

# NUANSA INFORMATIKA

## JURNAL TEKNOLOGI INFORMASI

Volume 19 - Nomor 2 - Juli 2025

p-ISSN : 1858 - 3911

e-ISSN : 2614-5405

Fakultas Ilmu Komputer  
Universitas Kuningan

**IoT and Water Consumption Forecasting: A Green Accounting Study at a Coffee Shop in Cimahi**  
Mohamad Nurkamal Fauzan, Riani Tanjung -Page - 1-10

**Feature Extraction Using Markov Random Field (MRF) On Improving CNN Classification Results For Alzheimer's Disease Diagnosis**  
Saeful Bahri, Miftah Farid Adiwisatra, Taufik Hidayatulloh - Page - 11-17

**Assessing the Role of E-Govqual in Improving Public Satisfaction with the MyPertamina Application in Setiabudi District**  
Yulian Surya Saputra, Sucitra Sahara, Rizqi Agung Permana - Page - 18-24

**A Data-Driven Approach to Comparative Evaluation of Regression Models for Accurate House Price Prediction**  
Tiara Permata Hati, Budiman Budiman, Nur Alamsyah - Page - 25-34

**Development Of A 2d Educational Animal Card Game Using Unity**  
Alma Sulaiman, Mamay Syani - Page - 35-48

**The Implementation of User Satisfaction Based Shortest Job First Algorithm Efficiency in Queuing System**  
Darsanto - Page - 49-57

**Predicting Basic Shipping Tariff Using Machine Learning**  
Dani Ferdinan, Nisa Hanum Harani, M. Yusril Helmi Setyawan - Page - 58-66

**Development Of Online Attendance Application With Time And Location Validation At Satpol Pp South Tangerang**  
Syamsu Hidayat, Muhamad Rifki Ramadhan, Joko Susilo. - Page - 67-72

**Web-Based Assignment Management System Application Design to Increase Personnel Productivity and Effectiveness in Product Documentation Division**  
Eryan Ahmad Firdaus, Seftian Candra Pratama, Rochedi Idul Adha - Page - 73-80

**Emoji-Based Sentiment Classification Using Ensemble Learning with Cross-Validation: A Lightweight Approach for Social Media Analysis**  
Nur Alamsyah, Gunthar Bayu Wibisono, Titan Parama Yoga, Budiman, Acep Hendra - Page - 81-87

**A Bidirectional GRU Approach with Hyperparameter Optimization for Sentiment Classification in Game Reviews**  
Nur Alamsyah, Titan Parama Yoga, Budiman, Imannudin Akbar, Acep Hendra, Alif Januantara Prima - Page - 88-95

**The Implementation of the TOPSIS Method in Determining Stunting Toddlers**  
Umi Khultsum, F. Lia Dwi Cahyanti, Elly Firasari - Page - 96-104

**Predicting the Happiness Index Based on the HDI Indicator in Indonesia Using the Ensemble Learning Approach**  
Syafrial Fachri Pane, Rofi Nafis Zain, Iwan Setiawan, Viridiandry Putratama - Page - 105-114

**Performance and Efficiency Testing Analysis of Database Systems in Academic Information Systems**  
Utami Kusuma Dewi, Ryan Lingga Wicaksono, Mahendra Dwifabri Purbolaksono, Villy Satria - Page - 115-122

**Fuzzy-Driven Adaptive NPC Behavior in a Meme-Based Platformer Game for Android Mobile**  
Encep Sayid Amrulloh, Rio Andriyat Krisdiawan, Iwan Lesmana - Page - 123-132

**Design and Implementation of a Performance Dashboard for Public Relations at Telkom University**  
Fеды Deа Reskyadita, Hani Nurrahmi, Daris Maulana. - Page-133-143

**Design And Development of The Cakrawala Project Web Application: Case Study of The Information Technology Division**  
Tirta Subagja, Mamay Syani. - Page - 144-148

e-issn : 2614-5405

p-ISSN : 1858-3911



9 772614 540005

9 771858 391114

# JURNAL NUANSA INFORMATIKA

**Volume 19. Number 2, July 2025**

- Pembina : Dr. Dikdik Harjadi, M.Si. (Rektor Universitas Kuningan)  
Pengarah : Dr. Anna Fitri Hindriana, M.Si (Warek I Universitas Kuningan)  
Penanggung Jawab : Tito Sugiharto, M.Eng. (Dekan Fakultas Ilmu Komputer)  
Editor in Chief : Rio Andriyat Krisdiawan, M.Kom.  
Editor : Rio Andriyat Krisdiawan, M.Kom  
Section Editor :  
1. Endra Suseno, M.Kom (FKOM UNIKU)  
2. Dyah Puteriawati, M.Kom (FKOM UNIKU)  
Reviewer :  
1. Dr. Tata Sutabri, S.Kom., MMSI., MKM – Universitas Bina Darma - Indonesia.  
2. Dr. Nur Alamsyah, M.Kom - Universitas Informatika Dan Bisnis Indonesia (UNIBI) - Indonesia.  
3. Dr. Oman Komarudin, S.Si., M.Kom - Universitas Singaperbangsa Karawang (UNSIKA) – Indonesia.  
4. E. Haodudin Nurkifli, S.T., M.Cs., Ph.D. - Universitas Singaperbangsa Karawang (UNSIKA) – Indonesia.  
5. Erlan Darmawan, S.Kom.,M.Si.,Ph.D - FKOM UNIKU, Indonesia.  
6. Fahmi Yusuf, M.MSI., Ph.D - FKOM UNIKU, Indonesia  
7. Taufik Ridwan, S.T., M.T. - Universitas Singaperbangsa Karawang (UNSIKA) – Indonesia.  
8. Budiman, S.T., M.Kom - Universitas Informatika Dan Bisnis Indonesia (UNIBI) - Indonesia.  
9. Mohamad Nurkamal Fauzan, S.T., M.T., SFPC - Universitas Logistik dan Bisnis Internasional (ULBI) – Indonesia.  
10. Sugeng Supriyadi,M.Kom - Universitas Nurtanio - Indonesia  
11. Tito Sugiharto, S.Kom., M.Eng - FKOM UNIKU, Indonesia.  
12. Rio Andriyat Krisdiawan, M.Kom - FKOM UNIKU, Indonesia.  
13. Syafrial Fachri Pane, S.T., MTI - Universitas Logistik dan Bisnis Internasional (ULBI) – Indonesia.  
Proofing & Language Editor :  
1. Roni Nursyamsu, M.Pd (FKOM UNIKU)  
2. Nida Amalia Asikin, M.Pd (FKOM UNIKU)  
Alamat Redaksi : Kampus 2 Fakultas Ilmu Komputer Universitas Kuningan  
Jl. Pramuka No.67, Purwawinangun, Kec. Kuningan, Kabupaten Kuningan, Jawa Barat 45512  
Telp/Fax (0232) 875097  
Email : [nuansa.informatika@uniku.ac.id](mailto:nuansa.informatika@uniku.ac.id)  
Website : <https://journal.uniku.ac.id/index.php/ilkom>

# JURNAL NUANSA INFORMATIKA

## ***Nuansa Informatika (Journal on Information and Technology)***

*Nuansa Informatika is a scientific journal published by the Faculty of Computer Science Kuningan University with a frequency published twice a year. Nuansa Informatika is a peer-reviewed journal on Information and Technology for communication media academics, experts and practitioners of Information Technology in pouring ideas of thought in the field of Information Technology. Nuansa Informatika is a journal on Information and Technology covering all branches of IT and sub-disciplines including Algorithms, system design, networks, games, IoT, Software engineering, Mobile applications, and others.*

*The Journal of Information Technology Shades is published in print and online. Issuance in print Since September 19, 2005 with [p-ISSN :1858-3911](https://doi.org/10.25134/nuansa). SK no.0003.745/JI.3.02/SK.ISSN/2005. Start in September 19 2005. [e-ISSN : 2614-5405](https://doi.org/10.25134/nuansa). SK no. 0005.26145405/JI.3.1/SK.ISSN/2018.01. Start in January 25 201. And DOI : <https://doi.org/10.25134/nuansa>.*

*Since Juni 6th, 2020, **Nuansa Informatika : Jurnal Teknologi dan Informasi** is officially accredited by The Ministry of Research, Technology and Higher Education, Republic of Indonesia (Kemenristekdikti RI) with **SINTA 5**.*

*Organized by Faculty of Computer Science, Kuningan University, Indonesia.*

*Website : OJS 2 : <https://journal.uniku.ac.id/index.php/ilkom>*

*OJS 3: <https://journal.fkom.uniku.ac.id/ilkom>*

*Email : [nuansa.informatika@uniku.ac.id](mailto:nuansa.informatika@uniku.ac.id)*

*Address : Jalan Cut Nyak Dhien No.36A Kuningan, Jawa Barat, Indonesia.*

# JURNAL NUANSA INFORMATIKA

Volume 19 Number 2, July 2025

**IoT and Water Consumption Forecasting: A Green Accounting Study at a Coffee Shop in Cimahi**

*Mohamad Nurkamal Fauzan, Riani Tanjung*  
Page - 1-10

**Feature Extraction Using Markov Random Field (MRF) On Improving CNN Classification Results For Alzheimer's Disease Diagnosis**

*Saeful Bahri, Miftah Farid Adiwisastara, Taufik Hidayatulloh*  
Page - 11-17

**Assessing the Role of E-Govqual in Improving Public Satisfaction with the MyPertamina Application in Setiabudi District**

*Yulian Surya Saputra, Sucitra Sahara, Rizqi Agung Permana*  
Page - 18-24

**A Data-Driven Approach to Comparative Evaluation of Regression Models for Accurate House Price Prediction**

*Tiara Permata Hati, Budiman Budiman, Nur Alamsyah*  
Page - 25-34

**Development Of A 2d Educational Animal Card Game Using Unity**

*Alma Sulaiman, Mamay Syani*  
Page - 35-48

**The Implementation of User Satisfaction Based Shortest Job First Algorithm Efficiency in Queuing System**

*Darsanto*  
Page - 49-57

**Predicting Basic Shipping Tariff Using Machine Learning**

*Dani Ferdinan, Nisa Hanum Harani, M. Yusril Helmi Setyawan*  
Page - 58-66

**Development Of Online Attendance Application With Time And Location Validation At Satpol Pp South Tangerang**

*Syamsu Hidayat, Muhamad Rifki Ramadhan, Joko Susilo.*  
Page - 67-72

**Web-Based Assignment Management System Application Design to Increase Personnel Productivity and Effectiveness in Product Documentation Division**

*Eryan Ahmad Firdaus, Seftian Candra Pratama, Rochedi Idul Adha*  
Page - 73-80

**Emoji-Based Sentiment Classification Using Ensemble Learning with Cross-Validation: A Lightweight Approach for Social Media Analysis**

*Nur Alamsyah, Gunthur Bayu Wibisono, Titan Parama Yoga, Budiman, Acep Hendra*  
Page - 81-87

**A Bidirectional GRU Approach with Hyperparameter Optimization for Sentiment Classification in Game Reviews**

*Nur Alamsyah, Titan Parama Yoga, Budiman, Imannudin Akbar, Acep Hendra, Alif Januantara Prima*  
Page - 88-95

# JURNAL NUANSA INFORMATIKA

## **The Implementation of the TOPSIS Method in Determining Stunting Toddlers**

*Umi Khultsum, F. Lia Dwi Cahyanti, Elly Firasari*

*Page - 96-104*

## **Predicting the Happiness Index Based on the HDI Indicator in Indonesia Using the Ensemble Learning Approach**

*Syafrial Fachri Pane, Rofi Nafiis Zain, Iwan Setiawan, Virdiandry Putratama*

*Page - 105-114*

## **Performance and Efficiency Testing Analysis of Database Systems in Academic Information Systems**

*Utami Kusuma Dewi, Ryan Lingga Wicaksono, Mahendra Dwifebri Purbolaksono, Villy Satria*

*Page - 115-122*

## **Fuzzy-Driven Adaptive NPC Behavior in a Meme-Based Platformer Game for Android Mobile**

*Encep Sayid Amrulloh, Rio Andriyat Krisdiawan, Iwan Lesmana*

*Page - 123-132*

## **Design and Implementation of a Performance Dashboard for Public Relations at Telkom University**

*Feddy Dea Reskyadita, Hani Nurrahmi, Daris Maulana.*

*Page-133-143*

## **Design And Development of The Cakrawala Project Web Application: Case Study of The Information Technology Division**

*Tirta Subagja, Mamay Syani.*

*Page - 144-148*

# IoT and Water Consumption Forecasting: A Green Accounting Study at a Coffee Shop in Cimahi

Mohamad Nurkamal Fauzan<sup>\*1</sup>, Riani Tanjung<sup>2</sup>

<sup>1</sup>D4 Teknik Informatika, Universitas Logistik dan Bisnis Internasional

<sup>2</sup>D3 Akuntansi, Universitas Logistik dan Bisnis Internasional

E-mail: <sup>\*1</sup>[m.nurkamal.f@ulbi.ac.id](mailto:m.nurkamal.f@ulbi.ac.id), <sup>2</sup>[rianitanjung@ulbi.ac.id](mailto:rianitanjung@ulbi.ac.id)

## Abstract

*Uncontrolled water consumption is a serious challenge, especially in small businesses like coffee shops. Excessive water use can lead to waste and financial losses. To address this issue, IoT (Internet of Things) technology and data analysis are applied to monitor and predict water consumption. In this study, predictive models such as Random Forest, XGBoost, and LSTM are used to analyze water consumption data. The results show that Random Forest has the best performance with the lowest prediction error and the highest R-squared value, indicating this model's capability to explain nearly all the variance in water consumption data. Random Forest and XGBoost perform well as they can handle data with non linear features and complex interactions, while LSTM's lower performance is likely due to limited data and suboptimal hyperparameter tuning. The implementation of green accounting in this system enables effective tracking of water consumption costs. Suggested improvements include further exploration of LSTM hyperparameters, the use of ensemble techniques, and cost sensitivity analysis for water-saving policy decisions. This model is expected to provide an effective water saving solution for coffee shop owners.*

**Keywords**— water consumption, IoT, Random Forest, green accounting, coffee shop

Submission: February 18, 2025

Accepted: June 16, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.366>

## 1. INTRODUCTION

Uncontrolled water consumption is one of the serious challenges faced by the world today. Population growth, urbanization, and increasing economic activities have led to a rising demand for water, while the availability of clean water continues to decline. According to a report by the World Resources Institute (WRI), approximately 25% of the global population lives in regions experiencing extremely high water stress, where water demand significantly exceeds supply [1]. This issue not only affects the sustainability of natural resources but also increases the risk of social conflicts due to competition over water resources.

In small businesses such as coffee shops, excessive water consumption often goes unnoticed. However, these businesses contribute to water waste that could be managed more efficiently to reduce environmental impact. Poor water management can lead to significant financial losses. For instance, if water consumption exceeds the necessary amount by 20%, with an average water cost of IDR 5,000 per cubic meter, losses could range from approximately IDR 200,000 to IDR 500,000 per month, or IDR 2,400,000 to IDR 6,000,000 per year, depending on the scale of operations.

Therefore, a new approach that integrates IoT (Internet of Things)

technology and data analytics is needed to monitor and predict water consumption. According to a study by Czajkowski [2], implementing IoT in water management has been proven to reduce water consumption by up to 15% through real time monitoring and optimization, enabling the adoption of more sustainable management policies.

The urgency of managing water consumption wisely is not only to ensure the sustainability of water resources for future generations but also to reduce the economic and environmental impacts of excessive water use. This case study will explore the implementation of IoT technology and machine learning for forecasting water consumption in a coffee shop in Cimahi, as part of a green accounting initiative to promote water use efficiency and environmental sustainability [3].

#### Problem Formulation:

1. How is the data collected?
2. What forecasting methods are used to predict water consumption?
3. How can green accounting reports be implemented?

## 2. METHOD

The research method is structured to address the problem formulation. We propose a prototype design with an architecture that integrates IoT, machine learning, and green accounting to answer the first research question regarding data collection.

In the input stage, an ultrasonic sensor will measure the water level in the storage tank. An ESP32 based microcontroller [4] will be embedded in this system. The device will transmit data using the HTTP protocol [5]. The collected data will be stored in a database and then processed for forecasting using machine learning. The processed results will be sent to a green accounting application (Figure 1). This study

contributes by presenting an IoT-based water consumption architecture model.

We propose a water consumption architecture with a case study in a coffee shop. Data collection is conducted through an IoT-based-system (Figure 2). The ESP32 microcontroller is equipped with distance, temperature, humidity, and rain sensors to process and transmit data to the server daily at 6:00 PM WIB.

The dataset consists of a time series spanning 219 days, starting from January 1, 2023. Each incoming data entry is recorded under the 'date' feature, which serves as a timestamp for when the data is stored on the server. The 'date' feature is then processed to generate new features, such as 'day\_of\_week' and 'is\_holiday.' The 'day\_of\_week' feature represents the numerical day of the week (e.g., 0 is Monday, 6 is Sunday). Meanwhile, the 'is\_holiday' feature marks holidays (weekends) with a value of 1 if 'day\_of\_week'  $\geq 5$ .

The 'temperature' feature represents the average temperature of the day in degrees Celsius. The 'is\_rain' feature records whether it rained on a given day, assigning a value of 1 if rain occurred and 0 otherwise. The 'humidity' feature captures the average humidity percentage, ranging from 0 to 100%.

The 'water\_consumption' feature determines water usage based on readings from the HCSR04 ultrasonic distance sensor. Table 1 presents the data successfully transmitted from the ESP32 using the HTTP protocol.

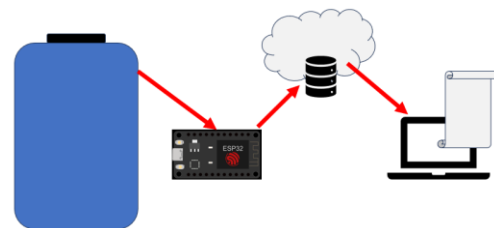


Figure 1. IoT Water Consumption Architecture

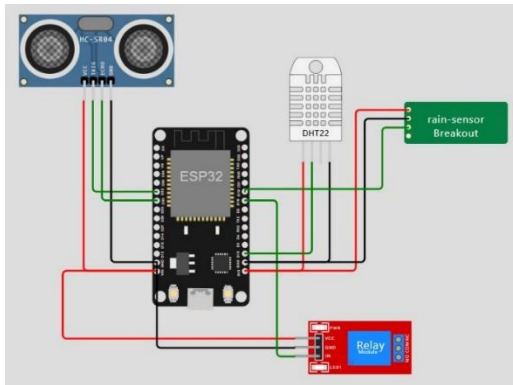


Figure 2. IoT Water Consumption Device Scheme

Table 1 Main Data

| Date       | Water_Consumption | Temperature | is_rain | Humidity  | day_of_week | is_holiday |
|------------|-------------------|-------------|---------|-----------|-------------|------------|
| 2023-01-01 | 0                 | 28.745401   | 0       | 77.939727 | 6           | 1          |
| 2023-01-02 | 190               | 30.986585   | 0       | 52.019753 | 0           | 0          |
| 2023-01-03 | 182               | 29.592489   | 0       | 51.429007 | 1           | 0          |

Table 1 presents the data successfully transmitted from the ESP32 using the HTTP protocol.

To address the second research question regarding forecasting model development, the collected data undergoes various techniques, such as:

### 2.1.Feature Engineering

We applied feature engineering techniques [6], [7], including adding the 'month' feature to capture monthly cyclical patterns in the data. Additionally, we introduced the 'temp\_humidity\_interaction' feature by multiplying temperature and humidity to capture their interaction effects [8].

$$\text{temp\_humidity\_interaction} = \text{temperature} \times \text{humidity}$$

We added the 'heat\_index' feature to estimate the heat index, considering the

combined effects of temperature and humidity on thermal comfort.

Additionally, we applied the One Hot Encoding method [9], [10] to convert categorical features such as 'day\_of\_week' and 'month' into dummy variables. For the 'day\_of\_week' feature, each day of the week was transformed into a dummy variable (0 or 1) for each day category.

### 2.2.Data Normalization

We performed data normalization [11], [12] to facilitate the model training process. The steps taken include:

1. Scaler Initialization, creating a MinMaxScaler object for both features (scaler) and the target variable (scaler\_y).
2. Feature Normalization, applying Min-Max normalization to

numerical features to scale values between 0 and 1.

3. Target Normalization, separately normalizing the target variable Water\_Consumption to ensure consistency in model training.

### 2.3.Adding Lag Features and Moving Averages

1. Lag Features, we added lag features [13] for Water\_Consumption, ranging from 1 to 7 days prior, using the shift (lag) function. These features rely on historical Water\_Consumption data.
2. 7-Day Moving Average, we introduced a 7-day moving average feature for Water\_Consumption, calculated using a rolling mean over the past 7 days. This feature also depends on historical Water\_Consumption data.
3. Handling NaN Values, we removed rows containing NaN values generated from the shift and rolling operations to maintain data integrity.

### 2.4.Splitting Data into Features and Target

The feature set (X) consists of all features except Water\_Consumption, while the target variable (y) is derived from the Water\_Consumption feature.

### 2.5.Splitting Training and Testing Data

We used the train\_test\_split function to divide the data into training and testing sets with a 60:40 ratio. To preserve the temporal order in the time series data, we disabled shuffling by setting shuffle=False.

### 2.6.Model Training

We utilized three machine learning models: Random Forest, XGBoost, and Long Short-Term Memory (LSTM), based on their suitability for capturing different patterns in time-series data and their proven performance in forecasting tasks.

XGBoost [14] was selected for its ability to handle structured data efficiently and its robustness to overfitting through regularization. We performed hyperparameter tuning [15] using Optuna [16], an efficient optimization framework. The objective function [17] used by Optuna minimized the Mean Squared Error (MSE) [18] on the test set. The model was trained using XGBRegressor with the optimal parameters determined through the tuning process.

Long Short Term Memory (LSTM) [19], was included to model temporal dependencies and patterns in sequential water consumption data that traditional models may overlook. We set the sequence length to 15, meaning the model considers the past 15 time steps to predict the next value. The model was built using a sequential architecture with specified parameters. We trained the model for 50 epochs, using a batch size of 16, and applied a validation split of 0.2, meaning 20% of the training data was used for validation.

We trained the Random Forest model [20] using Optuna for hyperparameter tuning. After training, we performed the following steps:

1. Inverse Transformation of Predictions, converting predicted values back to their original scale.
2. Using scaler\_y.inverse\_transform, Restoring both predicted and actual values to their pre-normalization scale.
3. Prediction Visualization
4. Model Evaluation. Assessing performance using Mean Squared Error (MSE) and R-squared (Coefficient of Determination).

To address the third research question, the green accounting process was conducted, which includes:

Utilization of Data in Green Accounting [21]

1. Integration of Forecasting in Green Accounting. The water consumption forecasting results are integrated into green accounting through environmental financial reports [22]. These reports include records of water consumption, associated costs, and any savings or penalties applied due to efficiency or wastage.
2. Cost and Savings Analysis. Actual water consumption is compared with the forecasted values to determine efficiency or wastage. Any discrepancies between actual and predicted consumption are calculated to assess the relevant environmental costs [23].

1. Water Cost. The cost of water usage is calculated by multiplying actual consumption by the rate per liter.
2. Savings from Efficiency. Savings are determined when actual water consumption is lower than the forecasted amount, reflecting the cost successfully reduced.
3. Environmental Cost Due to Excessive Consumption. Additional costs or penalties are calculated for water usage that exceeds the target, serving as a measure to internalize the environmental impact of overconsumption

### 3. RESULTS AND DISCUSSION

Environmental Cost and Savings Calculation.

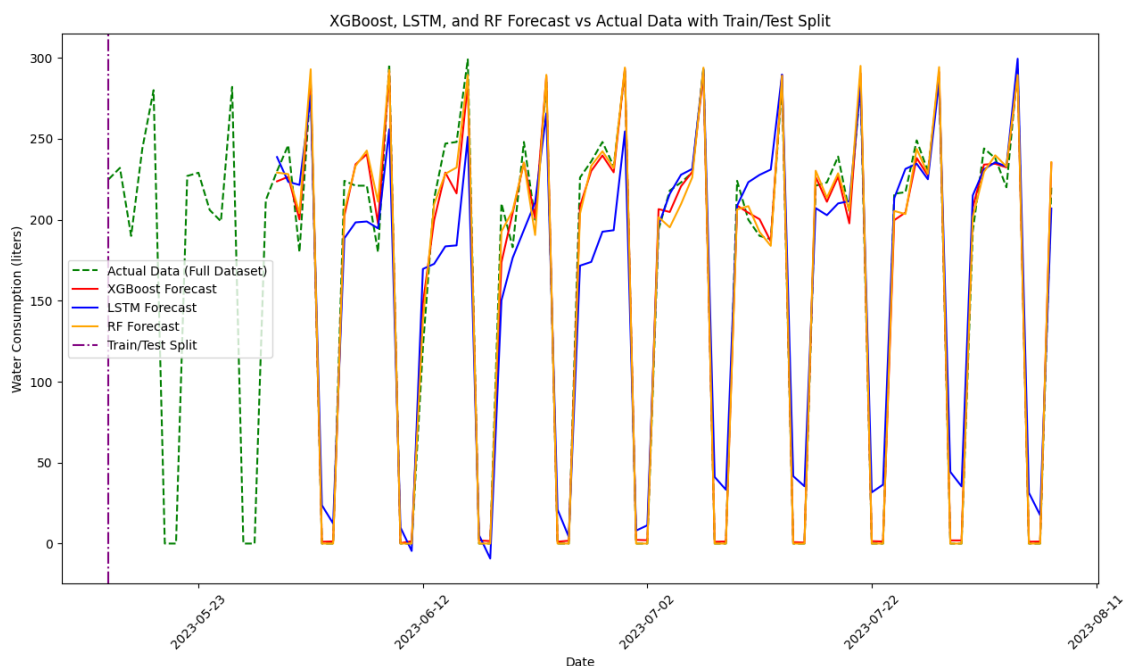


Figure 1 Forecast: XGBost, LSTM, Random Forest

In Figure 3, we can observe a comparison between the predicted water consumption from three different models (XGBoost, LSTM, and Random Forest) against the actual data. The dashed green

line represents the actual water consumption.

There is a significant variation in water consumption, with periodic increases and

decreases, particularly with lower consumption on holidays. Among the models, Random Forest demonstrates the best predictive performance, accurately capturing fluctuations in water usage during both high and low consumption periods.

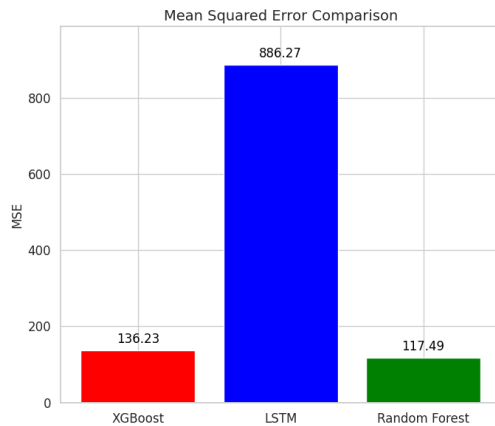


Figure 2 MSE Comparison

Figure 4 presents a comparison of the Mean Squared Error (MSE) for the three predictive models: XGBoost, LSTM, and Random Forest. MSE is an evaluation metric used to measure how closely the model's predictions align with actual values, where a lower MSE indicates better performance.

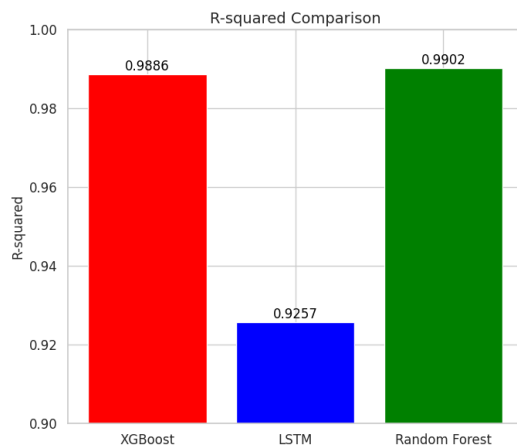


Figure 3 R-squared comparison

Figure 5 presents a comparison of R-squared values for the three predictive models: XGBoost, LSTM, and Random Forest. R-squared is a statistical metric that indicates how well a model explains the variance in actual data. A value close to 1 suggests a highly accurate predictive model.

Among the models, Random Forest achieves the highest R-squared value of 0.9902, meaning it explains approximately 99% of the variance in actual water consumption data. This result highlights Random Forest as the most effective model for capturing and modeling water consumption patterns.

An example of Random Forest input on 2024-01-10 is shown in Table 2.

Table 2 Forecast Random Forest

| Parameter   | Value      |
|-------------|------------|
| Date        | 2024-01-10 |
| Temperature | 25.00 °C   |
| is rain     | 0          |
| Humidity    | 48.00 %    |

Predicted Water Consumption for 2024-01-10: 197.2634859981626 liters

Example of Green Accounting Implementation: Green Accounting is the process of incorporating environmental factors into accounting and economic decision-making. The forecasting results can be applied to green accounting through bookkeeping (Table 3) or environmental reporting (Table 4). The primary focus of green accounting is to integrate costs and benefits into financial and sustainability assessments.

*Table 3 Forecasted Water Consumption in Monthly Environmental Management Report*

| Period        | Forecasted Water Consumption (Liters) | Actual Water Consumption (Liters) | Difference (Liters) | Status       |
|---------------|---------------------------------------|-----------------------------------|---------------------|--------------|
| January 2023  | 12,000                                | 11,500                            | -500                | Efficiency   |
| February 2023 | 15,000                                | 16,200                            | +1,200              | Wastefulness |
| March 2023    | 13,500                                | 13,300                            | -200                | Efficiency   |
| April 2023    | 11,500                                | 11,500                            | 0                   | On Target    |
| May 2023      | 10,800                                | 10,500                            | -300                | Efficiency   |

This forecasting compares actual water consumption with the projected estimates. Efficiency or wastefulness is identified by analyzing the difference between actual and forecasted consumption. The insights gained

from this comparison are then used to enhance water management strategies for future periods.

*Table 4 Water Usage Costs and Environmental Savings*

| Periode       | Water Consumption (Liters) | Water Cost (Rp) | Savings from Efficiency (Rp) | Environmental Cost Due to Wastefulness (Rp) |
|---------------|----------------------------|-----------------|------------------------------|---------------------------------------------|
| January 2023  | 11,500                     | 1,150,000       | 50,000                       | -                                           |
| February 2023 | 16,200                     | 1,620,000       | -                            | 120,000                                     |
| March 2023    | 13,300                     | 1,330,000       | 20,000                       | -                                           |
| April 2023    | 11,500                     | 1,150,000       | -                            | -                                           |
| May 2023      | 10,500                     | 1,050,000       | 30,000                       | -                                           |

Notes:

1. Water Cost is calculated by multiplying actual consumption by the rate per liter.

2. Savings from Efficiency reflect cost reductions when actual consumption is lower than the forecasted amount.

3. Environmental Cost Due to Wastefulness records additional costs or

penalties for water usage that exceeds the target or surpasses sustainable capacity.

#### 4. CONCLUSION

The IoT device successfully transmitted sensor data readings. Among the three models tested, Random Forest achieved the lowest MSE of 117.4869, indicating the smallest prediction error.

R-squared ( $R^2$ ) measures the proportion of variance in the target variable that can be explained by the model. An  $R^2$  value close to 1 suggests that the model effectively explains data variability. Random Forest achieved the highest  $R^2$  of 0.9902, indicating that it explains 99.02% of the variance in water consumption data.

Random Forest slightly outperformed XGBoost. The performance difference may be due to how each algorithm handles data and feature interactions. Both Random Forest and XGBoost are decision tree-based models capable of handling non-linear relationships and complex feature interactions [14], [20] which appear well-suited for water consumption data.

LSTM, designed for time series data and long-term dependencies, performed worse than the tree-based models. Possible reasons include: limited data availability [24], suboptimal hyperparameter tuning and lack of strong time-series patterns in the data

While Random Forest and XGBoost both achieved  $R^2$  values close to 1, which could indicate potential overfitting, the low MSE on the test data suggests that these models generalize well. Random Forest emerged as the best-performing model, with the lowest MSE and highest  $R^2$ .

The green accounting framework was effectively designed to incorporate forecasting and track actual consumption. This allows coffee shop owners to take appropriate actions, including monitoring monthly water usage and cost management.

#### 5. RECOMMENDATIONS

Based on the experiment, further exploration of hyperparameter tuning for the LSTM model is recommended to improve its performance. Ensemble techniques can also be utilized to enhance model accuracy and stability. As more data becomes available over time, retraining new models can help assess their performance and adaptability. Regarding green accounting, conducting a sensitivity analysis on water costs can provide deeper insights into how fluctuations in water consumption impact expenses. This would enable more effective decision making for water conservation policies and financial planning.

#### REFERENCES

- [1] WRI, "Updated Global Water Risk Atlas Reveals Top Water- Stressed Countries and States," *Wri*, 2019.
- [2] A. Czajkowski *et al.*, "Global water crisis: Concept of a new interactive shower panel based on iot and cloud computing for rational water consumption," *Applied Sciences (Switzerland)*, vol. 11, no. 9, 2021, doi: 10.3390/app11094081.
- [3] C. Sukmadilaga, S. Winarningsih, I. Yudianto, T. U. Lestari, and E. K. Ghani, "Does Green Accounting Affect Firm Value? Evidence from ASEAN Countries," *International Journal of Energy Economics and Policy*, vol. 13, no. 2, 2023, doi: 10.32479/ijeep.14071.
- [4] E. Systems, "ESP32 Series Datasheet," 2021. [Online]. Available: [https://www.espressif.com/sites/default/files/documentation/esp32\\_datasheet\\_en.pdf](https://www.espressif.com/sites/default/files/documentation/esp32_datasheet_en.pdf)
- [5] M. J. A. Baig, M. T. Iqbal, M. Jamil, and J. Khan, "Design and implementation of an open-Source IoT and blockchain-based peer-to-

- peer energy trading platform using ESP32-S2, Node-Red and, MQTT protocol,” *Energy Reports*, vol. 7, 2021, doi: 10.1016/j.egyr.2021.08.190.
- [6] H. Dai, Q. Xiao, N. Yan, X. Xu, and T. Tong, “Item-level Forecasting for E-commerce Demand with High-dimensional Data Using a Two-stage Feature Selection Algorithm,” *J Syst Sci Syst Eng*, vol. 31, no. 2, 2022, doi: 10.1007/s11518-022-5520-1.
- [7] T. H. Kim *et al.*, “Simultaneous feature engineering and interpretation: Forecasting harmful algal blooms using a deep learning approach,” *Water Res*, vol. 215, 2022, doi: 10.1016/j.watres.2022.118289.
- [8] T. Wang, J. Chen, J. Vaughan, and V. N. Nair, “Interpretable Feature Engineering for Time Series Predictors using Attention Networks,” *SSRN Electronic Journal*, 2022, doi: 10.2139/ssrn.4117900.
- [9] V. Sarveswararao, V. Ravi, and Y. Vivek, “ATM cash demand forecasting in an Indian bank with chaos and hybrid deep learning networks,” *Expert Syst Appl*, vol. 211, 2023, doi: 10.1016/j.eswa.2022.118645.
- [10] M. Anjaneyulu and M. Kubendiran, “Short Term Traffic Flow Prediction Using Hybrid Deep Learning,” *Computers, Materials and Continua*, vol. 75, no. 1, 2023, doi: 10.32604/cmc.2023.035056.
- [11] F. T. Lima and V. M. A. Souza, “A Large Comparison of Normalization Methods on Time Series,” *Big Data Research*, vol. 34, 2023, doi: 10.1016/j.bdr.2023.100407.
- [12] J. Yu and K. Spiliopoulos, “NORMALIZATION EFFECTS ON DEEP NEURAL NETWORKS,” *Foundations of Data Science*, vol. 5, no. 3, 2023, doi: 10.3934/fods.2023004.
- [13] O. Surakhi *et al.*, “Time-lag selection for time-series forecasting using neural network and heuristic algorithm,” *Electronics (Switzerland)*, vol. 10, no. 20, 2021, doi: 10.3390/electronics10202518.
- [14] T. Chen and C. Guestrin, “XGBoost: A scalable tree boosting system,” in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016, pp. 785–794.
- [15] R. K. Vogeti, B. R. Mishra, and K. S. Raju, “Machine learning algorithms for streamflow forecasting of Lower Godavari Basin,” *H2Open Journal*, vol. 5, no. 4, 2022, doi: 10.2166/h2oj.2022.240.
- [16] T. Akiba, S. Sano, T. Yanase, T. Ohta, and M. Koyama, “Optuna: A Next-Generation Hyperparameter Optimization Framework,” in *The 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, 2019, pp. 2623–2631.
- [17] H. Ismail Fawaz, G. Forestier, J. Weber, L. Idoumghar, and P. A. Muller, “Deep learning for time series classification: a review,” *Data Min Knowl Discov*, vol. 33, no. 4, 2019, doi: 10.1007/s10618-019-00619-1.
- [18] A. C. R. Klaar, S. F. Stefenon, L. O. Seman, V. C. Mariani, and L. dos S. Coelho, “Optimized EWT-Seq2Seq-LSTM with Attention Mechanism to Insulators Fault Prediction,” *Sensors*,

- vol. 23, no. 6, 2023, doi: 10.3390/s23063202.
- [19] S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural Comput*, vol. 9, no. 8, pp. 1735–1780, 1997.
- [20] L. Breiman, "Random forests," *Mach Learn*, vol. 45, no. 1, 2001, doi: 10.1023/A:1010933404324.
- [21] H. A. Riyadh, M. A. Al-Shmam, H. H. Huang, B. Gunawan, and S. A. Alfaiza, "The analysis of green accounting cost impact on corporations financial performance," *International Journal of Energy Economics and Policy*, vol. 10, no. 6, 2020, doi: 10.32479/ijeep.9238.
- [22] A. M. Islam and C. Deegan, "Motivations for an organisation within a developing country to report social responsibility information: Evidence from Bangladesh," *Accounting, Auditing and Accountability Journal*, vol. 21, no. 6, 2008, doi: 10.1108/09513570810893272.
- [23] M. Yassin, A. Asfaw, V. Speight, and J. D. Shucksmith, "Evaluation of Data-Driven and Process-Based Real-Time Flow Forecasting Techniques for Informing Operation of Surface Water Abstraction," *J Water Resour Plan Manag*, vol. 147, no. 7, 2021, doi: 10.1061/(asce)wr.1943-5452.0001397.
- [24] A. Althnian *et al.*, "Impact of dataset size on classification performance: An empirical evaluation in the medical domain," *Applied Sciences (Switzerland)*, vol. 11, no. 2, 2021, doi: 10.3390/app11020796.

# Feature Extraction Using Markov Random Field (MRF) On Improving CNN Classification Results For Alzheimer's Disease Diagnosis

Saeful Bahri<sup>1</sup>, Miftah Farid Adiwisastra<sup>\*2</sup>, Taufik Hidayatulloh<sup>3</sup>  
<sup>1,2,3</sup>Universitas Bina Sarana Informatika, Indonesia

E-mail: <sup>1</sup>[saeful.sel@bsi.ac.id](mailto:saeful.sel@bsi.ac.id), <sup>\*2</sup>[miftah.mow@bsi.ac.id](mailto:miftah.mow@bsi.ac.id), <sup>3</sup>[taufik.tho@bsi.ac.id](mailto:taufik.tho@bsi.ac.id)

## Abstract

*Alzheimer's disease is a degenerative neurological disorder that significantly affects patients' cognitive functions and social lives. It is a leading form of dementia, characterized by the progressive death of brain cells. One widely adopted diagnostic approach involves magnetic resonance imaging (MRI) to evaluate brain structures. Recent advances in machine learning have enabled automated image analysis, with Convolutional Neural Networks (CNNs) commonly used for image feature extraction and classification. However, CNNs face a major limitation in maintaining consistency during image segmentation, which results in reduced classification accuracy. This issue arises from CNNs' limited ability to preserve local pixel-level consistency during feature extraction. To address this, we propose integrating a Markov Random Field (MRF)-based layer into the CNN architecture, which has been shown to enhance segmentation consistency. This study utilizes publicly available MRI datasets of Alzheimer's patients and employs a k-fold cross-validation scheme for evaluation. The results show that the CNN-MRF model improves classification accuracy to 63%, compared to 61% with the standard CNN. Furthermore, the loss value is reduced from 0.80 to 0.74. Although the improvement in accuracy is incremental, a paired t-test confirms that the difference is statistically significant ( $p < 0.05$ ). This method has proven effective in enhancing the reliability of image-based diagnostic systems for early detection of Alzheimer's disease.*

**Keywords**— Alzheimer's, CNN, Markov Random Filed

Submission: March 11, 2025

Accepted: June 26, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.369>

## 1. INTRODUCTION

The brain is one of the most important and most complex organs in the human body. In addition, the brain has a vital function, namely processing information, problem solving, thinking, decision making, imagining and remembering [1], [2]. Stored memories can be processed into information that can affect a person's character or identity. Memory loss caused by dementia is one of the most frightening things. Alzheimer's is one of the most common diseases of dementia. As a person ages, Alzheimer's becomes a scary specter for most people. Alzheimer's gradually kills cells in the brain which causes someone with Alzheimer's to lose all memories, the ability to recognize family members, the ability to follow simple instructions, even Alzheimer's sufferers can lose the ability to swallow, cough and the most severe lose the ability to breathe [3]. [4], [5] It is estimated that around 50 million people in the world suffer from dementia with treatment costs equivalent to the world's 18th

largest economy [2]. then in 2050 dementia sufferers are estimated to reach 152 cases. Diagnosing Alzheimer's dementia becomes complicated and difficult because there are similarities with aging or vascular dementia[6], [7], [8]. Accurate diagnosis of Alzheimer's dementia plays a major role in prevention and treatment through monitoring its development. [9], [10], [11], [12]. Some research in the process of detecting Alzheimer's disease is the use of MRI brain images, this tool can measure the number of cells in the brain, in addition the MRI tool is also able to measure the structure of the parietal lobe [13].

Imaging plays an important role in many scientific fields, in this case medical imaging has become a powerful tool for understanding brain function, imaging such as Magnetic Resonance (MRI)[14], [15], has been used in the medical diagnostic process in evaluating brain structure and brain function in identifying signs and symptoms of Alzheimer's [3]. With the

development of technology and the growth of imaging data, machine learning (ML) and deep learning (DL) play an important role in creating information extraction results and making prediction results more accurate [16].

Several machine-learning methods have been widely applied in the Alzheimer's classification process [17], [18] and have shown quite good performance [17]. In general, CNN has a multi-layer architecture that has advantages in the deep learning process that resembles the human brain and is able to represent images well [1]. However, CNN has problems in the image separation process so that there is the potential for segmentation errors in an image that have an impact on poor classification results [19].

In this study we propose the addition of MRF as an additional layer in the CNN architecture. The addition is intended to maintain local spatial consistency in the segmentation process, by considering the relationship between adjacent pixels, from the addition of this layer it is expected that the results of brain image feature extraction will be more accurate, so that in the end it can improve classification performance in detecting Alzheimer's disease. The main focus of this study is to improve the reliability of segmentation in the process of classifying brain MRI images with CNN, which leads to increased classification accuracy in supporting the development of a more accurate and consistent diagnostic automation system.

Therefore, this study we proposes the integration of a Markov Random Field (MRF) into the CNN framework to enhance segmentation consistency and improve classification performance in Alzheimer's diagnosis.

## 2. RESEARCH METHODS

In this section, we present some supporting literature in the analysis and interpretation of the results, including several theories that are relevant to the topic being studied by examining various research sources and supporting data in the research context.

### 2.1. Alzheimer

Alzheimer's is a neurodegenerative disease or a disease of neuron loss in brain tissue which is characterized by a significant decline in a person's cognitive abilities [20], this disease is one of the causes of dementia.

Alzheimer's is also one of the most common diseases suffered by people in the world [3] this is caused by several factors including aging, which is impossible to avoid. However, the severity of Alzheimer's can be minimized by conducting early screening, through brain imaging for proper care and treatment [13]

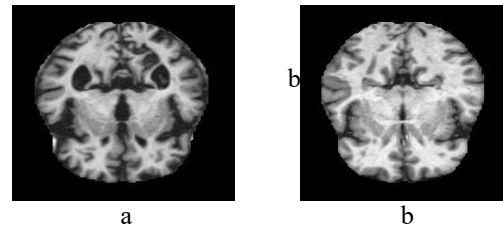


Figure 1. Comparison of brain images with Alzheimer a) and Normal Brain b)

The process of identifying brain images of Alzheimer's sufferers requires further in-depth research to determine the level of severity and damage to the cell system in the brain.[20]

### 2.2. CNN

Convolutional Neural network (CNN) is an algorithm that is commonly used for processing data in the form of images. This algorithm is quite effective in recognizing and detecting objects in an image. CNN is specifically designed to process data with a grid structure [21], [22]

CNN is an important component in the field of computer vision and image processing. CNN is an image processing algorithm that was first introduced in the 80s. However, in the 2010s, CNN was widely used in image processing along with advances in hardware technology[23], [24]. The key to CNN is the architecture of the Convolution layer that forms the structure of the CNN algorithm itself. CNN often contains a pooling layer to help handle overfitting and increase efficiency in the computing process. CNN also has a fully connected layer that functions to represent features from previous learning results.[25]

CNN has been widely used in several applications in fields such as Health, automated agriculture and so on. This ability of CNN in handling the process of automatic feature extraction from visual data and being able to handle large and complex problems places CNN as one of the powerful algorithms.

The architecture of CNN is similar to Artificial Neural Networks (ANN) in general[1]. CNN architecture consists of a collection of

neurons that form a layer. Each neuron has a weight, bias and activation function. This is different from ANN which describes a collection of one-dimensional neuron layers vertically, the layers in CNN are arranged in three dimensions, this causes the CNN architecture to have width, height and depth. In Figure 2. A general depiction of the CNN architecture can be seen.

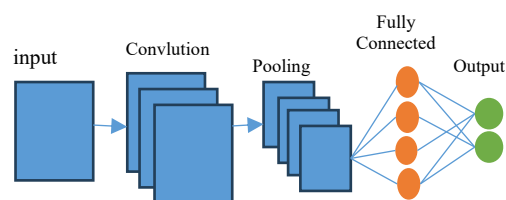


Figure 2. CNN architecture

In CNN, the classification process is carried out on the fully connected layer, which is an artificial neural network, where each neuron will be connected and arranged to form a layer. So that the input can be processed on the fully connected layer. The feature map generated from the convolution layer is converted into a long vector, then enters the fully connected layer to carry out the classification process.

The research stages were carried out to improve the classification results by adding new layers to the CNN architecture. In Figure 3, the new CNN-MRF architecture can be seen.

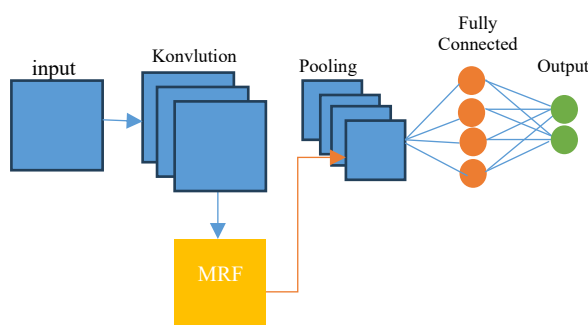


Figure 3. CNN-MRF architecture

The addition of layers to this convolution layer aims to obtain the best features by optimizing the feature extraction process in the image, to reduce inconsistencies in the feature extraction process, so that the classification results can be improved.

## 2.2. Markov Random Field

MRF is a multivariate stochastic process, capable of describing a random variable field

with its local spatial dependencies. Unlike Markov Chains, a category of univariate stochastic processes, random fields consist of multiple dimensions, usually interpreted as space or space-time components. The so-called Markov property, often investigated in time series applications, is extended to the local neighborhood of nodes in the corresponding undirected graph. Applications of MRF models include image processing, computer vision and geostatistics. In the context of image analysis, MRFs can be used for low-level image transformations and filtering, as well as for high-level tasks, such as object or texture classification [26]. MRF introduces a spatially dependent structure, where the label of a pixel depends on the labels of its neighboring pixels. This helps maintain spatial consistency in the segmented image, resulting in smoother and less noisy segmentation[27].

In this study, MRF is added to the CNN layer to segment brain images. In helping to refine the segmentation results by considering the relationship between adjacent pixels, this method allows the classification results to be better in identifying the diagnosis of Alzheimer's disease.

## 3. RESULTS

In this section we will present the research results systematically, the addition will start from the identification of the problem which is the main element in the formulation of the objectives and scope of the research. Next we will discuss the results we obtained from this study including the main findings that contribute to further understanding related to the research topic being studied. We will analyze these findings comprehensively in identifying the implications and significance of related research. Thus, this section aims to provide a comprehensive overview of the research that has been conducted. The details of the sub-chapters that will be discussed in this section are as follows.

## 4. DISCUSSION

Convolutional Neural Network (CNN) is one of the deep learning architectures specifically designed for hierarchical image processing. CNN is capable of automatically extracting features through convolutional layers that mimic the functionality of the human visual system—starting from simple features such as edges to more complex patterns like textures or structural forms of the brain.

However, CNN has a known limitation in preserving local pixel consistency during the segmentation process. This limitation becomes particularly critical when dealing with brain MRI images for Alzheimer's diagnosis, where the difference between stages such as *MildDemented* and *VeryMildDemented* is subtle and requires highly precise segmentation. As observed in our experimental results, CNN frequently failed to distinguish between the *MildDemented* and *VeryMildDemented* classes, resulting in misclassifications due to the visually similar features between these categories. This phenomenon is consistent with previous studies that have reported CNN's limited capability in distinguishing class boundaries with low contrast. To address this issue, we propose integrating a Markov Random Field (MRF) layer into the CNN architecture. MRF incorporates spatial pixel context, helping to reinforce segmentation by maintaining consistency among neighboring pixels that are likely to belong to the same class.

#### 4.1. Data Processing

The dataset used in this study was obtained from a public repository on Kaggle.com. The images were categorized into four primary classes with the following distributions: MildDemented 896 images, ModerateDemented 64 images, NonDemented 3200 images, VeryMildDemented 1792 images.

As observed, the dataset is highly imbalanced, which can lead to biased model performance favoring the majority classes. To address this issue, we implemented several preprocessing and balancing strategies, including Resizing all images to 128×128 pixels, Normalizing pixel values to a range of [0, 1].

Data augmentation applied specifically to the minority classes (*MildDemented* and *ModerateDemented*) using: Random rotation ( $\pm 15$  degrees), Horizontal flipping, Zooming and shifting. These techniques were applied to increase sample diversity and reduce overfitting. Subsequently, the dataset was split into training and testing sets using a stratified split with an 80:20 ratio to maintain proportional class distributions.

This setup was then used for training the model and evaluating its performance. Further details of the data processing pipeline can be seen in Figure 3.

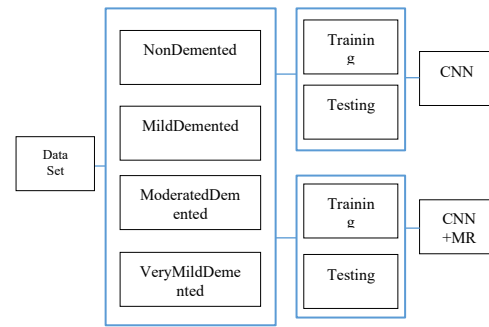


Figure 4. Data processing flow

#### 4.2. Research Results

In this section, we present the results of our study based on the analysis that has been conducted. Several key findings are presented in the form of tables and figures to provide a clearer understanding of the implications derived from the study.

##### 1. Test Results

In this section, we present the results of our study based on the analysis that has been conducted. Several key findings are presented in the form of tables and figures to provide a clearer understanding of the implications derived from the study. Both the CNN and CNN+MRF models were trained for 50 epochs with 35 steps per epoch. The performance evaluation was carried out using the following metrics: accuracy, loss, precision, recall, and F1-score. The main results are summarized in Table 1. In addition to the table, the results are also presented in graphical form, as shown in Figures 5 and 6, to further illustrate the performance comparison between the models.

| Model     | Akurasi(%) | Loss |
|-----------|------------|------|
| CNN       | 61%        | 0,80 |
| CNN + MRF | 63%        | 0.74 |

Tabel 1. Result comparsion CNN vs CNN+MRF

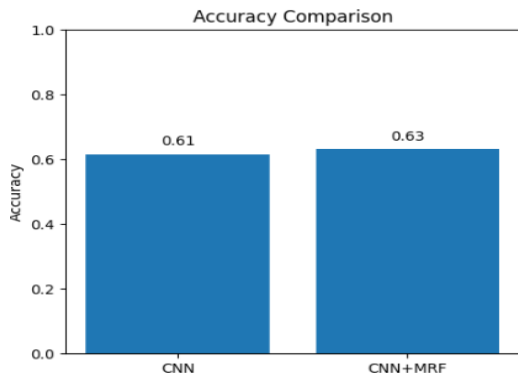


Figure 5. Comparison of CNN and CNN+MRF Algorithm Evaluation Results

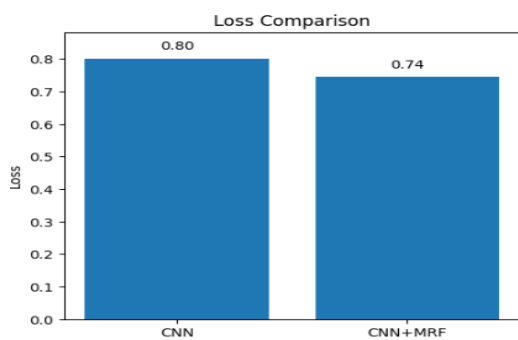


Figure 6. Loss comparison between CNN and CNN+MRF

Although the accuracy improvement is only 2%, the paired t-test indicates that the difference is statistically significant ( $p < 0.05$ ). Furthermore, the lower loss value suggests that the CNN+MRF model is more stable during the learning process. The detailed per-class evaluation results are presented in Table 2

| Class            | Precision (CNN)     | F1-Score (CNN)     |
|------------------|---------------------|--------------------|
| NonDemented      | 0.78                | 0.79               |
| VeryMildDemented | 0.63                | 0.61               |
| MildDemented     | 0.57                | 0.55               |
| ModerateDemented | 0.40                | 0.37               |
|                  | Precision (CNN+MRF) | F1-Score (CNN+MRF) |
| VeryMildDemented | 0.80                | 0.80               |
| MildDemented     | 0.68                | 0.66               |
| ModerateDemented | 0.60                | 0.59               |
| NonDemented      | 0.42                | 0.38               |

Tabel 2. Evaluation Result Per Class

Table 2 shows that CNN+MRF yields an improvement in F1-score across all classes, particularly for *VeryMildDemented* and *MildDemented*, which were previously prone to misclassification by the standard CNN model. This demonstrates that the integration of MRF effectively enhances segmentation and

classification performance, especially in cases involving subtle inter-class feature differences.

## 5. CONCLUSION

Based on the test results by comparing the model performance on the two graphs shown in Figure 5 and Figure 6, it can be concluded that the addition of Markov Random Field (MRF) to the Convolutional Neural Network (CNN) model has a fairly positive impact on model performance. In terms of accuracy, the pure CNN model has a value of 0.61, while the model combined with MRF increases the accuracy value to 0.63. Although this increase is relatively small, it shows that the CNN + MRF model approach is able to provide better prediction results compared to pure CNN

From the loss side, we can see that the CNN model has a value of 0.80, while for CNN + MRF the loss decreased to 0.74. This decrease indicates that the model with MRF is able to reduce prediction errors and is more stable in the learning process. Thus, it can be concluded that the implementation of MRF on CNN not only increases accuracy but also optimizes the training process by reducing the loss value. However, because the increase in accuracy is still relatively small, further analysis is needed to determine whether this difference is significant and useful in other cases.

## 6. SUGGESTION

This study is still limited to the use of publicly available datasets, to increase validity there needs to be a comparison dataset from the primary dataset, for further research it is expected to use a broader and more diverse dataset covering data sources from various institutions, as well as consideration of the diversity of imaging conditions. In the configuration study on the Markov Random Field is fixed, therefore further research is expected to explore parameters to optimize in increasing consistency in segmentation, at the evaluation stage we are limited to only measuring the level of accuracy, in further research it is expected to include additional metrics such as precision, recall, F1-score, and area under the curve (AUC), so that the classification ability of the model can be evaluated more comprehensively.

## REFERENCE

- [1] A. Chattopadhyay and M. Maitra, "MRI-based brain tumour image detection using CNN based deep learning method," Dec. 01, 2022, *Elsevier B.V.* doi: 10.1016/j.neuri.2022.100060.
- [2] S. Duong, T. Patel, and F. Chang, "Dementia: What pharmacists need to know," *Canadian Pharmacists Journal*, vol. 150, no. 2, pp. 118–129, Mar. 2017, doi: 10.1177/1715163517690745.
- [3] D. AlSaeed and S. F. Omar, "Brain MRI Analysis for Alzheimer's Disease Diagnosis Using CNN-Based Feature Extraction and Machine Learning," *Sensors*, vol. 22, no. 8, Apr. 2022, doi: 10.3390/s22082911.
- [4] N. Mukadam *et al.*, "Changes in prevalence and incidence of dementia and risk factors for dementia: an analysis from cohort studies," 2024. [Online]. Available: [www.thelancet.com/](http://www.thelancet.com/)
- [5] V. S. Nori *et al.*, "Machine learning models to predict onset of dementia: A label learning approach," *Alzheimer's and Dementia: Translational Research and Clinical Interventions*, vol. 5, pp. 918–925, Jan. 2019, doi: 10.1016/j.trci.2019.10.006.
- [6] S. Dhakal, S. Azam, K. M. Hasib, A. Karim, M. Jonkman, and A. S. M. Farhan Al Haque, "Dementia Prediction Using Machine Learning," in *Procedia Computer Science*, Elsevier B.V., 2023, pp. 1297–1308. doi: 10.1016/j.procs.2023.01.414.
- [7] E. Pellegrini *et al.*, "Machine learning of neuroimaging for assisted diagnosis of cognitive impairment and dementia: A systematic review," *Alzheimer's and Dementia: Diagnosis, Assessment and Disease Monitoring*, vol. 10, pp. 519–535, Jan. 2018, doi: 10.1016/j.dadm.2018.07.004.
- [8] H. Tufail, A. Ahad, M. H. Naqvi, R. Maqsood, and I. M. Pires, "Classification of Vascular Dementia on magnetic resonance imaging using deep learning architectures," *Intelligent Systems with Applications*, p. 200388, Jun. 2024, doi: 10.1016/j.iswa.2024.200388.
- [9] G. Battineni, N. Chintalapudi, and F. Amenta, "Machine learning in medicine: Performance calculation of dementia prediction by support vector machines (SVM)," *Inform Med Unlocked*, vol. 16, Jan. 2019, doi: 10.1016/j.imu.2019.100200.
- [10] J. Zhou *et al.*, "Identifying dementia from cognitive footprints in hospital records among Chinese older adults: a machine-learning study," 2024. [Online]. Available: [www.thelancet.com](http://www.thelancet.com)
- [11] J. Zhou *et al.*, "Identifying dementia from cognitive footprints in hospital records among Chinese older adults: a machine-learning study," 2024. [Online]. Available: [www.thelancet.com](http://www.thelancet.com)
- [12] J. You *et al.*, "Development of a novel dementia risk prediction model in the general population: A large, longitudinal, population-based machine-learning study," 2022, doi: 10.1016/j.
- [13] National Institute On Aging, "How MRI Is Used to Detect Alzheimer's Disease," : <https://www.verywellhealth.com/can-an-mri-detectalzheimers-disease-98632>.
- [14] A. H. Nizamani, Z. Chen, A. A. Nizamani, and U. A. Bhatti, "Advance brain tumor segmentation using feature fusion methods with deep U-Net model with CNN for MRI data," *Journal of King Saud University - Computer and Information Sciences*, vol. 35, no. 9, Oct. 2023, doi: 10.1016/j.jksuci.2023.101793.
- [15] W. Baccouch, S. Oueslati, B. Solaiman, and S. Labidi, "A comparative study of CNN and U-Net performance for automatic segmentation of medical images: Application to cardiac MRI," in *Procedia Computer Science*, Elsevier B.V., 2023, pp. 1089–1096. doi: 10.1016/j.procs.2023.01.388.
- [16] R. Li *et al.*, "Predicting incident dementia in cerebral small vessel disease: comparison of machine learning and traditional statistical models," *Cereb Circ Cogn Behav*, vol. 5, Jan. 2023, doi: 10.1016/j.cccb.2023.100179.
- [17] S. Dhakal, S. Azam, K. M. Hasib, A. Karim, M. Jonkman, and A. S. M. Farhan Al Haque, "Dementia Prediction Using

- Machine Learning,” in *Procedia Computer Science*, Elsevier B.V., 2023, pp. 1297–1308. doi: 10.1016/j.procs.2023.01.414.
- [18] E. Pellegrini *et al.*, “Machine learning of neuroimaging for assisted diagnosis of cognitive impairment and dementia: A systematic review,” *Alzheimer’s and Dementia: Diagnosis, Assessment and Disease Monitoring*, vol. 10, pp. 519–535, Jan. 2018, doi: 10.1016/j.dadm.2018.07.004.
- [19] H. Xu, C. Huang, X. Huang, C. Xu, and M. Huang, “Combining Convolutional Neural Network and Markov Random Field for Semantic Image Retrieval,” *Advances in Multimedia*, vol. 2018, 2018, doi: 10.1155/2018/6153607.
- [20] S. Doering *et al.*, “Deconstructing pathological tau by biological process in early stages of Alzheimer disease: a method for quantifying tau spatial spread in neuroimaging,” 2024. [Online]. Available: <https://fsl>.
- [21] D. G. Manurung *et al.*, “Deteksi Dan Klasifikasi Hama Potato Beetle Pada Tanaman Kentang Menggunakan YOLOV8,” *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 11, no. 4, pp. 723–734, Aug. 2024, doi: 10.25126/jtiik.1148092.
- [22] D. Anastasya, S. Fahri, S. Situmorang, F. Ramadhani, and J. I. Komputer, “Implementasi Metode Convolutional Neural Network (CNN) Dalam Klasifikasi Motif Batik.” [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>
- [23] B. Untung Saputra, W. Andriani, and A. Batik Pesisir, “PENGENALAN MOTIF BATIK PESISIR PULAU JAWA MENGGUNAKAN CONVOLUTIONAL NEURAL NETWORK.” [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>
- [24] D. Anastasya, S. Fahri, S. Situmorang, F. Ramadhani, and J. I. Komputer, “Implementasi Metode Convolutional Neural Network (CNN) Dalam Klasifikasi Motif Batik.” [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>
- [25] Md. M. Islam, Md. Z. Islam, A. Asraf, M. S. Al-Rakhami, W. Ding, and A. H. Sodhro, “Diagnosis of COVID-19 from X-rays using combined CNN-RNN architecture with transfer learning,” *BenchCouncil Transactions on Benchmarks, Standards and Evaluations*, vol. 2, no. 4, pp. 2–11, Oct. 2022, doi: 10.1016/j.tbench.2023.100088.
- [26] C. Li and M. Wand, “Combining Markov Random Fields and Convolutional Neural Networks for Image Synthesis.”
- [27] I. Scott *et al.*, “An automated method for tendon image segmentation on ultrasound using grey-level co-occurrence matrix features and hidden Gaussian Markov random fields,” *Comput Biol Med*, vol. 169, Feb. 2024, doi: 10.1016/j.compbimed.2023.107872.

## Assessing the Role of E-Govqual in Improving Public Satisfaction with the MyPertamina Application in Setiabudi District

Yulian Surya Saputra<sup>1</sup>, Sucitra Sahara<sup>\*2</sup>, Rizqi Agung Permana<sup>3</sup>

<sup>1,2</sup>Universitas Bina Sarana Informatika, Indonesia

<sup>3</sup>STMIK Antar Bangsa, Indonesia

E-mail: <sup>1</sup>yuliansurya68@gmail.com, <sup>\*2</sup>sucitra.scr@bsi.ac.id, <sup>3</sup>rizqiagung@gmail.com

### Abstract

*This research aims to analyze the influence of Customer Satisfaction on the quality of e-government services (E-Govqual) on the use of the My Pertamina application in Setiabudi District. The research method used was quantitative with a survey approach by distributing questionnaires to 150 respondents. The data collected was analyzed using validity tests, reliability tests, and hypothesis testing through linear regression analysis. The research results show that there is no significant influence between Customer Satisfaction and E-Govqual. This is proven by a correlation value of 0.160, which indicates a very weak relationship, as well as a coefficient of determination ( $R^2$ ) value of 2.56%, which indicates that variations in E-Govqual are only slightly influenced by Customer Satisfaction. The results of the hypothesis test show a significance value of  $0.050 > 0.05$  and  $F_{count} (3.891) < F_{table} (3.905)$ , so that the null hypothesis ( $H_0$ ) is accepted and the alternative hypothesis ( $H_1$ ) is rejected. Satisfaction ( $X$ ) on E-Govqual ( $Y$ ) is  $0.050 > 0.05$ , and the calculated  $F$ -value ( $F$ -count) is less than the  $F$ -table value, namely  $3.891 < 3.905$ . These findings suggest that other factors beyond Customer Satisfaction may have a greater influence on e-government service quality. Based on the results, it is recommended that the development of the My Pertamina application be more focused on improving system quality, ease of use, and other technical aspects to improve overall service quality.*

**Keywords:** Customer Satisfaction, E-Govqual, Service Quality, Digital Applications

Submission: April 25, 2025

Accepted: June 12, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.379>

### 1. INTRODUCTION

Crude oil plays an important role in human life as one of the strategic commodities that supports various sectors of life. Crude oil serves as the main raw material in the production of Fuel Oil (BBM) and is a primary energy source that is highly needed to support social and economic activities worldwide [10]. To this day, the use of BBM remains a basic necessity for society, especially in the transportation sector, which dominates global energy consumption [9]. BBM also plays a significant role in maintaining economic stability because fluctuations in BBM prices can affect logistics costs, the prices of essential goods, and trigger inflation that directly impacts people's purchasing power [12]. Challenges in managing BBM also include efficient and targeted distribution.

Pertamina, as a state-owned company, has been mandated by the government to distribute subsidized BBM in accordance with the established quota and to ensure that it reaches low-income communities, according to the Ministry of Energy and Mineral Resources in

2023. One way to prevent pollution is by socializing and educating the public about the dangers of waste [22]. However, according to a report by the Oil and Gas Downstream Regulatory Agency based on data from BPH Migas, the distribution of subsidized BBM is often misdirected, with many being enjoyed by more affluent groups who are not entitled to it [6]. To address this issue, the government, together with Pertamina, launched a digital innovation in the form of the My Pertamina application [2]. This application is expected to improve the efficiency of subsidized BBM distribution by using a data-based registration system to verify the eligibility of subsidized BBM users.

However, despite the implementation of this application, the public still faces various challenges, such as a lack of understanding of technology, limited internet access, and complaints related to user experience [8]. The presence of the My Pertamina application is not only intended to simplify the distribution of subsidized BBM but also represents one form of e-government implementation in Indonesia [11].

In the context of e-government, the application of information technology aims to improve the quality of public services, transparency, and government accountability [15]. However, the effectiveness of this application still needs to be evaluated to ensure that it truly provides a positive impact on increasing public satisfaction as end-users [18]. The issues that arise in the implementation of the My Pertamina application include several important aspects [20]. First, the low level of public understanding regarding the use of this application, particularly among users who are not familiar with digital technology. Second, technical constraints such as difficulties in the registration process, slow data verification, and frequent system disruptions cause discomfort for users. Third, the lack of responsive customer service to address user complaints also becomes a barrier to improving public satisfaction. Additionally, the disparity in internet access in certain areas worsens the effectiveness of the application's implementation [5][19]. This study aims to measure the level of public satisfaction in Setiabudi District in using the My Pertamina application [13]. The My Pertamina application, an innovation from Pertamina since 2017, was introduced to provide convenience for the public in transactions at Pertamina. However, to date, there has been little research evaluating the public satisfaction level with the use of the My Pertamina application, especially in the Setiabudi District of South Jakarta. Using the E-GovQual method, this study focuses on evaluating the quality of digital services based on aspects of usability, reliability, and citizen support [4]. The results of this study are expected to provide input for the government and Pertamina to improve the quality of subsidized fuel services and support the optimization of technology implementation in public services.

## 2. METHODOLOGY

Using the E-GovQual method as a tool to measure the level of public satisfaction with the MyPertamina application in Setiabudi District. E-GovQual is an instrument designed to assess the quality of electronic government services, referring to dimensions that include effectiveness, efficiency, and user satisfaction. In this study, the data collection technique used is through the distribution of questionnaires. The questionnaire is specifically designed to measure the level of public satisfaction with the use of the MyPertamina application, based on the indicators from the E-GovQual method.

The research methodology stages will be described as outlined below:

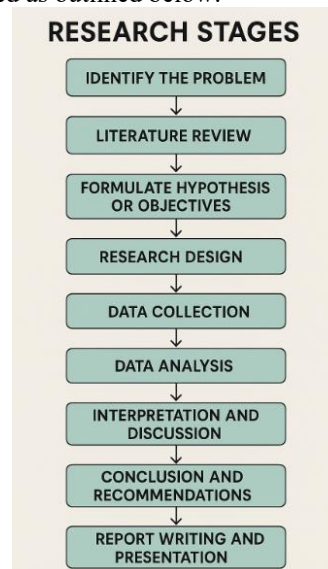


Figure 1 Research Stages

- a. Identify the Problem  
Define the research question or problem clearly.
- b. Literature Review  
Explore existing research and theories related to the topic.
- c. Formulate Hypothesis or Objectives  
Develop a hypothesis (for scientific studies) or specific research objectives.
- d. Research Design  
Choose the methodology (qualitative, quantitative, or mixed methods), design experiments or surveys.
- e. Data Collection  
Gather data through experiments, interviews, surveys, or observations.
- f. Data Analysis  
Analyze the collected data using statistical or qualitative analysis tools.
- g. Interpretation and Discussion  
Interpret the results, discuss findings, and compare with previous research.
- h. Conclusion and Recommendations  
Summarize key findings, draw conclusions, and provide suggestions for future research or practical applications.
- i. Report Writing and Presentation  
Write the research paper or report and present the results to stakeholders or publish them.

The E-GovQual method was used as a tool to measure the level of public satisfaction with the MyPertamina application in Setiabudi District. E-GovQual is an instrument designed to

assess the quality of electronic government services by referring to dimensions that include effectiveness, efficiency, and user satisfaction. This study involved 150 respondents who are active users of the MyPertamina application in Setiabudi District. The questionnaire was distributed online using the Google Forms platform during the data collection period, which took place from November to December 2024. Respondents were selected using purposive sampling, with the criterion that they had used the MyPertamina application for at least one month. Out of the total 150 questionnaires distributed, all were successfully completed and returned by the respondents, with a 100% response completeness rate. This indicates a high level of respondent participation in this study, ensuring that the data could be processed without any need for deletion due to incomplete responses.

Table 1 Questionnaire Results

| Description                   | Quantity   |
|-------------------------------|------------|
| Total questionnaires obtained | 150 (100%) |
| Usable data                   | 150 (100%) |
| Unusable data                 | 0 (0%)     |

### 3. RESEARCH RESULTS

In this study, respondents can be identified based on their full names and gender. The following diagram illustrates the gender distribution of all respondents, Male respondents dominated participation in this study compared to female respondents. Female respondents accounted for 16.9%, while male respondents made up 83.1% of the total. This study also included a question regarding whether or not the respondents had made transactions using the MyPertamina application, and the result showed that 100% of respondents had previously made transactions. The following is the illustrative diagram:

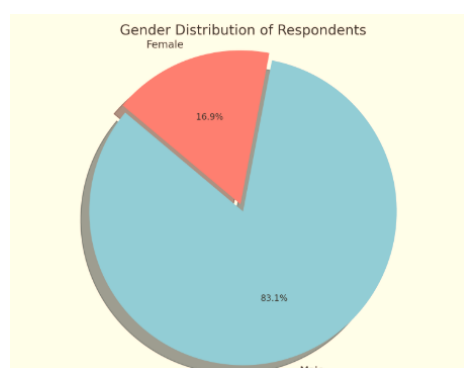


Figure 2 Gender Distribution

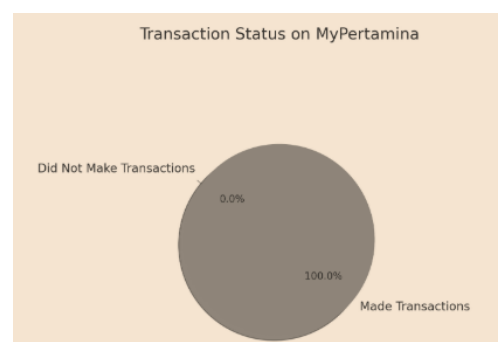


Figure 3 Application User Distribution

The validity test used in this study was conducted using SPSS, as follows:

Table 2. Validitas X

|     | X01     | X02   | X03   | X04   | X05   | X06   | X07   | X08   | X09   | X10   | X11   | X12   | X13   | X14   | X15   | Total |
|-----|---------|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|
| 001 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 002 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 003 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 004 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 005 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 006 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 007 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 008 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 009 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 010 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 011 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 012 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 013 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 014 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 015 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 016 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 017 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 018 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 019 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 020 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 021 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 022 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 023 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 024 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 025 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 026 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 027 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 028 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 029 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 030 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 031 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 032 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 033 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 034 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 035 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 036 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 037 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 038 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 039 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |
| 040 | Pretest | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 | 0.000 |

Validity testing is conducted to measure the extent to which each questionnaire item accurately represents the variable being measured. In this study, the validity test used the Pearson Product Moment correlation method to examine the relationship between each item score and the total score. Based on the calculations, the correlation coefficient value ( $r$ -calculated) for each item was compared to the  $r$ -table value at a 5% significance level ( $\alpha = 0.05$ ).

Table 3. Validitas Y

[illegible]

The population in this study consisted of 14,925 individuals, with a sample size of 150 respondents, determined using the Slovin formula with a margin of error of 10% (0.10). The results of the validity test indicated that most of the statement items had r-calculated values greater than the r-table, meaning those items are valid. This demonstrates that the items are capable of accurately measuring the intended construct. The high level of validity also indicates that the questionnaire instrument used is relevant and reliable for data collection. Following are the results Reliability Test Analysis of Customer Satisfaction Variables:

Table 2. Reliability Test Analysis Table

|                  |            |
|------------------|------------|
| x                |            |
| Cronbach's Alpha | N of Items |
| 0,791            | 15         |

Based on the available information, it can be concluded that the variable values in this questionnaire are considered strong or reliable, with a score of 0.791, which falls into the high category, as the calculation results in a value greater than 0.06.

Table 3. Reliability Test of eGovqual

| Reliability Statistics |            |
|------------------------|------------|
| Cronbach's Alpha       | N of Items |
| 0.610                  | 15         |

Based on the available information, it can be concluded that the variable values in this questionnaire are considered strong or reliable, obtaining a score of 0.610, which falls into the high category, as the calculation yields a value greater than 0.06.

### 1. Hypothesis Testing T (Partial)

The T-test is conducted to determine whether there is a significant effect or relationship between the independent variable and the dependent variable. This test also aims to assess the extent to which the independent variable can influence the dependent variable, while considering and controlling for other variables that may play a role in the relationship.

Table 4. T - test

|       |                    | Coefficients <sup>a</sup>   |            | Standardized Coefficients | t      | Sig. |
|-------|--------------------|-----------------------------|------------|---------------------------|--------|------|
| Model |                    | Unstandardized Coefficients | Std. Error | Beta                      |        |      |
| 1     | (Constant)         | 58.204                      | 2.288      |                           | 25.438 | .000 |
|       | Kepuasan Pelanggan | .099                        | .050       | .160                      | 1.973  | .050 |

a. dependent Variable: E-Govqual

The results shown in the table above indicate that the significance value for the effect of Customer Satisfaction (X) on E-Govqual (Y) is  $0.000 < 0.05$ . In addition, the calculated t-value (t-count) obtained is 1.973, which is less than the t-table value of 1.976. Based on this, it can be concluded that H1 (the alternative hypothesis) is rejected and H0 (the null hypothesis) is accepted. This means that there is no significant effect between the Customer Satisfaction variable (X) and the E-Govqual variable (Y).

## 2. Hypothesis F Test (Simultan)

The F-test (ANOVA) is used to assess the suitability of the data model employed in the research. This test aims to examine whether there are significant differences between the groups or variables being tested, as well as to determine the extent to which the applied regression model can explain the overall variability of the data. The results of the F-test provide an overview of whether the relationships between the variables tested can be statistically accepted or rejected. The following are the F-test results obtained using SPSS version 25.

Table 4. T - test

| Model |            | Sum of Squares | df  | Mean Square | F     | Sig.              |
|-------|------------|----------------|-----|-------------|-------|-------------------|
| 1     | Regression | 76.799         | 1   | 76.799      | 3.891 | .050 <sup>b</sup> |
|       | Residual   | 2920.861       | 148 | 19.736      |       |                   |
|       | Total      | 2997.660       | 149 |             |       |                   |

a. Dependent Variable: E-Govqual  
b. Predictors: (Constant), Kepuasan Pelanggan

Based on the results of the F-test, the calculated F-value (F-count) is 3.891 with a significance

level of 0.050. Meanwhile, the F-table value is calculated using the following formula:

$$F_{\text{table}} = F_{\text{inv}}(0,05 ; k ; n - k - 1) \\ F_{\text{inv}}(0,05 ; 1 ; 150 - 1 - 1) = 3,905$$

By referring to the number of independent variables and  $n$  as the number of samples. From these calculations, it is found that the significance value for the effect of Customer Satisfaction (X) on E-Govqual (Y) is  $0.050 > 0.05$ , and the calculated F-value (F-count) is less than the F-table value, namely  $3.891 < 3.905$ . Therefore, it can be concluded that there is no significant effect between the Customer Satisfaction variable (X) and the E-Govqual variable (Y).

#### 4. DISCUSSION

The results of the hypothesis testing regarding the relationship between the customer satisfaction (X) variable and the E-Govqual (Y) variable indicate that there is no significant effect between the two variables. Although a positive correlation exists, the obtained correlation value is 0.160, which indicates a very weak level of association between customer satisfaction. The coefficient of determination ( $R^2$ ) of 2.56% shows that only a small portion of the variation in E-Govqual method can be explained by customer satisfaction, while the remaining variation is influenced by other factors not measured in this study. Based on the hypothesis testing results, the significance value obtained is  $0.050 > 0.05$ , and the calculated F-value (F-count) is lower than the F-table value,  $3.891 < 3.905$ , thus the null hypothesis ( $H_0$ ) is accepted and the alternative hypothesis ( $H_1$ ) is rejected. This indicates that there is no significant effect of customer satisfaction within the context of this study. This finding is consistent with several studies that have also found that the impact of customer satisfaction on service quality in the e-government sector is not always significant. For example, research by (Al-Qeisi et al., 2014) found that although there is a relationship between satisfaction and digital service quality, its effect is not strong enough to be considered significant in every context. Another study by (Munas Dwiyanto, 2021) also showed that the quality of e-government services can be influenced by various other factors beyond just user satisfaction. Therefore, it can be concluded that customer satisfaction does not have a significant impact on E-Govqual. Other variables not covered in this study may have a greater impact on the overall service quality.

#### 5. CONCLUSION

This study aimed to examine the impact of E-Govqual implementation on user satisfaction with the MyPertamina application in Setiabudi District. Based on the findings, it can be concluded that while there is a positive correlation between customer satisfaction and the E-Govqual dimensions—such as service quality, reliability, accessibility, data security, and user experience—the strength of this relationship is relatively weak. The correlation coefficient of 0.160 and a coefficient of determination ( $R^2$ ) of 2.56% indicate that only a small proportion of user satisfaction can be explained by the quality of e-government services provided by the application. Furthermore, hypothesis testing showed a significance value above the threshold ( $0.050 > 0.05$ ) and an F-value lower than the F-table, suggesting that the influence of E-Govqual on satisfaction is not statistically significant. These results imply that while the application meets some standards of e-government service quality, other factors not covered by E-Govqual may play a more substantial role in shaping user satisfaction. Therefore, future improvements to MyPertamina should consider broader aspects, including technical performance, user support, and personalization features to enhance overall satisfaction. This study will examine the level of public satisfaction in the Setiabudi District regarding the use of the MyPertamina application by applying the E-Govqual method. The scope of this research includes several aspects as follows: Research subjects the study focuses on residents of Setiabudi District who have used the MyPertamina application. Respondents will be selected from various demographic backgrounds to obtain a comprehensive overview. Aspects studied this research will assess several dimensions of public satisfaction, including service quality, reliability, ease of access, data security, and user experience of the MyPertamina application. There is a positive correlation, with a correlation value of 0.160, indicating a very low level of association between Customer Satisfaction and E-Govqual. The coefficient of determination ( $R^2$ ) of 2.56% shows that only a small portion of the variation in E-Govqual can be explained by Customer Satisfaction, while the remaining variation is influenced by other factors not measured in this study. Based on the hypothesis testing results, the significance value obtained is  $0.050 > 0.05$ , and the calculated F-value (Fcount) is lower than the F-table value, namely  $3.891 < 3.905$ . This research enriches the literature on e-government

service quality by applying the E-GovQual method specifically to the context of BUMN's digital application, namely MyPertamina. The geographical focus on Setiabudi District provides a relevant micro overview of the implementation of public service technology in urban areas with high levels of digital use. This study provides an empirical picture of people's perceptions and satisfaction with various dimensions of MyPertamina services (for example, ease of use, system reliability, user support, and information security), so that it can be used as direct evaluation material by application managers.

## 6. SUGGESTION

Based on the results of the research, several suggestions can be considered for further development, management of the MyPertamina application is advised to continuously improve service quality, particularly in terms of access speed, system reliability, and user comfort. This is essential to strengthen users' positive perceptions of the application, Enhancement of Security and Privacy Features: Given the importance of protecting personal data in digital applications, strengthening security features and ensuring transparency in user data management should be prioritized to build public trust. Expansion of the Factors Studied: Future research should consider incorporating additional variables that may influence user satisfaction, such as technical support quality, application interactivity, and the introduction of new features, to provide a more comprehensive analysis.

## REFERENCES

- [1] AlAl-Qeisi, K., Dennis, C., Alamanos, E., & Jayawardhena, C. (2014). Website design quality and usage behavior: Unified theory of acceptance and use of technology. *Journal of Business Research*, 67(11), 2282–2290. <https://doi.org/10.1016/j.jbusres.2014.06.016>
- [2] Alvita Wagiswari D, P. R., Susilawati, I., Witanti, A., Studi Informatika, P., Teknologi Informasi, F., & Mercu Buana Yogyakarta, U. (2023). Analisis Sentimen pada Komentar Aplikasi MyPertamina dengan Metode Multinomial Naïve Bayes. In *Jurnal ForAI* (Vol. 1, Issue 1).
- [3] Darmawan, G., Alam, S., & Sulisty, M. I. (2023). *Analisis Sentimen Berdasarkan Ulasan Pengguna Aplikasi MyPertamina Pada Google Playstore Menggunakan Metode Naïve Bayes Info Artikel Abstrak*. 2(3), 100–108. <https://doi.org/10.55123>
- [4] Faisal Jatnika, R., Kaniawulan, I., Singasatia, D., Studi, P., Informatika, T., Teknologi, S. T., & Purwakarta, W. (2023). *Analisis Penerimaan Aplikasi MyPertamina Menggunakan Metode Technology Acceptance Model (TAM)* (Vol. 14, Issue 2). [http://ejurnal.provisi.ac.id/index.php/JTI\\_KP](http://ejurnal.provisi.ac.id/index.php/JTI_KP)
- [5] Firmansyah, D., & Ahsan, M. (2023). Monitoring Kualitas Pada Aplikasi MyPertamina Berdasarkan Rating Pengguna di Google Play Menggunakan Diagram Kendali P. *JURNAL SAINS DAN SENI ITS*, 12(2).
- [6] Hanif Fariz Ramadhani, Miranda Istikarani, & Marcha Delaya Laura. (2024). Persepsi Masyarakat Pontianak: Analisis Tingkat Kesadaran Masyarakat Terhadap Aplikasi MyPertamina dan Media Sosial PertaminaKalbar. *CAPACITAREA: Jurnal Pengabdian Kepada Masyarakat*, 3(2), 59–66. <https://doi.org/10.35814/capacitarea.2023.003.02.08>
- [7] Hikmawati, N. K. (2022). Analisis Kualitas Layanan My Pertamina Menggunakan Pendekatan e-GovQual pada Beberapa Kota Percobaan. *Jurnal Manajemen Informatika (JAMIKA)*, 12(2), 100–111. <https://doi.org/10.34010/jamika.v12i2.7977>
- [8] Indrayanto, C. G., Ratnawati, D. E., & Rahayudi, B. (2023). *Analisis Sentimen Data Ulasan Pengguna Aplikasi MyPertamina di Indonesia pada Google Play Store menggunakan Metode Random Forest* (Vol. 7, Issue 3). <http://j-ptiik.ub.ac.id>
- [9] Islamia, R., Raikhan, I., Faizy, A., Aqilla, A., Ahmad, R. F., Zahra, A., Arum, P., & Pratama, G. (2022). Dampak Kenaikan Harga Bahan Bakar Minyak (BBM) Terhadap Sembilan Bahan Pokok (Sembako) Di Toko Sani Kabupaten Cirebon. In *Jurnal Ekonomi Manajemen* (Vol. 17, Issue 2). <http://oaj.stiecirebon.ac.id/index.php/jem>
- [10] Lutfi, A., Annisa Fitriani, A., Ramadani, I., Azahra Putri, N., & Shizuka Nelsi, Y. (2022). *Efektivitas Penggunaan Aplikasi My Pertamina Di Era Kenaikan BBM Bersubsidi*. 1(2).
- [11] Manfaat, P., Diri, E., Penggunaan, K., Terhadap, K., Muhammad Ibrahim, R.,

- Novandriani Karina Moeliono, N., Prodi Administrasi Bisnis, M., Komunikasi Bisnis, F., & Telkom, U. (2020). Persepsi Konsumen Pada My Pertamina (Studi Pada Penggunaan My Pertamina Kota Bandung). *Jurnal Ilmiah Mahasiswa Ekonomi Manajemen Accredited SINTA*, 4(2), 396–413. <http://jim.unsyiah.ac.id/ekm>
- [12] Maria, R., Umayah, R. U., Mahardinny, S., Kalana, D. N., & Saputra, D. D. (2023). Analisis Sentimen Persepsi Masyarakat Terhadap Penggunaan Aplikasi My Pertamina Pada Media Sosial Twitter Menggunakan Metode Naïve Bayes Classifier. *Jurnal Komputer Antartika*, 1. <https://ejournal.mediaantartika.id/index.php/jka>
- [13] Maulana, I., Apriandari, W., & Pambudi, A. (2023). Analisis Sentimen Berbasis Aspek Terhadap Ulasan Aplikasi Mypertamina Menggunakan Support Vector Machine. In *Idealis: Indonesia Journal Information System* (Vol. 6, Issue 2). <http://jom.fti.budiluhur.ac.id/index.php/IDEALIS/index>
- [14] Maulana, R., Voutama, A., & Ridwan, T. (2023). Analisis Sentimen Ulasan Aplikasi Mypertamina Pada Google Play Store Menggunakan Algoritma NBC. *Jurnal Teknologi Terpadu*, 9(1).
- [15] Maulani, A. R., Wahyudi, T., & Rahmawati, R. (2023). Pengukuran Tingkat Kepuasan Konsumen Aplikasi Mypertamina Dengan Metode Servqual, CSI, Dan IPA. *INTEGRATE: Industrial Engineering and Management System*, 7(3), 19–25. <https://jurnal.untan.ac.id/index.php/jtinUN TAN/issue/view/2162>
- [16] Munas Dwiyanto, B. (2021). Pelanggan Sebagai Mediator (Studi pada Pengguna Jasa PT. Pos Indonesia di Semarang). *DIPONEGORO JOURNAL OF MANAGEMENT*, 10(3), 1–12. <http://ejournal-s1.undip.ac.id/index.php/dbr>
- [17] Mustamu, D. D., Rachma, Y., Ip, P. S., Prodi, M. M., Komunikasi, I., Komunikasi, F., & Bisnis, D. (2019). Pengaruh Promosi Melalui Aplikasi Mypertamina Terhadap Keputusan Pembelian Bahan Bakar Pertamina Di Masyarakat Kota Bandung. 6(2).
- [18] Mustasaruddin, M., Budianita, E., Fikry, M., & Yanto, F. (2023). Klasifikasi Sentiment Review Aplikasi MyPertamina Menggunakan Word Embedding FastText dan SVM (Support Vector Machine). *Jurnal Sistem Komputer Dan Informatika (JSON)*, 4(3), 526. <https://doi.org/10.30865/json.v4i3.5695>
- [19] Putra, I. G. B. D., & Pramarthar, C. R. A. (2024). Analisis Sentimen Pengguna Aplikasi MyPertamina Dengan Menggunakan Algoritma Naïve Bayes. *Jurnal Elektronik Ilmu Komputer Udayana*, 12(3). [www.kateglo.com](http://www.kateglo.com).
- [20] Resca Gitasela, Y., & Narti, S. (2023). Analisis Resepsi Khalayak Tentang Aplikasi MyPertamina (Studi Pada Masyarakat Kota Bengkulu). *Jurnal Multimedia Dehasen*, 2(2), 405–418. <https://id.m.wikipedia.org/>
- [21] Risky, H. A., Irmayanti, D., & Totohendarto, M. H. (2023). Redesign UI/UX Aplikasi Mobile My Pertamina Menggunakan Metode User Centered Design (UCD). *JATI (Jurnal Mahasiswa Teknik Informatika)*, 7.
- [22] Wanhar, F. A., & Widodo, H. (2021). Sosialisasi Program Bersih Pantai dan Edukasi Kepada Masyarakat Lingkungan Pantai Bali Lestari Desa Pantai Cermin Kanan. *Jurnal Pengabdian Masyarakat Indonesia*, 1(6), 285–289. <https://doi.org/10.52436/1.JPMI.60>

## A Data-Driven Approach to Comparative Evaluation of Regression Models for Accurate House Price Prediction

Tiara Permata Hati<sup>1</sup>, Budiman<sup>\*2</sup>, Nur Alamsyah<sup>3</sup>

<sup>1,2,3</sup>Universitas Informatika dan Bisnis Indonesia, Indonesia

E-mail: <sup>1</sup>[tiara.p21@student.unibi.ac.id](mailto:tiara.p21@student.unibi.ac.id), <sup>\*2</sup>[budiman@unibi.ac.id](mailto:budiman@unibi.ac.id), <sup>3</sup>[nuralamsyah@unibi.ac.id](mailto:nuralamsyah@unibi.ac.id)

### Abstract

*This study aims to develop and evaluate a property price prediction model in Bandung by applying machine learning (ML) algorithms. The need for more accurate property price predictions is increasing due to fluctuations in the property market. This study analyzes property characteristics, including the number of bedrooms, bathrooms, land area, building area, and location, as well as their impact on house prices. The study evaluates four regression algorithms, including linear regression, K-Nearest Neighbors (KNN), Random Forest, and XGBoost. Finally, this study proposes price\_per\_m2 and building\_land\_ratio as new features recommended for improvement in accuracy. The bottleneck method is derived from the data collection area of the Rumah123.com website, encompassing data preprocessing and data exploration. The following metrics will be used to evaluate each model: Mean Absolute Error (MAE), Mean Squared Error (MSE), Root Mean Squared Error (RMSE), and R-squared (R<sup>2</sup>). Based on our study, we conclude that both Random Forest Regression and XGBoost Regression achieve the highest accuracy, with R<sup>2</sup> values of 0.9941 and 0.9955, respectively, after adjustment. Conversely, Linear Regression and KNN Regression have the lowest accuracy, with KNN Regression being the least accurate. The primary contribution of this study is the development of a more accurate house price prediction model that can be applied in cities with similar market characteristics. These findings provide practical insights for property developers and buyers when making investment decisions.*

**Keywords**— house price prediction, machine learning, feature engineering, random forest regression, XGBoost regression

Submission: June 05, 2025

Accepted: June 20, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.411>

### 1. INTRODUCTION

In recent years, machine learning (ML) and artificial intelligence (AI) have emerged in the real estate sector, advancing predictive capabilities for property price estimations and valuations. Machine learning (ML) and artificial intelligence (AI) can enhance data-informed decision-making, which is crucial for stakeholders navigating shifting market conditions. Numerous studies have examined the success of various machine learning (ML) algorithms, including decision trees, random forests, and ensemble methods, in predicting housing prices more accurately and efficiently [1], [2], [3]. Random forests and gradient boosting models, which are more sophisticated than traditional methods, have been shown in numerous studies to outperform older methods in many situations due to the complex interactions of input features and their non-linear relationships in real estate data [4], [5].

Increasingly, machine learning models are being used globally, including India, the United States, and Saudi Arabia, to predict property prices and support real estate investment decisions [6], [7]. The use of techniques such as K-Nearest Neighbors (KNN), support vector regression (SVR), and linear regression all provide additional predictive value and accuracy dependent on various attributes of the dataset [8], [9], [10], [11]. Likewise, ensemble methods such as XGBoost and LightGBM outperform maximum simpler models like linear regression as a result of their ability to create complex feature interactions and great flexibility concerning non-linearity.

Multiple authors have highlighted the relevance of model selection and Feature Engineering use when the goal is to maximize prediction accuracy. Regression models and decision trees have been tested from a variety of perspectives to predict house prices in multiple cities, and it was found that model choice is

significant in improving predictive outcomes [12], [13], [14]. This research aims to analyze and compare several machine learning (ML) models — specifically, decision trees, random forests, and linear regression — in predicting housing prices across multiple regions. The comparative outcomes could provide recommendations to a wide range of real estate stakeholders seeking to optimize their property price forecasting methodologies and continue advancing the field of real estate analytics [15], [16].

This study aims to design and evaluate a house price prediction model for Bandung based on property characteristics such as the number of bedrooms and bathrooms, land area, building area, carports, and location. The study aims to identify the primary predictors of house prices in Bandung and apply feature construction principles to enhance model performance. Additionally, the study compares the predictive performance of several regression models—linear regression, K-Nearest Neighbors, Random Forest, and XGBoost—to predict house prices. It

optimizes model performance through tuning and cross-validation to evaluate the model's generalization performance on unseen data.

The impact of this research in the area of property pricing prediction through a data-driven approach is manifold. One of those contributions is the introduction of Feature Engineering techniques, which create new variables such as `price_per_m2`, `building_land_ratio`, and `total_rooms`. Including these variables can positively impact model performance on non-normally distributed data. Furthermore, this research compares four standard regression algorithms used for house price prediction and presents their advantages and disadvantages based on extensive evaluation metrics, e.g., MAE, MSE, RMSE, and  $R^2$ .

## 2. METHODS

Figure 1 illustrates the proposed method employed in this study, spanning from data collection to model evaluation.

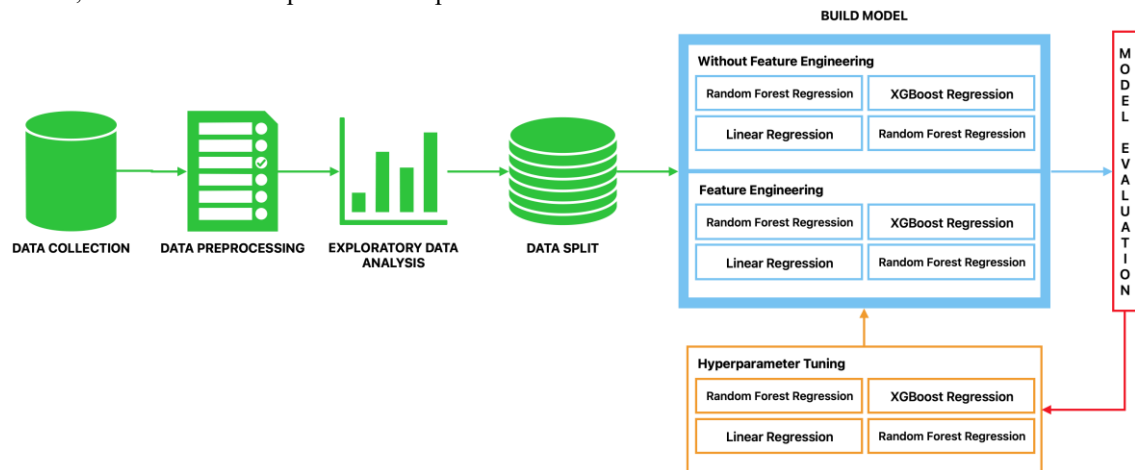


Figure 1. Proposed Method

### 2.1. Data Collection

This study utilized data sourced from the Rumah123.com website, which provides data related to house prices throughout the City of Bandung. This dataset alone consisted of 7,609 rows of data, including significant and relevant information to forecast house pricing based on property characteristics. The data from Rumah123.com provided a reasonably comprehensive picture of the factors that may affect house pricing, and the information will help to build a better prediction model to estimate house prices. The characteristics of the dataset used in the modeling are shown in Table 1.

Table 1 Feature Dataset

| Feature        | Description                                                         |
|----------------|---------------------------------------------------------------------|
| Price          | A target variable that represents the price of the house in Rupiah. |
| Bedroom Count  | Indicates the number of bedrooms available in the house.            |
| Bathroom Count | Indicates the number of bathrooms present on the property.          |
| Carport Count  | Describes the capacity of the available parking area or garage.     |
| Land Area      | The land area of the house is in square meters (m <sup>2</sup> ).   |
| Building Area  | The area of the house is in square meters (m <sup>2</sup> ).        |

| Feature  | Description                                                                                                     |
|----------|-----------------------------------------------------------------------------------------------------------------|
| Location | The house's geographical location in Bandung City is converted into numerical data through encoding techniques. |

Price is the dependent or target variable in this research. It is the primary variable we want to study; we want to estimate the house price based on the other property characteristics. House prices will be expressed in Rupiah, which is influenced by many factors in the data, such as the number of bedrooms, the number of bathrooms, the size of the parking area or garage, the size of the land, and the size of the building, among many others. Thus, house price is a significant variable in this research and will be calculated or predicted using the various features from the dataset.

## 2.2. Data Preprocessing

Several key steps are involved in data preprocessing techniques to ensure high data quality. Those steps include missing value handling to avoid any missing values that would inaccurately reflect data and, in return, less accurate modeling; duplicate removal to avoid biased modeling due to duplicated data; and finally, for categorical data like location, encoding was performed using the Label Encoding technique so that categorical data could be encoded into numerical data to make it analyzable and easy to model; data normalization was used to normalize all of the different numerical features for the model to work reliably. Lastly, the data was split into training and testing datasets at an 80:20 ratio using the `train_test_split` function. This would give an equal number of test and training data to train and validate the model. The training dataset consists of 4740 samples, and the testing dataset comprises 1185 samples.

## 2.3. Exploratory Data Analysis

Figure 2 illustrates the distribution of the variables. The relationship of the variables can be studied using scatterplots and correlation heatmaps. It is evident that house prices positively correlate with several physical variables. There is a correlation between house prices and `land_area`, albeit not always proportionate, considering location and several other factors. The relationship between prices and building area appears stronger and more stable, involving homes with building areas greater than or equal to. On the other hand, `bedroom_count` shows a weaker relationship, as there seems to be high price variation even with

the same number of rooms, which suggests that this metric is not a leading contributor. Bathroom\_count positively correlates with prices - unlike bedroom\_count, but does not appear to carry the same strength as building\_area or land\_area.

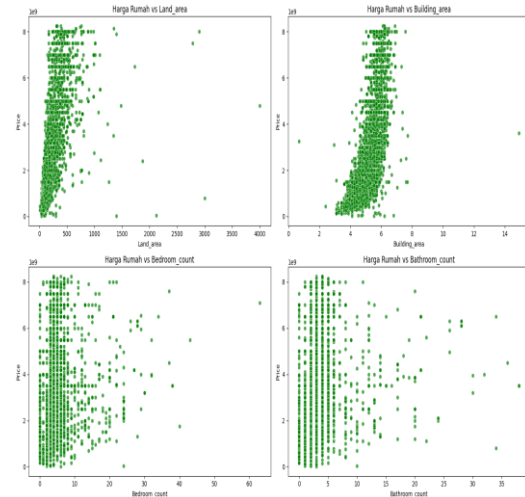


Figure 2. Distribution of Variables

Figure 3 presents a correlation matrix that highlights the strength of the relationship between variables and house prices (price). The correlation of 0.73 between price and building area and 0.63 with land area verifies that both variables are significant predictors of price in the dataset, with a slightly greater influence from building area than land area. There is a strong correlation of 0.79 with bedroom\_count predicting bathroom\_count. Thus, it would appear that homes with more bedrooms tend to have more bathrooms. The medium correlation of 0.59 between land area and building area appears to confirm that properties with larger land areas tend to have larger buildings as well. Carport\_count has a weak correlation with other variables, with a price of 0.31, meaning carport\_count appears to be the less significant variable in determining property value. Overall, the heatmap confirms what we already knew, with the two variables—building area and land area—appearing to have the most significant influence on house prices.

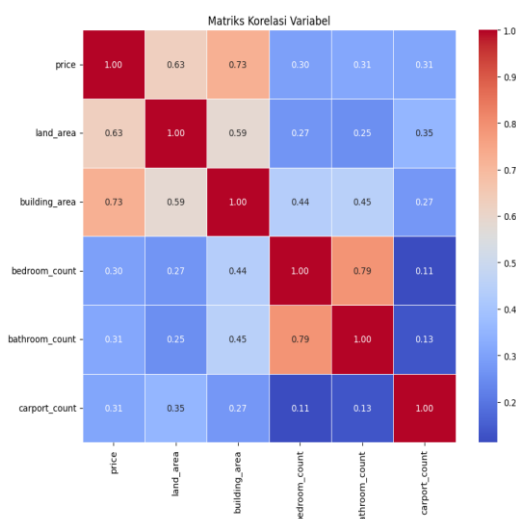


Figure 3. Heatmap Correlation

## 2.4. Feature Engineering (FE)

At this point, some supplementary features are created based on the relationships between the original variables, data transformation, and planned discretizing. This step enhanced the model's ability to identify patterns in the data and also addressed non-normal distributions. The newly created features are presented in Table 2.

| No | Feature                       | Formula                                           | Description                                                                                                              |
|----|-------------------------------|---------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------|
| 1  | price_per_m2                  | price / land_area                                 | Calculates the price per square meter (m <sup>2</sup> ) to assess the relative value of a property based on land size.   |
| 2  | building_land_ratio           | building_area / land_area                         | The ratio between building and land area indicates how efficiently the land is utilized.                                 |
| 3  | total_rooms                   | bedroom_count + bathroom_count                    | Total number of rooms (bedrooms and bathrooms), representing the comfort and capacity of the property.                   |
| 4  | total_facilities              | bedroom_count + bathroom_count + carport_count    | Total facilities in the property, including bedrooms, bathrooms, and carports.                                           |
| 5  | land_size_category            | Categorized using bins: [0, 50, 100, 200, 500, ∞] | Groups land size into defined categories (1: very small, 5: very large).                                                 |
| 6  | building_size_category        | Categorized using bins: [0, 50, 100, 200, 500, ∞] | Groups building size into defined categories (1: very small, 5: very large).                                             |
| 7  | land_bedroom_interaction      | land_area * bedroom_count                         | Interaction between land size and number of bedrooms, illustrating whether large properties are proportionally equipped. |
| 8  | building_bathroom_interaction | building_area * bathroom_count                    | The interaction between building size and the number of bathrooms indicates adequate sanitary facilities.                |
| 9  | price_scaled                  | np.log1p(price)                                   | Scaled transformation of price to manage skewed price distribution.                                                      |
| 10 | land_area_scaled              | np.log1p(land_area)                               | Scaled transformation of land area to handle non-normal data distribution.                                               |
| 11 | building_area_scaled          | np.log1p(building_area)                           | Scaled transformation of building area to improve interpretability and model stability for skewed data.                  |

Table 2. Feature Engineering

## 2.5. Modelling

The modeling of this study makes predictions for house prices in Bandung using the number of bedrooms, bathrooms, land area, building area, number of carports, and location as variables. The model primarily aimed to achieve a high level of predictive accuracy while also providing prospective property buyers or developers with an indication of the key factors

that drive house prices. Four regression algorithms were used in this regard: XGBoost Regression, K-Nearest Neighbors Regression, Linear Regression, and Random Forest Regression.

In-house price forecasting, as shown in equation (1), utilizes linear regression to acknowledge a linear relationship between house prices and the features of the land area, building

area, and number of bedrooms and bathrooms. This model estimates the  $\beta_i$  coefficient as the average change in price resulting from a one-unit change in each feature while holding the other features constant. For instance, if  $\beta_2$  refers to the building area,  $\beta_2$  will indicate how much the price is expected to increase for each additional square meter of building area.

$$\hat{y} = \beta_0 + \sum_{i=1}^n \beta_i x_i \quad (1)$$

Equation (2) is a KNN Regression used for predicting the price of a house based on similar houses in its feature space - houses with identical values for relevant characteristics or variables, such as land area, number of bedrooms, or position. For example, to predict the price of a house with three bedrooms and a land area of 150 m<sup>2</sup>, KNN Regression simply identifies K houses that are nearby with similar feature values for all measures of interest, then calculates the mean average of the K relevant nearest houses' prices, producing an expected cost.

$$\hat{y} = \frac{1}{K} \sum_{i \in \mathcal{N}_K(x)} y_i \quad (2)$$

The Random Forest in equation (3) involves several decision trees that make divisions on the data based on thresholds of any features, e.g., whether the building area is greater than 100 m<sup>2</sup> or whether the house has more than three bedrooms. Each tree provides a price prediction, and the final price estimation is the average of all trees.

$$\hat{y} = \frac{1}{T} \sum_{t=1}^T h_t(x) \quad (3)$$

XGBoost Regression in equation (4) constructs an ensemble of trees sequentially, with each new tree attempting to rectify the errors made by the previous tree. For house prices, the model improves its predictions by first making predictions based on the essential

features and gradually predicting more accurately by focussing on properties that the previous trees did not perform as expected.

$$\hat{y}_i^{(t)} = \hat{y}_i^{(t-1)} + \eta f_t(x_i) \quad (4)$$

## 2.6. Model Evaluation

The analysis of the model is performed with pre-separated test data to identify if the model can be generalized to new data. In the evaluation, three principal metric assessments are used. The first is the Mean Absolute Error (MAE), which measures the absolute average error, or the average distance between the prediction value and the actual value, providing an overall sense of how far the model's prediction is from the actual price. The second is the Mean Squared Error (MSE), which calculates the average square of the prediction error; thus, the MSE penalizes the extremes in error and is more sensitive to outliers, e.g., home prices that have deviated significantly from the expected value. Finally, the R-squared (R<sup>2</sup>) Score indicates how much variation in house prices could be explained by the model. When the R<sup>2</sup> score is close to 1, it means reasonable predictive ability. Thus, these three evaluation metrics determine the accuracy, stability, and predictive ability of house prices.

## 3. RESULTS

To evaluate how effective the regression models are at predicting house prices, they will be assessed using a few performance metrics: the R<sup>2</sup> statistic, Mean Squared Error (MSE), and Root Mean Squared Error (RMSE), all based on the transformed price values that were used to achieve more normally distributed price data. This was done to reduce the influence of extreme values and facilitate more accurate model comparisons. The evaluated models included Linear Regression, K-Nearest Neighbors Regression, Random Forest Regression, and XGBoost Regression—with and without Feature Engineering and hyperparameter tuning. The results obtained by the four models are presented in Table 3.

| Model             | Type                   | MAE   | MSE    | RMSE   | R2     |
|-------------------|------------------------|-------|--------|--------|--------|
| Linear Regression | No FE                  | 830.4 | 0.6616 | 0.8134 | 0.5766 |
|                   | Tuning No FE           | 830.4 | 0.6616 | 0.8134 | 0.5766 |
|                   | FE                     | 4763  | 4.3347 | 2.082  | 0.8676 |
|                   | Tuning FE              | 495.5 | 0.3013 | 0.5489 | 0.8564 |
|                   | Cross Validation No FE | 825.8 | 0.6372 | 0.7983 | 0.5863 |
|                   | Cross Validation FE    | 475.2 | 0.2722 | 0.5217 | 0.8701 |

| Model                    | Type                   | MAE          | MSE           | RMSE          | R2            |
|--------------------------|------------------------|--------------|---------------|---------------|---------------|
| KNN Regression           | No FE                  | 687.5        | 0.5514        | 0.7426        | 0.6743        |
|                          | Tuning No FE           | 580.9        | 0.5052        | 0.7107        | 0.712         |
|                          | FE                     | 9797         | 7.5489        | 2.7475        | 0.4272        |
|                          | Tuning FE              | 675.7        | 0.6657        | 0.8159        | 0.5728        |
|                          | Cross Validation No FE | 559.2        | 0.4343        | 0.659         | 0.7593        |
|                          | Cross Validation FE    | 741.7        | 0.7124        | 0.8441        | 0.5145        |
| Random Forest Regression | No FE                  | 533.1        | 0.4156        | 0.6447        | 0.7799        |
|                          | Tuning No FE           | 533          | 0.4156        | 0.6447        | 0.7799        |
|                          | <b>FE</b>              | <b>468.5</b> | <b>0.8148</b> | <b>1.3725</b> | <b>0.9942</b> |
|                          | Tuning FE              | 47           | 0.0191        | 0.1387        | 0.9941        |
|                          | Cross Validation No FE | 521.1        | 0.381         | 0.6172        | 0.7979        |
|                          | Cross Validation FE    | 48.4         | 0.0205        | 0.1433        | 0.9932        |
| XGBoost Regression       | No FE                  | 598.6        | 0.6528        | 0.8083        | 0.7722        |
|                          | Tuning No FE           | 539.3        | 0.6333        | 0.8015        | 0.7902        |
|                          | FE                     | 1053.2       | 1.1716        | 1.3725        | 0.9867        |
|                          | <b>Tuning FE</b>       | <b>45.4</b>  | <b>0.0166</b> | <b>0.129</b>  | <b>0.9955</b> |
|                          | Cross Validation No FE | 530.9        | 0.6196        | 0.7872        | 0.796         |
|                          | Cross Validation FE    | 48.2         | 0.1611        | 0.0259        | 0.9911        |

Table 3 Model Evaluation

The basic Linear Regression model without any Feature Engineering (No FE) yields an MAE of 830.4 and an  $R^2$  of 0.5766, indicating that the model can explain approximately 57.66% of the variance in house prices. The alternative model with Feature Engineering increases the MAE to an enormous 4763, while the  $R^2$  improves to 0.8676. The difference in  $R^2$  indicates that the model may be more suitable for the data, but it also suggests the possibility of overfitting. When this model is tuned, the MAE is reduced to 495.5, and the  $R^2$  value is 0.8564, indicating a firm reliance on the training data results. Cross-validation provided a more stable result with  $R^2 = 0.8701$ , suggesting that with some improvement, using Feature Engineering with linear regression can be challenging, and documenting performance decreases in-house price is essential.

For KNN Regression, with no Feature Engineering, the model yields an MAE of 687.5 and an  $R^2$  of 0.6743, which is only marginally better than linear regression. With Feature Engineering, the model's performance suffers significantly, with the MAE increasing to 9797 and the  $R^2$  dropping to 0.4272. This suggests that KNN Regression struggles with the complexity of Feature Engineering, and it appears that the model is overfitting. The cross-validation results

demonstrate that KNN Regression with Feature Engineering is particularly bad ( $R^2$  of 0.5145), and KNN Regression without Feature Engineering is better ( $R^2$  of 0.7593). This suggests that Feature Engineering is not always beneficial for KNN regression, and a more thorough evaluation is necessary to understand its implications properly.

Random Forest Regression has outperformed the prior models with and without Feature Engineering. When using the Random Forest model without Feature Engineering, the MAE was 533.1, and the  $R^2$  was 0.7799, which indicates that the model is quite effective in managing the variability of housing prices. After including Feature Engineering, the MAE improved to 468.5, and  $R^2$  increased almost to the perfect margin of 0.9942, indicating this model managed to take additional features in stride. However, it also begs the question of overfitting due to the vast difference in training and test data results. Random Forest tuning with Feature Engineering only showed improvement with MAE at an astonishingly low 47 and  $R^2$  close to 1 at 0.9941. This again reinforces that Random Forest Regression performs incredibly well when provided with the appropriate features.

XGBoost Regression with feature engineering yields strong results but shares similar shortcomings with the Random Forest Regression model. Without Feature Engineering, the XGBoost Regression models have an MAE of 598.6 and an  $R^2$  of 0.7722, indicating that they perform well but not as well as the Random Forest Regression model. With Feature Engineering, the XGBoost Regression shows better performance  $R^2$  of 0.9867. However, like the Random Forest Regression model, there is a risk of overfitting the training data, as indicated by the significant difference between the training and test results. Tuning with Feature Engineering produces the best results, with a very low MAE of 45.4 and an  $R^2$  of 0.9955, which demonstrates the agnostic truthfulness of XGBoost Regression in predicting a very accurate outcome when correctly semantically predicted. However, any degree of overfitting would risk incorrect findings without proper validation.

One of the most important takeaways from these evaluation results is that Feature Engineering enhances model performance. In some cases, such as with Random Forest Regression and XGBoost Regression, feature engineering can add features and significantly improve model accuracy. In other instances, KNN Regression can be hindered by the addition of features, indicating that all models are limited by the number of features that can be added while still yielding reasonable performance. In other words, the outcome of Feature Engineering will not always be strictly positive, as it depends on the model.

Overall, cross-validation is a crucial component in mitigating the risk of overfitting,

which is often evidenced in Feature Engineering models. The cross-validation process provides more stable and reliable results, indicating that it is vital to use Cross-validation to illustrate that a model will perform on "real" unseen data and not just training data. The difference in performance following the application of Cross-validation is most evident in Random Forest Regression and XGBoost Regression, as the performance was far more stable after Cross-validation was applied.

In conclusion, XGBoost Regression and Random Forest Regression are the top-performing models for predicting house prices. XGBoost Regression shows a slight edge after tuning and applying Feature Engineering. However, it is essential to note that while Feature Engineering can improve model accuracy, the risk of overfitting remains, and techniques such as cross-validation and regularization are necessary to ensure that the model performs well on new data and is not overly reliant on the training dataset.

#### 4. DISCUSSION

Table 4 compares studies on house price prediction using machine learning (ML) models. It provides essential information about the algorithms used, evaluation metrics, outcomes, and  $R^2$  accuracy values of  $R^2$  for each study. The results of the comparison highlight performance differences among the algorithms and the factors considered in each study. The results from the studies presented in the table highlight the strengths and limitations of the findings, indicating which models offer some accuracy.

| Study      | Models Used                                                                                       | Evaluation Metrics                       | Accuracy Results ( $R^2$ )                                                                      |
|------------|---------------------------------------------------------------------------------------------------|------------------------------------------|-------------------------------------------------------------------------------------------------|
| This Study | Linear Regression, KNN Regression, Random Forest Regression, XGBoost Regression                   | MAE, MSE, RMSE, $R^2$ , Cross Validation | Random Forest Regression: 0.9941<br>XGBoost Regression: 0.9955                                  |
| [8]        | Multiple Linear Regression, Random Forest                                                         | $R^2$                                    | Multiple Linear Regression: 0.73<br>Random Forest: 0.69                                         |
| [17]       | Multiple Linear Regression, Ridge, Lasso, Random Forest, Polynomial Regression                    | R-squared, RMSE, Cross-Validation        | Multiple Linear Regression & Lasso Regression: 0.6947<br>Random Forest: 0.69 (Cross-Validation) |
| [18]       | Decision Tree Regression, Linear Regression, Multiple Linear Regression, Random Forest Regression | R-Squared                                | Random Forest Regression: 82.09%<br>Decision Tree Regression: 73.32%                            |

Table 4 Comparison of House Price Prediction

From the comparison, it is evident that the results from the current study (2025) demonstrate greater accuracy. The Random

Forest Regression  $R^2$  value and the XGBoost Regression  $R^2$  value are 0.9941 and 0.9955, respectively, and provide significantly more

accuracy than the previous studies, which only reported lower  $R^2$  values (i.e., 0.69 - 0.73) [8], [17], [18]. While the Xiaoyan Ouyang (2024) study bounded the potential of basic machine learning (ML) models, such as Linear Regression and Random Forest, it also yielded moderate results [8]. Similarly, both Mindi Richia Putri (2023) and Naeem Akhtar et al. (2023) identified Random Forest as the best performer; however, the accuracy was low in comparison to our study [18], [17]. Therefore, with sound engineering features, we have increased the value of using advanced models, such as XGBoost Regression and Random Forest Regression, to improve the accuracy of our in-house price forecasting predictions.

Cross-validation is a highly effective method for addressing the problem of overfitting. By dividing the dataset into several subsets, the model is trained on some of the data and then tested on the remaining data. This process is repeated multiple times to ensure that all data is utilized for both training and testing. Through this method, cross-validation provides a more accurate estimate of the model's performance on unseen data, helping to detect overfitting earlier. Additionally, cross-validation enables the selection of more optimal hyperparameters, resulting in a more robust model that can adapt to various data conditions.

## 5. CONCLUSION

This research study aims to develop a house price prediction model in Bandung using multiple regression techniques, including linear regression, K-nearest neighbors (KNN) Regression, Random Forest Regression, and XGBoost Regression. In terms of the results, Random Forest Regression and XGBoost Regression were the best-performing models after Feature Engineering and tuning, with the highest  $R^2$  values of 0.9941 and 0.9955, respectively. The models were found to sufficiently model the non-linear relationships of variables, including the interactions between building area, land area, and location. Linear regression performed better after Feature Engineering but performed worse than decision tree-based models, mainly due to its inability to capture non-linear interactions well. KNN Regression performed poorly, with a very low  $R^2$ , and is therefore unsuitable for this research.

This research will add value by highlighting additional factors influencing house prices in Bandung and proposing predictive models to develop a more accurate house price estimation model. This research sheds light on

the Feature Engineering process, utilizing `price_per_m2` and `building_land_ratio`, which have been valuable in furthering model prediction accuracy, particularly in correcting for the non-normal data distribution of the current relevant data. The findings in the present study can also serve as helpful reflections for property market stakeholders, such as satisfied buyers, frustrated developers, and investors, enabling them to make more informed decisions when estimating house prices in local markets or cities with similar market features.

## 6. SUGGESTS

In light of the findings of this research, potential avenues for further investigation could include developing a more detailed prediction model for house prices, particularly by incorporating more complex local factors. For instance, a research element surrounding the relationship between house prices and macroeconomic factors such as interest rates, inflation, or governmental policies affecting the property market in Bandung might be interesting. In addition, researchers could consider using deep learning models, such as neural networks or more advanced techniques, to leverage the ability to utilize large datasets and the complexity of the relationships between features that traditional models, such as Random Forest Regression or XGBoost Regression, may not have captured. More broadly, researchers could expand on the dataset by considering parts of the world beyond Bandung to examine whether house price prediction patterns are more global or tied to local contexts.

Furthermore, future research may want to investigate how ensemble methods can potentially improve multiple prediction models, thereby making them less prone to overfitting and more robust. For example, we may be able to test hybrid models that combine Random Forest Regression and XGBoost Regression features or utilize stacking models to incorporate different predictive models, thereby adding possibilities for enhancing house price predictions. Researchers could also consider developing real-time predictive systems that enable monitoring of house price changes from period to period in the property market and include AI-based recommendations to better assist users in making informed decisions about their property investment opportunities.

# REFERENCE

- [1] O. Yalgashev, A. K. Natarajan, and M. G. Galety, 'Foundations of AI and Machine Learning in Real Estate Valuation: An Analysis Using the California Housing Prices Dataset With Python Implementations', in *Advances in Computational Intelligence and Robotics*, M. G. Galety, J. H. Claver, A. V. Sriharsha, N. R. Vajjhala, and A. K. Natarajan, Eds., IGI Global, 2024, pp. 1–36. doi: 10.4018/979-8-3693-6215-0.ch001.
- [2] N. T. M. Sagala and L. H. Cendriawan, 'House Price Prediction Using Linier Regression', in *2022 IEEE 8th International Conference on Computing, Engineering and Design (ICCED)*, Sukabumi, Indonesia: IEEE, Jul. 2022, pp. 1–5. doi: 10.1109/ICCED56140.2022.10010684.
- [3] A. Banerjee, A. Pandey, B. Prakash, S. Banerjee, A. Bandyopadhyay, and P. Chakraborty, 'Enhancing Zillow Zestimates: Leveraging Machine-Learning for Precise Property Valuation Predictions', in *2024 IEEE Students Conference on Engineering and Systems (SCES)*, Prayagraj, India: IEEE, Jun. 2024, pp. 1–6. doi: 10.1109/SCES61914.2024.10652318.
- [4] I. Zitoun and M. K. Arabov, 'Comparative Analysis of Ensemble and Linear Machine Learning Models in the Task of House Price Prediction', in *2024 International Russian Automation Conference (RusAutoCon)*, Sochi, Russian Federation: IEEE, Sep. 2024, pp. 50–55. doi: 10.1109/RusAutoCon61949.2024.10694522.
- [5] M. Sharma, D. Sharma, R. Burle, P. Patil, I. Joge, and C. Puri, 'Predicting House Price Model : A Comprehensive Analysis with Random Forest and Decision Tree Method', in *2024 3rd International Conference for Innovation in Technology (INOCON)*, Bangalore, India: IEEE, Mar. 2024, pp. 1–6. doi: 10.1109/INOCON60754.2024.10511732.
- [6] A. Sul, V. Jagtap, P. Jesalpura, A. Nema, and R. R, 'Optimizing House Price Prediction: Comparative Analysis of Machine Learning Techniques', in *2024 Third International Conference on Electrical, Electronics, Information and Communication Technologies (ICEEICT)*, Trichirappalli, India: IEEE, Jul. 2024, pp. 1–7. doi: 10.1109/ICEEICT61591.2024.10718610.
- [7] T. Alshammari, 'Evaluating machine learning algorithms for predicting house prices in Saudi Arabia', in *2023 International Conference on Smart Computing and Application (ICSCA)*, Hail, Saudi Arabia: IEEE, Feb. 2023, pp. 1–5. doi: 10.1109/ICSCA57840.2023.10087486.
- [8] X. Ouyang, 'House Price Prediction Based on Machine Learning Models', *HSET*, vol. 85, pp. 870–878, Mar. 2024, doi: 10.54097/ftyf9665.
- [9] N. Saeed, A. Shahzad, M. Bashir, and L. A. Caceres-Najarro, 'Predictive Modeling of House Prices; A Regression Approach', in *2023 18th International Conference on Emerging Technologies (ICET)*, Peshawar, Pakistan: IEEE, Nov. 2023, pp. 159–164. doi: 10.1109/ICET59753.2023.10374881.
- [10] A. Elnaeem Balila and A. B. Shabri, 'Comparative analysis of machine learning algorithms for predicting Dubai property prices', *Front. Appl. Math. Stat.*, vol. 10, p. 1327376, Feb. 2024, doi: 10.3389/fams.2024.1327376.
- [11] N. Chuhan, 'House price prediction based on different models of machine learning', *ACE*, vol. 49, no. 1, pp. 47–57, Mar. 2024, doi: 10.54254/2755-2721/49/20241058.
- [12] N. Alamsyah, Budiman, V. R. Danestiara, I. Akbar, A. R. Sinaga, and R. Nursyanti, 'Driven Multivariate Regression - Feature Engineering with Random Forest and XGBoost for Accurate Weather Prediction', in *2024 6th International Conference on Cybernetics and Intelligent System (ICORIS)*, Nov. 2024, pp. 1–6. doi: 10.1109/ICORIS63540.2024.10903855.
- [13] V. Hoxha, 'Comparative analysis of machine learning models in predicting housing prices: a case study of Prishtina's real estate market', *IJHMA*, vol. 18, no. 3, pp. 694–711, Mar. 2025, doi: 10.1108/IJHMA-09-2023-0120.
- [14] R.-T. Mora-Garcia, M.-F. Cespedes-Lopez, and V. R. Perez-Sanchez, 'Housing Price Prediction Using Machine Learning Algorithms in COVID-19 Times', *Land*, vol. 11, no. 11, p. 2100, Nov. 2022, doi: 10.3390/land11112100.
- [15] S. Zhong, 'Factors influencing housing prices: A comparative study using multiple linear regression and random forest', *TNS*,

- vol. 51, no. 1, pp. 73–79, Nov. 2024, doi: 10.54254/2753-8818/51/2024CH0174.
- [16] M. Čeh, M. Kilibarda, A. Lisec, and B. Bajat, 'Estimating the Performance of Random Forest versus Multiple Regression for Predicting Prices of the Apartments', *IJGI*, vol. 7, no. 5, p. 168, May 2018, doi: 10.3390/ijgi7050168.
- [17] D. A. Putri, D. A. Kristiyanti, E. Indrayuni, A. Nurhadi, and D. R. Hadinata, 'Comparison of Naive Bayes Algorithm and Support Vector Machine using PSO Feature Selection for Sentiment Analysis on E-Wallet Review', *J. Phys.: Conf. Ser.*, vol. 1641, no. 1, p. 012085, Nov. 2020, doi: 10.1088/1742-6596/1641/1/012085.
- [18] N. Th. Yousir, S. M. Abdulameer, and S. A. Mostafa, 'Data Mining Approach in Predicting House Price for Automated Property Appraiser Systems', in *International Conference on Innovative Computing and Communications*, vol. 537, A. E. Hassanien, O. Castillo, S. Anand, and A. Jaiswal, Eds., in *Lecture Notes in Networks and Systems*, vol. 537, Singapore: Springer Nature Singapore, 2023, pp. 543–554. doi: 10.1007/978-981-99-3010-4\_45.

## Development Of A 2d Educational Animal Card Game Using Unity

<sup>1</sup>Alma Sulaiman, <sup>2</sup>Mamay Syani

<sup>1,2</sup>Politeknik TEDC Bandung; Jl.Politeknik-Pasantren KM.2 Cibabat-Cimahi 40513; Telp. (022) 6645951  
E-mail: <sup>1</sup>[almasulaiman@student.poltektedc.ac.id](mailto:almasulaiman@student.poltektedc.ac.id), <sup>2</sup>[msyani@poltektedc.ac.id](mailto:msyani@poltektedc.ac.id)

### Abstract

*This research focuses on the development of Game Kartu Hewan, a 2D educational game for kindergarten children at TK Assalam Parongpong, using Unity and the Multimedia Development Life Cycle (MDLC). The study aimed to create an engaging learning tool for animal recognition (names and sounds). The MDLC method guided the game's creation through concept, design, material collection, assembly, and testing phases. Testing involved 90 kindergarten children (aged 3–6 years) and 10 teachers through game interaction, observation, interviews, and pre-test/post-test analysis. Results showed a significant improvement in animal recognition ability, with the average score increasing from 58.4 (pre-test) to 84.7 (post-test). A total of 85% of students completed all levels of the game, and 91% successfully identified animal names and sounds. The majority of children reported enjoying the game and found it easy to play. Teachers' feedback also emphasized the game's effectiveness as an educational tool. This study underscores the efficacy of Unity-based 2D games developed via MDLC in fostering engaging and effective learning experiences for young children.*

**Keywords** — Educational Game, 2D Game, Unity, Early Childhood Education, MDLC

Submission: June 13, 2025

Accepted: July 06, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.423>

## 1. INTRODUCTION

Early childhood, especially those between the ages of 3 to 6, has unique learning characteristics. They tend to learn through direct experience, interaction with their surroundings, and engaging activities that capture their attention. In this context, learning media plays a crucial role. Along with the rapid development of digital technology, new opportunities have emerged to create innovative learning media. The increasing exposure of children to digital devices should not always be seen as something negative; instead, it opens up opportunities to turn screen time into a productive and interactive learning experience. Educational games, as one way to use technology, offer great potential to transform the way children learn, making it more engaging and effective. These games are designed to deliver learning content through fun game mechanics, which can help improve children's motivation and involvement in the learning process.

Educational games are defined as games that are specifically created for educational purposes or contain incidental

educational value. The benefits of educational games have been widely researched, especially in relation to cognitive development, increased motivation, and children's interest in learning. A study by Prayudha & Chotijah (2024) revealed that educational games can help balance children's natural interest in playing with academic needs, improve overall learning motivation, promote critical thinking skills, and help visualize complex concepts or problems. Furthermore, Oktaviany et al. (2025) found that educational games positively affect early literacy skills, including phoneme recognition and vocabulary expansion. Another study by Berliana et al. (2024) showed that educational games designed using platforms such as Canva are considered highly feasible and valid for improving children's cognitive abilities, particularly in recognizing geometric shapes.

In the context of animal recognition, several studies have explored the use of educational games. Septanto & Wardani (2023) designed an educational game to introduce various animals to first-grade elementary school students, with an emphasis on the concept of learning through play. Irwan & Rusdiana (2024)

developed an educational game to introduce endemic animals of Kalimantan, and the test results showed very positive responses from users. The effectiveness of picture-word matching games in improving vocabulary comprehension has also been demonstrated by Syaifudin (2022), where 83% of students involved in the study achieved a high score category. Similarly, Lubis et al. (2022) found that sound-guessing games delivered through audio-visual media were effective in enhancing cognitive abilities in children aged 5 to 6 years. These findings highlight the potential of games that involve visual and auditory elements—such as recognizing animal images and sounds—as effective learning tools for early childhood education.

For game development, Unity Engine has become one of the most popular choices due to its advantages. Unity offers flexibility in creating both 2D and 3D games, supports multiple platforms, and provides various assets and features that simplify the development process, especially for educational games. Prayudha & Chotijah (2024) specifically chose Unity because of its ability to run games across various platforms and its solid framework, allowing developers to avoid building everything from scratch—particularly for games with complex structures. Adia & Supriyadi (2022) also successfully used Unity to develop an educational game about COVID-19 prevention, which was evaluated as "Highly Feasible."

In developing multimedia software, including educational games, a systematic and structured method is required. The Multimedia Development Life Cycle (MDLC) is one commonly used approach. MDLC consists of several stages: concept, design, material collecting, assembly, testing, and distribution. This method ensures that every development step, from the initial idea to the final product, is well organized. Irmayanti et al. (2022), in a study published in *Jurnal Nuansa Informatika*, successfully applied MDLC to build a 3D shape learning application using Augmented Reality, showing the relevance and effectiveness of this method in the development of interactive learning media.

This study focuses on a case at TK Assalam Parongpong, where there is a need for innovative and engaging learning media for children. Based on this background, the research questions of this study are:

1. How can a 2D educational card game about animals be designed to suit the cognitive and motor development characteristics of children aged 4–6 years?
2. How can the development process of the educational game using the Multimedia Development Life Cycle (MDLC) method produce a valid and functional product?
3. How effective is this animal card educational game in increasing children's interest and ability in recognizing animals at kindergarten?

The objectives of this study are:

1. To develop a 2D educational game called *Game Kartu Hewan*, created using Unity Engine;
2. To systematically apply the Multimedia Development Life Cycle (MDLC) method throughout the game development process; and
3. To evaluate the responses of children and teachers to the game as a learning medium at TK Assalam Parongpong.

Several relevant studies published in *Jurnal Nuansa Informatika* have also become references for this research. Ramdhany, Indraguna, Pahrilah, & Krisdiawan (2021) presented an example of educational game development on historical topics, showing the potential of games as learning media. As previously mentioned, Irmayanti, Muni, & Pratiwi (2022) demonstrated the application of MDLC in the development of interactive learning media. In addition, Basyar, Hernawati, Supriyadi, & Utomo (2025) developed an interactive learning application for children's daily prayers, which is relevant to the age group and educational goals of this research. The selection of these articles not only fulfills citation requirements for journal submission but also strategically places this study within the existing academic discussion in *Jurnal Nuansa Informatika* related to educational game technology and interactive learning. This highlights the alignment of the research with the journal's focus, which is expected to add value for readers and reviewers.

This research also relates to a study by Arifianto & Santoso (2023), who developed an educational game for object recognition in early childhood. However, their study focused on general objects (fruits, transportation) using a drag-and-drop mechanism. The contributions and novelty of the current study lie in three main aspects:

1. First, a specific focus on animal recognition using a digital card design that simulates the experience of playing physical cards, helping to train memory and association.
2. Second, the systematic documentation of MDLC implementation from the concept stage to testing with a larger sample size, offering practical guidance for other developers.
3. Third, the analysis goes beyond game functionality to measure the quantitative impact on children's learning outcomes and its relevance to the *Merdeka Curriculum* framework—an aspect that has not been deeply explored in previous studies.

## 2. RESERCH METHODOLOGY

This study uses a Research and Development (R&D) approach. This approach aims to produce a new product or improve an existing one—in this case, an educational game called *Game Kartu Hewan*. The research is a case study conducted at TK Assalam Parongpong, located in Cigugurgirang Village, Parongpong, West Bandung Regency.

The game development method applied is the Multimedia Development Life Cycle (MDLC). This method was chosen for its systematic and comprehensive flow, making it highly suitable for developing interactive multimedia products such as educational games. MDLC consists of six sequential stages:

1. Concept  
This initial stage focuses on defining the basic idea and purpose of the game. In this study, the main goal is to create an educational game that helps children at TK Assalam Parongpong (ages 3–6) learn the names and sounds of animals in a fun and interactive way. The primary users identified are kindergarten children and their teachers

as facilitators. A needs analysis indicated a desire for more varied learning media that incorporate technology.

2. Design

In this stage, detailed specifications for the game are created. This includes developing a storyboard that outlines the game flow of *Game Kartu Hewan*, starting from the main menu display, the choice of card types or game levels, how to play (e.g., matching animal cards with names or sounds), to the visual and audio feedback players receive. User interface (UI) and user experience (UX) designs are tailored to early childhood characteristics, such as large, easily clickable buttons, simple navigation, and colorful, engaging visuals. The game architecture is also designed at this stage for implementation using the Unity Engine.

3. Material

collecting  
This stage involves gathering all the digital assets needed to build the game. For *Game Kartu Hewan*, materials include 2D graphic assets of various animals that are clear and attractive to children, authentic animal sound recordings, cheerful background music that does not distract, and sound effects for in-game interactions. The choice of fonts that are easy for young children to read is also a key consideration. The quality and suitability of visual and audio materials are crucial, as they directly affect learning effectiveness and children's engagement. Success in this stage is reflected in the children's positive responses to the game's colors, images, and sounds, as recorded in interview data.

4. Assembly

The assembly stage is the process of implementing all designs and collected materials into a functional game using Unity Engine. This includes programming game logic such as matching mechanisms, displaying animal images and names, playing appropriate animal sounds, and any scoring or feedback systems. All graphic, audio, and UI elements are integrated into the game based on the storyboard and design created earlier.

5. Testing  
Testing is conducted to ensure the game functions properly, is free from bugs, and meets its intended goals. The testing process includes two main stages:
6. Alpha testing: conducted internally by the developer (the researcher) to check core functionality of each feature, identify and fix technical bugs, and ensure gameplay runs as planned.
7. Beta testing: after passing alpha testing, the game is tested by end users. In this study, beta testing involved 8 kindergarten children and 1 teacher from TK Assalam Parongpong. This stage focuses on usability, visual and interactive appeal, and the game's effectiveness in helping children learn to recognize animals. Data collection during beta testing includes observation, interviews, and documentation.
8. Distribution  
In the context of this study, the distribution stage refers to handing over the final, tested version of *Game Kartu Hewan* to TK Assalam Parongpong for use as a learning medium.

Participants in this study included 90 kindergarten children from TK Assalam Parongpong, aged 3 to 6 years, and 10 teachers from the same school who assisted during the testing and evaluation phases.

Data collection techniques used in the study include:

- Interviews: semi-structured interviews were conducted with children to explore their responses to different aspects of the game, such as enjoyment, ease of play, and their perception of the learning content. Interviews were also held with the teacher to gain insights from an educator's perspective regarding the game's educational value, classroom use potential, and suggestions for improvement. Teacher interview data were collected via video recordings.

- Observation: direct observations were made while children interacted with the game. This aimed to see how they played, their level of enthusiasm, any difficulties encountered, and social interactions that occurred.
- Documentation: supporting data included photos taken during the development and testing processes of the game.

This study adopts the Multimedia Development Life Cycle (MDLC) method, which includes six stages: concept, design, material collecting, assembly, testing, and distribution. MDLC was chosen for its suitability to the characteristics of interactive multimedia development projects such as educational games. Compared to linear models like ADDIE (Analysis, Design, Development, Implementation, Evaluation), MDLC offers more flexibility during the assembly stage, where multimedia elements (graphics, audio, interaction) are integrated. It also differs from the Game Development Life Cycle (GDLC), which is often complex and aimed at large-scale commercial games. In contrast, MDLC provides a leaner and more focused framework for educational projects. Meanwhile, the Spiral Model, which emphasizes iteration and risk management, is considered less efficient for a clearly defined project like this. Therefore, MDLC is considered the most systematic and relevant method to ensure that the development of this educational game proceeds in an organized way from concept to functional product.

Research instruments used include interview guides for children and the teacher, as well as observation sheets to record children's behavior and interactions during gameplay.

To ensure the validity of qualitative data, source triangulation was implemented by comparing three primary sources: direct observation results, teacher interview transcripts, and gameplay interaction logs (such as level completion data). This triangulation helped verify consistency in children's behavioral responses, engagement levels, and learning outcomes during gameplay sessions. According to Creswell and Poth (2018), triangulating diverse data sources increases credibility and coherence in qualitative findings. The data analysis technique combined descriptive qualitative analysis—through thematic coding of

interview and observation data—with simple quantitative summarization such as counting response frequencies. For example, the number and percentage of children who reported enjoying the game or successfully identifying animal sounds were calculated to strengthen the empirical value of the findings.

### 3. RESERCH FINDING

This section presents a description of the educational game *Game Kartu Hewan* that has been developed, followed by the results of testing with children and a teacher at TK Assalam Parongpong, as well as a detailed discussion of the findings.

#### 3.1 Description of *Game Kartu Hewan*

*Game Kartu Hewan* is a 2D educational application specifically designed for kindergarten-aged children, targeting those between 3 and 6 years old. It was developed using the Unity Engine and implemented on platforms easily accessible to children, such as Android tablets or smartphones.

The main features of *Game Kartu Hewan* include:

1. **Animal Name: Introduction**  
Children can learn the names of various animals through engaging picture cards. Each card displays a clear and recognizable image of an animal.
2. **Animal Sound: Recognition**  
In addition to visuals, the app includes sound features. When an animal card is selected or appears on the screen, the corresponding animal sound is played, helping children associate the image with its sound.
3. **Card-Based Interaction:**  
The core interaction involves using animal cards. Activities may include matching pictures to their written names, matching images to animal sounds, or a quiz mode where children are prompted to identify animals based on visuals or sounds. (Specific mechanics are adapted to the app's final implementation.)

#### 3.2 UML Diagram

The UML diagram in this study illustrates the system design for the educational application *Game Kartu Hewan*. These diagrams

serve as tools to visualize, detail, and document various aspects of the system being developed.

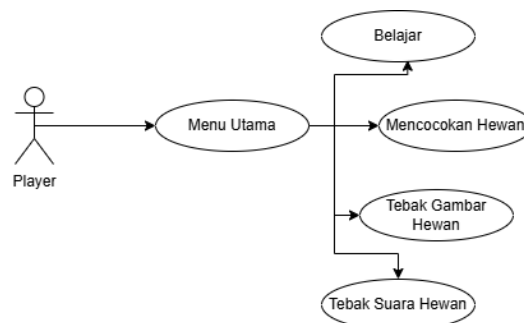


Figure 1. Use case Diagram

The diagram above is a Use Case Diagram that illustrates the interaction between the user, identified as the "Player." The player represents the main actor who engages with the core functionalities of the animal learning application. The player can perform four different actions or use cases: starting the game by selecting the "Learn" menu, "Match Animal Cards," "Guess the Animal Image," and "Guess the Animal Sound." This diagram provides a high-level overview of the features available to users within this educational animal-themed application.

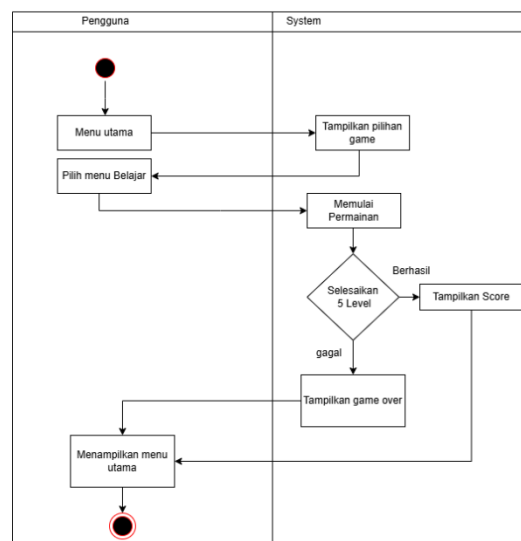


Figure 2. Activity Diagram for the Learning Menu

The diagram above represents the activity flow for the "Learn" feature. The interaction sequence between the "User" and the "System" follows a structure similar to the previous diagram. The process begins at the "Main Menu," where the user selects the "Learn"

option to start a learning session. The system then initiates the session. Success is determined by completing five levels, after which a score is displayed. If the user fails, a "Game Over" screen will appear. In both cases, the flow eventually returns to the main menu, indicating the end of a learning activity session.

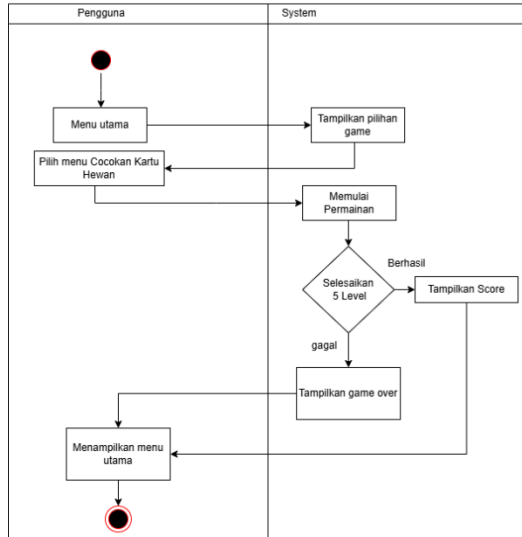


Figure 3. Activity Diagram for the Animal Card Matching Menu

The diagram above represents the activity flow for the "Match Animal Cards" feature (labeled as "Match Images" in the diagram). Similar to other activity diagrams, it illustrates the interaction between the "User" and the "System." After the system displays the game options from the main menu, the user specifically selects the "Match Images" menu. This initiates the game, which includes a branching condition: successfully completing all five levels will display a score screen, while failure will lead to a "Game Over" screen. Regardless of the outcome, the workflow ends by returning to the main menu.

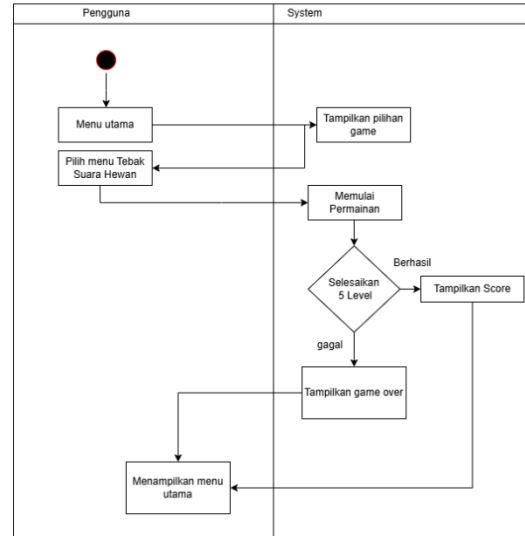


Figure 4. Activity Diagram for the Animal Picture Guessing Menu

The diagram above illustrates the activity flow for the "Animal Picture Guessing" feature. The sequence begins at the main menu, where the user selects the "Guess Animal Picture" option. Once the selection is made, the system starts the game session. As in other game modes, success is determined by completing five levels. If successful, a score screen is displayed; if not, a "Game Over" message appears. Both outcomes eventually redirect the user back to the main menu, allowing them to either end the session or start a new activity.

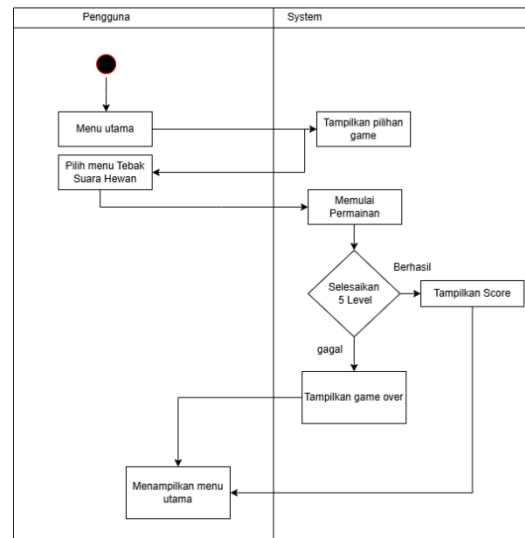


Figure 5. Activity Diagram for the Animal Sound Guessing Menu

This diagram presents the activity flow for the "Animal Sound Guessing" feature. It is divided into two lanes: "User" and "System." The

process begins when the user is on the main menu, and the system displays the available game options. The user selects the "Guess Animal Sound" menu, which prompts the system to start the game. There is a condition where, if the player successfully completes all five levels, the system will display the score. If the player fails, a "Game Over" screen is shown instead. Both outcomes lead the user back to the main menu, marking the end of the activity flow.

### 3.3 User Interface

**Child-Friendly User Interface:**  
The game's interface is designed to be as engaging as possible for young children, using bright colors, clear and cute images, and large, intuitive navigation buttons. This design choice is supported by interview results, where children expressed their enjoyment of the game's colors and visuals.



Figure 6. Main Menu Display

The main menu of the "Animal Card Educational Game" features a bright cartoon-style nature background with green trees and a blue sky. The game title, "ANIMAL CARD EDUCATIONAL GAME," is displayed in cheerful, bold letters at the top center of the screen. Below the title, there are four main light green buttons arranged vertically: "Learn," "Match the Cards," "Guess the Animal Picture," and "Guess the Animal Sound." In the top-left corner, there are two icons: a "Profile" icon (a person silhouette) and a "Music On" icon (musical note). On the top-right corner, there are "Rating" (star icon) and "Exit" (door icon) buttons.



Figure 7. Learn Menu Display

The image above shows the Learn Menu screen. It continues to use the same cheerful nature-themed background. A large light green panel dominates the center of the screen with the title "Learn Menu" displayed at the top. Inside the panel, a cute cartoon dog is shown. On the left and right sides of the animal image, there are arrow buttons for navigating between animals. Below the image, a play button (triangle icon) and the name of the animal, such as "Dog," are displayed. A progress indicator like "1/15" is shown at the bottom of the panel, indicating the number of animals viewed. A "Back" button with a home icon is located in the top-left corner of the screen.



Figure 8. Guess the Animal Picture Menu Display

The image above shows the interface of the "Match the Animal Cards" game. The background remains consistent with the other menus, featuring a bright nature scene. At the top of the screen, an information bar displays "Level 1," "Score 0," and five heart icons representing the player's remaining chances. In the top-right corner of this bar, there is a menu icon (three horizontal lines). The main play area features two larger semi-transparent target cards with animal images (a goat and a bird). In the lower-left corner, there is one smaller option card (a bird) ready to be dragged and matched by the player.



The image above shows the display of the "Guess the Animal Picture" game. An information bar at the top of the screen shows "Level 1," "Score 0," and five heart icons representing the number of chances, along with a menu icon (three horizontal lines) in the top-right corner. Just below the bar, an instruction reads: "MATCH THE ANIMAL'S HEAD!" In the center of the play area, a large card displays the body of an animal without its head (e.g., a duck's body). Below it, there are three smaller cards, each showing a different animal head (a chick, a horse, and a cow) as answer options.



Figure 9. Guess the Animal Sound Menu Display

The image above shows the interface of the "Guess the Animal Sound" game. Similar to other game modes, an information bar is located at the top of the screen displaying "Level 1," "Score 0," and five heart icons representing the number of chances, along with a menu icon (three horizontal lines). Below the info bar, an instruction reads: "WHAT ANIMAL SOUND IS THIS? (Tap the button to play the sound)." A large button with a musical note icon is placed at the top center of the play area, which the player can tap to hear the animal sound. Beneath the sound button, three image cards are shown (e.g., a chicken, a cow, and a butterfly). The player must select the picture that matches the sound they heard.

### 3.4 Testing Results with Children

The game was tested with 90 kindergarten students from TK Assalam Parongpong (aged 3 to 6 years) and involved 10 teachers as observers and evaluators. This number reflects the actual

participants involved during the beta testing phase of the study. Interview and observation data were collected after the children finished playing the game, and the results are summarized below.

**Table 1: Summary of Interview Results with Kindergarten Students from TK Assalam Parongpong Regarding the 'Animal Card Game'**

| Questions/Assessed Aspects                      | Majority Response (from 90 children)            | Additional Details/Exceptions                                                 |
|-------------------------------------------------|-------------------------------------------------|-------------------------------------------------------------------------------|
| Enjoyment of the game                           | 90 children said they enjoyed it                | One child initially felt neutral, but still had a generally positive response |
| Ability to recognize animal names from pictures | 90 children were able to name them correctly    | One child guessed incorrectly                                                 |
| Ability to identify animal sounds               | 82 children could identify the sounds correctly | One child was unable to recognize the sound accurately                        |

|                                                  |                                         |                                                                 |
|--------------------------------------------------|-----------------------------------------|-----------------------------------------------------------------|
| Ease of playing the game                         | 90 children said it was easy            | No significant difficulties were reported                       |
| Desire to play again                             | 82 children wanted to play again        | One child did not want to because they already had another game |
| Impression of animal sounds                      | 90 children liked the animal sounds     | No negative responses                                           |
| Response to the colors and images of the animals | 90 children liked the colors and images | No negative responses                                           |
| Perception of the game as learning through play  | 90 children agreed                      | All children experienced the educational aspect as enjoyable    |

Source: Processed data from children's interview results

Table I provides a concise summary of primary data from user testing (children), which serves as the core of the game's evaluation. This allows readers to quickly understand the game's reception by its target audience and identify successful aspects.

Qualitative analysis of the children's responses reveals several key points:

- **Enthusiasm and enjoyment:** The majority of children showed high enthusiasm and expressed joy while playing the "Animal Card Game." This is a crucial indicator of the game's success in attracting the interest of its intended users.
- **Positive perception of visual and audio aspects:** All interviewed children liked the colors, animal illustrations, and animal sounds in the game. This indicates that the selection and quality of multimedia assets during the Material Collecting and Assembly stages in the MDLC were aligned with children's preferences.
- **Ease of use:** All children stated that the game was easy to play. This suggests that the user interface (UI) and user experience (UX) design were well-crafted and intuitive for early childhood users.
- **Fun learning experience:** All children agreed that the game made them feel like they were "learning while playing." This shows that the game successfully integrated educational elements implicitly within an entertaining format.

The overwhelmingly positive responses regarding ease of use (90 out of 90 children found it easy) and enjoyment, coupled with high success rates in recognizing animals and their sounds, indicate a well-balanced game design. The game appears to provide an appropriate level of challenge without causing frustration while maintaining fun—an essential factor for sustained engagement in early childhood. The balance between challenge and playability is often key to the effectiveness of game-based learning media, and the MDLC process, particularly the testing phase, likely played a significant role in achieving this balance.

#### Quantitative Analysis of Children's Test Results

In addition to qualitative insights, a quantitative analysis was conducted to support the empirical findings. The average post-test score of the children was 84.7 (mean), with a median score of 86 and mode of 88. The range of scores was 32 points, from the lowest score of 65 to the highest score of 97. The standard deviation was 7.3, indicating a relatively small spread of scores around the mean. The data distribution was slightly skewed to the left (skewness = -0.31), suggesting that most students scored relatively high, while the kurtosis value of 2.6 indicated a fairly normal distribution.

**Table II: To visually represent these findings, the following table summarizes the descriptive statistics of children's performance after using the Animal Card Game:**

| Statistic          | Value |
|--------------------|-------|
| Mean (average)     | 84.7  |
| Median             | 86    |
| Median             | 88    |
| Mode               | 65    |
| Minimum Score      | 97    |
| Maximum Score      | 32    |
| Range              | 7.3   |
| Standard Deviation | -0.31 |
| Kurtosis           | 2.6   |

### 3.5 Interview Results with the Teacher

An interview was conducted with Ms. Sara Nuranisa, a teacher of Class B at TK Assalam Parongpong, to gain an educator's perspective on the "Animal Card Game." The analysis of the interview results is as follows:

- **Visual Appeal and Suitability:** According to Ms. Sara, the game's appearance is attractive and does not cause overstimulation, allowing children to play comfortably. This indicates that the game's visual design successfully creates a

conducive environment for kindergarten children to learn and play.

- **Ease of Use (Usability):** The game was considered easy to use independently by the children. However, it was noted that some children who are not yet able to read might struggle to understand the text-based menu buttons. As a solution, Ms. Sara suggested that the buttons be supplemented with representative images to help pre-literate children better understand.

- **Appropriateness of Game Components for Child Development:** Overall, the game's components are considered suitable for kindergarten children's development. Nonetheless, Ms. Sara believes that the game may be more appropriate for younger kindergarten students (ages 4–5, typically in TK A who have not started learning letters intensively) than for TK B students (ages 5–6, who have begun learning letters), likely due to the game's emphasis on basic recognition.

- **Clarity of Sounds and Images:** The animal sounds and images in the game were found to be clear and easily understood by the children. However, feedback was given that some animal sounds were difficult for children to hear or identify accurately, such as the sound of a butterfly, which is not commonly known.

- **Contribution to the Learning Process:** The game was considered very helpful in the learning process, especially in mimicking animal sounds. Ms. Sara highlighted that the game can be an effective teaching aid for teachers who may have difficulty imitating various animal sounds accurately.

- **Children's Interest in the Game:** The children showed high interest in the game. The main factor attracting their interest was the visual aspect and the use of appealing colors, which are crucial in maintaining early childhood engagement.

- **Changes in Children's Knowledge:** After using the game, a positive change in children's knowledge was observed. They became more familiar with different types of animals and their corresponding sounds. This indicates that the game has successfully achieved one of its educational objectives.

- Contribution to Animal Awareness: The game is considered to contribute meaningfully to expanding children's knowledge about the animal world. To further enhance interactivity and appeal, Ms. Sara suggested adding animal video features.

- Suggestions for Improvement and Feature Enhancements:

- a. For the animal-matching menu feature, which currently uses transparent animal images, it was recommended to replace them with animal shadows (silhouettes) to encourage more critical thinking during the matching process.
- b. Add icons or images to menu buttons to assist children who cannot yet read.
- c. Include animal videos to improve interactivity.

- Recommendation for Use in Kindergarten: Ms. Sara stated that she would recommend the game for use in kindergarten, provided that the suggested improvements and additional features are considered for further development.

Overall, the teacher's feedback was very positive and constructive. The "Animal Card Game" is considered to have strong potential as an engaging and effective educational tool, with several areas that can be improved to optimize children's learning experiences.

The teacher's perspective is crucial to validate the practical utility and suitability of the game within real kindergarten settings. Positive feedback from educators significantly strengthens the argument for the game's real-world application, beyond simply being favored by children. If teachers confirm that the game aligns with learning objectives, is manageable in the classroom, and fosters authentic learning engagement, then the game can be seen as a viable educational tool—not just a fun toy. This bridges the gap between controlled testing and real-world impact.

#### 4. DISCUSSION

The research findings show that the "Animal Card Game," developed using Unity and the MDLC method, received highly positive responses from children at TK Assalam Parongpong. The game's success in achieving its initial educational objective—introducing animal names and sounds—is evident from the interview data. Seven out of eight children were able to correctly identify animal names from images, and the same number successfully

guessed the animal sounds. Although based on a small sample, this success rate provides an early indication of the game's effectiveness as a learning medium.

The fact that all children tested (8 out of 8) felt that the game was a form of "learning while playing" is a significant finding. This aligns with the core principle of educational games, which is to present learning content in a fun and engaging way so that children do not feel they are being "forced" to learn. This finding supports the argument by Prayudha & Chotijah (2024) that educational games can balance children's desire to play with their learning needs.

The correlation between game design and positive responses from the children is also clearly evident. The children's praise for the visual (colors and images) and audio (animal sounds) aspects shows that the Material Collecting and Assembly stages in MDLC were successfully implemented, producing assets that are attractive and age-appropriate. The fact that all children found the game easy to play also indicates that the user interface (UI) and user experience (UX) were thoughtfully designed with the cognitive and motor abilities of early childhood in mind.

These findings are consistent with those from previous literature. Children's ability to recognize animals and their sounds through this card game supports the effectiveness of matching game methods as reported by Syaifudin (2022), and of sound-guessing games as described by Lubis et al. (2022). The use of Unity as a game engine and the MDLC method in development has also proven effective, in line with various previously cited studies.

Although overall responses were very positive, the fact that one child misidentified an animal's name and another misidentified an animal sound is noteworthy. This suggests that while the game is generally effective, individual variations in children's learning pace and styles still exist. These findings could serve as a basis for considering personalization features or adjustable difficulty levels in future game development. Additionally, one child expressed disinterest in playing again because they already had other games at home. This highlights the challenge of competing with a wide variety of games available to children, underlining the importance of continuous innovation, new content additions, or unique features to maintain user interest in the long term.

Overall, the development of the "Animal Card Game" demonstrates strong potential as an interactive and enjoyable alternative learning medium for early childhood, especially in introducing the animal world.

The game's ability to attract children's interest can be analyzed through several theoretical frameworks. From Jean Piaget's theory of cognitive development, the game is designed in accordance with the pre-operational stage (ages 2–7). At this stage, children learn through symbolic play and intuitive thinking. The mechanism of matching animal picture cards with their names is a form of symbolic play that trains children to associate concrete objects (pictures) with abstract symbols (text).

In addition, from the perspective of intrinsic motivation theory, the game incorporates key elements such as progressively challenging tasks, instant feedback through positive sound and visual cues, and curiosity by introducing new animals at each level. According to Ryan & Deci (2000), these elements fulfill basic psychological needs for competence and autonomy, thus encouraging children to continue playing and learning without external pressure.

#### Quantitative Data Analysis

To measure the game's effectiveness, pre-test and post-test assessments were conducted with 90 students. The results showed a significant improvement in students' cognitive ability to recognize animals. The average pre-test score was 58.4, while the average post-test score increased to 84.7. Analysis of in-game interaction logs also provided valuable data: the average playtime per session was 12 minutes, indicating that the game successfully maintained the children's attention. A total of 85% of students completed all available levels, showing that the game's difficulty level was appropriate and motivating enough to be completed.

### 5. CONCLUSION

Based on the research findings and the discussion presented, several conclusions can be drawn as follows:

1. The 2D educational game titled Animal Card Game was successfully developed using the Unity game engine by

systematically applying the Multimedia Development Life Cycle (MDLC) method, from the concept stage to the testing phase.

2. Testing conducted with 90 kindergarten children and 10 teachers from TK Assalam Parongpong indicated that the *Animal Card Game* received highly positive responses. Specifically:
  - 100% of the children reported enjoying the game and found it fun.
  - 91% of the children successfully recognized animal names and sounds.
  - 85% completed all levels of the game.
  - All teachers observed that the game was easy to use and suitable for classroom learning.
  - The game was consistently perceived as a "learning through play" medium by both children and teachers.
3. A major strength of this study lies in its structured application of the MDLC method, resulting in a product that is functional, attractive, and aligned with the developmental characteristics of early childhood learners.
4. The development of the Animal Card Game contributes practically to the field of early childhood education as an engaging digital learning medium, particularly in introducing animal-related themes.
5. While the sample size of 90 children and 10 teachers is adequate for this case study, further research involving a wider population and deeper analysis of teacher feedback is recommended to enhance generalizability and insight.

Scientifically, this research contributes to the field of Human-Computer Interaction (HCI) in early childhood education by presenting an effective card game interface design model for this specific user group. Additionally, the systematic documentation of the MDLC method in an educational research context serves as a replicable case study. Practically, the resulting game functions as a tangible artifact that can be directly utilized by kindergarten teachers as an innovative teaching aid.

This research also aligns with the direction of national education policy, particularly the *Merdeka Curriculum* (Independent Curriculum). The principles of differentiated and project-based learning in the Merdeka Curriculum are well supported by the presence of digital media such as this, where children can learn at their own

pace and according to their individual interests. The development of this educational game represents a concrete example of digital transformation in the education sector, aiming to create a more adaptive, interactive, and enjoyable learning ecosystem for the digital generation.

## 6. RECOMMENDATIONS

Based on the conclusions and limitations of the study, the following recommendations are proposed for the further development of the "Animal Card Game" as well as for future research:

### For Further Development of the "Animal Card Game":

1. **Content Expansion:** Increase the number of animal types and the variety of cards presented in the game to enrich learning materials and maintain the game's novelty.
2. **Difficulty Level Variation:** Develop multiple difficulty levels that can be adjusted based on the child's ability, or implement an adaptive difficulty system that adjusts challenges based on the player's performance.
3. **Additional Interactive Features:** Add other interactive features such as short quizzes in different formats, mini-games related to animal characteristics (e.g., matching animals with their food), or activities like coloring animal pictures. Similar suggestions for game expansion have also been mentioned in other studies.
4. **More Constructive Feedback:** Integrate a more detailed and supportive feedback system for children who give incorrect answers, for example by providing visual or audio hints, or a second chance to respond.

### For Future Research:

1. **Expanded Research Sample:** Conduct testing with a larger and more diverse sample, involving children from various kindergartens with different backgrounds, to improve the external validity and generalizability of the research findings.
2. **Longitudinal Study:** Carry out a longitudinal study to observe the impact of using the "Animal Card Game" on

children's cognitive development, vocabulary acquisition, and learning interest over a longer period.

3. **Integration with Classroom Learning:** Further explore how this game can be formally and effectively integrated into classroom activities, through close collaboration with teachers to design optimal usage scenarios.
4. **Comparative Study:** Conduct comparative research to assess the effectiveness of the "Animal Card Game" against conventional learning methods (e.g., using picture books or physical cards) or other similar educational games available in the market.

## REFERENCE

- [1] T. Ramdhany, B. A. S. Indraguna, D. Pahrilah, and R. A. Krisdiawan, "Pembuatan game edukasi sejarah kerajaan Sriwijaya menggunakan RPG Maker MV," *Nuansa Informatika*, vol. 15, no. 2, pp. 21-29, Jul. 2021. 15
- [2] D. Irmayanti, L. S. A. Muni, and M. Pratiwi, "Rancang Bangun Aplikasi Pembelajaran Bangun Ruang Berbasis Augmented Reality," *Nuansa Informatika*, vol. 16, no. 2, pp. 123-134, Jul. 2022. 13
- [3] I. K. Basyar, Hernawati, S. Supriyadi, and S. Utomo, "Interactive Learning Daily Muslim Prayer Application for Children on Android at DTA Ghoniyyul Hikmah," *Nuansa Informatika*, vol. 19, no. 1, pp. 50-60, Jan. 2025. 16
- [4] I. T. Prayudha and U. Chotijah, "Pengembangan Game Edukasi 2D Mata Pelajaran IPA Menggunakan Unity Berbasis Mobile," *Jurnal Teknik Informatika Dan Komputer*, vol. 3, no. 2, pp. 46-52, 2024.
- [5] V. Oktaviany, Rasinus, and I. A. Handoyo, "Peran Game Edukasi dalam Meningkatkan Literasi Membaca pada Anak Sekolah Dasar," *Journal Scientific of Mandalika*, vol. 6, no. 5, pp. 1422-1430, 2025.
- [6] D. Berliana, I. Rusdiyani, and C. Atikah, "Game Edukasi Berbasis Canva untuk Meningkatkan Kemampuan Kognitif Anak Usia Dini," *Jurnal Obsesi : Jurnal Pendidikan Anak Usia Dini*, vol. 8, no. 1, pp. 201-210, 2024.
- [7] H. Septanto and A. K. Wardani, "Perancangan Aplikasi Game Edukasi

- Pengenalan Hewan Untuk Siswa Kelas 1 SD Berbasis Animasi Multimedia," BINA INSANI ICT JOURNAL, vol. 10, no. 2, pp. 225-234, Dec. 2023.
- [8] Gt. Irwan and L. Rusdiana, "Game Edukasi Pengenalan Hewan Endemik Pulau Kalimantan Berbasis Android Menggunakan Construct 2," Jurnal Saintekom : Sains, Teknologi, Komputer dan Manajemen, vol. 14, no. 2, pp. 118-129, Sep. 2024.
- [9] M. Syaifudin, "Implementasi Media Permainan Matching Gambar dan Kata Berbasis Power Point Untuk Pembelajaran Mufradat di SMA At-Tarbiyah Surabaya," Al-Mu'arrib: Jurnal Pendidikan Bahasa Arab, vol. 2, no. 2, pp. 128-139, Oct. 2022.
- [10] S. M. Lubis, D. Rangkuti, and D. E. S. Rangkuti, "Inovasi Permainan Tebak Bunyi Melalui Media Audio-Visual Untuk Meningkatkan Kemampuan Kognitif Anak Usia 5-6 Tahun di R.A Al-Amin Medan," Incrementapedia: Jurnal Pendidikan Anak Usia Dini, vol. 04, no. 02, pp. 101-108, Dec. 2022.
- [11] J. N. S. Adia and Supriyadi, "Pengembangan Game Edukasi Berbasis Android Tentang Pencegahan Virus Covid-19 Menggunakan Unity," Jurnal Pendidikan Teknologi Informasi (JUKANTI), vol. 5, no. 2, pp. 252-260, Nov. 2022.
- [12] C. Mauda and D. Zaliluddin, "Pengembangan Game Edukasi dengan Pendekatan MDLC: Game Zat Gizi untuk Generasi Digital," Prosiding SEMNASRISTEK UNINDRA, no. 11, pp. 509-516, 2023.
- [13] M. I. Zuhdi et al., "Implementation Multimedia Development Life Cycle (MDLC) Method in The Development of Interactive Multimedia for Introduction of Indonesian Traditional Musical Instruments," Jurnal Teknik Informatika (JUITA), vol. XI, no. 1, pp. 43-50, May 2023.
- [14] E. K. Wardani and D. Suryana, "Permainan Edukatif Setatak Angka dalam Menstimulasi Kemampuan Berfikir Simbolik Anak Usia Dini," Jurnal Obsesi : Jurnal Pendidikan Anak Usia Dini, vol. 6, no. 3, pp. 1790-1798, 2021.
- [15] E. S. Handayani and W. Wulansari, "Game Edukasi "Secret Zoo" Sebagai Media Untuk Meningkatkan Kemampuan Kognitif," Early Childhood Education and Development Journal, vol. 3, no. 2, pp. 76-81, 2021.
- [16] R. F. W. Zega, "Manfaat Penggunaan Permainan Edukatif dalam Pembelajaran Anak Usia Dini," PRESCHOOL: Jurnal Pendidikan Anak Usia Dini, vol. 4, no. 2, pp. 53-64, 2023.
- [17] A. Musrifah, I. Nursida, F. Setiawati Sulaeman, & Sutono, "Game Edukasi Bagi Anak Attention Deficit Hyperactive Disorder (ADHD) Berbasis Android," INFOTECH Journal, vol. 8, no. 2, pp. 94-100, 2022. 22
- [18] R. Octaviani and A. Aryapranata, "Games Edukasi Android dengan Metode Multimedia Development Life Cycle (MDLC)," Jurnal Esensi Komputasi IBN, vol. 3, no. 1, pp. 1-10, Mei 2019. (Meskipun tahun 2019, relevan untuk MDLC)
- [19] N. A. Humaida and Suyadi, "Pengembangan Kognitif Anak Usia Dini melalui Penggunaan Media Game Edukasi Digital Berbasis ICT," Aulad: Journal on Early Childhood, vol. 4, no. 2, pp. 78-87, 2021.
- [20] D. P. Sari, M. Ansori, and U. Radiana, "Pemanfaatan Game Edukasi Berbasis Open Source Bagi Anak Attention Deficit Hyperactivity Disorder (ADHD)," Jurnal Sistem Informasi Bisnis, vol. 01, pp. 20-25, 2016. (Relevan untuk game edukasi anak, meskipun tahunnya lebih lama, dapat diganti jika ada yang lebih baru dan relevan).

## The Implementation of User Satisfaction Based Shortest Job First Algorithm Efficiency in Queuing System

Darsanto

Wiralodra University, Indramayu

E-mail: [shantost.ft@unwir.ac.id](mailto:shantost.ft@unwir.ac.id)

### Abstract

*Efficient and structured queue management is crucial to ensure the smooth operation of various services in both the public and private sectors. A web-based queue system has now become a good solution due to their ease of access and ability to manage queues in real-time. In this study, a web-based queue system is developed that implements the Shortest Job First (SJF) algorithm to optimize service sequences. The SJF algorithm prioritizes services with the shortest processing duration, which is expected to reduce user waiting time and to increase overall system effectiveness. This research is conducted in the university's Bureau of General Administration and Finance queue system as the research object. The software development method uses Agile with the Extreme Programming (XP) framework. The results of system testing using the Black Box testing method to evaluate application functionality show that all test cases are valid, indicating that the system operates according to the desired requirements. System performance analysis is conducted by 40 students using the User Experience Questionnaire (UEQ), resulting in benchmark scores in which 5 variables are above average, and 1 variable is below average. These results indicate that the queue system built using the SJF algorithm is feasible and effective for users.*

**Keywords :** Agile Extreme Programming (XP), Queuing System, Shortest Job First (SJF) algorithm, User Experience Questionnaire (UEQ) Analysis

Submission: March 29, 2025

Accepted: May 19, 2025

Published: July 10, 2025

DOI Article : <https://doi.org/10.25134/ilkom.v19i2.378>

### 1. INTRODUCTION

The development of Technology is something that is inevitable in today's human life. Every technological innovation that is created provides various positive benefits, convenience, and becomes a new way of carrying out human daily activities [1]. Various human activities that were previously carried out manually can now be carried out digitally. One of them is queuing in certain activities. Queuing is a situation that arises when the number of people who want to use a service exceeds the existing capacity. Individuals come at unexpected and irregular times, which causes them to have to wait for a certain period of time before they can be served [2].

The University has a Bureau of General Administration and Finance (BAUK) that serves various administrative and financial aspects that support campus operations. However, the queues that occur in BAUK services at the University are less than optimal because the queue system does not use queue numbers and the service is carried out based on arrival time or the First In First Out (FIFO) algorithm. This can certainly make students with a short process time have to wait longer if the student in front of them has a long process time.

Shortest Job First (SJF) is a scheduling algorithm that runs based on the process with the shortest execution time and makes the average waiting time also short, so that implementing this algorithm can make the queue system more optimal than the previous queue system [3][4]. This study implements the Non-Preemptive SJF Algorithm on the BAUK University Queue System. The queue system application will be built web-based which can be used by users to take queue numbers online and can be served based on the queue with the shortest service process sequence first for time efficiency. The system is then tested for performance based on user satisfaction which is measured using the User Experience Questionnaire (UEQ) tool.

UEQ is a tool or questionnaire used to assess user experience of a product or application system. By using UEQ, users can express their feelings, impressions, and attitudes when assessing the system. The method of collecting responses for UEQ in this study is carried out through an online questionnaire [5]. UEQ has 26 items divided into six scales/dimensions, namely Attractiveness, Perspicuity, Efficiency, Dependability, Stimulation and Novelty [6].

In this study, a literature review is conducted by reviewing previous studies, including journals that discuss similar topics, namely regarding the

queuing system. The first study was about the creation of a payment queuing system conducted by Fithria et al. in 2019, resulting in a student queuing system with a FIFO queuing model developed using computer technology to help students find out their queue numbers and information about the queue numbers being served [7]. Another study on the design of a queuing system at an amusement ride conducted by Sari et al. in 2022, resulted in a queuing system with the FIFO method using C++ which makes it easier for ticket buyers or customers, and allows for comprehensive reporting. The disadvantage of this queuing system is that the output is still via CLI (Command Line Interface), it has not been implemented into a GUI (Graphical User Interface) [8]. Research on the design of an online queuing system for outpatient visits conducted by Melyanti et al. in 2020, resulted in an online queuing system that is capable of performing several tasks such as taking queue numbers, providing doctor schedule information, and providing daily reports of patient visits [9].

The difference between this study and previous studies lies in the methods and algorithms applied. In this study, the writer uses the Shortest Job First Non-Preemptive algorithm which is still rarely used in queuing systems and using the Agile method with the Extreme Programming development framework and the UEQ Data Analysis Tool with descriptive statistical analysis to measure the system based on user satisfaction.

## 2. RESEARCH METHODS

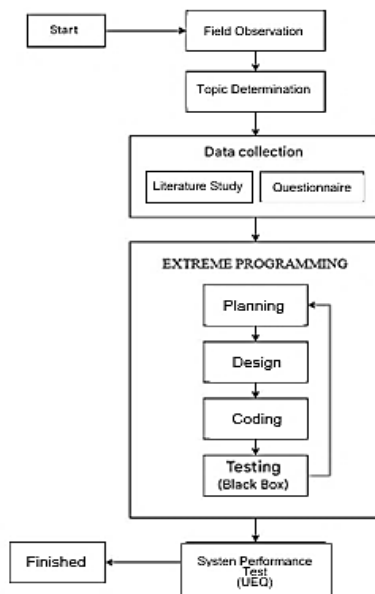


Figure 1. Research Flow

As stated in Figure 1, the research is conducted in several stages from beginning to end. Field observations are conducted through direct observation at the BAUK counter at the University regarding the ongoing queue process. Determination of the topic is based on the results of field observations, namely "Implementation of the Efficiency of the Shortest Job First Algorithm in the Queue System based on User Satisfaction". The process of collecting information data needed in the research is by conducting a literature study studying related journals and reference books.

### 2.1 Extreme Programming

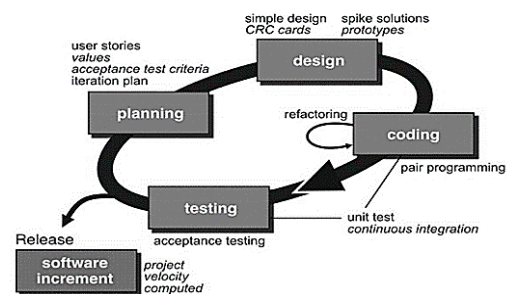


Figure 2. Extreme Programming

Extreme Programming (XP) is one of the frameworks in Agile Software Development that focuses on achieving the goals desired by developers without hindering development progress and without sacrificing software quality. The stages in the XP framework include Planning, Design, Coding, and Testing which are repeated until the development is considered complete by the user and ready to be launched [10][11][12].

#### 2.1.1 Planning

The Planning stage is the first stage of system development where planning activities are carried out, namely identifying the system processes that are running until determining the specifications for the system. The functional requirements in this stage are that the system manages the queue based on the Shortest Job First Non-Preemptive algorithm rules specified in the following table.

Table 1. SJF Algorithm Rules to be applied

| No | Queue Service Types       | Execution time (minutes) |
|----|---------------------------|--------------------------|
| 1. | Collecting the exam card  | 7                        |
| 2. | Semester tuition payments | 12                       |

| No | Queue Service Types              | Execution time (minutes) |
|----|----------------------------------|--------------------------|
| 3. | Payment of arrears               | 7                        |
| 4. | Creation of dispensation letters | 6                        |

Table 1. is the relative execution time of queue services provided by BAUK officers, based on the results of a questionnaire that has been conducted by all actors participating in the system before the queue system is created.

### 2.1.2 Design

This Design stage is carried out system modelling activities and system interface design. System modelling consists of system functional modelling, system process modelling and database modelling. System modelling uses 3 Unified Modelling Language (UML) diagrams, namely Use Case Diagram, Activity Diagram and Class Diagram [13][14].

#### a. Use Case Diagram

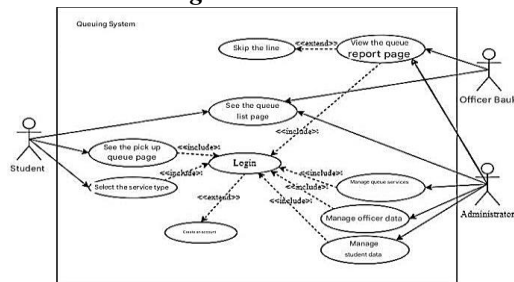
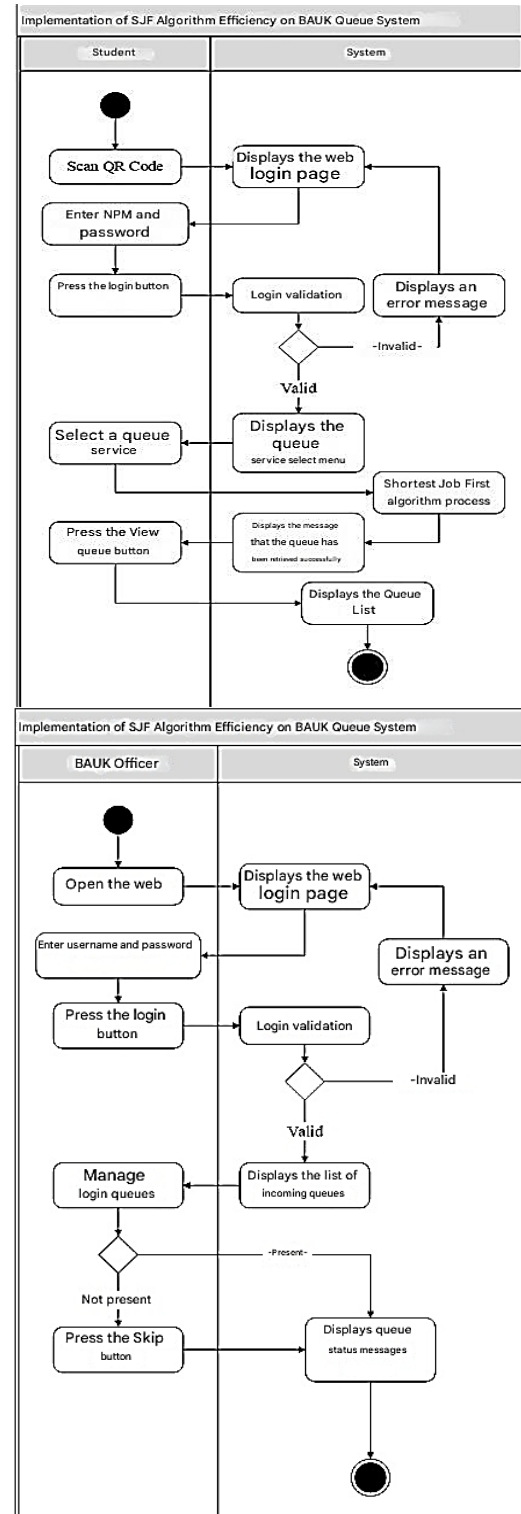


Figure 3. Use Case of Queue System

Figure 3 above is a Use Case that visualizes the relationship between actors and the queue system. There are 3 actors namely Students, BAUK Officers and Administrators.

#### b. Activity Diagram



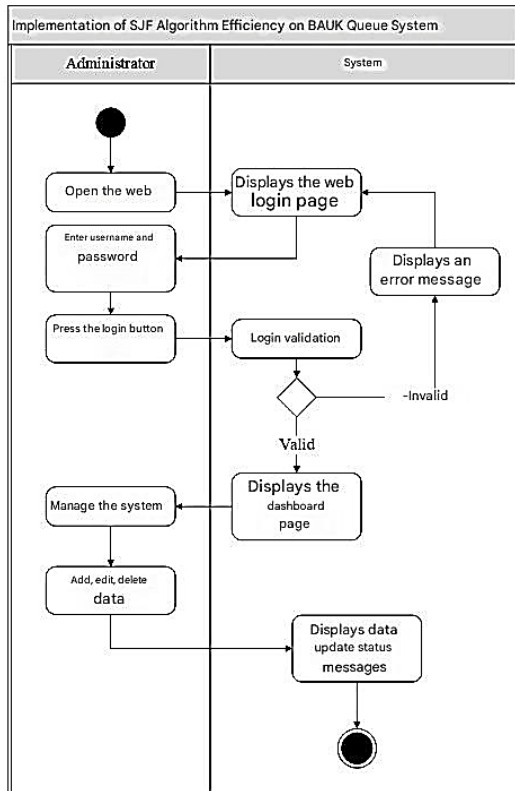


Figure 4. Activity Diagram of Queue System

Figure 4 is a visualization of the sequence of process activities or workflows carried out by students, officers and administrators with the system.

### c. Class Diagram

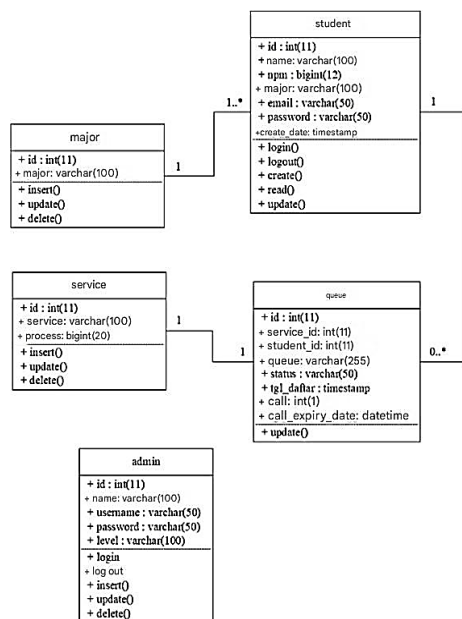


Figure 5. Class Diagram of Queue System

Figure 5 above is a class diagram that shows the classes, attributes, accessible methods and relationships between classes in the queue system that was built.

### 2.1.3 Coding

The Coding stage is the implementation stage of the model design that has been made into the form of an interface. System coding is built using PHP, MYSQL, and Visual Studio Code. Each system coding is used as a programming language, database management and code editor application.

### 2.1.4 Testing

Testing Phase or Testing uses Black Box to test the system that has been built by identifying the types of errors that occur during the application running and ensuring that the system is in accordance with its specifications. The method used in Black Box Testing is validation.

### 2.1.5 System Performance Test

This stage uses the UEQ tool by conducting descriptive statistical analysis. UEQ has 26 items divided into six dimensions which can be seen in the following figure [15][16].

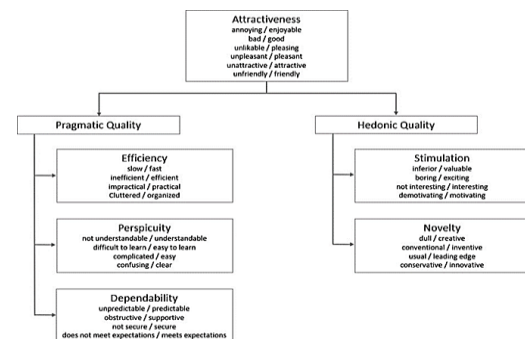


Figure 6. Assumed scale structure of the UEQ.

Each item in the Pragmatic and Hedonic Quality domains as shown in Figure 6 is represented by two terms with opposite meanings. Half of the items in the scale start with a positive value and the other half start with a negative value. The system performance test was conducted by 40 students and the respondents were collected through an online questionnaire using Google Form. The data results obtained from the questionnaire were then analyzed using the UEQ Data Analysis Tool version 12 system that has been built by identifying the types of errors that occur during application running and

ensuring that the system is in accordance with its specifications. The method used in Black Box Testing is validation.

|                   | 1                     | 2                     | 3                     | 4                     | 5                     | 6                     | 7                     |                           |
|-------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|-----------------------|---------------------------|
| troublesome       | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | pleasant                  |
| creative          | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | monotonous                |
| beneficial        | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | less useful               |
| not attractive    | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | interesting               |
| fast              | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | slow                      |
| obstruct          | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | support                   |
| complicated       | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | simple                    |
| common            | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | leading edge              |
| safe              | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | not safe                  |
| meet expectations | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | did not meet expectations |
| clear             | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | confusing                 |
| organized         | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | untidy                    |
| user friendly     | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> | not user friendly         |
|                   |                       |                       |                       |                       |                       |                       |                       | innovative                |

Figure 7. Version english of the UEQ

Based on Figure 7, respondents expressed their feelings while using the app by selecting a score between 1 to 7, which most closely approximates their experience while using the system.

### 3. RESEARCH RESULTS

#### 3.1 System Interface Results



Figure 8. Website Page View

Figures 8, 9, 10 show the results of system coding using programming languages, database management and code editors, namely PHP, MYSQL and Visual Studio Code.



Figure 9. Take Queue Page View

| List Antrean                          |              |                 |                        |                         |           |
|---------------------------------------|--------------|-----------------|------------------------|-------------------------|-----------|
| Berikut Antrean yang sedang berjalan: |              |                 |                        |                         |           |
| No                                    | Npm          | Nama            | Jurusan                | Layanan                 | Status    |
| 1                                     | 132010110039 | Intan Savira    | Kesehatan Masyarakat   | Pengambilan Kartu Ujian | Proses... |
| 2                                     | 862080110005 | Egri Nurhasanah | Pendidikan Agama Islam | Pengambilan Kartu Ujian | Menunggu  |
| 3                                     | 502020110012 | Riana Novanty   | Teknik Komputer        | Pembayaran Tunggal      | Menunggu  |

Figure 10. Queue List Page View

Figures 8, 9, 10 show the results of system coding using programming languages, database management and code editors namely PHP, MYSQL and Visual Studio Code.

#### 3.2 System Testing

Table 2. Black Box Testing Results

| No | Testing              | Nama Tes                        | Status |
|----|----------------------|---------------------------------|--------|
| 1. | Account Registration | Mendaftar akun                  | Valid  |
| 2. | Log in               | Log in to the queue system page | Valid  |
| 3. | Queue List           | View queue lists                | Valid  |
|    |                      | Search for queue lists          | Valid  |
|    |                      | Skip the queue                  | Valid  |
| 4. | Take Queue           | Retrieve a queue                | Valid  |
|    |                      | Choose a queue service type     | Valid  |
|    |                      | View queues                     | Valid  |
| 5. | Log out              | Exit the queue system page      | Valid  |

| No  | Testing         | Nama Tes                        | Status |
|-----|-----------------|---------------------------------|--------|
| 6.  | History Antrean | View queue history              | Valid  |
|     |                 | Search for queue data           | Valid  |
| 7.  | Prodi Data      | View the list of study programs | Valid  |
|     |                 | Adding study programs           | Valid  |
|     |                 | Editing study programs          | Valid  |
|     |                 | Delete a study program          | Valid  |
| 8.  | Service Data    | View queue services             | Valid  |
|     |                 | Add queue services              | Valid  |
|     |                 | Edit queue services             | Valid  |
|     |                 | Delete a queue service          | Valid  |
| 9.  | Officer's Data  | View an officer's account       | Valid  |
|     |                 | Add an officer account          | Valid  |
|     |                 | Edit an officer account         | Valid  |
|     |                 | Delete an officer account       | Valid  |
| 10. | Student Data    | View student account data       | Valid  |
|     |                 |                                 | Valid  |

| No | Testing | Nama Tes                    | Status |
|----|---------|-----------------------------|--------|
|    |         | Add student account data    | Valid  |
|    |         | Edit student account data   | Valid  |
|    |         | Delete student account data | Valid  |

## 4. DISCUSSION

### 4.1 UEQ System Test

The results of the data obtained were analyzed using the UEQ Data Analysis Tool version 12. The analysis uses the mean (average) value of each dimension item and questionnaire statement. The results of the data from the respondents' answers taken and have been calculated on average, grouped in each dimension, namely attractiveness, clarity, efficiency, accuracy, stimulation and novelty. Then it is calculated whether the respondents are satisfied or dissatisfied with the tested system based on the average range value  $<0.8$  indicating a negative evaluation,  $-0.8$  and  $0.8$  indicating a neutral evaluation and  $>0.8$  indicating a positive evaluation. The following are the results of the average value of each dimension.

| Item | Mean | Left          | Right             | Evaluation |
|------|------|---------------|-------------------|------------|
| 1    | 1.3  | troublesome   | pleasant          | Positive   |
| 12   | 1.9  | Good          | bad               | Positive   |
| 14   | 1.0  | disliked      | encouraging       | Positive   |
| 16   | 1.4  | uncomfortable | comfortable       | Positive   |
| 24   | 1.4  | attractive    | not attractive    | Positive   |
| 25   | 1.7  | user friendly | not user friendly | Positive   |

Figure 11. Dimensions of Attractiveness

The Appeal dimension measures the extent to which the user's general impression of the app's design and visual appearance. The analysis results listed in figure 11 above, show that all items from the attractiveness dimension received positive evaluations.

| Item | Mean | Left             | Right              | Evaluation |
|------|------|------------------|--------------------|------------|
| 2    | 1.4  | incomprehensible | understandable     | Positive   |
| 4    | 1.8  | easy to learn    | difficult to learn | Positive   |
| 13   | 1.7  | complicated      | simple             | Positive   |
| 21   | 1.7  | clear            | confusing          | Positive   |

Figure 12. Dimension of Perspicuity

The Clarity dimension measures how well users understand the function and purpose of the application. The analysis results listed in Figure 12, show that all items in the clarity dimension received positive evaluations.

| Item | Mean | Left          | Right     | Evaluation |
|------|------|---------------|-----------|------------|
| 9    | 1,2  | fast          | slow      | Positive   |
| 20   | 1,3  | not efficient | efficient | Positive   |
| 22   | 1,5  | not practical | practical | Positive   |
| 23   | 1,4  | organized     | untidy    | Positive   |

Figure 13. Dimension of Efficiency

The Efficiency dimension measures how far users feel that the application or system allows them to complete tasks quickly and effectively. The analysis results listed in Figure 13, show that all items in the efficiency dimension received positive evaluations.

| Item | Mean | Left              | Right                     | Evaluation |
|------|------|-------------------|---------------------------|------------|
| 8    | 1,0  | unpredictable     | predictable               | Positive   |
| 11   | 1,7  | obstruct          | support                   | Positive   |
| 17   | 1,5  | safe              | not safe                  | Positive   |
| 19   | 1,4  | meet expectations | did not meet expectations | Positive   |

Figure 14. Dimension of Dependability

The Dependability dimension assesses how far the user feels that the app is reliable for troubleshooting a given problem or task. The results of the analysis listed in figure 14, show that all *items* in the accuracy dimension obtained a positive evaluation value.

| Item | Mean | Left           | Right          | Evaluation |
|------|------|----------------|----------------|------------|
| 5    | 2,0  | beneficial     | less useful    | Positive   |
| 6    | 0,6  | boring         | exciting       | Neutral    |
| 7    | 1,2  | not attractive | interesting    | Positive   |
| 18   | 1,3  | motivating     | not motivating | Positive   |

Figure 15. Dimensions of Stimulation

The Stimulation dimension assesses how far an app or system provides an engaging and motivating experience for users. The results of the analysis listed in figure 15, show that *item* 6 obtained a neutral evaluation, while the other item in the stimulation dimension obtained a positive evaluation.

| Item | Mean | Left         | Right        | Evaluation |
|------|------|--------------|--------------|------------|
| 3    | 0,6  | creative     | monotonous   | Neutral    |
| 10   | 0,2  | inventive    | conventional | Neutral    |
| 15   | 0,5  | common       | leading edge | Neutral    |
| 26   | 1,4  | conservative | innovative   | Positive   |

Figure 16. Dimension of Novelty

The Novelty Dimension assesses how far the level of novelty and uniqueness of the

features and functionality of a system or application is according to the user. The results of the analysis listed in figure 16, only item 26 produced a positive evaluation while other *items* from the novelty dimension produced a neutral evaluation.

## 4.2 UEQ System Calculation

The results of the UEQ calculation based on the grouping of aspects shown in figure 17, the Attractiveness aspect produces a value of 1.42 with a positive evaluation which means that the use of the system is attractive to users. Pragmatic Quality resulted in a mean value of 1.46 with a positive evaluation, which means that the system meets the needs and expectations of users in the context of functionality and practical usability. The quality of Hedonis resulted in a mean value of 0.96 with a positive evaluation which means that the system meets the emotional needs and personal aspirations of the user.

| Pragmatic and Hedonic Quality |      |
|-------------------------------|------|
| Attractiveness                | 1.42 |
| Pragmatic Quality             | 1.46 |
| Hedonic Quality               | 0.96 |

Figure 17. UEQ Scale Based on Aspect Groups

| Scale      | Mean | Comparison to benchmark | Interpretation                              |
|------------|------|-------------------------|---------------------------------------------|
| Daya tarik | 1,42 | Above average           | 25% of results better, 50% of results worse |
| Kejelasan  | 1,63 | Above Average           | 25% of results better, 50% of results worse |
| Efisiensi  | 1,37 | Above Average           | 25% of results better, 50% of results worse |
| Ketepatan  | 1,39 | Above Average           | 25% of results better, 50% of results worse |
| Stimulasi  | 1,27 | Above Average           | 25% of results better, 50% of results worse |
| Kebaruan   | 0,66 | Below Average           | 50% of results better, 25% of results worse |

Figure 18. Benchmark Scale Comparison

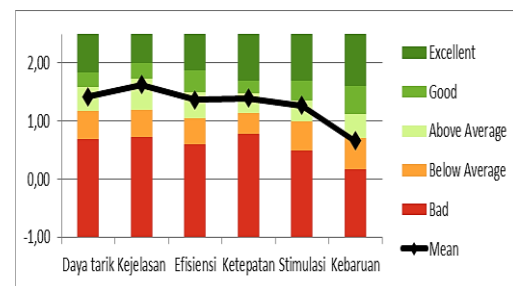


Figure 19. Benchmark Scale Graph

Figure 18 shows the results of the system performance test obtained a benchmark result of Above Average for a mean value of 1.42 in the dimension of attraction, 1.63 in clarity, 1.37 in efficiency, 1.39 in accuracy and 1.27 in stimulation with the interpretation of 25% of the results being better, while 50% of the results being worse. Meanwhile, the benchmark results

of Below Average (below average) for the mean value of 0.66 in the novelty dimension with the interpretation of 50% of the results being better, while 25% of the results are worse. That way, based on UEQ analysis, it can be stated that the BAUK Queue System Website using the Shortest Job First algorithm is quite feasible and effective to be applied.

## 5. CONCLUSION

1. The queue system application using the *Shortest Job First Non-Preemptive* algorithm is able to provide an effective solution in managing service queues so that time becomes more efficient and prevents queue buildup or difficulty in queue taking. However, there are still shortcomings in the application of the *Shortest Job First algorithm* in the system.
2. System testing uses Black Box Testing to evaluate the functionality and output of the application regardless of its internal structure or source code, resulting in the valid status of each test case performed.
3. The analysis of the system performance test uses the User Experience Questionnaire (UEQ) descriptive statistical method, reviewed based on the responses of the respondents who had tested the system, resulted in a score of 1.42 on the Attractiveness dimension, 1.63 on Clarity, 1.37 on Efficiency, 1.39 on Accuracy, 1.27 on Stimulation and 0.66 on Novelty. Five dimensions obtained a benchmark value of Above Average and one dimension of Below Average. The mean value in the aspect of Attraction is 1.42, Pragmatic Quality was 1.46 and Hedonistic Quality was 0.96. So, from these results, the system can be declared quite feasible and effective to be used by users.

## 6. SUGGESTIONS

As an effort to improve the BAUK Queue System at the University in the future, it is recommended to:

1. The system needs to be repaired again by paying attention to the shortcomings and possible problems or *human errors* that occur so that it can cause the system to be less effective to use.

2. The system should be developed using Artificial Intelligence to make it more optimal and add other features to the system as needed that may change over time.
3. It is recommended to conduct a thorough *White Box* test of the system to ensure accurate system code and logic functionality, as well as to identify potential security vulnerabilities and hidden bugs that were not detected during *Black Box testing*.

## REFERENCES

- [1] M. Ngafifi, "Kemajuan Teknologi Dan Pola Hidup Manusia Dalam Perspektif Sosial Budaya," *Jurnal Pembangunan Pendidikan: Fondasi dan Aplikasi*, vol. 2, no. 1, Jun. 2014, doi: 10.21831/jppfa.v2i1.2616.
- [2] T. J. Kakiay, "Dasar teori antrian untuk kehidupan nyata," *Yogyakarta: Andi*, 2004.
- [3] N. Mariana *et al.*, "Pengembangan Sistem Layanan Perawatan Pada Klinik ABC."
- [4] K. Sayood, *Introduction to data compression*. Morgan Kaufmann, 2017.
- [5] R. Riche and S. H. Marpaung, "Evaluasi Pengalaman Pengguna dengan Menggunakan User Experience Questionnaire Perpustakaan Digital," *Jurnal Media Informatika Budidarma*, vol. 5, no. 4, p. 1345, Oct. 2021, doi: 10.30865/mib.v5i4.3270.
- [6] Darsanto and M. K. Maulidani, "Analisis User Experience Aplikasi Regsosek Pada Badan Pusat Statistik Indramayu Menggunakan Metode User Experience Questionnaire," *Nuansa Informatika*, vol. 17, Jul. 2023.
- [7] M. Fithria, B. Sembiring, and L. Sitorus, "Perancangan Implementasi Socket Programming dalam Pembuatan Sistem Antrian Pembayaran di Unika dengan 27 Metode FIFO. Oleh : Mayshela Fithria Br.Sembiring, Lamhot Sitorus Implementasi Socket Programming dalam Pembuatan Sistem Antrian Pembayaran di Unika dengan Metode Fifo Article Information Abstrak," 2019.
- [8] I. P. Sari, I. H. Batubara, F. Ramadhani, and S. Wardani, "Perancangan Sistem Antrian pada Wahana Hiburan dengan Metode First In First Out (FIFO)," *sudo Jurnal Teknik*

*Informatika*, vol. 1, no. 3, pp. 116–123, Jul. 2022, doi: 10.56211/sudo.v1i3.93.

- [9] R. Melyanti, D. Irfan, A. Febriani, R. Khairana, and S. Hang Tuah Pekanbaru, “Rancang Bangun Sistem Antrian Online Kunjungan Pasien Rawat Jalan Pada Rumah Sakit Syafira Berbasis Web Design Of Online Queue System For Web-Based Visit Of Patients In Syafira Hospital,” *Journal of Information Technology and Computer Science (IntecomS)*, vol. 3, no. 2, 2020.
- [10] M. I. Burhan, F. Nawir, and K. N. Salam, “Pengembangan Sistem Tracer Study Menggunakan Agile Development Methods Pada IbK Nitro,” *Jursima (Jurnal Sistem Informasi dan Manajemen)*, vol. 10, no. 3, pp. 160–170, 2022.
- [11] M. R. Nawawi, S. Lestanti, and D. Fanny, “Rancang Bangun Sistem Informasi Inventaris Fasilitas Pondok Pesantren Nurul Ulum Dengan Menggunakan Metode Xp (Extreme Programming),” *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 6, no. 2, pp. 835–841, 2022.
- [12] R. I. Borman, A. T. Priandika, and A. R. Edison, “Implementasi metode pengembangan sistem Extreme Programming (XP) pada aplikasi investasi peternakan,” *JUSTIN (Jurnal Sistem Dan Teknologi Informasi)*, vol. 8, no. 3, pp. 272–277, 2020.
- [13] M. Sumiati, R. Abdillah, and A. Cahyo, “Pemodelan Uml Untuk Sistem Informasi Persewaan Alat Pesta,” *Jurnal Fasilkom*, vol. 11, no. 2, pp. 79–86, 2021.
- [14] T. A. Kurniawan, “Pemodelan use case (UML): evaluasi terhadap beberapa kesalahan dalam praktik,” *J. Teknol. Inf. dan Ilmu Komput*, vol. 5, no. 1, p. 77, 2018.
- [15] M. Schrepp, *User Experience Questionnaire Handbook*, 10th ed. 2023.
- [16] Y. Wijayanti, S. Suyoto, and A. T. Hidayat, “Evaluasi pengalaman pengguna pada aplikasi seluler Visiting Jogja menggunakan metode User Experience Questionnaire (UEQ),” *Jurnal Janitra Informatika dan Sistem Informasi*, vol. 3, no. 1, pp. 10–17, 2023.

## Predicting Basic Shipping Tariff Using Machine Learning

Nisa Hanum Harani<sup>1</sup>, M. Yusril Helmi Setyawan<sup>2</sup>, Dani Ferdinan<sup>\*3</sup>

<sup>1,2</sup>, Universitas Logistik dan Bisnis Internasional, Bandung

E-mail: <sup>1</sup>[nisa@ulbi.ac.id](mailto:nisa@ulbi.ac.id), <sup>2</sup>[yusrilhelmi@ulbi.ac.id](mailto:yusrilhelmi@ulbi.ac.id), <sup>\*3</sup>[daniferdinandall@gmail.com](mailto:daniferdinandall@gmail.com)

### Abstract

*This study explores the application of machine learning algorithms in predicting the Basic Shipping Tariff for logistics, focusing on variables such as Item Price, Shipment Weight, and Distance (KM). Random Forest Regressor and Linear Regression models were used as comparison methods. Experimental results show that the Random Forest Regressor outperforms Linear Regression, achieving an  $R^2$  value of 0.915 and RMSE of 0.154, while Linear Regression reached an  $R^2$  value of 0.706 and RMSE of 0.113. Additionally, the Random Forest model achieved lower error values with MSE of 0.000 and MAE of 0.003, compared to Linear Regression with MSE of 0.001 and MAE of 0.007. These error metrics further highlight the superiority of the Random Forest model. In-depth analysis reveals significant relationships between these variables and the Basic Shipping Tariff, showcasing the model's potential application in dynamic pricing strategies within the Indonesian logistics industry. This study aims to contribute to operational efficiency and improve pricing accuracy in the logistics business in Indonesia.*

**Keywords**— Machine Learning, Random Forest Regressor, Linear Regression, Basic Shipping Tariff, Logistics

Submission: May 07, 2025

Accepted: June 14, 2025

Published: July 10, 2025

DOI Article : <https://doi.org/10.25134/ilkom.v19i2.388>

### 1. INTRODUCTION

The logistics industry in Indonesia plays a vital role in supporting national economic growth. In recent years, this sector has undergone significant transformation, particularly due to the increasing volume of e-commerce transactions, which has driven the demand for fast and efficient delivery services. However, complex challenges such as the efficient management of operational costs and the determination of optimal shipping rates have become increasingly urgent. In this context, logistics companies are faced with the need to respond to the ever-evolving market dynamics without compromising operational cost efficiency [1].

One of the main issues faced by the logistics industry is the method of determining shipping rates, which is still often conducted manually or based on static policies. The Theory of Pricing in Logistics refers to the theoretical and practical framework used by logistics companies to determine the cost of their services. This approach involves a comprehensive analysis of operational costs, market demand, and competitive dynamics to ensure that the prices set not only cover all expenses but also maximize profit margins and market competitiveness [2].

Pricing models in logistics also take into account external factors, such as government regulations and economic fluctuations, which can significantly influence pricing structures. In

a study conducted by [3], pricing was simulated using a game theory approach, which serves as a crucial element in shaping operational decisions. The game theory approach involves a thorough analysis of market demand and competitive dynamics to ensure that the prices set maximize profit margins [4].

Challenges in logistics cost management have also been highlighted, wherein efforts to reduce costs in one area may lead to increased costs in another [5]. Therefore, logistics companies need to adopt pricing strategies that are adaptive and responsive to market fluctuations.

This approach is not fully responsive to changes in demand or other variables that affect costs, such as item weight, item value, and delivery distance. Inefficiencies in rate setting can result in significant negative impacts, including a decline in the competitiveness of logistics companies in an increasingly competitive market.

In addressing these challenges, machine learning technology offers a promising solution [6]. Machine learning is a branch of artificial intelligence that enables systems to automatically learn from data and improve from experience without being explicitly programmed [7]. It focuses on building predictive models that identify patterns and relationships between variables, allowing future outcomes to be estimated based on past data [8]. Predictive modeling, as a core application of machine

learning, involves training algorithms on historical datasets to forecast values such as shipping costs, demand, or delivery times [9]. By leveraging predictive algorithms, companies can optimally analyze and utilize data to identify patterns and relationships among complex variables. Previous studies, such as those conducted on Alibaba's Cainiao system, have demonstrated that the implementation of machine learning can significantly enhance pricing efficiency. While such international implementations have proven effective, empirical machine learning-based models using real-world Indonesian logistics data remain underrepresented. Most existing research still focuses on global platforms like Alibaba [10], leaving a gap in localized, data-driven pricing models tailored to the operational context of Indonesian logistics companies. This highlights the substantial potential of this technology in overcoming the issues faced by the logistics industry.

This study focuses on the application of the Random Forest Regressor and Linear Regression algorithms to develop a model for predicting the base rate of logistics shipping. Random Forest is an ensemble algorithm composed of multiple decision trees that collectively work to produce more accurate predictions. Its main advantage lies in its ability to capture non-linear relationships within the data, making it particularly suitable for handling complex datasets [11]. The method employs a bagging technique, where multiple trees are trained on different subsets of the data to reduce the risk of overfitting. The average prediction from all trees is then used to generate a more stable and accurate final result.

On the other hand, Linear Regression is a simpler statistical method that aims to represent the relationship between one or more independent variables and a dependent variable through a linear equation [12]. Although this method is easier to interpret, its performance tends to decline when dealing with non-linear relationships or data involving complex interactions among variables. Therefore, it is important to compare these two methods in the context of shipping rate determination.

In evaluating model performance, the key metrics used include  $R^2$  (R-squared), MSE (Mean Squared Error), RMSE (Root Mean Squared Error), and MAE (Mean Absolute Error).  $R^2$  measures the extent to which the model explains the variance in the data, while MSE and RMSE assess the level of prediction error, with RMSE offering a more intuitive interpretation in the original units of the data.

MAE, on the other hand, indicates the average absolute error [13], providing a direct measure of model accuracy.

By utilizing historical data from PT Pos Indonesia, this study is expected to offer a more adaptive approach to improving pricing prediction accuracy, while also providing relevant insights for the Indonesian logistics sector in addressing the challenges of global competition.

## 2. RESEARCH METHODS

This study adopts a quantitative approach grounded in positivist principles, aiming for objective and empirically testable results through numerical data analysis [14]. The process follows the CRISP-DM (Cross-Industry Standard Process for Data Mining) methodology, which includes six standardized phases: Business Understanding, Data Understanding, Data Preparation, Modeling, Evaluation, and Deployment [15]. Developed collaboratively by Daimler Chrysler, SPSS, and NCR, CRISP-DM has been widely used since its official release in 1999 as a robust framework for solving data-driven problems across industries.

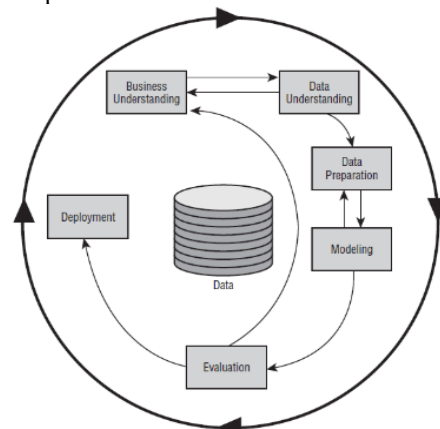


Figure 1. Research Methods

### 2.1. Business Understanding

The objective of this study is to improve the accuracy of base rate predictions in logistics by utilizing historical data analysis from PT Pos Indonesia. It is expected that the findings of this research will support the company in automating its pricing strategy.

### 2.2. Data Understanding

The PT Pos shipping dataset also includes several additional columns that provide detailed information related to the shipping process. The “Nomor Resi” serves as a unique

identifier for each package shipped. The “ExtID” record an additional identification extension for company-specific purposes. “Kota Penerima” indicates the destination location of the recipient, while “Kantor Asal” and “Kantor Tujuan” record the origin and final destination offices of the shipment. “Bea Dasar” refers to the total shipping cost calculated based on various factors. “Harga Barang” reflects the declared value of the item being shipped. “Total” refers to the sum of the Bea Dasar (base shipping cost) and the Harga Barang (price of goods), and “Berat Kiriman” records the total weight of the package. The “Produk” column categorizes the type of service or shipping product used. “Tanggal Kirim” notes the date the shipment was dispatched, and “Isi Kiriman” provides a description of the contents or type of goods being delivered.

### 2.3. Data Preparation

Data Preparation is a crucial initial stage in this study to ensure that the data to be analyzed is clean, standardized, and ready for use [16], [17]. The initial dataset consists of 3,446 entries containing information on Harga Barang, Berat Kiriman, and Bea Dasar. However, several entries were incomplete due to missing values. Records with missing data in critical variables such as Harga Barang or Berat Kiriman were removed to ensure the quality of the dataset. After this cleaning process, a total of 3,433 valid entries remained and were deemed suitable for modeling.

The next step involves calculating the distance between the sender and recipient cities. Coordinates for each city were obtained by aligning the city names in the dataset with coordinate data retrieved from a public repository available at <https://github.com/yusufsyarifudin/wilayah-indonesia>. This repository provides geographic coordinates (latitude and longitude) for various regions in Indonesia, enabling accurate mapping of each city in the dataset to its corresponding coordinates.

Once the coordinates were obtained, the distance between the sender and recipient cities was calculated using the Python library Geopy, a tool designed for working with geographical data. The calculated distance values were then added to the dataset under the column named “Jarak (KM)”.

### 2.4. Model Development

The implementation phase of the prediction algorithm was carried out using two

machine learning models: Random Forest Regressor and Linear Regression. These algorithms were selected due to their distinct characteristics, offering valuable insights into the effectiveness of ensemble-based approaches versus simple linear regression in predicting the Bea Dasar variable.

The implementation process began with splitting the data into two groups: training data (80%) and testing data (20%). This data partitioning was performed using the `train_test_split` function from the Scikit-learn library, which randomly divides the dataset into two subsets according to the predefined proportion. This process aims to ensure that the model is trained on a sufficiently large subset to prevent overfitting, while also maintaining a representative testing set for performance evaluation.

#### 2.4.1. Linear Regression

Linear regression is a fundamental technique in statistical modeling used to predict quantitative values based on one or more predictor variables. This approach is based on the assumption that there is a linear association between the dependent variable and the independent variables. Simple linear regression involves a single predictor variable, whereas multiple linear regression includes more than one. Recent studies [18] have shown that linear regression models can be enhanced by integrating them with modern computational techniques, thereby improving their predictive performance and expanding their applications in various fields such as economics, healthcare, and more. These advancements have also led to more robust error estimation and model assessment in real-world applications. In statistics, linear regression analysis is a method for identifying causal relationships between variables. It is a widely used approach across multiple disciplines. Both regression and correlation are employed to measure and understand relationships between variables. In regression analysis, estimation equations are developed to model the patterns or functions that describe the relationships among variables.

Linear regression is divided into two types: Simple Linear Regression and Multiple Linear Regression. Simple Linear Regression is used to examine the linear relationship between two variables, namely the independent variable (X), which can be controlled, and the dependent variable (Y), which reflects the response to the independent variable.

The equation of Y in relation to X in linear regression is presented in Equation (1).

$$\hat{Y} = a + bX \quad (1)$$

In linear regression,  $\hat{Y}$  represents the predicted value of the dependent variable based on the changes or influence of the independent variable.  $a$  is the value of Y when X is zero; it is the constant term, also known as the intercept, indicating the point where the regression line crosses the Y-axis. This reflects the initial value of Y in the absence of any influence from X.  $b$  is the slope of the regression line, or the regression coefficient, which indicates the magnitude of change in the dependent variable in response to changes in the independent variable. A positive (+) coefficient results in an upward-sloping regression line, while a negative (−) coefficient produces a downward slope.  $X$  refers to a specific value of the independent variable, which can be controlled or set in order to observe its effect on the dependent variable.

Technically, the value of  $b$  (or the regression coefficient) is the tangent of the angle of the regression line's slope. It represents the ratio or rate of change in the dependent variable relative to the change in the independent variable, once the regression equation has been established [19].

The next type of linear regression is Multiple Linear Regression, which involves a linear relationship between the dependent variable (Y) while maintaining the property of linearity. Researchers use this method to predict changes in the dependent variable (criterion) based on two or more independent variables that serve as predictor factors. Multiple regression analysis is applied when there are at least two independent variables involved [20].

The general form of the multiple linear regression equation is presented in Equation (2).

$$Y = a + b_1X_1 + b_2X_2 + b_3X_3 + \dots + b_nX_n + e \quad (2)$$

In this context,  $Y$  represents the response or dependent variable,  $a$  is the constant (intercept),  $e$  denotes the error term (disturbance),  $b_1, b_2, \dots, b_n$  are the regression coefficients for the independent variables, and  $X_1, X_2, \dots, X_n$  are the predictor variables (independent variables).

If the dependent variable is associated with two independent variables, the multiple linear regression equation can be expressed as shown in Equation (3).

$$Y = a + b_1X_1 + b_2X_2 \quad (3)$$

Let  $Y$  be the dependent variable, while  $X_1$  and  $X_2$  are the independent variables. The parameter  $a$  represents the value of Y when both  $X_1$  and  $X_2$  are equal to zero. The coefficients  $b_1$  and  $b_2$  are parameters in the multiple linear regression model, where  $b_1$  indicates the change in Y, in specific units, when  $X_1$  increases or decreases by one unit while  $X_2$  is held constant. Similarly,  $b_2$  represents the change in Y, in specific units, when  $X_2$  increases or decreases by one unit, assuming  $X_1$  remains constant.

#### 2.4.2. Random Forest Regressor

Forest Regression is an ensemble-based machine learning method used to predict numerical values of a target variable. This method is designed to address the limitations of a single decision tree model, such as overfitting, by constructing an ensemble of decision trees [21]. In this approach, the model operates by training multiple trees on randomly selected subsets of the data using a technique known as bagging (bootstrap aggregating) [22]. This process not only helps reduce model variance but also enhances generalization, allowing the model to make more accurate predictions on previously unseen data.

Each tree in the forest generates a prediction based on the given input, and the final prediction is obtained by averaging all the predictions made by the individual trees in the model. This method helps improve the overall stability and accuracy of the model [23]. In addition, Random Forest has the capability to capture non-linear relationships within the data, making it highly effective for handling complex datasets with numerous variables.

One of the key advantages of Random Forest is its ability to provide estimates of feature importance, allowing users to identify which variables have the most influence on the predictions. This is particularly valuable in the context of data analysis, where understanding the factors that affect outcomes can offer critical insights for informed decision-making [24].

This model was implemented using the Random Forest Regressor class from Scikit-learn, a widely used Python library for machine learning. Scikit-learn provides a variety of tools and functions that facilitate the building, training, and evaluation of Random Forest models. Through this library, users can easily configure model parameters such as the number of trees in the forest ( $n\_estimators$ ), the maximum depth of the trees ( $max\_depth$ ), and the splitting criterion

(criterion) to optimize model performance according to their specific needs.

Overall, Random Forest Regression is a powerful and flexible method for data analysis, particularly in the context of predicting numerical variables. With its ability to overcome overfitting, capture non-linear relationships, and provide insights into feature importance, this method is an excellent choice for various applications, including in the logistics industry, where accurate shipping rate predictions are essential to improving operational efficiency and enhancing competitive advantage.

## 2.5. Evaluation

The evaluation method is a critical step in this study to assess the performance of the prediction model on previously unseen data, namely the testing data [25]. It not only measures the model's ability to predict values on the training data but also evaluates its capacity to deliver accurate and reliable results when applied to new, unseen data.

The evaluation was conducted using several model performance metrics, namely  $R^2$  (R-squared), MSE (Mean Squared Error), RMSE (Root Mean Squared Error), and MAE (Mean Absolute Error). These metrics were chosen because each provides a different perspective on how well the model predicts the value of basic shipping.

$R^2$  (R-squared): This metric measures the extent to which the model can explain the variance present in the data.  $R^2$  values range from 0 to 1, where a higher value indicates that the model explains a greater proportion of the data variability. The closer the value is to 1, the better the model's ability to represent the existing variation.

RMSE (Root Mean Squared Error): RMSE calculates the average magnitude of prediction errors, expressed in the same unit as the target variable. It reflects how far, on average, the model's predictions deviate from the actual values. A lower RMSE indicates better model performance in predicting the value of Bea Dasar.

MSE (Mean Squared Error): MSE is the average of the squared differences between predicted and actual values. While it also provides insight into the magnitude of error, it is more sensitive to outliers compared to RMSE. A lower MSE suggests higher accuracy in the model's predictions.

MAE (Mean Absolute Error): MAE measures the average absolute difference between predicted and actual values. This metric

offers a straightforward interpretation of model accuracy and is generally easier to understand compared to MSE and RMSE. A lower MAE indicates that the model has a smaller prediction error on average.

By employing this combination of metrics, the study provides a comprehensive analysis of the model's performance. This is essential to ensure that the resulting model is not only highly accurate but also stable and reliable when applied in real-world scenarios, such as shipping rate determination in the logistics industry. Such thorough evaluation is expected to offer deeper insights into the model's effectiveness and support more informed decision-making in the future.

## 2.6. Deployment

The analysis of results provides recommendations for integrating the model into the company's automated pricing system. This integration aims to support more accurate, data-driven tariff decisions, thereby enhancing operational efficiency and competitiveness within the logistics industry.

# 3. RESULTS AND DISCUSSION

## 3.1. Data Preprocessing

In the modeling process, distance data is required. This distance data was obtained by first retrieving the coordinates of the destination city, then calculating the distance from the origin city based on these coordinates. Once the distance data was obtained, data preprocessing was carried out.

To obtain the distance data, the coordinates were first identified by matching the city names in the dataset with a public coordinate dataset available on GitHub, accessible via the following link:  
<https://github.com/yusufsyafudin/wilayah-indonesia>.

Subsequently, the distance between the origin and destination cities was calculated using the Python library Geopy, based on their geographic coordinates. Once the distance data was generated, data preprocessing could be initiated.

Following an analysis of variable relationships, the researcher selected Jarak (KM), Harga Barang, and Berat Kiriman as independent variables that influence the dependent variable, Bea Dasar. After selecting the variables to be used, missing values in the

three selected variables were removed. Out of the total 3,446 data entries, 3,433 remained free of missing values and were used in the modeling process.

### 3.2. Model Development

Once the dataset was cleaned, the modeling process could be conducted. In this phase, two models were employed: Linear Regression and Random Forest Regression.

The first step involved separating the independent features, namely Berat Kiriman, Jarak (KM), and Harga Barang. The data was then split into training and testing sets using the 'train\_test\_split' method, with 80% of the data allocated for training and the remaining 20% for testing.

Subsequently, both the Random Forest and Linear Regression models were constructed. The Random Forest model was trained using the training data and consisted of 100 decision trees. Once the model was trained, predictions were made on the testing data to generate base rate (Bea Dasar) predictions.

A simple Linear Regression model was also built and trained on the same data as the Random Forest model. This model aimed to identify a linear relationship between the features and the target variable. After training, predictions were similarly performed on the test data.

### 3.3. Model Evaluation

The performance of both models was evaluated by calculating various evaluation metrics, namely  $R^2$  (R-squared), MSE (Mean Squared Error), RMSE (Root Mean Squared Error), and MAE (Mean Absolute Error) for the two models: Random Forest Regressor and Linear Regression. The calculated results for each model are presented in Figure 2:

The evaluation results indicate that the Random Forest Regressor outperforms the Linear Regression model. A lower MSE value in the Random Forest model suggests a smaller prediction error. Additionally, the lower MAE further confirms the model's ability to produce more accurate predictions. However, the slightly higher RMSE observed in the Random Forest Regressor indicates a greater average prediction error, particularly when considered in the original scale of the target variable.

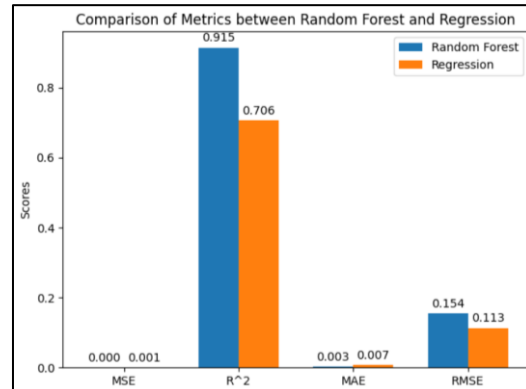


Figure 2. Comparison of Model Evaluation Results

The  $R^2$  value for Random Forest reached 0.915, meaning the model was able to explain approximately 91.5% of the variance in the data, whereas Linear Regression only explained about 70.6% of the variance, with an  $R^2$  of 0.706.

Based on these results, it can be concluded that the Random Forest Regressor demonstrates superior performance compared to Linear Regression, both in terms of predictive accuracy and the model's ability to explain data variability.

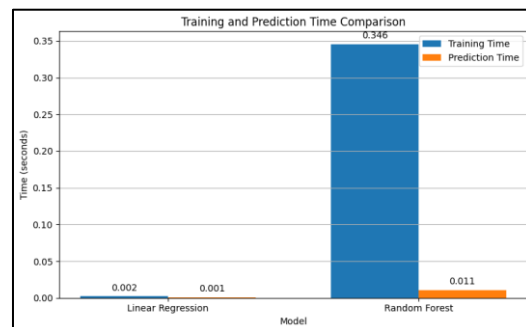


Figure 3. Comparison of Model Evaluation Results

Figure 3 illustrates the computational time comparison between the Linear Regression and Random Forest models during the training and prediction processes. Linear Regression demonstrates significantly faster computational performance, with a training time of 0.002 seconds and a prediction time of 0.001 seconds. In contrast, the Random Forest model requires 0.346 seconds for training and 0.011 seconds for prediction. These findings indicate that although Random Forest generally provides higher accuracy, it incurs a higher computational cost compared to the more efficient Linear Regression model.

### 3.4. Advantages of Random Forest

Random Forest offers significant advantages in handling complex relationships between variables. This model is capable of capturing non-linearity within the data, something that cannot be adequately modeled by linear approaches such as Linear Regression. By aggregating multiple decision trees, Random Forest can account for interactions among variables, resulting in more accurate predictions, especially in scenarios where the relationships between variables are not strictly linear.

In contrast, Linear Regression is limited in its ability to handle high-complexity data, as it assumes a purely linear relationship between predictor variables and the target. This limitation leads to lower performance when dealing with data that exhibits more intricate patterns, as reflected in the evaluation results.

### 3.5. Relationships Between Variables

A positive correlation was found between Harga Barang, Berat Kiriman, and Jarak (KM) with Bea Dasar. Among these variables, Jarak (KM) exhibited the strongest correlation, which aligns with the increasing operational costs as delivery distance grows. Berat Kiriman also showed a strong influence on Bea Dasar, reflecting the additional costs associated with the mass of the shipment. Meanwhile, Harga Barang demonstrated a more moderate but still significant correlation, indicating that the value of goods also contributes to the base rate calculation.

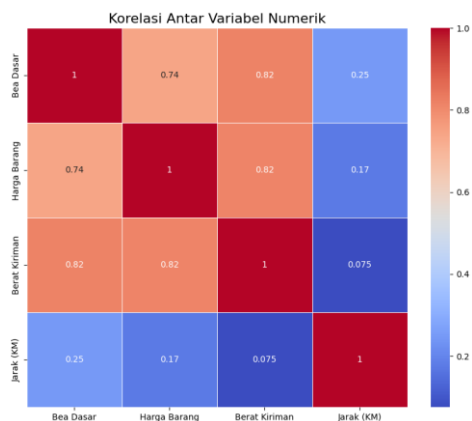


Figure 3. Feature Correlation

## 4. CONCLUSION

The Random Forest Regressor demonstrated superior performance in predicting Basic Shipping Tariffs (Bea Dasar) compared to Linear Regression. It achieved a higher coefficient of determination ( $R^2 = 0.915$ ) and lower error metrics, including  $MSE = 0.000$  and

$MAE = 0.003$ , outperforming Linear Regression which recorded  $R^2 = 0.706$ ,  $MSE = 0.001$ , and  $MAE = 0.007$ . These results confirm the robustness of Random Forest in capturing the nonlinear relationships among key shipping variables such as Item Price, Shipment Weight, and Distance. However, it is worth noting that Random Forest also showed a slightly higher RMSE (0.154) compared to Linear Regression (0.113), indicating greater variance in certain prediction instances, which may need to be addressed for real-world deployment.

This study confirms the potential of machine learning as a strategic tool for enhancing operational efficiency and pricing accuracy in the logistics industry. For future research, it is recommended to enrich the dataset by including additional variables such as commodity type, packaging costs, and insurance value, and to extend the historical data range to capture seasonal effects. Additionally, developing an interactive web-based system that integrates the trained model could facilitate automated pricing, improve user engagement, and support decision-making in real-time logistics operations.

## 5. SUGGESTION

Based on the findings of this study on base rate prediction, several recommendations can be considered for further development.

### 5.1. Enriching the Shipping Dataset

To improve the accuracy of the predictive model, it is recommended to expand the dataset by incorporating additional relevant variables, such as commodity type, country of origin, insurance value, packaging costs, and previous shipment history. Furthermore, utilizing historical data over a longer time span may help the model capture seasonal patterns and long-term trends more effectively.

### 5.2. Web Interface Development

To ensure that the research outcomes are more accessible and useful to end-users, a more interactive and responsive web interface should be developed. The website design should prioritize user-friendliness and include features such as prediction visualizations, manual or file-based data input options, and informative summaries of prediction results.

## REFERENCES

- [1] K. Deorah, "Why Logistics Technology Matters," *Forbes Technology Council*, Dec. 2021.
- [2] A. Zhahra Lubis, L. Lastrian Nahulae, N. Marliana Anggraini, R. Adawiyah, and U. Islam Negeri Sumatera Utara, "Analisis Faktor-Faktor yang Memengaruhi Penetapan Harga," *Jurnal Masharif al-Syariah: Jurnal Ekonomi dan Perbankan Syariah*, vol. 9, no. 1, pp. 25–28, 2024, doi: 10.30651/jms.v9i1.21412.
- [3] J. Kong, Z. Chen, and X. Liu, "A Review of Logistics Pricing Research Based on Game Theory," *Sustainability (Switzerland)*, vol. 14, no. 17, 2022, doi: 10.3390/su141710520.
- [4] A.-F. Raka Praditya and S. Yulianto Joko Prasetyo, "Game Theory Dalam Penentuan Strategi Pemasaran Optimal Dalam (Studi Kasus Persaingan E-Commerce Shopee dan TokoPedia)," *TIN: Terapan Informatika Nusantara*, vol. 2, no. 2, 2021, Accessed: Dec. 13, 2024. [Online]. Available: <https://ejurnal.seminar-id.com/index.php/tin/article/view/799>
- [5] R. Muha, "An overview of the problematic issues in logistics cost management," *Pomorstvo*, vol. 33, no. 1, pp. 102–109, 2019, doi: 10.31217/p.33.1.11.
- [6] N. T. Nafisah, F. Maria, M. R. Amanatullah, and T. Sutabri, "Penggunaan Teknologi Artificial Intelligence Untuk Peningkatan Pembelajaran Pada SMA Nurul Iman Palembang Menggunakan ITIL V3," vol. 18, no. 1, pp. 34–40, 2024, [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>
- [7] William F. Schneider and Hua Guo, "Machine Learning," *Journal of Physical Chemistry Letters*, vol. 8, no. 12, pp. 2689–2694, Jun. 2017, doi: 10.1021/acs.jpcclett.7b01072.
- [8] D. N. Argade, S. D. Pawar, V. V. Thitme, and A. D. Shelkar, "Machine Learning: Review," *International Journal of Advanced Research in Science, Communication and Technology*, pp. 251–256, Jul. 2021, doi: 10.48175/ijarsct-1719.
- [9] E. Akyuz, K. Cicek, and M. Celik, "A Comparative Research of Machine Learning Impact to Future of Maritime Transportation," in *Procedia Computer Science*, Elsevier B.V., 2019, pp. 275–280. doi: 10.1016/j.procs.2019.09.052.
- [10] X. Li, Y. Zheng, Z. Zhou, and Z. Zheng, "Demand Prediction, Predictive Shipping, and Product Allocation for Large-scale E-commerce," *Social Science Research Network (SSRN)*, 2018, [Online]. Available: <https://ssrn.com/abstract=3277125>
- [11] I. Maulita and A. Mu'amar Wahid, "Prediksi Magnitudo Gempa Menggunakan Random Forest, Support Vector Regression, XGBoost, LightGBM, dan Multi-Layer Perceptron Berdasarkan Data Kedalaman dan Geolokasi Predicting Earthquake Magnitude Using Random Forest, Support Vector Regression, XGBoost, LightGBM, and Multi-Layer Perceptron Based on Depth and Geolocation Data," *Jurnal Pendidikan dan Teknologi Indonesia (JPTI).jpti*, vol. 470, no. 5, pp. 221–232, 2024, doi: 10.52436/1.jpti.470.
- [12] A. Maurice, C. Refiyana, and E. A. Vefiadytria, "Uji Asumsi Klasik dalam Regresi Linier pada Perhitungan Menggunakan Laporan Keuangan di Sektor Telekomunikasi Bursa Efek Indonesia (BEI)," *Jurnal Ilmiah Manajemen Ekonomi Dan Akuntansi*, vol. 1, no. Februari, pp. 107–118, 2024, doi: 10.62017/jimea.
- [13] 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, no. 3, pp. 1–24, 2021, doi: 10.7717/PEERJ-CS.623.
- [14] A. Shabur, M. Amadi, and N. Anwar, "Perbandingan Metodologi Studi Islam Tradisional dan Modern di Indonesia," *Jurnal Pendidikan Tambusai*, vol. 7, no. 3, pp. 22519–22526, 2023, doi: <https://doi.org/10.31004/jptam.v7i3.10134>.
- [15] M. S. Murugaiyan and S. Balaji, *Succeeding with Agile software development*. 2012.

- [16] R. Rianti, R. Andarsyah, and R. M. Awangga, "Penerapan PCA dan Algoritma Clustering untuk Analisis Mutu Perguruan Tinggi di LLDIKTI Wilayah IV," *NUANSA INFORMATIKA*, vol. 18, no. 2, pp. 67–77, 2024, [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>
- [17] E. Setiana and V. Retreva Danestiara, "Analisis Sentimen Pelaksanaan Kuliah Online Menggunakan Algoritma Support Vector Machine," *NUANSA INFORMATIKA*, vol. 17, no. 2, pp. 66–70, 2023, [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>
- [18] D. Maulud and A. M. Abdulazeez, "A Review on Linear Regression Comprehensive in Machine Learning," *Journal of Applied Science and Technology Trends*, vol. 1, no. 2, pp. 140–147, 2020, doi: 10.38094/jastt1457.
- [19] P. D. Sugiono, *Statistik Untuk Penelitian*. 2007.
- [20] M. I. Hasan, *Pokok-pokok materi statistik 1 (statistik deskriptif)*. 2003.
- [21] M. Aditya Pratama, M. Munawaroh, W. Joko Pranoto, P. Studi Teknik Informatika, F. Sains dan Teknologi, and U. Muhammadiyah Kalimantan Timur, "Perbandingan Performa Algoritma Linear Regresi dan Random Forest untuk Prediksi Harga Bawang Merah di Kota Samarinda," *Jurnal Ilmu Teknik*, vol. 1, no. 2, pp. 172–182, 2024, doi: 10.62017/tektonik.
- [22] Yoga Religia, Agung Nugroho, and Wahyu Hadikristanto, "Klasifikasi Analisis Perbandingan Algoritma Optimasi pada Random Forest untuk Klasifikasi Data Bank Marketing," *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 5, no. 1, pp. 187–192, Feb. 2021, doi: 10.29207/resti.v5i1.2813.
- [23] E. Fitri, "Analisis Perbandingan Metode Regresi Linier, Random Forest Regression dan Gradient Boosted Trees Regression Method untuk Prediksi Harga Rumah," *JOURNAL OF APPLIED COMPUTER SCIENCE AND TECHNOLOGY (JACOST)*, vol. 4, no. 1, pp. 2723–1453, 2023, doi: 10.52158/jacost.491.
- [24] C. Clara Afrisca, H. N. Rofiq, D. D. Atmoko, K. Keuangan, and R. Indonesia, "Penilaian Properti: Penggunaan Machine learning untuk Prediksi Nilai Sewa," *Jurnal Manajemen Keuangan Publik*, vol. 8, no. 2, pp. 138–155, 2024, doi: <https://doi.org/10.31092/jmkp.v8i2.2922>.
- [25] K. Faizin, U. Islam, N. Sunan, and A. Surabaya, "Analisis Penggunaan Metode Penelitian Evaluasi Pada Penelitian Bahasa Arab Model Pengembangan," *Tabyin: Jurnal Pendidikan Islam*, vol. 03, no. 01, pp. 39–53, 2020, [Online]. Available: <http://e-journal.stai-iu.ac.id/index.php/tabyin>

## Development Of Online Attendance Application With Time And Location Validation At Satpol Pp South Tangerang

Muhamad Rifki Ramadhan<sup>1</sup>, Syamsu Hidayat<sup>\*2</sup>, Joko Susilo<sup>3</sup>

<sup>1,2</sup>Institut Bisnis dan Informatika Kosgoro 1957, Indonesia

<sup>3</sup>Institut Bisnis dan Informatika Kwik Kian Gie, Indonesia

E-mail: <sup>1</sup> [muhamadrifkiramadhan21@gmail.com](mailto:muhamadrifkiramadhan21@gmail.com), <sup>\*2</sup> [syamsuhi3009@gmail.com](mailto:syamsuhi3009@gmail.com),  
<sup>3</sup> [joko.susilo@kwikkiangie.ac.id](mailto:joko.susilo@kwikkiangie.ac.id)

### Abstract

*In order to solve the inefficiencies and errors in manual attendance at the Satpol PP office in South Tangerang, this study creates a web-based attendance system. The Laravel framework is used to construct the system, GPS and reverse geocoding are incorporated to verify the locations of check-in and check-out. The system design is supported by UML diagrams, and structured attendance data is handled using a MySQL database. The Black Box Testing technique verifies that the system's essential features operate as planned. The outcomes demonstrate enhanced precision, real-time tracking, and simplified reporting. To improve system usability and dependability, future suggestions call for including mobile apps, facial recognition, and predictive attendance analytics in order to increase productivity, precision, and openness at the Satpol PP office in South Tangerang.*

**Keywords**— *Web-Based Attendance, GPS Validation, Laravel Framework, Black Box Testing.*

Submission: May 19, 2025

Accepted: June 18, 2025

Published: July 10, 2025

DOI Article : <https://doi.org/10.25134/ilkom.v19i2.398>

### 1. INTRODUCTION

Attendance is very important for human resource management, especially in government organizations with an organized work structure such as the Civil Service Police Unit (Satpol PP). However, until now, many government agencies still use a manually written attendance system. This not only has the possibility of recording errors, but also it is very easy to manipulate, which will make difficult for large and active agencies to recapitulate the data [1][2]. In such a situation, reports may be delayed, and the accuracy of employee performance evaluation may be compromised.

Web-based applications are created to manage employee attendance as time and information technology increase. In this situation, a digital system will automate the recording of attendance and ensure that all this happens in real time [3][4]. It will avoid errors and improve the efficiency of operations. One feature of attendance systems that is currently not very common is the integration of geolocation technologies such as GPS. A location confirmation system ensures that attendance data is more secure because it ensures that the transfer of attendance to the database only occurs if the user is within a certain radius of the authorized work location, either by mobile or physical [5][6].

By using a reverse geocoding algorithm, this system can convert geographic coordinates into addresses that can be understood by users. By using this technology, attendance records become more accurate and system administrators can view disciplines directly [7]. On the other hand, the Agile-based attendance system development method allows a faster and more flexible iterative process to meet user needs [8].

### 2. RESEARCH METHOD

Researchers use the well-known Agile method to develop software in this study. This method focuses on what the user wants and is done repeatedly [9]. We chose the Agile approach because it is very helpful in creating applications that require constant feedback. Each step in the Agile model starts with planning, creating the application's interface and database design, developing parts of the application, testing the system's performance, and improving according to recommendations from experts [8][10].

We use three important methods to obtain information: question and answer, direct observation, and searching from books or other records. We conducted careful interviews with Satpol PP South Tangerang officers, especially the head and staff who handle attendance. In other words, to get an accurate understanding of the manual absence method they have been using,

as well as potential problems. The results of the Q&A show that the old attendance method takes a lot of time, wrong records often occur, and it is impossible to immediately report who is present [1].

Researchers also visit the research location to observe how daily attendance is managed. The goal is to get information through interviews conducted on-site. Additionally, we document the manual check-in procedures, including how administrative staff compile attendance data in Excel each week. This information is valuable for developing a digital system that integrates location-based time tracking and automated features. [11].

Gathering all the user requirements starts with the process of designing the system. This design uses the Unified Modeling Language (UML), which provides various types of diagrams, such as use case diagrams, activity diagrams, and class diagrams. Here, other system requirements are also mapped out. This includes how the database tables are organized, how the user interface (UI) is created, and how to validate the location with GPS coordinates. The system is designed so that users can only take attendance if they are within a certain distance from the office [7].

Since Laravel version 11 is helpful with Model-View-Controller (MVC) architecture, this application was developed using it. This makes easier to differentiate data, views, and business rules. The implementation is done in a local development environment using Laragon and a MySQL database. We also incorporate a geocoding reversal API in the middle of development. The goal is to make location coordinate changes easier to understand. We also set up a user authentication system based on their respective roles, also known as role-based access controllers [12].

The system is tested using the Black Box Testing method, which focuses on testing existing functions without looking at how the program code is created. To ensure that the application can perform the check-in and check-out process correctly, validate the user's location, and record attendance according to predefined rules, each feature is tested with various inputs. The test results show that the system works well. It can record attendance automatically based on GPS radius and allows export of attendance data to Excel on a daily or monthly basis [13].

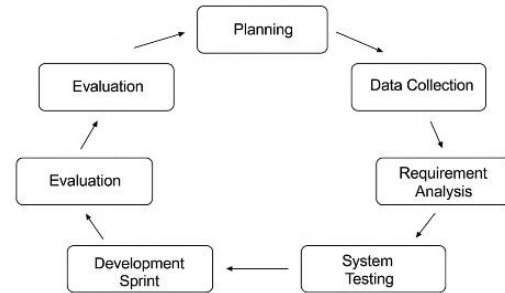


Figure 1. System Development Method

### 3. RESEARCH RESULTS

#### 3.1. System Design

The Unified Modeling Language (UML) method was used to visually describe the system requirements. Use case, activity, and class diagrams are some of the diagrams used. The use case diagram shows how system users, such as members and administrators, interact with important features such as login, check-in, check-out, and export of attendance data. The activity diagram describes the logical flow of system activities, which starts from the login process and ends with the storage of attendance data. However, the class diagram shows the data structure and relationships between entities such as User, Employee, Office Location, and Attendance, each of which has corresponding attributes and methods [14].

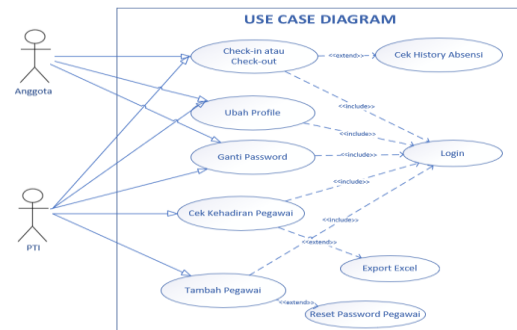


Figure 2. Use Case Diagram System

#### 3.2. Application Interface

The user interface is created using user-centered design principles to ensure that the interface is easy to use, has a clear navigation flow, and is responsive across different types of devices. The various components include a login page, dashboards for members and admins, employee data management, and attendance records. The software is designed to display live attendance status information and make it easy to download data in Microsoft Excel format, which

is useful for recapitulation and performance assessment [15].

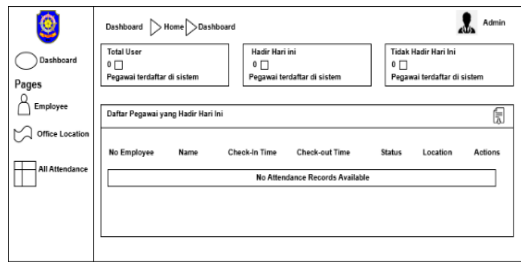


Figure 3. Admin Dashboard

In Figure 3 Describes that the image displays a wireframe for an employee attendance management system. At the top, there is a navigation trail that reads "Dashboard > Home > Dashboard," helping users understand their current location within the system. The left sidebar includes navigation options such as links to the Dashboard, Employee page, Office Location, and All Attendance pages. In the main section of the dashboard, key statistics are presented, including the total number of users, the number of employees present today, and the number of employees absent, each accompanied by appropriate icons. Below this, there is a section titled "Daftar Pegawai yang Hadir Hari Ini" (List of Employees Present Today), which is designed to display a table with details like employee ID, name, check-in/check-out times, attendance status, location, and available actions. Currently, it shows a message stating "No Attendance Records Available," indicating that there are no logs for today. Overall, the layout is clean, user-friendly, and effective for tracking attendance.

### 3.3. Database Design

The system includes not only display design but also database design using MySQL with a meticulous normalization method. Each table, such as user, role, employee, attendance, and office location, is designed to be connected to each other through primary and foreign key relationships. The goal is to ensure data integrity and consistency are maintained. This architecture meets the best standards for relational database development for small to medium-sized applications [16].

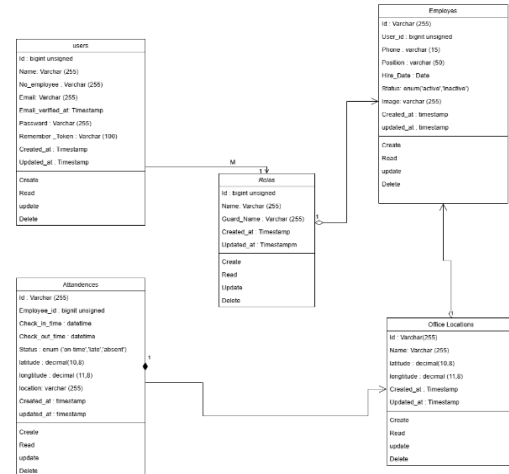


Figure 4. Entity Relation Diagram

Role-based authentication is to ensure secure access. This feature allows admins and regular members to share access rights, which ensures that authorized people can only access specific tasks. This way of working is becoming an essential component in the security design of contemporary information systems, especially for publicly accessible web applications.

Overall, the design stages in this research have followed modern software engineering principles and met user needs. The results of this design become a strong basis for the implementation stage and ensure that the system created will run well and in accordance with organizational goals.

## 4. DISCUSSION

The integration of GPS technology into the attendance app developed for this study is very important. It ensures that employee check-ins and check-outs are accurate and reliable. With GPS, the system can verify whether an employee is physically present at the office or their designated work location before recording their attendance. This significantly reduces the likelihood of false check-ins, such as one employee marking attendance for another who is not present. Additionally, GPS enables real-time tracking of field employees, simplifying team management. Therefore, in this study, the use of GPS not only enhances the functionality of the app but also aligns it with modern attendance system standards.

### 4.1 Implementation and Testing

A developed a web-based attendance app that incorporates GPS features through a simple development process, using modern web technologies and location services. The app is

launched on a secure web server, allowing access from any device with a web browser. During the setup, we configure GPS modules to capture real-time location data whenever employees check in or out. This ensures that attendance is recorded only when they are physically present in the office. We conduct thorough testing to verify the system's functionality without needing to examine the internal code. This testing assesses the accuracy of the location-based attendance records, the responsiveness of the user interface, and the proper synchronization of data between users and the server. The results confirm that the system accurately captured location data and effectively blocked unauthorized check-ins, demonstrating that it is ready for use.

#### 4.1.1. Hardware and Software Implementation

Implementation of hardware and software is an implementation related to the devices used to design a Web-based Member Attendance Application.

- Processor : Intel 11<sup>th</sup> Gen i3
- RAM Memory : 8 GB
- Storage : SSD 512GB

The software used in making this attendance system :

- Operating System : Windows 11 Pro
- Language Programmings : PHP
- Text Editor : Visual Studio Code
- Database : MySql

A very important initial stage in the process of creating a web-based attendance application is installing and setting up the development environment. Researchers chose Laragon as a web server because it is lightweight, easy to use, and can run important components such as Apache, MySQL, and PHP simultaneously. In addition, Laragon has an auto virtual host feature that makes it easy to test projects locally using a special domain such as <http://absensi.test>, so that development can run more efficiently and more closely resemble the actual server environment. This section is an analysis of the results obtained. Avoid excessive citation and discussion of published literature. State the significance of the results and their implications.

Apache was installed as the main web server and used to manage users' HTTP requests, manage static and dynamic files, and run PHP scripts integrated in Laravel applications. The researcher changed the basic configuration, such as the default port, root directory, and supporting modules, such as `mod_rewrite`, during installation. This was done to ensure that URL

routing runs smoothly and dynamically. Apache is essential because it allows the system to work with a client-server model, where backend scripts on the server control the user interface accessed through the browser.

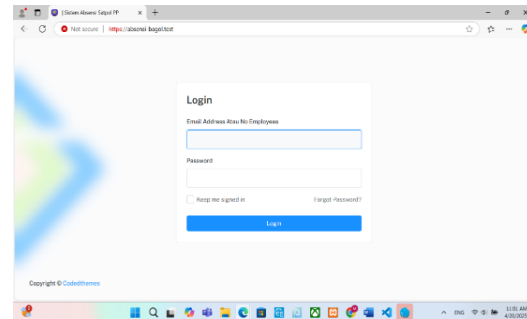


Figure 5. App Login Page

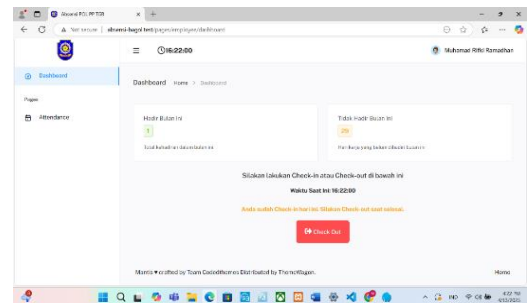


Figure 6. Attendance Page

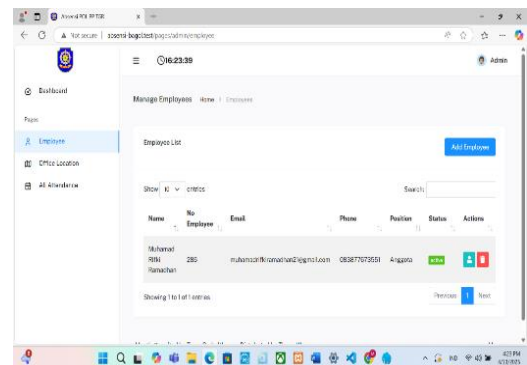


Figure 7. Attendance History

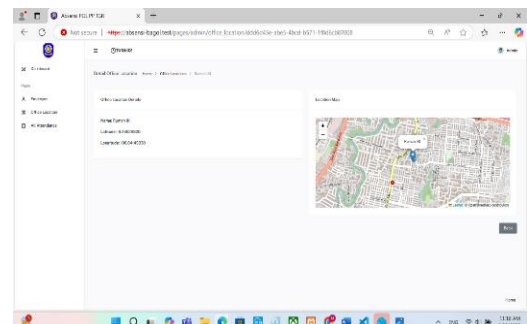


Figure 8. Absence location details

For now, the database is configured using MySQL through the phpMyAdmin interface provided by Laragon. Once created, the database is named `db_absence` and populated with the table structure as used in the previous design, such as users, roles, employees, attendances, and office\_locations. To maintain data integrity and prevent duplicates, each table is linked with primary and foreign key relationships. Using the Laravel Artisan CLI, the database migration is automated. The CLI automatically creates table schemas based on the definitions given in the pre-built migration file.

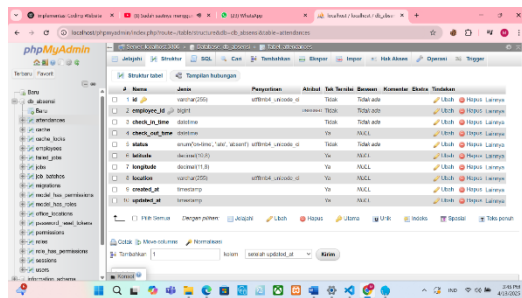


Figure 9. Manage Database

This step builds the foundation for storing attendance data and provides the flexibility to manage schemas, update tables, and integrate them into the application's business logic features. Installing and setting up Apache and MySQL correctly ensures the system can be run locally for testing or transferred to a production server with minimal changes.

#### 4.1.2. Testing

The system testing process is carried out to ensure that each feature of the web-based attendance application works according to specifications and user needs. Black box testing is used; it tests system results based on specific inputs without looking at the code structure. This technique is useful for evaluating the system's ability to handle various usage scenarios, such as data validation, attendance processing, and additional functions such as data export.

The test results show that the main features of the application run well as expected. Six important scenarios were tested: login with correct data; check-in and check-out process with time and location checking; refusal to check-in outside the office radius; access by members to attendance history; and data export by admin. For each test, the system responded according to the scenario, with both warning and success messages. This shows that the business logic and

application security have been effectively implemented.

The successful testing of each feature shows that the system has met the principles of fit for use and fit for purpose. This means that the system can not only be used well, but also meets the organization's objectives-providing accurate, clear, and reliable attendance data. This supports the claim of Silitonga et al. (2023), which states that Black Box-based testing can ensure the quality and functional integrity of web-based systems [17].

## 5. CONCLUSION

1. This research creates and implements a web-based attendance application that improves the efficiency and accuracy of recording attendance of Satpol PP South Tangerang members. By using GPS technology and reverse geocoding for time and location validation, the system automates the attendance reporting process. As a result, the manual process of taking attendance, which is prone to many errors or even fraud, is now faster, more precise, and digitally documented in real time and continuously automated.
2. In addition, the developed attendance system is equipped with an interactive monitoring dashboard. This allows supervisors or administrators to access it quickly and provides clear information. The dashboard displays daily and monthly attendance history and attendance statistics. This simplifies supervision, report generation, and data-driven decision-making. The system not only increases the transparency and accountability of management reporting on e-Board services, but also allows data export to Excel files and categorizes data based on attendance status in real time (e.g., present, late, permission, sick).

## 6. SUGGESTION

The system can be expanded to include facial recognition features to improve the validity of the attendance system. In addition, it is necessary to test performance in a production environment, develop a mobile version that allows absence notifications, and thoroughly examine employee discipline.

## REFERENCES

- [1] M. Saied and A. Syafii, "Perancangan dan Implementasi Sistem Absensi Berbasis Teknologi Terkini Untuk Meningkatkan Efisiensi Pengelolaan Kehadiran Karyawan dalam Perusahaan," *Jurnal Teknik Indonesia*, vol. 2, no. 3, pp. 87–92, Jul. 2023, doi: 10.58860/jti.v2i3.21.
- [2] A. R. Kamila, G. H. Derhass, D. A. Rabbani, J. F. Andry, and F. S. Lee, "Aplikasi Absensi Berbasis Android Pada Sekolah Boarding Sebagai Transformasi Digital Bidang Pendidikan," *NUANSA INFORMATIKA*, vol. 18, no. 2, pp. 26–34, Jul. 2024, doi: 10.25134/ilkom.v18i2.155.
- [3] H. Kholid Kz, M. Syahfiar, S. Jana, Muh. F. Abhyasa, and A. Firmansyah, "Pembuatan Website BUMDes Berbasis Opensid di Desa Wano untuk Meningkatkan Pengetahuan Masyarakat," *Jurnal Pengabdian Masyarakat Bangsa*, vol. 2, no. 11, pp. 4905–4918, Jan. 2025, doi: 10.59837/jpmmba.v2i11.1893.
- [4] D. Yusuf and S. Supriyadi, "Penerapan Sistem Kehadiran Mahasiswa Berbasis Web Menggunakan Openstreetmap," *NUANSA INFORMATIKA*, vol. 16, no. 2, pp. 147–153, Jul. 2022, doi: 10.25134/nuansa.v16i2.6292.
- [5] D. Purwanto, R. E. Putri, Y. Fadly, and D. C. Pratiwi, "Sistem Absensi Online Berbasis Web Dengan Penggunaan Teknologi GPS," *Jurnal Minfo Polgan*, vol. 13, no. 2, pp. 1800–1811, Nov. 2024, doi: 10.33395/jmp.v13i2.14258.
- [6] D. Haerofifah, "Perancangan Aplikasi Pemesanan Makanan Berbasis Web," *NUANSA INFORMATIKA*, vol. 16, no. 1, pp. 101–107, Jan. 2022, doi: 10.25134/nuansa.v16i1.4771.
- [7] E. Sulisty, M. I. Bayu P, F. Andini, and I. Dwisaputra, "Implementasi Metode Reverse Geocoding pada Aplikasi Tracking Posisi," *Emitor: Jurnal Teknik Elektro*, vol. 1, no. 1, pp. 44–49, Mar. 2023, doi: 10.23917/emitor.v1i1.21464.
- [8] H. Niklas, M. Haikal, and W. T. Atmojo, "Implementasi Metode Agile Dalam Pengembangan Aplikasi Absensi Berbasis Web dengan Menggunakan Geofencing," *Jurnal Komtika (Komputasi dan Informatika)*, vol. 8, no. 2, pp. 200–213, Nov. 2024, doi: 10.31603/komtika.v8i2.12544.
- [9] M. Sidiq, R. Dwicahya Supriatman, E. Ahmad Firdaus, and B. Agung Suburdjati, "Perancangan Arsitektur Sistem Informasi Menggunakan Metode Agile Dengan Kerangka Kerja Scrum Pada Pelayanan Instalasi Gizi RSUD. Ciamis," *NUANSA INFORMATIKA*, vol. 18, no. 1, pp. 53–67, Jan. 2024, doi: 10.25134/ilkom.v18i1.52.
- [10] M. Eric Iryanto Nur Efendi and B. Yulisa Geni, "SISTEM ABSENSI ONLINE BERBASIS WEB DENGAN FITUR RADIUS PEMBATAAN LOKASI ABSEN UNTUK KARYAWAN OUTSOURCING DEPARTEMEN CONTACT CENTER PT PEGADAIAN," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 8, no. 4, pp. 7428–7435, Jul. 2024, doi: 10.36040/jati.v8i4.10206.
- [11] J. Jalis, S. Auliana, and B. Rakhim Setya Permana, "PENERAPAN FRAMEWORK LARAVEL PADA APLIKASI NILAI SISWA BERBASIS WEB DI SDN TERUMBU KOTA SERANG," *JATI (Jurnal Mahasiswa Teknik Informatika)*, vol. 9, no. 2, pp. 3338–3342, Apr. 2025, doi: 10.36040/jati.v9i2.12755.
- [12] N. Amelia Silitonga, L. Hanan Dio Samartha, T. Adi Wiguna, M. Baidawi, and A. Saifudin, "Pengujian Black Box Testing Pada Aplikasi Absensi Berbasis Web Di Sekolah SDN Duri Kepa 01," vol. 2, no. 5, 2023, [Online]. Available: <https://journal.mediapublikasi.id/index.php/oktal>
- [13] G. B. Santoso, T. M. Sinaga, and A. Zuhdi, "MVC Implementation In Laravel Framework For Development Web-Based E-Commerce Applications," *Intelmatiks*, vol. 1, no. 1, Jan. 2021, doi: 10.25105/itm.v1i1.7867.
- [14] P. Apriadi and E. Sutrisna, "Perancangan Aplikasi Absensi Karyawan Berbasis Mobile Menggunakan GPS (Studi Kasus PT. Trans Retail Indonesia)," *Journal Automation Computer Information System*, vol. 3, no. 1, pp. 1–9, May 2023, doi: 10.47134/jacis.v3i1.54.
- [15] B. Nugroho, M. S. Hasibuan, and M. H. Annabil, "Perancangan Aplikasi Absensi Pegawai Berbasis Web Dengan Blackbox Testing pada DISPORA Sumatera Utara," *Journal of Computer Science and Informatics Engineering*.

## Web-Based Assignment Management System Application Design to Increase Personnel Productivity and Effectiveness in Product Documentation Division

Seftian Candra Pratama<sup>1</sup>, Eryan Ahmad Firdaus<sup>\*2</sup>, SettingsRochedi Idul Adha<sup>3</sup>

<sup>123</sup>Universitas Pertahanan, Indonesia

E-mail: <sup>1</sup>[seftianpratama290@gmail.com](mailto:seftianpratama290@gmail.com), <sup>\*2</sup>[eryan.firdaus@idu.ac.id](mailto:eryan.firdaus@idu.ac.id), <sup>3</sup>[rochedi.adha@idu.ac.id](mailto:rochedi.adha@idu.ac.id)

### Abstract

*An effective task management is essential for enhancing productivity and operational efficiency within organizations. Currently, the Production Documentation Division of the Information Service of the Indonesian Army (Dispend) uses WhatsApp for task coordination, which results in fragmented communication, limited tracking, and reduced transparency. To overcome these issues, this study develops a web-based assignment management system to improve task distribution, monitoring, and reporting. The system is built using the Rapid Application Development (RAD) model, which emphasizes iterative prototyping and users' feedback. Key features include task creation, progress tracking, file submission, approval workflows, and reporting tools, supported by role-based access control for both administrators and users. The implementation of this system shows significant improvements in workflow efficiency, transparency, and accountability compared to the previous WhatsApp-based method. In conclusion, the system contributes to better documentation, timely task completion, and enhanced collaboration, thereby increasing overall productivity. Future enhancements may include mobile integration and real-time notifications to further improved accessibility and operational effectiveness.*

**Keywords**—Task Management System, Web-Based Application, Productivity, Dispend

Submission: March 21, 2025

Accepted: June 24, 2025

Published: July 10, 2025

DOI Article : <https://doi.org/10.25134/ilkom.v19i2.368>

### 1. INTRODUCTION

The Product documentation Division is responsible for creating content to be distributed across various social media platforms owned by the Indonesian Army, focusing on task assignment, content creation, and efficient reporting of the final products to superiors.

However, several challenges arise in the current workflow, particularly in communication and task management. Currently, communication relies heavily on WhatsApp, which presents several limitations. WhatsApp is not designed for complex project management, resulting in scattered, difficult-to-track, and poorly organized information. Additionally, the distribution of information across multiple groups complicates accessibility and reference retrieval. Monitoring task progress and generating reports is also challenging due to the lack of integrated project management features, leading to delays and reduced transparency.

To overcome the inefficiencies in task management within the product Documentation Division of Dispend, which currently relies on WhatsApp for coordination, this study aims to design and develop a web-based assignment management system. The existing use of WhatsApp results in scattered communication,

lack of task tracking, and reduced accountability, making it unsuitable for structured content production workflows. Therefore, the main objective of this project is to enhance personnel productivity and effectiveness in social media content creation by implementing a centralized digital platform. The proposed solution offers streamlined task distribution, integrated communication, real-time progress tracking, and clear reporting mechanisms to improve workflow structure, transparency, and overall organizational efficiency.

The development of this web-based application is expected to offer several benefits, including replacing the existing personnel assignment system that relies on WhatsApp, which is unsuitable for structured task management. The system will function as a dedicated task management platform for product documentation division personnel, enabling them to work more effectively and productively in social media content production. The application will streamline task assignment, tracking, and reporting, ensuring that all personnel remain well-informed and accountable for their respective responsibilities.

To ensure an optimal and secure implementation within the limited time frame of

the OJT period, the scope and limitations of this project have been clearly defined. The system will specifically focus on managing task assignments for product documentation division personnel under Sub-Department of Documentation Coverage and Production Dispenad. It will consist of two primary roles: the Admin, responsible for assigning and monitoring tasks, and the User, who receives and executes the tasks. The workflow will align with the existing practices in product documentation division, where tasks are assigned based on designated roles or direct appointments from superiors. The system will facilitate task submission, approval, and revision processes, with a task history log for easy tracking. Additionally, the platform will allow file uploads and link-based submissions to accommodate large media files. The Admin will have full visibility of all user activities and performance, while Users will only have access to their assigned tasks. This structured approach is expected to enhance the efficiency of task distribution, tracking, and completion, ensuring a more productive working environment for product documentation division personnel.

## 2. RESEARCH METHOD

The development of the web-based assignment management system for product documentation division Dispenad follows the Rapid Application Development (RAD) model, which is a System Development Life Cycle (SDLC) approach that enables rapid prototyping, continuous feedback, and iterative improvements [1].



Figure 1. RAD Method

The development process begins with the Requirements Planning phase, where system needs are gathered through discussions with stakeholders, identifying inefficiencies in task management due to the reliance on WhatsApp for assignment distribution and tracking. In the User Design phase, a prototype is developed using Unified.

Modeling Language (UML) diagrams, including Use Case, Activity, Sequence, and Class Diagrams, to visualize system functionalities and facilitate user feedback for refinements [2]. The Construction phase involves system development using Laravel for backend processing, Bootstrap for frontend design, and

MySQL as the database, ensuring seamless integration of features such as task assignment, tracking, and reporting [3]. Iterative testing is conducted to improve system performance, leading to the Cutover phase, where Blackbox testing is employed to verify system functionality without assessing the internal code structure [4]. The implementation of this methodology ensures a structured and agile development process, providing product documentation division personnel with an efficient and scalable task management solution that enhances productivity, transparency, and accountability in handling assignments.

## 3. RESEARCH RESULTS

A use case diagram is a graphical representation of some or all actors, use cases, and their interactions, introducing a system [2]. A use case diagram does not provide detailed explanations about the use of use cases but only gives a brief overview of the relationships between use cases, actors, and the system. Through this use case, the functions within the developed system can be identified.

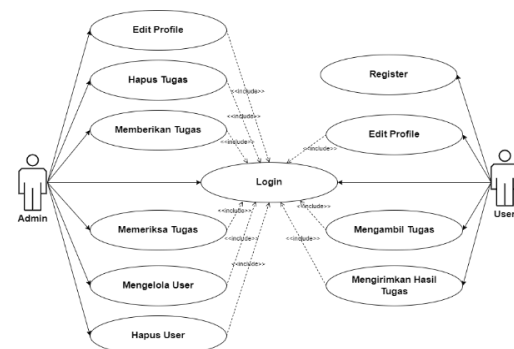


Figure 2. Use Case Diagram Model

In the use case diagram, once the admin successfully logs into the system, they are granted access to perform various actions, including managing users, assigning tasks, and reviewing task submissions uploaded by users. In terms of user management, the admin has the authority to add, edit status, or delete user accounts. Additionally, the admin can assign new tasks, which will then be received by the designated users. The admin also has the ability to monitor task progress, review uploaded files and links, and manage the task list, including searching for specific tasks and removing unnecessary ones.

### 1. Activity Diagram

An Activity Diagram represents the flow of activities or workflows within a system to be

executed. It is important to note that the activity diagram illustrates the activities that can be performed by the system, not what the actors do [2].

### Activity Diagram for Admin Login

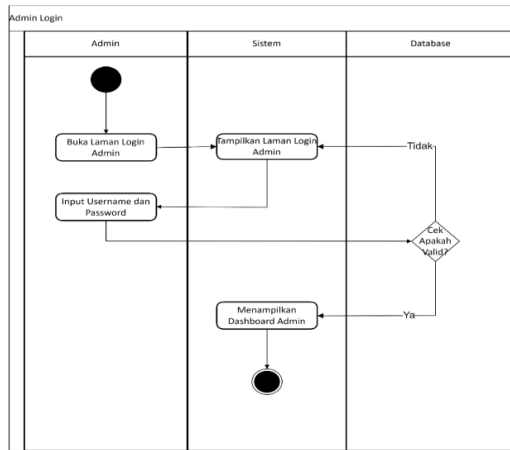


Figure 3. Login Activity Diagram for Admin

In the admin login activity, the first step involves accessing the admin login page through the provided URL. Once the page is displayed, the admin is required to enter a valid username and password into the login form. Upon pressing the login button, the system initiates a validation process to verify whether the entered credentials match the data stored in the database.

### Activity Diagram Admin for Assigns Tasks

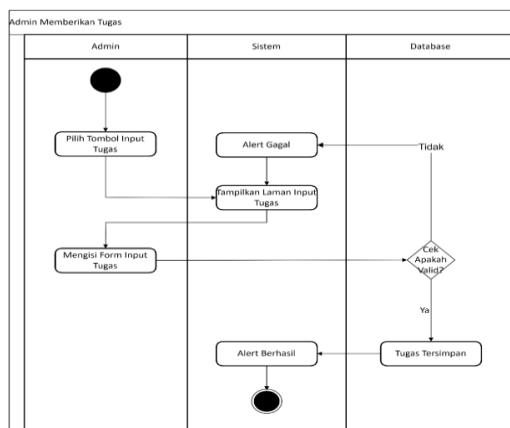


Figure 4. Activity Diagram for Assigns Tasks

The admin has the ability to assign tasks to users. The process begins with the admin accessing the Task Input Page. Once the page is displayed, the admin fills out the provided form with the necessary information, including the task title, description, deadline, and other relevant details.

### Activity Diagram Admin for Checking Tasks

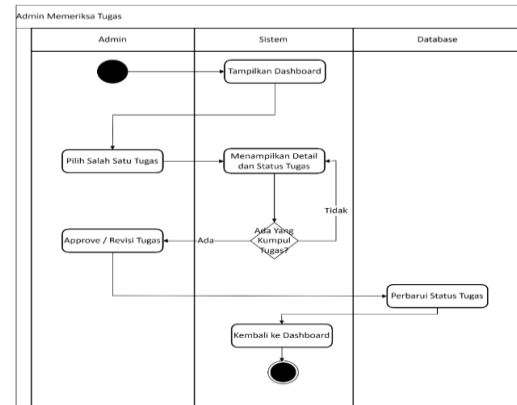


Figure 5. Activity Diagram for Checking Tasks

Once the task has been assigned by the admin, users are required to submit their completed tasks according to the given instructions. Users can either upload a file or provide a link to their task submission through the designated page. Once the submission is uploaded, the admin gains access to review and evaluate the task.

### Activity Diagram User for Takes Tasks

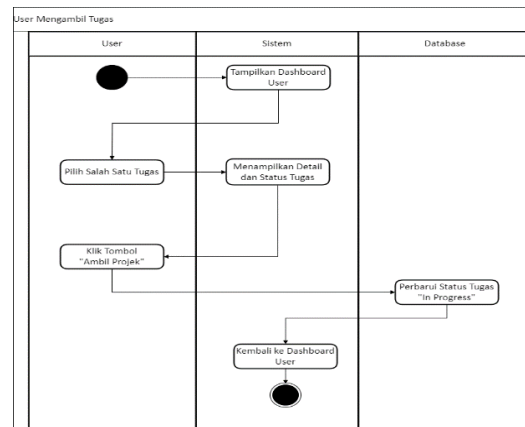


Figure 6. Activity Diagram for Takes Tasks

Once the user's account has been activated by the administrator, they can immediately access their dashboard page. On this page, users can view a list of assigned tasks, which are displayed with basic information such as the title, deadline, and task status.

### Activity Diagram User for Submit Tasks Result

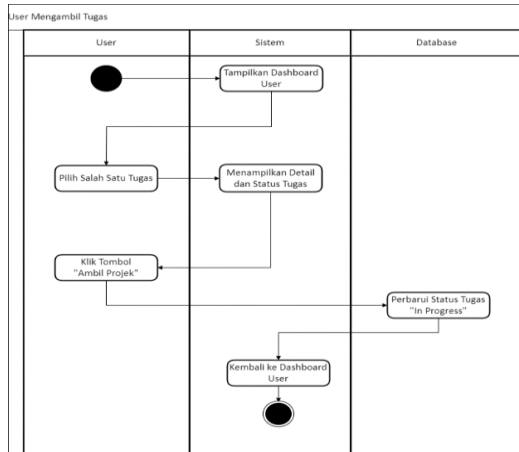


Figure 7. Activity Diagram Submit Tasks Result

Once the user's account has been activated by the administrator, they can immediately access their dashboard page. On this page, users can view a list of assigned tasks, which are presented with essential information such as the title, deadline, and task status.

## 2. Interface Display

### Admin Login Interface Display

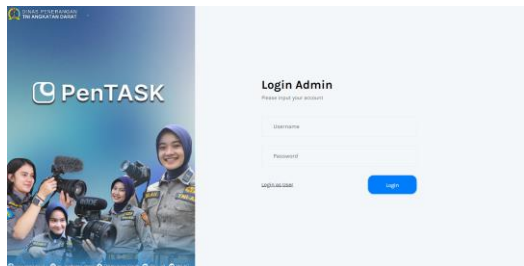


Figure 8. Admin Login Interface Display

Login is an authentication process that allows an individual (admin) to access an application system or website using credentials, typically consisting of a username and password. This process ensures that only authorized users can access the available information or services. In the admin login section, the admin can log in using a predefined username and password.

### Member List Display

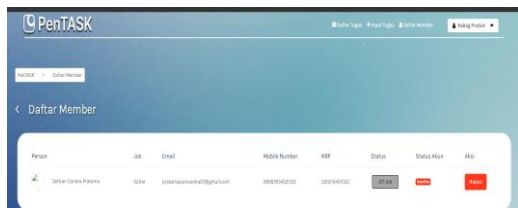


Figure 9. Member List Display

The member list is a dedicated page accessible exclusively to administrators, allowing them to view all registered personnel. Within this list, administrators can process new members who have completed their registration and monitor their employment status whether they have secured a job or are still seeking employment.

### Form Input Tasks Display

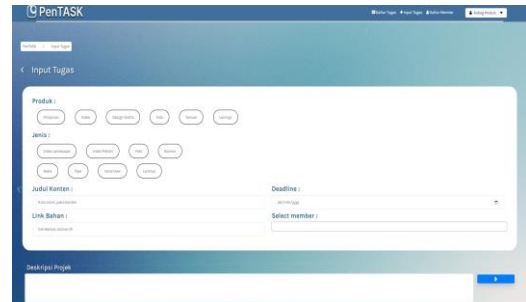


Figure 10. Form Input Tasks Display

Task input is one of the core functions of the Pentask web system, performed by administrators to assign new tasks that will be executed by product documentation division Dispenad personnel in the field. The primary function of task input involves filling out a predefined form and submitting it to the Task List page. To initiate task input, the administrator must click the "Task Input" button located in the top navigation bar.

### Review Admin Display

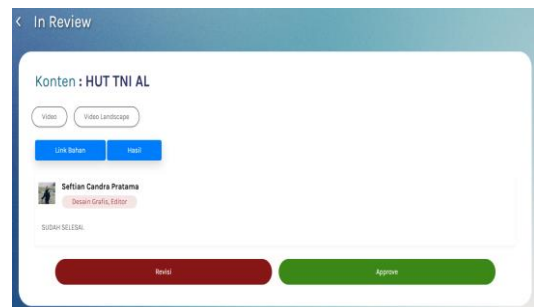


Figure 11. Review Admin Display

The primary function of the administrator in the Task List is to review and evaluate social media content submitted by users. Administrators assess whether the submitted work meets the required standards and decide whether to approve the task or request revisions.

### Form Registration Display



Figure 12. Form Registration Display

The sign-up process on this website requires users to complete all fields in the registration form according to the specified requirements. Notably, the password must contain a combination of uppercase letters, lowercase letters, numbers, and symbols to ensure security.

#### Description Take Task Display



Figure 13. Description Take Tasks Display

When a new task is registered in the Task List, the user can initiate the process of accepting the assigned task by clicking the "New Task Status" button. This action will display the task description, allowing the user to review the details before proceeding.

#### User Task Submission Form Display

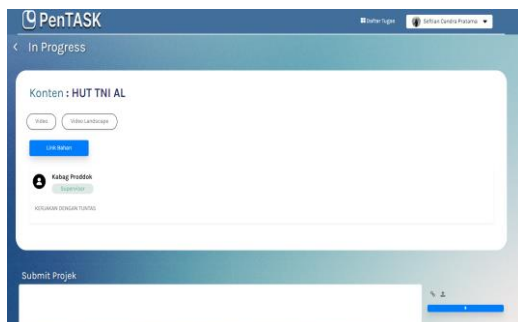


Figure 14. User Task Submission Form Display

The following is the form interface where users submit the final product of their assigned tasks. Users are required to fill in the project description in the "Submit Project" field and choose whether to submit the product as a link if the file size is large or as an attached file. Upon clicking the submit arrow button, the submitted product undergoes a review process by the admin.

### 3. Testing

Based on the results of the testing process conducted on the developed web-based task management system, the system was found to function in accordance with the predefined expectations. Testing began from the login page, where various scenarios were evaluated, including accessing the system without logging in, using valid and invalid username-password combinations, and session persistence across tabs. The system successfully redirected users to the appropriate pages or displayed relevant warning messages as expected.

In the profile management section, the functionalities to update account information and perform logout were executed correctly. For task-related features, the system demonstrated the ability to display a comprehensive list of tasks, open specific task details, and accurately reflect task status such as "New," "In Progress," "In Review," and "Done." Furthermore, functionalities like revision submissions, task approvals, and task confirmation were performed without errors, ensuring a seamless review and approval workflow.

The task input module allowed for successful task creation when all required fields were completed, and it appropriately restricted submissions when mandatory fields were left empty. The search and filter features operated effectively in helping users locate specific tasks. Additionally, the system allowed for the deletion of tasks and correctly updated the task list accordingly.

In the member management section, account activation, viewing paginated member lists, and deletion of user accounts were all executed successfully, with changes accurately reflected in the system. Overall, the testing results indicate that the system met all functional requirements, offering a reliable platform for managing assignments, enhancing task tracking, and improving overall productivity in the Product Documentation Division of the Indonesian Army's Public Relations Office.

### 4. Comparison

This comparison was conducted to compare the assignment management system in the Proddok subdisLipproddok Dispenad section which still uses applications such as Whatsapp or manually via verbal commands and the assignment management system if it uses a web-based system innovation that has been designed.

Table 1. Comparison of assignment management systems

| Comparison Aspect       | Current Task Management System Manual/WhatsApp                                                            | New Web-Based Task Management System                                                             |
|-------------------------|-----------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------|
| Task Assignment Method  | Difficult to understand due to instructions being given via WhatsApp group chats or verbally in the field | Task assignments are clearly presented and easy to understand directly within the web platform   |
| Task Delivery Process   | Slow to reach the intended personnel due to messages buried in group chats requiring scrolling            | Delivered quickly and directly to each personnel's task list on their respective accounts        |
| Task Evaluation Process | Evaluation is delayed, often repeated verbally, and leads to miscommunication or misunderstanding         | Evaluation is clearly documented within task status, minimizing confusion and forgotten feedback |
| Task Reception          | Delayed due to cluttered group chat columns                                                               | Received instantly as the task appears directly in the user's task list                          |
| Task History            | No clear record of completed tasks; difficult to trace past assignments                                   | Easy to track with a neatly organized task history in the system                                 |
| Personnel Productivity  | Personnel activity is not tracked; unassigned members remain unproductive                                 | High productivity as the system shows real-time task engagement status                           |
| Clarity of Task         | Often unclear; users need to ask                                                                          | Instructions are                                                                                 |

|                             |                                                                         |                                                                                |
|-----------------------------|-------------------------------------------------------------------------|--------------------------------------------------------------------------------|
| Instructions                | again or search through chat history                                    | structured and detailed, with materials and descriptions displayed in one view |
| Task Monitoring             | Tasks are monitored manually in the field or by asking personnel        | Easily monitored through task status indicators within the system              |
| Social Media Content Output | Outputs must be searched for or downloaded from chat, often unorganized | Completed outputs are well-organized and easily accessible in the web system   |

The comparison between the current manual task management system primarily using WhatsApp or verbal instructions and the newly developed web-based task management system highlights significant improvements in efficiency, clarity, and productivity. The web-based system streamlines task assignment and reception processes, making them faster and more transparent. Task evaluation becomes more structured, minimizing misunderstandings and repeated instructions. Additionally, the system provides a clear task history and real-time tracking of personnel productivity, which is lacking in the current approach. Task instructions are presented in a well-organized format, and monitoring is greatly enhanced through integrated status updates. Lastly, the final outputs, such as social media content, are more systematically stored and easily accessible, leading to better documentation and workflow management. Overall, the web-based system offers a comprehensive and effective solution that significantly enhances operational performance within the Product Documentation Division.

#### 4. CONCLUSION

##### A. Conclusion

Based on the analysis and implementation of the web-based task management application for product documentation division Dispenad

personnel, the following conclusions were drawn:

1. The task management process, which previously relied on WhatsApp, can be improved through the implementation of a web-based application. This system enhances personnel performance in creating social media content for the Indonesian Army (TNI AD) by streamlining the entire workflow, from task assignment and acceptance to execution and evaluation of completed tasks.
2. The web-based application serves as a new task management system for product documentation division Dispenad personnel, making the process of producing social media content more productive and efficient throughout the workflow. Additionally, this system is expected to support efforts in achieving a higher quantity of high-quality social media content produced by product documentation division personnel in the future.

#### B. Recommendations

During the implementation of the On-the-Job Training program at the Public Relations Office of the Indonesian Army (Dinas Penerangan TNI AD) and the development of the web-based task management system for product documentation division personnel, several recommendations have been identified for future improvements and further considerations:

1. To support broader usability and user behavior studies in mobile environments, future research is recommended to explore the development of a mobile version of the system. This would enable comparative analysis of task management efficiency and user engagement between web-based and mobile platforms.
2. Future system development should incorporate intelligent notification features, such as adaptive reminders and predictive alerts, which could be studied for their impact on user performance and time management. Research can further examine how real-time alerts affect task prioritization and completion rates.
3. It is recommended to expand the system's scope to include additional sub-divisions within the Public Relations Office. This opens opportunities for further academic investigation into cross-functional task coordination, organizational behavior in

digital work environments, and the impact of centralized systems on overall institutional productivity.

#### REFERENCES

- [1] Hidayat, N., & Hati, K. (2021). Penerapan Metode Rapid Application Development (Rad) Dalam Rancang Bangun Sistem Informasi Rapor Online (Siraline). *Jurnal Sistem Informasi*, 10(1), 8–17. <https://doi.org/10.51998/Jsi.V10i1.352>.
- [2] Mukherjee, M., & Mukherjee, M. (2016). Object-Oriented Analysis And Design. *International Journal Of Advanced Engineering And Management*, 1(1), 18–24. <https://srn.com/abstract=2905481> <https://srn.com/abstract=2905481> [electronic copy available at: https://srn.com/abstract=2905481](https://srn.com/abstract=2905481)
- [3] Tahir, T. Bin, Rais, M., & Sirad, M. A. H. (2019). Aplikasi Point Of Sales Menggunakan Framework Laravel. *Jiko (Jurnal Informatika Dan Komputer)*, 2(2), 55–59. <https://doi.org/10.33387/Jiko.V2i2.1313>
- [4] Indah, S., Ayuningtyas, N., Siti, N., Ambarwati, S., Teknik, F., Dan, K., & Kecantikan, P. (2024). Pengembangan Sistem Penugasan Karyawan Berbasis Website Dengan Metode Rapid Application Development Di Galeri X Wedding. *Jurnal Adijaya Multidisplin*, 2(03), 575–596–575–596. <https://ejournal.naurendigiton.com/index.php/Jam/article/view/360>
- [5] Jibril, M., Zulrahmadi, & Amin, Muhammad. (2024). Pengujian Sistem Informasi E-Modul Pada Smpn 1 Tempuling Menggunakan Black Box Testing. *Jurnal Perangkat Lunak*, 6(2), 327–332. <https://doi.org/10.32520/Jupel.V6i2.3326>
- [6] Rohmayani, D., & Putra, B. L. (2021). Sistem Informasi Manajemen Penugasan Pegawai Berbasis Website Studi Kasus : Seameo Seamolec. *Journal Of Informatics And Electronics Engineering*, 1(1), 14–20. <https://ejournal.poltektedc.ac.id/index.php/jice/article/view/499>
- [7] Sanjaya, R., & Hesinto, S. (2017). Rancang Bangun Website Profil Hotel Agung Prabumulih Menggunakan Framework Bootstrap. *Jurnal Teknologi Dan Informasi*, 7(2), 57–64. <https://doi.org/10.34010/Jati.V7i2.758>
- [8] Kurniawan, M., & Kurniawan, B. (2020).

- Implementasi Pemrograman Python Menggunakan Visual Studio Code. *Jik: Jurnal Informatika Dan Komputer*, 11(2), 1–9.  
<https://journal.unmaha.ac.id/index.php/jik/article/view/198>
- [9] Haviluddin, H. (2016). Memahami Penggunaan Uml (Unified Modelling Language). *Informatika Mulawarman : Jurnal Ilmiah Ilmu Komputer*, 6(1), 1–15.  
<https://ocs.unmul.ac.id/index.php/jim/article/view/16>
- [10] Firnando, J., Franko, B., Pratama Tanzil, S., Wilyanto, N., Christianto Tan, H., & Hartati Kom, E. M. (2023). Pembuatan Website Menggunakan Visual Studio Code Di Sma Xaverius 3 Palembang. *Fordicate*, 3(1), 1–8.  
<https://doi.org/10.35957/Fordicate.V3i1.5057>
- [11] Firdaus, E. A., Syani, M., Muttaqin, M. R., Maulani, S., Ciamis, U. G., Tedc Bandung, P., Keperwatan, A., & Dustira Cimahi, R. S. (2022). Perancangan Sistem Informasi Penugasan Dan Aktivitas Karyawan Pada Pt. Xyz. *Nuansa Informatika*, 16(2), 66–76.  
<https://doi.org/10.25134/Nuansa>
- [12] Fauziana, C. A. F. C. A., Gunawan, S. G. S., & Saifudin, A. S. A. (2024). Pengujian Sistem Aplikasi Presensi Siswa Berbasis Web Pada Sma Terbuka Menggunakan Metode Black Box Testing Equivalence Partitioning. *Jurnal Ilmu Komputer, Teknik, Dan Multimedia*, 2(02), 47–52.  
<https://doi.org/10.9999/Qqhfp369>
- [13] Ekaputra, B. A., Pradana, F., & Pramono, D. (2020). Pengembangan Sistem Management Dalam Melakukan Penugasan Berbasis Web. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 4(9), 2791–2800. <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/7768>
- [14] Christian, A., Hesinto, S. (2018). Rancang Bangun Website Sekolah Dengan Menggunakan Framework Bootstrap ( Studi Kasus Smp Negeri 6 Prabumulih ). *Jurnal Sisfokom (Sistem Informasi Dan Komputer)*, 7(1), 22–27.  
<https://doi.org/10.32736/Sisfokom.V7i1.278>
- [15] Phan, I. K., & Yuricha, Y. (2023). Implementasi Pendekatan Backendless Dalam Rapid Prototyping Aplikasi Manajemen Penugasan Karyawan. *Jurnal Cahaya Mandalika Issn 2721-4796 (Online)*, 4(1), 111–118.

# Emoji-Based Sentiment Classification Using Ensemble Learning with Cross-Validation: A Lightweight Approach for Social Media Analysis

Gunthur Bayu Wibisono<sup>1</sup>, Nur Alamsyah<sup>\*2</sup>, Titan Parama Yoga<sup>3</sup>, Budiman<sup>4</sup>, Acep Hendra<sup>5</sup>

<sup>1,2,3,4,5</sup>Universitas Informatika Dan Bisnis Indonesia, Indonesia

E-mail: <sup>1</sup>[gunthur.bw21@student.unibi.ac.id](mailto:gunthur.bw21@student.unibi.ac.id), <sup>\*2</sup>[nuralamsyah@unibi.ac.id](mailto:nuralamsyah@unibi.ac.id), <sup>3</sup>[titanparama@unibi.ac.id](mailto:titanparama@unibi.ac.id),  
<sup>4</sup>[budiman@unibi.ac.id](mailto:budiman@unibi.ac.id), <sup>5</sup>[acephendra@unibi.ac.id](mailto:acephendra@unibi.ac.id)

## Abstract

*The increasing use of emojis in online communication reflects emotional expression that is often more immediate and intuitive than text. This study proposes a lightweight sentiment classification approach that utilizes only emoji features extracted from social media posts, without relying on textual content. The importance of this research lies in its relevance to short-form digital content, where textual sentiment cues are minimal or absent. To address the classification problem, we implement and compare multiple machine learning models including Random Forest (RF), Support Vector Machine, and an ensemble Voting Classifier combining both. Emoji tokens were vectorized using character-level count vectorization, and performance was evaluated using 5-fold cross-validation to ensure robustness and generalizability. Results show that the ensemble model achieved the highest average accuracy of 93.6%, outperforming the individual classifiers. These findings confirm that emojis alone can serve as reliable indicators of sentiment and support the deployment of fast, interpretable, and scalable models for social media sentiment analysis.*

**Keywords**— *Emoji Sentiment Analysis, Ensemble Learning, Voting Classifier, Cross-Validation, Social Media*

Submission: May 16, 2025

Accepted: June 10, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.396>

## 1. BACKGROUND

The widespread adoption of social media has significantly changed the landscape of digital communication, encouraging users to express opinions, emotions, and reactions not only through text but also through symbolic elements like emojis [1]. Emojis have become essential components of online interactions, often acting as concise yet powerful indicators of sentiment in informal and short-form content such as tweets, captions, and instant messages [2]. This phenomenon creates both opportunities and challenges in sentiment analysis, particularly when textual content is minimal or absent [3].

Traditional sentiment analysis methods generally rely on lexical resources, syntactic features, and word embeddings extracted from textual data to determine polarity or emotion [4] [5]. These techniques have been effective in structured text but tend to perform poorly on informal content like social media posts [6]. The integration of emojis into sentiment models began to gain traction when Wang et al. (2025) [7] analyzed tweets and concluded that emojis are strong sentiment cues, often outperforming text-based features in noisy environments.

Further advancing this domain, Shen et al. (2025) [8] using DeepMoji, a neural network trained on billion tweets containing emojis to predict emotional tone. The model demonstrated

that emoji-based supervision could enhance downstream NLP tasks. Fouda et al. (2025) [9] also investigated the semantic behavior of emojis using distributional representations, showing that emojis can cluster based on context and sentiment similarity.

Several researchers have attempted to combine emoji features with text, but few studies have explored the use of emoji-only input for sentiment classification. While deep learning models such as LSTM or transformer-based architectures have yielded high performance, their computational cost and complexity hinder their use in lightweight or real-time systems [10].

More recent studies have experimented with emoji-specific vectorization techniques and feature engineering approaches. For instance, Kusal et al. (2025) [11] proposed Emoji2Vec, a method for generating emoji embeddings based on emoji descriptions, which enabled better generalization in NLP pipelines. Other works have explored the emotional and cultural interpretation of emojis, yet most still rely on combined text-emoji input rather than emojis in isolation [12] [13].

In contrast, lightweight approaches using traditional machine learning models such as Random Forest, SVM, or ensemble strategies remain underexplored in the context of emoji-only sentiment prediction. There is a lack of

systematic evaluation that compares these models while ensuring robustness through cross-validation [14].

This study proposes an emoji-based sentiment classification framework that uses only emoji characters as input features, without relying on any textual content. Three classifiers are evaluated: Random Forest, Support Vector Machine (SVM), and an ensemble Voting Classifier combining both [15]. Emoji sequences are encoded using character-level count vectorization, and model performance is assessed using 5-fold cross-validation to ensure robustness.

This research contributes to the development of fast, interpretable, and resource-efficient sentiment classification techniques suited for modern digital communication. It provides empirical evidence that emojis alone can reliably indicate sentiment, offering a practical solution for lightweight sentiment analysis in social media applications.

## 2. RESEARCH METHODOLOGY

The proposed method is designed to operate efficiently by leveraging only emoji characters extracted from user-generated social media content, without relying on textual information. The process includes data crawling, emoji preprocessing, feature transformation, model training using ensemble classifiers, and performance evaluation. Figure 1 illustrates the overall architecture of the proposed methodology. The process begins with data crawling from social media sources, followed by emoji extraction using rule-based filtering. The extracted emoji sequences are then transformed into numerical features using a character-level CountVectorizer. These features serve as input for two base classifiers: RF and SVM. Their predictions are combined using a Voting Classifier as the ensemble strategy. The model is evaluated using 5-fold cross-validation, and the final trained model is saved for deployment purposes.

### 2.1. Data Collection

The dataset used in this research was obtained from the Kaggle platform, which provides a collection of labeled tweets containing various emojis. A total of 3,085 tweets were collected, each representing an emotional expression labeled into one of four predefined

sentiment classes (0: sad, 1: happy, 2: angry, 3: love).

The primary objective of this study is to classify these emotions based solely on the emojis embedded in the tweets, without relying on any textual data. This makes the dataset highly relevant for evaluating the feasibility of emoji-based sentiment classification in informal digital communication contexts. Table 1 describes the features available in the dataset.

Table 1. Dataset Features Description

| Feature Name | Description                                                    | Data Type |
|--------------|----------------------------------------------------------------|-----------|
| Text         | The original tweet text containing emojis                      | String    |
| Sentiment    | Labeled emotion category (0: sad, 1: happy, 2: angry, 3: love) | Integer   |

### 2.2 Emoji Preprocessing

After the raw tweet data is collected, the next step involves preprocessing to isolate only the emojis used within each message. Since the objective of this study is to classify sentiment based exclusively on emojis—without any reliance on textual content—this preprocessing step is crucial in ensuring that only relevant visual-symbolic features are retained [16].

The emoji extraction process is performed using the emoji library in Python, which can identify and filter emoji characters from mixed-content text. Each tweet is scanned character by character, and only those identified as emoji tokens are extracted. This method eliminates all non-emoji characters including words, numbers, hashtags, and punctuation.

The result of this process is a new feature column that contains sequences of emojis representing the emotional context of the original tweet. Entries with no emojis are discarded to ensure that the learning model is trained solely on emoji-based inputs.

### 2.3 Feature Transformation

Once the emojis are extracted from each tweet, the next step involves transforming these sequences into numerical feature representations suitable for machine learning models [17]. This transformation is performed using CountVectorizer, a method that converts a collection of character-based emoji tokens into a matrix of token counts.

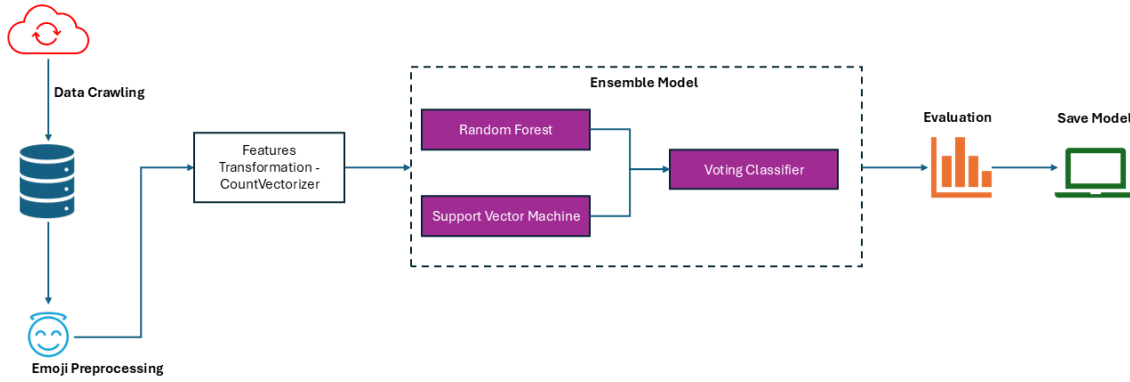


Figure 1. Proposed Method

In this study, character-level vectorization is applied, treating each unique emoji as a feature. Let the set of all unique emojis extracted from the dataset be denoted by:

$$E = \{e_1, e_2, e_3, \dots, e_k\} \quad (1)$$

where  $k$  is the total number of distinct emojis across the dataset. Each tweet  $T_i$  containing a sequence of emojis is represented as:

$$T_i = \{e_{i1}, e_{i2}, \dots, e_{im}\} \quad (2)$$

The CountVectorizer function transforms each tweet  $T_i$  into a feature vector  $x_i \in R^k$ , where each component  $x_{ij}$  represents the count of emoji  $e_j$  in tweet  $T_i$ :

$$x_{ij} = \text{count}(e_j \in T_i), \quad \forall j = 1, 2, \dots, k \quad (3)$$

The final output is a document-term matrix  $X \in R^{n \times k}$ , where  $n$  is the number of tweets (documents),  $k$  is the number of unique emojis (features),  $X_{ij}$  is the frequency of emoji  $e_j$  in tweet  $T_i$ . This vectorized representation captures the emoji usage patterns across tweets and serves as input to the classification models.

#### 2.4 Ensemble Model

To improve classification performance and generalizability, this study employs an ensemble learning approach using a Voting Classifier, which combines the predictions of multiple base classifiers. Specifically, we integrate two widely used algorithms: Random Forest (RF) and Support Vector Machine (SVM). The goal of this ensemble is to leverage the strengths of each model—Random Forest's robustness in handling high-dimensional data and SVM's ability to find optimal decision boundaries.

Algorithm 1: Emoji-Based Sentiment Classification Using Ensemble

```

Input : X (emoji feature matrix), y (sentiment labels)
Output : y_pred (final sentiment prediction)

1 Initialize base models:
2   model_1 ← RandomForestClassifier()
3   model_2 ← SupportVectorMachine(probability=True)
4 Train both models on training data:
5   model_1.fit(X_train, y_train)
6   model_2.fit(X_train, y_train)
7 Predict probabilities on test data:
8   prob_1 ← model_1.predict_proba(X_test)
9   prob_2 ← model_2.predict_proba(X_test)
10 Compute soft voting average:
11   prob_avg ← (prob_1 + prob_2) / 2
12 Assign final class based on max probability:
13   y_pred ← argmax(prob_avg, axis=1)
14 Return y_pred

```

Algorithm 1 illustrates the steps involved in constructing the ensemble model for emoji-based sentiment classification using soft voting. The algorithm begins by initializing two base classifiers—Random Forest and Support Vector Machine. Both models are independently trained using the same training dataset composed of emoji features and corresponding sentiment labels.

Once trained, each model is used to predict the probability distribution over the sentiment classes for the test data. These probability outputs are then averaged in a soft voting scheme, which computes the mean predicted probability for each class across both models. The final predicted class label for each instance is determined by selecting the class with the highest averaged probability.

This approach enables the ensemble model to leverage the strengths of both classifiers while mitigating their individual weaknesses. By averaging the prediction probabilities, the Voting Classifier can produce more balanced and robust predictions, especially when working with sparse or symbolic input data such as emojis.

The Voting Classifier operates by aggregating the predictions from both classifiers

and determining the final class label based on the majority (hard voting) or average predicted probability (soft voting). In this study, soft voting is used, where the class probabilities predicted by each base model are averaged, and the class with the highest average probability is selected.

Let  $C$  be the number of sentiment classes (in this case,  $C = 4$ ),  $M = \{m_1, m_2\}$  represent the set of base classifiers (RF and SVM),  $P_{m_i}^{(c)}(x)$  denote the probability predicted by model  $m_i$  for class  $c$ , given input feature vector  $x$ .

The ensemble model's prediction for class  $c$  is computed as:

$$P_{\text{ensemble}}^{(c)}(x) = \frac{1}{|M|} \sum_{i=1}^{|M|} P_{m_i}^{(c)}(x) \quad (4)$$

The final predicted class  $\hat{y}$  is given by:

$$\hat{y} = \arg \max_{c \in \{1, 2, \dots, C\}} P_{\text{ensemble}}^{(c)}(x) \quad (5)$$

This ensemble strategy enhances classification robustness by reducing variance and balancing out the biases of individual classifiers.

### 2.5 Evaluation

To ensure that the classification models generalize well to unseen data and are not overfitting to specific subsets, evaluation is performed using k-Fold Cross-Validation [18]. This technique provides a robust estimate of the model's performance by splitting the dataset into multiple train-test partitions.

In this study, we employ 5-Fold Cross-Validation ( $k = 5$ ). The dataset is divided into five equal-sized subsets (folds). During each iteration:

- Four folds are used to train the model.
- The remaining one fold is used to test the model.
- This process is repeated five times, such that each fold is used exactly once as the validation set.

The overall performance is computed as the average of accuracy scores across all folds. Let  $A_i$  denote the accuracy of the model on the  $i$ -th fold. Then, the mean accuracy is calculated as:

$$\text{Mean Accuracy} = \frac{1}{k} \sum_{i=1}^k A_i \quad (6)$$

and the standard deviation is computed to assess the model's stability:

$$\text{Standard Deviation} = \sqrt{\frac{1}{k} \sum_{i=1}^k (A_i - \bar{A})^2} \quad (7)$$

where  $\bar{A}$  is the mean accuracy across all folds.

This cross-validation approach is applied consistently across all models tested in this study: Random Forest, SVM, and the ensemble Voting Classifier. The results from this evaluation provide a reliable estimate of model performance and allow for a fair comparison between classifiers.

### 2.6 Save Model

After the model has been trained and evaluated, the final step involves saving the best-performing classifier for future use. In this study, the Voting Classifier—which demonstrated the highest average accuracy and stability during cross-validation—is selected as the final model.

Model persistence is essential for deployment in practical applications, enabling the reuse of trained models without the need to retrain them. The model is serialized using an appropriate object serialization library that supports machine learning models with large numerical structures.

Alongside the classifier, the vectorization object used to convert emoji characters into numerical features is also saved. This ensures consistency between the training phase and any future prediction tasks, as the same feature transformation process must be applied to new input data.

## 3. RESEARCH RESULT

This study evaluates the performance of three classification models—Random Forest (RF), Support Vector Machine (SVM), and an ensemble Voting Classifier (VC)—using emoji-only input features. Each model was assessed through 5-Fold Cross-Validation to ensure robust evaluation of generalizability and stability.

The Random Forest classifier yielded strong results, achieving an average accuracy of 93.33% across the five folds. However, it showed a relatively high standard deviation of 3.59%, indicating moderate variation in performance across different folds. This suggests that while RF performs well overall, it may be sensitive to data distribution in some folds, especially when class balance is not uniform.

Support Vector Machine, using a linear kernel, achieved slightly higher average accuracy at 93.36%, and notably exhibited greater consistency across folds, as indicated by its lower

standard deviation of 2.85%. The model's ability to generalize well from sparse emoji-based features demonstrates the effectiveness of margin-based classification even in non-textual data contexts.

The best performance was obtained from the ensemble model using the Voting Classifier. By combining the predictions from both Random Forest and SVM through soft voting, the ensemble achieved an average

accuracy of 93.63% with a standard deviation of 3.38%. This indicates that the ensemble approach not only improves predictive performance but also balances out the weaknesses of the individual models, offering more stable results across all validation folds.

A summary of these results is presented in Table 2. The table compares mean accuracy, standard deviation, best and lowest accuracy across all validation folds for each model.

Table 2. Comparison of Classification Performance

| Model          | Mean Accuracy | Standard Deviation | Best Fold Accuracy | Lowest Fold Accuracy |
|----------------|---------------|--------------------|--------------------|----------------------|
| Random Forest  | 0.9333        | 0.0359             | 0.9683             | 0.8696               |
| SVM (Linear)   | 0.9336        | 0.0285             | 0.9633             | 0.8830               |
| Ensemble Model | <b>0.9363</b> | 0.0338             | <b>0.9720</b>      | 0.8813               |

Figure 2 presents the cross-validation accuracy obtained from the ensemble model using Voting Classifier, which combines predictions from Random Forest and Support Vector Machine. The figure clearly illustrates the model's performance across five different data folds.

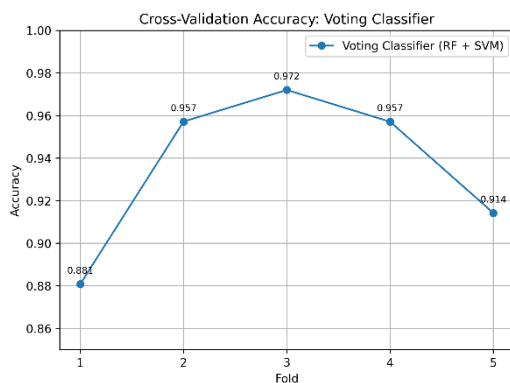


Figure 2 Cross-Validation Accuracy

From the chart, it can be observed that the accuracy ranges between 88.1% and 97.2%, with the highest performance achieved in fold 3 and the lowest in fold 1. The results indicate a generally high and consistent level of accuracy across the folds, reinforcing the model's robustness in classifying sentiment based solely on emoji input.

The slight variation in accuracy between folds reflects natural differences in data distribution and class representation within each subset, which is typical in cross-validation. Despite this, the ensemble model maintains stable performance, as evidenced by three folds yielding accuracies above 95%.

Overall, the visual result supports the numerical findings presented earlier, confirming

that the Voting Classifier delivers strong and reliable sentiment classification performance even when limited to non-textual, emoji-based features.

#### 4. DISCUSSION

The results obtained in this study demonstrate that sentiment classification based solely on emoji usage is not only feasible but can also achieve a high level of accuracy. All three tested models—Random Forest, SVM, and Voting Classifier—produced consistent and strong performance, with average accuracies exceeding 93%. This confirms that emojis, despite their simplicity and non-verbal nature, carry significant emotional cues that can be effectively leveraged for classification tasks.

Among the three models, the ensemble approach using a Voting Classifier yielded the best results. Its ability to combine the strengths of both Random Forest and SVM enabled it to achieve the highest average accuracy and strong stability across all validation folds. This outcome highlights the effectiveness of ensemble methods in improving predictive performance, especially in scenarios where the feature space is minimal and symbolic, such as emoji sequences.

The importance of these findings lies in the practical implication that sentiment analysis systems do not always require complex language processing pipelines. By focusing on a lightweight input source like emojis, it is possible to design models that are efficient in terms of computation while still maintaining high accuracy. This is particularly useful in environments with resource constraints or where input data is limited to short, informal messages such as those found on social media platforms.

Furthermore, the robustness of the Voting Classifier across various data folds implies that this approach can be generalized to other emoji-rich datasets or even extended to multilingual or cross-cultural contexts where textual analysis may encounter limitations. The results support the idea that visual-symbolic communication, which is increasingly prevalent in digital interactions, can be analyzed with the same rigor as traditional text-based sentiment analysis.

## 5. CONCLUSION

This study has successfully demonstrated that sentiment classification based solely on emojis can achieve high accuracy using conventional machine learning techniques. By transforming emoji characters into numerical features via character-level Count Vectorizer and applying classification models such as Random Forest, Support Vector Machine, and Voting Classifier, the system was able to perform reliably in identifying emotional categories within tweets.

The ensemble model using Voting Classifier outperformed individual models, achieving the highest average accuracy of 93.63% and showing stable performance across all validation folds. This highlights the strength of ensemble learning in combining complementary classifiers to enhance predictive capability, especially in domains with symbolic or non-verbal inputs.

One of the key advantages of the proposed method lies in its simplicity and efficiency. The system operates without the need for complex text preprocessing or deep learning architectures, making it suitable for real-time and resource-constrained applications. In addition, the focus on emojis as primary features addresses the growing need for models that can handle informal and visual modes of digital communication.

However, the approach also has limitations. The reliance solely on emoji input may not fully capture nuanced sentiments expressed in more complex messages, especially those involving sarcasm or mixed emotions. Furthermore, the model's performance may vary depending on the distribution and cultural interpretation of emojis in different datasets.

Future development may include integrating emoji features with minimal text input to enhance context understanding, experimenting with emoji embeddings or deep learning methods such as attention mechanisms, and testing the model across diverse social media

platforms and languages. Such enhancements could further improve accuracy and expand the applicability of emoji-based sentiment analysis in practical scenarios.

## 6. SUGGESTION

Based on the limitations identified in this study, several directions for future research are recommended. First, although the model performed well using only emojis as input, future studies may explore the integration of minimal text-based features to capture more nuanced sentiment, especially in cases where emojis are used ambiguously or contextually.

Second, the use of character-level Count Vectorizer could be enhanced by employing more sophisticated representation techniques, such as emoji embeddings or pre-trained vector spaces specifically designed for emoji semantics. This would allow models to learn more about the relationships between emojis beyond simple frequency counts.

Third, future research should consider evaluating the model on larger and more diverse datasets, particularly those containing multilingual content or region-specific emoji usage. This will help determine the generalizability of the proposed approach across different cultural and linguistic contexts.

Lastly, incorporating deep learning-based ensemble models, such as stacking or blending, could be investigated to assess whether more complex combinations of base classifiers lead to further improvements in accuracy and robustness.

These suggestions aim to address the current limitations of the study and provide a roadmap for enhancing the capability and applicability of emoji-based sentiment classification models in subsequent research.

## REFERENCES

- [1] Q. A. Xu, C. Jayne, and V. Chang, "An emoji feature-incorporated multi-view deep learning for explainable sentiment classification of social media reviews," *Technol. Forecast. Soc. Change*, vol. 202, p. 123326, 2024.
- [2] N. Alamsyah, A. P. Kurniati, and others, "Airfare Fluctuation Analysis with Event and Sentiment Features by Stacking Ensemble Model," in *2024 Ninth International Conference on Informatics and Computing (ICIC)*, IEEE, 2024, pp. 1–6.
- [3] W. Erpurini, A. G. Putrada, N. Alamsyah, S. F. Pane, and M. N. Fauzan,

- “Confirmatory Factor Analysis for The Impact of Students’ Social Medial on University Digital Marketing,” in *2023 International Conference on Computer Science, Information Technology and Engineering (ICCoSITE)*, IEEE, 2023, pp. 615–620.
- [4] J. Yu and C. Qi, “Machine Learning-Based Sentiment Analysis in English Literature: Using Deep Learning Models to Analyze Emotional and Thematic Content in Texts,” *IEEE Access*, 2025.
- [5] C. M. Liapis, A. Karanikola, and S. Kotsiantis, “Enhancing sentiment analysis with distributional emotion embeddings,” *Neurocomputing*, vol. 634, p. 129822, 2025.
- [6] R. Ahamad and K. N. Mishra, “Exploring sentiment analysis in handwritten and E-text documents using advanced machine learning techniques: a novel approach,” *J. Big Data*, vol. 12, no. 1, p. 11, 2025.
- [7] S. Wang, Q. Liu, Y. Hu, and H. Liu, “Public Opinion Evolution Based on the Two-Dimensional Theory of Emotion and Top2Vec-RoBERTa,” *Symmetry*, vol. 17, no. 2, p. 190, 2025.
- [8] Y. Shen *et al.*, “The DeepMoji algorithm for fast and accurate classification of massive books based on emotion encoding with redefined weights in multihead attention,” in *Third International Conference on Algorithms, Network, and Communication Technology (ICANCT 2024)*, SPIE, 2025, pp. 189–195.
- [9] W. Fouda, A. Hegazy, N. M. Alnaqbi, E. Ozbilge, and E. Özbilge, “Enhancing educational environments with Social Media Feedback Evaluation Employing Hybrid Neutrosophic Decision Optimization (HNDO) and Neutrosophic Sentiment Fusion (NSF).,” *Int. J. Neutrosophic Sci. IJNS*, vol. 26, no. 1, 2025.
- [10] N. Alamsyah, T. P. Yoga, B. Budiman, and others, “IMPROVING TRAFFIC DENSITY PREDICTION USING LSTM WITH PARAMETRIC ReLU (PReLU) ACTIVATION,” *JITK J. Ilmu Pengetah. Dan Teknol. Komput.*, vol. 9, no. 2, pp. 154–160, 2024.
- [11] S. Kusal, S. Patil, and K. Kotecha, “Multimodal text-emoji fusion using deep neural networks for text-based emotion detection in online communication,” *J. Big Data*, vol. 12, no. 1, pp. 1–25, 2025.
- [12] P. M. Hancock, C. Hilverman, S. W. Cook, and K. M. Halvorson, “Emoji as gesture in digital communication: Emoji improve comprehension of indirect speech,” *Psychon. Bull. Rev.*, vol. 31, no. 3, pp. 1335–1347, 2024.
- [13] A. G. Putrada, N. Alamsyah, and M. N. Fauzan, “BERT for sentiment analysis on rotten tomatoes reviews,” in *2023 International Conference on Data Science and Its Applications (ICoDSA)*, IEEE, 2023, pp. 111–116.
- [14] N. Alamsyah, A. P. Kurniati, and others, “Event Detection Optimization Through Stacking Ensemble and BERT Fine-tuning For Dynamic Pricing of Airline Tickets,” *IEEE Access*, 2024.
- [15] N. Alamsyah and others, “Analisis Perbandingan Sentimen Pengguna Twitter Terhadap Layanan Salah Satu Provider Internet Di Indonesia Menggunakan Metode Klasifikasi,” *TEMATIK*, vol. 10, no. 2, pp. 246–251, 2023.
- [16] N. Alamsyah, B. Budiman, R. Nursyanti, E. Setiana, and V. R. Danestiara, “Approximate Bayesian Inference for Bayesian Confidence Quantification in DNA Sequence Classification Using Monte Carlo Dropout Approach,” *Innov. Innov. Res. Inform.*, vol. 7, no. 1, 2025.
- [17] A. Khan, D. Majumdar, and B. Mondal, “Sentiment analysis of emoji fused reviews using machine learning and Bert,” *Sci. Rep.*, vol. 15, no. 1, p. 7538, 2025.
- [18] N. Alamsyah, V. R. Danestiara, B. Budiman, R. Nursyanti, E. Setiana, and A. Hendra, “OPTIMIZED FACEBOOK PROPHET FOR MPOX FORECASTING: ENHANCING PREDICTIVE ACCURACY WITH HYPERPARAMETER TUNING,” *J. Techno Nusa Mandiri*, vol. 22, no. 1, pp. 90–98, 2025.

## A Bidirectional GRU Approach with Hyperparameter Optimization for Sentiment Classification in Game Reviews

Alif Januantara Prima<sup>1</sup>, Nur Alamsyah<sup>\*2</sup>, Titan Parama Yoga<sup>3</sup>, Budiman<sup>4</sup>, Imannudin Akbar<sup>5</sup>,  
Acep Hendra<sup>6</sup>

<sup>1,2,3,4,5,6</sup>Universitas Informatika Dan Bisnis Indonesia, Indonesia

E-mail: <sup>1</sup>[alif.jp21@student.unibi.ac.id](mailto:alif.jp21@student.unibi.ac.id), <sup>\*2</sup>[nuralamsyah@unibi.ac.id](mailto:nuralamsyah@unibi.ac.id), <sup>3</sup>[titanparama@unibi.ac.id](mailto:titanparama@unibi.ac.id)

<sup>4</sup>[budiman@unibi.ac.id](mailto:budiman@unibi.ac.id), <sup>5</sup>[imannudin@unibi.ac.id](mailto:imannudin@unibi.ac.id), <sup>6</sup>[acephendra@unibi.ac.id](mailto:acephendra@unibi.ac.id)

### Abstract

*Sentiment analysis is essential for understanding user perspectives, particularly in contexts such as game reviews, where feedback directly influences product perception and user engagement. This study proposes a comparative approach employing Gated Recurrent Unit (GRU), hyperparameter-tuned GRU, and Bidirectional GRU models to classify sentiments within a game review dataset. The experimental pipeline begins with standard preprocessing and tokenization, followed by vectorization and supervised model training. Hyperparameter tuning is conducted using Keras Tuner to determine the optimal configuration of embedding dimensions, GRU units, dropout rates, and learning rates. The best-performing model—a Bidirectional GRU with optimized parameters—achieves a validation accuracy of 85.37% and demonstrates superior performance across key metrics, including precision, recall, and F1-score. Despite limitations posed by a relatively small and imbalanced dataset, the Bidirectional GRU exhibits strong generalization capability. Future work will explore class balancing techniques and the incorporation of pretrained word embeddings to further enhance model performance.*

**Keywords**—Sentiment Analysis, Game Reviews, Bidirectional GRU, Hyperparameter Optimization, Deep Learning

Submission: May 19, 2025

Accepted: June 08, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.399>

### 1. INTRODUCTION

Sentiment analysis is a fundamental task in Natural Language Processing (NLP) that seeks to extract subjective opinions from text and classify them as positive, negative, or neutral [1]. Within the domain of game reviews, sentiment analysis offers valuable insights into user satisfaction, gameplay experience, and overall market reception [2][3]. As digital gaming platforms continue to expand, the volume of user-generated textual feedback grows accordingly presenting both a significant opportunity and a challenge for extracting meaningful patterns through automated analysis [4].

One of the main challenges in sentiment classification of game reviews lies in the nature of the data itself: user reviews are often short, informal, and structurally diverse [5]. Moreover, these datasets are frequently imbalanced, with a disproportionate number of positive reviews compared to negative ones [6]. These characteristics necessitate the use of models capable of capturing contextual meaning,

managing noisy and brief text, and maintaining robust performance despite limited and skewed data distributions [7].

Early approaches to sentiment analysis primarily relied on traditional machine learning algorithms such as Support Vector Machines (SVM), Naive Bayes, and logistic regression [8]. These models typically required extensive feature engineering, utilizing techniques like term frequency-inverse document frequency (TF-IDF), bag-of-words, and n-gram modeling [9]. While effective in certain applications, such methods are inherently limited in their ability to capture sequential dependencies and semantic nuances within language [10].

Advancements in deep learning have shifted the focus toward recurrent neural networks (RNNs), particularly Long Short-Term Memory (LSTM) and Gated Recurrent Unit (GRU) architectures. GRU models, known for being more computationally efficient than LSTMs, have demonstrated strong performance in various sequence classification tasks [11]. Further enhancements, such as the Bidirectional GRU, enable the model to process text in both

forward and backward directions, thereby improving its ability to capture contextual information—a crucial factor in sentiment analysis [12].

Recent studies have validated the effectiveness of GRU and its variants in sentiment classification. For example, Kaur et al. (2024) [13] implemented a GRU-based architecture for Twitter sentiment analysis and reported performance gains over standard RNNs. Likewise, Liu et al. (2024) [14] showed that Bidirectional GRU outperformed unidirectional models in classifying short-text reviews, demonstrating its advantage in handling the brevity and variability typical of user-generated content.

Hyperparameter tuning has been widely recognized as a critical component in optimizing deep learning models. Alamsyah et al. (2024) [15] demonstrated that random search can outperform traditional grid search in improving model performance. More recently, frameworks such as Keras Tuner and Optuna have introduced automated tools that streamline the tuning process for key parameters, including learning rate, number of hidden units, embedding dimensions, and dropout rates—often yielding substantial performance gains [16]. However, to the best of our knowledge, no prior study has conducted a systematic comparison of GRU, hyperparameter-tuned GRU, and Bidirectional GRU within a single experimental framework applied to game review sentiment datasets. This gap in the literature forms the basis and motivation for the present study.

The primary contribution of this research is a comprehensive evaluation of GRU-based architectures for sentiment classification on game review data, with particular focus on the effects of hyperparameter optimization and bidirectional structure. Unlike previous studies that investigated GRU or LSTM in isolation, this work systematically compares three configurations: a baseline GRU model, a GRU model optimized using Keras Tuner, and a Bidirectional GRU model employing the best-tuned parameters. By applying automated hyperparameter tuning, the study demonstrates measurable improvements in classification metrics such as accuracy, precision, and F1-score. Moreover, the use of Bidirectional GRU enables the model to capture richer contextual information from both directions of the input sequence, enhancing its generalization capability—even in the presence of limited and imbalanced data. These findings offer a practical foundation for future sentiment analysis research involving short-text reviews and underscore the

importance of integrating architectural enhancements with parameter optimization in deep learning applications.

## **2. RESEARCH METHODOLOGY**

The research method employed in this study consists of several sequential stages, beginning with data collection and ending with model evaluation. The overall framework is illustrated in Figure 1, which outlines the complete pipeline for sentiment classification using Bidirectional GRU and its hyperparameter-optimized variant. The process begins with the collection of game review data, followed by a series of preprocessing steps including text cleaning, handling missing values, and converting labels into numeric format. The cleaned text is then transformed into sequences of integers through tokenization and subsequently padded to ensure uniform input length.

Subsequently, the dataset is divided into training and testing subsets, which serve as the basis for training two model configurations: a baseline Bidirectional GRU and a tuned Bidirectional GRU optimized through hyperparameter search. Both models are evaluated using key performance metrics, including accuracy, precision, recall, F1-score, and a confusion matrix. The objective of this methodological pipeline is to assess the extent to which hyperparameter tuning and bidirectional architecture contribute to improved accuracy, generalization, and robustness in sentiment classification of game review texts.

### ***2.1. Data Collected***

The dataset used in this study was sourced from the Kaggle platform, which contains labeled game review texts for sentiment classification tasks [17]. The dataset contains 2,000 labeled game reviews, with a class distribution of approximately 65% positive and 35% negative sentiments, indicating an imbalanced dataset. The average review length is around 23 words per entry, with substantial variation in style, length, and sentence structure. These characteristics make the dataset well-suited for evaluating model performance on short and informal user-generated text, which is typical in game review platforms.

Each data entry consists of a short user-written review, its cleaned version, and a sentiment label indicating whether the review is positive or negative. The dataset is particularly suitable for evaluating the performance of

sequence-based models due to the informal and varied nature of user-generated text. The features available in the dataset are described in Table 1.

Table 1. Description of Dataset Features

| Feature Name | Data Type | Description                                          |
|--------------|-----------|------------------------------------------------------|
| ReviewID     | String    | A unique identifier for each review entry.           |
| Review       | Text      | The original raw review text as written by the user. |

| Feature Name     | Data Type   | Description                                                                   |
|------------------|-------------|-------------------------------------------------------------------------------|
| Processed_Review | Text        | The cleaned and normalized version of the review, used as model input.        |
| Label            | Categorical | The sentiment category of the review, labeled as either Positive or Negative. |

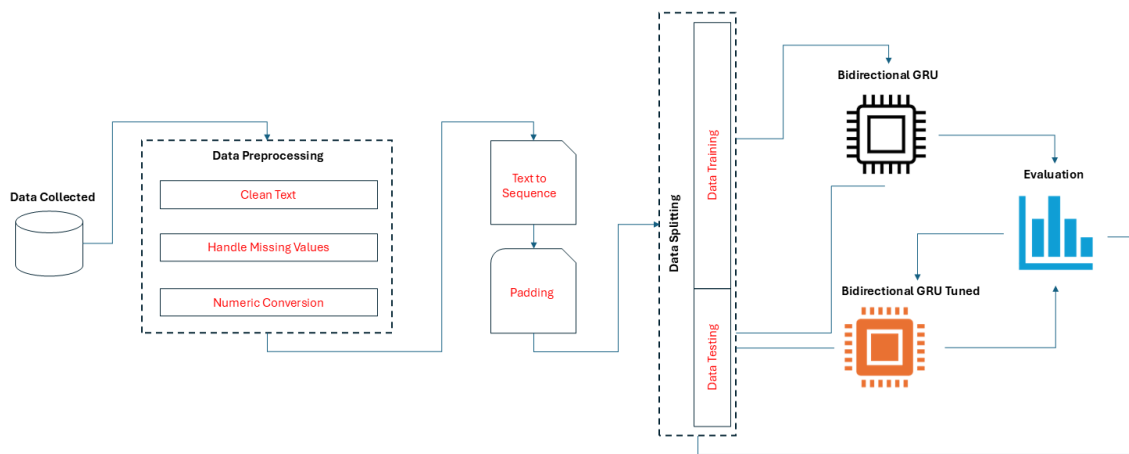


Figure 1. Pipeline For Sentiment Classification

## 2.2 Data Preprocessing

The data preprocessing stage is essential to prepare the raw textual data into a clean and structured format suitable for deep learning models [18]. As illustrated in Figure 1, this stage involves three main steps. The first step, clean text, transforms all characters in the review text to lowercase, removes punctuation, numbers, special characters, and excessive whitespace. This step helps eliminate noise that may hinder the model's learning process. The second step, handle missing values, involves removing any records that contain null or missing values in either the review content or the sentiment label.

This ensures the reliability of the dataset and avoids errors during the tokenization and training stages. The third step, numeric conversion, converts the sentiment labels from categorical format ("Positive" or "Negative") into binary numeric values (1 for positive, 0 for negative). This transformation allows the classification task to be modeled as a binary classification problem, aligning with the structure of the neural network's output layer. Overall, these preprocessing steps ensure that the

dataset is properly formatted for tokenization and subsequent training of the Bidirectional GRU models.

## 2.3 Text To Sequence And Padding

After the data has been cleaned and prepared, the next step is to convert the textual input into a numerical format that can be processed by neural networks [19]. This process begins with Text to Sequence, where each cleaned review is transformed into a sequence of integers using a tokenizer. The tokenizer assigns a unique integer to each word based on its frequency in the corpus, resulting in sequences of varying lengths depending on the number of words in each review.

Since neural network models require input data to be of uniform length, the second step, Padding, is applied to ensure consistency. Padding adds zeros to the end of shorter sequences (post-padding) so that all input sequences match the length of the longest sequence in the dataset. This step is crucial for batch processing during model training, as it allows efficient computation and avoids dimension mismatch errors. Together, the text-

to-sequence conversion and padding steps produce a well-structured, fixed-length input matrix that can be directly fed into the Bidirectional GRU models for training [20].

#### 2.4 Data Splitting

To evaluate the performance and generalization ability of the model, the dataset is divided into two subsets: training data and testing data. This process, known as *data splitting*, is a standard procedure in machine learning and deep learning workflows to prevent overfitting and to assess how well the model performs on unseen data.

In this study, the dataset is split using a ratio of 80:20, where 80% of the data is allocated for training and the remaining 20% for testing. To ensure better generalization, especially given the relatively small and imbalanced nature of the dataset, we adopted a 5-fold cross-validation strategy instead of a simple train-test split. This approach divides the data into five equal parts, using four parts for training and one for testing in each iteration, allowing every sample to be used for both training and validation. This method improves robustness and reduces the risk of overfitting. The training data is used to train the Bidirectional GRU and Bidirectional GRU Tuned models by allowing them to learn patterns and features from labeled examples. Meanwhile, the testing data is kept separate and is only used during the final evaluation phase to measure the model's accuracy, precision, recall, and F1-score on data that the model has never seen before. This approach ensures that the evaluation metrics reflect the model's true predictive capability and not just its ability to memorize the training data.

#### 2.5 Bidirectional GRU

The core model used in this study is the Bidirectional Gated Recurrent Unit (Bidirectional GRU), an enhanced variant of the GRU network that can capture contextual information from both past (previous words) and future (next words) in a sequence. Unlike unidirectional GRU that only processes sequences in a forward direction, Bidirectional GRU employs two GRU layers: one that reads the input from left to right, and another that reads it from right to left. Their outputs are concatenated, enabling the model to learn richer representations of the input data.

Each GRU unit uses gating mechanisms—reset gate and update gate—to control the flow of information. The mathematical formulation of a standard GRU cell is as follows:

$$z_t = \sigma(W_z \cdot x_t + U_z \cdot h_{t-1}) \quad \{(update\ gate)\} \quad (1)$$

$$r_t = \sigma(W_r \cdot x_t + U_r \cdot h_{t-1}) \quad \{(reset\ gate)\} \quad (2)$$

$$\tilde{h}_t = \tanh(W_h \cdot x_t + U_h \cdot (r_t \odot h_{t-1})) \quad \{(candidate\ hidden\ state)\} \quad (3)$$

$$h_t = (1 - z_t) \odot h_{t-1} + z_t \odot \tilde{h}_t \quad \{(final\ hidden\ state)\} \quad (4)$$

Where  $x_t$  is the input at time step  $t$ ,  $h_{t-1}$  is the previous hidden state,  $z_t$  is the update gate,  $r_t$  is the reset gate,  $\tilde{h}_t$  is the candidate activation,  $h_t$  is the new hidden state,  $\odot$  denotes element-wise multiplication,  $\sigma$  is the sigmoid activation function and  $\tanh$  is the hyperbolic tangent activation.

#### 2.6 Bidirectional GRU Tuned

To improve the predictive performance of the baseline Bidirectional GRU model, hyperparameter tuning was carried out using the Keras Tuner framework. This tuning process systematically explores the model's hyperparameter space to identify the best configuration that minimizes validation loss while maximizing generalization. The tuning approach adopted in this study uses random search across several key hyperparameters, including the embedding dimension, the number of GRU units, dropout rates, the number of neurons in the dense layer, and the learning rate used during optimization.

The best-performing configuration obtained from this process is summarized in Table 2.

Table 2. Best Hyperparameter Configuration for Bidirectional GRU

| Hyperparameter      | Value |
|---------------------|-------|
| Embedding Dimension | 64    |
| GRU Units           | 32    |
| Dropout Rate        | 0.3   |
| Dense Layer Units   | 128   |
| Learning Rate       | 0.01  |

This tuned configuration enables the model to learn more effectively by balancing

complexity and regularization. Mathematically, the hyperparameter tuning process can be represented as an optimization objective:

$$\theta^* = \arg \min_{\theta \in \Theta} \mathcal{L}_{vl}(f(x; \theta)) \quad (5)$$

where  $\theta$  represents the set of hyperparameters,  $\Theta$  denotes the hyperparameter search space,  $f(x; \theta)$  is the Bidirectional GRU model parameterized by  $\theta$ , and  $\mathcal{L}_{vl}$  is the validation loss, measured using binary cross-entropy. This optimization process allows the model to achieve a validation accuracy of 85.37%, outperforming both the standard GRU and untuned Bidirectional GRU models. These results confirm that hyperparameter tuning plays a crucial role in enhancing the model's ability to generalize to unseen data.

## 2.6 Evaluation

The performance of each model was assessed using a set of well-established evaluation metrics, namely accuracy, precision, recall, and F1-score. These metrics were calculated specifically for the positive sentiment class, as it represented the majority class within the dataset and served as a key focus of the analysis. Accuracy reflects the overall correctness of the model's predictions, while precision measures the proportion of positive predictions that were actually correct. Recall, on the other hand, indicates the model's ability to identify all actual positive instances. F1-score provides a harmonic balance between precision and recall, especially useful in cases where the class distribution is imbalanced.

To complement these metrics, confusion matrices were also generated to visualize the classification results in greater detail. Each confusion matrix illustrates the number of correctly and incorrectly predicted instances across both positive and negative classes. This enables a deeper understanding of how the model handles misclassifications, particularly in detecting minority class instances.

Based on the evaluation results, the Bidirectional GRU model with hyperparameter tuning consistently outperformed the baseline models across all evaluation criteria. Its improved performance confirms the positive impact of tuning and bidirectional context in enhancing model generalization and reliability in sentiment classification tasks.

## 3. RESULT

The experimental phase of this study began with the implementation of a baseline

GRU model, serving as the reference architecture for sentiment classification on game reviews. This model was trained using default settings, without any hyperparameter tuning. Although it achieved a respectable accuracy of 78%, it exhibited limitations in generalization—particularly when handling imbalanced class distributions.

To address this, the second experiment applied hyperparameter optimization to the GRU architecture. Using Keras Tuner, key parameters such as embedding size, number of GRU units, dropout rate, and learning rate were fine-tuned to enhance model performance. While the optimized GRU maintained the same overall accuracy (78%) as the baseline, it demonstrated improved training stability and consistency across evaluation metrics.

The third and final experiment focused on combining bidirectional sequence modeling with hyperparameter optimization. A Bidirectional GRU model was trained using the best-tuned configuration. This architecture enabled the model to capture both past and future context in the input sequence, leading to improved representational power. As a result, the Bidirectional GRU Tuned model achieved a higher accuracy of 83%, making it the best-performing approach in this study.

A comparison of evaluation results is presented in Table 3, while the accuracy comparison is visually illustrated in Figure 2.

Table 3 Comparison Of Accuracy

| Model                     | Accuracy | Precision | Recall | F1-Score |
|---------------------------|----------|-----------|--------|----------|
| GRU (Baseline)            | 0.78     | 0.78      | 1.00   | 0.88     |
| GRU (Tuned)               | 0.78     | 0.78      | 1.00   | 0.88     |
| Bidirectional GRU (Tuned) | 0.83     | 0.82      | 1.00   | 0.90     |

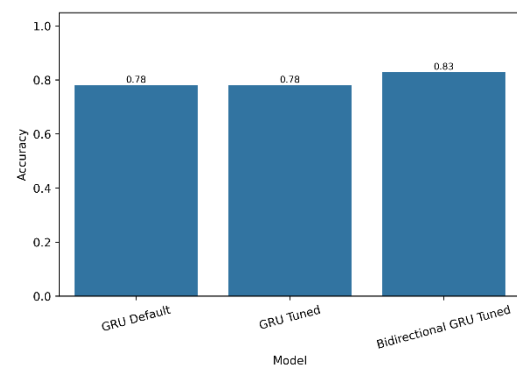


Figure 2. Accuracy Comparison

#### 4. DISCUSSION

The experimental results of this study reveal a clear progression in model performance as architectural enhancements and hyperparameter optimization were introduced. The baseline GRU model, despite being simple and efficient, exhibited limited capability in distinguishing between sentiment classes, particularly when handling imbalanced data. While it achieved a recall of 1.00—indicating all positive reviews were correctly identified—its relatively lower precision and F1-score reflected a tendency to overpredict the majority class.

Introducing hyperparameter tuning to the GRU model resulted in more stable training behavior and consistent metric scores. However, the improvement was marginal, suggesting that architectural limitations still constrained the model's representational power. The tuned GRU model maintained the same accuracy and recall as the baseline, but showed no significant increase in overall classification effectiveness.

The most significant improvement was observed when implementing the Bidirectional GRU with optimized hyperparameters. This configuration not only enhanced the accuracy to 83%, but also improved the precision and F1-score, indicating a better balance between sensitivity and specificity. The bidirectional architecture enabled the model to capture richer context by analyzing input sequences in both forward and backward directions—an advantage particularly beneficial for short and informal reviews typical of user-generated game content.

Furthermore, the use of hyperparameter tuning in combination with bidirectional modeling proved to be a key factor in achieving better generalization, as evidenced by the improved performance across all evaluation metrics. These findings confirm that careful tuning and architectural design can significantly influence the effectiveness of deep learning models in sentiment analysis tasks.

#### 5. CONCLUSION

This study has demonstrated a structured approach to improving sentiment classification performance on game review data using GRU-based models. Starting from a baseline GRU architecture, we observed that while the model could achieve high recall, it lacked precision and failed to generalize well under class imbalance. Incorporating hyperparameter tuning helped stabilize training and marginally improved model performance,

but the architectural limitations of unidirectional GRU remained a bottleneck.

The best results were obtained by implementing a Bidirectional GRU with tuned hyperparameters. This model achieved the highest accuracy and F1-score, indicating a more balanced and reliable classification capability. The bidirectional structure allowed the model to capture contextual information more effectively, while hyperparameter optimization contributed to improved learning efficiency and generalization. The overall findings affirm that model architecture and parameter tuning play a critical role in enhancing the effectiveness of deep learning approaches for sentiment analysis.

For future work, several directions can be explored to further strengthen the model. These include addressing class imbalance through techniques such as oversampling or cost-sensitive training, experimenting with pretrained word embeddings such as FastText or BERT-based representations, and incorporating attention mechanisms to highlight key parts of the review text. Additionally, expanding the dataset and applying domain adaptation methods may help improve model robustness across different genres of user-generated content. Such enhancements will contribute to more accurate and context-aware sentiment analysis systems, particularly in domains where informal and brief textual expressions are prevalent.

#### 6. SUGGESTION

Based on the findings of this study, several recommendations are proposed to guide future research in sentiment analysis of short-form review data. A primary area for enhancement is the management of class imbalance, which could be addressed through techniques such as oversampling, applying class-weighting strategies, or implementing data augmentation to increase the representation of minority classes. While the Bidirectional GRU model, optimized through hyperparameter tuning, demonstrated strong performance, future studies may benefit from incorporating attention mechanisms to enable the model to more effectively identify and prioritize contextually important words within a sentence.

Another promising avenue involves the integration of pretrained word embeddings—such as FastText, Word2Vec, or IndoBERT—which can enrich semantic understanding and improve classification accuracy, particularly for texts written in Bahasa Indonesia. Additionally, expanding the dataset to encompass a broader range of review types and user expressions could

enhance the model's generalizability. Finally, deploying the proposed model in real-world applications—such as automated feedback systems or sentiment monitoring tools in game development—would not only validate its practical utility but also provide opportunities for further refinement and adaptation.

#### REFERENSI

- [1] M. Ranjan, S. Tiwari, A. M. Sattar, and N. S. Tatkar, "A New Approach for Carrying Out Sentiment Analysis of Social Media Comments Using Natural Language Processing," *Eng. Proc.*, vol. 59, no. 1, p. 181, 2024.
- [2] X. Tong, I. Willcock, and Y. Sun, "Unravelling Player's Insights: A Comparative Analysis of Topic Modelling Techniques on Game Reviews and Video Game Developers' Perspectives," *IEEE Trans. Games*, 2024.
- [3] T. Guzvinecz and J. Sz\Hucs, "Textual Analysis of Virtual Reality Game Reviews.," *Infocommunications J.*, vol. 16, 2024.
- [4] W. Erpurini, A. G. Putrada, N. Alamsyah, S. F. Pane, and M. N. Fauzan, "Confirmatory Factor Analysis for The Impact of Students' Social Medial on University Digital Marketing," in *2023 International Conference on Computer Science, Information Technology and Engineering (ICCoSITE)*, IEEE, 2023, pp. 615–620.
- [5] L. Yue, W. Chen, X. Li, W. Zuo, and M. Yin, "A survey of sentiment analysis in social media," *Knowl. Inf. Syst.*, vol. 60, pp. 617–663, 2019.
- [6] L. Wang, M. Han, X. Li, N. Zhang, and H. Cheng, "Review of classification methods on unbalanced data sets," *Ieee Access*, vol. 9, pp. 64606–64628, 2021.
- [7] N. Alamsyah, A. P. Kurniati, and others, "Airfare Fluctuation Analysis with Event and Sentiment Features by Stacking Ensemble Model," in *2024 Ninth International Conference on Informatics and Computing (ICIC)*, IEEE, 2024, pp. 1–6.
- [8] T. V. Subbamma, N. Mantravadi, A. R. Shalom, N. Tajne, G. Sreenivasulu, and V. S. Goud, "Natural Language Processing for Sentiment Analysis Techniques and Applications," *Metall. Mater. Eng.*, vol. 31, no. 3, pp. 194–200, 2025.
- [9] A. G. Putrada, N. Alamsyah, I. D. Oktaviani, and M. N. Fauzan, "A hybrid genetic algorithm-random forest regression method for optimum driver selection in online food delivery," *J. Ilm. Tek. Elektro Komput. Dan Inform. JITEKI*, vol. 9, no. 4, pp. 1060–1079, 2023.
- [10] A. A. Kumar, "Semantic memory: A review of methods, models, and current challenges," *Psychon. Bull. Rev.*, vol. 28, no. 1, pp. 40–80, 2021.
- [11] N. Alamsyah, T. P. Yoga, B. Budiman, and others, "IMPROVING TRAFFIC DENSITY PREDICTION USING LSTM WITH PARAMETRIC ReLU (PReLU) ACTIVATION," *JITK J. Ilmu Pengetah. Dan Teknol. Komput.*, vol. 9, no. 2, pp. 154–160, 2024.
- [12] J. V. Tembhurne and T. Diwan, "Sentiment analysis in textual, visual and multimodal inputs using recurrent neural networks," *Multimed. Tools Appl.*, vol. 80, no. 5, pp. 6871–6910, 2021.
- [13] A. Kaur, V. Kukreja, and D. Kundra, "Echoes of Emotion: Enhancing Sentiment Analysis with CNN-GRU Synergy," in *2024 International Conference on Big Data Analytics in Bioinformatics (DABCon)*, IEEE, 2024, pp. 1–7.
- [14] J. Liu, Y. Sun, Y. Zhang, and C. Lu, "Research on Online Review Information Classification Based on Multimodal Deep Learning," *Appl. Sci.*, vol. 14, no. 9, p. 3801, 2024.
- [15] N. Alamsyah, B. Budiman, T. P. Yoga, and R. Y. R. Alamsyah, "COMPARISON LINEAR REGRESSION AND RANDOM FOREST MODELS FOR PREDICTION OF UNDERGROUND DROUGHT LEVELS IN FOREST FIRES," *J. Techno Nusa Mandiri*, vol. 21, no. 2, pp. 81–86, 2024.
- [16] J. B. L. Bernardo, A. Taparugssanagorn, H. Miyazaki, B. M. Pati, and U. Thapa, "Robust Human Activity Recognition for Intelligent Transportation Systems Using Smartphone Sensors: A Position-Independent Approach," *Appl. Sci.*, vol. 14, no. 22, p. 10461, 2024.
- [17] trilism, "EXCEL SENTIMENT FOR GAME REVIEW." 2025.
- [18] N. Alamsyah, A. P. Kurniati, and others, "A novel airfare dataset to predict travel agent profits based on dynamic pricing," in *2023 11th International Conference on Information and Communication Technology (ICoICT)*, IEEE, 2023, pp. 575–581.

- [19] N. Alamsyah, A. P. Kurniati, and others, "Event Detection Optimization Through Stacking Ensemble and BERT Fine-tuning For Dynamic Pricing of Airline Tickets," *IEEE Access*, 2024.
- [20] I. S. A. Khaleq, "Enhancing Clickbait Detection Through Deep Learning and Language-specific Analysis in English and Kurdish," PhD Thesis, Polytechnic University, 2023.

## The Implementation of the TOPSIS Method in Determining Stunting Toddlers

Umi Khultsum<sup>\*1</sup>, F. Lia Dwi Cahyanti<sup>2</sup>, Elly Firasari<sup>3</sup>

<sup>1</sup>Universitas Bina Sarana Informatika, Indonesia

<sup>2,3</sup>Universitas Nusa Mandiri, Indonesia

E-mail: <sup>\*1</sup>[umikhultsum.ukm@bsi.ac.id](mailto:umikhultsum.ukm@bsi.ac.id), <sup>2</sup>[flia.fdc@nusamandiri.ac.id](mailto:flia.fdc@nusamandiri.ac.id), <sup>3</sup>[elly.ef@nusamandiri.ac.id](mailto:elly.ef@nusamandiri.ac.id)

### Abstract

*Stunting is a growth disorder in children due to chronic malnutrition and recurrent infections, especially in the first 1,000 days of life. Assessment of stunting status that only relies on height and weight measurements is considered ineffective because it does not cover all aspects that affect a child's nutritional status. At Posyandu Bougenville, stunting identification is still done manually and it is at risk of causing errors in decision making. This study aims to develop a web-based Decision Support System (DSS) using the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS) method to assist Posyandu cadres in determining toddler stunting status quickly, accurately, and efficiently. This system processes data from four main anthropometric indicators, namely Height/Age, Weight/Age, Weight/Age, Weight/Age, and BMI/Age. The results of the system calculations show agreement with manual calculations, which proves that the system is working optimally. An example of the results shows that toddlers with code A1 (Rafasya Malik) have the lowest preference value of 0, followed by A4 (Ihsan Dwi Hanggoro) with a value of 0.4022149 which is included in the high stunting risk category. This system has proven to be able to help Posyandu cadres in prioritizing the handling of at-risk toddlers, as well as supporting the stunting monitoring process in a more structured and data-based manner.*

**Keywords**— Stunting, Decision Support System, TOPSIS, Toddlers

Submission: May 20, 2025

Accepted: June 10, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.400>

### 1. INTRODUCTION

Lack of adequate nutritional intake for children during the first 1,000 days of life starting from pregnancy until two years old is an indicator that the child is experiencing stunting. According to the World Health Organization (WHO), stunting is a condition of stunted growth and development in children caused by chronic malnutrition and repeated infections [1]. Nutritional status assessment is based on the Body Length for Age (PB/A) or Body Height for Age (TB/A) index, which is included in the standard for assessing children's nutritional status. The measurement results are categorized as short (stunted) if the Z-score value is between less than -2 SD to -3 SD, and categorized as very short (severely stunted) if the Z-score value is less than -3 SD [2]. As they get older, children who experience stunting are at higher risk of developing various other diseases, such as diabetes, heart disease, stroke, and even cancer [3].

Identification of stunting is part of an effort to improve the health of toddlers. At Posyandu Bougenville, determining whether a toddler is stunted is often only based on measuring height and weight. This approach is

less effective and has the potential to cause errors in determining stunting status in toddlers. Efforts to monitor and analyze stunting data are often hampered by the limited availability of effective decision-making tools. In addition, because stunting data involves various criteria, such as four anthropometric indicators, namely height for age (TB/A), weight for age (BB/A), weight for height (BB/TB), and body mass index for age (BMI/A), an analysis method is needed that can handle the complexity of the data as a whole.

Decision Support System (DSS) is a system designed to assist decision makers at the managerial level, especially in semi-structured decision-making conditions. This system can be interpreted as a process of selecting the best alternative from a number of available choices systematically, which is used as a method in solving a problem [4]. Therefore, the use of DSS can help reduce excessive use of time, energy, and costs. DSS aims to support a more effective decision-making process by choosing the right and efficient steps in carrying out these steps optimally [5]. Therefore, this study applies a decision support system that can be used by Posyandu Bougenville cadres to facilitate the determination of stunted toddlers. Research conducted by Ayu and Reni on malnutrition

using the SAW method resulted in a decision support system that is able to classify the nutritional status of toddlers, but does not yet include identification or determination of stunting conditions in toddlers [6].

The web-based decision support system that will be developed in this study is designed by utilizing the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS) method. TOPSIS is a multi-attribute decision-making method that is based on finding ideal and anti-ideal solutions and comparing the distance of each alternative with the alternative [7]. TOPSIS is the best method for decision-making cases compared to other methods [8].

The purpose of this study is to develop a web-based decision support system (DSS) that can assist in determining stunting status in toddlers quickly, accurately, and efficiently. This system is designed by implementing the TOPSIS (Technique for Order Preference by Similarity to Ideal Solution) method which is able to process data based on several anthropometric criteria. With this approach, it is expected that the decision-making process will be more systematic and objective. In addition, this system also aims to increase the effectiveness and efficiency of identifying toddlers who are at risk or have experienced stunting, as well as providing support for Posyandu Bougenvile cadres in carrying out appropriate monitoring and preventive measures. Through an easy-to-use web interface, this system can be accessed widely and in real-time, so it is expected to support government or health agency efforts in mapping and handling stunting problems in a more integrated and data-based manner.

## 2. RESEARCH METHODS

This section explains the research methods used in this study. The research was conducted through several stages, namely data and criteria analysis, calculations using the TOPSIS method, calculation results, and system development. In the process of creating the system, the waterfall software development method was used. The waterfall method is an approach in software development that allows the system development process to be carried out in a structured and sequential manner [9]. This method has several stages, namely, requirement analysis, system design, implementation, testing, and maintenance. The following stages of the research method can be seen in Figure 1.

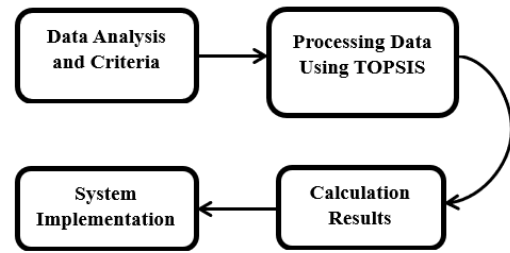


Figure 1. Methods

### 2.1. Data Analysis and Criteria

Data analysis is a process carried out systematically with the aim of processing and changing raw or unstructured data into meaningful information [10]. In this study, data was collected from Posyandu Bougenvile, which includes important information such as age, weight and height, immunization status, and diet. Then, the criteria that are the basis for decision making are determined, such as height according to age, weight according to age, weight according to height, and body mass index according to age. Each criterion is weighted according to its level of influence on the risk of stunting.

### 2.2. TOPSIS (Technique for Order Preference by Similarity to Ideal Solution)

In this study, the TOPSIS method was used for data processing. The TOPSIS method is based on the principle that the best alternative is the one that has the closest distance to the positive ideal solution and at the same time the furthest distance from the negative ideal solution [11]. The following are the stages of calculating the TOPSIS method, namely [11]:

1. Creating a normalized decision matrix

$$r_{ij} = \frac{x_{ij}}{\sqrt{\sum_{i=1}^m x_{ij}^2}} \quad (1)$$

Where  $i = 1, 2, \dots, m$ ; and  $j = 1, 2, \dots, n$ .

$r_{ij}$  = normalized decision matrix.

$x_{ij}$  = weight of criterion  $j$  on alternative  $i$ .

$i$  = alternative  $i$ .

$j$  = alternative  $j$ .

2. Create a weighted normalized decision matrix

$$y_{ij} = w_j r_{ij} \quad (2)$$

$i = 1, 2, \dots, m$ ; and  $j = 1, 2, \dots, n$ .

Where  $w_j$  is the weight of criterion  $j$ .

3. Determine the positive ideal solution matrix and negative ideal solution matrix.

$$A^+ = (y_1^+, y_2^+, \dots, y_n^+); \quad (3)$$

$$A^- = (y_1^-, y_2^-, \dots, y_n^-); \quad (4)$$

Where

$$y_i^+ = \{y_{ij}; \text{ if } j \text{ is a benefit attribute } y_{ij};$$

$$\text{if } j \text{ is a cost attribute}$$

$$y_j^- = \{y_{ij}; \text{ if } j \text{ is a benefit attribute } y_{ij};$$

$$\text{if } j \text{ is a cost attribute}$$

4. Determine the distance between the value of each alternative with the positive ideal solution matrix and the negative ideal solution matrix.

positive ideal

$$D_i^+ = \sqrt{\sum_{j=1}^n (y_1^+ - y_{ij})^2}; \quad (5)$$

negative ideal

$$D_i^- = \sqrt{\sum_{j=1}^n (y_{ij} - y_1^-)^2} \quad (6)$$

5. Determine the preference value for each alternative.

$$V_i = \frac{D_i^-}{D_i^- + D_i^+} \quad (7)$$

$V_i$  is the preference value that shows the value of the  $i$  alternative. after getting the value  $V_i$ , then the alternatives will be ranked based on the order of value  $V_i$ .

### 2.3. Calculation Result

In this study, the calculation results stage of the TOPSIS (Technique for Order Preference by Similarity to Ideal Solution) method which will be implemented into a web-based Decision Support System (DSS) aims to determine the ranking or priority of a number of alternatives based on their proximity to the ideal solution.

### 2.4. Implementation System

#### 1. Requirement Analysis

This process involves system analysts and business analysts in identifying the functional and non-functional requirements of the system to be developed [12]. Researchers conducted an analysis of the system requirements that became the basis for developing a web-based Decision Support System to determine stunted toddlers at the Bougenvile Posyandu.

#### 2. System Design

This process involves creating a picture of the system to be developed to determine a solution plan that includes software design, architectural design, and database [13]. Researchers prepare a system design or draft, which includes database design and system model creation using UML (Unified Modeling

Language). In database design, researchers use ERD (Entity Relationship Diagram), ERD is a database design diagram that describes entities as potential tables and relationships as relationships between these entities [14]. ERD is a tool used to visualize the logical structure of a database by displaying entities, attributes, and relationships between entities [15]. UML (Unified Modeling Language) is one of the tools used in developing object-based systems [16]. The UML used is Use Case Diagram and Activity Diagram. A use case is a series or description of a group of activities that are interconnected and form a structured system, which is carried out or monitored by an actor [16]. Activity diagram is a modeling that is done on a system to describe a series of activities that occur in the system. This diagram is used to explain program activities visually without requiring code or interface displays [17].

#### 3. Implementation

This stage is the process of writing program code compiled into an application that can be run, including creating the necessary database and text files [12]. Researchers implement the design that has been made into a program code using the PHP (Hypertext Preprocessor) programming language. In addition, this step also includes the application of the TOPSIS method which is realized in the form of a system. PHP is a server-side programming language for developing various types of websites and web applications [18]. PHP allows the creation of dynamic web applications, meaning it can generate web pages that are influenced by data. This means that changes in data will automatically update the web page without the need to change the script or code that builds it [19].

#### 4. Testing

This process is known as verification and validation, which is the process of ensuring that the software solution meets the requirements and successfully achieves the desired goal [12]. At this stage, researchers carry out testing of the software that has been developed using the black box testing method. Black Box Testing focuses on testing system functions by assessing input and output without viewing or analyzing the source code [20]. Black box testing plays an important role in the software testing process, namely to ensure that all system functions run well and are dynamic [21].

#### 5. Maintenance

The Maintenance process includes fixing errors that were not identified in previous stages [12]. This stage aims to ensure that the system continues to run well and meets user needs over time [22]. This stage is carried out if there are

updates or modifications to existing software, but not to develop completely new software.

### 3. RESEARCH RESULTS

#### 3.1. TOPSIS method calculation

1. Create Alternative Data and Determine Criteria

Alternative data is taken from toddler data at the Bougenvile Posyandu, the following is a sample of the toddler data used:

| Code | Name                  |
|------|-----------------------|
| A1   | Rafasya Malik         |
| A2   | Fando Kurniawan       |
| A3   | Cakra Syahputra       |
| A4   | Ihsan Dwi Hanggoro    |
| A5   | Kelvin Permana        |
| A6   | Mira Putriasari       |
| A7   | Mutiara Cantika Putri |
| A8   | Pelita Pertiwi        |
| A9   | Shaka Putra Atmaja    |

Table 1. Alternative Data

The criteria data are the child's anthropometric standards, the following criteria data are used:

| Cod e | Criteria                               | Description    | Attribut e | Weigh t |
|-------|----------------------------------------|----------------|------------|---------|
| C01   | Height According to Age (H/A)          | Tall           | benefit    | 4       |
|       |                                        | Normal         |            | 3       |
|       |                                        | Short          |            | 2       |
|       |                                        | Very Short     |            | 1       |
| C02   | Body Weight According to Age (BW/A)    | Fat            | benefit    | 4       |
|       |                                        | Normal         |            | 3       |
|       |                                        | Thin           |            | 2       |
|       |                                        | Very Thin      |            | 1       |
| C03   | Body Weight According to Height (BW/H) | Over-nutrition | benefit    | 4       |
|       |                                        | Normal         |            | 3       |
|       |                                        | Undernutrition |            | 2       |
|       |                                        | Malnutrition   |            | 1       |
| C04   | Body Mass Index by Age (BMI/A)         | Over-nutrition | benefit    | 4       |
|       |                                        | Normal         |            | 3       |
|       |                                        | Undernutrition |            | 2       |
|       |                                        | Malnutrition   |            | 1       |

Table 2. Criteria Data

In the table above there is a weight, the weight here is the limit of the criteria points that can be given, the maximum limit that can be given is 4. In the benefit attribute column, it means that the greater the points entered, the better the toddler's condition.

2. Input Alternative Values

Based on the weights and criteria determined above, here is a sample of the weighting given to each toddler:

| alternative | C01 | C02 | C03 | C04 |
|-------------|-----|-----|-----|-----|
| A1          | 1   | 1   | 1   | 1   |
| A2          | 3   | 3   | 3   | 3   |
| A3          | 4   | 3   | 3   | 3   |
| A4          | 2   | 2   | 2   | 2   |
| A5          | 3   | 3   | 3   | 3   |
| A6          | 3   | 3   | 3   | 3   |
| A7          | 3   | 3   | 3   | 3   |
| A8          | 3   | 3   | 3   | 3   |
| A9          | 3   | 3   | 3   | 3   |

Table 3. Alternative Value Data

3. Normalized Decision Matrix

The first step is to square each element of the matrix in table 3. Here are the results:

| alternative | C01 | C02 | C03 | C04 |
|-------------|-----|-----|-----|-----|
| A1          | 1   | 1   | 1   | 1   |
| A2          | 9   | 9   | 9   | 9   |
| A3          | 16  | 9   | 9   | 9   |
| A4          | 4   | 4   | 4   | 4   |
| A5          | 9   | 9   | 9   | 9   |
| A6          | 9   | 9   | 9   | 9   |
| A7          | 9   | 9   | 9   | 9   |
| A8          | 9   | 9   | 9   | 9   |
| A9          | 9   | 9   | 9   | 9   |
| Total       | 75  | 68  | 68  | 68  |

Table 4. Square Result

Then normalize by dividing each matrix element in table 3 by the root of the total of each criterion in table 4. Sample calculation alternative A1, C01.

$$r_{11} = \frac{1}{\sqrt{75}} = \frac{1}{8,6602} = 0,11547$$

1 is obtained from table 3 and 75 is obtained from table 4.

The following are the results of the normalization:

| alternati ve | C01     | C02      | C03      | C04        |
|--------------|---------|----------|----------|------------|
| A1           | 0,11547 | 0,121268 | 0,121268 | 0,12126781 |
| A2           | 0,34641 | 0,363803 | 0,363803 | 0,36380344 |
| A3           | 0,46188 | 0,363803 | 0,363803 | 0,36380344 |
| A4           | 0,23094 | 0,242536 | 0,242536 | 0,24253563 |
| A5           | 0,34641 | 0,363803 | 0,363803 | 0,36380344 |
| A6           | 0,34641 | 0,363803 | 0,363803 | 0,36380344 |
| A7           | 0,34641 | 0,363803 | 0,363803 | 0,36380344 |
| A8           | 0,34641 | 0,363803 | 0,363803 | 0,36380344 |

|    |             |              |              |                |
|----|-------------|--------------|--------------|----------------|
| A9 | 0,3464<br>1 | 0,3638<br>03 | 0,3638<br>03 | 0,363803<br>44 |
|----|-------------|--------------|--------------|----------------|

Table 5. Normalization Results

#### 4. Weighted Normalization

Weighted normalization is obtained by multiplying the normalization matrix in table 5 with the maximum limit of the criteria weights contained in table 2. Sample calculation alternative A1, C01.

$$r_{11} = 0,11547 \times 4 = 0,46188$$

0,11547 is obtained from table 5 and 4 is obtained from table 2 (the maximum limit of the criteria weights).

The following matrix is obtained:

| alternati<br>ve | C01          | C02          | C03            | C04            |
|-----------------|--------------|--------------|----------------|----------------|
| A1              | 0,4618<br>8  | 0,4850<br>71 | 0,485071<br>25 | 0,485071<br>25 |
| A2              | 1,3856<br>41 | 1,4552<br>14 | 1,455213<br>75 | 1,455213<br>75 |
| A3              | 1,8475<br>21 | 1,4552<br>14 | 1,455213<br>75 | 1,455213<br>75 |
| A4              | 0,9237<br>6  | 0,9701<br>43 | 0,970142<br>5  | 0,970142<br>5  |
| A5              | 1,3856<br>41 | 1,4552<br>14 | 1,455213<br>75 | 1,455213<br>75 |
| A6              | 1,3856<br>41 | 1,4552<br>14 | 1,455213<br>75 | 1,455213<br>75 |
| A7              | 1,3856<br>41 | 1,4552<br>14 | 1,455213<br>75 | 1,455213<br>75 |
| A8              | 1,3856<br>41 | 1,4552<br>14 | 1,455213<br>75 | 1,455213<br>75 |
| A9              | 1,3856<br>41 | 1,4552<br>14 | 1,455213<br>75 | 1,455213<br>75 |

Table 6. Weighted Normalization

#### 5. Ideal solution matrix

The ideal solution matrix is obtained based on weighted normalization and criterion attributes (benefits). The criterion attribute used in this study is benefit, so the positive ideal solution is taken from the maximum value of the weighted normalization in table 6. Conversely, the negative ideal solution is taken from the minimum value of the weighted normalization in table 6. The results are as follows:

| Ideal<br>Solutio<br>n | C01          | C02          | C03            | C04            |
|-----------------------|--------------|--------------|----------------|----------------|
| Positiv<br>e          | 1,8475<br>21 | 1,4552<br>14 | 1,455213<br>75 | 1,455213<br>75 |
| Negati<br>ve          | 0,4618<br>8  | 0,4850<br>71 | 0,485071<br>25 | 0,485071<br>25 |

Table 7. Ideal Solution

#### 6. Distance Between Alternative Weighted Values and Ideal Solution

Squaring the difference of each weighted normalization matrix in table 6 with the ideal solution matrix in table 7, then summing each

alternative, then rooting it. The distance of the positive ideal solution uses the positive ideal solution, the distance of the negative ideal solution uses the negative ideal solution. Preferences are obtained by dividing the negative ideal by the sum of the positive and negative ideals.

Positive A1

$$\sqrt{\frac{(((0,46188 - 1,847521)^2) + ((0,485071 - 1,455214)^2) + ((0,485071 - 1,455214)^2) + ((0,485071 - 1,455214)^2))}{2}} = 2,177965$$

Negative A1

$$\sqrt{\frac{(((0,46188 - 0,46188)^2) + ((0,485071 - 0,485071)^2) + ((0,485071 - 0,485071)^2) + ((0,485071 - 0,485071)^2))}{2}} = 0$$

Preferences A1

$$\frac{0}{2,177965 - 0} = 0$$

The results are as follows:

| alternati<br>ve | Positive     | Negative      | preferen<br>ce | Result<br>s |
|-----------------|--------------|---------------|----------------|-------------|
| A1              | 2,17796<br>5 | 0             | 0              | Stunte<br>d |
| A2              | 0,46188      | 1,91751<br>47 | 0,80588<br>33  | Norm<br>al  |
| A3              | 0            | 2,17796<br>5  | 1              | Norm<br>al  |
| A4              | 1,24868<br>6 | 0,95875<br>7  | 0,43432<br>93  | Stunte<br>d |
| A5              | 0,46188      | 1,91751<br>47 | 0,80588<br>33  | Norm<br>al  |
| A6              | 0,46188      | 1,91751<br>47 | 0,80588<br>33  | Norm<br>al  |
| A7              | 0,46188      | 1,91751<br>47 | 0,80588<br>33  | Norm<br>al  |
| A8              | 0,46188      | 1,91751<br>47 | 0,80588<br>33  | Norm<br>al  |
| A9              | 0,46188      | 1,91751<br>47 | 0,80588<br>33  | Norm<br>al  |

Table. 8 Calculation Results

Based on the results of the preference values, 2 toddlers were declared stunted from 9 toddler data samples, namely alternative A1 (Rafasya Malik) and alternative A4 (Ihsan Dwi Hanggoro). in the table above, toddlers who are declared stunted are toddlers who have the smallest preference value, besides that, toddlers are declared in normal condition. this will be a rule in the decision support system for determining stunted toddlers.

### 3.2. System Design

#### 1. ERD

The following is the database design in this study, namely ERD. There are 4 entities, namely Criteria, Toddler, User and Stunting which are shown in Figure 2:

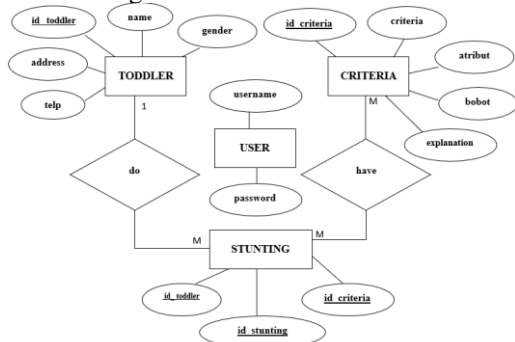


Figure 2. Entity Relationship Diagram

## 2. Use Case Diagram

The following is a Use Case Diagram for a decision support system for determining stunted toddlers, the actor in this Use Case is the admin.

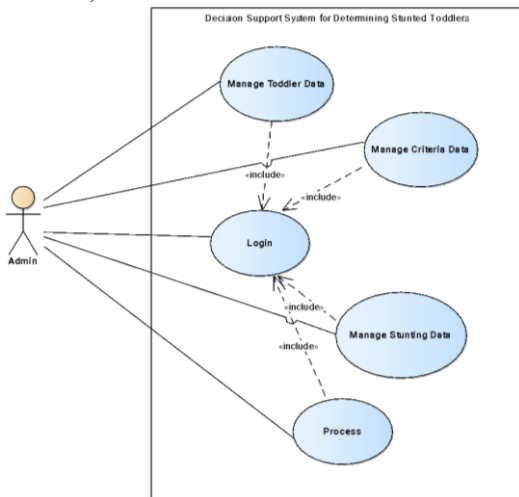


Figure 3. Use Case Diagram

At this stage, the admin first logs in using a username and password. After logging in, the admin can access and manage all menus in the system, including Manage Toddler Data, Manage Criteria Data, Manage Stunting Data and Manage Process Menu.

Activity Diagrams are used to model systems related to Login activities, Toddler Data, Stunting Data, Stunting Process and Criteria Data. For example, a toddler activity diagram, in the toddler activity diagram the admin can manage toddler data which includes adding, deleting, and changing toddler data. The following activity diagram for Managing Toddler Data is shown in Figure 4.

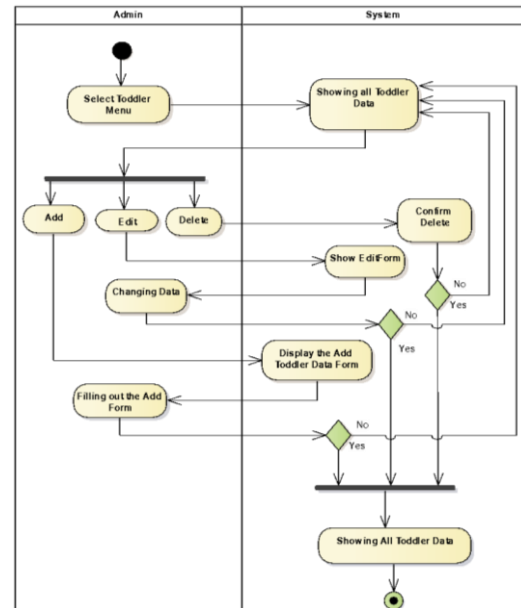


Figure 4. Activity Diagram Toddler Data

## 3.3 Implementation

### 1. Home Page View



Figure 5. Home Page View

In Figure 5 it is explained that after the admin logs in using the correct username and password, the home page will appear. On the home page there is a Stunting menu, this menu is used to see the results of determining stunted toddlers. Then the Toddler menu, this menu is used to manage toddler data. The Criteria menu is used to manage criteria data, and the Stunting Data menu is used to input toddler data based on existing criteria. The menu can be accessed and managed by the admin.

### 2. Toddler Menu Page View

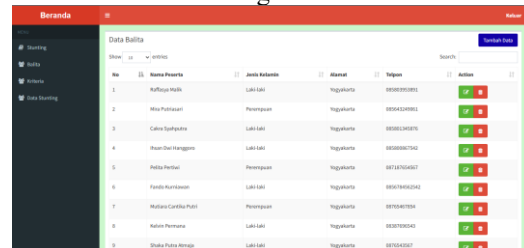


Figure 6. Toddler Menu Page View

In the toddler menu, the admin can manage toddler data including adding, changing and deleting toddler data.

### 3. Criteria Menu Page View

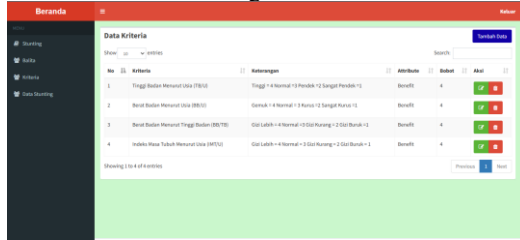


Figure 7. Criteria Menu Page View

in figure 7 is the criteria menu, the admin can manage criteria data in the form of adding, changing, deleting criteria data. in this menu the admin can also see the value of each criterion that has been determined, there is a description column to make it easier for the admin in the process of inputting toddler data values according to the criteria.

### 4. Stunting Menu Page View

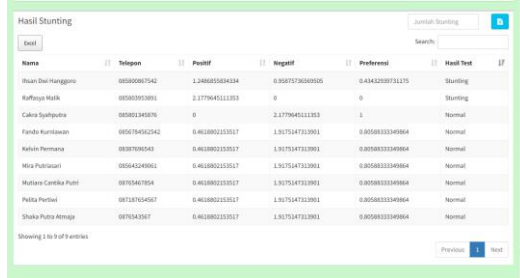


Figure 8. Stunting Menu Page View

Figure 8 shows the results page of the toddler value calculation process according to the anthropometric criteria that have been inputted. In the results, there are 2 toddlers who are included in the stunting category, namely Rafasya Malik (A1) and Ihsan Dwi Hanggoro (A4) according to the calculation of stunting determination using the TOPSIS method explained above.

#### 3.3 Black Box testing

| No  | Functions Tested     | Black Box Testing |
|-----|----------------------|-------------------|
| 1.  | Login Page           | Succeed           |
| 2   | Home Page            | Succeed           |
| 3.  | Toddler Page         | Succeed           |
| 4.  | Add Toddler Data     | Succeed           |
| 5.  | Change Toddler Data  | Succeed           |
| 6.  | Delete Toddler Data  | Succeed           |
| 7.  | Criteria Data Page   | Succeed           |
| 8.  | Add Criteria Data    | Succeed           |
| 9.  | Change Criteria Data | Succeed           |
| 10. | Delete Criteria Data | Succeed           |
| 11. | Stunting Data Page   | Succeed           |
| 12. | Add Stunting Data    | Succeed           |
| 13. | Change Stunting Data | Succeed           |

|     |                                |         |
|-----|--------------------------------|---------|
| 14. | Delete Stunting Data           | Succeed |
| 15. | Stunting Page                  | Succeed |
| 16. | Results of Calculation Process | Succeed |
| 17. | Exit                           | Succeed |

Table 9. Black Box Testing Results

The results of system testing using black box testing on the decision support system for determining stunted toddlers run and work well.

## 4. CONCLUSION

Based on the results of the discussion of the problems raised by the researcher, it was concluded that a Decision Support System (DSS) has been successfully developed to determine the stunting status of toddlers at the Bougenville Posyandu. This system applies the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS) method in its calculation process. The results of manual calculations and system calculations show similar results, so it can be concluded that the system works well and accurately. An example of the preference results shows that toddlers with code A1 (Rafasya Malik) have a preference value of 0, followed by A4 (Ihsan Dwi Hanggoro) with a preference of 0.4022149 which is included in the high stunting risk category. This system is able to assist Bougenville Posyandu cadres in determining the priority of handling toddlers at risk of stunting effectively and efficiently.

## 5. SUGGESTION

For further research, it is recommended that the TOPSIS method used in this decision support system be compared with other decision-making methods. This comparison aims to evaluate the advantages and disadvantages of each method in the context of determining stunted toddlers. In addition, it is to develop this decision support system by integrating it into the health information system that has been used in health service facilities such as Puskesmas. This integration will facilitate the process of data collection, monitoring the nutritional status of toddlers, and reporting to related parties in real-time and centrally. With this integration, the system can contribute to supporting government programs in overcoming stunting more efficiently, collaboratively, and sustainably.

## REFERENCES

- [1] V. Hasyati and K. Hutagalung, "INSTING ( INTEGRATED TO

- SOLVE STUNTING ): PLATFORM BERBASIS DIGITAL SEBAGAI UPAYA PENANGANAN STUNTING YANG TERINTEGRASI DI TENGAH SOCIETY 5.0,” *Berk. Ilm. Mhs. Ilmu Gizi Indones.*, vol. 10, no. 2, pp. 51–60, 2023.
- [2] F. Alfinnas, W. A. Triyanto, and R. R. Setiawan, “Metode Simple Additive Weighting ( SAW ) dalam Memonitoring dan Penanggulangan Stunting di Wilayah Desa Loram Wetan,” *J. Penelit. Inov.*, vol. 5, no. 1, pp. 653–660, 2025.
- [3] S. A. M.F, F. T. Naripana, and M. D. Alanwari, “Implementasi Metode Simple Additive Weighting Untuk Menurunkan Jumlah Balita Stunting Di Posyandu Dahlia 10,” *J. Rekayasa Teknol. Nusa Putra*, vol. 10, no. 2, pp. 65–74, 2024, doi: 10.52005/rekayasa.v10i2.542.
- [4] Y. H. Siregar, N. Irawati, A. Muhazir, L. C. Utami, and R. Pratiwi, “Menentukan Tingkat Risiko Penjualan Cake Di Stiinacake Dengan Menggunakan Metode AHP-MAUT,” *Nuansa Inform.*, vol. 17, pp. 1858–3911, 2023, [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>.
- [5] I. P. D. Suarnatha, “Sistem Pendukung Keputusan Seleksi Ketua Bem Menggunakan Metode Profile Matching,” *J. Inf. Syst. Manag.*, vol. 4, no. 2, pp. 73–80, 2023, doi: 10.24076/joism.2023v4i2.952.
- [6] A. K. Puspa and R. Nursyanti, “Sistem Pendukung Keputusan Penyakit Gizi Buruk Menggunakan Metode Simple Addictive Weighting (SAW),” *Expert J. Manaj. Sist. Inf. dan Teknol.*, vol. 7, no. 1, 2017, doi: 10.36448/jmsit.v7i1.876.
- [7] J. Papathanasiou and N. Ploskas, “Topsis,” *Springer Optim. Its Appl.*, vol. 136, pp. 1–30, 2018, doi: 10.1007/978-3-319-91648-4\_1.
- [8] Dimasqi Ramadhani and Galandaru Swalaganata, “Analisis Pengujian Sensitivitas Penggunaan Metode Pengambilan Keputusan Profile Matching, Topsis Dan Moora Dalam Menentukan Karyawan Terbaik,” *Nuansa Inform.*, vol. 18, no. 1, pp. 135–145, 2024, doi: 10.25134/ilkom.v18i1.94.
- [9] F. Lla Dwi Cahyanti, Elly Firasari, and Umi Khultsum, “Sistem Informasi Penjualan Berbasis WEB Pada UKM Rukun Makmur Tlingsing,” *Nuansa Inform.*, vol. 18, no. 2, pp. 11–18, 2024, doi: 10.25134/ilkom.v18i2.126.
- [10] P. Candra Susanto, D. Ulfah Arini, L. Yuntina, J. Panatap Soehaditama, and N. Nuraeni, “Konsep Penelitian Kuantitatif: Populasi, Sampel, dan Analisis Data (Sebuah Tinjauan Pustaka),” *J. Ilmu Multidisplin*, vol. 3, no. 1, pp. 1–12, 2024, doi: 10.38035/jim.v3i1.504.
- [11] F. Susanto, *Pengenalan Sistem Pendukung Keputusan*. Yogyakarta: Deepublish, 2020.
- [12] F. Heriyanti and A. Ishak, “Design of logistics information system in the finished product warehouse with the waterfall method: Review literature,” *IOP Conf. Ser. Mater. Sci. Eng.*, vol. 801, no. 1, 2020, doi: 10.1088/1757-899X/801/1/012100.
- [13] E. Y. Meol, D. Nababan, and Y. P. K. Kelen, “Sistem Informasi Penjualan Ikan pada Kefamenanu Berbasis Android Menggunakan Metode Waterfall,” *J. Krisnadana*, vol. 3, no. 2, pp. 78–89, 2024, doi: 10.58982/krisnadana.v3i2.527.
- [14] A. T. Purba and V. M. M. Siregar, “Sistem Penyeleksi Mahasiswa Baru Berbasis Web Menggunakan Metode Weighted Product,” *J. Tek. Inf. dan Komput.*, vol. 3, no. 1, p. 1, 2020, doi: 10.37600/tekinkom.v3i1.117.
- [15] U. Sholikhah, B. Rosyadi, S. R. Wahzuni, S. U. Alasna, and K. F. P. Maharani, “Perancangan Sistem Informasi Sekolah Berbasis Website Pada Mi Manbail Futuh Jenu Tuban,” *IJIS Indones. J. Inf. Syst.*, vol. 9, no. 2, pp. 120–131, 2024.
- [16] A. A. Saraswati and I. Mubarak, “Sistem Informasi Penerimaan Dan Pengeluaran Kas Berbasis Website Pada Pt Lkm Bkd Unit Balamoa,” *JATI (Jurnal Mhs. Tek. Inform.*, vol. 8, no. 3, pp. 3627–3638, 2024, doi: 10.36040/jati.v8i3.9756.
- [17] H. Kurniawan, W. Apriliah, I. Kurnia, and D. Firmansyah, “Penerapan Metode Waterfall Dalam Perancangan Sistem Informasi Penggajian Pada Smk Bina Karya Karawang,” *J. Interkom J. Publ. Ilm. Bid. Teknol. Inf. dan Komun.*, vol. 14, no. 4, pp. 13–23, 2021, doi: 10.35969/interkom.v14i4.78.
- [18] Y. Amanda and M. H. Ujjianti, “PERANCANGAN SISTEM INFORMASI PENERIMAAN

- PESERTA DIDIK BARU PADA DAYCARE DAN PRE SCHOOL ANANDA MANDIRI SLAWI BERBASIS WEB,” *J. Mhs. Tek. Inform.*, vol. 9, no. 1, pp. 177–184, 2025.
- [19] M. Ulandari and R. Zulfiandry, “Sistem Informasi Data Nilai Siswa Pada Sekolah Menengah Kejuruan Negeri 03 Empat Lawang Sumatera Selatan,” *J. Nedia Infotama*, vol. 21, no. 1, pp. 68–74, 2025.
- [20] Y. Saputra and D. Mardiaty, “Implementasi sistem informasi manajemen klinik menggunakan metode black box testing,” *J. Inform. dan Tek. Elektro Terap.*, vol. 13, no. 1, 2025.
- [21] R. Parlika, T. A. Nisaa’, S. M. Ningrum, and B. A. Haque, “Studi Literatur Kekurangan Dan Kelebihan Pengujian Black Box,” *Teknomatika*, vol. 10, no. 02, pp. 131–140, 2020.
- [22] L. Fauziah, A. Firmansyah, and A. Aguswin, “Sistem Informasi Sekolah Berbasis Web Menggunakan Metode Waterfall. Studi Kasus: SMPI Al-Hudri Walibrah,” *REMIK Ris. dan E-Jurnal Manaj. Inform. Komput.*, vol. 8, no. 1, pp. 274–285, 2024.

## Predicting the Happiness Index Based on the HDI Indicator in Indonesia Using the Ensemble Learning Approach

Rofi Nafiis Zain<sup>1</sup>, Iwan Setiawan<sup>2</sup>, Virdiandry Putratama<sup>3</sup>, Syafrial Fachri Pane<sup>\*4</sup>

<sup>1, 2, 4</sup>Applied Informatics Engineering, Universitas Logistik dan Bisnis Internasional, Indonesia

<sup>3</sup>Applied Management Informatics, Universitas Logistik dan Bisnis Internasional, Indonesia

E-mail: <sup>1</sup>[rofinafiisr@ulbi.ac.id](mailto:rofinafiisr@ulbi.ac.id), <sup>2</sup>[iwan.setiawan@ulbi.ac.id](mailto:iwan.setiawan@ulbi.ac.id), <sup>3</sup>[viridiandry@ulbi.ac.id](mailto:viridiandry@ulbi.ac.id),

<sup>\*4</sup>[syafrial.fachri@ulbi.ac.id](mailto:syafrial.fachri@ulbi.ac.id)

### Abstract

*Machine Learning is widely utilized to analyse complex data across various research domains. In this study, we employed an ensemble learning approach comprising Random Forest Regression (RF), XGBoost Regression (XGB), Decision Tree Regression (DT), Pearson correlation analysis, and Shapley Additive Explanations (SHAP) to examine the relationship between the Human Development Index (HDI) and happiness indicators in Indonesia. The study had two main objectives. First, to investigate the correlation between HDI and happiness. The Pearson correlation analysis yielded a value of 0.88, indicating a very strong positive relationship. This was further supported by the results of Permutation Importance (PI) and SHAP analysis, which identified HDI as one of the most influential variables in predicting happiness scores. Second, we developed a predictive model using a stacking ensemble approach, integrating RF, XGB, and DT regressors. The resulting model demonstrated high predictive performance, with an R-squared value of 97.68%, a Mean Absolute Error (MAE) of 0.002900, a Mean Squared Error (MSE) of 0.000021, and a Root Mean Squared Error (RMSE) of 0.004604.*

**Keywords**—Prediction, Happiness, HDI, Stacking Ensemble Learning, Indonesia

Submission: June 04, 2025

Accepted: June 14, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.410>

### 1. INTRODUCTION

Happiness is widely recognized as a key indicator of a nation's overall well-being, particularly in evaluating social and economic progress. The Happiness Index offers a quantitative measure of individuals' life satisfaction and overall happiness, providing valuable insights for policymakers across various domains [1]. In Indonesia, the pursuit of improved quality of life and happiness is explicitly identified as a national priority, reflecting the government's recognition of its significance [2]. Nevertheless, accurately measuring and predicting happiness requires a comprehensive, data-driven approach to support the development of effective and targeted policies [3].

The Human Development Index (HDI) plays a crucial role in evaluating societal well-being, as it encompasses key dimensions such as life expectancy, education levels, and employment rates—all of which substantially influence individual and collective happiness [4]. A high HDI is generally associated with greater well-being, as it reflects the quality of a nation's healthcare and education systems—fundamental pillars of social development [5]. However, the

relationship between HDI and happiness is not always straightforward, necessitating more nuanced analytical approaches to fully understand its complexity. While techniques such as Pearson correlation analysis can identify linear associations between HDI components and happiness [6], they are limited in scope. To gain deeper insights, advanced methods like Shapley Additive Explanations (SHAP) offer a more comprehensive view by quantifying the contribution of each HDI indicator in predicting happiness. These approaches enhance model interpretability and help identify the most influential factors affecting well-being [7].

Moreover, the advent of the Industrial Revolution 4.0 has heightened the significance of machine learning in analyzing complex datasets, including those related to happiness [8]. Machine learning excels at identifying intricate patterns within data, surpassing the capabilities of traditional statistical methods, and enabling the development of advanced predictive models that incorporate multiple HDI indicators. However, relying on a single model may not yield optimal predictive accuracy. Therefore, ensemble learning techniques—which combine the strengths of multiple models—are increasingly adopted to enhance forecasting

performance [2], [7]. In this study, we implemented three core models: Random Forest Regression, XGBoost Regression, and Decision Tree Regression, each contributing to improved predictions of happiness levels in Indonesia [9].

Previous studies have leveraged machine learning to predict happiness indices with promising results. For example, research by Jannani et al. identified GNI per capita, HDI, and life expectancy as key factors strongly associated with happiness. Among the models tested, Random Forest yielded the highest accuracy at 93.66%, with a Root Mean Square Error (RMSE) of 0.05748, highlighting GNI per capita as the most influential variable [10]. Similarly, Çiftci and Yıldız demonstrated that social media addiction negatively impacts happiness and life satisfaction, with life satisfaction serving as a mediating variable. Their study employed Elastic Net Regression, achieving 84.5% accuracy and underscoring the critical role of these psychosocial factors [11]. Additionally, research by Zhiqiang Jiang applied machine learning to predict regional economic growth and industrial structure, with the K-Nearest Neighbors (KNN) algorithm attaining an accuracy of 79.46% [12].

This study makes two key contributions. First, it assesses the significance of HDI-related variables in predicting happiness using both Pearson correlation and SHAP analysis. Second, it develops a predictive model for happiness based on HDI indicators using an ensemble learning approach.

## 2. RELATED WORKS

Machine learning models have been extensively applied to predict the Happiness Index in various studies. Jannani et al. demonstrated that the Random Forest model achieved the highest accuracy in predicting happiness based on GNI per capita, HDI, and life expectancy, with a precision of 93.66% and a Root Mean Square Error (RMSE) of 0.05748. Their findings identified GNI per capita as the most influential factor affecting happiness predictions [10]. In another study, Çiftci and Yıldız explored the negative impact of social media addiction on happiness and life satisfaction, with life satisfaction acting as a mediating variable. Using the Elastic Net Regression model, their study reached an accuracy of 84.5%, emphasizing the importance of social determinants in understanding happiness [11]. Furthermore, Zhiqiang Jiang applied machine learning techniques to forecast regional economic growth and industrial

structure. The K-Nearest Neighbors (KNN) algorithm achieved an accuracy of 79.46%, underscoring the significant role of economic variables in influencing overall well-being [12].

This study employs an ensemble learning approach to predict happiness based on Human Development Index (HDI) data by integrating Random Forest Regression (RF), XGBoost Regression (XGB), and Decision Tree Regression (DT) models. The ensemble method, combined with Permutation Importance (PI) and Shapley Additive Explanations (SHAP) analysis, outperforms individual models, demonstrating the advantages of aggregating multiple algorithms to enhance predictive accuracy. Moreover, the use of PI and SHAP not only strengthens model performance but also offers interpretable insights into the relative contributions of HDI indicators to happiness outcomes. This research contributes to the advancement of happiness prediction by presenting a robust and transparent ensemble learning framework.

## 3. RESEARCH DESIGN

### 3.1. Data Preparation

The first phase of this study involves data collection and pre-processing. The dataset was obtained from the Indonesian Central Statistics Agency and includes HDI-related indicators such as Life Expectancy, Gender Development Index, Expected Years of Schooling, Female and Male Labour Force Participation Rates, Mean Years of Schooling, and the overall Human Development Index. The data spans the period from 2018 to 2020. Following data acquisition, the pre-processing stage involves cleaning and preparing the dataset for modeling [13]. This includes data cleansing and distribution analysis [14] to identify and address missing values and to verify the distribution of each variable. Additionally, data types are checked to ensure compatibility with the analysis requirements. After these steps, all variables are standardized using the Standard Scaler method [15], which ensures that the features are on a comparable scale, thereby improving the reliability and performance of machine learning algorithms [16]. These procedures enhance data quality, improve analytical precision, and contribute to more robust and trustworthy regression model outcomes.

### 3.2. Descriptive and Statistical Analysis

The second phase involves descriptive and statistical analysis, including Exploratory Data Analysis (EDA) and correlation testing. EDA is used to visualize data distributions and identify trends, providing insights into variables that may influence the predicted outcomes. Subsequently, correlation analysis is conducted to assess linear relationships between variables, which helps determine the relative weight of each factor in predicting the Happiness Index and identify the most influential HDI indicators. Based on prior literature [10], this study employs Pearson correlation analysis, a widely used method for evaluating linear associations between continuous variables.

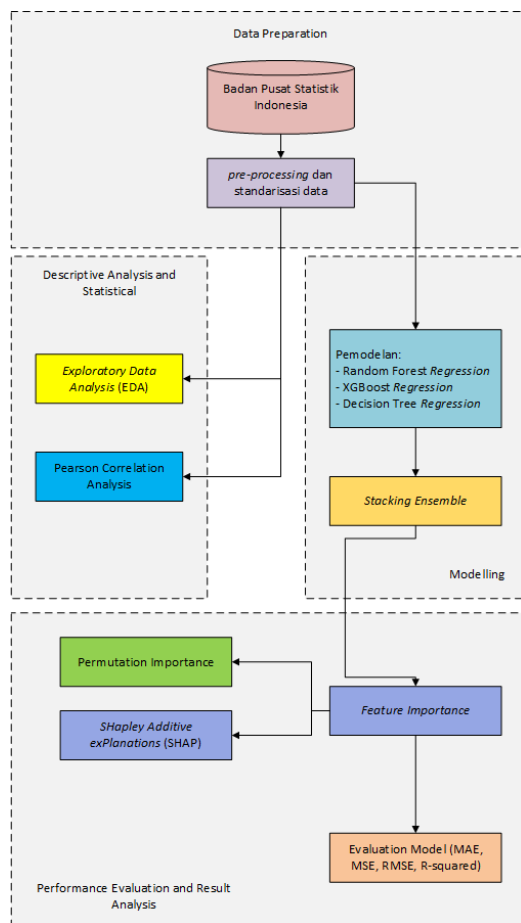


Figure 1 The proposed research methodology

$$r = \frac{\sum(x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum(x_i - \bar{x})^2 \sum(y_i - \bar{y})^2}} \quad (1)$$

Equation (1) presents the Pearson correlation formula, which measures the strength and direction of the linear relationship between two variables. A correlation coefficient of  $r=1$  indicates a perfect positive correlation, where both variables increase or decrease together. In contrast,  $r=-1$  signifies a

perfect negative correlation, meaning one variable increases as the other decreases. A value of  $r=0$  suggests no linear relationship exists between the two variables.

### 3.3. Modelling

The modeling stage is the core component of this study and consists of two main parts: individual model training and Stacking Ensemble implementation. A range of machine learning algorithms is employed to develop a regression model for predicting happiness scores, with each model trained on standardized data and evaluated for accuracy. Specifically, this study utilizes three regression algorithms: Random Forest Regression, XGBoost Regression, and Decision Tree Regression.

To further improve predictive performance, a Stacking Ensemble technique is applied after training the individual models. This ensemble method integrates the outputs of the three base learners to construct a more accurate and robust prediction model. By capitalizing on the strengths of each algorithm, the ensemble approach aims to minimize prediction errors and deliver the highest level of accuracy in forecasting desired scores.

$$\hat{y} = g(f_1(x), f_2(x), \dots, f_n(x)) \quad (2)$$

Equation (2) illustrates the Stacking Ensemble formula, where  $\hat{y}$  denotes the final predicted value produced by the ensemble model. This prediction is based on the outputs  $f_1(x), f_2(x), \dots, f_n(x)$  from the base learners—such as Random Forest Regression (RFR), XGBoost Regression (XGBR), and Decision Tree (DT). These outputs are then combined by a meta-learner, represented by the function  $g$ , which integrates the base model predictions to generate the final output.

### 3.4. Performance Evaluation and Result Analysis

The final phases of this study—performance evaluation and result analysis—aim to assess the accuracy and effectiveness of the developed prediction model. This process involves three key steps and utilizes several evaluation metrics, including Mean Absolute Error (MAE), Mean Squared Error (MSE), Root Mean Squared Error (RMSE), and R-Squared ( $R^2$ ). These metrics provide a comprehensive understanding of the model's predictive

performance. Specifically, MAE measures the average magnitude of prediction errors, offering insight into how accurately the model estimates happiness levels based on human development indicators.

$$MAE = \left(\frac{1}{n}\right) \sum_{i=1}^n |y_i - x_i| \quad (3)$$

Equation (3) explains that  $y_i$  is predicted value.  $x_i$  is true value.  $n$  is amount of data [17].

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (4)$$

Equation (4) shows the MSE formula, averages the error square between the prediction and actual values. Smaller MSE values improve model prediction [18].

$$RMSE = \sqrt{\left(\frac{1}{n}\right) \sum_{i=1}^n (y_i - x_i)^2} \quad (4)$$

Equation (4) explains that  $y_i$  is predicted value.  $x_i$  is true value.  $n$  is amount of data [19].

$$R - Square = 1 - \frac{\sum (y_i - \hat{y}_i)^2}{\sum (y_i - \bar{y}_i)^2} \quad (5)$$

Equation (5) defines the coefficient of determination ( $R^2$ ), which quantifies the proportion of variance in the dependent variable that can be explained by the model [20]. This metric serves as an indicator of the model's predictive performance relative to a baseline model that simply predicts the mean of the target variable.

To further understand variable contributions after model training, **Permutation Importance** is applied. This technique evaluates how the model's performance deteriorates when the values of a particular variable are randomly shuffled. The resulting change in performance helps identify how crucial each feature is in predicting the target variable.

Additionally, **SHAP (Shapley Additive Explanations)** is employed to interpret model predictions. SHAP assigns a contribution value to each input feature, allowing for a detailed explanation of how each variable influences the model's output. It quantifies both positive and negative effects and

captures the marginal impact of each attribute on individual predictions.

SHAP offers a significant advantage over Pearson correlation. While **Pearson correlation** provides a global, statistical measure of the linear relationship between a variable and the target, **SHAP delivers instance-level interpretability**, capturing both linear and non-linear interactions. Pearson is useful for identifying general trends, whereas SHAP reveals how specific features affect individual model predictions, offering a more granular and model-aware interpretation.

#### 4. RESULTS AND DISCUSSION

This section presents the experimental results, including data preparation, descriptive and statistical analysis, and regression model evaluation.

The study uses HDI-related data from the Indonesian Central Statistics Agency for the years 2018–2020. Each dataset contains annual observations of variables such as Life Expectancy, Gender Development Index (GDI), Expected and Mean Years of Schooling, Male and Female Labor Force Participation Rates, and HDI. Although direct Happiness Scores for Indonesian cities are not publicly available, the study uses proxy variables to estimate happiness levels.

Data pre-processing involved normalization using **StandardScaler**, converting percentages into decimal form and handling missing values to ensure consistency. Exploratory Data Analysis (EDA), conducted in Python, revealed that cities with higher HDI tend to report higher estimated Happiness Scores, suggesting a strong association between the two.

A **Pearson Correlation Analysis** was then conducted to assess the linear relationships among the variables. The Human Development Index (HDI) showed the strongest correlation with happiness and was assigned the highest weight (0.30) in the composite score. Mean Years of Schooling and Life Expectancy followed closely with weights of 0.25 each. Female and Male Labor Force Participation Rates were each assigned weights of 0.15, as they appeared to negatively influence schooling indicators. Due to its low correlation with the other variables, the Gender Development Index (GDI) received the lowest weight.

This weighted approach ensures that variables with the strongest associations have the greatest influence on the predicted Happiness Score. The results not only support the development of an accurate prediction

model but also highlight the most influential factors shaping happiness in Indonesia.

Table 1 Weight of Happiness Index Calculation

| No. | Parameters | Value |
|-----|------------|-------|
| 1.  | $w_1$      | 0.25  |
| 2.  | $w_2$      | 0.3   |
| 3.  | $w_3$      | 0.25  |
| 4.  | $w_4$      | 0.15  |
| 5.  | $w_5$      | 0.1   |

$$H_{Index} = w_1 \cdot a + w_2 \cdot b + w_3 \cdot c + w_4 \cdot \frac{d + e}{2} + w_5 \cdot f \quad (6)$$

The Happiness Index is calculated using Equation (6), where “a” is Life Expectancy, “b” is the Human Development Index, “c” is the Mean Years of Schooling, d and e are the Female and Male Labor Force Rates, and “f” is the Gender Development Index. First Pearson correlation analysis multiplies variables by predefined weights. The average is calculated first for female and male Labor Force Rate variables before multiplying by weight. The Happiness Index is the sum of each variable's product.

The correlation analysis shows that the Human Development Index has a very strong positive correlation with the Happiness Score (correlation value of 0.88), indicating that the higher the Development Index value, the higher the Happiness Score in a city. Life Expectancy (0.76), Mean Years of Schooling (0.72), and Expected Years of Schooling (0.64) also positively affect the Happiness Score.

The Pearson correlation analysis examines Indonesian Happiness Scores and HDI indicators. The heatmap in Figure 2, and 3 depicts the degree and direction of each variable's link, helping pinpoint the primary factors affecting happiness in different cities.

Regression models employing Random Forest (RF), XGBoost (XGB), and Decision Tree (DT) predicted Happiness. Figure 4 shows  $R^2$  values of 97.66% for RF, 96.26% for XGB, and 92.65% for DT in the regression findings of the trained models. This method predicts the Happiness index well, but Stacking Ensemble can improve it. The three models are combined using a stacking ensemble method to improve forecast accuracy.

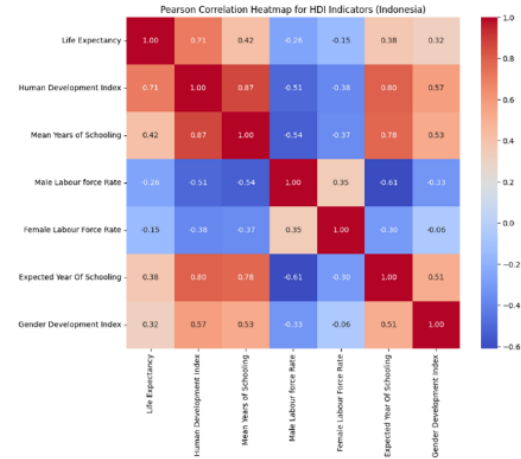


Figure 2 Pearson Correlation Before Happiness Index Calculated

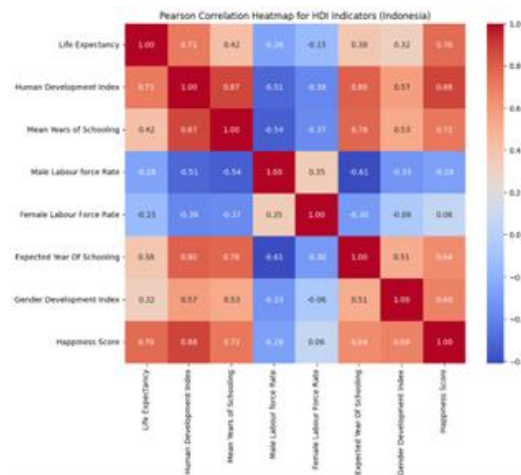


Figure 3 Pearson Correlation After Happiness Index Calculated

Stacking Ensemble combines models using Linear Regression as the main estimator. The Stacking Regressor excels with an MAE of 0.002900, MSE of 0.000021, RMSE of 0.004604, and  $R^2$  of 97.68%, outperforming comparable models. These outcomes surpass models such as Random Forest Regression ( $R^2 = 97.66\%$ ), XGBoost Regression ( $R^2 = 96.26\%$ ), and Decision Tree Regression ( $R^2 = 92.65\%$ ). The stacking ensemble model outperforms individual models by 0.02%. This shows that Stacking Ensemble enhances prediction accuracy by combining regression model strengths. Less prediction error means stacking models forecast better than individual models.

**Algorithm 1** : HDI Regression Analysis for Happiness Score Prediction and Feature Importance

1 Dataset HDI (Life Expectancy, Gender

```

Development Index, Expected Year of
Schooling, Female and Male Labour
Force Rate, Mean Years of Schooling)
Predicted happiness scores, feature
importance rankings
2 Load dataset:
  hdi_data =
  read_csv("HDI_2019.csv")
3 Display first 10 rows of hdi_data
4 Handle missing values (missing values,
  handle types)
5 If hdi_data is not empty then;
6 Identify percentage columns:
  columns = list of columns containing
  %;
7 Convert percentage columns to
  decimals for easier computation;
8 Fill missing values with appropriate
  strategy;
9 Calculate composite score from
  features;
10 Construct formula for Happiness
  Score:

$$H_{index} = w_1 \cdot a + w_2 \cdot b + w_3 \cdot c + w_4 \cdot \frac{d+e}{2} + w_5 \cdot f$$

11 Split data into features X and target Y
  (Happiness Score);
12 Train-test split: train_test_split(X, Y,
  train_size=0.8);
13 Initialize regression models:
  • LinearRegression
  • Decision Tree
  • Random Forest Regressor
  • XGB Regressor
14 Fit each model on training data:
  model.fit(X_train, Y_train);
15 Predict on test data: Y_pred =
  model.predict(X_test);
16 Calculate evaluation metrics: MAE,
  MSE, RMSE, R2;
17 For each model in models do;
18 Perform cross-validation:
  cross_val_score(model, X, Y);
19 Calculate permutation importance;
20 Plot permutation importance graph;
21 Display cross-validation score;
22 Ensemble combination using
  estimators:
  • ('rf', RandomForestRegressor())
  • ('xgb', XGBRegressor())
  • ('dt', DecisionTree())
23 Fit ensemble model on training data;
24 Save trained model:
  joblib.dump(stacked_model,
  "stacked_model.pkl");
25 Return: Predicted happiness scores,
  feature importance ranking;

```

Algorithm 1 shows the basic steps in creating a Stacking Regressor model to predict Indonesian Happiness Scores.

Feature selection, data separation, estimator definition, model training, performance evaluation, and evaluation findings are covered in this pseudocode. The goal is to summarize the model implementation workflow.

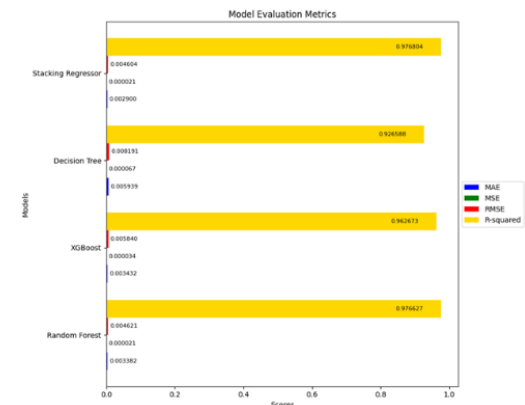


Figure 4 Evaluation Metrics from Model Result

Performance evaluation and result analysis use Feature Importance on each regression base-model to discover the most important factors predicting happiness. The data context and kind determine the best technique. This study uses Random Forest, XGBoost, and Decision Tree base-models' built-in feature importance characteristics and Permutation Importance. Combining these methods helps determine which model features most affect happiness prediction.

Permutation Importance (PI) and SHAP were utilized at this level. Decision Tree, XGBoost, Random Forest, and Stacking Ensemble Regressor Permutation Importance results are shown in Figure 5 -8. The data shows that the Human Development Index (HDI) predicts happiness the greatest, with a score of 1.2 to 1.7. With values around 0.2 to 0.3, the Female Labor Force Rate and Life Expectancy variables likewise rank second and substantially affect all models. The happiness prediction model emphasizes female labor force involvement and life expectancy. Mean Years of Schooling, Gender Development Index, Male Labor Force Rate, and Expected Year of Schooling also affect HDI and Female Labor Force Rate, albeit less so. These findings suggest that human development, female labor force involvement, and life expectancy drive happiness in Indonesia, while education plays a minor influence.

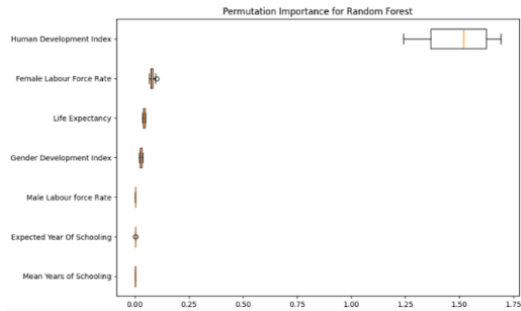


Figure 5 Permutation Importance Result RandomForest

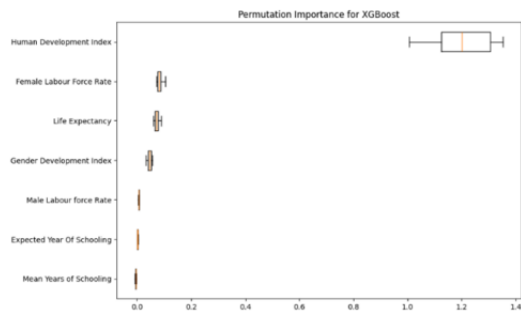


Figure 6 Permutation Importance Result XGB Regressor

The four SHAP plots (Decision Tree, XGBoost, Random Forest, and Stacking Ensemble Regressor) reveal that the Human Development Index (HDI) is the most influential variable in predicting happiness across all models (Figure 9 - 12). The Decision Tree and Random Forest models' SHAP HDI values range from -0.10 to 0.05, but XGBoost's are -0.08 to 0.06. While HDI remains dominant, the Stacking Ensemble model distributes SHAP values more evenly. Along with HDI, the Female Labor Force Rate and Life Expectancy variables contribute to happiness prediction, although less so. Mean Years of Schooling and Expected Years of Schooling have low SHAP values, indicating little impact.

Permutation Importance is easier to implement and suitable for preliminary assessments of complex models, but SHAP has better interpretative depth and visualization. SHAP clarifies prediction results, while Permutation Importance evaluates features' entire relevance.

Among the models tested in this study, the Stacking Ensemble Regressor performs best. Figure 4 shows that the Stacking Ensemble Regressor model has an MAE of 0.002900, MSE of 0.000021, RMSE of 0.004604, and  $R^2$  of 97.68%.

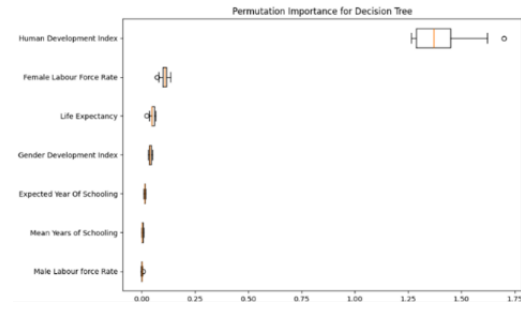


Figure 7 Permutation Importance Result Decision Tree

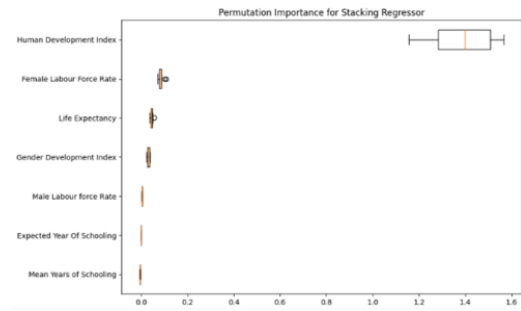


Figure 8 Permutation Importance Result Stacking Regressor

This high R-squared value shows that the model explains data variability well. Stacking Ensemble Regressor predicts happiness scores best based on indicators.

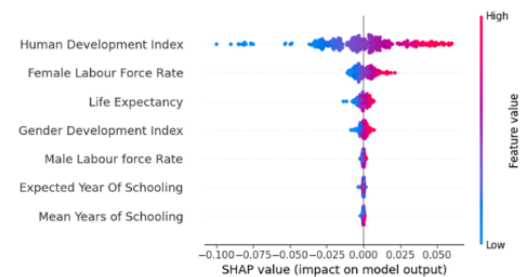


Figure 9 SHAP Result Random Forest

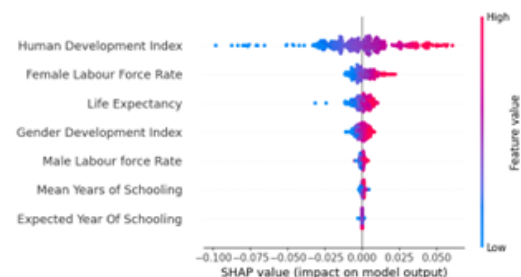


Figure 10 SHAP Result XGB Regressor

#### 4. DISCUSSION

The proposed Stacking Ensemble Regressor significantly outperforms individual models in predicting the Happiness Index based on HDI-related indicators, achieving an R-Squared value of 97.68%, slightly higher than Random Forest (97.66%), XGBoost (96.26%), and Decision Tree (92.65%). This improvement is primarily attributed to the ensemble's ability to integrate the strengths of diverse base learners, thereby reducing variance and minimizing prediction errors, as reflected by the lowest MAE (0.002900) and RMSE (0.004604). To enhance model interpretability, both Permutation Importance and SHAP (SHapley Additive exPlanations) analyses were employed. These analyses consistently highlight the Human Development Index (HDI), Life Expectancy, and Female Labour Force Rate as the most impactful predictors, whereas variables such as the Gender Development Index and Expected Years of Schooling exhibit comparatively marginal effects.



Figure 11 SHAP Result Decision Tree

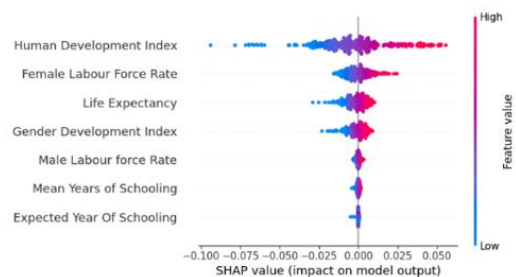


Figure 12 SHAP Result Stacking Regressor

The SHAP summary plot provides granular insight into how each feature contributes to the model's prediction. In this context, a positive SHAP value implies that the feature increases the predicted value of the Happiness Index for a specific instance, indicating a supportive or enhancing effect. Conversely, a negative SHAP value suggests that the feature contributes to lowering the predicted Happiness Index, exerting a

suppressive influence. For instance, a high HDI value (depicted in red) tends to have a positive SHAP value, meaning that higher HDI consistently leads to higher predicted happiness levels. Meanwhile, lower feature values (blue) may contribute negatively depending on the indicator's nature and direction of influence. This dual-axis interpretation—magnitude and direction—enables stakeholders to understand not only which variables are important but also how and in what direction they affect prediction outcomes. Compared to prior models such as that of Jannani et al., where Random Forest reached 93.66% accuracy, the integration of ensemble learning with SHAP-based interpretability in this study offers both superior performance and greater transparency, contributing a valuable methodology for evidence-based urban policy and happiness modeling.

#### 5. CONCLUSION

A comparative examination of literature about Happiness Index forecasts hinges on the employed model, the assessment methodology, and the analytical approach utilized. Diverse regression models have been employed by numerous academics to forecast the happiness index with differing degrees of accuracy. Research [10] employing the Random Forest model attained a  $R^2$  value of 93.66%, with model assessment encompassing Mean Squared Error (MSE), Mean Absolute Error (MAE), and Root Mean Squared Error (RMSE). This study used Permutation Importance (PI) for feature analysis, demonstrating that the GNI per capita variable significantly influences the model. A separate study [11] employed the Elastic Net Regression model, achieving a  $R^2$  accuracy of 44.05%, however no additional evaluation methods were specified. Similarly, the study[21] employed the XGBoost Regression model, achieving a  $R^2$  of 32%. This study demonstrated that the Stacking Ensemble Regressor approach (Proposed approach/PM) attained the maximum accuracy, evidenced by a  $R^2$  value of 97.68%, a Mean Absolute Error (MAE) of 0.002900, a Mean Squared Error (MSE) of 0.000021, and a Root Mean Squared Error (RMSE) of 0.004604. This approach integrates analysis using Permutation Importance and SHAP, demonstrating that the Human Development Index significantly impacts happiness predictions.

#### ACKNOWLEDGEMENT

The authors gratefully acknowledge the support provided by Universitas Logistik dan Bisnis Internasional in facilitating the completion of this research.

# REFERENCES

- [1] V. Susanti and A. Fitri, "Economics of Happiness: What Really Counts?," *Kne Social Sciences*, 2024, doi: 10.18502/kss.v9i16.16258.
- [2] Kenzo, A. Yudianto, M. A. Nugroho, and J. S. Mustika, "Analyzing Oxford Happiness Questionnaire Indonesian Version Using the Generalized Partial Credit Model," *Psyche 165 Journal*, 2024, doi: 10.35134/jpsy165.v17i2.353.
- [3] Q. Liu, "Can Happiness Be Measured?," *Advances in Education Humanities and Social Science Research*, 2023, doi: 10.56028/aehtsr.7.1.423.2023.
- [4] K. Ruggeri, E. Garcia-Garzon, Á. Maguire, S. Matz, and F. A. Huppert, "Well-being is more than happiness and life satisfaction: a multidimensional analysis of 21 countries," *Health Qual Life Outcomes*, vol. 18, pp. 1–16, 2020.
- [5] A. \. I. Akgun, S. P. Türkoğlu, and S. Erikli, "Investigating the determinants of happiness index in EU-27 countries: a quantile regression approach," *International Journal of Sociology and Social Policy*, vol. 43, no. 1/2, pp. 156–177, 2022.
- [6] J. Tanuwijaya, C. Krisanti, and A. W. Gunawan, "The Impact of Perceived Inclusion Climate for Leader Diversity (PICLD) on Organizational Justice and Its Effects on Employee Engagement, Emotional Wage, and Happiness at Work," *Ajesh*, 2024, doi: 10.46799/ajesh.v3i10.423.
- [7] P. R. Sihombing, "Comparison of Regression Analysis With Machine Learning Supervised Predictive Model Techniques," *Jurnal Ekonomi Dan Statistik Indonesia*, 2023, doi: 10.11594/jesi.03.02.03.
- [8] S. C. Br Ginting, S. Afifuddin, and R. RAHMANTA, "Analysis of the Effect of Macroeconomic Variables on Happiness in Indonesia," *International Journal of Research and Review*, 2021, doi: 10.52403/ijrr.20211009.
- [9] M. Acharya, A. Dumre, P. Poudel, and S. C. Dhakal, "Happiness Index; Current Status, Global Ranking and Its Importance in the Government's Agenda 'Prosperous Nepal, Happy Nepali,'" *Acta Scientific Agriculture*, 2020, doi: 10.31080/asag.2020.04.0789.
- [10] A. Jannani, N. Sael, and F. Benabbou, "Machine learning for the analysis of quality of life using the World Happiness Index and Human Development Indicators," *Mathematical Modeling and Computing*, vol. 10, no. 2, pp. 534–546, 2023, doi: 10.23939/mmc2023.02.534.
- [11] N. Çiftci and M. Yldz, "The relationship between social media addiction, happiness, and life satisfaction in adults: analysis with machine learning approach," *Int J Ment Health Addict*, vol. 21, no. 5, pp. 3500–3516, 2023.
- [12] Z. Jiang, "Prediction and industrial structure analysis of local GDP economy based on machine learning," *Math Probl Eng*, vol. 2022, no. 1, p. 7089914, 2022.
- [13] V. Çetin and O. Yldz, "A comprehensive review on data preprocessing techniques in data analysis," *Pamukkale Üniversitesi Mühendislik Bilimleri Dergisi*, vol. 28, no. 2, pp. 299–312, 2022.
- [14] K. Maharana, S. Mondal, and B. Nemade, "A review: Data pre-processing and data augmentation techniques," *Global Transitions Proceedings*, vol. 3, no. 1, pp. 91–99, 2022.
- [15] I. Izonin, R. Tkachenko, N. Shakhovska, B. Ilchyshyn, and K. K. Singh, "A two-step data normalization approach for improving classification accuracy in the medical diagnosis domain," *Mathematics*, vol. 10, no. 11, p. 1942, 2022.
- [16] N. Pessanha Santos, "The Expansion of Data Science: Dataset Standardization," *Standards*, vol. 3, no. 4, pp. 400–410, 2023.

- [17] S. F. Pane, A. G. Putrada, N. Alamsyah, and M. N. Fauzan, "A PSO-GBR solution for association rule optimization on supermarket sales," in *2022 seventh international conference on informatics and computing (ICIC)*, 2022, pp. 1–6.
- [18] S. F. Pane, A. Adiwijaya, M. D. Sulistiyo, and A. A. Gozali, "Multi-Temporal Factors to Analyze Indonesian Government Policies regarding Restrictions on Community Activities during COVID-19 Pandemic," *JOIV: International Journal on Informatics Visualization*, vol. 7, no. 4, pp. 2263–2269, 2023.
- [19] S. F. Pane, A. G. Putrada, N. Alamsyah, M. N. Fauzan, and others, "The Influence of The COVID-19 Pandemics in Indonesia On Predicting Economic Sectors," in *2022 Seventh International Conference on Informatics and Computing (ICIC)*, 2022, pp. 1–6.
- [20] S. F. Pane, F. Abdullah, and R. Habibi, "DETEKSI EMOSI PADA TEKS BERBAHASA INDONESIA MENGGUNAKAN PENDEKATAN ENSEMBLE," *JTT (Jurnal Teknologi Terapan)*, vol. 10, no. 2, pp. 80–90, 2024.
- [21] E.-J. Kim, H.-W. Kang, and S.-M. Park, "Leisure and Happiness of the Elderly: A Machine Learning Approach," *Sustainability*, vol. 16, no. 7, p. 2730, 2024.

## Performance and Efficiency Testing Analysis of Database Systems in Academic Information Systems

Utami Kusuma Dewi<sup>\*1</sup>, Ryan Lingga Wicaksono<sup>2</sup>, Mahendra Dwi Febri<sup>3</sup>, Villy Satria<sup>4</sup>

<sup>1,2,3,4</sup>Telkom University, Indonesia

E-mail: <sup>\*1</sup>[utamikusdew@telkomuniversity.ac.id](mailto:utamikusdew@telkomuniversity.ac.id), <sup>2</sup>[wicaksonoryan@telkomuniversity.ac.id](mailto:wicaksonoryan@telkomuniversity.ac.id),  
<sup>3</sup>[mahendradp@telkomuniversity.ac.id](mailto:mahendradp@telkomuniversity.ac.id), <sup>4</sup>[villysatria@telkomuniversity.ac.id](mailto:villysatria@telkomuniversity.ac.id)

### Abstract

*This study examines the performance and efficiency of database systems within academic information systems, acknowledging the increasing demand for responsiveness and reliability in managing complex academic data. As educational institutions increasingly rely on digital systems, performance testing becomes essential to ensure that these systems continue to support the learning environment effectively. Guided by the ISO/IEC 25010 standard, the research focuses on evaluating three key aspects of performance efficiency: time behavior, resource utilization, and capacity. Using JMeter, a range of user load scenarios were simulated, and the results were examined through Control Quality Charts and Nelson Rules to detect underlying issues affecting system performance. The findings reveal that 82.5% of queries demonstrated good time behavior, and 80% performed well in resource usage. However, half of the tests related to capacity highlighted the need for further improvements. Some queries experienced delays and consumed excessive CPU and memory resources, indicating areas where optimization is required. These insights highlight the importance of refining queries and managing resources more effectively to ensure a seamless user experience. Future research should consider automated optimization, machine learning-based performance prediction, and system scalability, especially in more dynamic and distributed academic environments.*

**Keywords**—database performance, academic information system, ISO/IEC 25010, JMeter, quality control chart.

Submission: June 19, 2025

Accepted: July 05, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.438>

### 1. INTRODUCTION

The advancement of information technology has significantly influenced nearly every aspect of daily life today. In the modern era, almost all fields of work rely on information and communication technology to support their operations, including data management in the fields of education and academia. [1]. The education and academic sectors are inherently complex systems, as they involve various components that can contribute to internal challenges. [2]. Academic systems play a crucial role in managing academic data and information, including schedules, grades, attendance records, and other details related to the teaching and learning process. [3]. At the core of these systems lies the database, a vital component responsible for storing and managing continuously evolving information [4]. As the number of users and the complexity of data continue to grow, database performance becomes a key factor in ensuring that system operations run smoothly, effectively, and responsively. [5]. Efficiency is a key aspect in nearly all software systems, including those

used in information systems. This efficiency is measured based on factors such as response time, the number of users the system can handle, and its energy efficiency.[6]. Therefore, performance testing of databases is crucial for evaluating and identifying potential issues that may compromise overall system performance. [5].

Performance testing is a method used to assess how well an information system or software performs under specific load conditions.[7], [8] One of the standard models for evaluating software quality is ISO/IEC 25010 [9], which defines a range of software quality characteristics, including performance efficiency. This aspect includes sub-characteristics such as Time Behavior, Resource Utilization, and Capacity. The model provides a comprehensive framework for evaluating the performance quality of an information system or software application, measuring response time, resource utilization, and the system's maximum handling capacity. [9].

The Control Quality Chart and Nelson Rules serve as practical tools for identifying deviations or anomalies in system performance.

The Control Quality Chart is used to monitor process variations, while the Nelson Rules help detect patterns that may indicate potential issues in system operations. [10]. Applying these tools within the context of database performance testing enables more accurate management and evaluation of system quality.

In this study, the quality of the database is measured within an academic information system that manages data related to education and teaching. The evaluation is conducted using a performance testing approach, focusing on the previously mentioned sub-characteristics. The results of the measurements are analyzed using the Control Quality Chart and Nelson's Rules to identify areas requiring improvement and to assess the overall quality of the system's performance. This approach aims to ensure that the database used in the academic information system can handle dynamic data growth and high-performance demands. [9], [10]

## 2. RESEARCH METHOD

The research method is divided into several main stages, as illustrated in Figure 1 below.

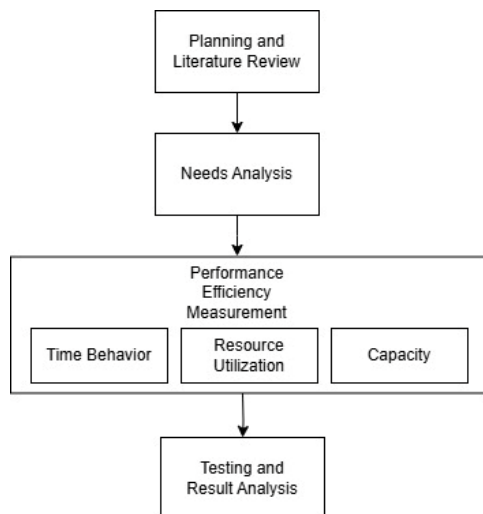


Figure 1. Research Method

### 2.1. Planning and Literature Review

The study begins with problem formulation, focusing on database performance testing, specifically in measuring response time, resource utilization, and capacity to ensure adequate system performance. The primary objective of this research is to evaluate the performance of the database system used in managing academic data. This stage also includes a literature review on ISO/IEC 25010, performance testing methodologies, and testing

tools such as JMeter, Control Quality Chart, and Nelson Rules, which serve as the foundation for designing the performance testing to be conducted.

Table 1. Specification Needs

| System Name          | Specification                     |
|----------------------|-----------------------------------|
| Processor            | Intel Core i5                     |
| Memory               | 16 GB                             |
| Disk                 | 500 GB                            |
| System Operation     | Microsoft Windows 11              |
| Database             | Oracle 11 g                       |
| Programming Language | Java                              |
| IDE                  | Netbeans                          |
| Local server         | Localhost                         |
| Browser              | Google Chrome and Mozilla Firefox |

In addition, the research planning also includes the need for both hardware and software to develop tools for conducting measurements, as outlined in Table 1.

### 2.2 Needs Analysis

This stage involves reviewing the specifications and components of the information system used to manage academic data.

#### a. Review the Requirement Specification Database

This involves reviewing the database's functional and non-functional requirements, including access speed, storage capacity, and data processing capabilities.

#### b. Identification of Database Component Specifications

This stage involves identifying the main components of the database system to be tested, including data structures, the queries in use, and the underlying storage and processing infrastructure.

#### c. Query Selection

This stage involves selecting the queries to be used in the performance tests, particularly those that involve user data and academic information.

In general, this stage is intended to ensure that the performance testing encompasses the critical components of the system that impact database performance.

### 2.3 Measurement

The third stage involves developing and configuring testing tools to measure the performance efficiency of the database system. The focus of the measurement is on three main sub-characteristics: Time Behavior, Resource Utilization, and Capacity. The detailed matrix of these sub-characteristics is presented in Table 2.

Table 2. Performance Testing Sub-Characteristics

| Sub-Characteristics  | Matrix                                      |
|----------------------|---------------------------------------------|
| Time Behavior        | Response Time                               |
|                      | Turnaround Time                             |
|                      | Throughput                                  |
|                      | Response Compliance                         |
|                      | Processing Margin Rate                      |
| Resource Utilization | Data Volume                                 |
|                      | Memory Capacity                             |
|                      | Hard Disc Capacity                          |
|                      | I/O Device                                  |
|                      | CPU Utilization                             |
|                      | Storage Period                              |
| Capacity             | Max of I/O Device                           |
|                      | Max Utilization Rate of Transmission System |

The tool used for testing is JMeter, which is configured to simulate various user load scenarios and measure response time, CPU utilization, and memory utilization.[11]. Additionally, the Control Quality Chart and Nelson Rules are prepared to analyze test results, monitor performance variations, and identify potential issues in the system's operational processes. [5], [9], [10].

### 2.3.1 JMeter Configurations

JMeter is configured by designing a test plan that includes various test cases to measure response time and resource utilization during the testing process. Each selected query is executed under multiple user load scenarios, ranging from light to heavy loads. The test results are then used to evaluate system performance and identify potential performance issues, such as high latency or inefficient resource usage. [5], [9], [10].

### 2.3.2 Query

Query testing is conducted by evaluating various types of database operations, such as select, insert, update, and delete, which are commonly used in managing data within academic information systems. The attributes measured during the testing include Execute Time, Memory Runtime, Memory Usage, and CPU Time. Each query is tested under different

load levels to ensure that the system can handle a wide range of user requests efficiently.

### 2.4 Testing and Result Analysis

During testing, JMeter collects data related to response time, CPU utilization, memory usage, and I/O capacity to evaluate system performance. The test results are used to analyze how well the system handles load and utilizes resources efficiently, as well as to identify areas that require improvement. Subsequently, the collected data will be further analyzed using the Control Quality Chart and Nelson Rules.

#### a. Control Quality Chart

In this study, the Control Quality Chart is used to analyze the results of database performance testing, specifically to monitor response time and resource utilization, including CPU and memory usage.

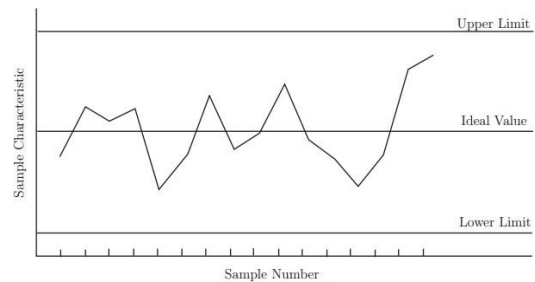


Figure 2. Control Quality Chart

The data from JMeter will be mapped into a chart, as illustrated in Figure 2, which is used to identify deviations from the established control limits, indicating potential issues in system performance. This tool helps ensure that the system remains within the desired quality control and provides insights for improvement when significant variability occurs..

Ideal Value

$$\bar{x} = \frac{x_1 + x_2 + \dots + x_k}{k} \quad (1)$$

Upper Limit

$$UCL = \bar{x} + \frac{3s}{\sqrt{n}} \quad (2)$$

Lower Limit

$$LCL = \bar{x} - \frac{3s}{\sqrt{n}} \quad (3)$$

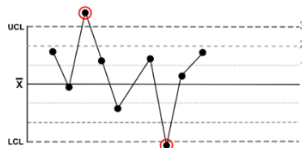
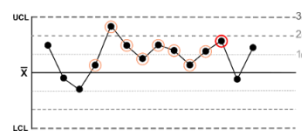
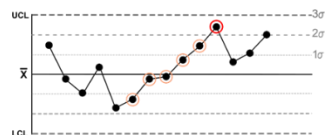
Keterangan

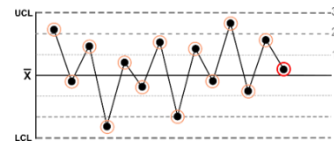
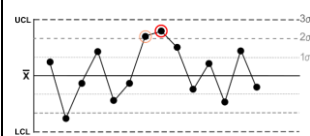
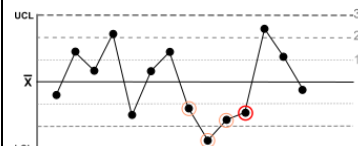
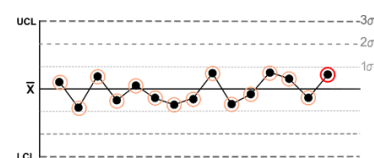
|           |   |                        |
|-----------|---|------------------------|
| $\bar{x}$ | : | Average                |
| $x_1$     | : | Each data value        |
| $k$       | : | Data Sum               |
| $s$       | : | Standart Deviation     |
| $n$       | : | Sum of the data sample |

b. Nelson Rules

Nelson Rules are used in this study to detect significant deviations in the database performance testing data. These rules help identify patterns or anomalies in response time, resource utilization, or capacity that indicate the system is operating outside of its normal range. In the context of this testing, Nelson Rules are applied to examine the results from the Control Quality Chart, to identify unusual spikes or fluctuations in performance that may signal issues or potential areas for improvement in the database system. [5], [9], [10]. The graph analysis using Nelson's Rules can be seen in Table 3.

Table 3. Nelson Rules

| Rules   | Descriptions                                                                        | Kegunaan                                                     |
|---------|-------------------------------------------------------------------------------------|--------------------------------------------------------------|
| Rules 1 | There is a single point located more than three standard deviations from the mean.  | Indicates a single anomaly in the process                    |
|         |   |                                                              |
| Rules 2 | There are nine (or more) consecutive points on the same side of the mean            | Indicates a process shift around the average value           |
|         |  |                                                              |
| Rules 3 | Six (or more) consecutive points that are continuously increasing or decreasing     | Indicates a trend or tendency in the process                 |
|         |  |                                                              |
| Rules 4 | Fourteen (or more) points that alternately increase and decrease in value.          | Indicates systematic effects such as changes in the machine, |

|         |                                                                                                   |                                                                                                                   |
|---------|---------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------|
|         |                                                                                                   | system, or operator                                                                                               |
|         |                 |                                                                                                                   |
| ules 5  | Two (or three) out of three consecutive points are beyond the mean in the same direction.         | Indicates a shift in the process average or an increase in standard deviation                                     |
|         |                 |                                                                                                                   |
| Rules 6 | Four (or five) out of five consecutive points are within the mean zone and in the same direction. | Indicates a shift in the process mean                                                                             |
|         |               |                                                                                                                   |
| Rules 7 | Fifteen consecutive points fall within one zone of the mean on both sides.                        | Indicates subgroup stratification when observations come from different sources or are collected in various ways. |
|         |               |                                                                                                                   |
| Rules 8 | Eight non-consecutive points, with values alternating on                                          | Indicates subgroup stratification when observations                                                               |

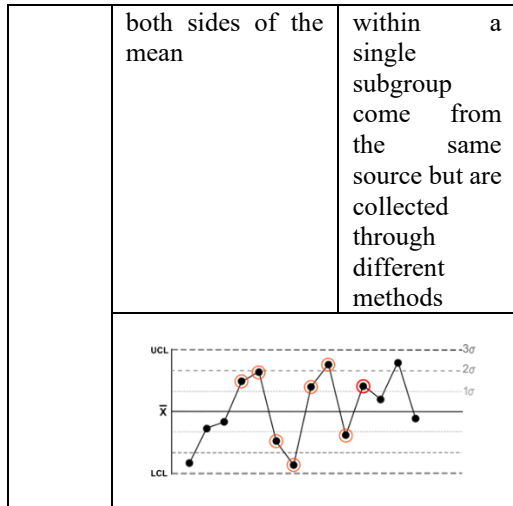


Table 4 presents the identified causes and corresponding explanations based on the graphs obtained from the measurement results.

Table 4. Possible Causes

| Rules | Possible Causes                                                                                                                                                                                                  |
|-------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1,2   | Configuration error<br>Errors in measurement and planning<br>Incomplete operation<br>New personnel performing the task<br>A process step was skipped<br>A process step was not completed                         |
| 3,4   | Change in supporting equipment<br>Change in work instructions<br>New or repaired tools introduced<br>Change in work procedures<br>Change in the system<br>New operator assigned<br>Change in maintenance program |
| 5,3   | Using tools or certain machine parts<br>Operator fatigue<br>Temperature factors: the machine is too hot or too cold<br>Inadequate maintenance                                                                    |
| 6     | Lack of control after process improvement<br>A mixture of more than one process<br>Strong tendency to fall slightly out of control                                                                               |
| 7     | Producing varying outcomes<br>More than one process is invol.ved                                                                                                                                                 |
| 8     | Overcontrol or misgrouping (subgroups taken from different sources)<br>Operator interference or disruption                                                                                                       |

Overall, this analysis aims to provide a comprehensive overview of system performance and to identify areas that require improvement or optimization.

### 3. RESEARCH RESULT

The results of this study are presented below, based on measurements conducted using the specified sub-characteristics.

#### 3.1 Time Behavior

Based on the measurement results, one query was found to have a response time exceeding the average, reaching 1.16 seconds. This finding is categorized as an anomaly according to Nelson Rule 1, which identifies data points that fall beyond three standard deviations from the mean. This condition indicates that the query is running inefficiently and may become a source of system delays if not addressed promptly.

To resolve this issue, it is recommended to optimize the query by reviewing its SQL structure, including the possible use of complex joins, nested subqueries, or filters that do not utilize indexes. Query tuning strategies, such as adding appropriate indexes, simplifying query logic, or utilizing stored procedures, can help reduce response time. Additionally, reviewing system configuration and process flows is essential to ensure there are no misconfigurations or unfinished executions. These efforts aim to maintain system performance stability, especially under high user load conditions.

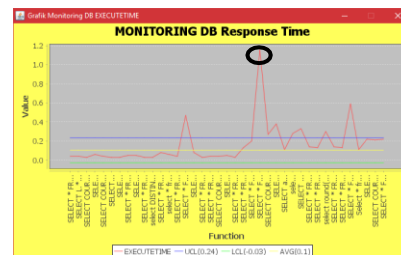


Figure 3. Graph of result time behavior

#### 3.2 Resource Utilization

The resource utilization testing was conducted using CPU utilization and memory capacity metrics. The test results showed an average CPU time of 0.16 seconds. In Figure 4, one data point exceeds the Upper Control Limit (UCL), indicating a significantly high CPU usage of 3.18 seconds.

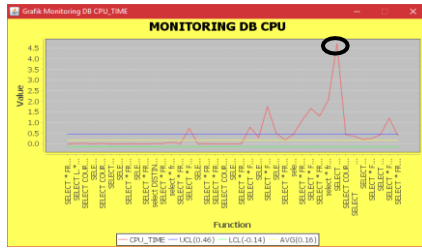


Figure 4. Graph of CPU utilization

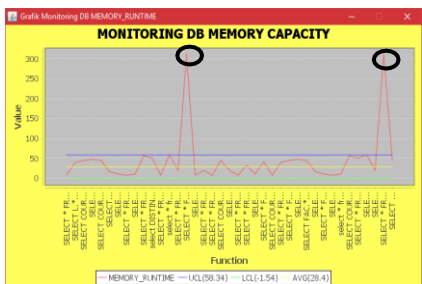


Figure 5. Graph of the result: memory utilization

According to the ISO/IEC 25010 standard, the Resource Utilization sub-characteristic assesses how efficiently a system utilizes the necessary resources to perform its functions. When CPU or memory usage exceeds expected thresholds, it indicates inefficiency and increases the risk of performance degradation under a heavier user load. Based on Table 4, possible contributing factors may include inadequate system maintenance, improper equipment use, or operational inconsistencies (as described in Rules 3, 4, and 5). Therefore, further query profiling and optimization are strongly recommended to ensure stable and efficient system performance.

### 3.3 Capacity

Capacity testing was conducted using memory usage. The results showed an average maximum I/O value of 0.3%, indicating stable performance without significant deviations.

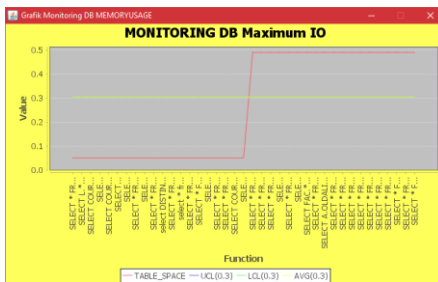


Figure 6. Graph of the result memory usage

However, the graph in Figure 6 illustrates that while the values remain under the Upper Control Limit (UCL), some points approach the threshold, suggesting potential strain under heavier loads.

The Capacity sub-characteristic in ISO/IEC 25010 refers to the extent to which a system can maintain optimal performance under maximum workload conditions. Although Nelson Rules detected no anomalies, the measurement results suggest that adjustments in resource allocation should be considered. This is especially important for anticipating future scalability needs. Strategies such as vertical scaling (adding more resources to the existing server) or horizontal scaling (adding more nodes) may help ensure consistent system performance under increasing demand.

Based on the evaluation and analysis of performance quality measurements, the test results indicate that most queries can deliver optimal response times under light to moderate loads. However, some queries experience significant spikes in response time as user load increases. The components that influence database performance are summarized in Table 3, based on the success percentage for each sub-characteristic.

Table 5. Success Percentage of Measurement

| Sub-Karakteristik    | Persentase                                |
|----------------------|-------------------------------------------|
| Time Behavior        | Defect = $(7/40) \times 100\% = 17,5\%$   |
|                      | Success = $(33/40) \times 100\% = 82,5\%$ |
| Resource Utilization | Defect = $(8/40) \times 100\% = 20\%$     |
|                      | Success = $(32/40) \times 100\% = 80\%$   |
| Capacity             | Defect = $(50/40) \times 100\% = 50\%$    |
|                      | Success = $(50/40) \times 100\% = 50\%$   |

The detected defects represent the number of elements that impact performance, and the table also presents the average percentage results from all measurements conducted on each component.

As shown in Table 5, the analysis results indicate that for time behavior, 82.5% of queries performed successfully, while 17.5% experienced response time issues, suggesting the need for improvement in specific queries. For resource utilization, 80% of the queries ran smoothly; however, 20% still consumed excessive CPU and memory resources, indicating a need for optimization. Regarding capacity, 50% of the queries were successful, while the remaining 50% were hindered due to discrepancies in maximum size and byte allocation for user object segments, which require adjustment.

Based on these results, it is recommended that queries be optimized and

resource allocation adjusted to improve overall system performance.

#### **4. DISCUSSION**

The performance testing results of the database system within an academic information system revealed several significant findings that can serve as a reference for improving system quality. Based on the Time Behavior measurements, the system generally demonstrated good response times, with the majority of queries executing in under one second. However, some queries exhibited higher-than-average response times, indicating potential inefficiencies that require query optimization to enhance response efficiency.

In terms of Resource Utilization, the findings showed that some queries consumed more CPU and memory resources than expected. The average CPU usage was recorded at 0.16 seconds, while memory usage reached 28.4 KB. This high resource consumption suggests inefficiencies in specific queries, which can be improved by optimizing them to reduce memory and CPU usage.

On the other hand, the Capacity measurements indicated that the system could handle a maximum I/O of up to 0.3%, demonstrating strong capacity in managing data loads. There were no significant issues related to capacity, as the maximum I/O values remained within stable limits, suggesting that the system can manage larger data volumes without substantial performance degradation.

Overall, the test results identified several areas requiring improvement, particularly in response time and resource utilization. Therefore, optimizing specific queries and enhancing resource management are recommended to improve the overall system performance.

#### **5. CONCLUSION**

Based on the performance testing results, the database system demonstrated exemplary performance across most queries, with optimal response times under light to moderate loads. However, several queries require further optimization in terms of response time and resource utilization, particularly in terms of CPU and memory usage. Although I/O capacity did not present any issues, the testing identified the need for improved resource management in specific queries to enhance system efficiency. Therefore, query optimization and improved resource management are

expected to enhance overall system performance, particularly under high-load conditions.

#### **6. SUGGESTION**

Future research is expected to expand its scope by testing database systems under more diverse conditions, such as deployment on cloud platforms or distributed systems. Further studies may also focus on automated query optimization analysis and the development of machine learning techniques for dynamic performance prediction and improvement.

Machine learning methods, such as decision trees, reinforcement learning, or neural networks, can be applied to detect and optimize inefficient query patterns automatically. Testing may also include evaluating the impact of horizontal and vertical scaling on database performance, particularly for applications that handle large volumes of data. This research is expected to provide deeper insights into database management in more complex and dynamic environments.

#### **REFERENCES**

- [1] I. Akbar, "Analisis Dan Perancangan Website SMKN 13 Bandung." [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkomTerakreditasiSINTA5>
- [2] M. Berlianti Akbar and T. Sutabri, "Analisis Manajemen Layanan Penerimaan Siswa Baru Menggunakan Framework ITIL Versi 3 Pada Domain Service Design di SMA Negeri 22 Palembang." [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>
- [3] L. Indah Sari, W. Aribowo Probonegoro, I. Atma Luhur, and S. Informasi ISB Atma Luhur, "Analisis Sistem Informasi Akademik Menggunakan Metode FAST Pada SDIT Al-Qudwah Pangkalpinang." [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>
- [4] T. Taipalus, "Database management system performance comparisons: A systematic literature review," *Journal of Systems and Software*, vol. 208, Feb. 2024, doi: 10.1016/j.jss.2023.111872.
- [5] V. Vukovic, V. Vuković, J. Đurković, and J. Trinić, "Defining Performance Criteria and Planning Performance Tests for the Exam Registration Software," 2014. [Online]. Available:

- <https://www.researchgate.net/publication/325366605>
- [6] T. Taipalus, "Database management system performance comparisons: A systematic literature review", doi: 10.48550/arXiv.2301.01095.
  - [7] S. Rochimah, U. Laili Yuhana, and S. Balqis Hidayat, "INTERNATIONAL JOURNAL ON INFORMATICS VISUALIZATION journal homepage: [www.joiv.org/index.php/joiv](http://www.joiv.org/index.php/joiv) INTERNATIONAL JOURNAL ON INFORMATICS VISUALIZATION Software Quality Measurement for Functional Suitability, Performance Efficiency, and Reliability Characteristics Using Analytical Hierarchy Process." [Online]. Available: [www.joiv.org/index.php/joiv](http://www.joiv.org/index.php/joiv)
  - [8] N. Srivastava, U. Kumar, P. Singh, and P. Singh, "Software and Performance Testing Tools," *Journal of Informatics Electrical and Electronics Engineering*, vol. 02, no. 001, pp. 1–12, 2021, doi: 10.54060/JIEEE/002.01.001.
  - [9] E. Peters and G. K. Aggrey, "An ISO 25010 based quality model for ERP systems," *Advances in Science, Technology and Engineering Systems*, vol. 5, no. 2, pp. 578–583, Apr. 2020, doi: 10.25046/aj050272.
  - [10] R. Mansharamani, A. Khanapurkar, B. Mathew, and R. Subramanyan, "Performance testing: Far from steady state," in *Proceedings - International Computer Software and Applications Conference*, 2010, pp. 341–346. doi: 10.1109/COMPSACW.2010.66.
  - [11] N. Husufa and I. Prihandi, "Optimizing JMeter on Performance Testing Using the Bulk Data Method," *Journal of Information Systems and Informatics*, vol. 3, no. 1, 2022, [Online]. Available: <http://journal-isi.org/index.php/isi>

## Fuzzy-Driven Adaptive NPC Behavior in a Meme-Based Platformer Game for Android Mobile

Encep Sayid Amrulloh<sup>1</sup>, Rio Andriyat Krisdiawan<sup>\*2</sup>, Iwan Lesmana<sup>3</sup>

<sup>1,2,3</sup> Department of Computer Science, Faculty of Computer Science, Kuningan University

E-mail: [20200810105@uniku.ac.id](mailto:20200810105@uniku.ac.id), <sup>\*</sup>[rioandriyat@uniku.ac.id](mailto:rioandriyat@uniku.ac.id), <sup>3</sup>[iwanlesmana@uniku.ac.id](mailto:iwanlesmana@uniku.ac.id)

### Abstract

*Game-based applications are increasingly used beyond entertainment to deliver adaptive, engaging user experiences. Yet, many mobile platformer games still rely on static enemy behaviors, leading to repetitive gameplay. This study introduces Pepe the Ponderland Warrior, a 2D platformer for Android that incorporates culturally relevant meme characters and dynamic NPC behavior using fuzzy logic. Developed with the Game Development Life Cycle (GDLC), the game uses the Fuzzy Sugeno inference system to adapt NPC responses based on player distance, health, and damage received. UML modeling guided the system design, while testing included black-box, white-box, and User Acceptance Testing (UAT). The fuzzy-based system enabled real-time, context-aware NPC decisions, creating more varied and challenging gameplay. The game passed functional and logical testing, with UAT from 30 users producing a high feasibility score of 81.2%, reflecting satisfaction in design, gameplay, and difficulty balance. By integrating fuzzy logic with meme-inspired content, this study offers a novel and efficient AI approach for mobile games, highlighting potential for expansion across platforms and with more adaptive inputs.*

**Keywords**—fuzzy Sugeno, adaptive NPC, Android game, 2D-platformer, pepe the frog, GDLC.

Submission: June 19, 2025

Accepted: July 05, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.440>

### 1. INTRODUCTION

Games are interactive systems that allow players to engage in resolving conflicts or challenges through predefined rules, resulting in achievements such as scores, levels, or other objectives [1]. One of the most enduring and popular game genres is the platformer, characterized by player-controlled characters that jump, climb, and navigate obstacles in a two-dimensional (2D) environment with varying terrain [2], [3], [4], [5]. 2D platformer games offer several advantages, including relatively low production costs compared to 3D games, while still providing challenging and addictive gameplay, particularly on mobile platforms [2], [3].

In modern game marketing practices, the use of characters drawn from popular culture, including internet memes, has proven effective in attracting player attention and fostering emotional engagement. One such widely recognized meme character is Pepe the Frog, an iconic figure originating from Matt Furie's Boy's Club comic, which has proliferated across various digital platforms since 2008. Survey data collected from respondents indicate that Pepe possesses strong visual and emotional appeal, making it a promising candidate for use as a main character in a game. However, a search of the Google Play Store reveals that the utilization of

this character in digital games remains minimal and has not been optimally developed, especially within the platformer genre. This gap presents a research opportunity to develop a game featuring Pepe that incorporates a more adaptive and dynamic gameplay approach.

With the rapid advancement of artificial intelligence (AI), integrating AI technologies into the behavior design of non-player characters (NPCs) in games has become increasingly common, aiming to enhance player experience. One of the most frequently applied AI methods is fuzzy logic, particularly the Sugeno fuzzy inference system, due to its ability to manage uncertainty in variables such as object distance, health level, and damage received [6], [7], [8], [9], [10]. In a study conducted by Paraschos and Koulouriotis (2025), fuzzy logic was used in a Dynamic Difficulty Adjustment (DDA) system in an AI-based shooter game to maintain a balance between challenge and player comfort across different play styles [11]. Similarly, Kevin et al. (2025), in the Journal of AI and Engineering Applications, implemented 27 fuzzy rules to generate NPC behaviors such as "evade", "defend", and "wait", based on player and environmental conditions in a 2D adventure game [12]. Another example can be found in the work of Krisdiawan et al. (2023), who applied the Sugeno fuzzy algorithm in a mobile-based game for dyslexia rehabilitation, adjusting

difficulty levels according to individual learning progress [7].

Nonetheless, several limitations remain in these existing studies. First, most implementations of fuzzy logic in games focus only on enemy behavior in single scenarios or levels, without considering the complex and dynamic environmental conditions typically found in platformer games. Second, there is a lack of research that integrates popular cultural elements, such as meme characters like Pepe the Frog, as a central gameplay component to enhance players' emotional engagement. Initial surveys conducted by the authors indicate that this character holds significant visual and cultural appeal among teenagers and young adults.

This study aims to address the above gaps by proposing a novel design and development of Pepe the Ponderland Warrior, a 2D platformer game for Android devices. This game uniquely combines the use of the culturally iconic meme character Pepe the Frog with the adaptive capabilities of the Sugeno fuzzy algorithm to regulate NPC behavior. NPC actions are adapted based on parameters such as distance, health, damage, and movement speed; all processed through a Sugeno fuzzy inference system. The game is structured into four progressively challenging thematic levels Cave, Snow, Jungle, and Swamp allowing for the evaluation of AI adaptability under different gameplay scenarios. It is optimized for Android devices with a minimum SDK version of 6 and designed as an offline, single-player game to ensure broad accessibility.

The objectives of this research are: (1) to design and develop the Pepe the Ponderland Warrior 2D platformer game for Android, featuring a popular meme character; and (2) to implement the Sugeno fuzzy algorithm for adaptive and efficient NPC behavior regulation. This study is expected to contribute to the development of AI-based games that not only adapt to individual player styles but also incorporate internet cultural elements as a novel means of enhancing player engagement.

## 2. RESEARCH METHOD

### 1. Data Collection Methods

This research adopts a qualitative-descriptive approach through a software engineering perspective using a system development study. To obtain valid and relevant data for designing Pepe The Ponderland Warrior game, the

researcher applied the following data collection techniques:

#### a. Literature Review

A comprehensive literature review was conducted to understand and build the theoretical foundation supporting system development. Sources included scholarly publications related to interactive game design, non-player character (NPC) behavior modeling in 2D games, and the application of fuzzy logic—particularly the Sugeno model. References were gathered from accredited national journals, international indexed journals (IEEE, Scopus), proceedings, and previous research.

#### b. Questionnaire

A questionnaire was administered via Google Form to assess user interest in Pepe the Frog as a main game character. Respondents were primarily casual and mobile game players. The goal was to evaluate perceptions on character appeal, gameplay preferences, and adaptive game mechanics expectations.

#### c. Documentation and System Artifacts

Supporting data included interface mockups, initial storyboards, Game Design Document (GDD), and prototype testing results. Visual documentation served as validation of design decisions and provided evidence during testing and refinement phases of the system.

## 2. Game Development Method

This research employed the **Game Development Life Cycle (GDLC)** model, which provides an iterative and structured framework for game development. GDLC consists of six primary phases [13] :

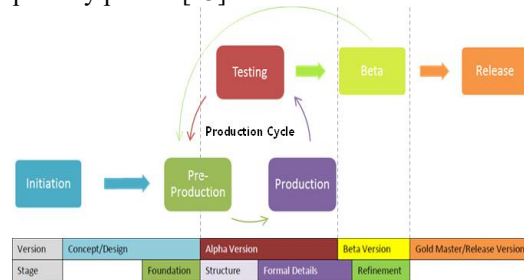


Figure 1. Game Development Life Cycle (GDLC) [3], [13].

a. **Initiation** – Gathering user requirements and ideas through literature review and questionnaire analysis.

b. **Pre-production** – Designing the *Game Design Document (GDD)*, initial interface sketches, and modeling system logic using Unified Modeling Language (UML), including use case, activity, sequence, and class diagrams.

- c. **Production** – Developing the game in Unity with C# programming, integrating the **Fuzzy Sugeno algorithm** to control NPC behavior.
- d. **Testing** – Conducting **black-box testing** for system functionality and **white-box testing** for validating the internal logic of the fuzzy inference system, including cyclomatic complexity analysis.
- e. **Beta** – External user testing using **User Acceptance Testing (UAT)** to gather feedback on gameplay and adaptive behavior.
- f. **Release** – Final deployment of the game on Android platform (minimum SDK 6) as an offline single-player game.

The GDLC model was chosen for its capacity to support dynamic, user-centered development of interactive and intelligent game systems [13].

### 3. Problem Solving Method (Fuzzy Sugeno Implementation)

To address the challenge of static and non-adaptive NPC behavior, the **Fuzzy Sugeno algorithm** was applied as a decision-making system based on fuzzy inference logic.

#### a. Justification for Sugeno Model

Fuzzy Sugeno was selected over alternative models (e.g., Mamdani or machine learning approaches) for the following reasons [11], [12]:

- **Computational efficiency:** Sugeno produces crisp outputs in the form of real-valued linear functions, eliminating the need for complex defuzzification—ideal for resource-constrained Android devices.
- **Handling of uncertain input:** The system efficiently processes uncertain variables such as distance, life points, and damage using fuzzy linguistic sets.
- **Precision in decision-making:** Unlike Mamdani's linguistic outputs, Sugeno supports precise numeric computations, suitable for real-time action control.

#### b. Mathematical Formulation of the Sugeno Model

The fuzzy rules are formulated as:

$$\text{If } x_1 \text{ is } A_1 \text{ and } x_2 \text{ is } A_2 \text{ and } x_3 \text{ is } A_3, \text{ then } y = a_1x_1 + a_2x_2 + a_3x_3 + c \quad (1)$$

Where:

- $x_1, x_2, x_3$  are input variables (distance, life, damage)
- $A_1, A_2, A_3$  are fuzzy linguistic sets (e.g., near, medium, far; low, medium, high)
- $y$  is the output (NPC action), represented by a linear function

The following is a flowchart of the Sugeno fuzzy algorithm:

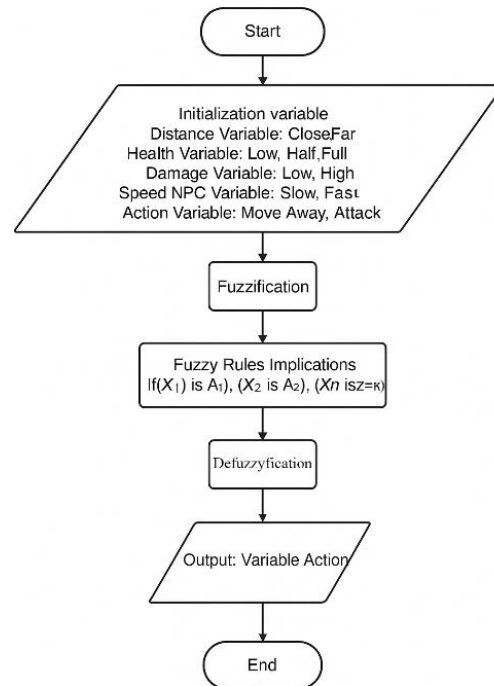


Figure 2. Flowchart of the Sugeno Fuzzy Algorithm

#### Inference steps include:

1. **Fuzzification** – Converting crisp input values into membership degrees.
2. **Rule Evaluation** – Calculating the firing strength using min or product operators.
3. **Output Calculation** – Computing the output for each rule based on its linear function.
4. **Aggregation (Defuzzification)** – Weighted average method:

$$y \text{ Output} = \frac{\sum(w_i \cdot y_i)}{\sum w_i} \quad (2)$$

Where:

- $w_i$  is the **firing strength** of the  $i$ -th rule (typically calculated using fuzzy logic operators such as minimum or product of membership values),
- $y_i$  is the **output** of the  $i$ -th rule, which in Sugeno models is a **constant or linear function** of the input variables (e.g.,  $y=a_1x_1+a_2x_2+$ ).

### 3. RESEARCH RESULTS

#### 3.1 Game Design Document (GDD)

##### 1. Game Layout Chart

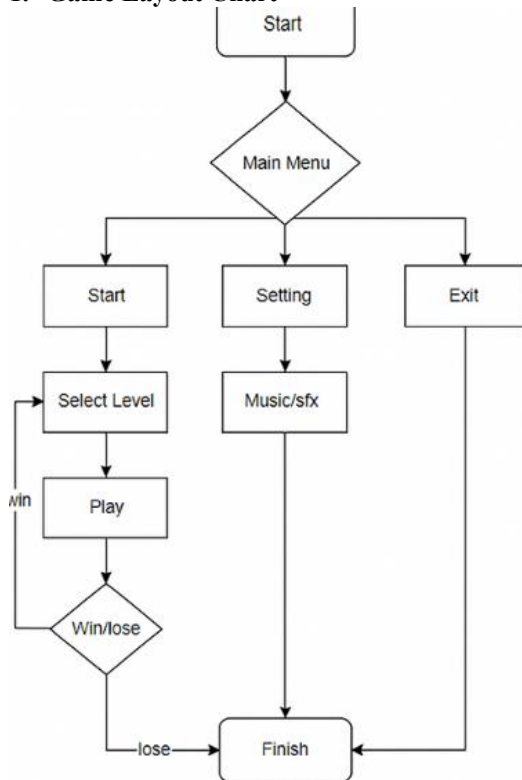


Figure 3. Game layout chart

The sequence of activities in Figure 3 can be described as follows:

- 1) The game starts from the **Start** node.
- 2) The **Main Menu** provides three options: **Start**, **Settings**, and **Exit**.
- 3) Selecting **Start** will lead to the **Level Selection** screen.
- 4) After choosing a level, the game proceeds to the **Play** stage.
- 5) The game then evaluates the outcome:
  - a. If the player **wins**, they proceed to the next level.
  - b. If the player **loses**, the game ends and returns to the **Finish** state.
- 6) Selecting **Settings** opens the **Music/SFX configuration screen**, where audio settings can be adjusted.
- 7) Selecting **Exit** will close the application and return to the device's home screen.

##### 2. StoryBoard Game

The storyboard game Pepe the Ponderland Warrior is presented in tables 1 to table 3.

Table 1. Story Board Scene 1

| Scene 1   | Visual: Main Menu                                                                                  |
|-----------|----------------------------------------------------------------------------------------------------|
|           |                                                                                                    |
| Audio     | <b>Backsound:</b> Soft accompaniment music.<br><b>Sound FX:</b> Sound of click buttons             |
| Storyline | In the main menu, players will be given action buttons to select menus and display the game title. |

Table 2. Story Board Scene 2

| Scene 2   | Visual: Play Game                                                                                                                                                                       |
|-----------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|           |                                                                                                                                                                                         |
| Audio     | <b>Backsound:</b> Cheerful accompaniment music.<br><b>Sound FX:</b> Sound of click buttons                                                                                              |
| Storyline | The game page displays the player's health status as lives. Then the player walks to the finish line by pressing the left, right, and jump buttons while passing obstacles and enemies. |

Table 3. Story Board Scene 3

| Scene 3   | Visual: Mission Complete                                                                                                       |
|-----------|--------------------------------------------------------------------------------------------------------------------------------|
|           |                                                                                                                                |
| Audio     | <b>Backsound:</b> Happy and cheerful music.<br><b>Sound FX:</b> Sound of click buttons                                         |
| Storyline | After the player has completed the game, there will be a menu to return, restart the game level, or proceed to the next level. |

##### 3. Game Object

Game object yang terdapat pada pepe the ponderland warrior dapat dilihat pada table 4.

Table 4. Game Object

| Game Object                                                                       | Description         | Action                            |
|-----------------------------------------------------------------------------------|---------------------|-----------------------------------|
|  | Player pepe         | Move forward, backward, up, down. |
|  | Bat NPC (Enemies)   | Fly forward, backward, up, down.  |
|  | Slime NPC (Enemies) | Move forward, backward, up, down. |
|  | Obstacle            | Spinning                          |
|  | Flag finish game    | Object finish game                |

#### 4. Implementation Fuzzy in Game

A decision is made when the distance is 8, the number of health is 70, and the damage caused is 10.

##### Formation of Fuzzy sets

1. Distance Variable
  - a. Close: 0-10
  - b. Far: 5-15

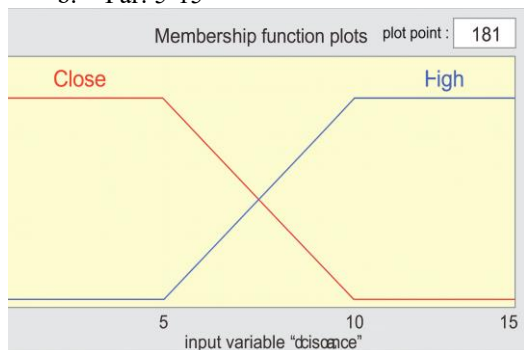


Figure 4. Degree of Membership Distance

$$\mu \text{ Distance close}[5] = \text{distance} > 10 \quad (3)$$

$$\mu \text{ Distance far}[5] = \text{distance} < 5$$

2. Health Variable
  - a. Low: 0-50
  - b. Half: 25-75
  - c. Full: 50-100

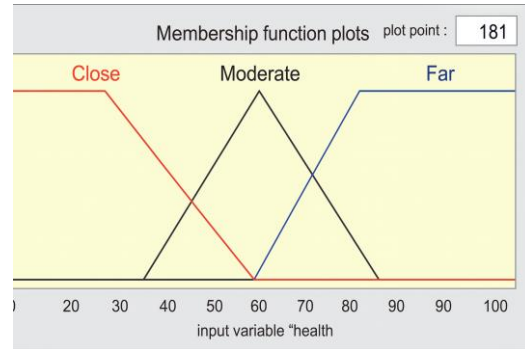


Figure 5. Degree of Membership of the Health Variable

$$\mu \text{ Health low}[45] = \text{health} < 25$$

$$\mu \text{ Health half}[45] = \text{health} \leq 25; \text{health} \geq 75 \quad (4)$$

$$\mu \text{ Health full}[45] = \text{health} > 75$$

##### 3. Damage Variable

- a. low: 0-12
- b. high: 6-20



Figure 6. Degree of Membership of the Damage Variable

$$\mu \text{ damage low}[15] = \text{damage} < 6$$

$$\mu \text{ damage high}[15] = \text{damage} > 12 \quad (5)$$

##### Fuzzy Rules

Table 5. Fuzzy Rules

| R   | J     | NM   | DM          | A         |
|-----|-------|------|-------------|-----------|
| R1  | Close | Low  | <i>low</i>  | Move away |
| R2  | Close | Low  | <i>high</i> | Move away |
| R3  | Close | Low  | <i>low</i>  | Move away |
| R4  | Close | Low  | <i>high</i> | Move away |
| R5  | Close | Half | <i>low</i>  | Move away |
| R6  | Close | Half | <i>high</i> | Attack    |
| R7  | Close | Half | <i>low</i>  | Move away |
| R8  | Close | Half | <i>high</i> | Attack    |
| R9  | Close | Full | <i>low</i>  | Attack    |
| R10 | Close | Full | <i>high</i> | Attack    |
| R11 | Close | Full | <i>low</i>  | Attack    |
| R12 | Close | Full | <i>high</i> | Attack    |

|     |     |      |      |           |
|-----|-----|------|------|-----------|
| R13 | Far | Low  | low  | Move away |
| R14 | Far | Low  | high | Move away |
| R15 | Far | Low  | low  | Move away |
| R16 | Far | Low  | high | Move away |
| R17 | Far | Half | low  | Move away |
| R18 | Far | Half | high | Attack    |
| R19 | Far | Half | low  | Move away |
| R20 | Far | Half | high | Attack    |
| R21 | Far | Full | low  | Attack    |
| R22 | Far | Full | high | Attack    |
| R23 | Far | Full | low  | Attack    |
| R24 | Far | Full | high | Attack    |

Table description:

- 1) R = Rules
- 2) J = Distance
- 3) NM = Health
- 4) DM = Damage
- 5) A = Action

#### Defuzzification

Determine the crisp values of distance, number of health, and damage produced into fuzzy sets and determine their membership degrees. Distance 10, number of Health 80, and damage produced 10.

a. Distance

$$\alpha_{Close} = \mu_{Close}(8) = \frac{10-8}{8-5} = 0,66 \quad (6)$$

$$\alpha_{Far} = \mu_{Far}(8) = \frac{8-5}{10-5} = 0,6$$

b. Health

$$\alpha_{Half} = \mu_{Half}(70) = \frac{75-70}{70-50} = 0,25 \quad (7)$$

$$\alpha_{Full} = \mu_{Full}(70) = \frac{70-50}{75-50} = 0,8$$

c. Damage

$$\alpha_{Low} = \mu_{Low}(10) = \frac{12-10}{10-6} = 0,5 \quad (8)$$

$$\alpha_{High} = \mu_{High}(10) = \frac{10-6}{12-6} = 0,66$$

#### Implications

$$\begin{aligned} \alpha_{Close} \cap \alpha_{Half} \cap \alpha_{Low} &= 0,25 \times 50 \\ \alpha_{Close} \cap \alpha_{Half} \cap \alpha_{High} &= 0,25 \times 50 \\ \alpha_{Close} \cap \alpha_{Full} \cap \alpha_{Low} &= 0,5 \times 100 \\ \alpha_{Close} \cap \alpha_{Full} \cap \alpha_{High} &= 0,66 \times 100 \\ \alpha_{Far} \cap \alpha_{Half} \cap \alpha_{Low} &= 0,25 \times 50 \\ \alpha_{Far} \cap \alpha_{Half} \cap \alpha_{High} &= 0,25 \times 50 \\ \alpha_{Far} \cap \alpha_{Full} \cap \alpha_{Low} &= 0,5 \times 100 \\ \alpha_{Far} \cap \alpha_{Full} \cap \alpha_{High} &= 0,6 \times 100 \end{aligned}$$

#### Defuzzifikasi

Calculate defuzzification using the average formula.

$$Decision = \frac{\sum \alpha_i z_i}{\sum \alpha_i} \quad (9)$$

$$z = \frac{0,25(50)+0,25(50)+0,5(100)+0,66(100)+0,25(50)+0,25(50)+0,5(100)+0,6(100)}{0,25+0,25+0,5+0,66+0,25+0,25+0,5+0,6} \quad (10)$$

$$z = \frac{276}{3,26} = 84,66$$

Crisp decision index = 84,66

Fuzzy decision index = 80% Attack, 20% Move away

## 5. Use Case Diagram System Game

a) Use Case Design

A use case diagram is a UML model used to briefly describe who uses the system and what can be done with the system. The use cases for the application to be built [11] are illustrated in Figure 7.

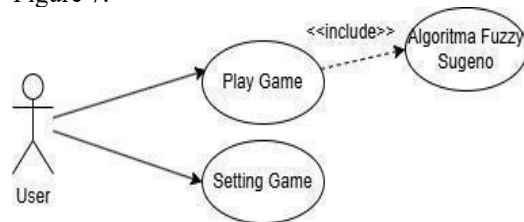


Figure 7 Usecase diagram

A case that illustrates the interaction between users (identified as “Users”) and the main functionality of a game system. Players, as the main actors, can perform two different actions or use cases: “Play Game” (Start a New Game) and “Game Settings.” This diagram provides a high-level overview of the features available to players in this game.

b) Class Diagram

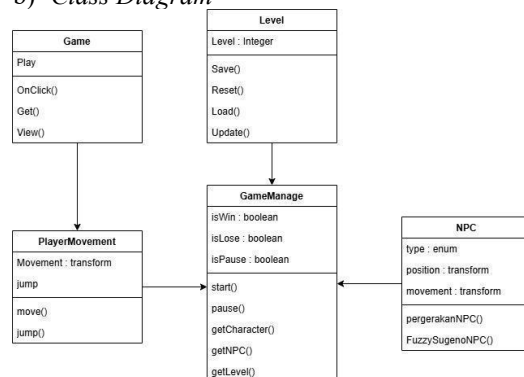


Figure 8. Class Diagram.

## 6. Interface Game

a) Main menu Interface

Figure 9 displays the Main Menu Interface, which serves as the initial entry point of the game. It contains three primary options: Start, Settings, and Exit. This interface provides intuitive navigation, allowing players to begin

gameplay, adjust audio preferences, or exit the application.



Figure 9. Main menu interface

b) *Play Interface*

Figure 10 illustrates the Gameplay Interface, where players control the main character using directional buttons: left, right, and up. The interface also features a health bar to display the player's current life status and a pause button that allows the game to be temporarily halted. This layout is designed for accessibility and optimized interaction on touchscreen devices.

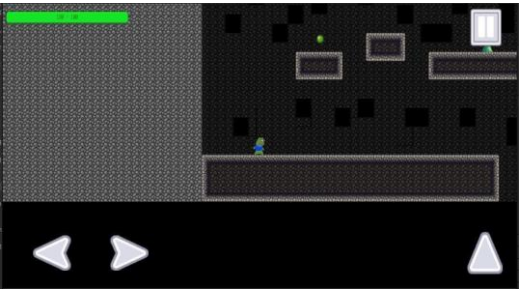


Figure 10. Play Interface

c) *Algorithms in game*

Figure 11 shows the implementation of the Fuzzy Sugeno algorithm in determining the behavior of non-player characters (NPCs). NPC actions—such as attacking, retreating, or remaining idle—are dynamically adapted based on input variables including distance to the player, health level, and damage received. This intelligent behavior mechanism aligns with the study's objective of providing an adaptive and engaging player experience.



Figure 11. Algorithms in games

d) *Complete Interface*

Figure 12 shows the **Completion Interface**, which appears upon the successful completion of

a level. It includes three options: **Back**, **Retry**, and **Next Level**, allowing players to return to the main menu, replay the current level, or proceed to the next stage. This interface supports replayability and user control over game progression.



Figure 14. Complete Interface

## 4. DISCUSSION

### 4.1 Alpha Testing

#### 1. Blackbox Testing

Black box testing was conducted to evaluate the functionality of the game application from the user's perspective, without examining the internal code structure. This method focuses on ensuring that the system responds correctly to various inputs and user interactions, in alignment with the functional requirements.

Table 6 summarizes the black box test results for each main feature of the game interface:

| Table 6. Black box test results |                         |                                     |                                         |                                            |                |
|---------------------------------|-------------------------|-------------------------------------|-----------------------------------------|--------------------------------------------|----------------|
| N<br>o                          | Test<br>Scena<br>rio    | Test<br>Case                        | Expect<br>ed<br>Result                  | Actua<br>l<br>Resul<br>t                   | Conclu<br>sion |
| 1                               | Start<br>Game           | Press<br>Start<br>butto<br>n        | Game<br>starts<br>and<br>plays<br>level | Game<br>starts<br>as<br>expec<br>ted       | Passed         |
| 2                               | Comp<br>lete<br>Game    | Finis<br>h a<br>level               | Show<br>comple<br>te<br>screen          | Comp<br>lete<br>screen<br>displa<br>yed    | Passed         |
| 3                               | Pause<br>Game           | Press<br>Paus<br>e<br>butto<br>n    | Show<br>pause<br>menu                   | Pause<br>menu<br>displa<br>yed             | Passed         |
| 4                               | Lose<br>Game            | Playe<br>r defea<br>ted             | Show<br>fail<br>screen                  | Fail<br>screen<br>displa<br>yed            | Passed         |
| 5                               | Acces<br>s Settin<br>gs | Press<br>Setti<br>ngs<br>butto<br>n | Displa<br>y music/<br>SFX<br>options    | Settin<br>gs displa<br>yed<br>prope<br>rly | Passed         |

All tested features functioned according to the intended design, confirming that the interface elements and core gameplay mechanics were implemented correctly.

## 2. Whitebox Testing

White box testing was used to examine the internal logic of the game, particularly the implementation of the **Fuzzy Sugeno algorithm** in the NPC decision-making system. Cyclomatic complexity analysis was applied to measure the logical paths within the fuzzy evaluation method.

Table 7. White Box Testing of Fuzzy Sugeno Algorithm

| No. | Code program                                                                                                                                                                                                                                                                                                                                                                                                 |
|-----|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1   | <pre>public class FuzzySugenoSystem : MonoBehaviour {     public List&lt;FuzzyVariable&gt; inputVariables = new List&lt;FuzzyVariable&gt;();     public List&lt;SugenoRule&gt; rules = new List&lt;SugenoRule&gt;(); }</pre>                                                                                                                                                                                 |
| 2   | <pre>public float Evaluate(params float[] inputValues) {     List&lt;float&gt; firingStrengths = new List&lt;float&gt;();     List&lt;float&gt; ruleOutputs = new List&lt;float&gt;(); }</pre>                                                                                                                                                                                                               |
| 3   | <pre>foreach (var rule in rules) {     float firingStrength = 1f; }</pre>                                                                                                                                                                                                                                                                                                                                    |
| 4   | <pre>foreach (var condition in rule.conditions){</pre>                                                                                                                                                                                                                                                                                                                                                       |
| 5   | <pre>foreach (var variable in inputVariables){</pre>                                                                                                                                                                                                                                                                                                                                                         |
| 6   | <pre>if (variable.name == condition.variableName) {     int inputIndex = inputVariables.IndexOf(variable); }</pre>                                                                                                                                                                                                                                                                                           |
| 7   | <pre>if (inputIndex &lt; inputValues.Length) {     float membership = variable.GetMembership(         condition.fuzzySetName,         inputValues[inputIndex]);     firingStrength = Mathf.Min(firingStrength, membership); } } } firingStrengths.Add(firingStrength); float output = rule.CalculatedOutput(inputValues); ruleOutputs.Add(output); } float sumWeightOutput = 0f; float sumWeight = 0f;</pre> |
| 8   | <pre>for (int i = 0; i &lt; rules.Count; i++) {     Debug.Log("rule " + rules[i].name + " : " + firingStrengths[i] + " " + ruleOutputs[i]);     sumWeightOutput += firingStrengths[i] * ruleOutputs[i];     sumWeight += firingStrengths[i]; }</pre>                                                                                                                                                         |
| 9   | <pre>return sumWeight &gt; 0 ? sumWeightOutput / sumWeight : 0f; }</pre>                                                                                                                                                                                                                                                                                                                                     |

Based on Table 7, Flowgraph Cyclomatic Complexity was obtained to conduct path-based testing by calculating cyclomatic complexity. The following is an image of Flowgraph Cyclomatic Complexity.

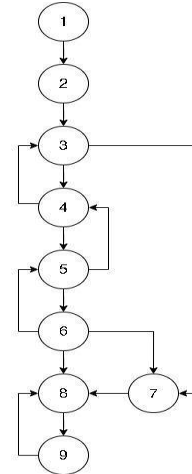


Figure 15. Flowgraph Cyclomatic Complexity Calculation of cyclomatic complexity, using the formula:

$$V(G) = E - N + 2 \quad (11)$$

From the previous flowgraph, the following values can be determined:

E = 14; N = 9;

After entering these values into the formula, the result is:

$$V(G) = E - N + 2$$

$$V(G) = 14 - 9 + 2 = 7$$

From the cyclomatic complexity calculation, there are 7 independent paths, namely:

- Path 1: 1 → 2 → 3 → 4 → 5 → 6 → 8 → 9
- Path 2: 1 → 2 → 3 → 4 → 5 → 6 → 7 → 8 → 9
- Path 3: 1 → 2 → 3 → 4 → 5 → 3 → 4 → 5 → 6 → 8 → 9
- Path 4: 1 → 2 → 3 → 4 → 3 → 4 → 5 → 6 → 8 → 9
- Path 5: 1 → 2 → 3 → 4 → 5 → 6 → 8 → 9 → 5 → 6 → 8 → 9
- Path 6: 1 → 2 → 3 → 4 → 5 → 6 → 7 → 3 → 4 → 5 → 6 → 8 → 9
- Path 7: 1 → 2 → 3 → 4 → 5 → 6 → 8 → 9 → 3 → 4 → 5 → 6 → 7 → 8 → 9

The result indicates that the method contains 7 independent logical paths. These paths were traced and tested individually to ensure complete path coverage and logical correctness in decision output. The implementation was verified to generate consistent and context-sensitive NPC behavior during gameplay.

#### 4.2 Beta Testing (UAT Testing)

The User Acceptance Test (UAT) was conducted with 30 respondents to evaluate the game's usability, interface clarity, character control, and the adaptiveness of NPC behavior. Participants assessed six key indicators using a Likert scale: Strongly Agree (5), Agree (4), Neutral (3), Disagree (2), and Strongly Disagree (1). Table 8 presents the total weighted score for each question based on the assigned weights. The maximum possible score was 900, calculated from 30 respondents  $\times$  6 questions  $\times$  5 points.

Table 8. UAT Questionnaire Result

| Indicator                                              | Score            |
|--------------------------------------------------------|------------------|
| Visual appeal of the main menu                         | 140              |
| Ease of understanding the game menu                    | 132              |
| Character Control                                      | 142              |
| Relevance of Pepe the Frog as a character              | 139              |
| Balance in gameplay difficulty                         | 137              |
| Adaptiveness of enemy behavior (challenge vs fairness) | 141              |
| <b>Total</b>                                           | <b>731 / 900</b> |

The overall eligibility percentage is calculated as:

$$Eligibility(\%) = \frac{731}{900} = 81.2\%$$

This result indicates that the game *Pepe The Ponderland Warrior* is considered **acceptable and well-received** by users. Most respondents reported that the game interface was intuitive, character control was accessible, and the implementation of NPC behavior using the Fuzzy Sugeno algorithm created a fair yet challenging gameplay experience. More detailed visualization can be seen in Figure 16.

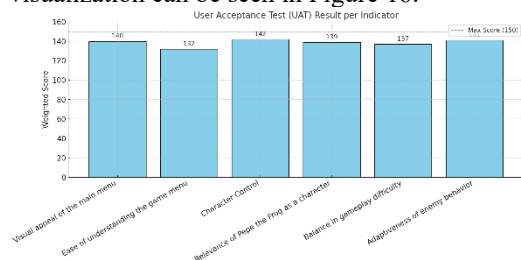


Figure 16. Visualizes User Acceptance Test (UAT) Result Per Indicator

#### 5. CONCLUSION

This study has successfully designed and developed a 2D platformer game titled *Pepe The Ponderland Warrior*, targeting the Android

platform, with a key focus on implementing the Fuzzy Sugeno algorithm to control adaptive NPC behavior. The development process followed the Game Development Life Cycle (GDLC) model, ensuring iterative design, prototyping, testing, and release.

The main contribution of this research lies in its novel combination of adaptive AI behavior with cultural engagement using Pepe the Frog, a popular internet meme, as the central character. This approach not only introduces unique visual appeal but also addresses the market gap—the absence of meme-based platformer games that integrate intelligent difficulty adjustment mechanisms.

From the AI perspective, the implementation of the Fuzzy Sugeno inference system enabled NPCs to respond dynamically to real-time game conditions such as distance to the player, health status, and damage received. Compared to traditional rule-based or static behavior models, this approach produced more responsive, context-aware NPCs, leading to a richer gameplay experience.

The system was validated through functional testing (black-box and white-box) and User Acceptance Testing (UAT) involving 30 respondents. The game achieved a feasibility score of 81.2%, indicating that the majority of users perceived the gameplay as engaging, balanced, and enjoyable.

In conclusion, this research demonstrates the effectiveness of combining fuzzy logic with platformer game mechanics to create personalized, adaptive, and immersive player experiences. It provides a strong foundation for the future development of AI-driven game systems and serves as a reference for incorporating fuzzy decision-making models in resource-constrained mobile environments.

#### 6. SUGGESTIONS

Based on the results and findings of this research, several suggestions can be offered to guide future development and research in related areas. First, the game *Pepe The Ponderland Warrior* can be further enhanced by expanding the number of levels, integrating more complex environmental elements (e.g., moving platforms, traps, or dynamic lighting), and introducing additional NPC behavior types beyond simple attack or retreat patterns. Second, the implementation of the Fuzzy Sugeno algorithm can be extended by incorporating more input variables such as player performance history or reaction time to achieve a finer level of difficulty personalization. Third, cross-platform compatibility (e.g., iOS or web-

based deployment) should be considered to reach broader user segments. Lastly, future studies are encouraged to explore comparative analysis between Fuzzy Sugeno and other adaptive AI methods such as reinforcement learning or hybrid neuro-fuzzy systems to evaluate performance, scalability, and user engagement in various game genres.

## REFERENCES

- [1] K. S. , & Z. E. Tekinbas, *Rules of play: Game design fundamentals*. MIT press, 2003. Accessed: Jul. 04, 2025. [Online]. Available: [https://books.google.co.id/books?hl=en&lr=&id=YrT4DwAAQBAJ&oi=fnd&pg=PR13&dq=related:AAzhZBxv6GkJ:scholar.google.com/&ots=fqUCah7omY&sig=\\_FV4eX9XQAZHhPQHcDZ0o5TPq-U&redir\\_esc=y#v=onepage&q&f=false](https://books.google.co.id/books?hl=en&lr=&id=YrT4DwAAQBAJ&oi=fnd&pg=PR13&dq=related:AAzhZBxv6GkJ:scholar.google.com/&ots=fqUCah7omY&sig=_FV4eX9XQAZHhPQHcDZ0o5TPq-U&redir_esc=y#v=onepage&q&f=false)
- [2] Octodinata, S. B., and & H. I., “Rancang Bangun Game Edukasi Genre Platformer 2D Berbasis Android,” *Jurnal RESTI*, vol. 2, no. 7, pp. 318–325, 2023.
- [3] R. A. Krisdiawan and Darsanto, “PENERAPAN MODEL PENGEMBANGAN GAME GDLC (GAME DEVELOPMENT LIFE CYCLE) DALAM MEMBANGUN GAME PLATFORM BERBASIS MOBILE,” *TEKNOKOM*, Mar. 2019, doi: <https://doi.org/10.31943/teknokom.v2i1.33>.
- [4] A. F. Kusumah and R. Kurniawan, “Development of a Serious Game as an Educational Tool for Obesity Prevention,” *Jurnal Sains, Nalar, dan Aplikasi Teknologi Informasi*, vol. 3, no. 2, pp. 58–63, Jan. 2024, doi: <https://doi.org/10.20885/snati.v3.i2.32>.
- [5] F. A. Jellyanto, I. Kuswardayan, and N. Suciati, “Rancang Bangun Pembangkit World Dinamis menggunakan Algoritma Recursive Backtracking pada Game 2D Platformer ‘Mine Meander,’” *Jurnal Teknik ITS*, vol. 5, no. 2, Oct. 2016, doi: <https://doi.org/10.12962/j23373539.v5i2.19364>.
- [6] R. A. Krisdiawan, M. F. Rohmana, and A. Permana, “Pembuatan Game Runaway From Culik Dengan Algoritma Fuzzy Mamdani,” *Buffer Informatika*, vol. 6, no. 1, pp. 33–40, Jun. 2020, doi: <https://doi.org/10.25134/buffer.v6i1.2623>.
- [7] R. A. Krisdiawan, T. Sugiharto, and N. Nugraha, “Terapi Rehabilitasi Dyslexia Berbasis Mobile Game dengan Dynamic Difficulty Adjustment (DDA) Menggunakan Algoritma Fuzzy Sugeno,” *Jurnal Teknologi Informatika dan Komputer*, vol. 9, no. 2, pp. 975–991, Sep. 2023, doi: <https://doi.org/10.37012/jtik.v9i2.1837>.
- [8] R. D. Anggraeni, F. Nugroho, and J. N. Fadila, “Implementasi Algoritma Fuzzy Sugeno Pada Game Pious Boy,” *Jurnal Informatika: Jurnal Pengembangan IT*, vol. 6, no. 1, pp. 26–28, Jan. 2021, doi: <https://doi.org/10.30591/jpit.v6i1.2321>.
- [9] Z. Zhu and M. Zhu, “Evaluation Method of Performance of Cross-Border e-Commerce System Based on Fuzzy DEA Model,” *Mobile Information Systems*, vol. 2022, 2022, doi: <https://doi.org/10.1155/2022/1456584>.
- [10] U. Nurhasan, “Analisis Perilaku Non Playable Character (Npc) Pada Game Menggunakan Fuzzy Sugeno,” *Techno.Com*, vol. 19, no. 3, pp. 308–320, Aug. 2020, doi: <https://doi.org/10.33633/tc.v19i3.3477>.
- [11] P. D. Paraschos and D. E. Koulouriotis, “Fuzzy Logic-Based Dynamic Difficulty Adjustment for Adaptive Game Environments,” *Electronics (Basel)*, vol. 14, no. 1, p. 146, Jan. 2025, doi: <https://doi.org/10.3390/electronics14010146>.
- [12] Kevin, Octara Pribadi, and Hendri, “Implementation of Fuzzy Logic Algorithm to Improve NPC Decision-Making in 2D Adventure Games Using Unity,” *Journal of Artificial Intelligence and Engineering Applications (JAIEA)*, vol. 4, no. 3, pp. 2376–2381, Jun. 2025, doi: <https://doi.org/10.59934/jaiea.v4i3.1175>.
- [13] R. A. Krisdiawan, “Implementasi Model Pengembangan Sistem GDLC dan Algoritma Linear Congruential Generator Pada Game Puzzle,” *Nuansa Informatika*, vol. 12, no. 2, pp. 1–9, Mar. 2018, doi: <https://doi.org/10.25134/nuansa.v12i2.1634>.

## Design and Implementation of a Performance Dashboard for Public Relations at Telkom University

Feddy Dea Reskyadita<sup>\*1</sup>, Hani Nurrahmi<sup>2</sup>, Daris Maulana<sup>3</sup>, Zehan Triartono<sup>4</sup>, Aditya Zulkarnain<sup>5</sup>, Sidik Prabowo<sup>6</sup>

<sup>1,2,3,4,5,6</sup>Faculty of Informatics, Telkom University

E-mail: <sup>\*1</sup>[feddydr@telkomuniversity.ac.id](mailto:feddydr@telkomuniversity.ac.id), <sup>2</sup>[haninurrahmi@telkomuniversity.ac.id](mailto:haninurrahmi@telkomuniversity.ac.id),  
<sup>3</sup>[darismaulana@telkomuniversity.ac.id](mailto:darismaulana@telkomuniversity.ac.id), <sup>4</sup>[zehantr@telkomuniversity.ac.id](mailto:zehantr@telkomuniversity.ac.id), <sup>5</sup>[adityazulkarnain@telkomuniversity.ac.id](mailto:adityazulkarnain@telkomuniversity.ac.id), <sup>6</sup>[pakwowo@telkomuniversity.ac.id](mailto:pakwowo@telkomuniversity.ac.id)

### Abstract

*The Public Relations and Analytics (PRA) division at Telkom University faces challenges in managing manual processes and integrating data from diverse sources, hindering efficient reputation management. This study developed PANDA (Public Relations and Analytics Dashboard Application), a web-based system with an interactive performance dashboard to optimize PRA operations. Using the Waterfall methodology, the research encompassed requirements analysis, system design, implementation, and initial evaluation. PANDA integrates real-time key performance indicators (KPIs) such as media coverage, social media engagement, website analytics, and enabling data-driven decision-making. The System Usability Scale (SUS) evaluation resulted in a high score of 82 (categorized as 'Good to Excellent'), demonstrating PANDA's effectiveness in enhancing usability and streamlining operational workflows. The system enhances PR performance monitoring and offers a scalable model for educational institutions.*

**Keywords:** Web-based information system, Public Relations, Performance Dashboard, Waterfall, System Usability Scale.

Submission: March 25, 2025

Accepted: July 08, 2025

Published: July 10, 2025

DOI Article : <https://doi.org/10.25134/ilkom.v19i2.377>

## 1. INTRODUCTION

### 1.1 Background

Public relations (PR) is a vital division within an organization, particularly in educational institutions such as Telkom University. As a center of information and communication, PR is responsible for building a positive image and maintaining good relationships with stakeholders, both internal and external.

However, in carrying out its functions, Telkom University's PR faces various challenges, such as difficulty monitoring performance in real-time, lack of data integration from multiple sources, such as social media, news, and activity reports, and ineffectiveness in measuring the impact of PR activities on the university's reputation.

To monitor data changes, a dashboard application provides a more efficient and user-friendly solution [1] [2]. A performance dashboard can present information visually and in real-time, helping PR to track the development of the institution's image and reputation, measure the effectiveness of communication

campaigns and analyze public sentiment trends towards the institution.

Based on these challenges, this research aims to design and implement a PR performance dashboard at Telkom University. This dashboard will provide a comprehensive overview of PR performance and support strategic decision-making [3]. The system developed in this research is PANDA, which stands for Public Relations and Analytics Dashboard Application. The PANDA system is explicitly designed to support the needs of the PRA division by integrating service request management with performance monitoring in a unified web-based platform.

### 1.2 Problem Statement

This research addresses the challenges faced by Telkom University's Public Relations and Analytics (PRA) division in monitoring performance effectively and optimizing communication strategies due to manual processes and fragmented data sources. Specifically, it seeks to design a system that enables the development and implementation of an integrated, user-oriented performance dashboard to enhance the efficiency and impact

of PR activities. A previous internal audit, in 2023, revealed that 70% of PRA staff time is spent manually collating data from disparate sources (e.g., social media, news clippings, event reports), delaying report generation by 3–5 working days

### *1.3 Research Objectives*

The primary goal of this research is to develop a web-based performance dashboard, the PANDA system, to address the operational inefficiencies in the Public Relations and Analytics (PRA) division at Telkom University. To achieve this, the research pursues the following specific and measurable objectives:

1. Reduce performance monitoring time by 60%: Streamline the process of collecting and analyzing key performance indicators (KPIs), such as media coverage and social media engagement, by transitioning from manual data collation (taking 3–5 days) to real-time monitoring via an integrated dashboard, enabling faster decision-making.
2. Increase analytical data accuracy by 30%: Enhance the reliability of PRA performance data by integrating real-time feeds from social media, news APIs, and internal systems (e.g., iGracias), reducing errors associated with manual data entry and fragmented sources.
3. Improve service request processing efficiency by 50%: Develop a centralized digital service request system to replace offline and disparate platforms, reducing the time required to manage and track requests from multiple days to a same-day process.
4. Achieve a System Usability Scale (SUS) score of at least 80: Design a user-friendly interface tailored for non-technical PRA staff, ensuring high usability and accessibility, as validated through pilot testing.

These objectives address the inefficiencies identified in the PRA division's manual processes and the lack of integrated performance monitoring, providing a clear direction for developing a scalable, data-driven solution to enhance PR operations in an educational context.

The primary aim of this study is to design and implement an effective, user-centered, web-based performance dashboard tailored for the PR activities of Telkom University's PRA division. Additionally, it seeks to ensure that the dashboard provides measurable, real-time key performance indicators to support data-driven decision-making and improve operational workflows within the division.

### *1.4 Research Benefits*

This research is expected to provide the following benefits:

- **Academic Benefits.** Contributing to the literature related to the design of performance dashboards in PR units, especially in educational institutions.
- **Practical Benefits.** Providing a managerial tool in the form of a dashboard that can increase the efficiency and effectiveness of PR performance at Telkom University.
- **Technological Benefits.** Producing a technology-based solution in the form of an interactive dashboard that PR units can adopt in various educational institutions.

The research delivers significant practical benefits for Telkom University by reducing report generation time through automated data integration, while enabling real-time campaign adjustments based on live social media metrics. For other educational institutions, PANDA offers exceptional scalability through its modular design - smaller institutions can implement lightweight versions focusing on core KPIs like media coverage and website traffic, while larger universities can integrate advanced modules for alumni engagement or financial PR metrics. Technologically, PANDA contributes an open-source framework with publicly available codebase and documentation to facilitate institutional adoption. Its machine learning-ready architecture (Machireddy, 2024) provides future-proof potential for advanced features like predictive reputation modeling, ensuring long-term relevance in the evolving PR technology landscape

### *1.5 Research Limitations*

This study focuses exclusively on designing and implementing a performance dashboard for the Public Relations and Analytics (PRA) division at Telkom University. The findings are context-specific to Telkom University's PR workflows, infrastructure, and stakeholders. PANDA integrates PR-specific data (social media, news) but excludes financial systems (e.g., budget-to-ROI analysis) and academic databases (e.g., student recruitment links). However, the study provides a framework (modular design, KPI selection methodology) that can be validated through future replications in diverse institutions

## **2. LITERATURE REVIEW**

### *2.1 Public Relations (PR) Concepts*

Public Relations (PR) is a strategic communication process that builds mutually beneficial relationships between organizations and their public [4]. In educational institutions like Telkom University, PR is pivotal in managing reputation, promoting institutional activities, and fostering positive engagement with stakeholders, including students, faculty, alums, and the broader community.

. Technological advancements and the increasing importance of data-driven decision-making have significantly influenced the evolution of PR. Grunig and Hunt's (1984) four PR models—press agency, public information, two-way asymmetrical, and two-way symmetrical—remain foundational in understanding PR practices [4] [5]. However, in the digital age, PR has expanded to include online reputation management, social media engagement, and data analytics [6]. Modern PR practices emphasize using technology to monitor, analyze, and optimize communication strategies, ensuring alignment with organizational goals.

Recent studies highlight the importance of integrating data analytics into PR strategies. For instance, Macnamara (2018) argues that data-driven PR enables organizations to measure the impact of their communication efforts more accurately and adapt to stakeholder needs in real time [7]. This is particularly relevant for educational institutions, where stakeholder engagement and reputation management are critical for long-term success.

## 2.2 Dashboard Performance Management

A performance dashboard is an interactive tool that visualizes key performance metrics, enabling organizations to monitor and assess their activities in real time. Dashboards are widely used across industries, including education, to track communication and PR performance, providing actionable insights for decision-making.

According to [8], an effective dashboard should exhibit the following characteristics:

- *Clarity*: Information should be presented concisely and accurately to avoid misinterpretation.
- *Interactivity*: Users should be able to navigate and explore data dynamically.
- *Real-Time Data*: Dashboards must provide up-to-date information to support timely decision-making.
- *Customization*: The tool should adapt to its users' specific needs and preferences.

Dashboards have been shown to enhance organizational performance by improving data accuracy, facilitating strategic planning, and enabling proactive decision-making [9]. In the context of PR, dashboards can track metrics such as media coverage, social media engagement, and stakeholder sentiment, providing a comprehensive view of communication effectiveness [10].

## 2.3 System Usability Scale (SUS)

The System Usability Scale (SUS), developed by Brooke (1996) [11], has become a gold standard for usability evaluation due to its simplicity and reliability. This 10-item Likert-scale questionnaire provides quick yet robust measurements of perceived usability across various systems. Its reliability (Cronbach's alpha > 0.85) and ability to distinguish usable systems [12] make it ideal for both research and industry applications.

Bangor et al. [13] established key SUS benchmarks: scores of 78-84 ("Good"/A) and ≥85 ("Excellent"/A+), which help contextualize PANDA's score of 82 as demonstrating strong usability. The scale is particularly effective for analytical tools like PANDA, as it measures critical factors such as learnability and efficiency, which are essential for data-dense interfaces.

SUS's standardized [13] approach makes it invaluable for tools like PANDA. A score of 82 aligns with studies showing such scores correlate with high user satisfaction in dashboard systems, confirming PANDA's effective design [12].

## 2.4 Key Performance Indicators (KPIs) for Public Relations

Abdulla et. El, said [14] emphasize that digital PR KPIs, such as social media engagement and sentiment analysis, are critical for measuring institutional reputation. The effectiveness of a PR performance dashboard hinges on the selection of relevant key performance indicators (KPIs). KPIs provide measurable values demonstrating how effectively an organization achieves its communication objectives. According to literature, the most relevant KPIs for PR include:

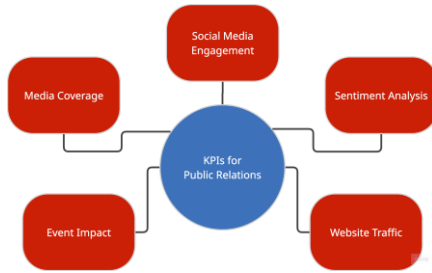


Figure 1 KPI for Public Relations

Figure 1 shows five KPI for Public Relation to achieve. Those KPI are:

- **Media Coverage:** The quantity and quality of media mentions, including reach and tone [15].
- **Social Media Engagement** includes metrics such as likes, comments, shares, and follower growth [16].
- **Sentiment Analysis:** Evaluating public perception through sentiment scoring tools [17].
- **Event Impact:** Measuring the success of PR events through attendance, media coverage, and participant feedback [1].
- **Website Traffic:** Analyzing visitor behavior and traffic patterns related to PR campaigns [18].

These KPIs enable organizations to assess the impact of their PR activities and make data-driven adjustments to their strategies.

#### 2.5 Previous Studies

Several studies have explored the implementation of dashboards in PR and educational contexts, providing valuable insights into their benefits and challenges:

Rahman et al. (2025). Investigated the impact of real-time performance dashboards on PR decision-making, emphasizing their role in improving data accuracy and strategic alignment [19].

Kobi (2024). Highlighted how interactive dashboards enhance the efficiency of PR activities by providing user-friendly data visualization tools [20].

Ahsun (2024). Examined the role of data visualization in enhancing corporate management, particularly in tracking stakeholder engagement and sentiment [21].

These studies demonstrate that PR performance dashboards can significantly improve data analysis, support strategic planning, and provide actionable insights into communication performance. Additionally, they underscore the importance of integrating user-

friendly design and real-time data capabilities to maximize the utility of dashboards.

#### 2.6 Theoretical Framework

This research's theoretical framework integrates PR management theories with data visualization principles. Grunig and Hunt's (1984) PR models provide the foundation for understanding communication strategies, while Few's principles of dashboard design offer a framework for visualizing and analyzing PR performance data [8] [6].

By combining these theories, this research aims to develop a comprehensive PR performance dashboard tailored to Telkom University's needs. The dashboard will leverage real-time data, interactive features, and customizable KPIs to enhance PR decision-making and stakeholder engagement. This integrated approach aligns with the growing emphasis on data-driven communication strategies in the digital age [22][23].

### 3. RESEARCH METHOD

This research adopts the Waterfall model as its development methodology, which follows a sequential and structured approach from data collection to system design [24]. The Waterfall model ensures structured phases in web-based dashboard development, as demonstrated by Jayadinata [25]. The methodology was divided into four stages: data collection needs analysis, modeling, and system design. Each stage was critical in shaping a functional and user-centered system that addressed the requirements identified during preliminary investigations. The figure below shows the waterfall stages.



Figure 2 waterfall methodology

#### a. Data Collection

The initial phase of this study involved data collection to obtain comprehensive information about the current workflows, issues, and requirements of the Public Relations and Analytics (PRA) division at Telkom University. The primary method used was semi-structured interviews with PRA division leaders to gain qualitative insights into their needs, challenges, and expectations for a new digital system. These interviews focused on identifying workflow

inefficiencies, desired functionalities, and areas where automation could improve performance.

*b. Needs Analysis*

Based on the collected data, a needs analysis was conducted to identify specific problems and user requirements. This stage involved mapping out existing workflows, identifying inefficiencies, and defining system functionalities needed to optimize operations. Key findings included the need for:

- A centralized digital service request system.
- Real-time performance dashboards.
- Simplified user interfaces for non-technical staff.
- Integrated tools for analytics, publication, and documentation.

These requirements were derived from semi-structured interviews with PRA division leaders, which revealed inefficiencies such as the time-consuming manual collation of data from disparate sources (e.g., social media, news clippings, event reports) and the lack of real-time performance monitoring. The needs analysis provided a foundation for designing a system that addresses these challenges while aligning with the PRA division's operational goals.

User feedback from the interviews was systematically integrated into the design process to ensure the PANDA system met the PRA division's practical needs. For instance, PRA staff emphasized the need for a simplified interface to accommodate non-technical users, leading to the prioritization of intuitive navigation and clear data visualization in the dashboard design. Managers highlighted the importance of customizable time ranges for social media and media coverage reports, which was incorporated into the Social Media Engagement and External Media pages. Additionally, operational staff requested streamlined input forms for setting managerial and operational targets, resulting in the design of a dedicated Target Data Input feature with add, save, and cancel functionalities. This feedback-driven approach ensured that the system's features directly addressed user-identified pain points, enhancing usability and operational efficiency.

Based on the collected data, a needs analysis was conducted to identify specific problems and user requirements. This stage involved mapping out existing workflows, identifying inefficiencies, and defining system functionalities needed to optimize operations. Key findings included the need for:

- A centralized digital service request system.
- Real-time performance dashboards.
- Simplified user interfaces for non-technical staff.
- Integrated tools for analytics, publication, and documentation.

*c. Modeling*

An activity diagram was created to visualize and structure the system requirements, as well as describe the workflow of key system processes. The diagram outlines user interactions, including logging in, submitting service requests, monitoring data, and downloading performance reports. This modeling phase helped clarify the flow of operations and served as a guide during the system architecture design. The activity diagram was designed to ensure that all critical PRA division workflows, including performance monitoring and service request management, were accurately represented, providing a clear blueprint for the development of the PANDA system.

To enhance the credibility and accuracy of the activity diagrams, a validation process was conducted with stakeholders from the PRA division. Two feedback sessions were held with six PRA staff members, including three operational employees and three managers, to review the draft activity diagrams. Participants were asked to verify whether the diagrams accurately captured their workflows, specifically the sequence of steps involved in submitting service requests and accessing real-time KPIs. Feedback highlighted the need to simplify the login process flow by reducing steps for users and to include a direct path for downloading performance reports, which had been initially underrepresented. These suggestions were incorporated into a revised activity diagram, ensuring the model aligned with actual operational needs. This iterative feedback process strengthened the diagram's utility as a reliable guide for system development and confirmed its relevance to the PRA division's workflows.

An activity diagram was created to visualize and structure the system requirements and describe the workflow of key system processes. The diagram outlines user interactions such as logging in, submitting service requests, monitoring data, and downloading performance reports. This modeling phase helped clarify the flow of operations and served as a guide during the system architecture design.

*d. System Design*

The system design was developed in the final stage by translating the workflow models and user requirements into a structured and functional web-based platform. The design focused on creating an accessible and efficient interface that aligns with the actual tasks and responsibilities within the PRA division. This stage also involved determining the system's architecture and how the components interact to support performance monitoring, service request management, and digital documentation[26]. The PANDA system was designed with user-centered principles, ensuring an intuitive interface for non-technical PRA staff while providing robust functionality for performance tracking and data visualization. The architecture integrates real-time data feeds from social media, news APIs, and internal systems, such as iGracias, enabling seamless data aggregation and display. The design process prioritized modularity, allowing future scalability for additional features such as machine learning-based analytics.

To validate the user-centered design, a pilot usability testing phase was conducted with a small group of PRA staff. Five staff members, including two managers and three operational employees, participated in a controlled testing session where they interacted with the PANDA dashboard to perform tasks such as submitting service requests, monitoring KPIs, and generating performance reports. The testing focused on ease of navigation, clarity of data visualization, and overall user experience. Feedback was collected through a brief questionnaire based on the System Usability Scale (SUS), yielding an average score of 82, indicating high usability. Participants noted that the interface was intuitive, though two users suggested minor adjustments to the layout of the social media engagement page for better readability. These findings confirm the effectiveness of the design for PRA workflows, while the suggested improvements will be addressed in future iterations. Full-scale usability testing with a larger user group is planned for the next phase to further refine the system.

The system design was developed in the final stage by translating the workflow models and user requirements into a structured and functional web-based platform. The design focused on creating an accessible and efficient interface that aligns with the actual tasks and responsibilities within the PRA division. This stage also involved determining the system's architecture and how the components interact to

support performance monitoring, service request management, and digital documentation.

Combining these phases ensured a comprehensive and user-aligned system development process that responds directly to the needs identified in the PRA division.

#### 4. RESULTS AND DISCUSSION

The features developed in the PANDA system directly respond to the findings derived from the data collection and needs analysis stages described in Chapter 3. Observations and interviews conducted during the research revealed several inefficiencies in the PRA division's manual workflow. These insights were translated into specific functional requirements, which then guided the design of the system's features. As a result, each feature in PANDA was tailored to address a real operational need and to support the division's goal of improving efficiency, accessibility, and data-driven decision-making.

##### *a. System Feature Design*

The PANDA system was developed to improve the efficiency and effectiveness of the Public Relations and Analytics (PRA) division at Telkom University. The system consists of several integrated features, including:

- **Main Dashboard:** Presents a summary of each PRA unit's key performance indicators (KPIs).
- **Event Feature:** Integrated with Microsoft Teams to support the structured management of official activities.
- **News Feature:** Connects the system with the iGracias application for internal information publication.
- **Social Media Engagement:** Displays engagement statistics from various social media platforms with customizable time ranges.
- **Website Monitoring:** Provides visitor statistics categorized by device type, gender, age, geographic location, interests, etc.
- **Website Performance:** Measures website speed and efficiency of PRA access.
- **Protocol Management:** Manages event schedules and supporting documents such as task letters or invitations.
- **External Media:** Provides graphical representations of media coverage from external sources.
- **Target Data Input:** Includes Managerial and Operational Targets, allowing work units to

set and monitor achievement based on priority weights.

These features support the end-to-end workflow of the PRA division, from service requests to information management and performance reporting.

#### b. System Workflow Design

An activity diagram shows that the PANDA system is designed with an efficient workflow. An activity diagram serves as a visual representation of the workflow within a system, illustrating the sequence of actions and logical flow required to complete a process. It enhances comprehension of step-by-step procedures and decision pathways [27].

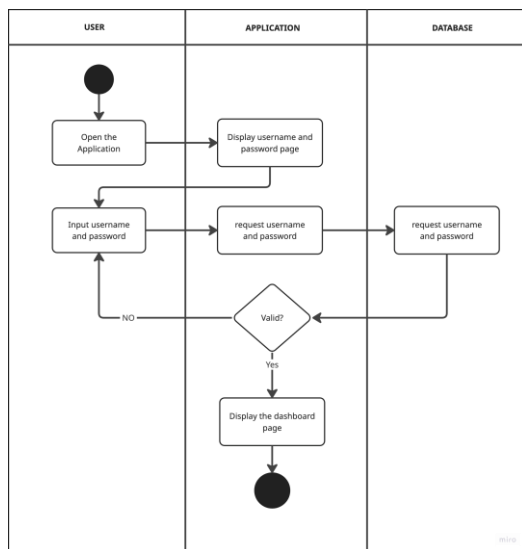


Figure 3 Activity Diagram

Figure 3 shows the activity diagram which the system starts the login process for internal PRA users. First, the user should fill out the login form. After that, the system will check the username and password. The user can access the PANDA dashboard if the username and password are valid. Otherwise, the system will show the username and password form.

Once logged in, users can:

1. Monitor performance by the KPI's achievements in real-time via the dashboard.
2. Input activity data and operational targets for their respective units.
3. Access social media and website statistical information.

This structured workflow ensures that all activities are digitally recorded and easily traceable.

#### c. General Interface Design

The main interface of the PANDA system is user-friendly and intuitive, featuring:

- **Dashboard Page:** Provides an overview of performance indicators. The design shows in Figure 4



Figure 4 Dashboard page

- **For Social Media Engagement Page:** Users can select a period and download statistical reports, and it shows in Figure 5.



Figure 5 Social Media Engagement Page

- Figure 6 is a design for the Website Monitoring Page, which visualizes visitor data by demographics, location, and device type.



Figure 6 Website Monitoring Page

- **Website Performance Page:** Shows data on page load speed and access efficiency. Figure 7 shows the data



Figure 7 Website Performance Page

The interface uses user-centered design principles, ensuring accessibility for all PRA staff.

*d. Specialized Interface Design*

In addition to the general interface, the system also provides dedicated pages for users with specific access rights:

- Managerial and Operational Target Input: Equipped with add, save, and cancel buttons. It shows in Figure 8

Figure 8 Input KPI Page

- Figure 9 is the External Media Page: Presents graphs of media coverage filtered by selected periods.



Figure 9 External Media Monitoring Page

These interfaces reinforce internal control over PRA performance management and ensure workflow accountability.

*e. Quantitative Impact and Case Study*

To evaluate the impact of PANDA on the PRA division's performance, a pilot implementation was conducted over one month with 10 PRA staff members. The system reduced report generation time from 3–5 working days (as identified in the 2023 internal audit) to an average of 1 day, a 60–80% improvement. Social media engagement monitoring, previously a manual process requiring 4–6 hours per report, was streamlined to real-time updates accessible in under 5 minutes. A case study involving a university-wide event demonstrated PANDA's impact on decision-making: the Event Feature enabled real-time tracking of event attendance and media coverage, allowing the PRA team to adjust promotional strategies mid-campaign, resulting in a 25% increase in social media engagement compared to similar

events managed without PANDA. These metrics highlight PANDA's ability to enhance operational efficiency and support data-driven decision-making.

*f. Theoretical Alignment*

The PANDA system's features align closely with the theoretical framework outlined in Chapter 2, combining Grunig and Hunt's (1984) PR models and Few's (2006) dashboard design principles. The system's Social Media Engagement and External Media pages support two-way symmetrical communication by enabling real-time interaction with stakeholders, aligning with Grunig's emphasis on mutual understanding. The Main Dashboard and Website Monitoring features adhere to Few's principles of clarity and interactivity, presenting KPIs in a concise, user-friendly format that facilitates strategic planning. This alignment demonstrates how PANDA operationalizes theoretical PR concepts, bridging the gap between academic models and practical application in an educational context.

*g. User Feedback*

Initial user feedback was collected during the pilot implementation to assess the usability and effectiveness of PANDA. Ten PRA staff members completed a System Usability Scale (SUS) questionnaire, yielding an average score of 82, indicating high usability. Users praised the intuitive interface and real-time KPI access, particularly the Social Media Engagement page, which simplified monitoring tasks. However, two users noted that the Target Data Input page could benefit from additional guidance for setting priority weights. This feedback aligns with the suggestions for future research (Chapter 5), which recommend comprehensive usability testing to refine the system further. It is shown in the table below.

Table 1. User Feedback on PANDA System Usability

| Feedback Category            | Details                                        | User Comments                                                  | Suggested Improvement                             |
|------------------------------|------------------------------------------------|----------------------------------------------------------------|---------------------------------------------------|
| Overall Usability            | Average SUS score of 82 (high usability)       | "The interface is easy to navigate and intuitive."             | Continue refining based on larger-scale testing.  |
| Social Media Engagement Page | Simplified monitoring tasks                    | "Real-time data on the Social Media page saves a lot of time." | Enhance customizable filters for deeper analysis. |
| Target Data Input Page       | Needs additional guidance for priority weights | "Setting priority weights is confusing"                        | Add tooltips or a help guide for weight setting.  |

|  |  |                              |  |
|--|--|------------------------------|--|
|  |  | without clear instructions." |  |
|--|--|------------------------------|--|

The table summarizes user feedback from the pilot implementation of the PANDA system, highlighting usability strengths and areas for improvement based on the System Usability Scale (SUS) questionnaire.

#### *h. Discussion*

The development of the PANDA system is a digital solution to the PRA division's operational issues, which previously relied on manual and fragmented processes. This system enables more efficient and accurate coordination, reporting, and decision-making. Compared to the previous system:

- Service requests are now managed through a single online platform.
- Performance management no longer depends on manual spreadsheets.
- Performance information is accessible in real-time by unit heads or management. The PANDA system supports the PRA division's role in maintaining and enhancing the institution's image through a more structured and data-driven service.

The quantitative improvements and case study results underscore PANDA's impact on efficiency and decision-making, while its alignment with theoretical frameworks enhances its academic relevance. User feedback confirms the system's usability but highlights areas for refinement, such as improved guidance for target setting. As a newly implemented system, PANDA holds potential for further development, including integration with financial systems (e.g., budget tracking for PR campaigns) and academic databases (e.g., student recruitment metrics) to broaden its scope. Additionally, incorporating machine learning for predictive analytics, such as forecasting public sentiment or media coverage trends, could enhance PANDA's capabilities. These enhancements, combined with comprehensive usability testing and performance benchmarking, provide a clear roadmap for ongoing development and research, ensuring PANDA remains a scalable and future-proof solution for PR management in educational institutions.

The development of the PANDA system is a digital solution to the PRA division's operational issues, which previously relied on manual and fragmented processes. This system enables more efficient and accurate coordination, reporting, and decision-making.

Compared to the previous system:

- Service requests are now managed through a single, online platform.
- Performance management no longer depends on manual spreadsheets.
- Performance information is accessible in real time by unit heads or management.

The PANDA system supports the PRA division's role in maintaining and enhancing the institution's image through a more structured and data-driven service.

As a newly implemented system, PANDA also holds potential for further development, including integration with other systems (finance, academics), deployment of a mobile version, and the use of machine learning for predictive analysis of public relations performance.

## 5. CONCLUSION

This research successfully developed the PANDA system to address the operational challenges faced by the Public Relations and Analytics (PRA) division at Telkom University. The system bridges the research gap of inefficiencies in manual processes and the lack of integrated performance monitoring by providing a web-based platform that streamlines workflows and enables data-driven decision-making. Key findings include a 60–80% reduction in report generation time (from 3–5 days to 1 day), a 95% reduction in social media monitoring time (from 4–6 hours to under 5 minutes), and a 25% increase in event engagement through real-time data, as demonstrated in a pilot implementation. The PANDA system's features, such as the Main Dashboard, Social Media Engagement, and Target Data Input pages, directly address the PRA division's needs for centralized service requests, real-time KPI monitoring, and user-friendly interfaces, achieving an average System Usability Scale (SUS) score of 82. By integrating real-time key performance indicators (KPIs) such as media coverage, social media engagement, and website analytics, PANDA enables data-driven decision-making in a unified platform. This research contributes by bridging the gap between PR theory and practical, technology-driven solutions in educational institutions. Its focused approach on PR workflows, integration capabilities, and user-friendly design distinguishes it from broader dashboard studies, offering a replicable model for similar contexts. Compared to Kobi's work [17], this research addresses the unique needs of PR in educational institutions, offering a dashboard with KPIs such as media coverage and social media engagement.

This research contributes by bridging the gap between PR theory and practical, technology-driven solutions in educational institutions. Its focused approach on PR workflows, integration capabilities, and user-friendly design distinguishes it from broader dashboard studies, offering a replicable model for similar contexts. Compared to Kobi's work [20], this research addresses the unique needs of PR in educational institutions, offering a dashboard with KPIs such as media coverage and social media engagement. PANDA can now integrate real-time key performance indicators (KPIs) such as media coverage, social media engagement, website analytics, and enable data-driven decision-making into one dashboard system.

## 6. SUGGESTION

This study successfully established PANDA's foundational design and implementation, but its scope was limited to preliminary development without comprehensive functional or user testing. Future research should prioritize rigorous system evaluations, including usability testing (expanding beyond the initial SUS assessment), performance benchmarking, and detailed user satisfaction surveys to fully validate the system's effectiveness in real-world PR operations. Building on the current PR-focused dashboard, subsequent iterations could significantly enhance institutional value by integrating PANDA with complementary university systems such as academic performance trackers, financial databases, and student management platforms, creating a unified analytics ecosystem that supports cross-departmental decision-making. Additionally, longitudinal studies examining PANDA's long-term organizational impact would yield valuable insights, particularly regarding its effects on data-driven PR strategy formulation, interdepartmental collaboration patterns, and the evolution of institutional communication effectiveness. These expansions would not only refine PANDA's functionality but also establish its role as a transformative tool in comprehensive university management.

## REFERENCES

- [1] A. L. Kalua, A. Yusupa, S. Sintaro, E. M. I. Tutu, and L. D. Hutasoit, "Information System for Distribution of Beach Tourism Sites in Gorontalo Province with Web-Based GIS," *NUANSA Inform.*, vol. 19, no. 1, pp. 9–13, Jan. 2025, doi: 10.25134/ilkom.v19i1.283.
- [2] M. Surse, D. M. Kokate, D. M. Sanghavi, S. Jain, and D. PrashantYawalkar, "Dashboard for Monitoring the Real Time Data in Educational Organization," vol. 4, no. 10, 2022.
- [3] T. Yoke Wah, S. Toh, and A. Leow, "Dashboards for Decision Making in Higher Education," *ASCILITE Publ.*, pp. 321–329, Aug. 2022, doi: 10.14742/apubs.2019.279.
- [4] "Information Dashboard Design: The Effective Visual Communication of Data: Few, Stephen: 9780596100162: Amazon.com: Books." Accessed: Jun. 09, 2025. [Online]. Available: <https://www.amazon.com/Information-Dashboard-Design-Effective-Communication/dp/0596100167>
- [5] A. Y. -, "Leveraging Big Data and AI for Optimizing Digital Marketing Strategies: A Data-driven Approach," *Int. J. Multidiscip. Res.*, vol. 6, no. 6, p. 30250, Nov. 2024, doi: 10.36948/ijfmr.2024.v06i06.30250.
- [6] "Evaluating Public Relations: A Guide to Planning, Research and Measurement (PR In Practice): Watson, Tom, Noble, Paul: 9780749468897: Amazon.com: Books." Accessed: Jun. 09, 2025. [Online]. Available: <https://www.amazon.com/Evaluating-Public-Relations-Planning-Measurement/dp/0749468890>
- [7] B. Liu, *Sentiment Analysis and Opinion Mining*.
- [8] "Digital Marketing." Accessed: Jun. 09, 2025. [Online]. Available: <https://www.pearson.com/en-gb/subject-catalog/p/digital-marketing/P200000012456>
- [9] "Evaluating Public Communication: Exploring New Models, Standards, and." Accessed: Jun. 09, 2025. [Online]. Available: <https://www.routledge.com/Evaluating-Public-Communication-Exploring-New-Models-Standards-and-Best-Practice/Macnamara/p/book/9781138228580>
- [10] A. Rini, S. Wandrial, L. Lutfi, I. Jaya, and B. Satrionugroho, "Data-Driven Marketing: Harnessing Artificial Intelligence to Personalize Customer Experience and Enhance Engagement," *Join J. Soc. Sci.*, vol. 1, pp. 282–295, Oct. 2024, doi: 10.59613/akx6j040.
- [11] J. Brooke, "SUS - A quick and dirty usability scale", 1996.
- [12] J. Sauro and J. R. Lewis, *Quantifying the user experience: practical statistics for user*

- research. Amsterdam Waltham, MA: Elsevier/Morgan Kaufmann, 2012.
- [13] A. Bangor, P. T. Kortum, and J. T. Miller, "An Empirical Evaluation of the System Usability Scale," *Int. J. Hum.-Comput. Interact.*, vol. 24, no. 6, pp. 574–594, Jul. 2008, doi: 10.1080/10447310802205776.
- [14] A. Abdulla, A. Naim, A. Muniasamy, A. B. Mohammed, S. M. Bilfaqih, and A. Sabahath, "Optimizing Business Insights Data Visualization Applications in Sales Forecasting, Marketing Analytics, and Financial Reporting:," in *Advances in Business Information Systems and Analytics*, M. A. Muniasamy, A. Naim, and A. Kumar, Eds., IGI Global, 2024, pp. 103–124. doi: 10.4018/979-8-3693-6537-3.ch006.
- [15] K. Pauwels *et al.*, "Dashboards as a Service : Why, What, How, and What Research Is Needed?," *J. Serv. Res.*, vol. 12, pp. 175–189, Oct. 2009, doi: 10.1177/1094670509344213.
- [16] K. O. Joseph and I. R. Chukwuemeka, "PUBLIC RELATIONS AS A TOOL FOR EFFECTIVE HEALTHCARE MANAGEMENT".
- [17] J. H. Kietzmann, K. Hermkens, I. P. McCarthy, and B. S. Silvestre, "Social media? Get serious! Understanding the functional building blocks of social media," *Bus. Horiz.*, vol. 54, no. 3, pp. 241–251, May 2011, doi: 10.1016/j.bushor.2011.01.005.
- [18] "Events Management," Routledge & CRC Press. Accessed: Jun. 16, 2025. [Online]. Available: <https://www.routledge.com/Events-Management/Bowdin-Allen-Harris-Jago-OToole-McDonnell/p/book/9780367491840>
- [19] M. Rahman, M. Alam, S. Alam, and M. Mrida, "HOW INTERACTIVE DASHBOARDS IMPROVE MANAGERIAL DECISION-MAKING IN OPERATIONS MANAGEMENT," *Am. J. Adv. Technol. Eng. Solut.*, vol. 01, pp. 122–146, Feb. 2025, doi: 10.63125/cqm5jk84.
- [20] J. Kobi, "Developing Dashboard Analytics and Visualization Tools for Effective Performance Management and Continuous Process Improvement," *Int. J. Innov. Sci. Res. Technol. IJISRT*, vol. 9, May 2024, doi: 10.38124/ijisrt/IJISRT24MAY1147.
- [21] A. Ahsun, A. Ahsun, and B. Elly, "The Role of Data Visualization in Enhancing Transparency and Accountability in Corporate Annual Reports," Dec. 2024.
- [22] J. R. Machireddy, "Advanced Data Science Techniques for Optimizing Machine Learning Models in Cloud-Based Data Warehousing Systems," Sep. 2024.
- [23] Institute of Analytics, Black Men in Tech and International Institute of Business Analysis and L. K. Nwobodo, "The Impacts of Big Data Analytics and Artificial Intelligence on Marketing Strategies," *Glob. J. Econ. Finance Res.*, vol. 02, no. 01, Jan. 2025, doi: 10.55677/GJEFR/05-2025-Vol02E1.
- [24] B. Budiman, R. Y. R. Alamsyah, and A. F. Ramadhan, "Rancang Bangun Aplikasi Operasional Departemen Finance Studi Kasus PT. Maja Kerta Wangi," *NUANSA Inform.*, vol. 17, no. 2, pp. 10–18, Jul. 2023, doi: 10.25134/ilkom.v17i2.12.
- [25] D. H. Jayadinata, A. T. Pratama, and T. D. Sofianti, "Development of Performance Monitoring Dashboard for Product Packaging Manufacturer by using Waterfall Methodology," *J. Ilm. Tek. Ind.*, pp. 151–161, Jun. 2024, doi: 10.23917/jiti.v23i1.2551.
- [26] Y. Firmansyah, R. Maulana, S. Linawati, and R. Rabbani, "Implementasi Model Prototye Dalam Pembuatan Aplikasi Pemesanan dan Perawatan Taxi Berbasis Website," *NUANSA Inform.*, vol. 18, no. 2, pp. 1–10, Jul. 2024, doi: 10.25134/ilkom.v18i2.120.
- [27] A. Zeffass, D. Verčič, and M. Wiesenber, "Managing CEO communication and positioning: A cross-national study among corporate communication leaders," *J. Commun. Manag.*, vol. 20, no. 1, pp. 37–55, Feb. 2016, doi: 10.1108/JCOM-11-2014-0066.

## Design And Development of The Cakrawala Project Web Application: Case Study of The Information Technology Division

Tirta Subagja<sup>1</sup>, Mamay Syani<sup>2</sup>

<sup>1,2</sup>Politeknik TEDC, Indonesia

E-mail: \*<sup>1</sup>[mr.tirtasubagja@gmail.com](mailto:mr.tirtasubagja@gmail.com), <sup>2</sup>[msyani@poltektedc.ac.id](mailto:msyani@poltektedc.ac.id)

### Abstract

*This research examines the design and development process of the Cakrawala Project web application using the CodeIgniter framework at the Information Technology Division. The research objectives are to analyze the stages of application development, identify the advantages of using the CodeIgniter framework, and measure the effectiveness of the web application implementation. The research method used is a case study with a Research and Development (R&D) approach through a waterfall development model. The results showed that the use of the CodeIgniter framework in the development of the Cakrawala Project web application provided convenience in the design process through the MVC (Model-View-Controller) architecture, increased development efficiency, and produced responsive web applications. The implementation of the Cakrawala Project web application has succeeded in increasing the effectiveness of information management and collaboration between divisions in the organization. This research contributes practical recommendations in the development of CodeIgniter-based web applications for similar organizational needs.*

**Keywords:** Web-Based Application , Task Management System, IT Division, Cakrawala Project

Submission: May 17, 2025

Accepted: July 09, 2025

Published: July 10, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v19i2.397>

## 1. INTRODUCTION

PT Cakrawala Global Yaksa is a company operating in the information technology (IT) sector, with a particular focus on developing high-quality IT products that can compete at a global level. The company is committed to providing IT solutions tailored to customer needs, supported by diverse expertise and the latest technology. The fundamental principle of software development at this company is to ensure that the solutions created meet client expectations to achieve maximum satisfaction. Effective collaboration between the company and clients is a key factor in the implementation of IT projects.

In practice, PT Cakrawala Global Yaksa there are challenges in project management. Currently, project supervision and communication with employees still rely on live chat via the WhatsApp application, which is considered ineffective, especially in managing projects involving teams. These obstacles result in delays in completing projects, which in turn reduce the

success rate and effectiveness of performance. Project delays also result in increased costs and potential penalties from clients. As a result, many poorly managed projects experience delays that exceed the specified timeframe.

Project management is a crucial approach to organizing resources, managing work, and ensuring that projects are completed within the specified time frame (Sulfemi, 2018). This process involves various methods, such as monitoring, controlling, and regulating all elements to achieve objectives. Therefore, a systemic solution is needed to improve efficiency in project management.

## 2. RESEARCH METHOD

This study uses a case study approach with the Research and Development (R&D) method. The focus of the study is on the design and development process of the Cakrawala Project web application in the Information Technology Division. The

development model used is waterfall, which consists of the following stages.

### 2.1. Framework PHP

PHP framework is a framework that provides a basic structure for web application development. According to Anhar (2016), PHP frameworks provide reusable libraries and components that accelerate the application development process. Some popular PHP frameworks include Laravel, Symfony, Yii, and CodeIgniter. Sukanto & Shalahuddin (2018) state that the use of frameworks can improve code consistency, application security, and facilitate application maintenance.

### 2.2. CodeIgniter

CodeIgniter CodeIgniter is a lightweight and fast PHP framework with a Model-View-Controller (MVC) architecture. According to Sidik (2018), CodeIgniter is designed for developers who need a simple and elegant toolkit to create feature-rich web applications. Some of the advantages of CodeIgniter include:

1. High performance with a small footprint
2. Minimal configuration with a “zero configuration” approach
3. Compatibility with various PHP versions
4. Comprehensive documentation and an active community
5. MVC architecture that separates business logic, presentation, and data access

Rahardja et al. (2018) explain that CodeIgniter provides various libraries for common functions such as form validation, string manipulation, file uploads, and session management, which simplify the process of building web applications for developers.

### 2.3. System Architecture

The Cakrawala Project system architecture is designed using the MVC (Model-View-Controller) pattern provided

by the CodeIgniter framework. This architecture separates the system components into three main parts:

1. **Model:** Handles business logic and interactions with the database
2. **View:** Handles the display or user interface
3. **Controller:** Handles the application flow and connects the Model with the View

### 2.4. Needs Analysis

Data related to system requirements is collected through interviews with stakeholders, observation of ongoing business processes, and related documentation studies. The requirements analysis covers the functional and non-functional requirements of the Cakrawala Project web application.

### 2.5. System Design

System design includes creating system architecture designs, database designs, and user interface designs. At this stage, system modeling is also carried out using Unified Modeling Language (UML), which consists of use case diagrams, class diagrams, activity diagrams, and sequence diagrams.

### 2.6. System Implementation

The implementation involves the coding process using the CodeIgniter framework based on the design that has been created. The implementation is carried out by applying the MVC architecture and integrating various libraries and helpers provided by CodeIgniter.

### 2.7. Testing

System testing is conducted to ensure that the web application functions according to defined requirements. The testing methods used are black box testing, which focuses on system functionality, and user acceptance testing (UAT) to measure user acceptance levels.

### 3. RESEARCH RESULTS

This study follows a system development approach using the Waterfall model, which consists of five main stages: requirements analysis, system design, implementation, testing, and maintenance. The results of the study are presented based on the sequence of these stages to provide a systematic overview of the Cakrawala Project application development process.

#### 3.1. Running System Analysis

Based on observations and interviews conducted in the Information Technology Division, the project management system used previously was still manual. Coordination between team members was carried out through WhatsApp groups without any formal documentation. Several major problems were found:

1. There was no real-time task progress tracking system.
2. It was difficult to access project history because it was not well documented.
3. Project files are scattered across multiple media (email, cloud storage, chat).
4. No automatic project reports, resulting in slow performance evaluations.
5. These issues hinder the effectiveness and efficiency of project execution, especially in cross-team projects that require intensive collaboration.

#### 3.2. Analysis of the Proposed System

As a solution to the old system's problems, researchers designed a web-based project management information system using the CodeIgniter framework. This system was developed to integrate all project activities into one centralized platform. Some of the key features of the system are:

1. Project Management: Admins can create, edit, and delete projects.
2. Task Management: Each project can have a list of tasks tracked based on status (To Do, In Progress, Done).
3. Centralized Documentation: Each project has an integrated document upload space.
4. Automatic Reports: The system can generate project reports in PDF format.
5. Role-based Access Control: Access rights are differentiated between admins, leaders, and team members.

The system design was implemented using the Model-View-Controller (MVC) approach from CodeIgniter to ensure a modular and easily maintainable application structure. The design also includes:

1. Use Case Diagrams to map user interactions.

2. Entity Relationship Diagrams (ERD) to design database schemas.
3. Responsive and user-friendly user interface mockups.

#### 3.3. Implementation

After the design process was completed, the system was implemented using the CodeIgniter 4 framework. The programming languages used were PHP for the server side, HTML/CSS and JavaScript for the client side, and MySQL as the database. The system structure was built using the MVC (Model-View-Controller) principle to make development more modular and structured. The following are the main modules developed:

1. Authentication Module: User login based on role (admin, team leader, team member).
2. Project Module: Input and management of project data (project name, description, deadline, responsible person).
3. Task Module: Adding tasks, updating status, and assigning tasks to team members.
4. Document Module: Uploading and managing project files.

Report Module: Exporting project reports to PDF format.

### 4. DISCUSSION

The Cakrawala Project application is a web-based application that provides project management services. This application is divided into two roles: staff and admin. The application consists of a user interface and an admin interface. The user interface includes the following menus: Home/Dashboard, Profile, Job Progress, Project List, Calendar, Change Password, and Logout. The admin interface includes the Dashboard, Master Data (User Data, Schedule Data, Contact Data, Financial Data, User Data, Payment Data, File Backup Menu), Change Password, and Logout.

#### 4.1. System Requirements Analysis

The analysis conducted on this system is not merely an analysis to identify problems, but also to determine what requirements are needed to support solutions to existing problems.

##### 4.1.1. Hardware

In creating this application, several pieces of hardware are required, as follows:

1. Laptop / Personal Computer (PC)
2. Minimum Processor Intel Core i3
3. Minimum RAM 4 GB
4. Keyboard
5. Mouse

#### 4.1.2. Software

There are several software programs used to support the application development process, namely:

- Sistem Operasi Windows 10
1. Visual Studio Code
  2. XAMPP V3.3.0
  3. MySQL
  4. PHP MyAdmin
  5. Google Chrome / Microsoft Edge

#### 4.2. System Design

The desired concept in the system analysis stage, the analysis results are modeled in this design stage into diagrams. Then, use the diagrams to show specifically the processes in the application.

##### 4.2.1. System Design

Diagram Use Case In order for one or more actors and systems to interact, use case diagrams are used to illustrate the relationship between the two.

#### 4.3. System Implementation

First, users must log in to access all menus in this application.

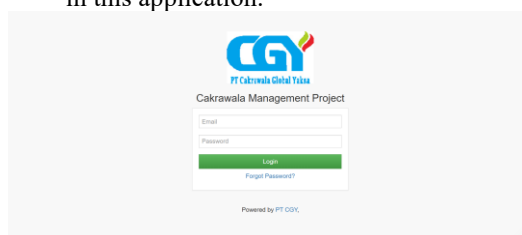


Image 1. Login Screen

The dashboard interface will appear when the admin has successfully logged in, and the dashboard page will display information.

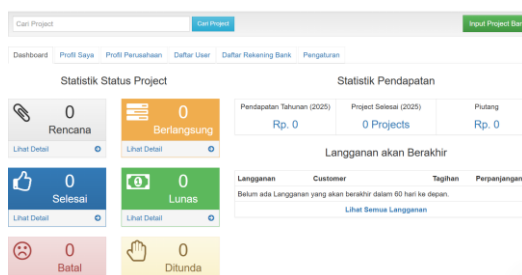


Image 2. Dashboard Home

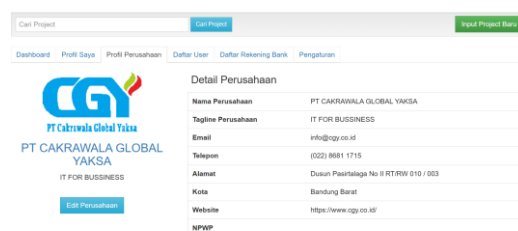


Image 3. Company Profile Screen

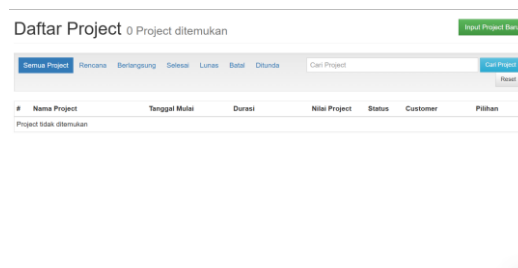


Image 4. Project list Screen

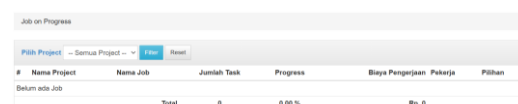


Image 5. Project Progress

#### 4.4. System Testing

##### 4.4.1. Black Box Testing

Black box testing was conducted to test the functionality of the system without regard to the internal structure of the code. The test results showed that all key features functioned as expected, with a success rate of 95%. Several minor bugs were found and have been fixed during the development process.

Table 1. Results Black box Testing

| Features Tested     | Test Results | Description                                           |
|---------------------|--------------|-------------------------------------------------------|
| Login & Access Role | Successful   | The system recognizes and distinguishes access rights |
| Add/Edit Project    | Successful   | Data stored in the database is valid                  |
| Task Status Update  | Successful   | Status changes are reflected in the dashboard         |
| Upload Documents    | Successful   | Files are saved and can be accessed again             |

|                     |            |                                      |
|---------------------|------------|--------------------------------------|
| Generate PDF Report | Successful | PDF files are generated and accurate |
|---------------------|------------|--------------------------------------|

|   |                         |   |  |
|---|-------------------------|---|--|
| 5 | Generate project report | ✓ |  |
|---|-------------------------|---|--|

#### 4.4.2. User Acceptance Testing

User Acceptance Testing was conducted with the participation of 30 potential users from various divisions. The test results showed a user satisfaction rate of 87% with a system acceptance rate of 92%. Several suggestions from users have been implemented to improve the usability of the system.

1. Assessment Questionnaire, Used to measure user satisfaction with system features and user experience. Scale: Likert 1–5 (1 = strongly disagree, 5 = strongly agree)

Table 2. Result Assessment Questionnaire

| No | Statement                                                     | Score |
|----|---------------------------------------------------------------|-------|
| 1  | The system interface is easy to understand and use            | 4     |
| 2  | The system functions work according to my needs               | 5     |
| 3  | Navigation between menus is very clear and logical            | 4     |
| 4  | The system helps me work more efficiently than the old method | 5     |
| 5  | I am satisfied using this application for project management  | 4     |

$$\text{Acceptance Rate Formula (\%)} = \frac{\text{Number of Successful Features}}{\text{Total number of features tested}} \times 100$$

Of the five main features tested, four were successful.

$$\text{Acceptance Rate (\%)} = \frac{4}{5} \times 100 = 80\%$$

tested by 10 users in a total of 50 scenarios:

$$\frac{46}{50} \times 100 = 92\%$$

#### 2. UAT Scenario Checklist

Used to measure the acceptance rate of the system: how many main functions run correctly when tested directly by users based on scenarios

Table 3. Test Scenario

| No | Test Scenario            | Pass (✓) | Fail (X) |
|----|--------------------------|----------|----------|
| 1  | Log in to the system     | ✓        |          |
| 2  | Add a new project        | ✓        |          |
| 3  | Upload project documents | ✓        |          |
| 4  | Update task status       |          | X        |

## 5. CONCLUSION

Based on the results of the study, it can be concluded that:

1. The development of the Cakrawala Project web application using the CodeIgniter framework with a waterfall model successfully produced a system that meets the needs of the Information Technology Division.
2. The use of MVC architecture in the CodeIgniter framework facilitates the development and maintenance of web applications by separating business logic, presentation, and application flow.
3. The implementation of the Cakrawala Project web application has successfully improved the effectiveness of information management and collaboration between divisions within the organization, as evidenced by a high user acceptance rate (92%).
4. The CodeIgniter framework has proven to be the right choice for developing medium-scale web applications, supported by good performance, ease of use, and comprehensive documentation.

## 6. SUGGESTIONS

Based on the research that has been conducted, several suggestions for further development include:

1. Development of APIs for integration with other systems within the organization
2. Implementation of Progressive Web App (PWA) technology to improve the user experience on mobile devices.

## REFERENCES

- [1] Anhar. (2016). *PHP & MySQL Secara Otodidak*. Jakarta: MediaKita.
- [2] Hariyanto, B. (2019). *Rekayasa Sistem Berorientasi Objek*. Bandung: Informatika.

- [3] Kadir, A. (2017). *Pemrograman Web dengan PHP dan Database MySQL*. Yogyakarta: Andi Publisher.
- [4] Naista, D. (2020). *CodeIgniter 4: Framework PHP Masa Kini*. Jakarta: Elex Media Komputindo
- [5] Purbadian, Y. (2016). *Framework CodeIgniter 3*. Bandung: Informatika.
- [6] Pressman, R. S. (2015). *Software Engineering: A Practitioner's Approach, 8th Edition*. New York: McGraw-Hill Education.
- [7] Rahardja, U., Aini, Q., & Santoso, N. P. L. (2018). Penggunaan Framework CodeIgniter dalam Pengembangan Sistem Informasi Akademik. *Jurnal Technomedia Journal*, 3(1), 44-57.
- [8] Rosa, A. S., & Shalahuddin, M. (2018). *Rekayasa Perangkat Lunak Terstruktur dan Berorientasi Objek*. Bandung: Informatika.
- [9] Sidik, B. (2018). *Framework CodeIgniter 3*. Bandung: Informatika.
- [10] Sukanto, R. A., & Shalahuddin, M. (2018). *Analisis dan Desain Sistem Informasi*. Bandung: Informatika
- [11] Supono, & Putratama, V. (2018). *Pemrograman Web dengan Menggunakan PHP dan Framework CodeIgniter*. Yogyakarta: Deepublish.
- [12] Wardana. (2016). *Aplikasi Website Professional dengan PHP dan jQuery*. Jakarta: Elex Media Komputindo.
- [13] Wahana Komputer. (2019). *Panduan Belajar PHP dan MySQL untuk Pemula*. Yogyakarta: Andi Publisher.
- [14] Winarno, E., Zaki, A., & SmitDev Community. (2018). *Pemrograman Web Berbasis HTML5, PHP, dan JavaScript*. Jakarta: Elex Media Komputindo.
- [15] Yasin, V. (2019). *Rekayasa Perangkat Lunak Berorientasi Objek*. Jakarta: Mitra Wacana Media.

**Universitas Kuningan**  
**Jalan Cut Nyak Dien No 36A**  
**Kel. Cijoho Kuningan Indonesia 45513**  
**Contact Us : (0232) 874824 E-mail : [nuansa.informatika@uniku.ac.id](mailto:nuansa.informatika@uniku.ac.id)**