3	0.0E-4
	0.0006	2.9608	2.9482	2.9608	1.3E-2	0.0E-4
	0.0008	2.9610	2.9439	2.9608	1.7E-2	2.0E-4
	0.001	2.9607	2.9396	2.9607	2.1E-2	0.0E-4

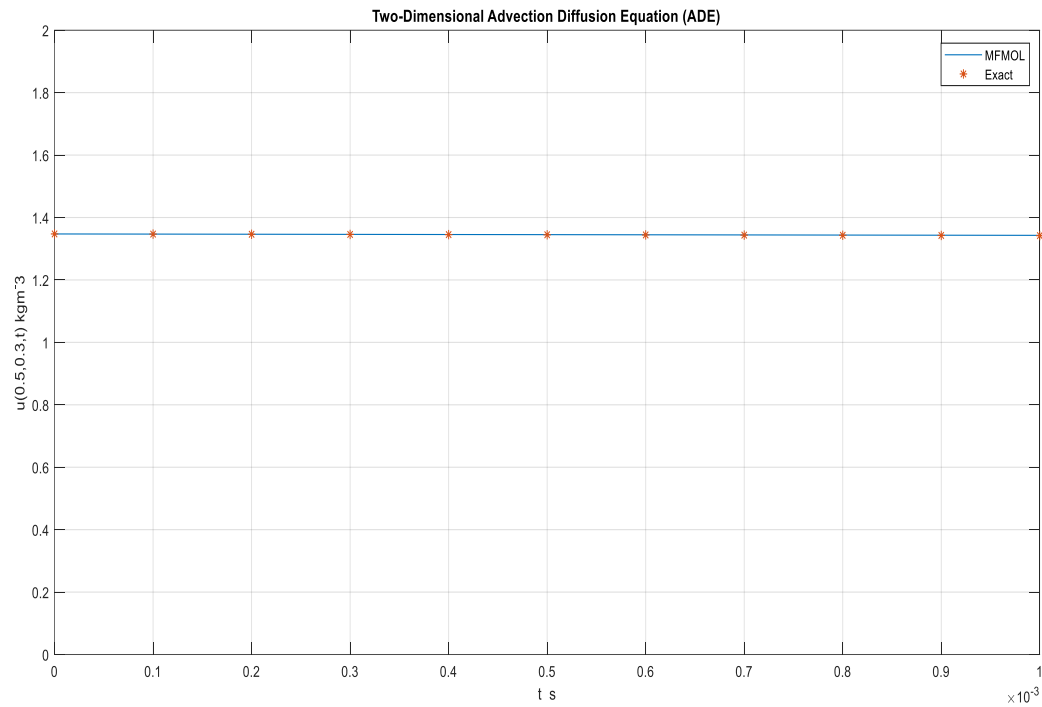


Figure 1. Plot of Example 1, for the flux averaged solute concentration corresponding to fixed point $(x, y) = (0.5, 0.3)$ over $0 \leq t \leq 0.001$ with time steps $\Delta t = 0.0001$

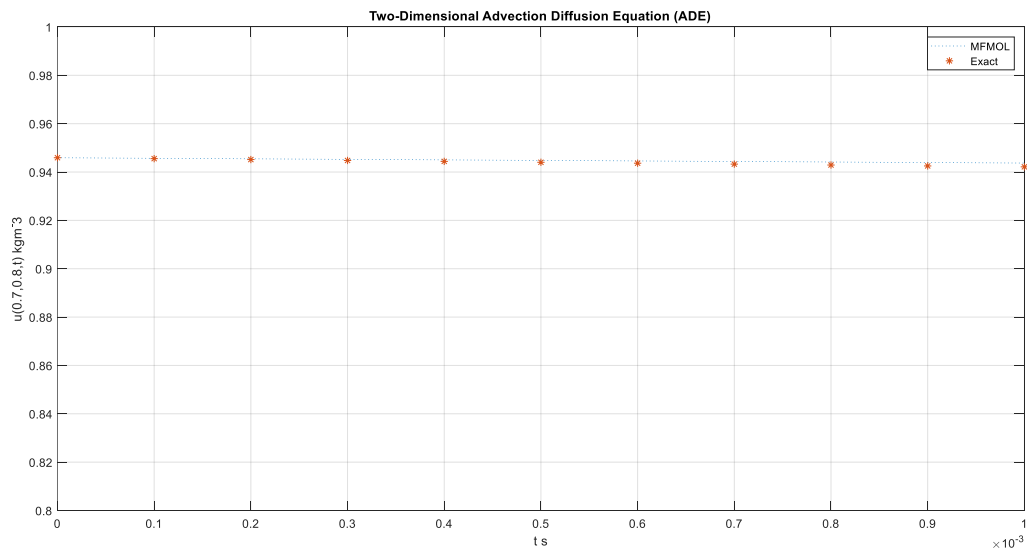


Figure 2. Plot of Example 1, for the flux averaged solute concentration corresponding to fixed point $(x, y) = (0.7, 0.8)$ over $0 \leq t \leq 0.001$ with time steps $\Delta t = 0.0001$

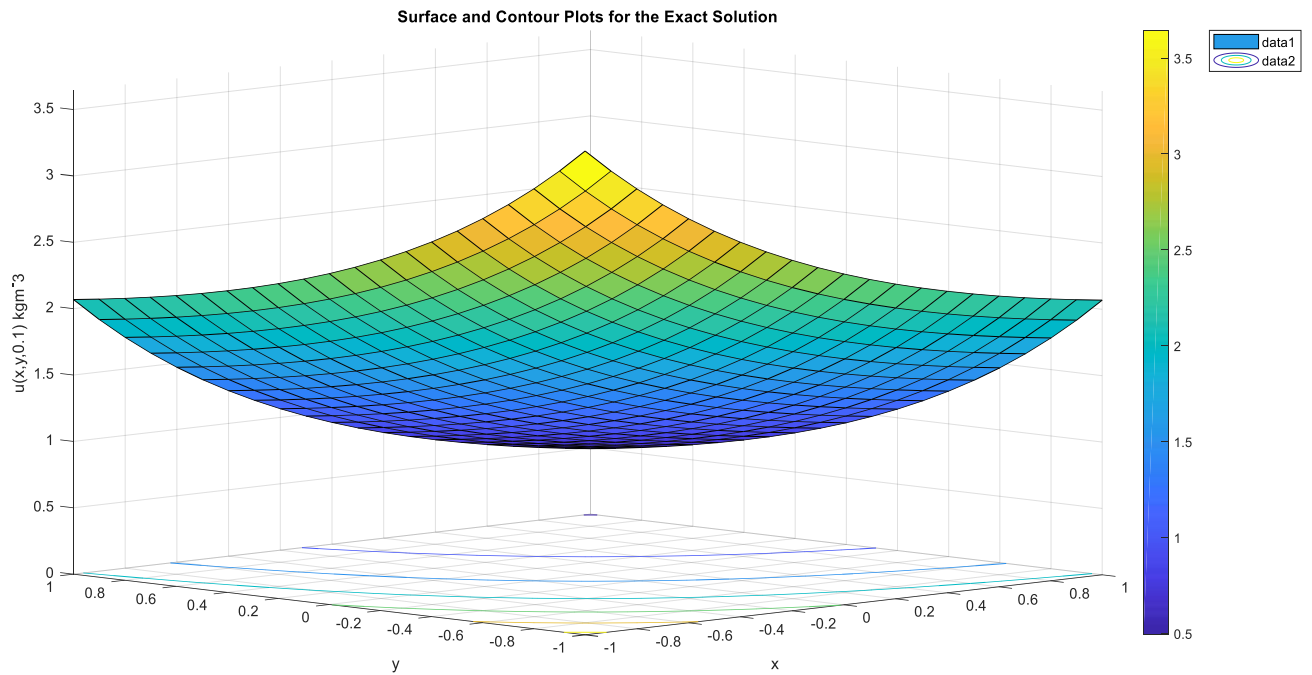


Figure 3. Surface and Contour plots of Example 1 for the exact solutions of the spatial distributions of the flux-averaged concentration $u(x,y,t)$ at $t = 0.1$

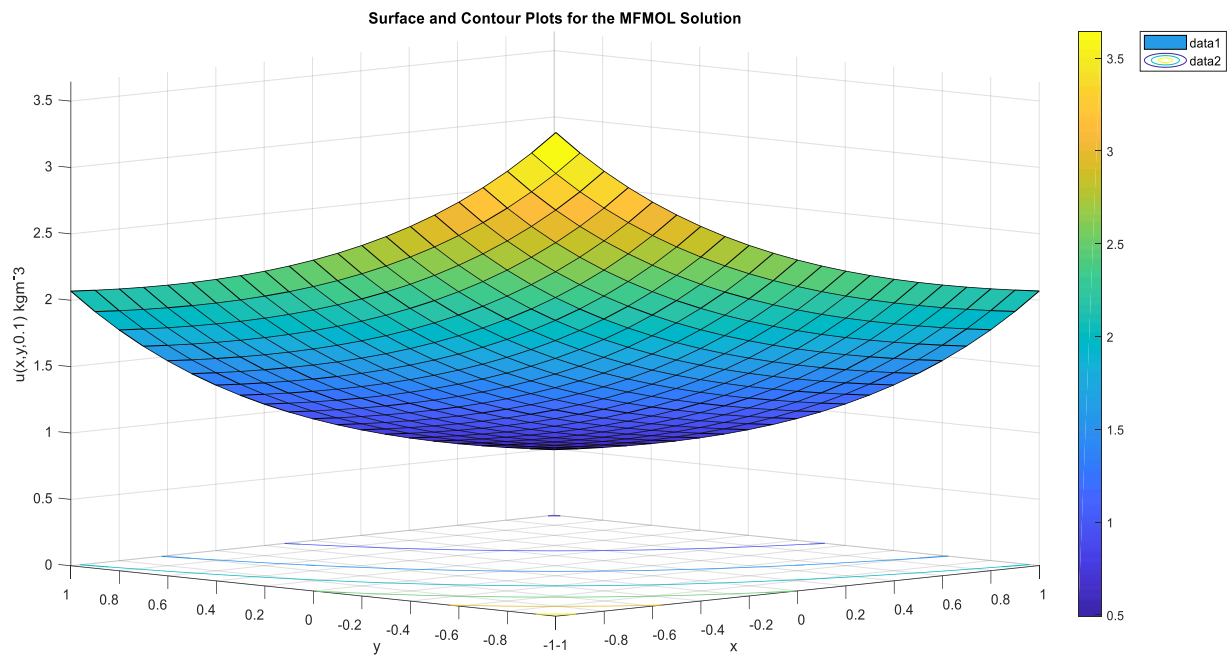


Figure 4. Surface and Contour plots of Example 1 for the MFMOL solutions of the spatial distributions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.1$

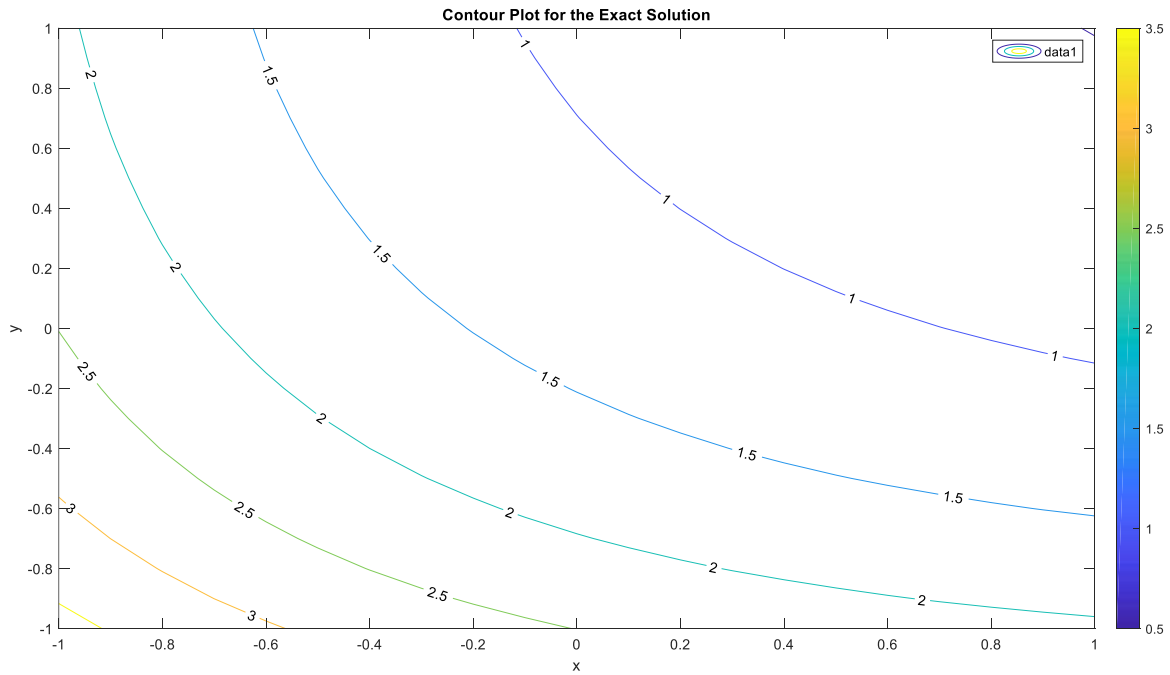


Figure 5. Contour plots of Example 1, for the exact solutions of the spatial distributions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.1$

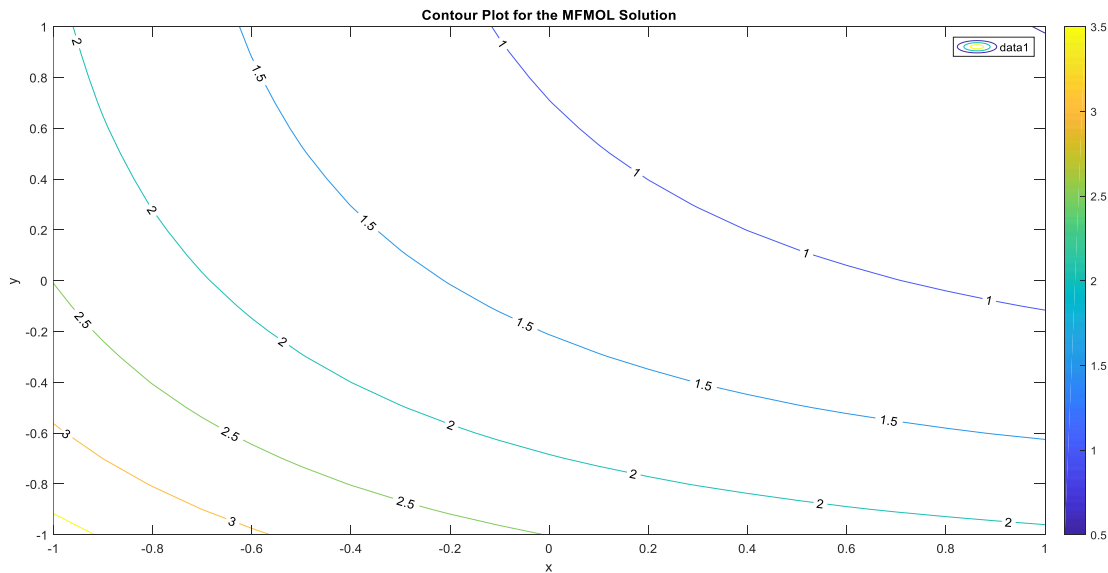


Figure 6. Contour plots of Example 1, for the MFMOL solutions of the spatial distributions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.1$

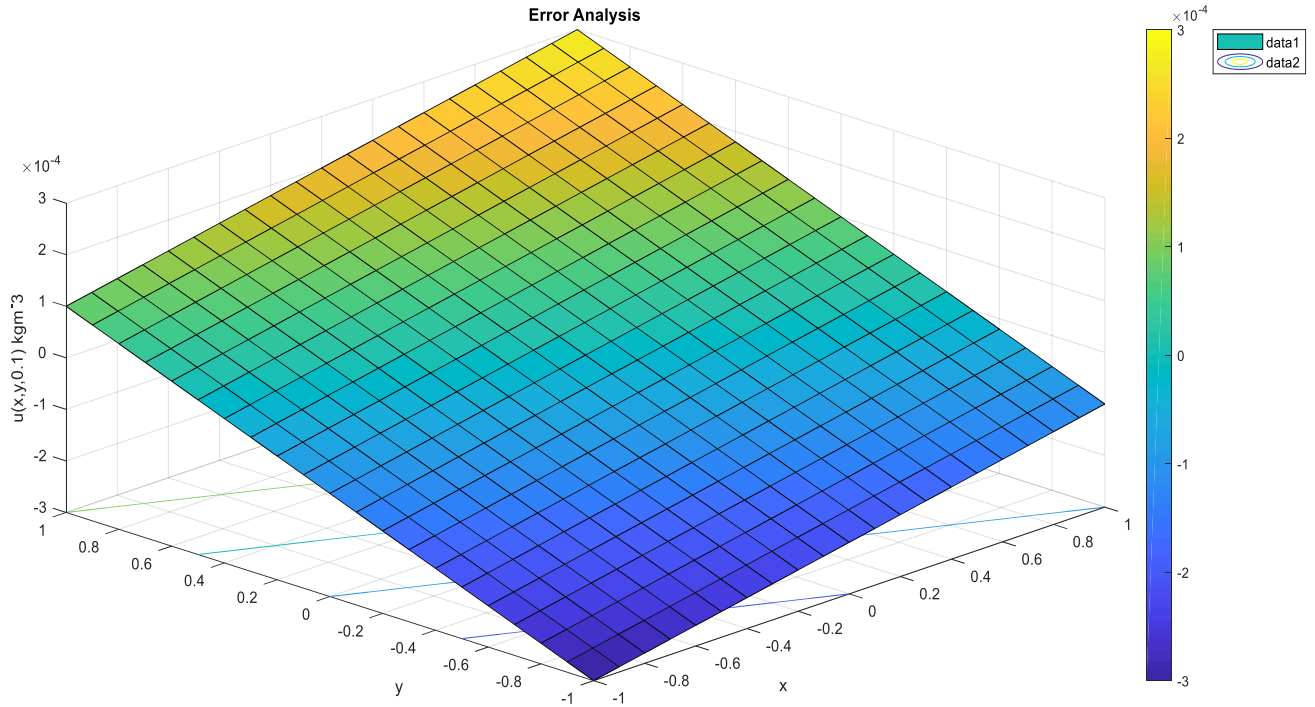


Figure 7. Surface and Contour plots of the Error between the exact and MFMOL solutions of Example 1 for the spatial distributions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.1$

In Table 1, the numerical results of flux-averaged (flowing) concentrations of example 1 using the proposed method MFMOL were compared with the exact solutions and the Variational Iteration Method (VIM) solution for effluent concentrations at varied temporal step size Δt and fixed position (x, y) for diffusion-dominated flow ($Pe = 1$), with initial and boundary conditions, in homogeneous porous medium. The proposed method MFMOL gave good approximations and stability than VIM solutions as shown with the computational values and the absolute errors e_1 and e_2 respectively. This established convergence of the new method MFMOL in simulating flowing concentration for 2D solute transport models in homogeneous and isotropic porous media. The sensitivity of the MFMOL values to spatial nodes (x, y) is observed in Table 1, although negligible, where the point $(0.5, 0.3)$ gave more

accuracy than the point $(0.7, 0.8)$ as shown by absolute errors e_1 . Figure 1 and Figure 2 showed the graphical displays of the flux-averaged concentrations of the MFMOL solutions and the exact solutions at fixed spatial points $(0.5, 0.3)$ and $(0.7, 0.8)$ respectively over various temporal step size Δt . Figure 3 and Figure 4 showed the combined surface and contour plots for the exact solution and the MFMOL solution respectively, for spatial distributions $u(x, y, t)$ of flux-averaged solute concentration at fixed $t = 0.1$. Figure 5 and Figure 6 showed the explicit contour plots for the exact solutions and the MFMOL solutions respectively for spatial distributions $u(x, y, t)$ of flux-averaged concentrations at fixed time $t = 0.1$. Figure 7 showed the combined surface and contour plots of the Error between the MFMOL solutions and the exact solutions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.1$.

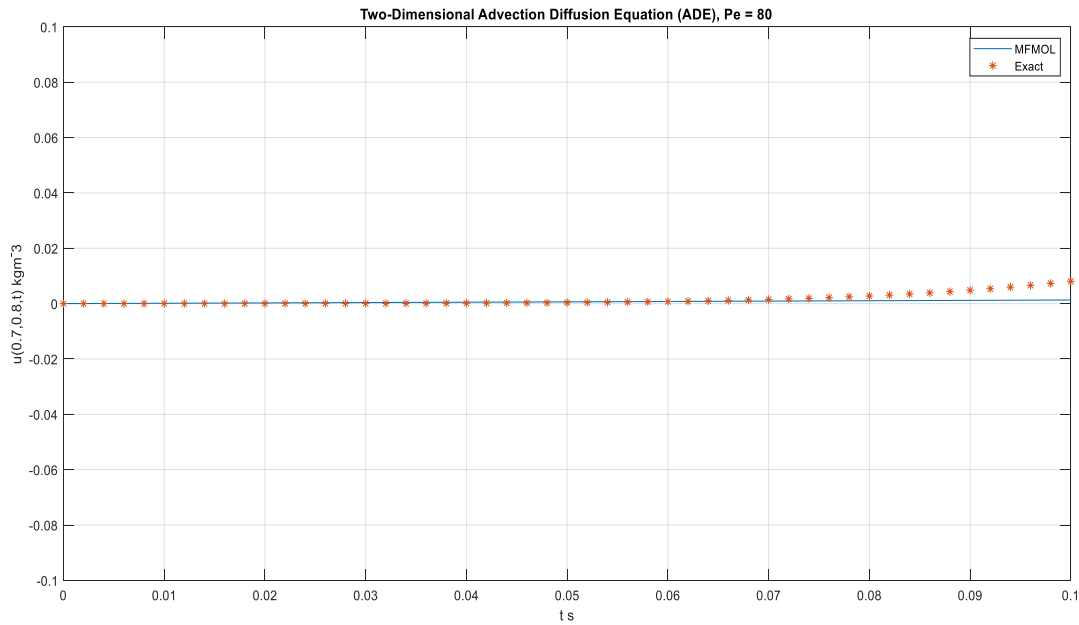


Figure 8. Plot of Example 2 corresponding to $(0.7, 0.8)$ over $0 \leq t \leq 0.1$ with $\Delta t = 0.002$

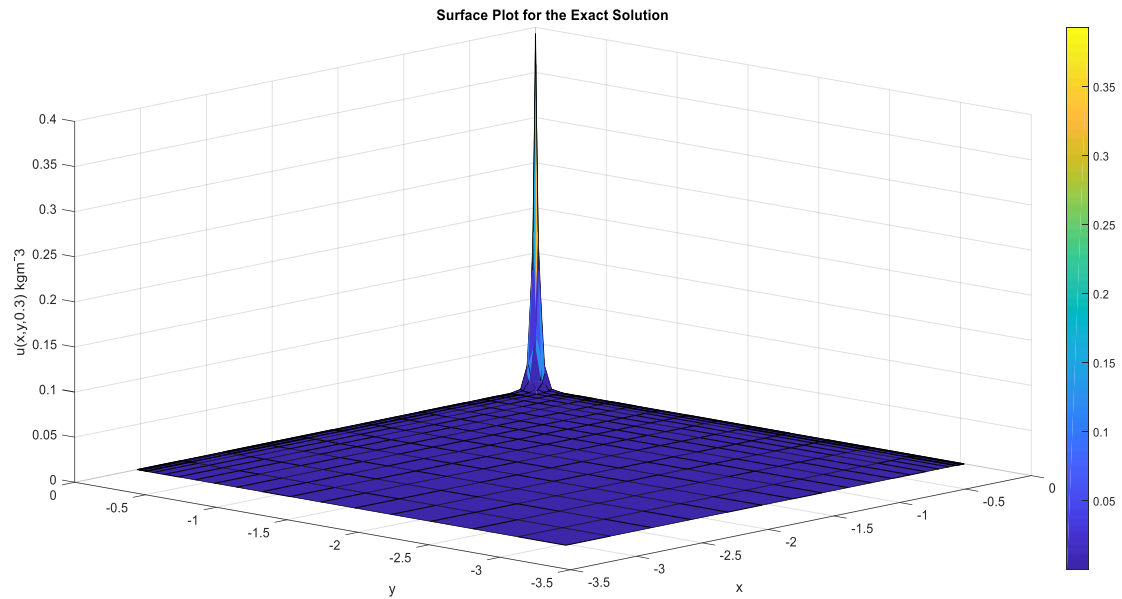


Figure 9. Surface Plots of Example 2 for the exact solutions of the spatial distributions of the flux-averaged solute concentration $u(x, y, t)$ at $t = 0.3$

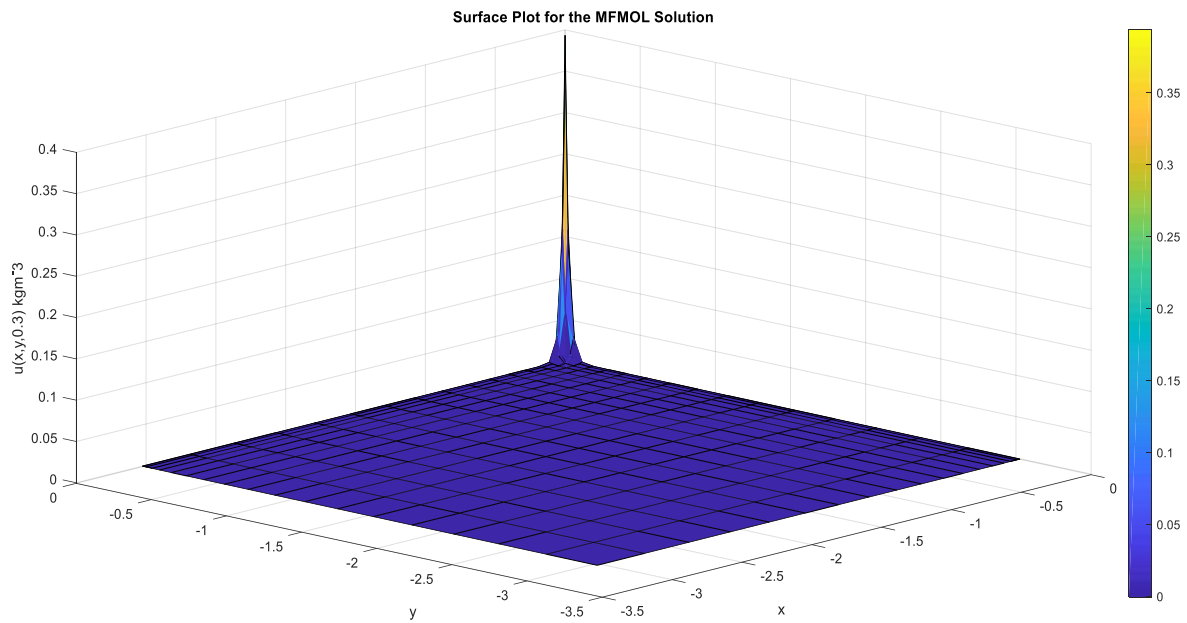


Figure 10. Surface Plots of Example 2 for the MFMOL solutions of the spatial distributions of the flux-averaged solute concentration $u(x, y, t)$ at $t = 0.3$

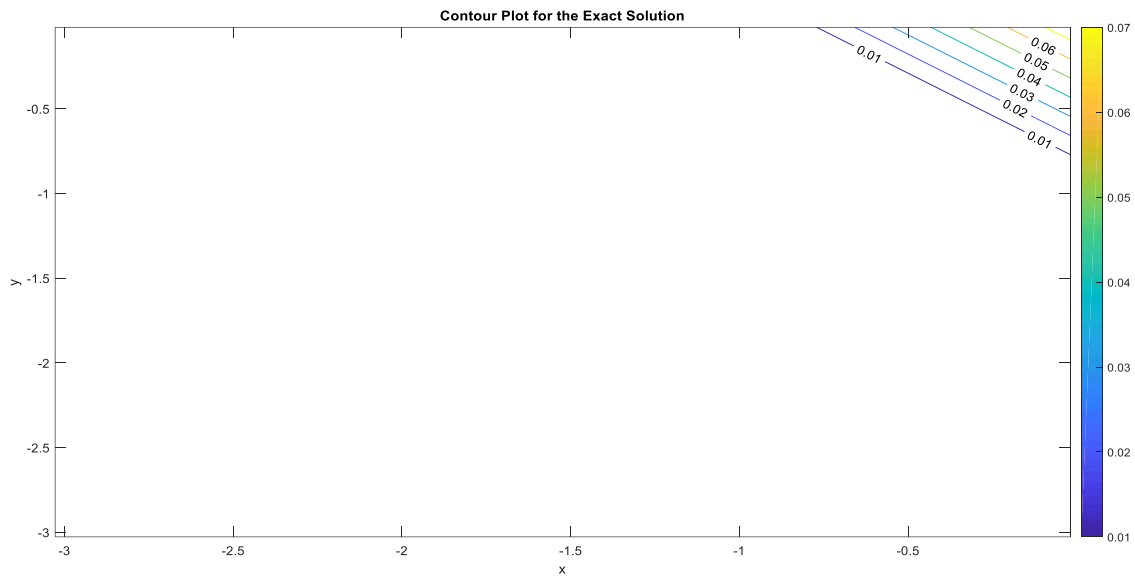


Figure 11. Contour plots of Example 2, for the exact solutions of spatial distributions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.3$

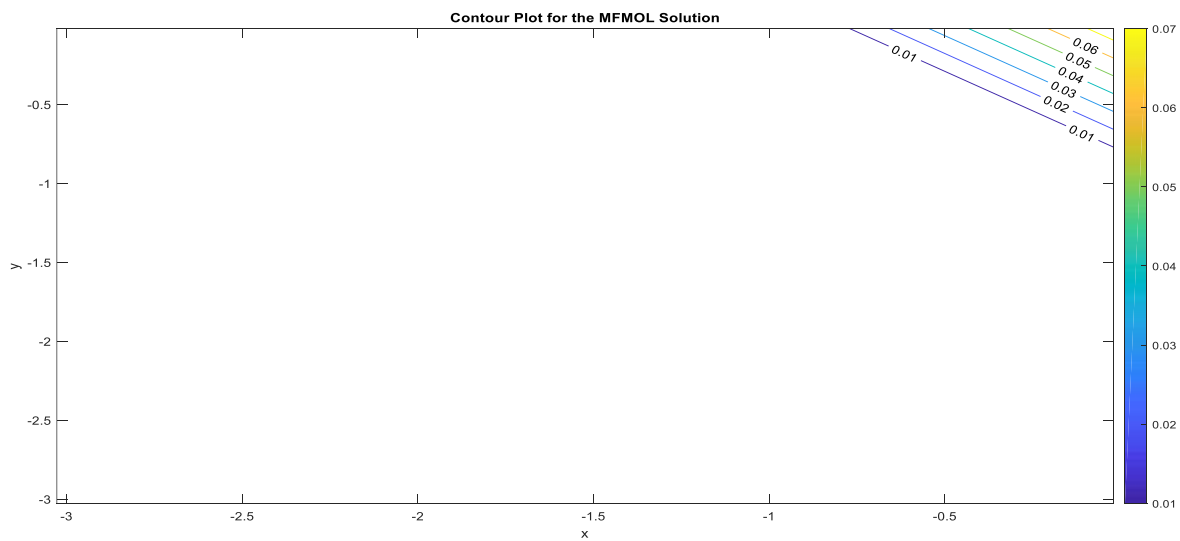


Figure 12. Contour plots of Example 2, for the MFMOL solutions of spatial distributions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.3$

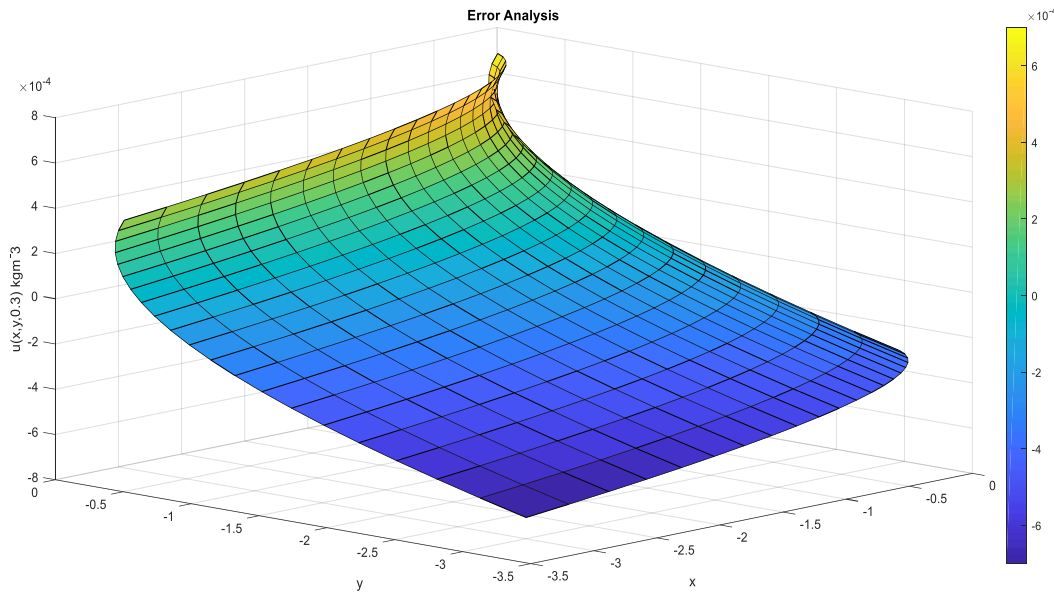


Figure 13. Surface plots of the errors between the exact and MFMOL solutions of Example 2 for the spatial distributions of the flux-averaged concentrations $u(x, y, t)$ at $t = 0.3$

In Table 2, the computational values of the proposed method MFMOL were compared with known exact solutions and the VIM solutions for the flowing concentrations at varied temporal step size Δt and fixed position (x, y) for advection-dominated flow ($Pe = 80$) subject to complex initial and boundary conditions in the complex domain. The proposed method resulted in good approximations and stability with minimal numerical oscillations without the use of a stabilization technique. Its accuracy here is almost the same or slightly higher than the VIM solutions as shown via the computational values and the absolute errors e_1 and e_2 respectively. The absolute errors computed also showed that the MFMOL is convergent for advection-controlled 2D solute transport models where the numerical oscillations or instabilities used to be significant. The sensitivity of the MFMOL values to spatial nodes (x, y) is observed in Table 2, although negligible, where the point $(0.7, 0.8)$ gave more accuracy than the point $(0.5, 0.3)$ as shown by the absolute errors e_1 . Also, higher time steps $\Delta t = 0.02$ were used for computational optimal accuracy instead of lower time steps often

used for better accuracy in transient flow models. The use of lower time steps is analogous to mesh refinement for optimal accuracy in mesh-based methods. Fig. 8 showed the graphical display of the flux-averaged concentrations of the MFMOL solutions and the exact solutions at fixed spatial points $(0.5, 0.3)$ over various temporal step size Δt . Figs 9 and 10 showed the surface plots of the exact and the MFMOL solutions respectively for spatial distribution $u(x, y, t)$ of the flux-averaged concentration fixed at $t = 0.3$. Figs. 11 and 12 showed the contour plots for the exact solutions and the MFMOL solutions respectively, for spatial distribution $u(x, y, t)$ of flux-averaged concentration fixed at $t = 0.3$. Fig 13 showed the surface plots of the error between the exact and the MFMOL solutions of the spatial distribution $u(x, y, t)$ for the flux-averaged concentration fixed at $t = 0.1$.

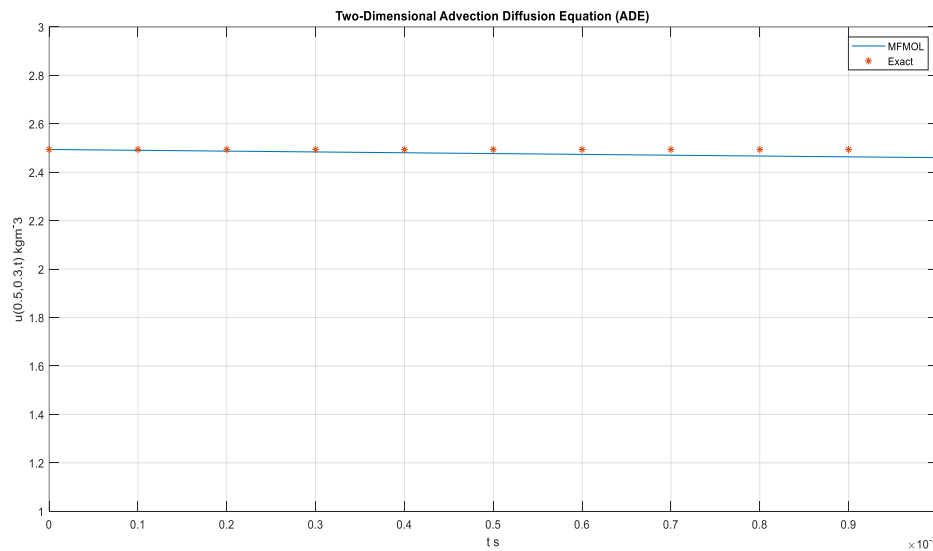


Figure 14. Plot of Example 3 corresponding to $(0.5, 0.3)$ over $0 \leq t \leq 0.001$ with $\Delta t = 0.0001$

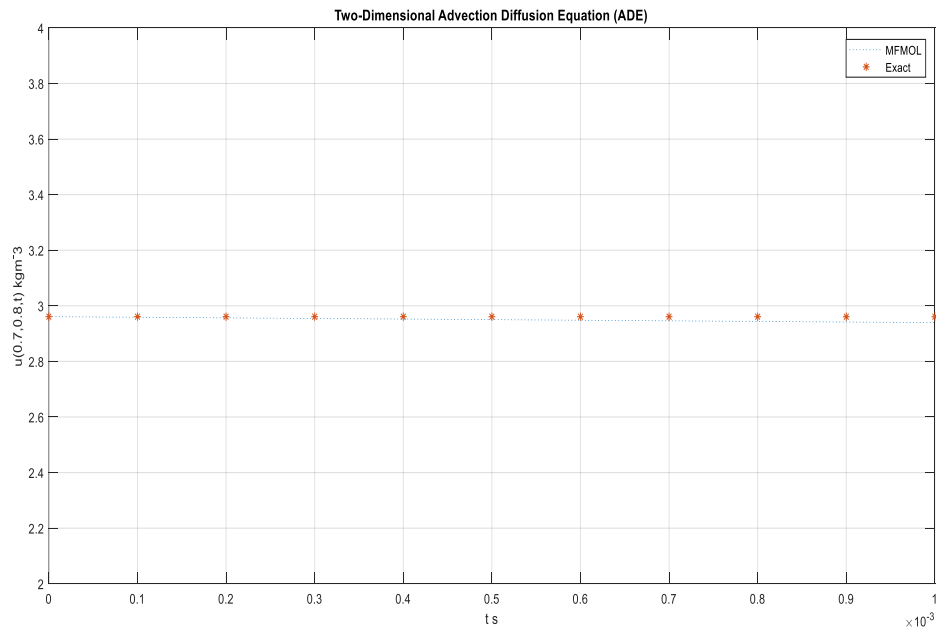


Figure 15. Plot of Example 3 corresponding to $(0.7, 0.8)$ over $0 \leq t \leq 0.001$ with $\Delta t = 0.0001$

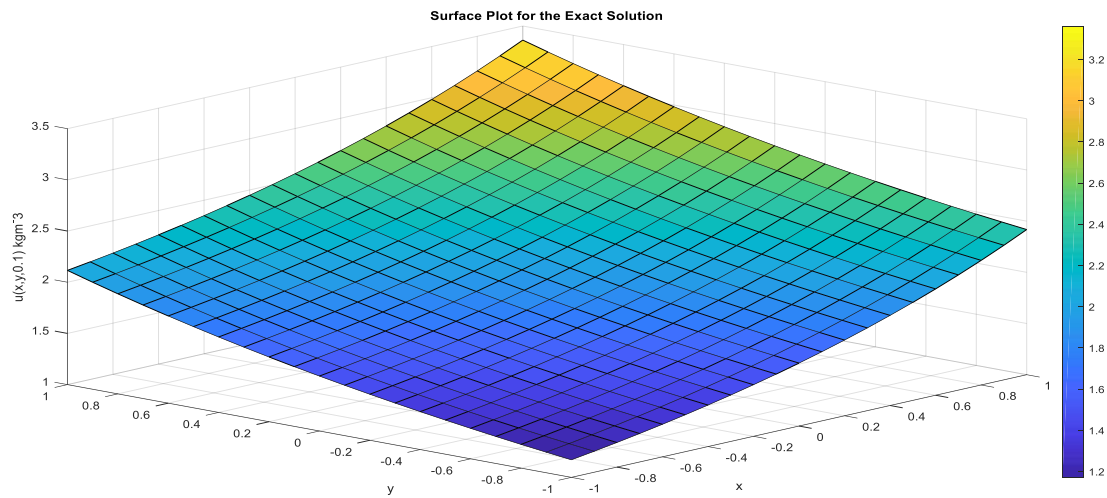


Figure 16. Surface plot of example 3 for the exact solution of the spatial distributions of the flux-averaged solute concentration $u(x, y, t)$ at $t = 0.1$

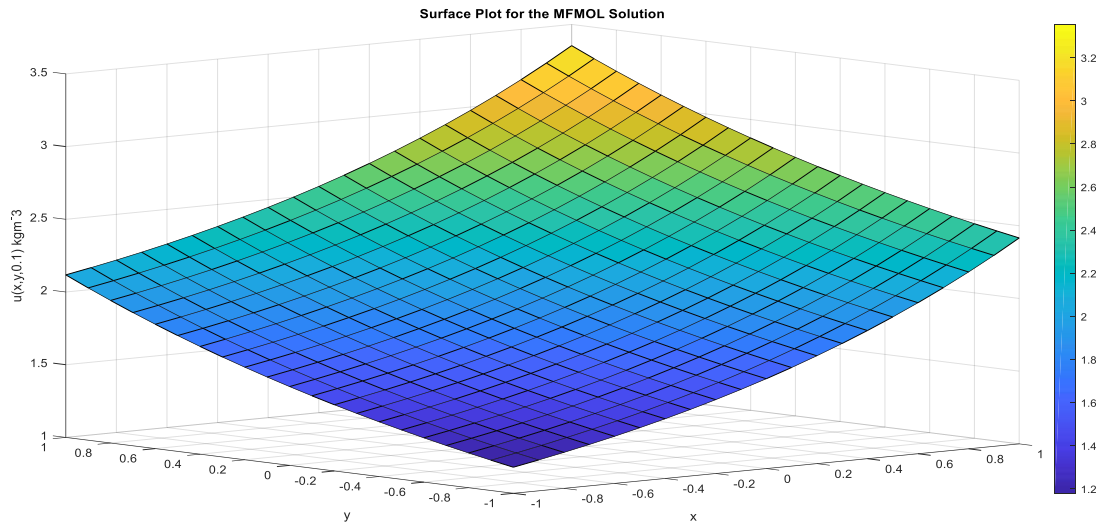


Figure 17. Surface plot of example 3 of the MFMOL solution of the spatial distributions of the flux-averaged solute concentration $u(x, y, t)$ at $t = 0.1$

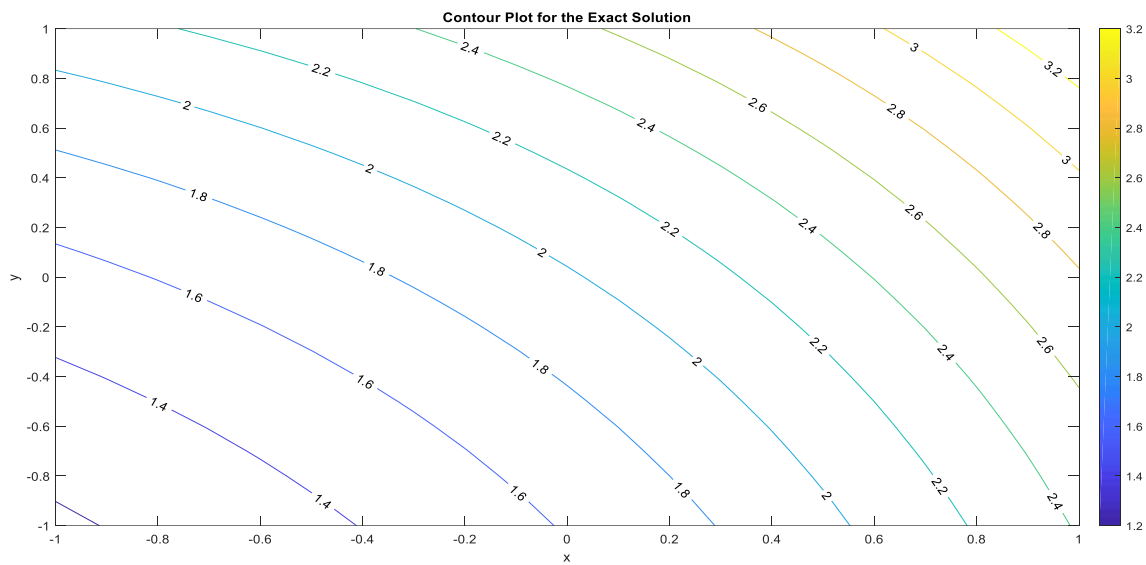


Figure 18. Contour plot of example 3, for the exact solutions of the spatial distributions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.1$

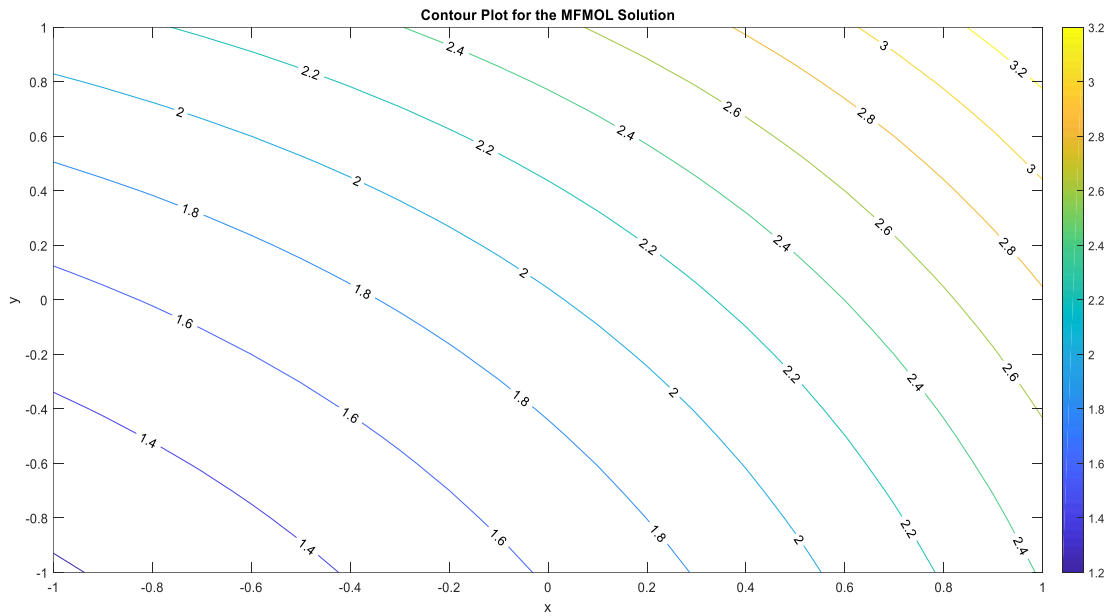


Figure 19. Contour plot of example 3, for the MFMOL solutions of the spatial distributions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.1$

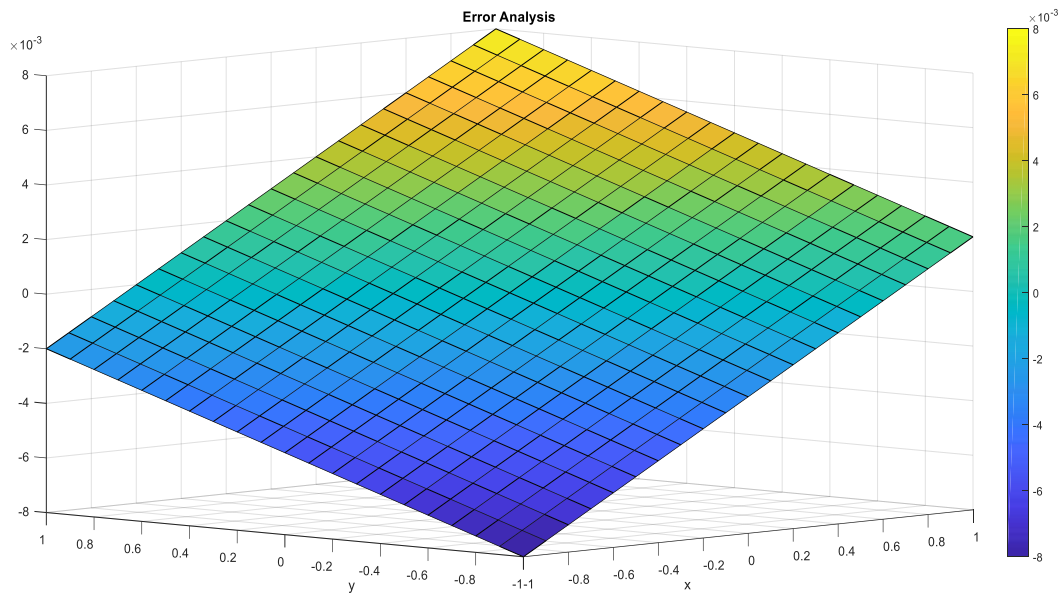


Figure 20. Surface plots of the error between the exact and MFMOL solutions of Example 3 for the spatial distributions of the flux-averaged concentration $u(x, y, t)$ at $t = 0.1$

In Table 3, the approximate values of the proposed MFMOL were compared with existing exact solutions and the VIM solutions for the effluent concentration at varied temporal step size Δt and constant position (x, y) for the case of diffusion-dominated flow ($Pe < 1$) with the complex initial and boundary conditions in the complex domain of porous structure. It is observed that the proposed method gave a good approximation, as validated by the exact solutions and VIM solutions. Here, the VIM solutions are similar to the exact solutions. Also, the sensitivity of the solutions is independent of both spatial positions (x, y) and the lower time steps Δt ; rather, it depends on the nature of the subsidiary conditions. The absolute errors computed also showed that the MFMOL is convergent for diffusion-controlled 2D solute transport models. Figs. 14 and 15 showed the graphical displays of the flux-averaged concentrations of the MFMOL solutions and the exact solutions at fixed spatial points $(0.5, 0.3)$ and $(0.7, 0.8)$, respectively, over various temporal step sizes Δt . Figs. 16 and 17 showed the surface plots of the exact and the MFMOL solutions for spatial distribution $u(x, y, t)$ of the flux-averaged concentrations respectively at $t = 0.1$. Figs. 18 and 19 showed the contour plots for the exact solutions and the MFMOL solutions for spatial distribution $u(x, y, t)$ of flux-averaged concentration respectively at $t = 0.1$. Fig. 20 showed the surface plots of the error between the exact and the MFMOL solutions of the spatial distribution $u(x, y, t)$ for the flux-averaged concentration fixed at $t = 0.1$.

The Variational Iteration Method (VIM) is a powerful computational method for various problems in the applied sciences and engineering, due to its stability and fast convergence to closed-form solutions or approximate solutions in series form [34], [35]. However, it has limited capacity in solving boundary value problems as its zeroth approximation depends only on the initial conditions, and its solution may not satisfy the boundary conditions. This setback is overcome by the proposed method MFMOL, which utilises the Kronecker delta function property to easily impose boundary conditions, making the present method superior to VIM in solving a wide range of boundary value problems.

4 Conclusion

In this paper, a novel Meshfree Method of Lines (MFMOL) was proposed in strong form formulations for solving 2D Advection Diffusion Equations (ADEs) modelling flux-averaged (flowing) concentrations of solute transport in nonreactive, source-free, homogeneous and isotropic porous media. The method was applied to both advective ($Pe > 1$) and diffusive ($Pe < 1$) problems without any stabilization technique for the convective terms of the models, and the results agreed with both existing exact solutions and the Variational Iteration Method (VIM) solutions, with negligible numerical instabilities as shown in Tables 1-3 and Figs. 1–20. This shows the accuracy and efficiency of the new method, not only in simulating 2D solute transport in homogeneous complex domains without a stabilization technique, but also in overcoming the various challenges ranging from low-quality meshes, numerical instabilities, discontinuities, massive computational efforts, high costs, to substantial setup time in generating quality meshes, which are encountered by mesh-based methods due to mesh interpolation for the complex domains. Regarding the easy imposition of boundary conditions, which VIM may not always satisfy in solving solute transport models, the present method shows superior performance over VIM by utilizing the Kronecker delta function property to impose boundary conditions for numerous models easily.

This research acknowledges some limitations for optimal computational accuracy of the proposed method. These include the geometrical location of star nodal points in the domain of the problem and its boundary, the time steps and the complexities of the initial and boundary conditions. The sensitivity of the limit factors to the accuracy of the solution is trivial with negligible effects on computational errors.

Despite these limitations, various features of the proposed method involved in generating acceptable computational results have established its solid foundation for handling 2D models with constant parameters in homogeneous and isotropic porous domains. Considering the effectiveness of the proposed method for 2D models, future research is to

explore its capacity for solving 3D models, coupled transport-reaction systems and various time-dependent models such as series RC and LC circuits in electricity and magnetism [36], Magnetohydrodynamic models (MHD) [37], and others. Additionally, with the series of transformation [38], [39], appropriate discretisation of variable-dependent parameters in the models [40] or other approaches, future work will examine the significant potential of the proposed method in solving transport models in heterogeneous and anisotropic porous media for practical applications such as groundwater flow, reservoir flow, shallow water equation, and others.

Acknowledgement

The work presented in this paper was supported by the Federal Government of the Republic of Nigeria through PTDF/ED/ISS/PHD/JAK/2045/21.

References

- [1] M. Zhinjuan, C. Xiaofei, and M. Lidong, "A Hybrid Interpolating Meshless Method for 3D Advection Diffusion Problems," *Mathematics*, MDPI, 2022.
- [2] Z. Yongou, D. Sina, L. Wei, and C. Yingbin, "Performance of the radial point interpolation method (RPIM) with implicit time integration scheme for transient wave propagation dynamics," *Computers & Mathematics with Application*. Elsevier. Vol. 114, pp. 95-111, 2022.
- [3] E. Oñate and M. Manzan, "A general procedures for deriving stabilized space-time finite elements method for advective-diffusive problems," *Int. J. Numer. Meth. Fluids*. Vol. 31, pp. 203-221, 2015.
- [4] J. Carrer, C. Curha and W. J. Mansur, "The boundary element method applied to the solution of two-dimensional diffusion-advection problems for non-isotropic materials," *J. Braz. Soc. Mech. Sci*. Vol. 39, pp. 4533-4545, 2017.
- [5] E. Barbieri and N. Petrinic, "Three-dimensional crack propagation with distance-based discontinuous kernels in meshfree methods," *Comput. Mech*. Vol. 53, pp. 325-342, 2014.

- [6] P. N. Vinh, R. Timon, B. Stephane, and D. Mac, “Meshless methods: A review and computer implementation aspects,”. *Mathematics and Computers in Simulation*. Elsevier. Vol. 79. pp. 763-813, 2008.
- [7] Q. Wu, P. P. Peng and Y. M. Cheng, “The Interpolating element-free Galerkin method for elastic large deformation problems,”. *Sci. China*. Vol. 64, pp. 364-374, 2021.
- [8] H. Sirajul, H. Arshad and U. Maryam, “RBFs Meshless Method of Lines for the Numerical Solution of Time-Dependent Nonlinear Coupled Partial Differential Equations,” *Applied Mathematics*, Scientific Research, 2001.
- [9] T. Belytschko, Y. Y. Lu, and L. Gu, “Element-free Galerkin Methods. *International Journal for Numerical Methods in Engineering*,” vol. 37, no. 2, pp. 229–256, 1994.
- [10] M. Zhijuan, C. Xiaofei and M. Lidong, A Hybrid Interpolating Meshless Method for 3D Advection-Diffusion Problems. *Mathematics*, MDPI, 2022.
- [11] B. Nayroles, G. Touzot and P. Villon, “Generalizing the finite element method: Diffuse Approximation and diffuse elements.” *Comput. Mech*. Vol. 10, no. 5, pp 307-318, 1992.
- [12] G.R. Liu and Y.T. Gu, “A point interpolation method for two-dimensional solids,” *Int. J. Numerical Methods Eng.*, Vol. 50, pp. 937-951, 2001.
- [13] C. Heng and Z. Guodong, “Analyzing 3D Advection-Diffusion Problems by Using the Improved Element-Free Galerkin Method,” *Mathematical Problems in Engineering*, Hindawi, 2020.
- [14] M. Mategaonkar and T.T. Eldho, “Two dimensional contaminant transport modelling using meshfree point collocation method (PCM),” *Engineering Analysis with Boundary Elements*. Elsevier, 2012.
- [15] B. Nagina, I. A. T. Syed, and H. Sijarul, “Meshless Method of Lines for Numerical Solution of Kawahara Type Equations,” *Applied Mathematics*. Scientific Research, 2011.
- [16] S. N. Atluri and T. Zhu, “A new meshless local Petrov–Galerkin (MLPG) approach in Computational mechanics,” *Comput. Mech.*, Vol. 22, pp. 117–127, 1998.
- [17] E. Oñate, S. Idelsohn, O. C. Zienkiewicz, and R. L. Taylor, “A finite point method in computational mechanics applications to convective transport and fluid flow,” *Int. J.*

Numerical Methods Eng., Vol. 39, pp. 3839–3866, 1996.

- [18] G.R. Liu, “*Meshfree Methods: Moving beyond the Finite Element Method*,” Taylor and Francis Group, LLC, New York. 2010.
- [19] M. B. Liu, G. R. Liu, and Z. Zong, “An overview on smoothed particle Hydrodynamics,” *Int. J. Comput. Methods*, Vol. 5, no 1. pp. 135–188, 2008.
- [20] [20] W. K. Liu, S. Jun and Y. Zhang, “Reproducing kernel particle methods,” *Int. J. Numerical Methods Fluids*, Vol. 20, pp. 1081–1106, 1995.
- [21] J. G. Wang and G. R. Liu, “A point interpolation meshless method based on radial basis Functions,” *International Journal for Numerical Methods in Engineering*, 2002.
- [22] M. J. D. Powell, “*The theory of radial basis function approximation in 1990*, in *Advances in Numerical Analysis*,” F. W. Light, ed., Oxford University Press, Oxford, pp. 303–322, 1992.
- [23] H. Wendland, “Error estimates for interpolation by compactly supported radial basis functions of minimal degree,” *J. Approximation Theory*. Vol. 93. pp 258–396. 1998
- [24] E. J. Kansa, “A scattered data approximation scheme with application to computational Fluid dynamics I & II, *Comput. Math. Appl.* Vol. 19, pp. 127–161, 1990.
- [25] Z. Yandaizi, W. Tielin and Z. Jesse, “Development of gas-solid fluidization: Particulate and aggregative,” *Powder Technology*. Elsevier. Vol. 421, 2023.
- [26] S. Das and T. I. Eldho, “A Meshless Weak-Strong Form Method for the Simulation of Coupled Flow and Contaminant Transport in an Unconfined Aquifer,” *Transport in Porous Media*. Springer. Vol. 143, pp. 703 – 737, 2022.
- [27] E. Oñate, S. Idelson, O. C. Zienkiewicz, R. L. Taylor R. L. and C. Saco, “A Stabilized finite point method for analysis of fluid mechanics problems,” *Comp. Methods Appl. Mech and Eng*, Elsevier. Vol. 139. pp. 315 – 346, 1996.
- [28] E. Oñate, “Finite Increment Calculus (FIC): a framework for deriving enhanced Computational methods in mechanics,” *Advanced Modelling and Simulation in Engineering Science*, CIMNE, Barcelona, Spain, 2016.

- [29] A. Aatish and T. I. Eldho, "Simulation of flows and Transport process – Scope of Meshless Methods: Recent Advances in Computational Mechanics," *Lecture Notes in Mechanical Engineering*. Springer. Singapore, 2021.
- [30] Q. Xinqiang, H. Gang and P. Gaosheng, "Finite Point Method of Nonlinear Convective Diffusion Equation," *Filomat*. Vol. 5, no 34, pp. 1517 – 1533, 2020.
- [31] A. R. Appadu and H. H. Gidey, "Time-Splitting Procedures for the Numerical Solution of the 2D Advection-Diffusion Equation," *Mathematical Problems in Engineering*, Hindawi Publishing Corporation. Article ID 634657, 2013.
- [32] J. A. Kazeem, "A Meshfree Method for the Solutions of Multiphase Flow Problems in Porous Media," Unpublished doctoral thesis, University of Abuja, F.C.T., Nigeria, 2025.
- [33] A. M. S. Muhannad and F. J. Borhan, "Numerical Solutions Based on Finite Difference Technique for Two-Dimensional Advection-Diffusion Equation," *British Journal of Mathematics and Computer Science*. Article no: BJMCS.25464, 2016.
- [34] W. Abdul Majid, "*Linear and Nonlinear Integral Equation: Methods and Applications*," Springer-Verlag, New York, (2011).
- [35] N. Okiotor, F. Ogunfiditimi, and M. O. Durojaye, "On the computation of the lagrange multiplier for the variational iteration method (VIM) for solving differential equations," *Journal of Advances in Mathematics and Computer Science*, Vol. 35, no 3, pp. 74-92, 2020.
- [36] H. P. Albertus, "The Lagrangian and Hamiltonian for RLC Circuit: Simple Case," *International Journal of Applied Sciences and Smart Technologies*. Vol. 2, no 2, pp. 169–178, 2020.
- [37] D. R. Kirubaharan, A. D. Subhashini, N. V. N. Babu, and G. Murali, "Quantitative Analysis of Magnetohydrodynamic Sustained Convective Flow via Vertical Plate," *International Journal of Applied Sciences and Smart Technologies*. Vol. 6, no 2, pp. 407– 416, 2024.
- [38] M. K. Singh, P. Singh and V. P. Singh, "Analytical Solution for Two-Dimensional Solute Transport in Finite Aquifer with Time-Dependent Source Concentration," *Journal of Engineering Mechanics*. Vol 136, no 10, 2010.

- [39] P. Singh, "One Dimensional Solute Transport Originating from a Exponentially Decay Type Point Source along Unsteady Flow through Heterogeneous Medium," *Journal of Water Resource and Protection*. Vol. 3, pp. 590-597, 2011.
- [40] G. D. Hutomo, J. Kusuma, A. Ribal, A. G. Mahie and N. Aris, "Numerical solution of 2-dadvection-diffusion equation with variable coefficient using du-fort frankel method," *J. Phys: Conf. Ser.*1180 012009, 2019.

This page intentionally left

Coconut Shell-Based Briquettes for Sustainable Energy: A Bibliometric Study on Biomass Mixtures and Binder Materials

Ridho Darmawan¹, Awaly Ilham Dewantoro^{1,2},
Efri Mardawati^{1,2*}

¹*Department of Agro-Industrial Technology, Faculty of Agro-Industrial
Technology, Universitas Padjadjaran, Jatinangor 45363, Indonesia*

²*Research Collaboration Center for Biomass and Biorefinery between BRIN
and Universitas Padjadjaran, Jatinangor 45363, Indonesia*

**Corresponding Author: efri.mardawati@unpad.ac.id*

(Received 11-06-2025; Revised 23-06-2025; Accepted 08-07-2025)

Abstract

Coconut shell is one of the potential biomass resources that has been widely developed as a raw material for briquettes to support renewable energy initiatives and circular economy practices. This study aimed to explore the development, research focus, and future directions of coconut shell briquette through a systematic literature review. The *Methodi Ordinatio* approach was employed for analysis, resulting in a final portfolio of 134 selected documents, which were further examined to identify trends and research gaps. The findings showed that mixing coconut shell with other biomass such as from wood-based and agricultural-based residues could enhance the briquette performance. Moreover, alternative binders such as lignocellulosic carbohydrates and its derivatives, plant sap, and waste cooking oil offered promising substitutes for food-based materials. Oily biomass, such as eucalyptus wastes and pine resin, was also found to improve briquette performance due to its volatile content. In addition, the integration of automation technologies based on microcontrollers and the Internet of Things (IoT) has been applied to improve production efficiency and consistency. The findings of this study can serve as a foundation for future development focused on material formulation and technological innovations for coconut shell-based briquette production that are more efficient, sustainable, and responsive to future energy needs.

Keywords: Alternative binders, Bibliometric analysis, Biomass mixtures, Coconut shell briquettes.

1 Introduction

Global energy demand continues to rise in line with population growth and the pace of industrialization [1]. The reliance on fossil fuel sources has resulted in various negative environmental impacts, thereby encouraging the search for more sustainable energy



alternatives. Renewable energy sources such as solar, wind, hydro, and biomass are being widely developed to reduce the strain on finite fossil resources. Among these, biomass holds a strategic position due to its renewable and widespread availability, particularly in the form of organic residues [2]. One notable application of biomass that has gained increasing attention is the production of biomass briquettes—commonly referred to as bio-briquettes—an alternative solid fuel that is environmentally friendly, cost-effective, and suitable for replacing conventional fossil-based energy sources [3].

The development of biomass briquettes has shown significant progress through the diversification of raw materials. Various types of biomass—such as rice husk, sugarcane bagasse, wood residue, and coconut shell—have been utilized as primary feedstocks for briquette production [4–6]. Among these options, coconut shell stands out as a leading biomass source due to its high density, favorable fixed carbon content, and relatively high calorific value [7–8]. Moreover, the utilization of coconut shell supports the principles of biorefinery and circular economy by transforming organic waste into value-added alternative energy, while simultaneously reducing carbon emissions and waste volume [9–10].

Recent research trends on coconut shell biomass briquettes have become increasingly diverse and innovative. The focus of studies has expanded beyond basic characterization to include the exploration of production technologies, biomass blend formulations, and the development of sustainable natural binders [11–13]. Several recent studies have evaluated the use of agricultural and plantation waste—such as oil palm empty fruit bunches, corn stalk, and rice straw—as additives in coconut shell briquettes [14–16]. In addition, emerging research has begun to integrate technological approaches, such as microcontroller-based automation and the Internet of Things (IoT), into the briquette production process [17–20]. These developments highlight the need for a comprehensive literature synthesis that can consolidate scientific findings and identify research gaps that remain open for further investigation.

A systematic approach such as *Methodi Ordinatio* is considered effective in addressing the need to evaluate the continuously evolving research landscape [21–22]. This technique combines criteria of quality, citation count, and publication recency to construct a representative portfolio of documents [23]. By applying *Methodi Ordinatio*,

the review findings can be developed into a visual map network that illustrates research cluster, dominant topics, and interrelated concepts. The use of this technique offers significant advantages in identifying research gaps and formulating strategic directions for future research development [24].

The aim of this study was to systematically explore and sythesize the developmet of research on coconut shell-based briquettes using the *Methodi Ordinatio* approach. This review also integrated bibliometric analysis based on a visual map network to obtain a comprehensive overview of research trends, scientific contributions, and potential future research directions. The findings of this study are expected to provide a strong scientific foundation for the development of more efficient, sustainable, and applicable coconut shell briquette innovations.

2 Material and Methods

2.1 Data Collection

Methodi Ordinatio, was used to structure systematic literature review, with data collection limited to the Scopus database [21, 23]. The main procedural steps are presented in Table 1. Data collection was conducted at the middle of January 2025 employing search query as “**briquette**” and “**coconut**”, which outlined 157 documents with a total of 63 of keywords. The collected documents were the sorted by publication year—limited to the last decade (2015 to 2024)—and duplicate entries were removed using Mendeley as the reference manager, outlined as 142 documents. A thorough screening and reading of all documents were counducted to filter those relevant to this study. The final portfolio consisted of 134 documents and 25 keywords, which were further evaluated through bibliometric analysis and visualized using a visual map network.

Table 1. Main steps to perform the systematic review through *Methodi Ordinatio*.

Steps	Description	
Database search	I	Initial portfolio
	Database	Scopus
	Keywords	63

		Number of documents	157
Filtering procedures	II	Elimination of duplicates and limiting to recent 10 years publication (2015 to 2024)	
	III	Screening title and keywords	
	IV	Reading of abstracts	
	V	Reading of full texts	
		Final portfolio	134
		Number of documents	25
		Keywords	
Content analysis	VI	Year of publication	
		Types of document	
		Keywords	

2.2 Bibliometric Analysis and Visual Maps Generation

Bibliometric analysis was conducted to synthesize the information obtained from the final portfolio (comprising 134 documents and 25 keywords), following a modified set of procedures [23–24]. The final portfolio was used to generate a visual map network with the assistance of VOSviewer v1.6.20, and the results served as a reference for conducting the systematic review in this study. the subsequent discussion was focused on key and critical parameters (based on visual map network), examined comprehensively to map potential directions for future development.

3 Results and Discussions

3.1 Visual Map Network and Bibliometric Results

Recent studies on the development of coconut shell-based briquettes has shown an increasing trend each year, as illustrated in Fig. 1. Discussion on combustion and the application of briquettes for stoves or cookers began in 2015 [25]. In subsequent years, studies began to focus on the engineering and optimization of the production process [26–27]. Starting in 2017, the addition of other biomass—such as tropical fruit waste, rice husks, water hyacinth, and wood residues—was explored to enhance the performance and

characteristics of coconut shell-based briquettes [28–31]. From 2020 onwards, intensive studies were conducted on substituting binder materials, introducing alternatives such as plant sap and waste cooking oil (WCO) [13, 32–33]. More recent developments, beginning in 2023, have shifted toward techno-economic studies and environmental impact analyses, aimed at promoting more efficient practices and supporting the implementation of a circular economy [16, 34–38]. This is evidenced by the integration of internet- and microcontroller-based automation, as well as the emergence of agro-industrial and animal manure as alternative briquette binders [17–18, 39–41].

Fig. 2 presents the distribution of publication types recorded over the past decade (2015–2024). Based on the final portfolio, experimental studies dominated the landscape with 67 documents, followed by conference papers with 60 documents. These figures indicate that research on the development of coconut shell-based briquettes has been predominantly conducted through direct experimentation to obtain primary data as sources of new scientific findings. Meanwhile, publication types such as reviews (4 documents) and book chapters (3 documents), highlight a lack of interest in summarizing and synthesizing research gaps related to coconut shell-based briquettes development. This remains a critical need, as expert perspective are essential in guiding future directions for the development of environmentally friendly, high-performance, and economically valuable coconut shell-based briquettes.

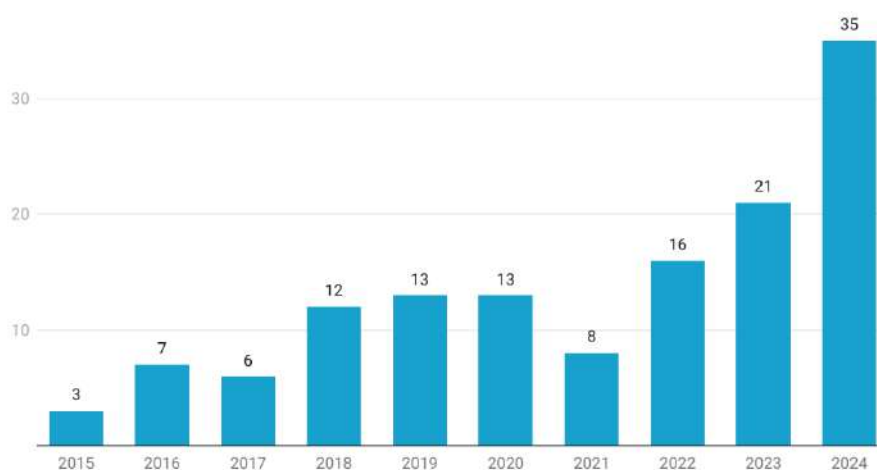


Figure 1. Distribution of the number of published documents based on years of publication over the past of decade (2015 to 2024).

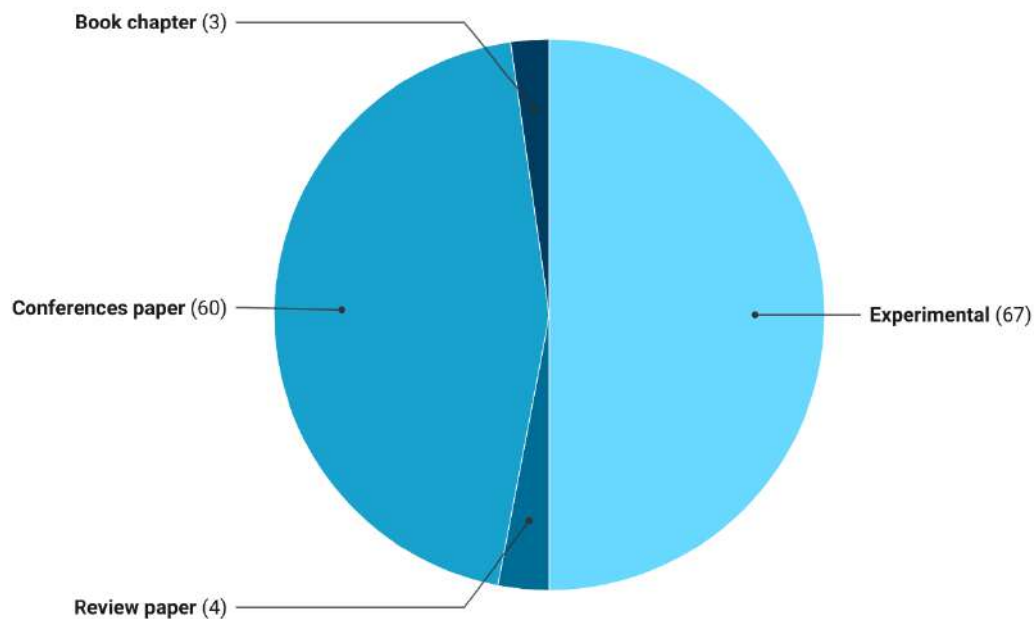


Figure 2. Distribution of the number of published documents based on publication categories over the past of decade (2015 to 2024).

Subsequently, the research trends in coconut shell-based briquettes developments—after filtered and screened into the final portfolio of 134 documents—were visualized as a visual map network, as shown in Fig. 3. Five distinct clusters were identified, each representing key research over the past decade: (i) briquette characteristics and performance; (ii) sustainable development; (iii) energy sources and utilization; (iv) adhesives or binders; and (v) recent developments. A noteworthy insight revealed through bibliometric approach is the growing emphasis, particularly in the past five years, on investigating biomass mixtures and binder materials used in the production of coconut shell-based briquettes. Several binder materials identified from the final portfolio include plat sap, WCO, agro-industrial waste, and animal manure—all of which have been shown to enhance the calorific value of coconut shell-based briquettes [28, 32–33, 40, 42–43].

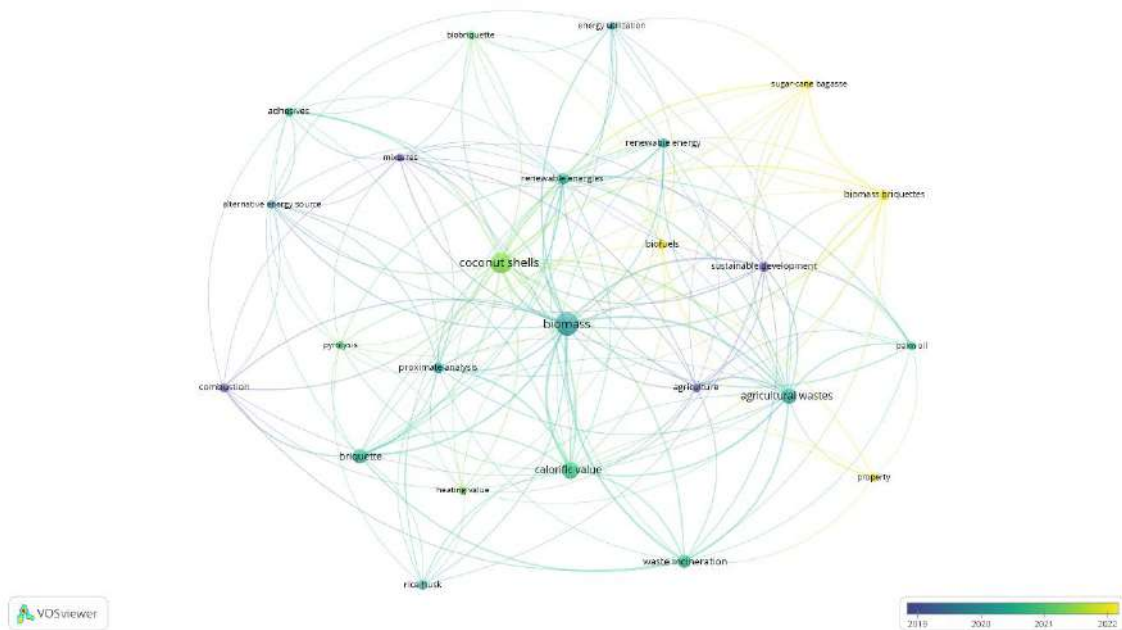


Figure 3. Visual map network of the developments on coconut shell-based briquettes over the past decade (2015 to 2024).

3.2 An Overview on Biomass Mixtures for Coconut Shell-Based Briquettes

Various types of biomass have been successfully combined in the development of coconut shell-based briquettes, including rice husks, sugarcane bagasse, and oil palm residues. As shown in Fig. 4, new findings have emerged involving additional biomass combinations such as rice straw, wood residues, and agro-industrial waste. Coconut shell-based briquettes blended with wood residues—such as rubber rods, sawdust, and eucalyptus wastes—demonstrated high performance, with calorific value ranging from 6,300–8,000 cal/g [44–46]. This is attributed to the denser structure of wood residues compared to other agro-industrial biomass, particularly in terms of lignin content [47]. Various wood residues are known to contain at least 25% lignin within their lignocellulosic structure, whereas other agro-industrial biomass sources typically contain lower lignin content, ranging from 10–25% [48]. This phenomenon is further supported by data in Table 2, which shows that coconut shell-based briquettes combined with agro-industrial residues—such as straw, husks, and bagasse—have calorific values below 6,000 cal/g.



Figure 4. Various biomass mixtures on the developing of coconut shell-based briquettes based on visual map network and bibliometric analysis findings.

Overall, the proportion of coconut shell charcoal is generally higher than that of other biomass-derived charcoals. A greater proportion of coconut shell charcoal, particularly when combined with rice husk charcoal, has been shown to enhance briquette performance. For instance is the results from Rodiah [32], that reported a higher calorific value of 5,257 cal/g with a 2:1 ratio of coconut shell-to-rice husk charcoal, compared to 1:1 ratio which yielded a calorific value below 5,000 cal/g. However, despite its high calorific value, the proximate analysis showed an excessively high volatile matter, approximately 48.99%. In addition, several studies that employed higher ratios of coconut shell charcoal—combined with rice straw, eucalyptus wastes, or coffee grounds—did not improve the proximate characteristics or overall performance of the briquette [11, 45, 54]. These findings highlight the need for further investigation based on the compiled data in this study

Table 2. Biomass mixtures on the recent developments of coconut shell-based briquettes.

Types of Biomass	Ratio ^a	MC ^b	VM ^b	AC ^b	FC ^b	CV ^c	Ref.
Rice straw	1:4	n.d.	n.d.	n.d.	n.d.	4,580	[11]
Rice husks	1:1	7.63	23.12	20.21	49.04	4,966	[49]
	1:1	7.52	23.40	21.16	47.92	4,886	[50]
	1:2	3.55	48.99	3.41	44.05	6,257	[32]
Sawdust	3:1	2.73	t.d.	1.75	t.d.	7,054	[51]
	3:1	4.23	t.d.	1.76	t.d.	7,561	[12]
	1:1	5.63	38.77	1.87	59.36	6,673	[52]
	1:1	1.00	23.44	8.28	67.28	7,564	[3]
Madan wood wastes	1:1	11.6	13.4	4.2	70.8	6,345	[44]
Rubber rods	1:1	5.87	4.76	3.56	86.81	8,082	[46]
Eucalyptus wastes	1:4	13.99	82.50	2.66	14.02	7,814	[45]
Corn cob, Sugarcane bagasse	6:10:20	4.14	n.d.	n.d.	n.d.	5,945	[53]
Coffee grounds	1:3	7.84	36.24	9.38	46.54	4,637	[54]
Durian peel	3:2	13.83	1.69	11.67	72.81	5,284	[55]
Rambutan peel	1:9	4.05	t.d.	4.83	52.36	6,673	[56]

^a The ratio of other biomass as mixture to coconut shell charcoal.

^b Unit expressed in percent (%); MC = moisture content; VM = volatile matter; AC = ash content; and FC = fixed carbon.

^c Calorific value (CV) expressed in calory per gram (cal/g).

n.d. means the parameter was not determined on the relevant studies.

Interestingly, brquettes with other types of biomass that did not mentioned yet showed higher performance, with calorific value exceeding 7,000 cal/g, particularly in the case of sawdust charcoal addition at a 3:1 ratio [12, 51].

Based on studies on briquette performance modelling, volatile matter and fixed carbon are the proximate characteristics that significantly influence the calorific value of briquettes [57–59]. This has been demonstrated in study bu Handayani *et al.* [52], which reported lower briquette performance compared to the findings of Dewantoro *et al.* [3]. This phenomenon occurs because a high volatile matter leads to a lower fixed carbon

(calculated by difference method), thereby reducing the potential energy released during oxidation—combustion reactions [60]. Further supporting this theory are the results of Kongpraset *et al.* [44], who reported a calorific value of 6,345 cal/g with a fixed carbon of 70.8%, which is lower than the results of Hamzah *et al.* [46], who obtained a calorific value of 8,082 cal/g and a fixed carbon of 86.81%. Both studies used wood residues as biomass mixtures on coconut shell-based briquettes formulation. However, a contrasting finding emerged from the combination of coconut shell with eucalyptus waste—also a wood-based biomass—which exhibited a relatively low fixed carbon only 14.02%, yet a high calorific value of 7,814 cal/g [45]. This anomaly is attributed to the presence of residual volatile oils, which are highly flammable and differ in nature from conventional particulate volatile matter. These oils enhance the energy performance of the developed briquettes despite the lower fixed carbon [61].

3.3 Various Types of Binder Materials for Coconut Shell-Based Briquettes

One of the key aspects in the development of coconut shell-based briquettes highlighted in the visual map network is the type of binder materials. Table 3 lists various types of binders that have been employed in the production of coconut shell-based briquettes, with tapioca flour being the most dominant. This raises a concern, as tapioca serves as a staple food in many countries and its extensive use in briquette production may potentially conflict with food access and security [62–63]. This issue must be addressed promptly, as increasing the proportion of tapioca has been shown to improve the calorific performance of coconut shell-based briquettes. Based on the results of the bibliometric study, several alternative binders have been explored, including carbohydrate derivatives, plant sap, animal manure, and even microorganisms such as fungi. These findings are crucial, as they offer the potential to substitute tapioca with more environmentally friendly and sustainable binders that do not compromise food resources [64–65].

Table 3. The effect of various binder materials toward calorific value in the development of coconut shell-based briquettes.

Briquettes Formulation	Binders	CV ^a	Ref.
Rice husks and coconut shell (1:1)	6% Tapioca flour	4,966 kal/g	[49]
Rice husks and coconut shell (1:1)	8% Tapioca flour	4,886 kal/g	[50]
Sawdust and coconut shell (3:1)	6% Tapioca flour	7,561 kal/g	[12]
Sawdust and coconut shell (3:1)	10% Tapioca flour	7,054 kal/g	[51]
Coconut shell	25% Cassava peel	6,266 kal/g	[66]
Eucalyptus wastes and coconut shell (1:4)	0,5% carboxymethyl cellulose (CMC)	7,814 kal/g	[45]
Coconut shell	7% EFB-based hydrogel with citric acid addition	7,878 kal,g	[3]
Rambutan peels and coconut shell (1:9)	20% Molasses	6,297 kal/g	[56]
Rice husks and coconut shell (1:2)	4% Starch and 4% Mango sap	6,257 kal/g	[32]
Rice husks and coconut shell (n.d.)	Pine resin	9,352 kal/g	[42]
Sugar palm dregs fiber, cassava dregs, and coconut shell (7:13:4)	<i>Ganoderma lucidum</i> (filamentous fungi)	4,522 kal/g	[67]

^a CV = calorific value of developed briquettes.

n.d. means the parameter was not determined on the relevant studies.

In addition to the type of binder, the binder-to-briquette ratio also plays a critical role in determining the performance of coconut shell-based briquettes. As shown in Table 3, carbohydrate-based binders and their derivatives—such as tapioca flour and carboxymethyl cellulose (CMC)—are among the most commonly used materials in coconut shell-based briquette development. Tapioca flour, when applied at a ratio of 6–10%, has been shown to produce a calorific value as 4,800 to 7,500 cal/g [12, 49–51]. Meanwhile, CMC, a modified carbohydrate compound, resulted in a higher calorific value of 7,814 cal/g, even with a much lower addition ratio of just 0.5% [45]. This indicates that carbohydrate modification can significantly reduce the required binder ratio

while maintaining excellent briquette performance. Another noteworthy finding is the potential of binder blending, such as combination of starch with mango sap, which can reduce starch dependency [32]. Furthermore, the use of biomass-derived binders presents a renewable and circular approach [68]. Several types of biomass have been applied as binders in coconut shell-based briquette production, including cassava peel (25% addition yielding 6,266 cal/g) [66], modified empty fruit bunches (modified EFB, 7% addition yielding 7,870 cal/g) [3], and molasses (20% addition yielding 6,297 cal/g) [56]. These results emphasize that even at relatively high ratios, biomass-based binders can enhance the value of coconut shell-based briquettes, supporting the implementation of sustainable circular economy practices.

In addition to carbohydrate-based binders and their derivatives, several alternative binders have been developed over the past decade, including plant sap and filamentous fungi. As previously discussed, plant sap can be used as a complementary binder to carbohydrate compounds, effectively reducing the required ratio while still maintaining good briquette performance [32]. Furthermore, plant sap has also been applied as a standalone binder, such as pine resin, which is known to contain volatile oils—similar to the case of combining coconut shell-based briquettes with eucalyptus waste. Coconut shell-based briquettes bound with pine resin have demonstrated high calorific performance, reaching up to 9,352 cal/g [42]. This phenomenon highlights the potential application of binders containing volatile and combustible oils as a strategy to further improve briquette energy performance. However, the use of such alternative binders requires more comprehensive investigation, especially regarding techno-economic aspects, including production cost and resource availability. A recent finding from the bibliometric study also reveals the emerging use of animal manure as a binder, which presents a promising renewable alternative for the future development of coconut shell briquettes [40, 43].

3.4 Future Novel Development

The combination of coconut shell-based briquette matrix with other biomass types and the use of alternative binders beyond tapioca flour presents a promising direction for future briquette development, as shown on Fig. 5. The utilization of wood-based biomass

(or residues derived from agroforestry activities) has been shown to enhance both the physical characteristics and performance of briquettes. However, fruit peels have also demonstrated encouraging potential, as evidenced by the use of durian and rambutan peels [55–56]. These findings offer opportunities to increase the value chain of agro-industrial residues and serve as tangible examples of circular economy practices. Moreover, the substitution of food-based binders, such as tapioca flour, is urgently needed to address emerging concerns related to food security [3]. One of the most promising binder alternatives identified in the literature is plant sap, which possesses high potential due to its natural adhesive properties and renewable origin as a plant exudate. This potential has been validated through the use of mango sap and pine sap, which have demonstrated calorific values of 6,257 cal/g and 9,352 cal/g, respectively, when used as binders in coconut shell briquette formulations [32, 42].

An interesting finding that emerged from this study is the utilization of oil-containing biomass in coconut shell-based briquette development. The presence of oil in briquettes can enhance both performance and energy potential, as demonstrated by the use of eucalyptus waste as a biomass mixture and pine resin as a binder [42, 45]. These materials are known to contain essential oils, which are volatile and highly flammable, thereby contributing to the increased calorific value of the briquettes [61]. This phenomenon suggests the potential application of WCO, a by-product of the food agroindustry that is increasing in volume annually [33]. Although the use of WCO in coconut shell-based briquette development has been widely studied since 2020, its actual application remains limited. In addition, recent advances in production technologies—particularly the implementation of integrated systems with microcontroller-based automation and Internet of Things (IoT) capabilities—present new avenues for further innovation in the sustainable development of coconut shell-based briquettes [17–19].

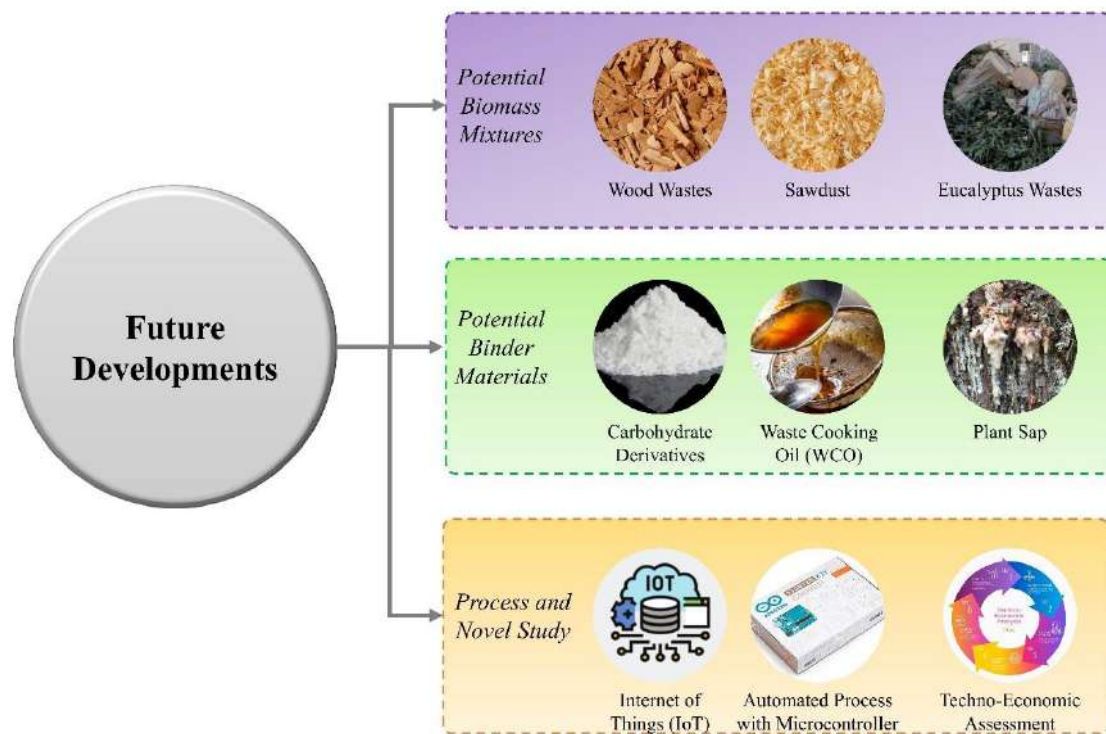


Figure 5. Future potential developments outlined from this study.

4 Conclusions

The development of biomass briquettes based on coconut shell has shown significant progress in terms of raw materials, energy characteristics, and production efficiency. In this context, the combination of various types of biomass and the exploration of alternative binders have emerged as promising aspects for future innovation. Biomass mixtures such as sawdust, straw, and agro-industrial residues have been proven to enhance the calorific value and combustion performance of briquettes. Alternative binders—including plant sap, WCO, animal manure, carbohydrate derivatives from lignocellulosic biomass, and even microorganisms such as filamentous fungi—offer substantial potential to reduce reliance on food-based materials. Furthermore, the utilization of biomass containing oil compounds—such as eucalyptus waste and pine resin—has shown positive effects on briquette performance due to their flammable volatile content. Innovations in automated production systems using microcontrollers and Internet of Things (IoT) technologies also signal a new direction

toward greater efficiency and process standardization. Overall, these findings provide a strong foundation for the advancement of coconut shell briquettes that are not only technically superior but also aligned with circular economy principles and long-term sustainability.

References

- [1] J. Scheffran, M. Felkers, and R. Froese, "Economic Growth and the Global Energy Demand," in *Green Energy to Sustainability: Strategies for Global Industries*, A. A. Vertès, N. Qureshi, H. P. Blaschek, and H. Yukawa, Eds., John Wiley and Sons Ltd, 2020, ch. 1, pp. 1–44. doi: 10.1002/9781119152057.ch1.
- [2] M. Antar, D. Lyu, M. Nazari, A. Shah, X. Zhou, and D. L. Smith, "Biomass for a sustainable bioeconomy: An overview of world biomass production and utilization," *Renew. Sustain. Energy Rev.*, vol. 139, no. December 2020, p. 110691, 2021, doi: 10.1016/j.rser.2020.110691.
- [3] A. I. Dewantoro, M. Fauzan, M. A. R. Lubis, D. Nurliasari, and E. Mardawati, "Carboxymethyl holocellulose as alternative carbohydrate-based binder for biomass briquette development," *Adv. Food Sci. Sustain. Agric. Agroindustrial Eng.*, vol. 7, no. 4, pp. 292–301, 2024, doi: 10.21776/10.21776/ub.afssaae.2024.007.04.2.
- [4] A. A. Adeleke *et al.*, "A Review on Biomass Briquettes as Alternative and Renewable Fuels," in *2023 2nd International Conference on Multidisciplinary Engineering and Applied Science, ICMEAS 2023*, 2023. doi: 10.1109/ICMEAS58693.2023.10429785.
- [5] N. T. S. Saptadi, A. Suyuti, A. A. Ilham, and I. Nurtanio, "Modeling of Organic Waste Classification as Raw Materials for Briquettes using Machine Learning Approach," *Int. J. Adv. Comput. Sci. Appl.*, vol. 14, no. 3, pp. 577–585, 2023, doi:

10.14569/IJACSA.2023.0140367.

- [6] I. S. Utami *et al.*, “How to Make Biomass Briquettes with Their Characteristics Analysis,” *Indones. J. Sci. Technol.*, vol. 9, no. 3, pp. 585–610, 2024, doi: 10.17509/ijost.v9i3.72170.
- [7] G. V. Prasetyadi and J. P. G. Sutapa, “Utilizing Merbau Wood and Coconut Shell Wastes as Biofuel in the Form of Pellets,” *J. Ecol. Eng.*, vol. 24, no. 1, pp. 172–178, 2023, doi: 10.12911/22998993/156057.
- [8] M. M. Ashfaq, G. Bilgic Tüzemen, and A. Noor, “Exploiting agricultural biomass via thermochemical processes for sustainable hydrogen and bioenergy: A critical review,” *Int. J. Hydrogen Energy*, vol. 84, no. July, pp. 1068–1084, 2024, doi: 10.1016/j.ijhydene.2024.08.295.
- [9] F. Vieira *et al.*, “Coconut Waste: Discovering Sustainable Approaches to Advance a Circular Economy,” *Sustain.*, vol. 16, no. 7, 2024, doi: 10.3390/su16073066.
- [10] S. U. Yunusa, E. Mensah, K. Preko, S. Narra, A. Saleh, and S. Sanfo, “A comprehensive review on the technical aspects of biomass briquetting,” *Biomass Convers. Biorefinery*, vol. 14, no. 18, pp. 21619–21644, 2024, doi: 10.1007/s13399-023-04387-3.
- [11] T. Sagdinakiadtikul and N. Supakata, “The application of using rice straw coconut shell and rice husk for briquette and charcoal production,” *Int. J. Energy, Environ. Econ.*, vol. 24, no. 2–3, pp. 283–292, 2016.
- [12] R. P. Dewi and M. Kholik, “The effect of adhesive concentration variation on the characteristics of briquettes,” in *Journal of Physics: Conference Series*, 2020. doi: 10.1088/1742-6596/1517/1/012007.

- [13] R. A. M. Napitupulu, S. Ginting, W. Naibaho, S. Sihombing, N. Tarigan, and A. Kabutey, "The effect of used lubricating oil volume as a binder on the characteristics of briquettes made from corn cob and coconut shell.," in *IOP Conference Series: Materials Science and Engineering*, 2020. doi: 10.1088/1757-899X/725/1/012010.
- [14] M. Amrullah, E. Mardawati, R. Kastaman, and S. Suryaningsih, "Study of bio-briquette formulation from mixture palm oil empty fruit bunches and palm oil shells," *IOP Conf. Ser. Earth Environ. Sci.*, vol. 443, no. 1, 2020, doi: 10.1088/1755-1315/443/1/012079.
- [15] P. Hwangdee, C. Jansiri, S. Sudajan, and K. Laloan, "Physical Characteristics and Energy Content of Biomass Charcoal Powder," *Int. J. Renew. Energy Res.*, vol. 11, no. 1, pp. 158–169, 2021.
- [16] B. V Bot, J. G. Tamba, and O. T. Sosso, "Assessment of biomass briquette energy potential from agricultural residues in Cameroon," *Biomass Convers. Biorefinery*, vol. 14, no. 2, pp. 1905–1917, 2024, doi: 10.1007/s13399-022-02388-2.
- [17] S. Anis *et al.*, "Arduino-Based Automatic Cutting Tool for Coconut Shell Charcoal Briquettes," *ARPN J. Eng. Appl. Sci.*, vol. 19, no. 9, pp. 536–540, 2024, doi: 10.59018/052473.
- [18] K. Jaito, S. Panya, S. Sripan, T. Suparos, and A. Pichaicherd, "Semi-automatic Charcoal Mixer and Compression," in *Lecture Notes in Networks and Systems*, 2024, pp. 76–87. doi: 10.1007/978-3-031-59164-8_7.
- [19] A. Prasetyadi, R. Sambada, and P. K. Purwadi, "Alternative method for stopping the coconut shell charcoal briquette drying process," in *E3S Web of Conferences*, 2024. doi: 10.1051/e3sconf/202447501007.

-
- [20] A. Brunerová *et al.*, “Manual wooden low-pressure briquetting press: An alternative technology of waste biomass utilisation in developing countries of Southeast Asia,” *J. Clean. Prod.*, vol. 436, 2024, doi: 10.1016/j.jclepro.2024.140624.
- [21] R. N. Pagani, J. L. Kovaleski, and L. M. Resende, “Methodi Ordinatio: a proposed methodology to select and rank relevant scientific papers encompassing the impact factor, number of citation, and year of publication,” *Scientometrics*, vol. 105, no. 3, pp. 2109–2135, 2015, doi: 10.1007/s11192-015-1744-x.
- [22] A. Arias *et al.*, “Recent developments in bio-based adhesives from renewable natural resources,” *J. Clean. Prod.*, vol. 314, no. September 2021, 2021, doi: 10.1016/j.jclepro.2021.127892.
- [23] A. I. Dewantoro, M. A. R. Lubis, D. Nurliasari, and E. Mardawati, “Emerging technologies on developing high-performance and environmentally friendly carbohydrate-based adhesives for wood bonding,” *Int. J. Adhes. Adhes.*, vol. 134, 2024, doi: 10.1016/j.ijadhadh.2024.103801.
- [24] A. I. Dewantoro, A. R. Alifia, T. Handini, L. Z. Qolbi, D. A. Ihsani, and D. Nurliasari, “Recent Developments in The Influencing Variables of Hydrodistillation for Enhancing Essential Oil Yields in Indonesia: A Brief Review,” *Int. J. Appl. Sci. Smart Technol.*, vol. 06, no. 02, pp. 301–320, 2024, doi: 10.24071/ijasst.v6i2.9191.
- [25] G. Murali, P. Goutham, I. Enamul Hasan, and P. Anbarasan, “Performance study of briquettes from agricultural waste for wood stove with catalytic combustor,” *Int. J. ChemTech Res.*, vol. 8, no. 1, pp. 30–36, 2015.
- [26] E. R. D. Padilla, I. C. S. A. Pires, F. M. Yamaji, and J. M. M. Fandiño, “Production and physical-mechanical characterization of briquettes from coconut fiber and

-
- sugarcane straw,” *Rev. Virtual Quim.*, vol. 8, no. 5, pp. 1334–1346, 2016, doi: 10.21577/1984-6835.20160095.
- [27] H. Saptoadi and M. A. Wibisono, “Optimization of temperature and time for drying and carbonization to increase calorific value of coconut shell using Taguchi method,” in *AIP Conference Proceedings*, 2016. doi: 10.1063/1.4943430.
- [28] A. Brunerová, H. Roubík, M. Brožek, D. Herák, V. Šleger, and J. Mazancová, “Potential of tropical fruit waste biomass for production of bio-briquette fuel: Using Indonesia as an example,” *Energies*, vol. 10, no. 12, 2017, doi: 10.3390/en10122119.
- [29] S. Suryaningsih, O. Nurhilal, Y. Yuliah, and C. Mulyana, “Combustion quality analysis of briquettes from variety of agricultural waste as source of alternative fuels,” in *IOP Conference Series: Earth and Environmental Science*, 2017. doi: 10.1088/1755-1315/65/1/012012.
- [30] T. Akbari, “Economic and environmental feasibility study of water hyacinth briquette in Cirata Reservoir,” in *E3S Web of Conferences*, 2018. doi: 10.1051/e3sconf/20187401001.
- [31] T. Syarif, “Biobriquette Characteristics of Mixture of Coal-Biomass Solid Waste Agro,” in *IOP Conference Series: Earth and Environmental Science*, 2018. doi: 10.1088/1755-1315/175/1/012031.
- [32] S. Rodiah, J. L. Al Jabbar, A. Ramadhan, and E. Hastati, “Investigation of mango (*Mangifera odorata*) sap and starch as organic adhesive of bio-briquette,” in *Journal of Physics: Conference Series*, 2021. doi: 10.1088/1742-6596/1943/1/012185.
-

-
- [33] A. Tumma, P. Dujjanutat, P. Muanruksa, J. Winterburn, and P. Kaewkannetra, "Biocircular platform for renewable energy production: Valorization of waste cooking oil mixed with agricultural wastes into biosolid fuels," *Energy Convers. Manag. X*, vol. 15, 2022, doi: 10.1016/j.ecmx.2022.100235.
- [34] B. V Bot, P. J. Axaopoulos, E. I. Sakellariou, O. T. Sosso, and J. G. Tamba, "Economic Viability Investigation of Mixed-Biomass Briquettes Made from Agricultural Residues for Household Cooking Use," *Energies*, vol. 16, no. 18, 2023, doi: 10.3390/en16186469.
- [35] B. V Bot, P. J. Axaopoulos, O. T. Sosso, E. I. Sakellariou, and J. G. Tamba, "Economic analysis of biomass briquettes made from coconut shells, rattan waste, banana peels and sugarcane bagasse in households cooking," *Int. J. Energy Environ. Eng.*, vol. 14, no. 2, pp. 179–187, 2023, doi: 10.1007/s40095-022-00508-2.
- [36] N. S. Kamal Baharin *et al.*, "Impact and effectiveness of Bio-Coke conversion from biomass waste as alternative source of coal coke in Southeast Asia," *J. Mater. Cycles Waste Manag.*, vol. 25, no. 1, pp. 17–36, 2023, doi: 10.1007/s10163-022-01539-x.
- [37] E. Hambali and H. Hardjomidjojo, "Investigation of Conceptual Supply Chain Design Process of Sustainable Shisha Briquette Production System," in *IOP Conference Series: Earth and Environmental Science*, 2024. doi: 10.1088/1755-1315/1358/1/012007.
- [38] K. Khaeso *et al.*, "Sustainable conversion of agricultural waste into solid fuel (Charcoal) via gasification and pyrolysis treatment," *Energy Convers. Manag. X*, vol. 24, 2024, doi: 10.1016/j.ecmx.2024.100693.

- [39] S. Ahmad, K. Winarso, R. Yusron, and S. Amar, "Optimization of Calorific Value in Briquette made of Coconut Shell and Cassava Peel by varying of Mass Fraction and Drying Temperature," in *E3S Web of Conferences*, 2024. doi: 10.1051/e3sconf/202449901009.
- [40] A. Ashwini, R. Arulvel, M. Shakila, W. S. Fan, and G. H. Kaur, "Characterization of coconut husk briquettes using cow dung as a binder with comparison to *Elaeis guineensis*," in *AIP Conference Proceedings*, 2024. doi: 10.1063/5.0229369.
- [41] S. Anis *et al.*, "Effect of Adhesive Type on the Quality of Coconut Shell Charcoal Briquettes Prepared by the Screw Extruder Machine," *J. Renew. Mater.*, vol. 12, no. 2, pp. 381–396, 2024, doi: 10.32604/jrm.2023.047128.
- [42] A. Setiawan, K. Khairil, S. Nurjannah, N. Nurmali, and Z. Fona, "Pine resin utilization as a binding agent for densification of coconut shells and rice husks at various pressures," in *AIP Conference Proceedings*, 2023. doi: 10.1063/5.0133273.
- [43] V. Sampathkumar, S. Manoj, V. Nandhini, N. J. Lakshmi, and S. Janani, "Briquetting of biomass for low cost fuel using farm waste, cow dung and cotton industrial waste," *Int. J. Recent Technol. Eng.*, vol. 8, no. 3, pp. 8349–8353, 2019, doi: 10.35940/ijrte.C6616.098319.
- [44] N. Kongprasert, P. Wangphanich, and A. Jutilarptavorn, "Charcoal briquettes from Madan wood waste as an alternative energy in Thailand," in *Procedia Manufacturing*, 2019, pp. 128–135. doi: 10.1016/j.promfg.2019.02.019.
- [45] E. Z. Nunes, A. M. de Andrade, and A. F. Dias Júnior, "Production of briquettes using coconut and eucalyptus wastes," *Rev. Bras. Eng. Agric. e Ambient.*, vol. 23, no. 11, pp. 883–888, 2019, doi: 10.1590/1807-1929/agriambi.v23n11p883-888.

- [46] F. Hamzah, A. Fajri, N. Harun, and A. Pramana, "Characterization of charcoal briquettes made from rubber rods and coconut shells with tapioca as an adhesive.," in *IOP Conference Series: Earth and Environmental Science*, 2023. doi: 10.1088/1755-1315/1182/1/012071.
- [47] D. S. Nawawi, A. Carolina, T. Saskia, D. Darmawan, and S. L. Gusvina, "Chemical Characteristics of Biomass for Energy," *J. Ilmu dan Teknol. Kayu Trop.*, vol. 16, no. 1, pp. 44–51, 2018, doi: <https://doi.org/10.51850/jitkt.v16i1.441.g367>.
- [48] D. Watkins, M. Nuruddin, M. Hosur, A. Tcherbi-Narteh, and S. Jeelani, "Extraction and characterization of lignin from different biomass resources," *J. Mater. Res. Technol.*, vol. 4, no. 1, pp. 26–32, 2015, doi: <https://doi.org/10.1016/j.jmrt.2014.10.009>.
- [49] Y. Yuliah, M. Kartawidjaja, S. Suryaningsih, and K. Ulfi, "Fabrication and characterization of rice husk and coconut shell charcoal based bio-briquettes as alternative energy source," in *IOP Conference Series: Earth and Environmental Science*, 2017. doi: 10.1088/1755-1315/65/1/012021.
- [50] S. Suryaningsih and O. Nurhilal, "Sustainable energy development of bio briquettes based on rice husk blended materials: An alternative energy source," in *Journal of Physics: Conference Series*, 2018. doi: 10.1088/1742-6596/1013/1/012184.
- [51] R. P. Dewi, "Utilization of sawdust and coconut shell as raw materials in briquettes production," in *AIP Conference Proceedings*, 2019. doi: 10.1063/1.5098179.
- [52] S. A. Handayani, K. Y. Widiati, and F. Putri, "Quality of Charcoal Briquettes from Sawdust Waste of Ulin (*Eusideroxylon zwageri*) and Coconut Shell (*Cocos nucifera*) Based on Variations in Ratio and Particle Size," in *IOP Conference Series: Earth and Environmental Science*, 2023. doi: 10.1088/1755-

1315/1282/1/012049.

- [53] N. Pawaree, S. Phokha, and C. Phukapak, "Multi-response optimization of charcoal briquettes process for green economy using a novel TOPSIS linear programming and genetic algorithms based on response surface methodology," *Results Eng.*, vol. 22, 2024, doi: 10.1016/j.rineng.2024.102226.
- [54] A. S. Erdiyanto, M. H. Asshidiqi, and G. Syachrir, "Bio-Briquettes Production from Spent Coffee Grounds, Composite Organic Waste, and Coconut Shells by Using Carbonization," in *IOP Conference Series: Earth and Environmental Science*, 2024. doi: 10.1088/1755-1315/1395/1/012010.
- [55] N. Herlina, S. A. Rahayu, Y. A. Sari, and H. Monica, "Utilization of durian peel waste and young coconut waste into biobriquettes as a renewable energy source," in *IOP Conference Series: Earth and Environmental Science*, 2024. doi: 10.1088/1755-1315/1352/1/012014.
- [56] G. S. Utami and E. Ningsih, "The Impact of Composition and Type of Material on the Characteristics of Fuel Briquettes," *Chem. Eng. Trans.*, vol. 108, pp. 1–6, 2024, doi: 10.3303/CET24108001.
- [57] U. B. Deshannavar, P. G. Hedge, Z. Dhalayat, V. Patil, and S. Gavvas, "Production and characterization of agro-based briquettes and estimation of calorific value by regression analysis: An energy application," *Mater. Sci. Energy Technol.*, vol. 1, no. 2, pp. 175–181, 2018, doi: 10.1016/j.mset.2018.07.003.
- [58] R. A. Ibikunle, A. F. Lukman, I. F. Titiladunayo, and A.-R. Haadi, "Modeling energy content of municipal solid waste based on proximate analysis: R-k class estimator approach," *Cogent Eng.*, vol. 9, no. 1, 2022, doi: 10.1080/23311916.2022.2046243.

- [59] J. J. S. Dethan, "Evaluation of an empirical model for predicting the calorific value of biomass briquettes from candlenut shells and kesambi twigs," *Adv. Food Sci. Sustain. Agric. Agroindustrial Eng.*, vol. 7, no. 3, pp. 253–264, 2024, doi: 10.21776/ub.afssaae.2024.007.03.6.
- [60] N. Arifin and R. Noor, "Effect of Composition The Mixture of Charcoal Briquettes Made frm Reeds (*Imperata cylindrica*) to Increase Calory Value," *Jukung J. Tek. Lingkungan.*, vol. 2, no. 2, pp. 61–72, 2016, doi: 10.20527/jukung.v2i2.2315.
- [61] K. Vershinina, V. Dorokhov, D. Romanov, and P. Strizhak, "Ignition, Combustion, and Mechanical Properties of Briquettes from Coal Slime and Oil Waste, Biomass, Peat and Starch," *Waste and Biomass Valorization*, vol. 14, pp. 431–445, 2022, doi: 10.1007/s12649-022-01883-x.
- [62] C. A. Monteiro, G. Cannon, J. C. Moubarac, R. B. Levy, M. L. C. Louzada, and P. C. Jaime, "The un Decade of Nutrition, the NOVA food classification and the trouble with ultra-processing," *Public Health Nutr.*, vol. 21, no. 1, pp. 5–17, 2018, doi: 10.1017/S1368980017000234.
- [63] P. Hemalatha *et al.*, "Multi-faceted CRISPR-Cas9 strategy to reduce plant based food loss and waste for sustainable bio-economy – A review," *J. Environ. Manage.*, vol. 332, no. December 2022, p. 117382, 2023, doi: 10.1016/j.jenvman.2023.117382.
- [64] P. Duarah, D. Haldar, A. K. Patel, C. Di Dong, R. R. Singhanian, and M. K. Purkait, "A review on global perspectives of sustainable development in bioenergy generation," *Bioresour. Technol.*, vol. 348, no. January, p. 126791, 2022, doi: 10.1016/j.biortech.2022.126791.

-
- [65] S. Nonsawang, S. Juntahum, P. Sanchumpu, W. Suaili, K. Senawong, and K. Laloon, “Unlocking renewable fuel: Charcoal briquettes production from agro-industrial waste with cassava industrial binders,” *Energy Reports*, vol. 12, pp. 4966–4982, 2024, doi: 10.1016/j.egy.2024.10.053.
- [66] B. Rudiyanto *et al.*, “Utilization of Cassava Peel (*Manihot utilissima*) Waste as an Adhesive in the Manufacture of Coconut Shell (*Cocos nucifera*) Charcoal Briquettes,” *Int. J. Renew. Energy Dev.*, vol. 12, no. 2, pp. 270–276, 2023, doi: 10.14710/ijred.2023.48432.
- [67] W. Agustina, P. Aditiawati, and S. S. Kusumah, “Myco-briquettes from sugar palm dregs fibre, cassava dregs and coconut shell charcoal with solid substrate fermentation technology,” in *IOP Conference Series: Earth and Environmental Science*, 2022. doi: 10.1088/1755-1315/963/1/012016.
- [68] N. Tripathi, C. D. Hills, R. S. Singh, and C. J. Atkinson, “Biomass waste utilisation in low-carbon products: harnessing a major potential resource,” *NPJ Clim. Atmos. Sci.*, vol. 2, no. 35, 2019, doi: 10.1038/s41612-019-0093-5.

This page intentionally left

Fine-Tuned IndoBERT-Based Sentiment Analysis for Old Indonesian Songs Using Contextual and Generating Augmentation

Gilang Ramdhani¹, Siti Yuliyanti^{1*}

¹*Department of Informatics, Faculty of Engineering, Siliwangi University,
Tasikmalaya, West Java, Indonesia*

**Corresponding Author: sitiyluyanti@unsil.id*

(Received 05-06-2025; Revised 10-07-2025; Accepted 26-07-2025)

Abstract

This study examines sentiment analysis of traditional Indonesian songs using a fine-tuned IndoBERT model, which has been enhanced through the incorporation of contextual and textual data augmentation. The dataset consists of user comments related to classic Indonesian songs, categorized into positive, negative, and neutral sentiments. Two augmentation strategies were applied: textual augmentation using text generation techniques and contextual augmentation leveraging semantic similarity. Evaluation was conducted using accuracy, precision, recall, and F1-score metrics. Results show that the model trained on the original dataset achieved balanced and stable performance (accuracy: 0.86). Textual augmentation, despite generating high data variation, reduced model accuracy (0.63) and introduced a bias toward negative sentiment. In contrast, contextual augmentation-maintained performance stability and even slightly improved precision (0.87). These findings indicate that contextual augmentation is more effective for enriching sentiment datasets without compromising model performance. The findings highlight the effectiveness of integrating pre-trained language models and data augmentation strategies for sentiment analysis in low-resource.

Keywords: IndoBERT, Sentiment Analysis, Old Indonesian Songs, Contextual Augmentation, Textual Augmentation.

1 Introduction

Sentiment analysis is a widely applied natural language processing (NLP) technique that seeks to identify and classify emotional tones expressed in text[1], [2]. In recent years, the increasing availability of user-generated content, such as comments on music platforms and social media, has opened new opportunities to analyze public perception toward songs, including those from past decades[3], [4]. Old Indonesian songs,



such as "Tuhan yang Aneh – Apa Elo Tega", often carry deep emotional narratives and have regained popularity in online spaces. However, the informal, poetic, and sometimes metaphorical nature of song-related commentary poses challenges for traditional sentiment analysis methods. The criteria for what we consider "old Indonesian songs" were implicitly based on temporal and stylistic attributes, but we acknowledge that this was not explicitly stated in the background section. In our study, *old Indonesian songs* are defined as songs that were released before the 1990s and are recognized as part of the classic or nostalgic repertoire of Indonesian music, typically characterized by lyrical richness, poetic language, and traditional musical arrangements. These songs are often still referenced in modern media and evoke a certain cultural sentiment tied to past eras. In future revisions, we will clarify this definition in the introduction to provide better contextual grounding.

Pretrained language models like IndoBERT, developed specifically for the Indonesian language, have shown great promise in capturing the linguistic nuances required for high-quality text classification[5], [6]. Yet, the performance of such models depends heavily on the diversity and richness of training data. To address this, data augmentation techniques can be employed to enhance model generalization [7], [8]. This study explores two augmentation strategies—textual augmentation using synonym and paraphrase generation, and contextual augmentation leveraging semantic embeddings—to fine-tune IndoBERT for sentiment analysis tasks related to "Tuhan yang Aneh – Apa Elo Tega".

Despite the growing interest in analyzing public sentiment toward music, especially emotionally intense or controversial songs like "*Tuhan yang Aneh – Apa Elo Tega*", sentiment analysis in the Indonesian language remains underexplored. Most existing models are trained on general-purpose datasets and often fail to capture the informal, figurative, or context-dependent expressions commonly found in song-related commentary. These limitations reduce model accuracy and can lead to biased or misleading sentiment predictions.

This is an important point, and we appreciate the critique. The reason for focusing on a single song in the experimental stage was to create a controlled case study to assess the effectiveness of IndoBERT when fine-tuned with contextual and generative data

augmentation methods. While the title implies a broader scope, we acknowledge that the data was limited to one song due to constraints in labeled sentiment datasets for old Indonesian lyrics. We do not claim that one song can fully represent all old Indonesian songs, but rather that this study serves as a proof of concept. Future work will expand to include a larger and more diverse corpus of old Indonesian songs to validate the generalizability of the findings. We will revise the title and discussion to reflect this scope limitation better.

Furthermore, high-quality annotated datasets for Indonesian sentiment analysis, particularly in the context of music and culture, are limited in both size and variety[9]. This scarcity of training data can limit the ability of even powerful models, such as IndoBERT, to generalize across diverse expressions and sentiment patterns. While data augmentation techniques offer a potential solution, not all augmentation methods are equally effective[10], [11]. Textual augmentation using random synonym replacement or text generation may introduce noise or distort sentiment polarity, while contextual augmentation strategies remain underutilized in Indonesian-language NLP tasks [12], [13].

The goal of this research is to evaluate the effectiveness of contextual and textual data augmentation in improving sentiment classification performance, while also understanding how public sentiment toward emotionally charged old songs is represented online. The findings are expected to contribute to better sentiment analysis practices in the domain of cultural and musical content, particularly in low-resource language settings like Bahasa Indonesia.

2 Material and Methods

The stages of this research include data collection, pre-processing, data sharing, IndoBERT Modeling, Fine-Tuning using Contextual and Textual Augmentation, and model evaluation as shown in Fig. 1.

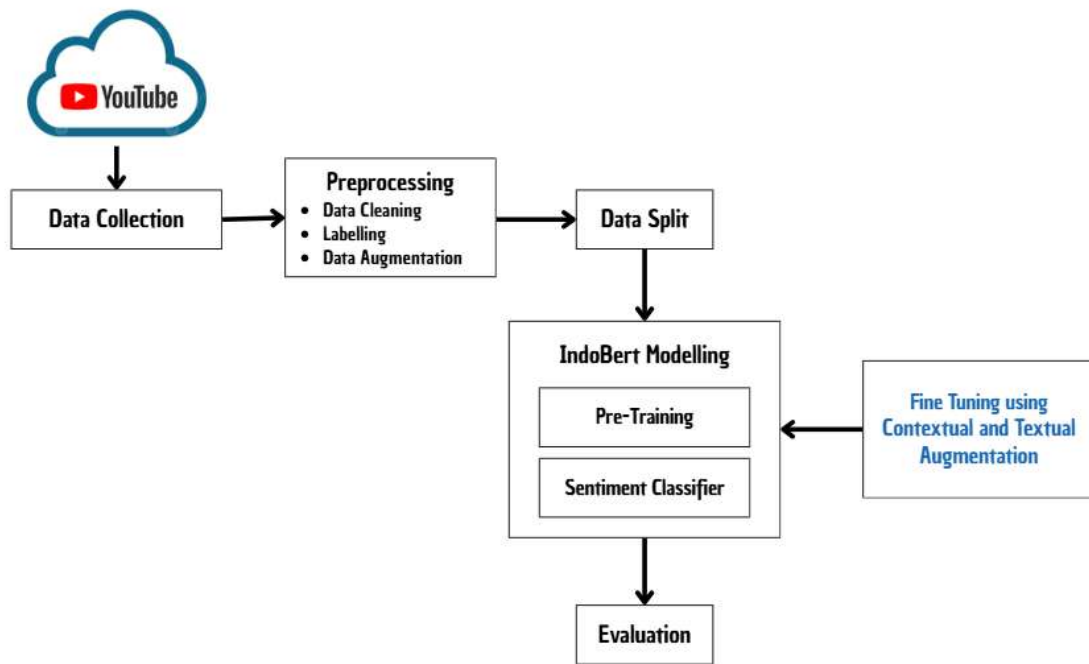


Figure 1. Research Framework

Data Collection

Comment data is gathered from the YouTube video for the song "Tuhan yang Aneh – Apa Elo Tega." This data consists of public comments extracted via web scraping and saved in CSV format. These comments represent users' genuine thoughts and are divided into three categories: positive, negative, and neutral [8], [14], [15]. Alongside the original dataset, this project incorporates additional data to enhance and balance the collection. The comments considered for this project are sourced from the YouTube video titled "Tuhan Yang Aneh - Apa Elo Tega (Old Song 2012)," which was re-uploaded by the Pepstun YouTube channel. The video has a total of 13,509 comments, of which 4,683 were successfully gathered, excluding replies to individual comments.

Preprocessing

After the data scraping process, data cleaning is carried out. These steps include: Removing links, usernames starting with the symbol "@" (@username) or hashtags, and numeric characters[16], [17]. In addition, symbols and unnecessary punctuation are also

removed to produce cleaner data that is ready to be labeled. Next comes the stage of labeling, in which remarks are sorted by hand or partially automatically into sentiment groups (positive, negative, neutral) [3], [18], [19]. Following that is the phase of data enhancement, which involves two forms of augmentation: textual augmentation that includes synonyms or paraphrases, and contextual augmentation utilizing embedding methods to preserve the underlying meaning.

Data Split

The data is divided into 70% training data to train the model, and 30% test data to measure the performance of the model [20], [21], [22], and the model is trained with 3 epochs.

IndoBERT using Contextual and Textual Augmentation

The model used is indobenchmark/indobert-large-p1, which is a large version of Indobert provided by IndoNLU and accessed through the Hugging Face Transformers Library[11]. Model training is carried out using a fine-tuning approach, namely retraining the pre-trained IndoBERT model on the YouTube comment dataset that has gone through the preprocessing and labeling process. The fine-tuning process is carried out separately for three data scenarios, namely:

1. Original data (without augmentation),
2. Data with Text Generating augmentation,
3. Data with Contextual Augmentation augmentation.

Each scenario uses 70% of the total dataset for training data, and the model is trained with 3 epochs. This training uses the Google Colab Platform which uses the Tesla T4 GPU.

3 Results and Discussions

IndoBERT's effectiveness was evaluated across three different dataset scenarios: the original data (without any modifications), data enhanced with text generation, and data that has been supplemented with contextual augmentation. The evaluation criteria include accuracy, precision, recall, and f1-score, with a confusion matrix providing a

breakdown of predictions for each category. The model excels in identifying the positive category (TP: 2039), with only 102 and 24 instances misidentified as neutral and negative, respectively. In terms of the neutral category, the model still performs well (TP: 1603), although it misclassifies 367 instances as positive, suggesting an inclination to overestimate positive sentiments. Regarding the negative category, the model demonstrates strong performance with TP: 382, and only a small number are incorrectly categorized as positive or neutral.

The model performs best in identifying **positive** sentiment, with high accuracy and relatively few misclassifications. It is also fairly accurate with **neutral** sentiment, though it often confuses it with positive, as shown in Fig. 2. The **negative** class has the lowest count and highest misclassification rate, suggesting the model struggles more with identifying negative sentiment correctly. The color intensity in the heatmap helps visually represent the number of occurrences, with darker shades indicating higher values.

The model shows stable and balanced performance without augmentation. The results are evenly distributed across the three classes, with a slight tendency towards the positive class, as shown in Fig. 3.

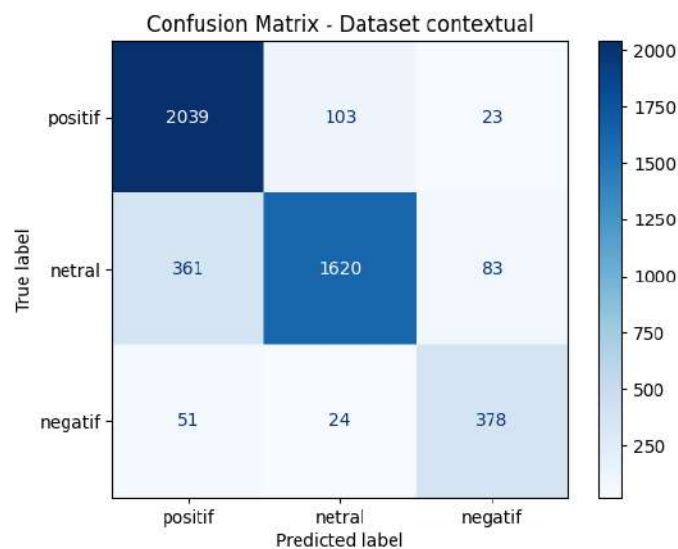


Figure 2. Confusion matrix contextual dataset

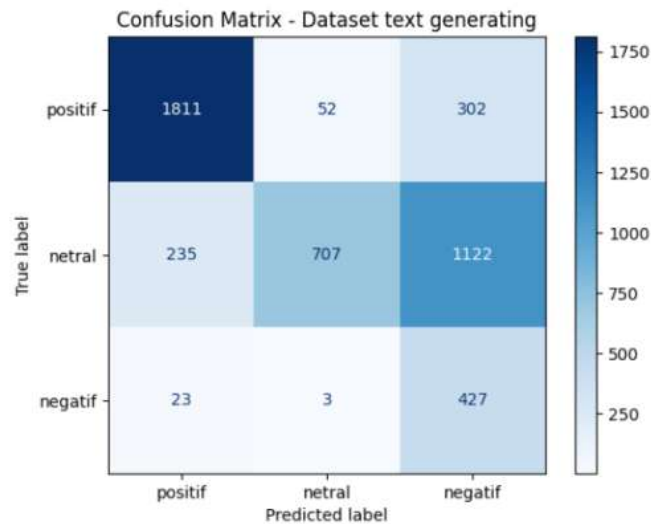


Figure 3. Confusion matrix text generating

Fig. 3 shows that the positive class is still classified dominantly (TP: 1811), but there is a spike in errors to negative by 302 comments that should be positive. The neutral class is significantly confused, with only 707 correct (TP), while 1122 are misclassified as negative. This shows the weakness of the model in recognizing neutrality when the data comes from text-generating augmentation. The negative class remains strong with TP: 427, showing the model's resilience to negative comments even after augmentation. Although the precision and F1-score values are high, the accuracy drops drastically. This indicates that text-generating augmentation increases the variation of the data but makes it difficult for the model to accurately distinguish neutral and negative comments.

The prediction distribution shown in Fig. 4 closely resembles the initial dataset. The model effectively identifies positive remarks, achieving 2039 true positives, while only misclassifying 103 as neutral and 23 as negative. Regarding neutral remarks, there is an enhancement in performance when compared to text generation. The misclassification numbers are lower, with 361 mistakenly labeled as positive and 83 as negative, resulting in 1620 true positives. The negative category is also fairly well-balanced, with 378 true positives, suggesting that contextual enrichment does not negatively impact the model's efficiency.

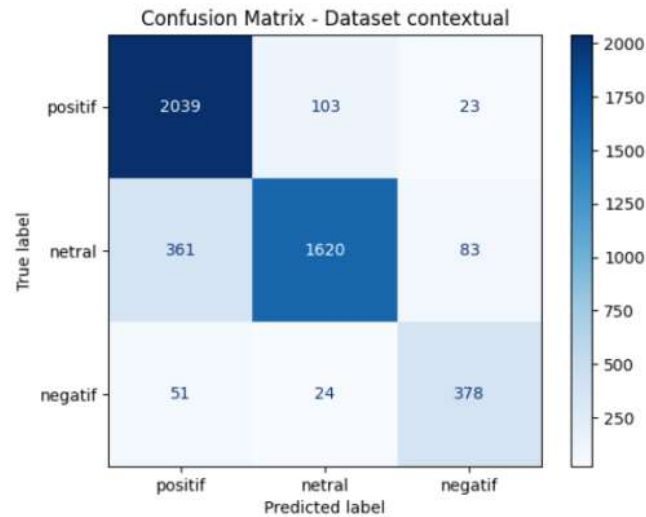


Figure 4. Confusion matrix text contextual

Assessment of model effectiveness following fine-tuning is presented in Table 1, which utilizes accuracy, precision, recall, and F1-score metrics, along with an evaluation of the confusion matrix across the three dataset scenarios. It can be inferred that the dataset enhanced with contextual augmentation yields the highest and most even performance.

Although the original dataset (without any enhancements) demonstrated strong results (both accuracy and F1-score of 0.86), this method faces constraints due to the limited quantity and diversity of the initial data. Conversely, contextual augmentation effectively enhanced the dataset without compromising the stability of the model, yielding evaluation outcomes that matched those of the original dataset (accuracy: 0.86, F1-score: 0.86), while maintaining a balanced classification

Table 1. Performance comparison table with augmentation

Dataset	Accuray	Precission	Recal	F1-Score	Description
Initial Dataset	0.86	0.86	0.86	0.86	Stable, no augmentation, balanced prediction
Text Generating	0.63	0.84	0.83	0.85	High variation but lower accuracy, biased to negative
Text Contextual	0.86	0.87	0.86	0.86	Variation maintained, performance remains stable

across sentiment categories.

In contrast, augmentation through text generation achieved a relatively high F1-score (0.85), but accuracy dropped significantly to 0.63. This suggests that the model struggles to differentiate between neutral and negative feedback, a challenge that is evident in the confusion matrix. This decline might stem from the generated text being less reflective of the original comments' structure and context.

Therefore, employing contextual augmentation techniques appears to be the most effective method for this project, as it enhances data variety without diminishing the quality of the model's classifications. This strategy is ideal for expanding the dataset while preserving both the accuracy and the generalization capabilities of the IndoBERT model.

4 Conclusions

This study proposed a sentiment analysis approach for old Indonesian songs using a fine-tuned IndoBERT model, incorporating both contextual and generative data augmentation techniques. Based on the evaluation results, the model trained with generative augmentation outperformed the one using contextual augmentation, achieving an accuracy of 87%, compared to the contextual model's performance. This indicates that generative augmentation provides more diverse and effective training data, which enhances the model's ability to generalize and classify sentiments more accurately in this specific domain. For future work, the following recommendations are proposed: incorporate emotion classification, further improve accuracy, develop a web-based interactive dashboard, and apply to other social media platforms.

Acknowledgements

The author would like to express his deepest gratitude to all parties who have assisted in this research.

References

- [1] S. Yuliyanti and Rizky, “Implementasi Algoritma Rabin Karp Untuk Mendeteksi Kemiripan Dokumen Stmik Bandung,” *J. Bangkit Indones.*, vol. 10, no. 02, p. 1, 2020, doi: 10.52771/bangkitindonesia.v10i02.124.
- [2] S. Yuliyanti, E. Nur Fitriani Dewi, and A. Nur Rachman, “Optimasi Rabin Karp dengan Rolling Hash dan k-Gram pada Similarity Check Dokumen Abstrak Jurnal,” vol. 12, no. 1, 2024, doi: 10.26418/justin.v12i1.71224.
- [3] H. Jayadianti, W. Kaswidjanti, A. T. Utomo, S. Saifullah, F. A. Dwiyanto, and R. Drezewski, “Sentiment analysis of Indonesian reviews using fine-tuning IndoBERT and R-CNN,” *Ilk. J. Ilm.*, vol. 14, no. 3, pp. 348–354, 2022, doi: 10.33096/ilkom.v14i3.1505.348-354.
- [4] G. Hakim, T. N. Fatyanosa, and A. W. Widodo, “Analisis Sentimen Masyarakat terhadap Kereta Cepat Whoosh pada Platform X menggunakan IndoBERT,” vol. 1, no. 1, pp. 1–10, 2023.
- [5] K. Ayu Pradani, L. Hulliyyatus Suadaa, and P. Korespondensi, “Automated Essay Scoring Menggunakan Semantic Textual Similarity Berbasis Transformer Untuk Penilaian Ujian Esai Automated Essay Scoring Using Transformer-Based Semantic Textual Similarity for Essay Assessment,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 10, no. 6, pp. 1177–1184, 2019, doi: 10.25126/jtiik.2023107338.
- [6] T. Wahyuningsih *et al.*, “Text Mining an Automatic Short Answer Grading (ASAG), Comparison of Three Methods of Cosine Similarity, Jaccard Similarity and Dice’s Coefficient,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 1, no. 2, pp. 343–348, 2021, doi: 10.17762/turcomat.v12i3.938.

-
- [7] F. Koto, A. Rahimi, J. H. Lau, and T. Baldwin, “IndoLEM and IndoBERT: A Benchmark Dataset and Pre-trained Language Model for Indonesian NLP,” *COLING 2020 - 28th Int. Conf. Comput. Linguist. Proc. Conf.*, pp. 757–770, 2020, doi: 10.18653/v1/2020.coling-main.66.
 - [8] B. V. Kartika, M. J. Alfredo, and G. P. Kusuma, “Fine-Tuned IndoBERT based model and data augmentation for indonesian language paraphrase identification,” *Rev. d’Intelligence Artif.*, vol. 37, no. 3, pp. 733–743, 2023, doi: 10.18280/ria.370322.
 - [9] M. Polignano, P. Basile, M. de Gemmis, G. Semeraro, and V. Basile, “AlBERTo: Italian BERT language understanding model for NLP challenging tasks based on tweets,” *CEUR Workshop Proc.*, vol. 2481, 2019.
 - [10] R. Silva Barbon and A. T. Akabane, “Towards Transfer Learning Techniques—BERT, DistilBERT, BERTimbau, and DistilBERTimbau for Automatic Text Classification from Different Languages: A Case Study,” *Sensors*, vol. 22, no. 21, 2022, doi: 10.3390/s22218184.
 - [11] M. A. Jahin, M. S. H. Shovon, M. F. Mridha, M. R. Islam, and Y. Watanobe, “A hybrid transformer and attention based recurrent neural network for robust and interpretable sentiment analysis of tweets,” *Sci. Rep.*, vol. 14, no. 1, p. 24882, 2024, doi: 10.1038/s41598-024-76079-5.
 - [12] T. Bey Kusuma, I. Komang, and A. Mogi, “Implementasi BERT pada Analisis Sentimen Ulasan Destinasi Wisata Bali,” *J. Elektron. Ilmu Komput. Udayana*, vol. 12, no. 2, pp. 409–420, 2023.
 - [13] Y. A. Singgalen, “Performance Analysis of IndoBERT for Sentiment Classification in Indonesian Hotel Review Data,” vol. 6, no. 2, pp. 976–986, 2025, doi: 10.47065/josh.v6i2.6505.

- [14] I. A. Oktariansyah, F. R. Umbara, and F. Kasyidi, "Klasifikasi Sentimen Untuk Mengetahui Kecenderungan Politik Pengguna X Pada Calon Presiden Indonesia 2024 Menggunakan Metode IndoBert," *Build. Informatics, Technol. Sci.*, vol. 6, no. 2, pp. 636–648, 2024, [Online]. Available: <https://ejurnal.seminar-id.com/index.php/bits/article/view/5435>
- [15] C. Sun, X. Qiu, Y. Xu, and X. Huang, "How to Fine-Tune BERT for Text Classification?," *Lect. Notes Comput. Sci. (including Subser. Lect. Notes Artif. Intell. Lect. Notes Bioinformatics)*, vol. 11856 LNAI, no. 2, pp. 194–206, 2019, doi: 10.1007/978-3-030-32381-3_16.
- [16] S. Redhu, "Sentiment Analysis Using Text Mining: A Review," *Int. J. Data Sci. Technol.*, vol. 4, no. 2, 2018, doi: 10.11648/j.ijdst.20180402.12.
- [17] S. Yuliyanti, T. Djatna, and H. Sukoco, "Sentiment mining of community development program evaluation based on social media," *Telkomnika (Telecommunication Comput. Electron. Control.)*, vol. 15, no. 4, pp. 1858–1864, 2017, doi: 10.12928/TELKOMNIKA.v15i4.4633.
- [18] M. G. Villar, J. B. Ballester, I. De, T. Diez, and I. Ashraf, "Analyzing Sentiments Regarding ChatGPT Using Novel BERT : A Machine Learning Approach," pp. 1–29, 2023.
- [19] D. H. Yuliyanti, Siti;Ula, "Essay Answer Detection System Uses Cosine Similarity and Similarity Scoring in Sentences," vol. 06, no. 02, pp. 337–348, 2024.
- [20] C. A. Bahri and L. H. Suadaa, "Aspect-Based Sentiment Analysis in Bromo Tengger Semeru National Park Indonesia Based on Google Maps User Reviews," *IJCCS (Indonesian J. Comput. Cybern. Syst.)*, vol. 17, no. 1, p. 79, 2023, doi: 10.22146/ijccs.77354.

- [21] M. Chiny, M. Chihab, and Y. Chihab, “LSTM , VADER and TF-IDF based Hybrid Sentiment Analysis Model,” vol. 12, no. 7, pp. 265–275, 2021.

- [22] J. Asian, M. Dholah Rosita, and T. Mantoro, “Sentiment Analysis for the Brazilian Anesthesiologist Using Multi-Layer Perceptron Classifier and Random Forest Methods,” *J. Online Inform.*, vol. 7, no. 1, p. 132, 2022, doi: 10.15575/join.v7i1.900.

This page intentionally left

Design of a Speed Sensorless Control System on a DC Motor using a PID Controller

Bernadeta Wuri Harini^{1,*}, Theodore Galeno Gunadi¹, Shakuntala
Gema Mahardika¹, Sirilus Praditya Pangestu¹, Regina Chelinia
Erianda Putri¹, Stefan Mardikus¹

¹*Faculty of Science and Technology, University of Sanata Dharma,
Paingan, Maguwoharjo, Depok, Sleman, Yogyakarta, Indonesia*

**Corresponding Author: wuribernard@usd.ac.id*

(Received 29-07-2025; Revised 23-08-2025; Accepted 25-08-2025)

Abstract

The speed sensorless control system on a DC motor is a DC motor speed control system that does not use a speed sensor to measure the motor speed. The motor speed value is estimated by an observer from the stator current and voltage that are measured using a sensor. This study uses an R observer method. The difference between the estimated speed and the reference speed is then used by the PID controller to adjust the motor speed to match the desired reference speed. PID parameter tuning using heuristic method. With a setpoint of 6800 RPM and using a combination of $K_p = 0.5$, $K_i = 0.05$, and $K_d = 0.34$, a speed value of 6792.76 RPM was obtained. Sensorless motor speed control using a PID controller produces an optimal system with a low Steady State Error (SSE) value of around 0.1%, very small oscillations of 0.39%, a fast rise time of 4 seconds, and a fast-settling time of 6 seconds.

Keywords: DC motor, Heuristic tuning, Observer, PID controller, speed sensorless

1 Introduction

Speed control of a DC motor is an important aspect in various control system applications, such as robotics, industrial automation, and electric vehicles. Typically, speed control systems use sensors such as encoders to obtain speed feedback. However, the use of sensors increases the system's cost, complexity, and vulnerability to physical or environmental disturbances. Therefore, alternative approaches are needed that can accurately estimate speed without directly using speed sensors. The system that doesn't use a sensor to measure the controlled variable is a sensorless control[1][2].

Observers or estimators are commonly used solutions to replace sensors in closed-loop control systems. To estimate motor speed, there are several observers that can be used, for example, Model Reference Adaptive System (MRAS)[3], Luenberger [4], Extended Kalman Filter (EKF)[5], and Sliding Mode Observer (SMO) [6]. These methods use complex algorithms. By utilizing a mathematical model of the system and input data such as voltage and current, the observer can estimate internal variables such as motor speed more simply [7]. In previous research, two observer methods have been proposed to estimate the speed of a DC motor. The two methods are the *R*-method observer and the *L-R* method[8]. In that study, the performance of the *L-R* method is better than the *R*-method. In this study, both observers will be applied to estimate the speed of a DC motor that is different from the DC motor in the previous study. From the two observer methods, the method that can estimate the speed of the DC motor with the least error will be selected.

The estimated results are then compared to a reference speed. The difference between the two values is then controlled using a controller. This study uses a PID (Proportional-Integral-Derivative) controller [9]. PID was chosen because it is a mature controller. PID controllers have proven reliable in controlling closed-loop control systems. The purpose of a PID controller is to maintain a stable motor speed at a desired value (set point). The combination of an observer and a PID controller is expected to provide responsive, accurate, and efficient control performance without the need for physical sensors. Some researchers use PID controllers to control DC motor speed in sensorless speed control systems. However, these systems use EKF observers[10]. Others use sensorless motor speed control using Integral Proportional (IP) speed controllers[11].

The PID controller has three parameters that must be determined. They are Proportional gain (K_P), Integral gain (K_I), and Derivative gain (K_D). In this study, the heuristic tuning method will be used to determine the PID controller parameter values. This method was chosen because it has been proven to be able to determine the controller parameters well [12]. In this study, it will be investigated whether the PID controller is able to control the speed of a DC motor in a sensorless control system that uses the *L-R* or *R* method observer properly.

2 Materials and Methods

The block diagram of the sensorless speed control system on the DC motor used in this study is shown in Fig. 1. The system consists of a DC motor, INA 219 current and voltage sensors, an observer, a controller, and a driver. The main circuit of this sensorless speed control system is shown in Fig. 2. The INA 219 current and voltage sensors measure the stator current and voltage. These currents and voltages are used to estimate the motor speed. This estimated speed is then compared to the reference speed. The difference between the two is called the error. A PID controller will calculate the controller output from this error input, which will be used to regulate the motor speed so that the motor speed will be the same or close to the reference speed. All calculations are performed by the Arduino Uno microcontroller. This study used variable DC power supply.

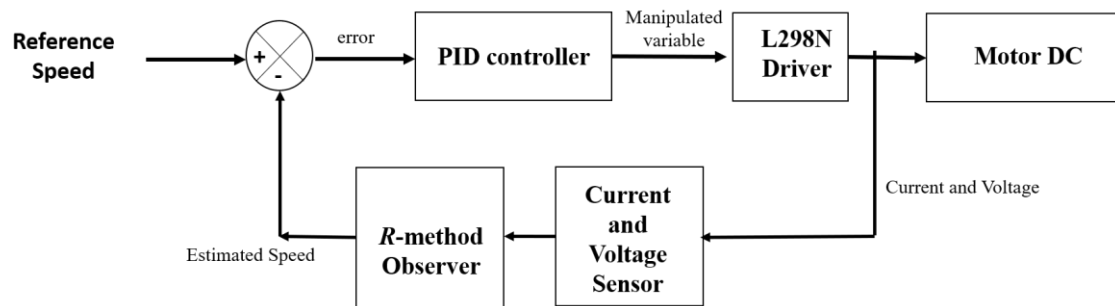


Figure 1. The block diagram of the sensorless speed control system on the DC motor using a PID controller

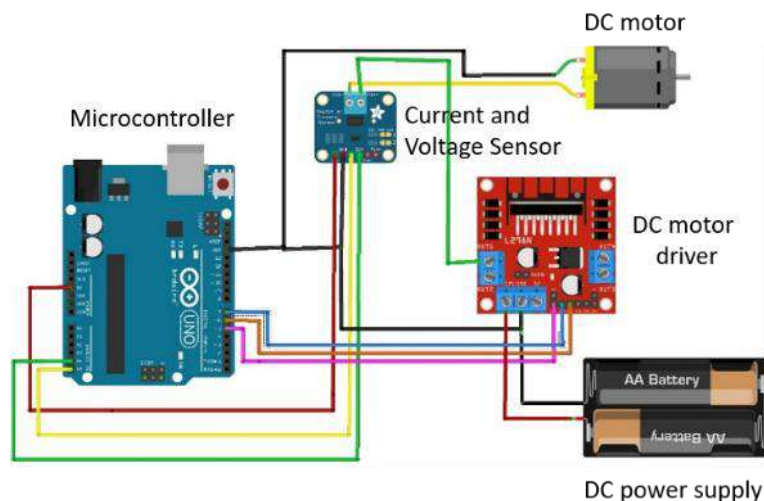


Figure 2. The main circuit of speed sensorless control of DC motor[8]

This sensorless speed control system does not use a speed sensor to measure the DC motor speed directly from the rotor rotation. The speed is estimated from the results of stator current and voltage measurements. Therefore, this system does not use a speed sensor, but uses current and voltage sensors, which are then calculated by the observer into an estimated speed signal.

Two observer techniques were put out in earlier studies to measure a DC motor's speed without the need for a speed sensor. The *R*-method observer and the *L-R* method observer are the suggested approaches [8]. Both of these techniques use the DC motor electrical equation to estimate a DC motor's speed. Equation 1 expresses the DC motor electrical circuit as the current passing through the resistance and inductance of the armature winding[13].

$$v_a = R_a i_a + L_a \frac{di_a}{dt} + e_a \quad (1)$$

where e_a is the back emf, L_a is the armature self-inductance brought on by the armature flux, R_a is the armature resistance, i_a is the armature current, and v_a is the armature supply voltage. Equation 1 can be used to determine the motor's resistance and inductance values, and equation 2 can be used to determine the armature voltage and current values to determine the back emf value (e_a).

$$e_a = v_a - (R_a i_a + L_a \frac{di_a}{dt}) \quad (2)$$

From the equation above, if the motor parameters (resistance and inductance) are known from previous tests, and the current and voltage are measured at any time using current and voltage sensors, then the back emf value can be calculated. The back emf value is proportional to the angular velocity of the rotor. It can be calculated using equation 3.

$$e_a = k_E \omega_m \quad (3)$$

Therefore, if the value of k_E has been obtained from the previous test, then the motor speed can be calculated from the equation. The difference between the *R* and *L-R* observer methods is that if *L-R* calculates the back emf value with equation 2, then the back emf value for the *R* observer method is obtained with equation 2, but the *L* value is ignored.

In this study, the difference between the reference speed and the estimated speed, called the error, will be investigated and then processed by the controller. The controller output, called the control signal, is then used to regulate the speed of the DC motor. All calculations are carried out with an Arduino Uno Microcontroller. The DC motor used in this study has specifications of 12 Vdc input voltage, 150 mA current, and a maximum speed of 7400 RPM, as shown in Fig. 3. The motor has a small current of 150 mA, so an L298N-type DC motor drive is needed, which has a maximum current of 2A.

The controller used is a Proportional Integral Derivative (PID) controller, as shown in Fig.4. The structure of the PID controller is parallel. The value of the controller output is

$$u(t) = K_P e(t) + K_I \int e(t) dt + K_D \frac{de(t)}{dt} \quad (4)$$

where K_P , K_I , and K_D are P , I , and D parameters, respectively.



Figure 3. Motor DC

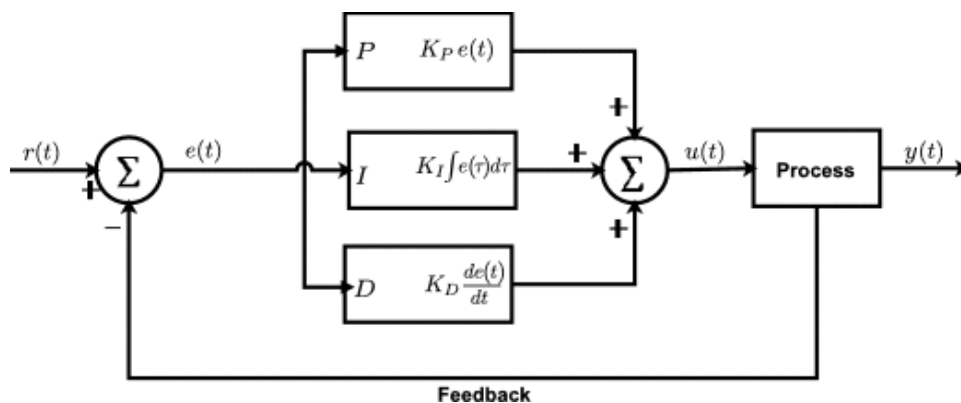


Figure 4. Block diagram of PID controller[14]

With practical experience, the heuristic approach to PID tuning is advanced, involving the human selection of controller variables based on the experimental expertise of a qualified designer who makes use of data on controlled variable estimations[14]. In the heuristic tuning procedure, there are general rules for obtaining approximate or qualitative results[12]. The first amplifier to be adjusted is K_P , starting from the lowest value until a stable response is obtained. In this first step, all K_I and K_D values are set to 0 (disabled). If a response with a steady-state error is obtained, the K_I constant is adjusted starting from a small K_I . If there is still a steady-state error, K_I is increased until the steady-state error is equal to 0. Increasing K_I usually causes a slower response. Therefore, to obtain the desired value, the differential amplifier constant can be increased starting from a small K_D , then increased until the optimum response is achieved. The system has an optimum response if it if it meets the following requirements

- a. fast response

The system produces the intended result in a short amount of time.

- b. minimal overshoot

The output stays within a reasonable range of the goal value. Overshoot is limited to 20%.

- c. short settling time

The output quickly stabilizes around the target value.

- d. small steady-state error

This means that there is little to no variation between the desired value and the final output. The SSE is limited to 2%.

- e. good stability

The system does not fluctuate or become unstable.

The flowchart of the DC motor speed sensorless control system using a PID controller is shown in Fig. 5. The flowchart includes several parts. They are initialization, reading the stator current and voltage values from the current and voltage sensors, calculating the estimated speed carried out by the observer, calculating speed calibration, calculating errors, and calculating the controller output. The moving average method is used to average the current and voltage measured by the current and voltage sensor in order to remove noise.

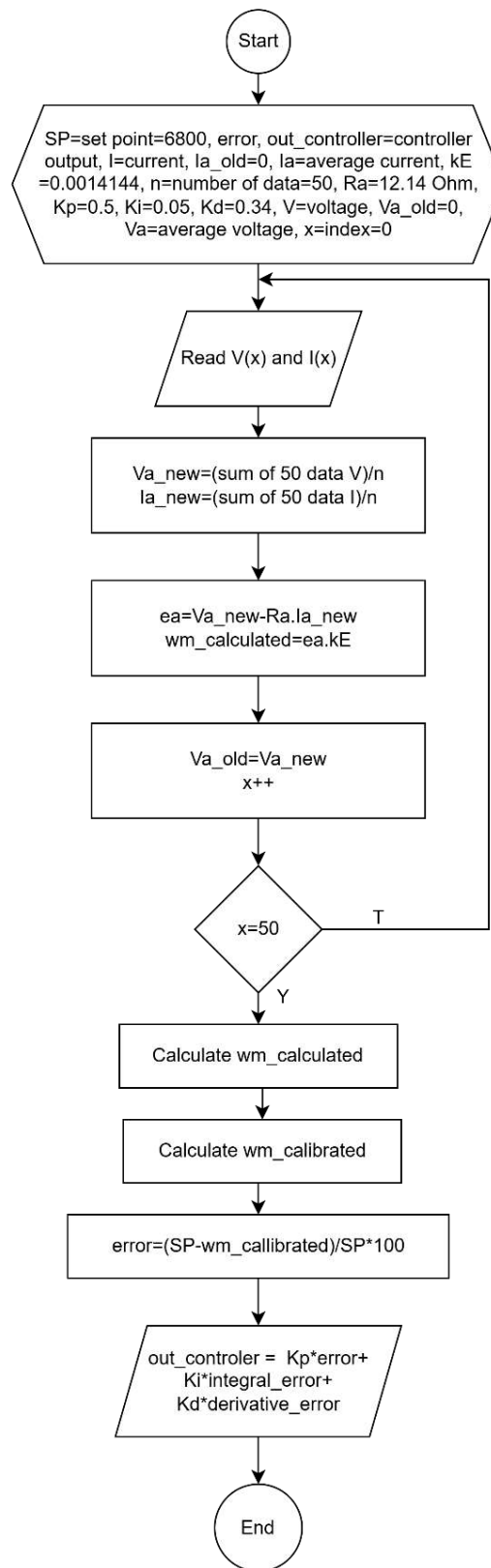


Figure 5. Flowchart of the speed sensorless control using a PID controller

3 Results and Discussions

By the methodology explained above, in this section, the research results will be explained, including the results of measuring motor parameters, observer design, and controlling the speed sensorless of a DC motor using a PID controller, including tuning controller parameters.

3.1. Result of Measuring Motor Parameters

The DC motor parameters measured are motor resistance (R_a) and motor inductance (L_a). Both parameters are measured with an LCR meter. The results of resistance measurements can be seen in Table 1, while the results of inductance measurements can be seen in Table 2.

Table 1. Motor Resistance Measurement

Position	Resistance Values (Ω)
1	12.14
2	12.14
3	12.14
4	12.14
$R_{a_average}$	12.14

Table 2. Motor Inductance Measurement

Position	Inductance Values (mH)
1	93.20
2	93.26
3	93.25
4	93.26
$L_{a_average}$	93.24

Resistance is measured with an RDC variable at several rotor positions, and then the results are averaged. Motor inductance is measured with an LCR meter at a frequency of 100 Hz at several rotor positions, and then the results are averaged. From the two tables, it can be seen that the DC motor used in this study has average $R_a = 12.14\Omega$ and $L_a = 93.24$ mH.

The R_a and L_a measurements above are then used to find the k_E value. Table 3 shows the calculation of the k_E value for the L - R method, while Table 4 shows the calculation of the k_E value for the R method. The k_E value is obtained from equation 3. The e_a value in Table 3 is obtained from equation 2, while the e_a value in Table 4 is obtained from equation 2, but the L_a value is ignored. The average k_E value in the L - R method was 0.001342, while in the R method it was 0.0014144.

Table 3. Calculation of k_E using the L - R method

v_a (V)	ω_m (RPM)	i_a (mA)	e_a	k_E
11.57	1665.0	160.5	9.6067	0.005770
11.62	7285.9	147.9	9.8264	0.001349
11.65	7285.9	174.8	9.5239	0.001307
11.65	7299.3	147.7	9.8602	0.001351
11.65	7301.8	147.4	9.8606	0.001350
k_E average				0.001561

Table 4. Calculation of k_E using the R method

v_a	ω_m (RPM)	i_a (mA)	e_a	k_E
4.88 V	2777.4	63.4	4.11	0.0014798
6.81 V	4081.1	87.5	5.75	0.0014089
11.68 V	7346.3	142.1	9.95	0.0013544
k_E average				0.0014144

3.2. Result of Observer Design

The results of the k_E calculation are then used to calculate the estimated speed, as shown in Tables 5 and 6. When compared with the actual speed measured by the tachometer, the error will be obtained as shown in both tables. It appears that the speed error calculated using the R method is smaller than the L - R method, which is 5.38%. Therefore, the observer that will be used is the R method observer.

Table 5. Comparison of estimated speed and actual speed with the L - R observer method

i_a	v_a	e_a	$\omega_{m_estimated}$	ω_m	Error
(mA)	(V)	(V)	(RPM)	(RPM)	(%)
147.9	11.62	9.8264	6294.44	7285.9	13.60
147.7	11.65	9.8602	6316.1	7299.3	13.47
147.4	11.65	9.8606	6316.35	7301.8	13.49
143.7	11.66	9.9155	6351.52	7323.1	13.27
143.5	11.66	9.9179	6353.08	7317.7	13.18
143.3	11.67	9.9303	6361.03	7314.7	13.04
Average error					13.34

Table 6. Comparison of estimated speed and actual speed with the R observer method

i_a	v_a	e_a	$\omega_{m_estimated}$	ω_m	Error
(mA)	(V)	(V)	(RPM)	(RPM)	(%)
136.2	11	9.35	6608.12	6978.9	5.31
136.3	11	9.35	6607.27	6981	5.35
137.3	11	9.33	6598.68	6973	5.37
137.32	11	9.33	6598.51	6973	5.37
138.02	11	9.32	6592.5	6973	5.46
139.52	11	9.31	6579.63	6972.8	5.64
140.4	11	9.3	6572.08	6972.8	5.75
141.82	11	9.28	6559.89	6972.8	5.92
Average error					5.38

As shown in Table 6, although the error obtained is small, it is still above 5%. To minimize the error, calibration is necessary, as shown in Fig. 6. From the graph, it appears that the relationship between estimated speed and actual speed is linear. From the figure, it appears that the calibration equation follows equation 5.

$$y = 1.0567x + 0.7262 \quad (5)$$

where x is the estimated speed ($\omega_{m_estimated}$) and y is the calibrated speed ($\omega_{m_calibrated}$). After calibration, the estimated speed value is close to the actual speed value with an error of 0.2019%, as shown in Table 7. The actual speed value is measured using a tachometer.

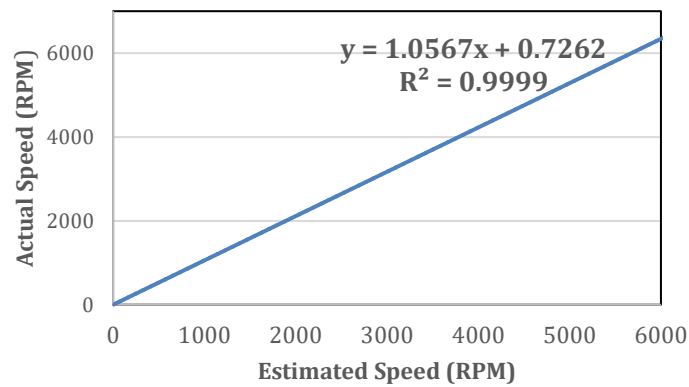


Figure 6. Calculation of calibrated speed

Table 7. The result of calibration

$\omega_{m_calibrated}$ (RPM)	$\omega_{m_tachometer}$ (RPM)	Error (%)
6842.59	6872.5	0.437
6852.45	6872.5	0.293
6853	6872.5	0.285
6856.63	6872.5	0.231
6857.35	6872.5	0.221
6868.48	6872.5	0.059
6870.66	6872.5	0.027
6872.11	6872.5	0.006
6872.66	6872.5	0.002
Average error		0.2019

3.3. Result of the Speed Sensorless of a DC Motor using a PID Controller

The implementation of a sensorless speed control system to control the speed of a DC motor can be seen in Fig. 7. The system consists of a power supply, current and voltage sensors, a microcontroller, an L298N driver, and a laptop to record the measurement results. The system is also equipped with an LCD to monitor the speed. The PID controller parameter values are input through three potentiometers, which are inputs for K_P , K_I , and K_D .

Fig. 8 shows the results of tuning K_P , K_I , and K_D using the heuristic tuning method. When tuning the K_P value with a setpoint of 6800 RPM and $K_P = 0.5$ (Fig. 8(a)), the system response graph shows quite a high overshoot. The motor speed reaches a stable condition at 6710.74 RPM, which indicates a steady state error (SSE) of 1.3% from the setpoint. In addition, the system experiences a speed spike to 6872.8 RPM, which means an overshoot (%OS) of approximately 2.41%. This occurs because the use of proportional gain (K_P) alone causes the system to respond aggressively to errors, but without sufficient damping. As a result, the system tends to exceed the setpoint before finally stabilizing, resulting in overshoot on the response graph.

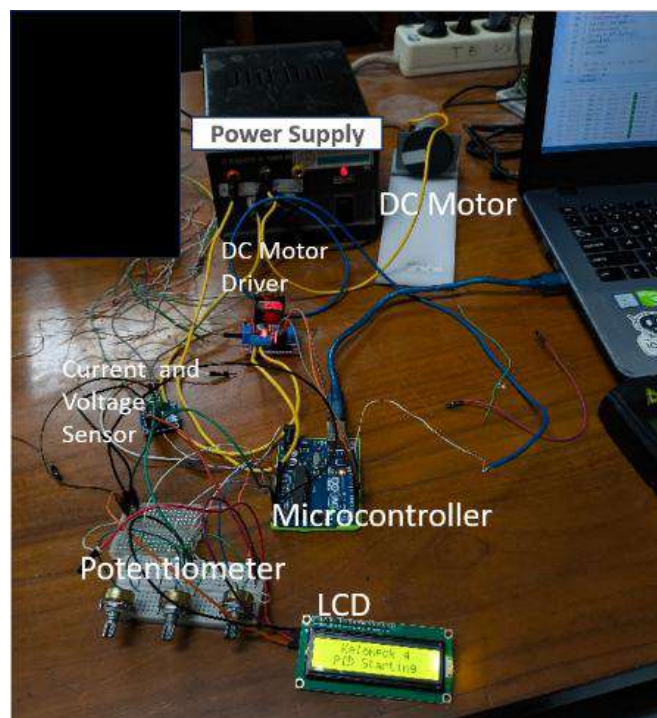


Figure 7. The implementation of a speed sensorless control system

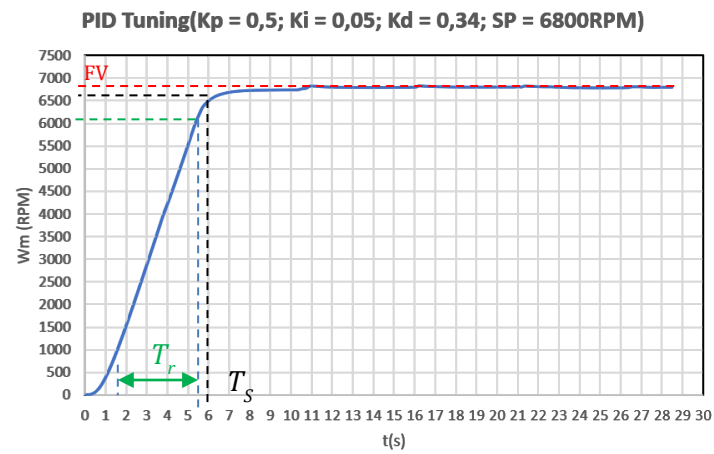
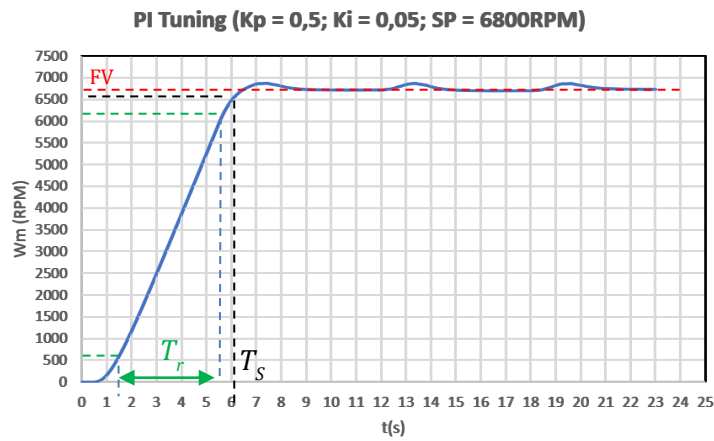
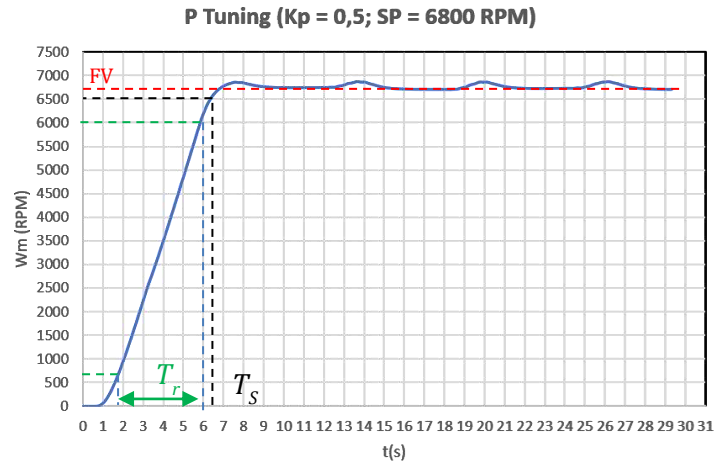


Figure 8. PID tuning using heuristic method

When tuning the K_I value with Set point = 6800RPM, $K_P = 0.5$, and $K_I = 0.05$ (Fig. 8(b)), it produces a system response graph with a slightly lower overshoot value than the K_P tuning process alone. The motor speed reaches a stable condition at 6733 RPM, which indicates a slightly smaller steady state error of 1% of the set point. However, the system experiences a speed spike up to 6874 RPM, which results in an overshoot of about 2.1% which is greater than the K_P tuning. This occurs because the addition of the integral component (K_I) makes the system more aggressive in correcting errors, including errors that accumulate over time. Without damping from the derivative component (K_D), the control action becomes excessive and tends to cause a response spike, resulting in a larger overshoot.

The last is the tuning of the K_D value with a setpoint of 6800 RPM using a combination of $K_P = 0.5$; $K_I = 0.05$; and $K_D = 0.34$ values (Fig. 8(c)) can provide or produce a more optimal system response graph because it produces a very small oscillation of 0.39% or can be considered no oscillation (equal to 0). The results of the system response show that the resulting DC motor speed response can reach a value of 6792.76 RPM stably and quickly approaching the set point value of 6800 RPM, with a very small level of deviation. The Steady State Error (SSE) value produced in this tuning is also the lowest, which is around 0.11%, indicating that the error between the actual value and the target speed value is at a minimum level. The system has the smallest settling time value (T_s) of the other tunings, i.e., 6 seconds. Therefore, this PID value tuning was chosen as the best configuration because it can provide stability, accuracy, and response that meet the needs of the sensorless DC motor speed control system implemented in this study. The performance results can be seen in Table 8.

Table 8. Comparison of Controller Performance

Controller	Final	T_r (s)	T_s (s)	%OS (%)	SSE (%)
	Value (FV)				
P	6710.74	4.3	6.4	1.31	1.3
PI	6733.95	4	6.2	0.97	0.97
PID	6792.76	4	6	0.11	0.11

4 Conclusions

From the results of the study, it can be concluded that the R method of the DC motor speed control system with a PID controller can work well. The use of observer method 2 was chosen because it was proven to provide the most accurate speed estimation with an average error of only 0.2019%, smaller than method 1. This speed estimation is then used as feedback for the PID controller, which is tuned using the heuristic method. The PID parameter tuning results with a configuration of $K_P = 0.5$, $K_I = 0.05$, and $K_D = 0.34$ demonstrated optimal system performance, producing an actual speed of 6792.76 RPM from a setpoint of 6800 RPM, with a very small steady-state error (SSE) of 0.1%, a rise time (T_r) of 4 seconds, and a settling time (T_s) of 6 seconds. Therefore, the combination of observer method 2 and a PID controller proved capable of controlling DC motor speed accurately, responsively, and efficiently without relying on a physical speed sensor.

References

- [1] P. Vas, *Sensorless Vector and Direct Torque Control*. Oxford University Press, 1998.
- [2] A. Glumineau and J. de León Morales, *Sensorless AC electric motor control*. Springer, 2015.
- [3] B. Wuri Harini, A. Subiantoro, and F. Yusivar, "Stability of the Rotor Flux Oriented Speed Sensorless Permanent Magnet Synchronous Motor Control," *IEEE Int. Symp. Ind. Electron.*, vol. 2018-June, pp. 283–289, 2018, doi: 10.1109/ISIE.2018.8433862.
- [4] B. W. Harini, "The Effect of Motor Parameters on the Induction Motor Speed Sensorless Control System using Luenberger Observer," vol. 4, no. 1, pp. 59–74, 2022.

- [5] G. Tian, Y. Yan, W. Jun, Z. Y. Ru, and Z. X. Peng, "Rotor Position Estimation of Sensorless PMSM Based on Extended Kalman Filter," *2018 IEEE Int. Conf. Mechatronics, Robot. Autom. ICMRA 2018*, pp. 12–16, 2018, doi: 10.1109/ICMRA.2018.8490558.
- [6] J. C. Gamazo-Real, E. Vázquez-Sánchez, and J. Gómez-Gil, "Position and speed control of brushless dc motors using sensorless techniques and application trends," *Sensors*, vol. 10, no. 7, pp. 6901–6947, 2010, doi: 10.3390/s100706901.
- [7] E. Engineering, U. Indonesia, and K. U. I. Depok, "Jurnal Teknologi A Synchronization Loss Detection Method," vol. 4, pp. 47–54, 2020.
- [8] Bernadeta Wuri Harini, Martanto, and Tjendro, "Comparison of Two DC Motor Speed Observers on Sensorless Speed Control Systems," *J. Nas. Tek. Elektro dan Teknol. Inf.*, vol. 11, no. 4, pp. 267–273, 2022, doi: 10.22146/jnteti.v11i4.5019.
- [9] K. Ogata, *Teknik Kontrol Automatik*. Jakarta: Erlangga, 1997.
- [10] K. Al-Maliki, Abdullah Y.; IQBAL, "FLC-based PID controller tuning for sensorless speed control of DC motor," in *2018 IEEE International Conference on Industrial Technology (ICIT)*, 2018, pp. 169–174.
- [11] F. Yusivar and S. Wakao, "Minimum requirements of motor vector control modeling and simulation utilizing C MEX S-function in MATLAB/SIMULINK," *Proc. Int. Conf. Power Electron. Drive Syst.*, vol. 1, pp. 315–321, 2001, doi: 10.1109/peds.2001.975333.
- [12] Bernadeta Wuri Harini, "Pengaruh Parameter Motor pada Sistem Kendali tanpa Sensor Putaran," *J. Nas. Tek. Elektro dan Teknol. Inf.*, vol. 10, no. 3, pp. 236–242, 2021, doi: 10.22146/jnteti.v10i3.1848.

-
- [13] M. Vidlak, L. Gorel, P. Makys, and M. Stano, "Sensorless speed control of brushed dc motor based at new current ripple component signal processing," *Energies*, vol. 14, no. 17, 2021, doi: 10.3390/en14175359.
- [14] S. B. Joseph, E. G. Dada, A. Abidemi, D. O. Oyewola, and B. M. Khammas, "Metaheuristic algorithms for PID controller parameters tuning: review, approaches and open problems," *Heliyon*, vol. 8, no. 5, p. e09399, 2022, doi: 10.1016/j.heliyon.2022.e09399.

This page intentionally left

In Vitro Regeneration of Dendrobium Through Somatic Embryogenesis from Leaf Explants

Rahmadyah Hamiranti^{1*}, Nanang Wahyu Prajaka¹, Yeni¹,
Desi Maulida¹, Lisa Erfa¹

¹*Department of Food Plant Cultivation Lampung State Polytechnic,
Jln Soekarno Hatta Rajabasa Bandar Lampung, Indonesia*

**Corresponding Author: rahmahamiranti@polinela.ac.id*

(Received 13-06-2025; Revised 21-08-2025; Accepted 25-08-2025)

Abstract

The purpose of this study was to determine the best basic medium type and TDZ concentration for the development of somatic embryos from Dendrobium orchid leaf explants. Three replications were arranged factorial (2x3) in a completely randomized design for this study. First, there were two types of basic media, ½ MS and MS. The second factor was the concentration of cytokinin thidiazuron 1, 2 and 3 mg/l. Each experimental unit consisted of 5 culture bottles; each contained 5 explants. The research results showed that 1) The use of both types of basic media was able to induce callus formation on Dendrobium 'Gradita 31' orchid from leaf explants. 2) The use of 3 mg/l thidiazuron which combined with ½ MS or MS media was able to form primary callus faster than other treatments. 3) The higher percentage for embryo somatic and shoots formation were also found in 3 mg/l thidiazuron.

Keywords: Dendrobium, embryogenesis, thidiazuron

1 Introduction

Dendrobium orchids have high economic value by contribute for 34% of the orchid business, followed by Phalaenopsis, Vanda and other genera [1]. The high market, due to interest and selling price for Dendrobium seedlings is the reason for increasing productivity in the availability of true to type seedling with a large number. Propagation of Dendrobium orchids can be done with generatively (through seeds) or vegetatively. Orchid seeds cannot be grown directly on planting media in the field because orchid seeds do not have endosperm.



The process of fruit formation naturally rare occurs. Moreover, the formation of *Dendrobium* orchid fruit takes quite a long time, around 3-4 months until the seeds are ready to be sown as in vitro culture [2]. Vegetative propagation through seedlings cannot be used as a mass propagation method because the planting material produced is very limited (2-4 seedlings per year). Meanwhile, keiki (seedling at the tip of the bulb) will only appear when mature *Dendrobium* orchids undergo stress condition [3].

Numerous studies have been conducted on orchid somatic embryogenesis in vitro. Hapsoro and Yusnita [4] stated that somatic embryogenesis is defined as the process of creating an embryo (a plant structure that already has root and shoot formation) from nonzygotic parts of the plant body (leaves, roots, stems, hypocotyl, etc.). The composition of the basic media is an important factor in the somatic embryogenesis of orchid plants. MS media is most widely used in the regeneration process through somatic embryogenesis in several types of orchids such as *Dendrobium*, *Spathoglottis*, *Cattleya*, *Rhynchostylis* and *Grammatophyllum* [5]–[11]. Several research results [12]–[15] showed that the use of half strength of MS media was able to supply sufficient nutrients to encourage the development of somatic embryos in *Dendrobium* and *Phalaenopsis*.

The ability of plant regeneration to respond on nutrients and growth regulators is species-specific, meaning that plants from different genotypes will respond differently in their plant regeneration patterns even though they have been given the same treatment[16] . In vitro plant regeneration is regulated by the balance and interaction of plant growth regulators (PGR) contained in the explant (endogenous) and with the PGR absorbed from the growth medium (exogenous). In certain proportions, the use of PGR can stimulate embryo formation [17]. Therefore, it is necessary to carry out a research regarding regeneration procedures through somatic embryogenesis which appropriate to the genotype of *Dendrobium* orchid plants. This research aimed to obtain the best type of basic media and concentration of TDZ for the formation of somatic embryos from *Dendrobium* orchid leaf explants.

2 Material and Methods

This research was carried out at the Tissue Culture Laboratory, Department of Food Plant Cultivation, Lampung State Polytechnic, Rajabasa, Bandar Lampung, from April to September 2023. The plant material was seedlings of *Dendrobium* 'Gradita 31' that obtained from Balai Penelitian Tanaman Hias. The explant was leaf segments with a size of $\pm 1 \text{ cm}^2$. Explant was planted as adaxial faces on medium (the top surface of the leaf touched the medium). Planting was carried out in LAFC under aseptic conditions. Explants were cultured in a culture room at a temperature of $25^\circ\text{C} \pm 2^\circ\text{C}$ with a subculture interval of 4 weeks on the same medium.

The media were MS [18] and half strength of MS ($\frac{1}{2}$ MS). $\frac{1}{2}$ MS was modified from MS into half concentration of macro elements. These two basic media were enriched with thiamine-HCL (0.1 mg/l), nicotinic acid (0.5 mg/l), pyridoxine-HCL (0.5 mg/l) and myo inositol (100 mg/l). Other ingredients added were sucrose 30 g/l, ascorbic acid 200 mg/l, citric acid 150 mg/l and plant growth regulator (as needed). Then adjust the pH to 5.8. The addition of 1 N NaOH was carried out if the initial pH was less than 5.8, conversely if the initial pH was more than 5.8 then 1 N HCL was given. After adjusting the pH to 5.8 then the media was cooked by adding 7 g/l of agar powder. After its boiled then the media was poured into culture bottles at a rate of 20 ml per bottle. The media-filled bottles were sealed with transparent plastic, fastened with rubber bands, and autoclaved for 15 minutes at 121°C with a pressure of $1,2 \text{ kg/cm}^2$.

The sterilizing tools for culture bottle sterilized were detergent, 0.5% sodium hypochlorite solution, and 70% alcohol. A variety of instruments are used, such as pipettes, pH meter, analytical scales, culture bottles, aluminum foil, plastic wrap, rubber bands, label paper, cameras, hand sprayers, magnetic stirrers, beakers, tweezers, scalpels, and laminar air flow cabinets (LAFC).

A completely randomized design with three replications arranged factorial (2×3) was used in this study. The first factor was the type of basic media, $\frac{1}{2}$ MS, and full MS. The second factor was the TDZ concentration 1, 2 and 3 mg/l combined with 2,4-D 0.5 mg/l.

Each experimental unit consisted of 5 culture bottles, each contained 5 explants. The observation data for each observation variable was analyzed for diversity using the F test, and if there are significant differences between treatments, then the middle value will be separated using the tukey test at a significance level of 0.05. Observation variables included callus performing periode, percentage of explants forming callus and somatic embryo, and percentage of shoot formation.

3 Results and Discussion

Planting were used the part of the leaf with leaf vein which the tip, base and sides of the leaf were not used. After that, the explants were cultured in a culture room under dark room conditions. The first primary callus appeared on day 77 after planting in the 3 mg/l TDZ treatment on both $\frac{1}{2}$ MS and MS media. Furthermore, primary callus appeared on day 85 in the TDZ 2mg/l treatment and day 87 in the TDZ 1mg/l treatment. (Table 1).

The primary callus was subcultured into the same media for forming somatic embryo. At week 16, somatic embryo formation appeared. It started from the embryogenic callus which forms a green globular structure. In this phase, globular embryos were transferred to a culture room with bright room conditions. Two weeks after being transferred, the globular embryo then elongates (coleoptile) and began to form shoots after 1 week. Physiological and morphological changes in shoot formation through somatic embryos from leaf explants can be seen in Fig 1.

According to [17] the indirect embryogenesis phase started from the induction phase where the explants run into dedifferentiation of cells that form callus and become competent, in response to the addition of TDZ. Afterwards, the competent cells respond to become embryogenic callus which then transformed to somatic embryo phase with the globular stage then the coleoptile stage and finally the regeneration phase to become a shoot. The results of research [19] state that the embryogenesis process in *Phalaenopsis amabilis* plants started from the formation of pro-embryos which appear close to the cut side of leaf. The pro-embryo then enlarges and forms a globular embryo. Mature embryo does not simultaneously undergo

development into coleoptiles, procambium and then shoots emerge. But there are some cells that are left behind in their development so they are still in the globular stage. Hence, in one clump of somatic embryo, we can see all the various stages from globular, coleoptile and then developing into plantlets.

The data from Table 1 showed that the TDZ concentration had a different effect, but the two basic media had the same effect on the time of callus performed, somatic embryo formation and shoots formation. Giving TDZ 3mg/l gave the highest percentage of somatic embryo formation and shoots formation. This was in line with several studies [12], [20], [21] that the use of TDZ 3 mg/l was able to provide the highest response in each phase of orchid embryogenesis. The performance of shoots can be seen on Fig 2.

The ability of TDZ to induce the formation of somatic embryo is thought to be because TDZ could attach to one side of the Cytokinin Binding Protein (CBP) that found in the cell membrane so that it can stimulate endogenous cytokinin production. Thus, increasing cytokinin levels in cells will stimulate cell division and stimulate tissue morphogenesis [22]. The use of ½ MS basic media was able to provide the same effect as MS media. This means that in the process of forming somatic embryo for *Dendrobium* 'Gradita 31' orchid, ½ MS media was able to supply sufficient nutrients during the embryogenesis process. ½ MS media can be used as a cheaper alternative in the mass multiplication process through somatic embryogenesis, because the concentration of chemical ingredients was only half of the composition from MS media. This is also supported by several studies [12]–[15] that the use of ½ MS media is able to induce the formation of somatic embryos in *Dendrobium* and *Phalaenopsis*. Low nitrogen concentrations in the media were also reported able to induce higher embryo growth in coffee plants [23], *Santalum album* [24] and sorghum [25].

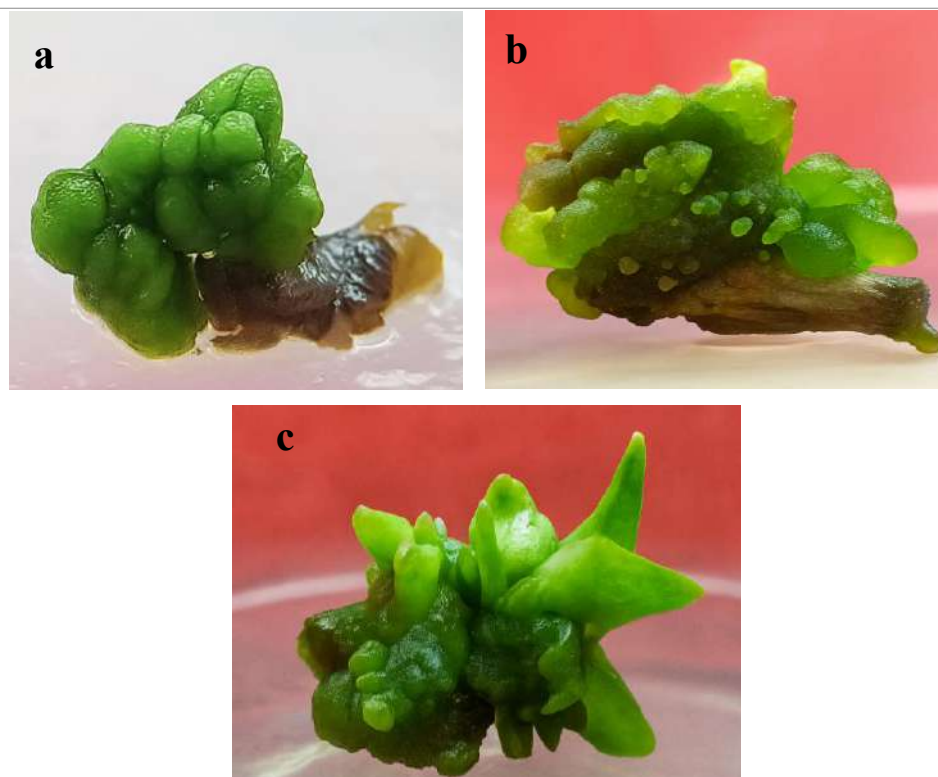


Figure 1. Somatic embryogenesis development of *Dendrobium* ‘Gradita 31’ (a) globular phase, (b) elongated embryo (coleoptile), (c) shoots formation.

Many studies have reported that auxin and cytokinin play an important role in the process of plant embryogenesis. On this research, we used 0.5mg/l 2,4-D as the auxin. The use of low concentrations of auxin combined with high levels of cytokinin has also been widely reported in orchid plants such as *Dendrobium malones* [11], *P. amabilis* [19], *Vanda tessellates* [26] and *Phalaenopsis* ‘Sogo Vivien’ [27]. On the other hand, another research for embryogenesis of *Phalaenopsis* ‘Hong kong’ single used of 3mg/l TDZ for whole stages was able to support embryogenesis process [12]. Other report [28] single auxin (3mg/l 2,4-D) was able to induce somatic embryo of *Vanda sumatrana*. The differences in the response of each orchid species to the type and concentration, cytokinin-auxin alone or in combination are very clearly visible, even though they are still in the same plant family. This is because plants have different genetic characteristics and their ability to respond the signals from the same growth regulator [16], [29], [30].

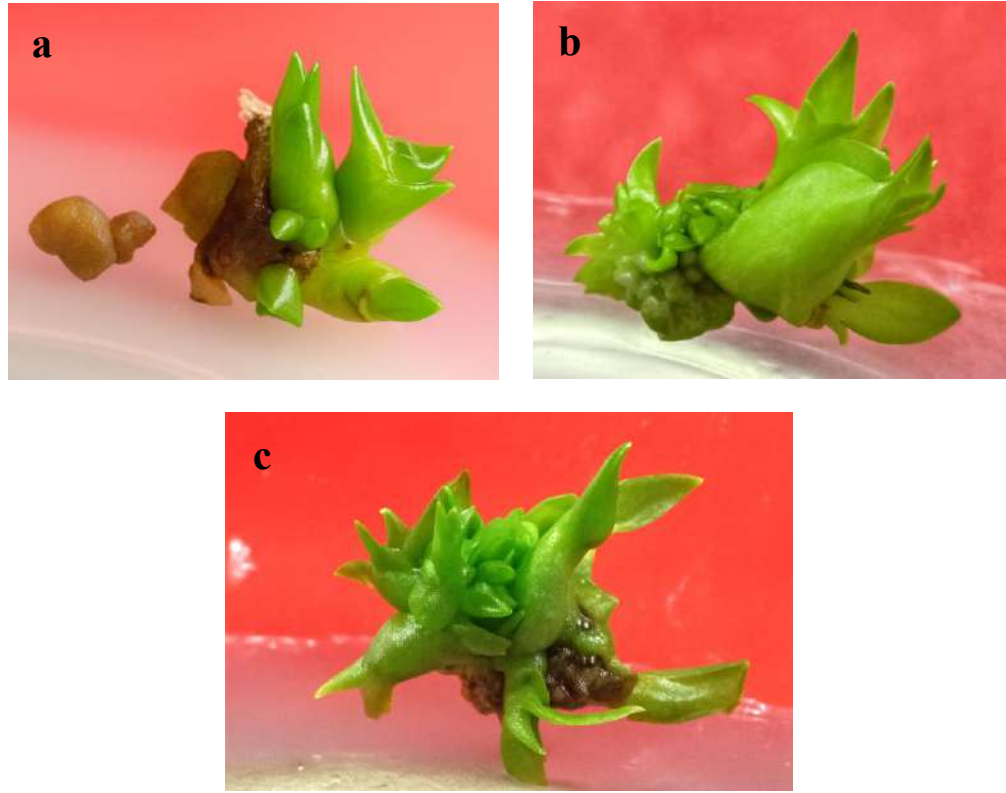


Figure 2. Shoots performance on different concentration of TDZ (a) 1mg/l, (b) 2mg/l, (c) 3mg/l.

Table 1. Effect of basic media and different concentration of TDZ on somatic embryogenesis *Dendrobium* 'Gradita 31'

Treatments	Callus performing period (days)	Primary callus (%)	Embryo somatic (%)	Shoots forming (%)
$\frac{1}{2}$ MS+ 1mg/l TDZ	86.00b	76.67a	58.33b	71.67b
$\frac{1}{2}$ MS+ 2mg/l TDZ	85.33b	83.33a	73.33ab	85.00ab
$\frac{1}{2}$ MS+ 3mg/l TDZ	77.00a	91.67a	91.67a	91.67a
MS+ 1mg/l TDZ	87.00b	75.00a	53.33b	70.00b
MS+ 2mg/l TDZ	85.67b	81.67a	66.67ab	83.33ab
MS+ 3mg/l TDZ	77.67a	91.67a	93.33a	90.00a

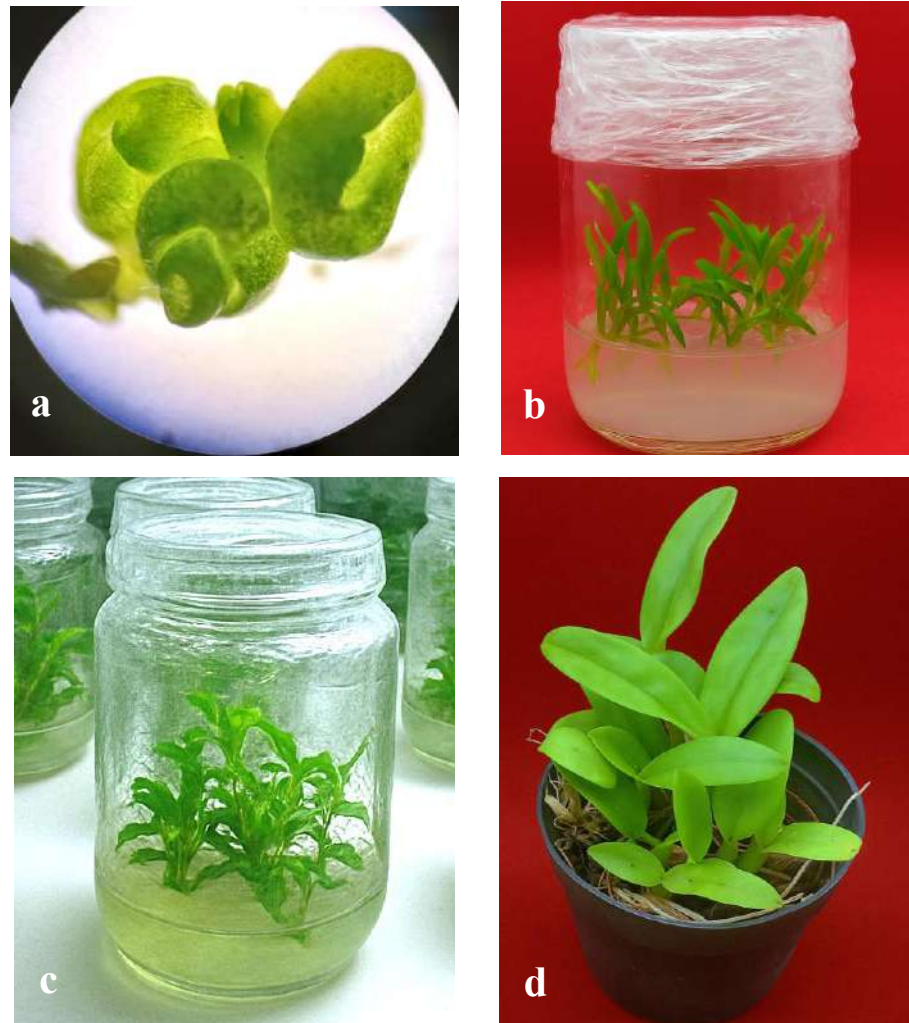


Figure 3. Performance seedling of *Dendrobium* 'Gradita 31' (a) globular phase microscopic (40x), (b) seedling (24 week), (c) seedling for acclimatization (30 week), (d) seedling 5 month after acclimatization

4 Conclusions

Based on the research on the effect of basic media and different concentration of TDZ on somatic embryogenesis *Dendrobium* 'Gradita 31', the conclusions were drawn:

1. The use of both types of basic media was able to induce callus formation on *Dendrobium* 'Gradita 31' orchid from leaf explants.
2. The use of 3 mg/l thidiazuron which combined with $\frac{1}{2}$ MS or MS media was able to form primary callus faster than other treatments.
3. The higher percentage for embryo somatic and shoots formation were also found in 3 mg/l thidiazuron.

Acknowledgments

The authors would like to acknowledge the Lampung State Polytechnic, Indonesia, for its financial support and also for the opportunity to carry out the research in plant tissue culture laboratory.

References

- [1] Kumparan Bisnis, 80% bibit anggrek di RI masih impor. [Online]. Available: <https://kumparan.com/@kumparanbisnis/80-bibit-anggrek-di-ri-masih-impor>, 2018.
- [2] D. Widiastoety, N. Solvia, and M. Soedarjo, "Potensi anggrek dendrobium dalam meningkatkan variasi dan kualitas anggrek bunga potong", *Jurnal Litbang Pertanian*, vol. 29 no.3, pp: 101–106, 2010, <https://garuda.kemdikbud.go.id/documents/detail/1672781>.
- [3] A. W. Purwanto, *Anggrek Budi Daya dan Perbanyakan*. 2016. [Online]. Available: www.upnyk.ac.id

-
- [4] D. Hapsoro and Yusnita, “Kultur Jaringan “Teori dan Praktik”, *Andi Yogyakarta Publisher*, 167 hlm, 2018.
- [5] A. N. Pyati, “In vitro Propagation of orchid (*Dendrobium ovatum* (L.) Kraenzl.) through Somatic Embryogenesis,” *Plant Tissue Cult Biotechnol*, vol. 32, no. 1, 2022, doi: 10.3329/ptcb.v32i1.60472.
- [6] M. Manokari, S. Priyadharshini, and M. S. Shekhawat, “Direct somatic embryogenesis using leaf explants and short term storage of synseeds in *Spathoglottis plicata* Blume,” *Plant Cell Tissue Organ Cult*, vol. 145, no. 2, pp. 321–331, May 2021, doi: 10.1007/s11240-021-02010-9.
- [7] T. Van Minh, “Micropropagation of *Rhynchostylis gigantea* orchid by somatic embryogenic cultureS,” *CBU International Conference Proceedings*, vol. 7, pp. 969–974, Sep. 2019, doi: 10.12955/cbup.v7.1486.
- [8] A.O. Melisa, “Pemberian Kombinasi 2,4-D dan Kinetin Terhadap Induksi Kalus dan Protocorm Like Bodies (PLB) Anggrek *Grammatophyllum scriptum* Secara In Vitro”, *Skripsi Jurusan Biologi Fakultas Matematika dan Ilmu Pengetahuan Alam Universitas Sebelas Maret Surakarta*, 2011.
- [9] J. Reddy St Joseph and J. Reddy, “In Vitro Organogenesis and Micropropagation of The Orchid Hybrid, *Cattleya Naomi* Kerns. *European Journal of Biomedical and Pharmaceutical Sciences*”, vol. 3, no.9, pp: 388–393, 2019. [Online]. Available: www.ejbps.com388

- [10] P. Sujjaritthurakarn and K. Kanchanapoom, "Efficient Direct Protocorm-Like Bodies Induction of Dwarf Dendrobium using Thidiazuron", [Online]. Available: www.notulaebiologicae.ro
- [11] S. Anjum, M. Zia, and M. F. Chaudhary, "Investigations of different strategies for high frequency regeneration of *Dendrobium malones* 'Victory,'" *Afr J Biotechnol*, vol. 5, no. 19, pp. 1738–1743, 2006, [Online]. Available: <http://www.academicjournals.org/AJB>
- [12] H. N. Boldaji, S. D. Dylami, S. Aliniaiefard, and M. Norouzi, "Efficient method for direct embryogenesis in phalaenopsis orchid," *Int J Horti Sci Technol*, vol. 9, no. 2, pp. 37–50, Dec. 2021, doi: 10.22059/ijhst.2020.296696.339.
- [13] K. Juntada, S. Taboonmee, P. Meetum, S. Poomjae, and P. N. Chiangmai, "Somatic Embryogenesis Induction from Protocorm-like Bodies and Leaf Segments of *Dendrobium sonia* 'Earsakul,'" *Silpakorn U Science & Tech J*, vol. 9, no. 2, pp. 9–19, 2015.
- [14] P. L. Lee and J. T. Chen, "Plant regeneration via callus culture and subsequent in vitro flowering of *Dendrobium huoshanense*," *Acta Physiol Plant*, vol. 36, no. 10, pp. 2619–2625, Sep. 2014, doi: 10.1007/s11738-014-1632-7.
- [15] F. Rachmawati, D. Permanik, R. B. Mayang, and B. Winarto, "Protokol Perbanyakan Masal Dendrobium 'Balithi CF22-58' secara In Vitro Melalui Embriogenesis Somatik Tidak Langsung (In Vitro Propagation Protocol of Dendrobium 'Balithi CF22-58' via Indirect Somatic Embryogenesis)," *Jurnal Hortikultura*, vol. 29, no. 2, p. 137, Jun. 2020, doi: 10.21082/jhort.v29n2.2019.p137-146.

- [16] Yusnita, “Kultur Jaringan Tanaman: Sebagai Teknik Penting Bioteknologi Untuk Menunjang Pembangunan Pertanian”, *Aura Bandar Lampung Publisher*, 86 hlm, 2015.
- [17] D. Hapsoro and Yusnita, “Embriogenesis Somatik In Vitro Untuk Perbanyak Klonal dan Pemuliaan Tanaman”, *Aura Yogyakarta Publisher*, 65 hlm, 2022.
- [18] T. Murasnige and F. Skoog, “A Revised Medium for Rapid Growth and Bio Assays with Tohaoco Tissue Cultures”, *Physiologia Plantarum*, vol. 15, 1962.
- [19] W. Mose, B. S. Daryono, A. Indrianto, A. Purwantoro, and E. Semiarti, “Direct Somatic Embryogenesis and Regeneration of an Indonesian orchid *Phalaenopsis amabilis* (L.) Blume under a Variety of Plant Growth Regulators, Light Regime, and Organic Substances,” *Jordan Journal of Biological Sciences*, vol. 13, no.4, pp: 509–518, 2020, [Online]. Available: <https://jjbs.hu.edu.jo/files/vol13/n4/Paper%20Number%2013.pdf>.
- [20] H. H. Chung, J. T. Chen, and W. C. Chang, “Cytokinins induce direct somatic embryogenesis of Dendrobium Chiengmai Pink and subsequent plant regeneration,” *In Vitro Cellular and Developmental Biology - Plant*, vol. 41, no. 6, pp. 765–769, Nov. 2005, doi: 10.1079/IVP2005702.
- [21] H. H. Chung, J. T. Chen, and W. C. Chang, “Plant regeneration through direct somatic embryogenesis from leaf explants of Dendrobium,” *Biol. Plant*, vol. 51, no.2, pp: 346–350, June 2007, doi: 10.1007/s10535-007-0069-x.
- [22] B. Guo, B. H. Abbasi, A. Zeb, L. L. Xu, and Y. H. Wei, “Thidiazuron: A multi-dimensional plant growth regulator,” *African Journal of Biotechnology*, vol. 10, no. 45. Academic Journals, pp. 8984–9000, 2011. doi: 10.5897/ajb11.636.

- [23] N. P. Samson, C. Campa, L. Le Gal, M. Noirot, G. Thomas, T. S. Lokeswari and A. de Kochko, "Effect of primary culture medium composition on high frequency somatic embryogenesis in different *Coffea* species," *Plant Cell Tissue Organ Cult*, vol. 86, no. 1, 2006, doi: 10.1007/s11240-006-9094-2.
- [24] S. Das, S. Ray, S. Dey, and S. Dasgupta, "Optimisation of sucrose, inorganic nitrogen and abscisic acid levels for *Santalum album* L. somatic embryo production in suspension culture," *Process Biochemistry*, vol. 37, no. 1, 2001, doi: 10.1016/S0032-9592(01)00168-6.
- [25] L. A. Elkonin and N. V. Pakhomova, "Influence of nitrogen and phosphorus on induction embryogenic callus of sorghum," *Plant Cell Tissue Organ Cult*, vol. 61, no. 2, pp. 115–123, 2000, doi: 10.1023/A:1006472418218.
- [26] M. Manokari, R. Latha, S. Priyadharshini, P. Jogam, and M. S. Shekhawat, "Short-term cold storage of encapsulated somatic embryos and retrieval of plantlets in grey orchid (*Vanda tessellata* (Roxb.) Hook. ex G. Don)," *Plant Cell Tissue Organ Cult*, vol. 144, no. 1, pp. 171–183, Jan. 2021, doi: 10.1007/s11240-020-01899-y.
- [27] P. D. Kasi and E. Semiarti, "Pengaruh Thidiazuron dan Naphtalene Acetic-Acid untuk Induksi Embriogenesis Somatik dari Daun Anggrek *Phalaenopsis* 'Sogo Vivien,'" 2016. [Online]. Available: <https://www.researchgate.net/publication/316644331>
- [28] A. T. Astuti, Z. A. Noli and S. Suwirmen, "Induksi Embriogenesis Somatik Pada Anggrek *Vanda Sumatrana* Schltr. dengan Penambahan Beberapa Konsentrasi Asam 2,4-Diklorofenoksiasetat (2,4-D) (Induction Somatic Embryogenesis of Orchid *Vanda Sumatrana* Schltr. with 2,4-Diclorofenoxiacetad Acid (2,4-D) Addition)," *Jurnal Biologi Universitas Andalas (J. Bio. UA.)*, vol. 7, no. 1, pp. 6–13.

- [29] E. J. Naranjo, O. Fernandez Betin, A. I. Urrea Trujillo, R. Callejas Posada, and L. Atehortúa Garcés, “Effect of genotype on the in vitro regeneration of *Stevia rebaudiana* via somatic embryogenesis,” *Acta Biolo Colomb*, vol. 21, no. 1, 2015, doi: 10.15446/abc.v21n1.47382.
- [30] T. W. Yam and J. Arditti, *Micropropagation of Orchids: Third Edition*, vol. 1–3. 2017. doi: 10.1002/9781119187080.

Braille Pattern Detection Modeling Using Inception V3 Architecture Using Median Filter Implementation and Segmentation

Abdul Latip^{1*}, Siti Yuliyanti¹, Muhammad Al-Husaini¹

¹ *Department of Informatics, Faculty of Engineering, Siliwangi University, Tasikmalaya, West Java, Indonesia*

**Corresponding Author: abdhul.latyp@gmail.com*

(Received 19-06-2025; Revised 27-07-2025; Accepted 25-08-2025)

Abstract

This study aims to detect Braille letter patterns using the InceptionV3 architecture combined with the application of median filter and image segmentation. The dataset consists of 4,160 Braille images, with an average of 160 images for each letter from A to Z. The data is divided into 3,900 images for training, which are then split into 3,120 images for training and 780 images for validation, and 260 images are used for testing. Each image is resized to 299x299 pixels before being fed into the model. This study uses 100 epochs and applies early stopping to avoid overfitting. Two learning rate values are tested, namely 0.001 and 0.0001. The results show that the application of a median filter and segmentation significantly improves model performance, producing better accuracy, precision, recall, and F1 values compared to models without these techniques. At a learning rate of 0.001, the model achieves 99.65% accuracy, 99.62% precision, and 99.61% recall. On the other hand, without a median filter and segmentation at a learning rate of 0.0001, although accuracy and precision decreased, the values still reached 99.65% and 99.62%.

Keywords: Brailier, Inception V3, Learning Rate, Median Filter, Segmentation

1 Introduction

Braille is a standard writing system for the blind, where each character is represented by a combination of dots in a specific pattern that can be felt with the fingertips [1]. Given the challenges in reading Braille, especially the need for high finger sensitivity, deep learning technology can play an important role in assisting automatic Braille translation [2].



Deep learning, especially Convolutional Neural Networks (CNN), has proven effective in image and pattern recognition [3]. Deep learning uses artificial neural networks with many layers to automatically learn data representations. CNN, inspired by the structure of human neural networks, works with two main phases: backpropagation for training and feedforward for image classification. Before classification, preprocessing processes such as bagging and cropping are used to focus on the objects to be classified [3]. This CNN has many architectures, one of which is InceptionV3. InceptionV3 is one of the most famous CNN architectures. This architecture is known for being easy to train and very accurate in image classification. InceptionV3 is capable of capturing features at different scales, from the smallest to the most complex, by using a combination of convolutional layers with various kernel sizes [3].

Previous studies have shown that the Inception V3 architecture can achieve high accuracy in Braille recognition, with accuracy values reaching 95.03% to 99.87% [4]. However, the detection performance of Braille models is often affected by noise in images. Therefore, preprocessing techniques such as median filtering and image segmentation become very important [5]. The median filter, which is a nonlinear filter, is effective in reducing noise without losing important image details by replacing the processed pixel value with the median value of a group of pixels. Image segmentation is also important to separate important elements from the background, thereby increasing focus on Braille patterns [6].

However, the performance of Braille detection models can be affected by noise in images. Therefore, preprocessing techniques such as median filtering and image segmentation become crucial [4]. Median filtering is effective in reducing noise without destroying important details, while segmentation helps separate Braille patterns from irrelevant backgrounds. Previous studies have shown that the combination of median filtering and segmentation can improve detection accuracy [5]. This study will explore the use of the Inception V3 architecture in Braille character recognition. Inception V3, as one of the effective CNN architectures in complex object recognition, can be enhanced by the integration of median filtering and segmentation to further improve the accuracy of Braille characters.

2 Related Work

In several studies, researchers have made great progress in developing BCR systems. To improve the accuracy and efficiency of these systems, they have used various techniques, such as machine learning, artificial neural networks (ANNs), and CNNs. These techniques allow Braille readers to interact more easily with digital devices, such as computers and smartphones.

In previous studies, such as [6], this study proposes a new approach for automatic recognition of Braille characters, which consists of two main stages: image alignment and enhancement using preprocessing techniques, and character recognition using a lightweight convolutional neural network (CNN). This approach replaces some modules in CNN with IRB blocks to reduce computational costs, resulting in an efficient and accurate model. Experiments on English Braille and Chinese double-sided Braille (DSBI) datasets show prediction accuracies of 95.2% and 98.3%. This method is more robust and effective than current approaches, with a prediction time of 0.01s for English Braille images and 0.03s for DSBI.

Meanwhile, [7] discusses the conversion of Braille to English text using deep learning. In this study, 26 English Braille images are used as the dataset that has gone through the segmentation process. The proposed method converts Braille visuals to English text using convolutional neural network (CNN) models such as LeNet, VGG-16, DenseNet121, ResNet50, and Inceptionv3. Among these models, the Inceptionv3-based system showed a high prediction accuracy rate, reaching 92%. Experimental results showed that the proposed Braille character recognition method produced accurate results.

Research [2] showed that combining multiple models can improve accuracy, but this method is often limited to a specific combination of models and datasets. In recent studies, models such as DarkNet-53, GoogleNet, SqueezeNet, and DenseNet-201 were incorporated into a more general transfer learning-based ensemble approach. This group approach achieved high F1 scores (89.42%, 99.58%, and 97.11% on HBO, BC, and AB datasets), and lower error values compared to a single model, such as DarkNet-53, which only achieved an F1 score of 87.54%. In addition, this method helps the development of more efficient Braille-based assistive technologies.

3 Methods

This research methodology is designed to develop a Braille pattern detection model using the Inception V3 architecture by implementing median filter and segmentation as part of the data preprocessing process. The steps of this research methodology are explained as follows, shown in Fig. 1.

This research includes several stages, namely:

Step 1: Literature Study

Literature study was conducted from data collection to evaluation to guide the research process. The sources include books, scientific journals, and other scientific works, related to detection using Inceptio-V3, median filter, segmentation and other related topics.

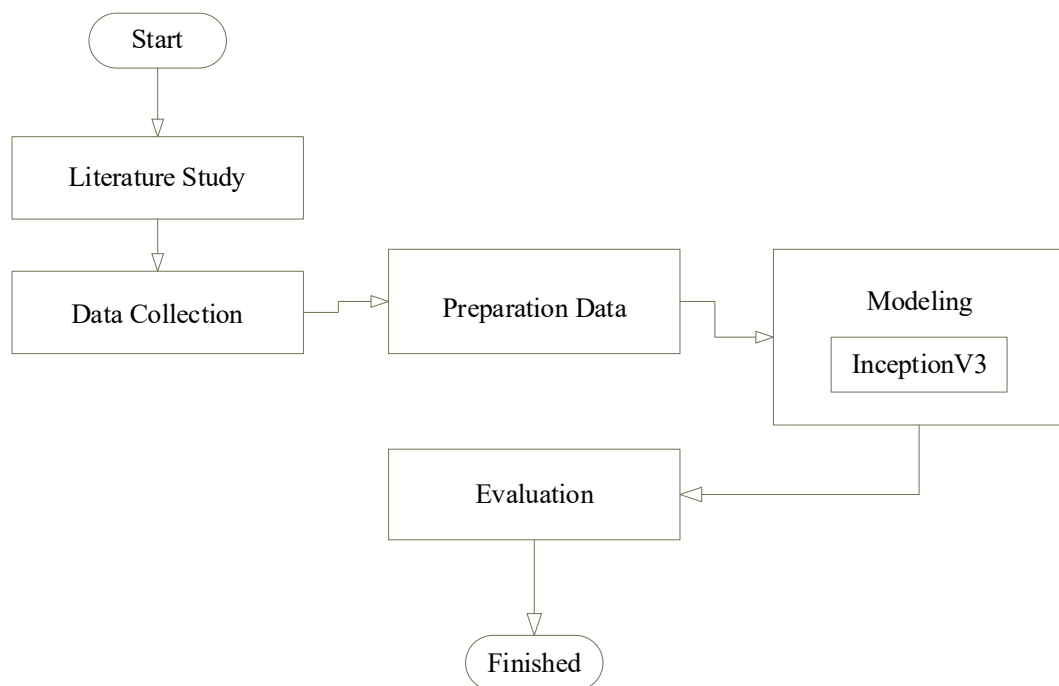


Figure 1. Research stages

Step 2: Data Collection

This study collects datasets using datasets taken from the github website belonging to user HelenGezahegn [4]. The data used has been filtered to select good quality data, so that it has been reduced and consists of 4,160 data samples, with an average of 160 samples for each character from A to Z, the image will be divided into 2 parts 3900 training data and 260 testing data, and for 3900 data divided into 3120 as training data and 780 as validation data [8].

Step 3: Data Preparation

Data preparation in this study includes several processes to ensure that the images used in the Convolutional Neural Network model using the inception V3 architecture have the appropriate quality and format. The first step taken is to resize the image. The image with the original size $H \times W$ is resized to the standard dimensions $H' \times W'$. This process can be formulated as:

$$I'(x, y) = I\left(\frac{H'}{H} \times x, \frac{W'}{W} \times y\right) \quad (1)$$

After resizing, the image is converted from color format (RGB) to grayscale using the following formula [9]:

$$Y = 0.299R + 0.587G + 0.144B \quad (2)$$

Where R , G , and B are the red, green, and blue channel values at pixel (x, y) .

To remove noise, a median filter is applied to the grayscale image. Median filtering is done by taking the median of the intensity values in the filter window $k \times k$ [10]:

$$I_{filtered}(x, y) = median\{I_{gray}(x + i, y + j)\} \quad (3)$$

for all $i, j \in [-k/2, k/2]$

The segmentation process is carried out using k-means clustering [11] with the following steps. Determine the number of k clusters. Randomly select k data points as the initial centroids. Once the initial centroids are determined, calculate the distance of each data point to the nearest centroid using the Euclidean Distance formula:

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (4)$$

Where:

$d(x, y)$ = Euclidean distance between points x and y .

x and y = Two points in n -dimensional space whose distances will be calculated.

Group members are identified based on the smallest data distance to the centroid.
Then calculate the new centroid using the centroid finding formula:

$$\text{centroid}_i = \frac{\sum_{j=1}^n x_j}{n} \quad (5)$$

Where:

Centroid = Centroid for the i -th cluster.

n = Total number of data points in the i -th cluster.

Repeat steps 2 through 4 until no more data changes cluster.

Step 4. Modeling

This study uses Convolutional Neural Network (CNN) with Inception-V3 model architecture, shown in Fig. 2. In the Inception-V3 architecture layer [13], all layers before the fully connected layer in each architecture are frozen first to maintain the weight parameters obtained from ImageNet. The fully connected layer is removed because transfer learning will be applied to the model to train a new fully connected layer that will be used to classify datasets with different numbers of categories. After freezing the layers before the fully connected layer, there is a flatten layer followed by two fully connected layers [14]. Braille pattern detection model architecture using InceptionV3 is shown in Fig. 3.

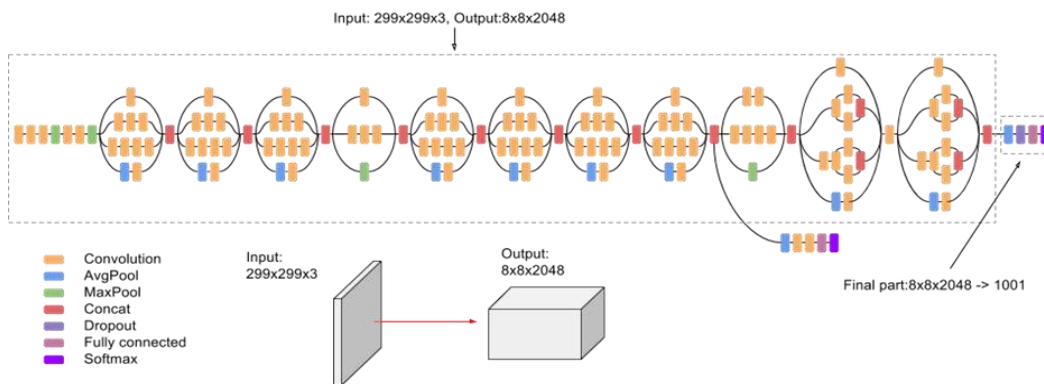


Figure 2. Inception V3 Architecture

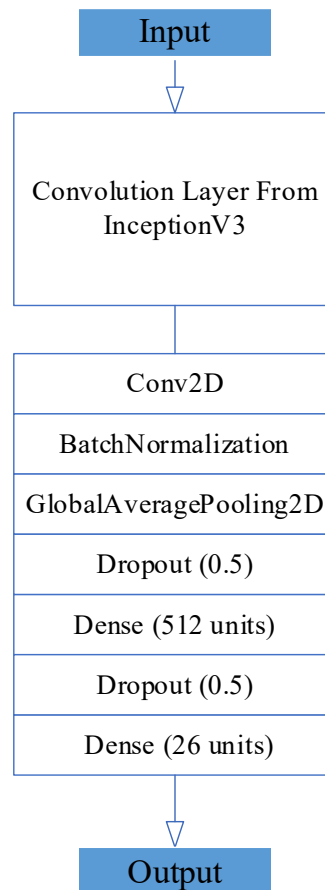


Figure 3. Braille Pattern Detection Model Architecture Using InceptionV3

Fig. 2 is a convolutional neural network (CNN) model architecture that uses convolution layers from InceptionV3 for feature extraction from input images. After that, Conv2D is applied to the image or spatial data. followed by BatchNormalization to stabilize and speed up training. And GlobalAveragePooling2D, which converts the convolution output into a feature vector. Next, a dropout layer is applied to reduce overfitting, followed by several dense layers with 512 neuron units. Then there is another Dropout layer and finally, a thick layer with 64 and 26 neuron units. 26 neuron units are used for classification into 26 output classes. This architecture combines the powerful feature extraction capabilities of InceptionV3 with regularization and normalization techniques to improve model performance and generalization.

Step 5. Evaluation

At this stage, various scenarios of Recall, Precision, F1 Score, and Accuracy settings are evaluated to improve the modeling results [17]. Finally, it will be known which scenario produces the best accuracy value and the lowest error rate for the braille pattern detection model using the Inception V3 architecture by applying median filter and segmentation. The formulas for Recall, Precision, F1 Score, and Accuracy are as follows:

$$\text{Akurasi} = \frac{TP + TN}{TP + FP + FN + TN} \quad (6)$$

$$\text{Presisi} = \frac{TP}{TP + FP} \quad (7)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (8)$$

$$\text{F1 Score} = 2 \times \frac{\text{Recall} \times \text{Presisi}}{\text{Recall} + \text{Presisi}} \quad (9)$$

Explanation of the equations TP (True Positive), TN (True Negative), FP (False Positive), and FN (False Negative).

4 Results and Discussions

In this study, the InceptionV3 model pre-trained on the ImageNet dataset was modified by adding new layers. The model was then optimized by specifying an image size of 299x299 [18], a batch size of 32, a number of epochs of 100, and a learning rate of 0.001 and 0.0001. Early stopping was also used to prevent overfitting. These parameters are important to ensure the model learns well from the data and achieves rapid convergence.

In the training phase, the InceptionV3 model implemented with median filter and segmentation showed quite good performance. During the training process, accuracy and loss metrics were observed to evaluate the model performance. The training and validation accuracy graphs showed a significant increase at the beginning of the epoch and stabilization at the end of the training process as shown in Fig. 4.

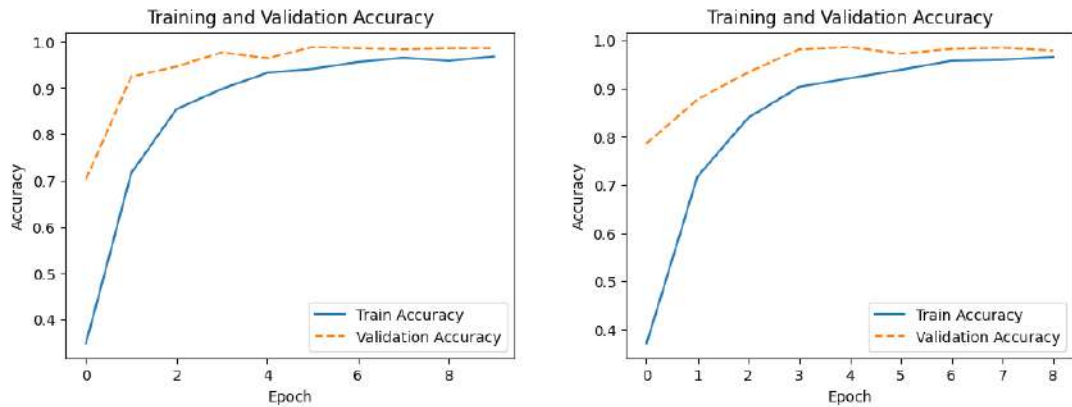


Figure 4. Accuracy graph trains the model with learning rate 0.001

The Fig. 4 above show the results of the training process with a learning rate of 0.001, presented in the form of graphs. In both graphs, dots represent the training data, while lines represent the validation data. The graph on the right demonstrates that the training accuracy reaches above 0.9664, with the validation accuracy showing similar results. Similarly, the graph on the left shows that both the training and validation accuracy exceed 0.9669.

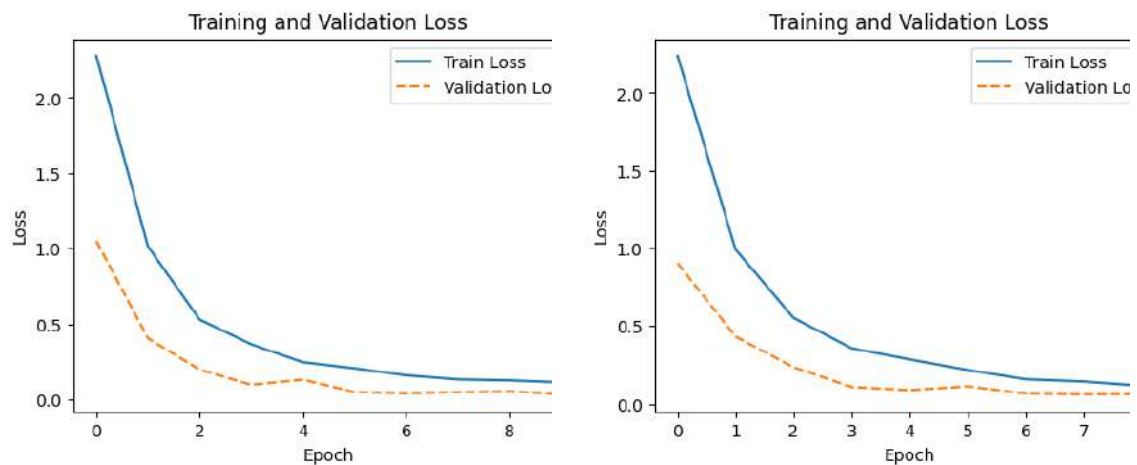


Figure 5. Loss graph of training the model with a learning rate 0.001

The Fig. 5 depict the results of the training process with a learning rate of 0.001, shown in graph form. In both graphs, dots represent the training data, and lines represent the validation data. The graph on the right shows that the training loss reached 0.1144. Similarly, the graph on the left indicates that the training loss reached 0.1289, then training was also carried out using a learning rate of 0.0001.

The Fig. 6 illustrate the results of the training process, presented as graphs where dots represent training data and lines represent validation data. The graph on the right shows that both training and validation accuracy exceed 0.9779. Similarly, the graph on the left indicates that the training accuracy reaches 0.9822.

The Fig. 7 depict the results of the training process with a learning rate of 0.001, shown in graph form. In both graphs, dots represent the training data, and lines represent the validation data. The graph on the right shows that the training loss reached 0.1400. Similarly, the graph on the left indicates that the training loss reached 0.1497.

After the training process is complete, the next step is to test the model. At this stage, the trained data is compared with the data that has been prepared during the preprocessing process. The dataset used for testing is 260 datasets.

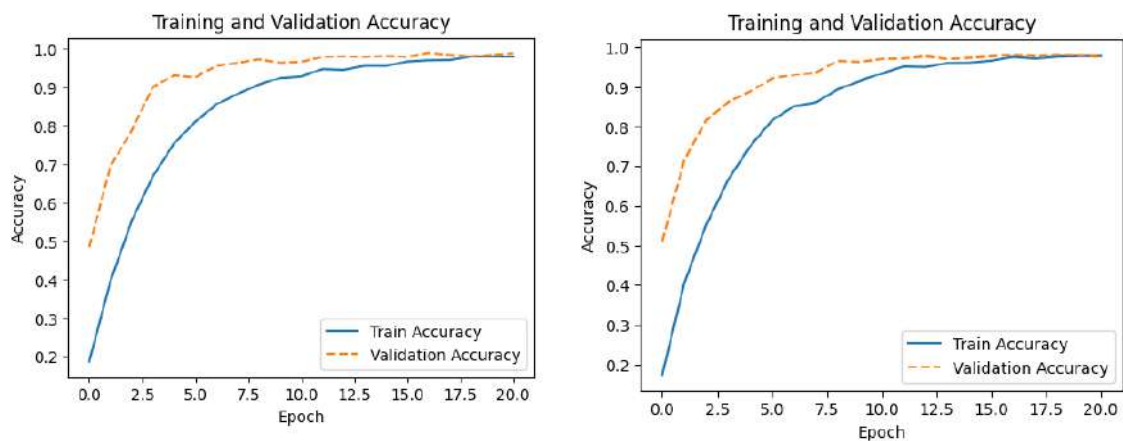


Figure 6 Accuracy graph train the model with learning rate 0.0001

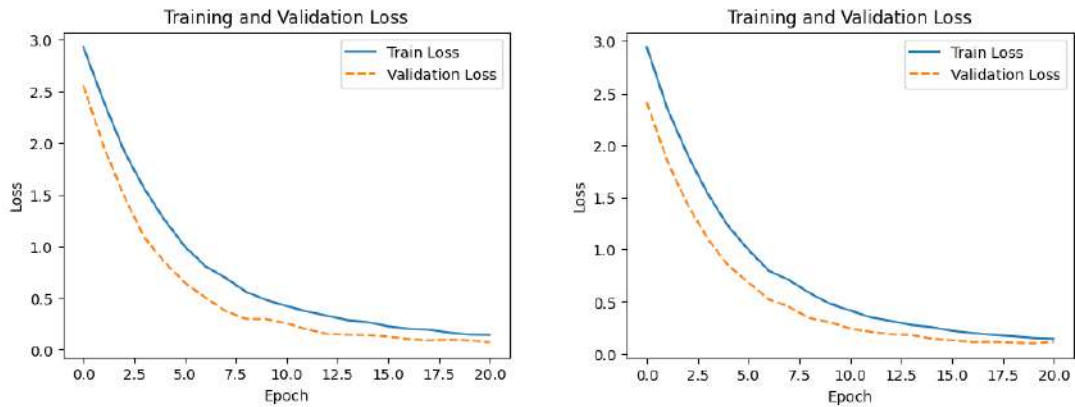


Figure 7. loss graph of training the model with learning rate 0.0001

Table 1. Performance evaluation metrics

Learning rate	Accuracy	precision	recall	f1 Score
Using Median Filter and Segmentation				
0.001	0.9965	0.9962	0.9961	0.9962
0.0001	0.9930	0.9923	0.9923	0.9923
Not Using Median Filter and Segmentation				
0.001	0.9936	0.9923	0.9924	0.9923
0.0001	0.9865	0.9846	0.9845	0.9846

Based on Table 1, it can be seen that the use of median filter and segmentation results in higher accuracy, precision, recall, and F1 score compared to not using both techniques, especially at a learning rate of 0.001. At this learning rate, the highest accuracy is achieved, namely 99.65% with precision and recall of 99.62% and 99.61%, respectively. In contrast, without median filter and segmentation, model performance decreases, especially at a learning rate of 0.0001, where accuracy drops to 98.65%, and precision and recall reach 98.46% and 98.45%. This shows that median filter and segmentation play an important role in improving model performance, especially at higher learning rates.

Table 2. Comparison table with previous research

Research	Methods	Accuracy
[1]	EfficientNetV2M and InceptionV3	82.07%
[2],[4]	Cnvolutional Neural Network (CNN).	81.54%
[3]	CNN using Segementation	95.77%
Our Reserach	InceptionV3 using Median Filter and Segmentation	99.65%

Table 2 presents a comparison between the proposed method and previous research in \terms of classification accuracy. The first referenced study employed EfficientNetV2M and InceptionV3 architectures, achieving an accuracy of 82.07%. The second study utilized a Convolutional Neural Network (CNN) model with image segmentation, resulting in an accuracy of 95.77%. In contrast, our proposed approach, which combines InceptionV3 with median filtering and Segmentation preprocessing, outperformed the others with a significantly higher accuracy of 99.65%. This demonstrates the effectiveness of the proposed method in enhancing classification performance.

5 Conclusions

The InceptionV3 model shows excellent performance with high training accuracy on all tested configurations, in this study the use of median filter and segmentation produces higher accuracy, precision, recall, and F1 scores compared to when not using both techniques, especially at a learning rate of 0.001, where the accuracy reaches 99.65%, and precision and recall reach 99.62% and 99.61%, respectively. Conversely, without median filter and segmentation, the model performance decreases, especially at a learning rate of 0.0001, where the accuracy reaches 99.65%, and precision reaches 99.62%.

Acknowledgements

The author would like to express his deepest gratitude to Siti Yuliyanti and Muhammad Al-Husaini, as academic supervisors, for their invaluable guidance and support during this research. Gratitude is also extended to the Informatics Study Program, Siliwangi University, which has facilitated the research environment. Special gratitude to

all colleagues and friends who have contributed in the form of corrections, feedback, and moral support during the preparation of this work.

References

- [1] K. M. Farrand, K. Koehler, and A. Vasquez, “Literary Braille Instruction: A Review of University Personnel Preparation Programs,” <https://doi.org/10.1177/0145482X221130356>, vol. 116, no. 5, pp. 617–628, Nov. 2022, doi: 10.1177/0145482X221130356.
- [2] N. Elaraby, S. Barakat, and A. Rezk, “A generalized ensemble approach based on transfer learning for Braille character recognition,” *Inf Process Manag*, vol. 61, no. 1, p. 103545, Jan. 2024, doi: 10.1016/J.IPM.2023.103545.
- [3] Jahandad, S. M. Sam, K. Kamardin, N. N. Amir Sjarif, and N. Mohamed, “Offline Signature Verification using Deep Learning Convolutional Neural Network (CNN) Architectures GoogLeNet Inception-v1 and Inception-v3,” *Procedia Comput Sci*, vol. 161, pp. 475–483, Jan. 2019, doi: 10.1016/J.PROCS.2019.11.147.
- [4] M. F. Herlambang, A. N. Hermana, and K. R. Putra, “Pengenalan Karakter Huruf Braille dengan Metode Convolutional Neural Network,” *Systemic: Information System and Informatics Journal*, vol. 6, no. 2, pp. 20–26, 2021, doi: 10.29080/systemic.v6i2.969.
- [5] H. Tang, R. Ni, Y. Zhao, and X. Li, “Median filtering detection of small-size image based on CNN,” *J Vis Commun Image Represent*, vol. 51, pp. 162–168, Feb. 2018, doi: 10.1016/J.JVCIR.2018.01.011.
- [6] R. H. Hasan, I. S. Aboud, and R. M. Hassoo, “Braille Character Recognition System : Review,” vol. 18, no. 1, pp. 30–39, 2024.

- [7] A. Al-Salman and A. AlSalman, “Fly-LeNet: A deep learning-based framework for converting multilingual braille images,” *Heliyon*, vol. 10, no. 4, p. e26155, Feb. 2024, doi: 10.1016/J.HELIYON.2024.E26155.
- [8] H. J. Andriansyah, Marindo, “Pengenalan Karakter Braille Memanfaatkan Convolutional Neural Network,” vol. 4, no. 1, pp. 41–54, 2021.
- [9] M. Minarni, K. Rizka, Y. Yuhendra, D. W. T. Putra, and I. Warman, “Penerapan Deteksi Sobel Berbasis Algoritma Backpropagation pada Pengenalan Pola Huruf Vokal,” *Jurnal Minfo Polgan*, vol. 12, no. 2, pp. 1829–1839, 2023, doi: 10.33395/jmp.v12i2.13019.
- [10] Mardiah et al, “Application of Image Median Filter and Histogram Equalization on Old Building Images Pendahuluan,” vol. 1, no. September, pp. 0–7, 2023.
- [11] M. A. Rajab and L. E. George, “Stamps extraction using local adaptive k- means and ISODATA algorithms,” *Indonesian Journal of Electrical Engineering and Computer Science*, vol. 21, no. 1, pp. 137–145, 2021, doi: 10.11591/ijeecs.v21.i1.pp137-145.
- [12] A. C. O. Reddy and K. Madhavi, “Hierarchy based firefly optimized k-means clustering for complex question answering,” *Indonesian Journal of Electrical Engineering and Computer Science*, vol. 17, no. 1, pp. 264–272, 2019, doi: 10.11591/ijeecs.v17.i1.pp264-272.
- [13] Jahandad, S. M. Sam, K. Kamardin, N. N. Amir Sjarif, and N. Mohamed, “Offline Signature Verification using Deep Learning Convolutional Neural Network (CNN) Architectures GoogLeNet Inception-v1 and Inception-v3,” *Procedia Comput Sci*, vol. 161, pp. 475–483, Jan. 2019, doi: 10.1016/J.PROCS.2019.11.147.

- [14] S. Raghuwanshi, A. Sukhad, A. Rasool, V. K. Meena, A. Jadhav, and K. Shivakarthik, "Early Detection of Brain Tumor from MRI Images Using Different Machine Learning Techniques," *Procedia Comput Sci*, vol. 235, pp. 3094–3104, Jan. 2024, doi: 10.1016/J.PROCS.2024.04.293.
- [15] Maad M. Mijwel, "Artificial Neural Networks Advantages and Disadvantages," *Mesopotamian journal of BigData*, vol. Vol. (2021, pp. 29–31, 2021.
- [16] S. Samidin and A. Fadjeri, "Klasifikasi Gambar Batu-Kertas-Gunting Menggunakan Convolutional Neural Network dengan Fungsi Callback untuk Mencegah Overfitting," *Jurnal Penelitian Inovatif*, vol. 4, no. 2, pp. 785–794, 2024, doi: 10.54082/jupin.413.
- [17] N. Saqib, K. F. Haque, V. P. Yanambaka, and A. Abdelgawad, "Convolutional-Neural-Network-Based Handwritten Character Recognition: An Approach with Massive Multisource Data," *Algorithms*, vol. 15, no. 4, 2022, doi: 10.3390/a15040129.
- [18] V. C. P. A. M. L. S. Shilpi Yadav, "Recognition of Pandemic virus Covid-19 using Inception-V3," *International Journal of Advanced Science and Technology*, vol. 29, no. 12s, pp. 1675–1679, Jun. 2020, Accessed: Sep. 10, 2024. [Online]. Available: <http://sersc.org/journals/index.php/IJAST/article/view/23906>

This page intentionally left

Sign Language Detection Models using Resnet-34 and Augmentation Techniques

Rizki Ramdhan Hilal¹, Aradea¹, Vega Purwayoga^{1*}

¹*Department of Informatics, Faculty of Engineering,
Siliwangi University, Indonesia*

**Corresponding Author: vega.purwayoga@unsil.ac.id*

(Received 26-06-2025; Revised 31-07-2025; Accepted 25-08-2025)

Abstract

For deaf or hard of hearing people, sign language is a primary means of communication, but low public understanding makes social engagement difficult. Researchers now use computer vision technology and Convolutional Neural Network (CNN) to detect sign language movements. Problems such as overfitting and missing gradients still exist. Using CNN and ResNet-34 architecture, as well as image augmentation to overcome this problem, this research builds a deep learning-based sign language detection model. The Indonesian Sign Language System (SIBI) dataset was used to test the model. The test results show that the model with image augmentation trained for more than 50 epochs obtained an accuracy of 99.4%, precision of 99.5%, recall of 99.5%, and an F1 score of 99.5%. The model without image augmentation produced an accuracy of 99.4%, recall of 99.3%, F1 score of 99.3%, and precision of 99.4%. ResNet-34 architecture overcomes the problem of missing gradients, while image augmentation avoids overfitting and improves model accuracy.

Keywords: Augmentation, Convolutional Neural Network, Overfitting, ResNet-34, Vanishing Gradient

1 Introduction

Humans utilize language as a tool to communicate with one another. Language can take many different forms, including signs, symbols, codes, and noises that are given meaning after being converted into human language according to predetermined rules [1], [2]. Using hand movements that follow the PUEBI Guidelines, people with hearing and speech disabilities can communicate through sign language [3].

The issue with using sign language as a communication tool is that not everyone fully understands this system of communication because there is a dearth of knowledge and resources about learning sign language, including books, courses, teachers, and other resources that can be a barrier for those who wish to take sign language classes [4]. Using

deep learning technology to create a photo image identification model that can assist in sign language translation is one way to address this issue. The aid of deep learning technology a model that is capable of self-learning computational procedures [5], to help those who are unable to communicate through sign language, this technology is intended to recognize hand motions and translate them into a language that they can understand sign language [4]. In earlier studies carried out by [6], Convolutional Neural Network (CNN) were used to create a sign language categorization system, and the accuracy rate was 99.82%.

Increasing the depth of the network in deep convolutional neural networks does not always lead to improved training accuracy, there are instances when training accuracy decreases [7], [8], [9], [10]. This is the vanishing gradients problem, which is a common impediment in CNN. This is because not all networks are simple to optimize, which leads to network degradation. Using residual networks, or ResNet, to add identity mapping will lessen degradation in deeper networks ResNet [11]. Previous research on the application of the ResNet architecture by [9], [12], [13], [14], [15] has successfully overcome the vanishing gradient problem in the CNN algorithm.

The possibility of overfitting arises from CNN models' memorization of specific, non-generalizable details of training images, which is another issue frequently associated with CNN algorithms digeneralisasi [16], [17], [18], [19], [20]. Therefore, by adjusting the dimensional transformation of the photos, image augmentation techniques were applied to the training samples in this study to increase the variety of images [21]. In earlier studies by [22], [23], [24], [25] on the application of augmentation techniques on datasets, has effectively addressed CNN models overfitting issue.

The primary research gaps in this study are related to the CNN algorithm's potential for vanishing gradient and overfitting in the final model, as indicated by the research references pertaining to the development of digital image detection models using this algorithm that were mentioned in the previous paragraph. Therefore, the goal of this research is to use the ResNet-34 architecture to create a detection model that can solve the disappearing gradients problem [26] because it may be applied to validation data and has a low error rate value, to produce the best accuracy outcomes, and enhancing the dataset with augmentation methods to prevent the model from becoming overfit along

with incorporating dataset augmentation methods to prevent the model from becoming overfit [21], [22], [23].

The contribution of this research is to develop a ResNet-34 model combined with the application of augmentation techniques on the dataset to overcome the problems of overfitting and vanishing gradients in the model in the process of detecting the Indonesian Sign Language System (SIBI). Compared to other studies, which are limited in their use of model architecture and dependent on limited datasets, this study provides evaluation and improvement using deeper residual networks.

2 Material and Methods

The steps or procedures utilized in research are referred to as the research stages [27]. To guarantee that every step of the research is conducted in an orderly fashion, this stage offers direction and arranges every step from start to finish. The phases of the research that will be done are based on earlier studies that [19], the research stages is shown in Fig. 1.

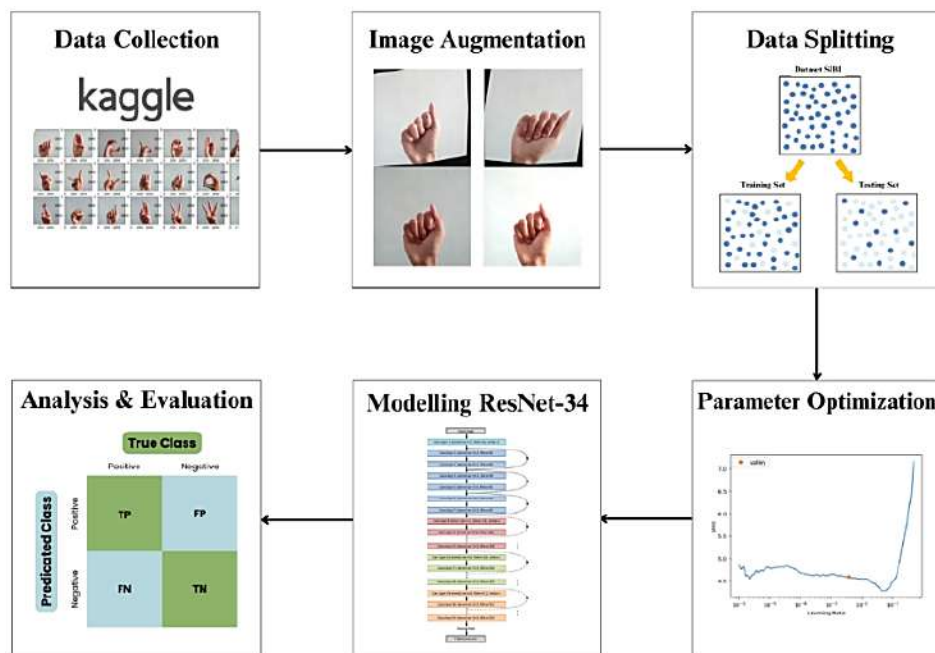


Figure 1. Research Stages

2.1 Data Collection

The collected dataset consists of digital photos of Indonesian sign language systems that were published to Kaggle in 2022 by Alvin Bintang. The dataset consists of 5280 images in total, grouped into 24 alphabetic classes (A–Y, except J and Z). There are 220 Sistem Isyarat Bahasa Indonesia (SIBI) photos per class.

2.2 Data Augmentation

The original dataset used in this study consisted of 5280 images, with 220 images in each class. To increase the variety of the dataset, augmentation techniques were applied. The augmentation techniques are based on earlier studies by [25] and [28] that employed rotation, shear horizontal and vertical, grayscale, saturation, brightness, and exposure. The arrangement of the augmentation values is shown in Table 1.

Based on the augmentation process in Table 1, the number of datasets increased to 7582 with the distribution of each class increasing by 96 images, so that the total number of datasets in each class was 316 images after the augmentation process. The augmentation technique was applied evenly to each class with the aim of maintaining class balance and avoiding class imbalance issues.

Table 1. Augmentation Technique

No.	Techniques	Description
1	Rotation	The rotation technique used is -10° and $+10^\circ$.
2	Shear horizontal & vertical	The horizontal and vertical shear techniques used are $\pm 10^\circ$ horizontal and $\pm 10^\circ$ vertical.
3	Grayscale	Grayscale technique used is 15% of the total dataset.
4	Saturation	Saturation technique used is -30% and +30%.
5	Brightness	The brightness technique used is -15% and +15%.
6	Exposure	The exposure technique used is -15% and +15%.

2.3 Data Splitting

Following the preprocessing phase, datasets are separated into training and test subsets. The `split_indices` function from the FastAI library is used to divide the dataset. The technique is based on earlier research by [29] utilizing a ratio of 80% for training data and 20% for test data.

2.4 Modelling ResNet-34

ResNet-34 architecture, which has several convolution layers and residue layers, is used for modeling in Fig. 2. The method of entering a 224 by 224 pixel picture is the first step. With an output of 128 x 128 pixels, the input image is passed through 64 filters into the first convolution layer in the second stage. Furthermore, there are four residual layers with the following output sizes 28 x 28 pixels with 4 residual blocks using 128 filters, 14 x 14 pixels with 6 residual blocks using 256 filters, and 7 x 7 pixels with 3 residual blocks using 512 filters. The size of the residual layers is 56 x 56 pixels with 3 residual blocks using 64 filters. After the residual layer, an average pooling layer reduces the output size to 1 x 1 pixels, followed by a fully connected layer with 256 and 128 neurons. The final stage is the softmax classifier layer which produces a classification with accuracy ($E = 0.99$). Each residual layer has skip connections that allow direct information flow, helping to overcome the vanishing gradient problem [12], [30].

2.5 Model Evaluation

Model evaluation is crucial for identifying the optimal model combination. It involves using matrices, considering factors like accuracy and the confusion matrix. The confusion matrix serves as a visual evaluation tool in machine learning [31].

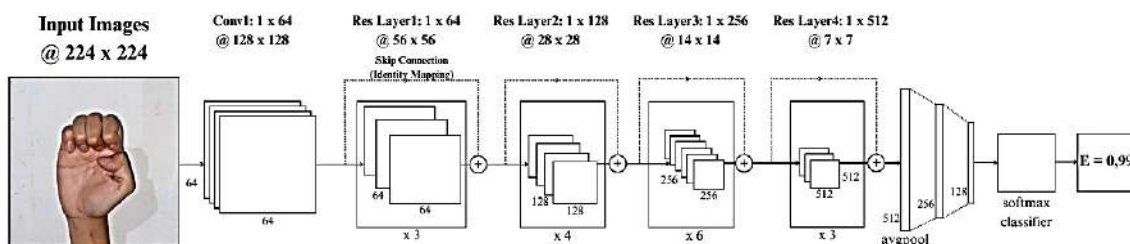


Figure 2. ResNet-34 Architecture

Its columns represent predicted class outcomes, while its rows depict actual class outcomes, facilitating the examination of all potential cases in classification problems [28]. Various metrics such as precision, recall, and F1 score are utilized within the confusion matrix. This matrix encompasses 4 key terms [32]:

- a. True Positive (TP): correctly classified positive data.
- b. True Negative (TN): the number of correctly classified negative data.
- c. False Positive (FP): the number of negative data classified as positive.
- d. False Negative (FN): the number of positive data classified as negative.

Accuracy and the confusion matrix can be formulated as shown in Equation (1)-(4) [31].

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

Accuracy is the ratio of correct predictions to the total data.

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

Precision is the ratio of true positive predictions to the total number of positive predictions.

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

Recall is the ratio of true positive predictions to the total number of actual positive instances.

$$F1\ Score = \frac{2 \times Recall \times Presisi}{Recall + Presisi} \quad (4)$$

The F1 score is the weighted average of precision and recall.

3 Results and Discussions

3.1 Data Augmentation

This method generates a final dataset of 7582 datasets by using 12 picture data from each class. The augmentation techniques are expected to test the model's robustness against diverse data. An example of the application of augmentation techniques can be seen in Fig. 3.

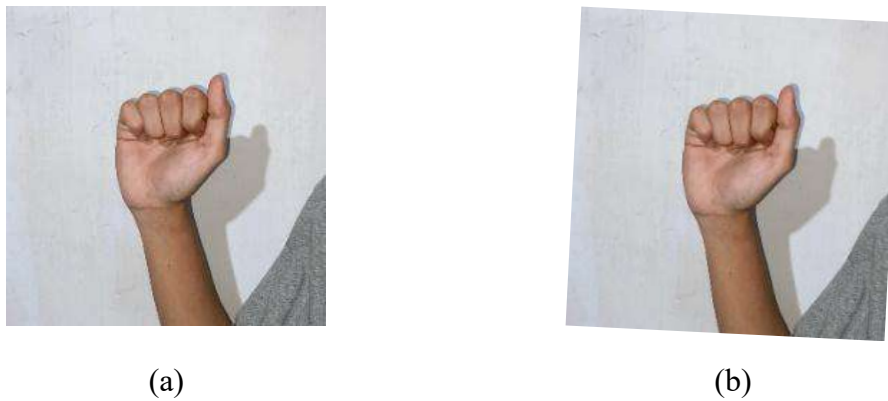


Figure 3. (a) Rotation augmentation technique, (b) Rotation augmentation technique.

3.2 Data Splitting

Data utilized in the model training phase is referred to as training data or train data. Up to 80% of the entire dataset was randomly selected for the training data utilized in this study, yielding a training dataset of 6056 data in total. Data used in the models testing phase is called test data or testing data. With a final test dataset consisting of 1526 data, the test data used in this study accounted for up to 20% of the entire dataset.

3.3 Modelling ResNet-34

Before entering the modeling stage, the model is subjected to hyperparameter tuning by considering the most optimal learning rate. Analyzing the model for vanishing gradient and overfitting is done by comparing the visualization when tuning the model by comparing the values of train loss, valid loss and error rate.

This research uses the one-cycle policy optimization method developed by Leslie Smith in 2018. By modifying the learning rate to hasten convergence and avoid overfitting, this technique seeks to train the model effectively [33]. The best weight decay value is first $1e-6$ (0.000001) and eventually $1e1$ (10). Fig 4 illustrates how to use the learning rate finder approach to get the ideal learning rate. The process of finding the optimal learning rate using the learning rate finder method is shown in Fig. 4.

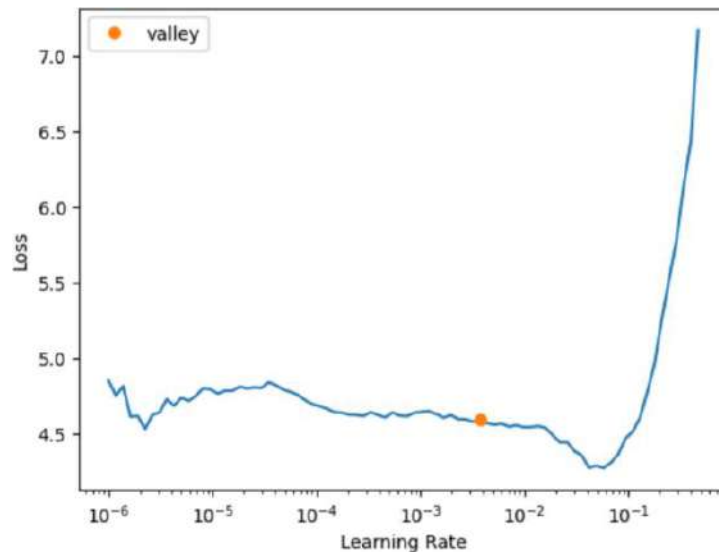


Figure 4. The Most Optimal Learning Rate Value

The Learning Rate Finder produced the curve shown in Fig. 3, where loss is shown on the y-axis and learning rate is plotted on the x-axis. The “valley” point on the loss curve denotes the ideal learning rate and is the lowest point. The optimal point shown as “valley” in the figure, is used to determine the lr_max value. Based on Fig. 3. Lr_max value of about 1.13×10^{-3} was chosen, as it shows a sharp increase in loss after the optimal point, indicating it is the best before the loss increases again.

Visualizations of train loss, valid loss and error rate at 25, 50 and 100 epochs are shown in Fig. 5 shows the training using 25 epochs, (a) without augmentation (b) with augmentation, Fig. 6 shows the training using 50 epochs, (a) without augmentation (b) with augmentation, Fig. 7 shows the training using 100 epochs, (a) without augmentation (b) with augmentation.

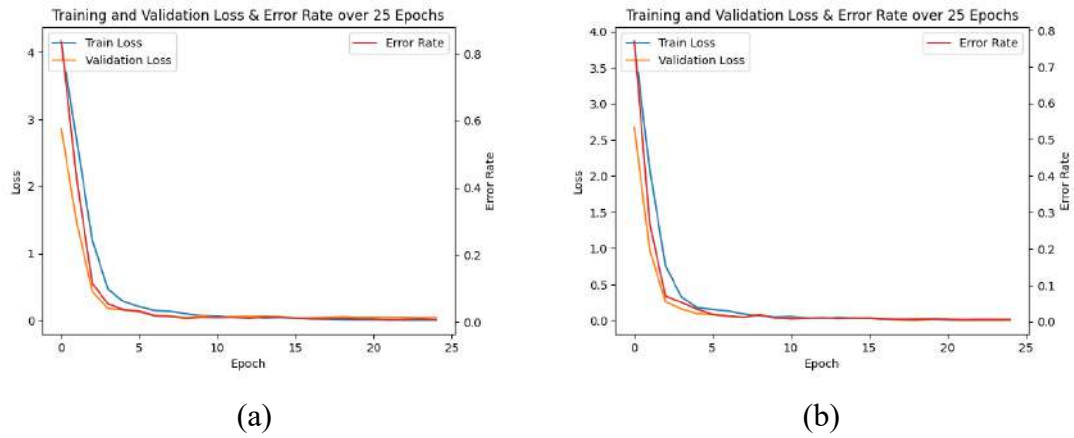


Figure 5. Training and validation loss and error rate at 25 epoch for (a) without augmentation and (b) with augmentation

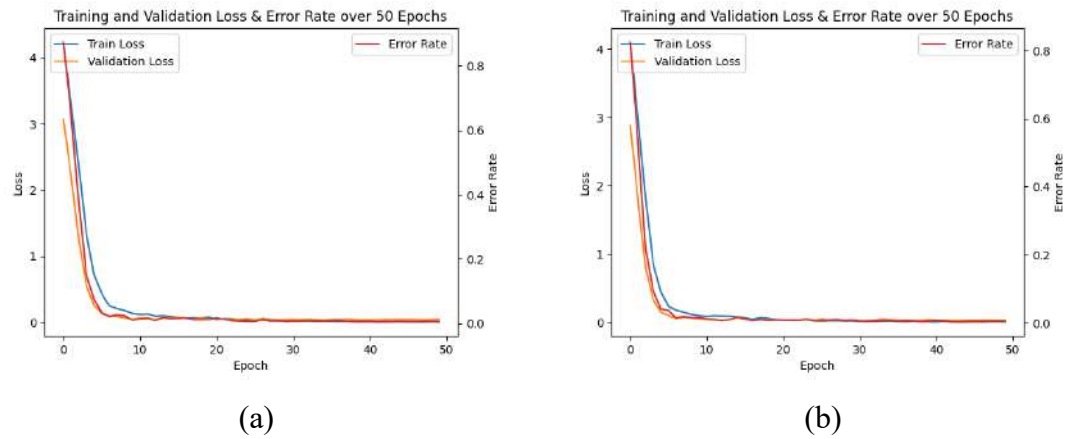


Figure 6. Training and validation loss and error rate at 50 epoch for (a) without augmentation and (b) with augmentation

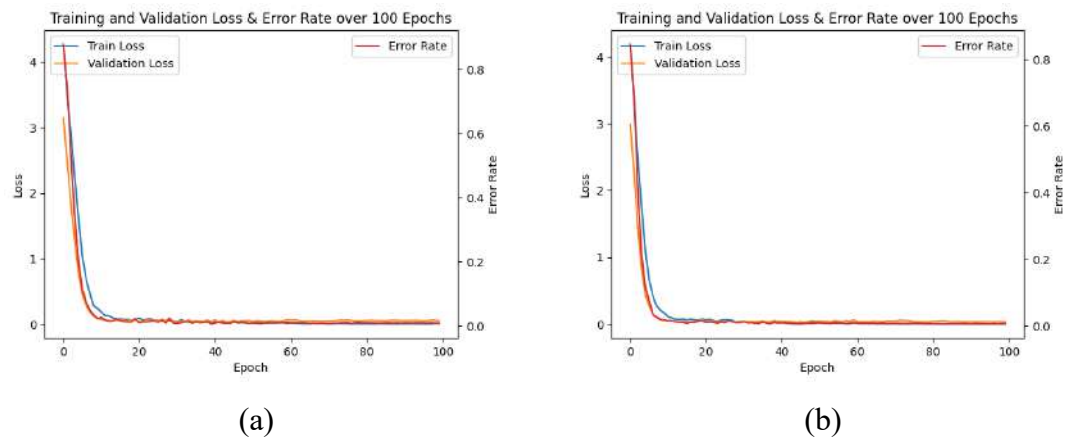


Figure 7. Training and validation loss and error rate at 100 epoch for (a) without augmentation and (b) with augmentation

Comparison of visualization train loss, valid loss, and error rate at 25, 50, and 100 epochs is shown in Table 2. From the results of the comparison value in Table 2. It can be concluded that the model using augmentation techniques at 50 epochs is the best performing model by looking at the results of train loss and valid loss which tend to be balanced, indicating that the model can generalize new data well and the model also does not occur overfitting.

3.4 Model Evaluation

To determine the ideal model combination, model evaluation is required. A matrix can be used to evaluate a model by taking into account many factors, including accuracy and confusion matrix findings. A visual assessment tool for machine learning is the confusion matrix [31]. To calculate all potential cases of classification difficulties, the confusion matrix's columns reflect the predicted class results, while its rows represent the actual class results [34]. The calculation results and Comparison of precision, recall, f1-score, and accuracy is shown using confusion matrix are shown in Fig. 8.

Table 2. Comparison of Visualization Result Values

Variable	Without augmen. (25 epoch)	With augmen. (25 epoch)	Without augmen. (50 epoch)	With augmen. (50 epoch)	Without augmen. (100 epoch)	With augmen. (100 epoch)
Train loss	0.008626	0.007812	0.002607	0.002599	0.008523	0.006596
Valid loss	0.048414	0.014932	0.045147	0.021599	0.054004	0.033828
Error rate	0.006629	0.004617	0.006629	0.004617	0.005120	0.000175

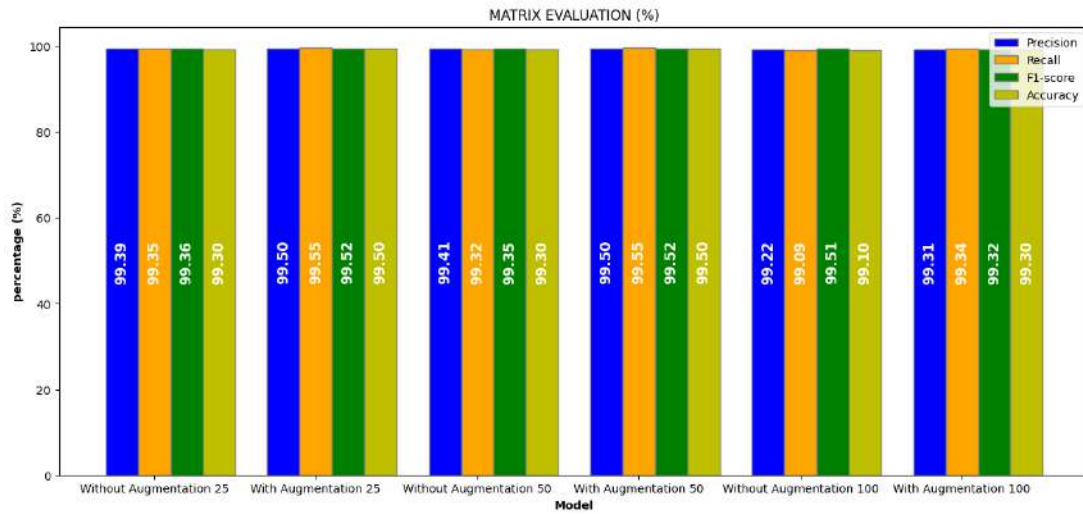
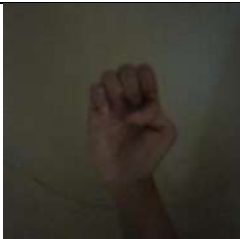
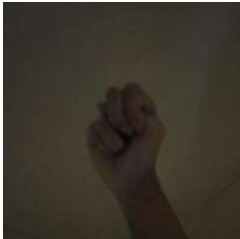


Figure 8. Visualization of Confusion Matrix

3.5 Error Analysis

Error analysis was conducted to identify limitations in the classification model. This analysis aimed to identify factors causing the model to fail in making predictions, particularly for letters of the alphabet that are similar in shape. Examples of images that were incorrectly classified by the model are shown in Table 3.

Table 3. Misclassified Alphabet Images

No.	Techniques	Actual Label	Predicted Label
1		E	S
2		N	S

3



C

X

Table 3 shows cases where the model failed to classify image objects. Most classification errors occurred because identical alphabet letters in the images indicated that the model tended to make mistakes when distinguishing between similar object structures. Minimal lighting on objects also affected the generalisation of the model's classification results.

The results of this analysis indicate that the model still struggles to generalise under low-light conditions, with objects that have high noise levels, and characters with similar shapes. To address these shortcomings, the model can be improved by adding more representative training data for alphabet classes with identical shapes and enhancing the image quality of the dataset.

3.6 Benchmarking with Previous Studies

The model evaluation was conducted by comparing the performance of the developed ResNet-34 model with the results of other relevant research models. The model evaluation is shown in Table 4.

Table 4. Benchmarking Model Results

Model / Study	Dataset Type	Architecture	Accuracy	Augmentation	Year
Arisandi et al. 2022 [6]	BISINDO	CNN 5-layer	93.00%	No	2022
Niswati et al. [11]	Cervical cancer image	ResNet-50	91.00%	No	2022
Ridhovan et al. 2022 [31]	Leaf disease	ResNet-152V2	95.00%	Yes	2022

This study	SIBI (24 classes)	ResNet-34	99.50%	Yes	2025
-------------------	-------------------	-----------	---------------	-----	------

Based on the model comparison results in Table 4, the developed model shows competitive accuracy compared to previous studies. These comparison results show that the use of augmentation techniques on the dataset contributes significantly to improving model performance.

4 Conclusions

The best model performance CNN was found in the 50 epoch process using augmentation techniques with results precision 99.5%, recall 99.5%, F1-score 99.5%, and accuracy 99.5%. From these results it can be concluded that the ResNet-34 architecture used by the CNN algorithm effectively prevents vanishing gradient, and image augmentation approaches effectively prevent overfitting and increase the accuracy of the final model. The system could be a communication aid for the general public, especially those with speech and hearing impairments, provided it undergoes more real-world testing.

Acknowledgements

The author would like to thanks Allah SWT, Siliwangi University, family, friends and lectures gave me provided guidance and motivation in carrying out this research.

References

- [1] Saleha and M. R. Yuwita, A. [Nama Penulis], "Semiotic Analysis of Dead End Traffic Sign Symbols," *MAHADAYA: Jurnal Ilmu Komunikasi*, vol. 3, no. 1, pp. 65-72, [2023]. [Online]. Available: <https://ojs.unikom.ac.id/index.php/mahadaya/article/view/7886>
- [2] F. S. Pandiangan and M. Rosadi, "Analisis Dialek Dalam Bentuk Bahasa Percakapan Dalam Film 'Imperfect' Karya Meira Anastasia," *Journal of*

-
- Educational Research and Humaniora (JERH)*, vol. 1, no. September, pp. 47–58, 2023.
- [3] Nasha Hikmatia A.E. and M. I. Zul, “Aplikasi Penerjemah Bahasa Isyarat Indonesia menjadi Suara berbasis Android menggunakan Tensorflow,” *Jurnal Komputer Terapan*, vol. 7, no. 1, pp. 74–83, 2021, doi: 10.35143/jkt.v7i1.4629.
 - [4] I. Sari, Fivrenodi, E. Altiarika, and Sarwindah, “Sistem Pengembangan Bahasa Isyarat Untuk Berkomunikasi dengan Penyandang Disabilitas (Tunarungu),” *Journal of Information Technology and society*, vol. 1, no. 1, pp. 20–25, 2023, doi: 10.35438/jits.v1i1.21.
 - [5] Nofal Anam, “Sistem Deteksi Simbol Pada Sibi (Sistem Isyarat Bahasa Indonesia) Menggunakan Mediapipe Dan ResNet-50,” 2022.
 - [6] L. Arisandi and B. Satya, “Sistem Klarifikasi Bahasa Isyarat Indonesia (Bisindo) Dengan Menggunakan Algoritma Convolutional Neural Network,” *Jurnal Sistem Cerdas*, vol. 5, no. 3, pp. 135–146, 2022, doi: 10.37396/jsc.v5i3.262.
 - [7] G. Latif, D. A. Alghmgham, R. Maheswar, J. Alghazo, F. Sibai, and M. H. Aly, “Deep learning in Transportation: Optimized driven deep residual networks for Arabic traffic sign recognition,” *Alexandria Engineering Journal*, vol. 80, no. July 2022, pp. 134–143, 2023, doi: 10.1016/j.aej.2023.08.047.
 - [8] S. Mekruksavanich, N. Hnoohom, and A. Jitpattanakul, “A Hybrid Deep Residual Network for Efficient Transitional Activity Recognition Based on Wearable Sensors,” *Applied Sciences (Switzerland)*, vol. 12, no. 10, 2022, doi: 10.3390/app12104988.
 - [9] P. A. Pattanaik, M. Mittal, M. Z. Khan, and S. N. Panda, “Malaria detection using deep residual networks with mobile microscopy,” *Journal of King Saud University*
-

-
- *Computer and Information Sciences*, vol. 34, no. 5, pp. 1700–1705, 2022, doi: 10.1016/j.jksuci.2020.07.003.
- [10] B. Tasci, M. R. Acharya, M. Baygin, S. Dogan, T. Tuncer, and S. B. Belhaouari, “InCR: Inception and concatenation residual block-based deep learning network for damaged building detection using remote sensing images,” *International Journal of Applied Earth Observation and Geoinformation*, vol. 123, no. August, p. 103483, 2023, doi: 10.1016/j.jag.2023.103483.
- [11] Z. Niswati, R. Hardatin, M. N. Muslimah, and S. N. Hasanah, “Perbandingan Arsitektur ResNet50 dan ResNet101 dalam Klasifikasi Kanker Serviks pada Citra Pap Smear,” *Faktor Exacta*, vol. 14, no. 3, p. 160, 2021, doi: 10.30998/faktorexacta.v14i3.10010.
- [12] H. Imaduddin, F. Y. A’la, A. Fatmawati, and B. A. Hermansyah, “Comparison of transfer learning method for COVID-19 detection using convolution neural network,” *Bulletin of Electrical Engineering and Informatics*, vol. 11, no. 2, pp. 1091–1099, Apr. 2022, doi: 10.11591/eei.v11i2.3525.
- [13] X. Ma, W. Chen, and Y. Xu, “ERCP-Net: a channel extension residual structure and adaptive channel attention mechanism for plant leaf disease classification network,” *Sci Rep*, vol. 14, no. 1, pp. 1–14, 2024, doi: 10.1038/s41598-024-54287-3.
- [14] Z. Bin Niu, S. Y. Jia, and H. H. Xu, “Automated graptolite identification at high taxonomic resolution using residual networks,” *iScience*, vol. 27, no. 1, p. 108549, 2024, doi: 10.1016/j.isci.2023.108549.
- [15] D. Sarwinda, R. H. Paradisa, A. Bustamam, and P. Anggia, “Deep Learning in Image Classification using Residual Network (ResNet) Variants for Detection of

- Colorectal Cancer,” *Procedia Comput Sci*, vol. 179, no. 2019, pp. 423–431, 2021, doi: 10.1016/j.procs.2021.01.025.
- [16] L. Ali and S. A. C. Bukhari, “An Approach Based on Mutually Informed Neural Networks to Optimize the Generalization Capabilities of Decision Support Systems Developed for Heart Failure Prediction,” *IRBM*, vol. 42, no. 5, pp. 345–352, 2021, doi: <https://doi.org/10.1016/j.irbm.2020.04.003>.
- [17] M. A. A. Fawwaz, K. N. Ramadhani, and F. Sthevani, “Klasifikasi Ras pada hewan peliharaan menggunakan Algoritma Convolutional Neural Network (CNN),” vol. 8, no. 1, pp. 715–730, 2020.
- [18] Rima Dias Ramadhani, A. Nur Aziz Thohari, C. Kartiko, A. Junaidi, T. Ginanjar Laksana, and N. Alim Setya Nugraha, “Optimasi Akurasi Metode Convolutional Neural Network untuk Identifikasi Jenis Sampah,” *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 5, no. 2, pp. 312–318, Apr. 2021, doi: 10.29207/resti.v5i2.2754.
- [19] J. Sanjaya and M. Ayub, “Augmentasi Data Pengenalan Citra Mobil Menggunakan Pendekatan Random Crop , Rotate , dan Mixup,” vol. 6, pp. 311–323, 2020.
- [20] Y. Vita Via, I. Yuniar Purbasari, and A. Putra Pratama, “Analisa Algoritma Convolution Neural Network (Cnn) Pada Klasifikasi Genre Musik Berdasar Durasi Waktu,” *SCAN Jurnal Teknologi dan Informasi*, vol. 17, no. 1, pp. 35–41, 2022, [Online]. Available: <http://ejournal.upnjatim.ac.id/index.php/scan/article/view/3251/2003>
- [21] R. Z. Fadillah, A. Irawan, M. Susanty, and I. Artikel, “Data Augmentasi Untuk Mengatasi Keterbatasan Data Pada Model Penerjemah Bahasa Isyarat Indonesia (BISINDO),” *Jurnal Informatika*, vol. 8, no. 2, pp. 208–214, 2021, [Online]. Available: <https://ejournal.bsi.ac.id/ejurnal/index.php/ji/article/view/10768>

- [22] N. E. Khalifa, M. Loey, and S. Mirjalili, “A comprehensive survey of recent trends in deep learning for digital images augmentation,” *Artif Intell Rev*, vol. 55, no. 3, pp. 2351–2377, 2022, doi: 10.1007/s10462-021-10066-4.
- [23] W. M. Pradnya D and A. P. Kusumaningtyas, “Analisis Pengaruh Data Augmentasi Pada Klasifikasi Bumbu Dapur Menggunakan Convolutional Neural Network,” *Jurnal Media Informatika Budidarma*, vol. 6, no. 4, p. 2022, 2022, doi: 10.30865/mib.v6i4.4201.
- [24] D. Putri Ayuni, Jasril, M. Irsyad, F. Yanto, and S. Sanjaya, “Augmentasi Data Pada Implementasi Convolutional Neural Network Arsitektur Efficientnet-B3 Untuk Klasifikasi Penyakit Daun Padi,” *ZONasi: Jurnal Sistem Informasi*, vol. 5, no. 2, pp. 239–249, 2023, doi: 10.31849/zn.v5i2.13874.
- [25] T. B. Sasongko, H. Haryoko, and A. Amrullah, “Analisis Efek Augmentasi Dataset dan Fine Tune pada Algoritma Pre-Trained Convolutional Neural Network (CNN),” *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 10, no. 4, pp. 763–768, 2023, doi: 10.25126/jtiik.20241046583.
- [26] Hendri Candra Mayana and Desmarita Leni, “Deteksi Kerusakan Ban Mobil Menggunakan Convolutional Neural Network dengan Arsitektur ResNet-34,” *Jurnal Surya Teknika*, vol. 10, no. 2, pp. 842–851, 2023, doi: 10.37859/jst.v10i2.6336.
- [27] N. I. Sanusi, S. Ramadhani, and M. Irsyad, “Analisa Gambar X-Ray Mammography dengan Convolution Neural Network pada Deep Learning dengan Arsitektur Resnet,” *Jurnal Sistem Komputer dan Informatika (JSON)*, vol. 4, no. 4, p. 604, 2023, doi: 10.30865/json.v4i4.6365.
- [28] M. F. Gunardi, “Implementasi Augmentasi Citra pada Suatu Dataset,” *Jurnal Informatika*, vol. 9, no. 1, pp. 1–5, 2023.

-
- [29] W. Maulana Baihaqi, C. Raras, A. Widiawati, D. P. Sabila, and A. Wati, “Analisis Gambar Sel Darah Berbasis Convolution Neural Network untuk Mendiagnosis Penyakit Demam Berdarah Convolution Neural Network-Based Image Analysis of Blood Cells to Diagnose Dengue Fever,” *Cogito Smart Journal* |, vol. 7, no. 1, 2021.
- [30] R. Cendekia Vandara, S. A. Wibowo, and K. Usman, “Performance Analysis of Face Alignment On 3-Dimensional (3D) Face Reconstruction Using Modified Position Map Regression Network.” Thesis, Telkom University, 2021.
- [31] A. Ridhovan and A. Suharso, “Penerapan Metode Residual Network (Resnet) Dalam Klasifikasi Penyakit Pada Daun Gandum,” *JIPi (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*, vol. 7, no. 1, pp. 58–65, 2022, doi: 10.29100/jipi.v7i1.2410.
- [32] I. Ariawan *et al.*, “Extraction of Morphometric Features the shape of mangrove leaves based on digital images and classification using the Support Vector Machine,” *Karbala International Journal of Modern Science*, vol. 10, no. 2, May 2024, doi: 10.33640/2405-609X.3349.
- [33] L. N. Smith, “A disciplined approach to neural network hyper-parameters: Part 1 - learning rate, batch size, momentum, and weight decay,” Mar. 2018, [Online]. Available: <http://arxiv.org/abs/1803.09820>
- [34] J. Xu, Y. Zhang, and D. Miao, “Three-way confusion matrix for classification: A measure driven view,” *Inf Sci (N Y)*, vol. 507, pp. 772–794, 2020, doi: <https://doi.org/10.1016/j.ins.2019.06.064>.

SVM and Ensemble Majority Voting Algorithm on Sentiment Analysis of Using ChatGPT in Education

Hildegardis Yayukristi Weko¹, Hari Suparwito^{1*}

¹*Informatics Department, Sanata Dharma University*

^{*}*Corresponding Author: shirsj@jesuits.net*

(Received 11-06-2025; Revised 18-10-2025; Accepted 18-10-2025)

Abstract

The pros and cons of using ChatGPT in education have caused academic debate as it has influenced current educational praxis. Discussions about the possibility of ChatGPT for writing manuscripts or doing assignments are rife on social media, one of which is Twitter. The purpose of this study is to understand the public perception of the use of ChatGPT in education. The proposed method is sentiment analysis with SVM and Majority Voting algorithms. SVM is one of the superior algorithms in pattern recognition and is suitable for use in classification. The Majority Voting ensemble algorithm combines independent algorithms' prediction results. In this research, majority voting uses three base classifiers, namely Naïve Bayes, Random Forest, and KNN. The results of the study showed that the accuracy of SVM is 83.6% and Majority Voting is 85.4%, with the accuracy of the NB, RF, and KNN base classifiers of 76.82%, 80.91%, and 74.5%, respectively. This proved that the Majority Voting Ensemble is superior to individual algorithms with higher accuracy values. This follows the results of previous research, where the ensemble performs better than the individual algorithm. The accuracy values of SVM and the Ensemble Majority Voting models showed that both models could successfully classify sentiment on tweet data for using ChatGPT in education.

Keywords: Education, Ensemble Majority Voting, Sentiment Analysis, ChatGPT, SVM

1 Introduction

ChatGPT is a natural language model developed by OpenAI to understand and provide natural responses in human conversations [1]. The development of ChatGPT began in 2018, and it was one of the most significant language models at the time. Since then, OpenAI has launched several more prominent versions of ChatGPT, including GPT-2, GPT3, and GPT-4. The general public has widely used ChatGPT to write various



education manuscripts. Therefore, it is common for students to use ChatGPT to complete school or university assignments [2].

The advent of ChatGPT has created pros and cons. Teachers revealed that the development of AI seems to have changed the current educational praxis [3], [4]. Nadiem Makarim questioned the impact of ChatGPT, as it makes teachers fearful due to the assessment of quantity and quality in the teaching-learning process. For example, in Texas, students' class certificates were held up due to the use of ChatGPT. However, ChatGPT, which can help improve learning effectiveness by providing access to a broader range of materials, is undeniable [5].

These pros and cons are often discussed on social media, especially Twitter. Twitter is one of the most important social media platforms in social interaction because users for audience management favour it with one of its "#" hashtag features. The hashtag #ChatGPT is always there every day, which indicates that there is high public attention to ChatGPT. This draws attention to sentiment analysis on tweets using ChatGPT [6]. Sentiment analysis is extracting opinions to analyze a person's opinion, sentiment or feelings towards a particular topic [7]. Previous studies have shown that #chatGPT is trending and makes for exciting research. Ananya Sarker et al. [8] conducted sentiment analysis of Twitter data with six different algorithms and found that SVM has the highest accuracy of 84%. Other studies were conducted by Rajani et al. [9]. They did a sentiment analysis of ChatGPT with four different algorithms and found that SVM had the best accuracy of 81.4% compared to NB, RF, and KNN. Abdullah Alsaeedi [10] has also conducted a study on sentiment analysis, which compares several algorithms in sentiment analysis, finding that the ensemble and hybrid algorithms are 85% superior to the SVM and Naïve Bayes algorithms. In the previous research, the use of Ensemble Majority Voting classifiers was also able to improve the model accuracy of single classifier models, especially for Decision Tree and KNN classifier models [11].

The purpose of this study is to analyze public sentiment towards the use of ChatGPT in education using machine learning approaches. Support Vector Machine (SVM) and Ensemble Majority Voting algorithms were used to perform sentiment analysis

on tweets with the hashtag #chatGPT. The Ensemble Majority Voting algorithm was built from Random Forest, KNN, and Naïve Bayes to optimize low-accuracy results from the three single algorithms.

2 Material and Methods

The overall work process carried out in this study can be divided into data collection, data preprocessing, data splitting, feature extraction, data modelling and performing metrics, as shown in Fig. 1 below.

2.1 Data

The first process is Data collection. This process consisted of three steps: data crawling, labelling and sampling. The data were crawled from Twitter with the keyword #ChatGPT from 24/03/2023 to 26/07/2023. The number of data is 8517 data tweets. The data were retrieved through the crawling process using Twitter API with the help of the *tweepy library* from Python. Data retrieval was based on #ChatGPT with several search keywords such as academic, school, assignment, exam, thesis, quiz, journal, seminar, and paper. The next is data labelling. Labelling or sentiment on the data was undertaken using *HuggingFace Transformers*, an open-source library that facilitates users to access large-scale models in building and experimenting [12]. The model used is the Indonesian RoBERTa Sentiment Classifier. The RoBERTa model (Robustly optimized BERT Pretraining approach) is a BERT (Bidirectional Encoder Representations from Transformers) pre-train model that has been optimized to exceed the performance of all post_BERT methods [13]. Indonesian RoBERTa Sentiment Classifier is a sentiment classification model based on the enhanced RoBERTa model on the Indonesian dataset with an accuracy of 94.36% [14].

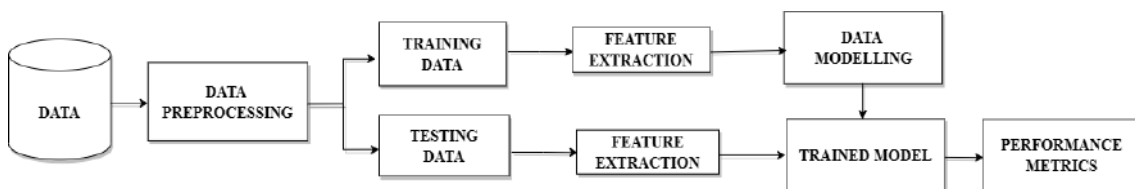


Figure 1. The overall work processes

The dataset that has been obtained is labelled using this model with three sentiment labels, namely positive sentiment, neutral sentiment and negative sentiment. The last step is balancing data. Unbalanced data sets in machine learning can result in incorrect predictions. Machine learning algorithms are designed to work best with balanced data [15]. Undersampling techniques were applied to the number of samples in the majority class to be balanced with the minority class.

2.2 Preprocessing

The data obtained from Twitter is unstructured and still contains noise. It is necessary to do preprocessing with several stages according to Fig. 2, namely Cleaning, to remove characters in tweets except the alphabet and emoticons.

To convert the emoticon in the tweet into a word that represents the emoticon, Emoticon Conversion is used. Next, Case Folding is applied to convert all letters in the tweet to lowercase. Sentences in tweets are truncated based on the words that make up the sentence using the Tokenize process. To convert ambiguous words such as abbreviations and acronyms into the original form of words that are considered to be normal language, Normalisation is used. Stopword Removal is used to eliminate words that have no meaning, for example, “yang”, “dan”, and “di”. The next step is Stemming. In this step, words would be converted into basic words.

2.3 Training and Testing Dataset

After Stemming, the data is ready to be divided into training and testing data with a ratio of 90:10, respectively.

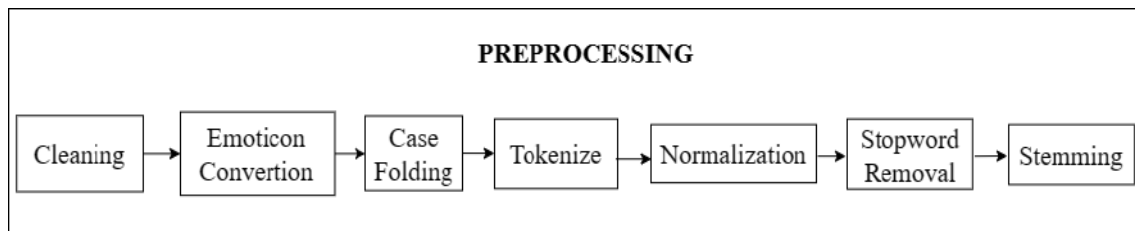


Figure 2. Preprocessing Flowchart

2.4 Feature extraction

Next, in each dataset, the words are weighted using the TF-IDF technique. Term Frequency-Inverse Document Frequency (TF-IDF) is a word weighting method used to measure the significance of features in a document [16]. Term Frequency (TF) determines how important a word is by how often it appears in a document; Inverse Document Frequency (IDF) considers a word critical in a document if it does not seem too often in other documents [17].

2.5 Modelling

Modelling is done using the data from feature extraction. Each model is validated using the K-Fold Cross Validation method on the training data. This method divides the data set into several parts and tests them individually while training is performed on the remaining data [18]. On training each fold, its output error is estimated; finally, the average of all mistakes is the estimated true error [19]. Fig. 3 shows all the various processes that occur in data processing.

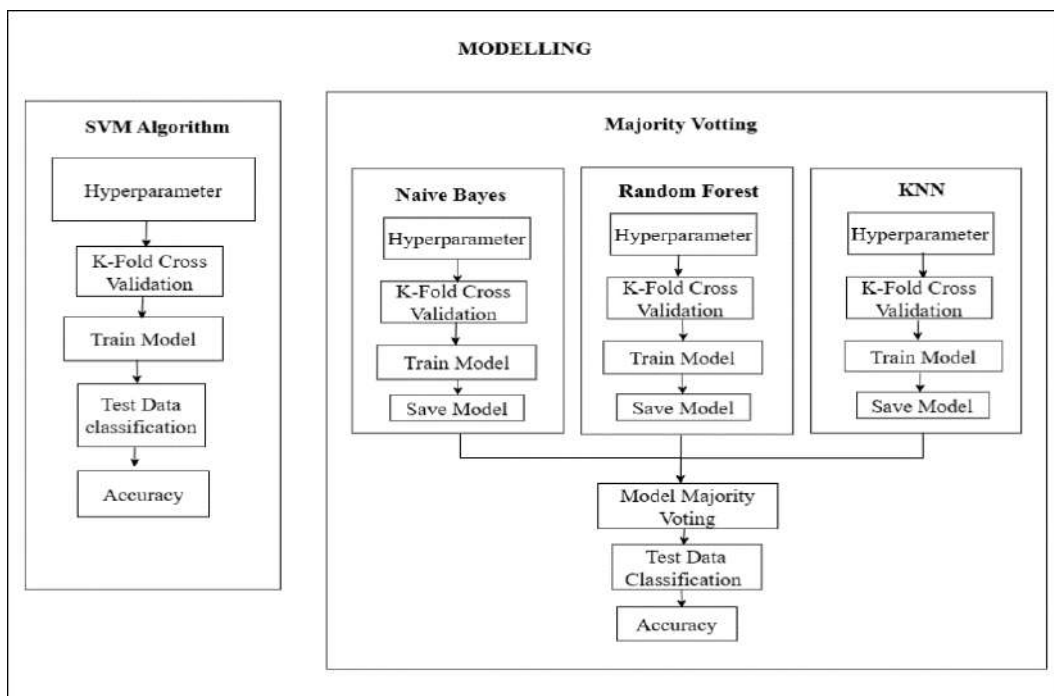


Figure 3. The modelling workflow

2.5.1 Support Vector Machine

SVM is a learning method that works according to the SRM (Structural Risk Minimisation) principle that aims to find the best hyperplane separating two classes in the input space [20]. A *hyperplane* is a dividing line if the number of input features is 2; a 2D hyperplane is needed to separate the number of input features 3 [21]. SVM uses kernels to map data to a higher dimensional space for data classification that cannot be done linearly [22], one of which is by using RBF (Radial Basis Function) kernels, which can be used in Non-linear SVM models [21]. Because of its ability that does not depend on the number of features, SVM is widely applied in classification problems and can produce good performance [20]. This research searches for the best kernel between RBF and Polynomial kernels with supporting parameters. The parameter combinations that will be performed are shown in Table 1.

2.5.2 Ensemble Majority Voting

The majority voting uses one or more classification algorithms, and the output results are given based on selecting predictors (model values from all algorithms) [23]. Each classifier, called the base classifier, chooses one class label, and the final output class label is the class label that receives more than half of the votes [24]. We combined three different algorithms: Naïve Bayes, Random Forest, and KNN. These three algorithms are optimized by finding the best parameters using GridSearchCV. The parameter combinations for all three are shown in Table 2.

Table 1. Combination of SVM parameters

Combination	Parameters	Value
1	Kernel C <i>gamma</i>	RBF [0.1, 1, 10, 100] [1, 0.1, 0.01, 0.001, 0.0001]
2	Kernel C <i>degree</i> <i>coefficient</i>	<i>Polynomial</i> [0.1, 1, 10, 100] [2, 3, 4] [0.0, 1.0, 2.0]

Table 2. Combination of Base Classifier Parameters

Algorithm	Parameters	Value
Naïve Bayes	nb__alpha	[1, 0.1, 0.01, 0.001, 0.0001]
	fit_prior	['True', 'False']
Random Forest	N_estimators	[5, 50, 100]
	Max_dept	[2, 10, 20, None]
KNN	n_neighbors	[3, 5, 7, 9, 11]
	weights	['uniform', 'distance']
	algorithm	['auto', 'ball_tree', 'kd_tree', 'brute']
	p	[1, 2]

2.6 Model Evaluation

Measurement of algorithm performance is measured using the Confusion Matrix method. The classification results obtained from each model will form a confusion matrix so that it can be used to calculate accuracy, recall, and precision.

3 Result and Discussions

3.1 Crawling Data

Tweets were crawled 11 times according to the keywords followed after #ChatGPT. For example, the first crawl used "#ChatGPT Academic", the second crawl used "#ChatGPT School", and so on. The data retrieved were tweets, id, dates, and usernames with a total of 8,517 crawled data. An example of crawled Twitter data is shown in Table 3.

Table 3. Example of tweet crawling results

Tweet	Id	Date/Time	Username
baru kali ini gw ngerjain ujian pake chatGPT wkwk.. enak bgt tyt, lovvvv	1553693089512194049	02-05-2023 10:33:54+00:00	kmjeongsgf
@RyomenRogue Ngantri chatgpt dulu agaknya 🤔	884334345900630016	02-05-2023 10:12:02+00:00	nakaharaRogue
GW CAPEK BANGET NUGAS SAMA CHATGPT PAKE ACARA NGAMBEK SEGALA NI AI	1380431790410584064	08-04-2023 09:29:35+00:00	ambivaIens

3.2 Labelling Data

The Indonesian RoBERTa Sentiment Classifier was used for labelling. Table 4 shows an example of the labelling results. The data distribution for each sentiment on the 8,517 tweets after labelling is shown in Fig. 4.

3.3 Undersampling

The labelled tweet data has an unbalanced amount of data, where the number of negative and neutral tweets is more than the number of positive tweets (Fig. 4). Undersampling technique balances the amount of tweet data in each sentiment classification. The amount of tweet data after undersampling is 5000 data, with the number of positive tweets = 1653, negative = 1705 and neutral = 1642. The sentiment distribution after undersampling is shown in Fig. 5.

Table 4 Example of labelling results

Tweet	Label	Score
baru kali ini gw ngerjain ujian pake chatGPT wkwk.. enak bgt tyt, lovvvv	POSITIVE	0,993
@RyomenRogue Ngantri chatgpt dulu agaknya 😊	NEUTRAL	0,917
GW CAPEK BANGET NUGAS SAMA CHATGPT PAKE ACARA NGAMBEK SEGALA NI AI	NEGATIVE	0,937

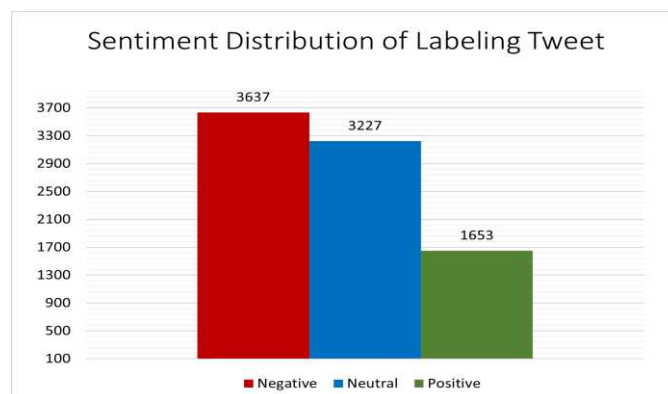


Figure 4. Sentiment Distribution of labelling tweets

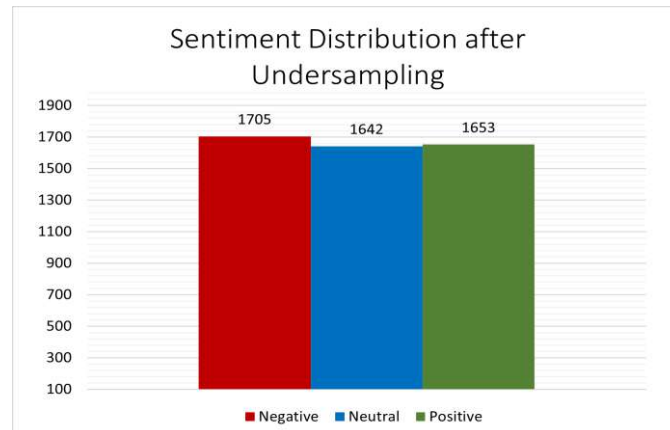


Figure 5. Sentiment Distribution of Tweets After undersampling

3.4 Preprocessing

Preprocessing is done to clean and prepare the tweet data. For example, the tweet data "@RyomenRogue Ngantri chatgpt dulu agaknya 😂" is cleaned through the cleaning stage and produces tweet data as in Table 5. Then, the emoticon conversion stage is applied and produces tweet data, as in Table 6. Next is to shrink all the letters with case folding. The results of case folding are shown in Table 7. The tokenizing stage follows this, and the tokenizing results are shown in Table 8. The next stage is Normalisation. The normalization results are shown in Table 9. The normalization results are then continued with the stopwords removal stage. The tweet data after experiencing stopwords removal is shown in Table 10. The last stage in preprocessing is stemming. The results of stemming can be seen in Table 11.

Table 5. Cleaning Result

Tweet	Cleaning
@RyomenRogue Ngantri chatgpt dulu agaknya 😂	Ngantri chatgpt dulu agaknya 😂

Table 6. Emoticon Conversion Result

Cleaning	Emoticon Conversion
Ngantri chatgpt dulu agaknya 😂	Ngantri chatgpt dulu agaknya Senang

Table 7. The Case Folding Result

Emoticon Conversion	Case Folding
Ngantri chatgpt dulu agaknya Senang	ngantri chatgpt dulu agaknya senang

Table 8. Tokenize Result

Case Folding	Tokenize
ngantri chatgpt dulu agaknya senang	['ngantri', 'chatgpt', 'dulu', 'agaknya', 'senang']

Table 9. The Normalization Result

Tokenize	Normalization
['ngantri', 'chatgpt', 'dulu', 'agaknya', 'senang']	['mengantri', 'chatgpt', 'dulu', 'agaknya', 'senang']

Table 10. Stopword Removal Result

Normalization	Stopword removal
['mengantri', 'chatgpt', 'dulu', 'agaknya', 'senang']	['mengantri', 'chatgpt', 'senang']

Table 11 Stemming Result

Stopword removal	Stemming
['mengantri', 'chatgpt', 'senang']	['antri', 'chatgpt', 'senang']

After the Stemming process, 5000 tweet data were divided into two datasets, namely the training and testing dataset, with a ratio of 90:10. So, there are 4500 data tweets in the training dataset and 500 data tweets in the testing data. Furthermore, from 4500 training data, a k-fold cross-validation process is carried out, which divides the training data and validation data according to the k value.

The last process before modelling was feature extraction. Feature extraction is done to convert the data into numerical form and is formed into a matrix for each word. The unigram function in the TF-IDF module calculates the weight of a single word. The results of feature extraction using TF-IDF for the 20 words with the highest weights are shown in Fig. 6.

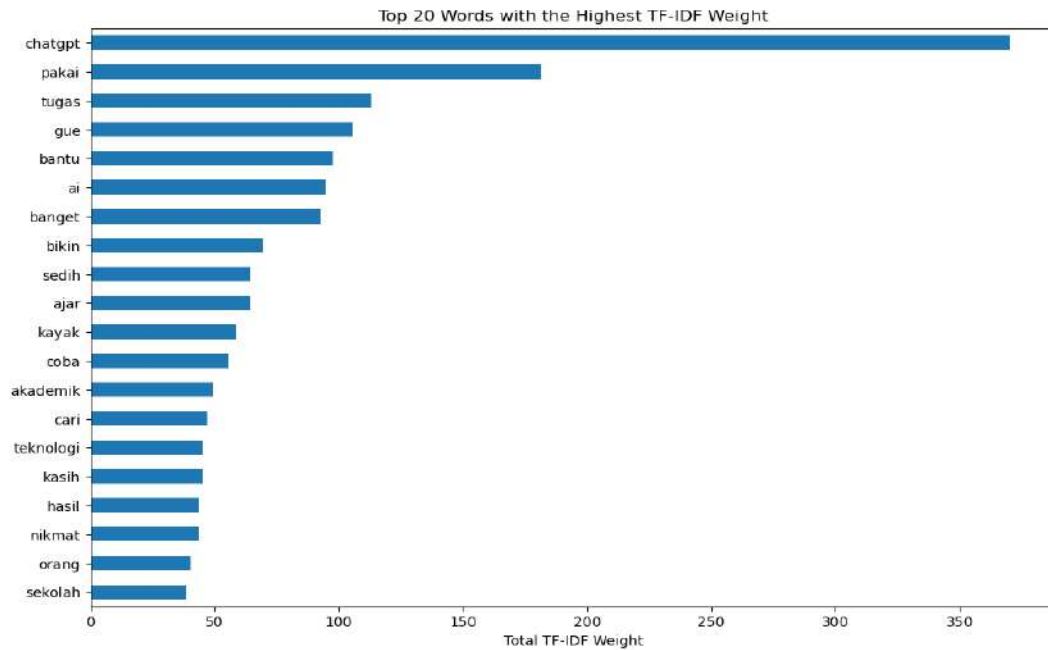


Figure 6. Top 20 Words with Highest TF-IDF Weight

3.5 Modelling Analysis

3.5.1 Support Vector Machine Model (SVM)

The accuracy results of the Support Vector Machine (SVM) on the training dataset were carried out by finding the best parameters between the RBF and Polynomial kernels (see Table 12). The results on the RBF and Polynomial kernels are shown with various Cost values. In contrast, the gamma, coefficient and degree values are fixed because these parameter values are stable and provide good accuracy compared to other parameter values. The best parameters of the SVM model for training data were the 'RBF' kernel with C value = 100 and gamma value = 1. With the K-Fold Cross Validation method K = 10 the SVM model got the best accuracy of 82.94%.

Table 12. The optimal SVM parameter training data results

	C	Accuracy K=5	Accuracy K=7	Accuracy K=10
Combination of RBF <i>gamma</i> = 1	0,1	0.4532	0.4612	0.4686
	1	0.8161	0.8189	0.8274
	10	0.8177	0.8233	0.8284
	100	0.8178	0.8245	0.8294

Combination of <i>Polynomial</i> <i>coefficient</i> = 1, <i>degree</i> = 4	C	Accuracy K=5	Accuracy K=7	Accuracy K=10
	0.1	0.8184	0.8202	0.8283
	1	0.8189	0.8225	0.8283
	10	0.8199	0.8235	0.8284
	100	0.8195	0.8235	0.8284

3.5.2 Ensemble Majority Voting Model

The majority voting model is built with three base classifiers, namely Naïve Bayes, Random Forest and KNN. The three models were optimized by finding the best parameters. The search results for each base classifier are shown in Table 13. The best parameters of the three algorithms are then used in model training with K-Fold cross-validation. It can be seen that the three algorithms get their best accuracy at K = 10, where the accuracy for Naïve Bayes, KNN and Random Forest algorithms are 77.1%, 74.1%, and 80.6%, respectively. The accuracy of the three base classifier compared to SVM (see Fig. 7) showed that these three individual algorithms have a lower accuracy.

Table 13. Base Classifier Parameter Results

Algorithm	K=5		K=7		K=10	
	Parameter	Accuracy	Parameter	Accuracy	Parameter	Accuracy
Naïve Bayes	<i>alpha:0.5</i> <i>fit_prior:False</i>	0.763	<i>alpha:0.5</i> <i>fit_prior:False</i>	0.765	<i>alpha:0.5</i> <i>fit_prior:true</i>	0.771
KNN	<i>algorithm:auto</i> <i>n_neighbors:9</i> <i>p:2</i> <i>weights:distance</i>	0.725	<i>algorithm:auto</i> <i>n_neighbors:7</i> <i>p:2</i> <i>weights:distance</i>	0.72889	<i>algorithm:auto</i> <i>n_neighbors:11</i> <i>p:2</i> <i>weights: distance</i>	0.741
Random Forest	<i>max_depth:</i> <i>None</i> <i>n_estimators:100</i>	0.789	<i>max_depth:</i> <i>None</i> <i>n_estimators:100</i>	0.802	<i>max_depth:</i> <i>None</i> <i>n_estimators:100</i>	0.806

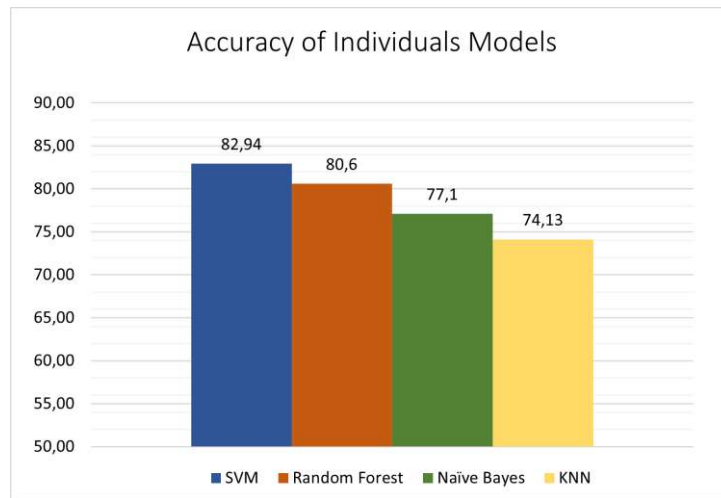


Figure 7. Comparison of individual algorithms

Next, the three base classifier algorithms with lower accuracy than SVM are used as the Majority Voting model builder in testing the test data.

3.5.3 SVM and Majority Voting on Testing Dataset

The SVM and Majority Voting models were tested on 500 testing datasets. The results of both models on the classification of test data are shown in Table 14. From the test results, it can be seen that the Majority Voting results now becoming higher than the Support Vector Machine (SVM) model with precision, recall and accuracy values of 86.04%, 86.09%, and 85.6%, respectively. This proves that the ensemble model works better than the individual classification model. In several previous studies [25]–[27], the ensemble classifiers performed better than the single classifier models. The main reason the ensemble classifier is better than the single model is that it provides a way to reduce the prediction variance, i.e. the amount of error in the prediction made by the single model forming the ensemble. When this occurs, this reduction in variance, in turn, leads to improved prediction performance [28].

Table 14. SVM and Majority Voting Model Results

Performing	SVM	Majority Voting
Precision	0,8392	0,8604
Recall	0,8363	0,8609
Accuracy	0,836	0,856

4 Conclusions

The SVM algorithm provides the best results in performing sentiment classification compared to 3 other algorithms, namely Naïve Bayes, Random Forest, and KNN. The Majority Voting model built with three low-accuracy algorithms (Naïve Bayes, Random Forest and KNN) can produce better accuracy than the SVM model with a difference of 2% where SVM accuracy is 83.6% and Majority voting accuracy is 85.6%. This proved the ensemble model that combines several individual models has successfully improved accuracy in sentiment classification on tweet data using ChatGPT in education.

References

- [1] M. Dowling and B. Lucey, “ChatGPT for (Finance) research: The Bananarama Conjecture,” *Financ Res Lett*, vol. 53, May 2023, doi: 10.1016/j.frl.2023.103662.
- [2] Ö. Aydın and E. Karaarslan, “OpenAI ChatGPT Generated Literature Review: Digital Twin in Healthcare,” *SSRN Electronic Journal*, Dec. 2022, doi: 10.2139/ssrn.4308687.
- [3] D. Baidoo-Anu and L. Ansah, “Education in the Era of Generative Artificial Intelligence (AI): Understanding the Potential Benefits of ChatGPT in Promoting Teaching and Learning,” Oct. 2023.
- [4] U. Bukar, M. S. Sayeed, S. F. Abdul Razak, S. Yogarayan, and O. Amodu, “Text Analysis of Chatgpt as a Tool for Academic Progress or Exploitation,” *SSRN Electronic Journal*, Oct. 2023, doi: 10.2139/ssrn.4381394.
- [5] I. Arifdarma, “Pengaruh Teknologi ChatGPT Terhadap Dunia Pendidikan: Potensi dan Tantangan,” *Jurnal Agriwidya*, vol. 4, no. 1, Mar. 2023, Accessed: Oct. 18, 2023. [Online]. Available: <https://repository.pertanian.go.id/handle/123456789/20278>
- [6] G. Koca, “Sentiment analysis with twitter data on bitcoin,” *Anadolu Üniversitesi İktisadi ve İdari Bilim. Fakültesi Derg.*, vol. 22, no. 4, pp. 19–30, 2021.
- [7] B. Pang and L. Lee, “Opinion mining and sentiment analysis,” *Foundations and Trends® in information retrieval*, vol. 2, no. 1–2, pp. 1–135, 2008.

-
- [8] A. Sarker, M. S. U. Zaman, and M. A. Y. Srizon, "Twitter data classification by applying and comparing multiple machine learning techniques," *International Journal of Innovative Research in Computer Science & Technology (IJIRCST) ISSN*, pp. 2347–5552, 2019.
 - [9] M. Tubishat, F. Al-Obeidat, and A. Shuhaiber, "Sentiment Analysis of Using ChatGPT in Education," *Wseas Transactions On Advances In Engineering Education*, vol. 20, pp. 60–66, Oct. 2023, doi: 10.37394/232010.2023.20.9.
 - [10] A. Alsaeedi and M. Z. Khan, "A study on sentiment analysis techniques of Twitter data," *International Journal of Advanced Computer Science and Applications*, vol. 10, no. 2, pp. 361–374, 2019, doi: 10.14569/ijacsa.2019.0100248.
 - [11] A. K. Putri and H. Suparwito, "Uji Algoritma Stacking Ensemble Classifier pada Kemampuan Adaptasi Mahasiswa Baru dalam Pembelajaran Online," *KONSTELASI: Konvergensi Teknologi dan Sistem Informasi*, vol. 3, no. 1, pp. 1–12, 2023.
 - [12] T. Wolf *et al.*, "Transformers: State-of-the-Art Natural Language Processing," Oct. 2020, pp. 38–45. doi: 10.18653/v1/2020.emnlp-demos.6.
 - [13] Y. Liu *et al.*, "RoBERTa: A Robustly Optimized BERT Pretraining Approach." Oct. 2019.
 - [14] Wilson Wongso, "indonesian-roberta-base-sentiment-classifier (Revision e402e46)." Hugging Face, Indonesia, 2023. Accessed: Oct. 18, 2023. [Online]. Available: <https://huggingface.co/w11wo/indonesian-roberta-base-sentiment-classifier>
 - [15] A. Indrawati, H. Subagyo, A. Sihombing, and S. A. Wagiyah, "Analyzing the Impact of Resampling Method for Imbalanced Data Text in Indonesian Scientific Articles Categorization," *BACA: Jurnal Dokumentasi dan Informasi*, vol. 42, no. 2, pp. 133–141, 2020.

- [16] C. Zhai and S. Massung, *Text data management and analysis: a practical introduction to information retrieval and text mining*. Morgan & Claypool, 2016.
- [17] F. Alzami, E. D. Udayanti, D. P. Prabowo, and R. A. Megantara, "Document Preprocessing with TF-IDF to Improve the Polarity Classification Performance of Unstructured Sentiment Analysis," *Kinetik: Game Technology, Information System, Computer Network, Computing, Electronics, and Control*, pp. 235–242, Aug. 2020, doi: 10.22219/kinetik.v5i3.1066.
- [18] M. U. Muhammad *et al.*, "Sentiment Analysis of Students' Feedback before and after COVID-19 Pandemic," *International Journal on Emerging Technologies*, vol. 12, no. 2, pp. 177–182, 2021, [Online]. Available: www.researchtrend.net
- [19] S. Shalev-Shwartz and S. Ben-David, *Understanding machine learning: From theory to algorithms*. Cambridge university press, 2014.
- [20] A. L. Hananto, B. Priyatna, A. Hananto, and A. P. Nardilasari, *DATA MINING : Penerapan Algoritma (SVM, Naive Bayes, KNN) Dan Implementasi Menggunakan Rapid Miner*. Bandung, Jawa Barat: Media Sains Indonesia, 2023.
- [21] N. Pavitha *et al.*, "Movie recommendation and sentiment analysis using machine learning," *Global Transitions Proceedings*, vol. 3, no. 1, pp. 279–284, 2022.
- [22] F. Rahmadayana and Y. Sibaroni, "Sentiment analysis of work from home activity using SVM with randomized search optimization," *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 5, no. 5, pp. 936–942, 2021.
- [23] J. Sadhasivam and R. B. Kalivaradhan, "Sentiment analysis of Amazon products using ensemble machine learning algorithm," *International Journal of Mathematical, Engineering and Management Sciences*, vol. 4, no. 2, p. 508, 2019.
- [24] Z.-H. Zhou, *Ensemble methods: foundations and algorithms*. CRC press, 2012.
- [25] M. Khalid, I. Ashraf, A. Mehmood, S. Ullah, M. Ahmad, and G. S. Choi, "GBSVM: Sentiment classification from unstructured reviews using ensemble classifier,"

Applied Sciences (Switzerland), vol. 10, no. 8, Apr. 2020, doi: 10.3390/APP10082788.

- [26] M. Hosni, I. Abnane, A. Idri, J. M. Carrillo de Gea, and J. L. Fernández Alemán, “Reviewing ensemble classification methods in breast cancer,” *Comput Methods Programs Biomed*, vol. 177, pp. 89–112, 2019, doi: <https://doi.org/10.1016/j.cmpb.2019.05.019>.
- [27] R. M. K. Saeed, S. Rady, and T. F. Gharib, “An ensemble approach for spam detection in Arabic opinion texts,” *Journal of King Saud University - Computer and Information Sciences*, vol. 34, no. 1, pp. 1407–1416, 2022, doi: <https://doi.org/10.1016/j.jksuci.2019.10.002>.
- [28] J. Brownlee, “Why Use Ensemble Learning?,” machinelearningmastery.com. Accessed: Nov. 06, 2023. [Online]. Available: <https://machinelearningmastery.com/why-use-ensemble-learning/#:~:text=There%20are%20two%20main%20reasons,the%20predictions%20and%20model%20performance>

This page intentionally left

AUTHOR GUIDELINES

Author guidelines are available at the journal website:

<http://e-journal.usd.ac.id/index.php/IJASST/about/submissions#authorGuidelines>

This page intentionally left blank