

NUANSA INFORMATIKA

JURNAL TEKNOLOGI INFORMASI

Volume 20 - Nomor 1 - Januari 2026

p-ISSN : 1858 - 3911

e-ISSN : 2614-5405

Fakultas Ilmu Komputer
Universitas Kuningan

Development of a VR Metaverse System for Campus Navigation Using the Dijkstra Algorithm

Mochammad Nur Fikri Denintyo Denintyo, Tito Sugiharto, Yulyanto - Pages : 1-10

Analysis of Location Determination Using Wireless Signals at Low Range Using Third Order Polynomial Regression

Sutiyo, Yogi Yudo - Pages : 11-16

Development of a Transparent Vehicle Administration System Using Web-Based Information Architecture

Ido Jaya Gainal, Eryan Ahmad Firdaus - Pages : 17-25

Implementation of the e-Voting System for the Election of the Chair of the Student Executive Board Using QR Codes

Sintya Devi Januari, Darsanto, Muh Pauzan - Pages : 26-35

A Study on the Application of the Iterator Pattern to Enhance Software Reusability

Siti Herawati - Pages : 36-41

Web-Based Asset Maintenance Monitoring System for CV Tartila Group

Riyan Setiawan, Yuga nata praja, Wasis Haryono - Pages : 42-49

User Interface and User Experience (UI/UX) Design for the Ummu Micco Salon Website Using the Design Thinking Approach

Fathan Aqsha Ghasani Aufa, Febriyansyah Ramadhan, Debi Irawan - Pages : 50-58

Operationalizing Model Context Protocol for Multi-Platform Academic Search: A Comparative Study of LLM Grounding

Rizqullah Aryaputra Piliang, Anindito, Eryan Ahmad Firdaus - Pages : 59 - 71

Sentiment Analysis for Delivery Pricing Optimization : Naïve Bayes Evaluation on Gojek Reviews

Nisa Hamum, Woro Isti Rahayu , Ruth Diana Purnamasari - Pages : 72-83

SI-PANDA: Web-Based Personnel Data and Dossier Management System for Pusinfolahta TNI

Muhammad Sulthan Nasyira, Eryan Ahmad Firdaus - Pages : 84 - 91

High Performance Yolov4 with Different Optimizers and Backbones to Determine Distance

Irma Amelia Dewi, Mikael Prapaskalis G, Muhammad Ichwan - Pages : 92 - 98

Important Strategies in Mitigating Cyber Threats Using SLR in the Student Environment

Herdianto Sukmana, Fathoni, Dani - Page : 99-108

Deep Learning Evaluation for Interactive Dashboard-Base Mail Classification

Ruth Diana Purnamasari, Nisa Hamum Harani - Pages : 109-120

Real-Time Bayesian Knowledge Tracing for Adaptive Vocabulary Practice in an Adventure Educational Game: Architecture, Verification, and Reproducibility

Ikhsan Khaeruddin, Rio Andriyat Krisdianwan, Nida Amalia Asikin - Pages : 121-130

e-issn : 2614-5405



p-ISSN : 1858-3911



JURNAL NUANSA INFORMATIKA

Volume 20. Number 1, Januari 2026

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. **Prof. Dr. Azham Hussain** – Universiti Utara Malaysia – Malaysia
2. **Ts. Dr. Zahian Ismail** - Universiti Utara Malaysia – Malaysia
3. **Rohaida Romli (Prof. Madya Dr.)** - Universiti Utara Malaysia – Malaysia
4. **Ts. Dr. Mohamed Ali Saip** - Universiti Utara Malaysia – Malaysia
5. **Tata Sutabri, S.Kom., MMSI., MKM** – Universitas Bina Darma - Indonesia.
6. **Oman Komarudin, S.Si., M.Kom** - Universitas Singaperbangsa Karawang (UNSIKA) – Indonesia.
7. **Nur Alamsyah, M.Kom** - Universitas Informatika Dan Bisnis Indonesia (UNIBI) - Indonesia.
8. **Haodudin Nurkifli, S.T., M.Cs., Ph.D.** - Universitas Singaperbangsa Karawang (UNSIKA) – Indonesia.
9. **Erlan Darmawan, S.Kom., M.Si., Ph.D** - FKOM UNIKU, Indonesia.
10. **Fahmi Yusuf, M.MSI., Ph.D** - FKOM UNIKU, Indonesia.
11. **Taufik Ridwan, S.T., M.T.** - Universitas Singaperbangsa Karawang (UNSIKA) – Indonesia.
12. **Betha Nurina Sari, S.Kom., M.Kom.** - Universitas Singaperbangsa Karawang (UNSIKA) – Indonesia.
13. **Budiman, S.T., M.Kom** - Universitas Informatika Dan Bisnis Indonesia (UNIBI) - Indonesia.
14. **Mohamad Nurkamal Fauzan, S.T., M.T., SFPC** - Universitas Logistik dan Bisnis Internasional (ULBI) – Indonesia.
15. **Syafrial Fachri Pane, S.T., MTI** - Universitas Logistik dan Bisnis Internasional (ULBI) – Indonesia.
16. **Intan Purnamasari, S.Kom., M.Kom.** - Universitas Singaperbangsa Karawang (UNSIKA) – Indonesia.
17. **Sugeng Supriyadi, M.Kom** - Universitas Nurtanio - Indonesia
18. **Tito Sugiharto, S.Kom., M.Eng** - FKOM UNIKU, Indonesia.
19. **Rio Andriyat Krisdiawan, M.Kom** – FKOM UNIKU, Indonesia.

Proofing & Language Editor :

1. Roni Nursyamsu, M.Pd (FKOM UNIKU)
2. Nida Amalia Asikin, M.Pd (FKOM UNIKU)

JURNAL NUANSA INFORMATIKA

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
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

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

JURNAL NUANSA INFORMATIKA

Volume 20 Number 1, Januari 2026

Development of a VR Metaverse System for Campus Navigation Using the Dijkstra Algorithm

Mochammad Nur Fikri Denintyo Denintyo, Tito Sugiharto, Yulyanto

Page : 1-10

Analysis of Location Determination Using Wireless Signals at Low Range Using Third Order Polynomial Regression

Sutiyo, Yogi Yudo

Page : 11-16

Development of a Transparent Vehicle Administration System Using Web-Based Information Architecture

Ido Jaya Gainal, Eryan Ahmad Firdaus

Page : 17-25

Implementation of the e-Voting System for the Election of the Chair of the Student Executive Board Using QR Codes

Sintya Devi Januari, Darsanto, Muh Pauzan

Page : 26-35

A Study on the Application of the Iterator Pattern to Enhance Software Reusability

Siti Herawati

Page : 36-41

Web-Based Asset Maintenance Monitoring System for CV Tartila Group

Riyan Setiawan, Yuga nata praja, Wasis Haryono

Page : 42-49

User Interface and User Experience (UI/UX) Design for the Ummu Micco Salon Website Using the Design Thinking Approach

Fathan Aqsha Ghasani Aufa, Febriyansyah Ramadhan, Debi Irawan

Page : 50-58

Operationalizing Model Context Protocol for Multi-Platform Academic Search: A Comparative Study of LLM Grounding

Rizqullah Aryaputra Piliang, Anindito, Eryan Ahmad Firdaus

Page : 59 - 71

Sentiment Analysis for Delivery Pricing Optimization : Naïve Bayes Evaluation on Gojek Reviews

Nisa Hanum, Woro Isti Rahayu , Ruth Diana Purnamasari

Page : 72-83

SI-PANDA: Web-Based Personnel Data and Dossier Management System for Pusinfolahta TNI

Muhammad Sulthan Nasyira, Eryan Ahmad Firdaus

Page : 84 - 91

High Performance Yolov4 with Different Optimizers and Backbones to Determine Distance

Irma Amelia Dewi, Mikael Prapaskalis G, Muhammad Ichwan

JURNAL NUANSA INFORMATIKA

Page : 92 - 98

Important Strategies in Mitigating Cyber Threats Using SLR in the Student Environment

Herdianto Sukmana, Fathoni, Dani

Page : 99-108

Deep Learning Evaluation for Interactive Dashboard-Base Mail Classification

Ruth Diana Purnamasari, Nisa Hanum Harani

Page : 109-120

Real-Time Bayesian Knowledge Tracing for Adaptive Vocabulary Practice in an Adventure Educational Game: Architecture, Verification, and Reproducibility

Ikhsan Khaeruddin, Rio Andriyat Krisdiawan, Nida Amalia Asikin

Page : 121-130

Development of a VR Metaverse System for Campus Navigation Using the Dijkstra Algorithm

Mochammad Nur Fikri Denintyo^{*1}, Tito Sugiharto², Yulyanto³

^{1,2,3}Universitas Kuningan, Indonesia

E-mail: ^{*1}20210810144@uniku.ac.id, ²tito@uniku.ac.id, ³yulyanto@uniku.ac.id

Abstract

Advancements in Virtual Reality (VR) and Metaverse technologies have created new opportunities to enhance user experiences, particularly in education. This study presents the design and development of a VR-based Metaverse application for the Faculty of Computer Science at Universitas Kuningan. The application functions as an interactive promotional tool and campus navigation aid, utilizing the Dijkstra algorithm to compute the shortest routes within the virtual environment. Key features include realistic 3D campus visualizations, integrated information on essential locations, and route optimization from the main gate to user-selected destinations. Testing was conducted through black-box, white-box, and User Acceptance Testing (UAT) methods. User feedback indicated that navigation was effective, with directional clarity and usability scored at 3.59 and 3.18, respectively, while the 3D visualization received scores above 3.4 for realism. Interactivity and immersion achieved a score of 3.35, and the interface was deemed intuitive and accessible across devices (3.38 and 3.82). Despite minor issues with system responsiveness (averaging 2.65), users expressed high overall satisfaction and a strong intention to reuse and recommend the application, with a recommendation score of 3.76. This implementation highlights the potential of Metaverse technology in education and serves as a foundation for its wider adoption in Indonesian higher education.

Keywords— Virtual Reality, Metaverse, Campus Navigation, Dijkstra Algorithm, Educational Technology.

Submission: June 21, 2025

Accepted: December 1, 2025

Published: January 20, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.445>

1. INTRODUCTION

The digital era has significantly transformed various aspects of life, including education. The rapid development of information and communication technology has enabled the emergence of innovative learning methods and new forms of interaction. Among these, Virtual Reality (VR) and the Metaverse have emerged as cutting-edge technologies offering immersive and interactive experiences in digital environments[1], [2][3].

Metaverse, as a shared virtual space, has drawn increasing attention across different sectors, including higher education[4]. It presents a promising research direction due to its potential to create more engaging and collaborative learning environments. In the context of education, Metaverse can overcome traditional learning limitations by providing enriched and interactive experiences[5].

As an institution focused on information technology, the Faculty of Computer Science at Universitas Kuningan has a unique opportunity to adopt VR and Metaverse technologies to enhance educational quality and the student

experience. Integrating VR in learning environments can also attract prospective students by offering innovative ways to explore the campus[6]. Despite the availability of supporting hardware such as Oculus VR devices, the use of these technologies remains limited and underexplored, especially in the context of promotional media.

Currently, the promotional strategies employed for campus introduction rely primarily on social media and printed brochures. These media lack interactivity and immersive qualities, preventing prospective students from experiencing the campus environment directly. The information provided is typically static and constrained by 2D formats, with limited capacity for real-time updates. Additionally, printed brochures require physical distribution, which limits their outreach.

Particularly in Campus 2 of the Faculty of Computer Science, there has been no application of Virtual Reality for promotional purposes, even though technology offers the potential to deliver a more engaging and realistic campus tour. This research aims to address this gap by developing a Metaverse-based virtual

campus environment to introduce the faculty's facilities interactively.

To support effective exploration within the virtual environment, this study integrates the Dijkstra algorithm, which is known for its efficiency in determining the shortest path between two points[7]. This implementation will facilitate navigation within virtual space, helping users to move from the main gate to various destinations within the campus.

Therefore, this study focuses on the design and development of a Metaverse VR application for the Faculty of Computer Science, implementing the Dijkstra algorithm to enable efficient navigation within the virtual campus environment. The final prototype aims to provide a virtual space that accurately represents the faculty's facilities while offering a functional, immersive, and educationally valuable user experience.

The objective of this research is to develop an interactive and immersive VR-based Metaverse prototype that not only serves as an innovative promotional tool but also enables accurate campus navigation using the Dijkstra algorithm. Through this system, prospective students can explore the faculty environment virtually, gain spatial awareness of important locations, and experience a realistic campus atmosphere without needing to be physically present. It is expected that this research will contribute to the advancement of VR and Metaverse-based educational technologies and serve as a foundation for further exploration and adoption within higher education in Indonesia, particularly in computer science faculties.

2. RESEARCH METHODS

The research method outlines the systematic steps undertaken in this study, presented in the form of a flowchart. This flowchart illustrates the overall research process, starting from problem identification, requirement analysis, system design and development to system testing and evaluation of the research outcomes. Each stage is structured to ensure a clear and logical progression toward achieving the research objectives.

The research process shows in Figure 1, beginning with system requirement identification through problem analysis and literature review. Relevant references from books, journals, and existing Metaverse implementations in higher education were examined. Data collection was conducted via observations and interviews with faculty leaders, lecturers, and students to identify user needs.

The next phase involved system and interface design, including use case diagrams, navigation flowcharts, 3D object models, and Metaverse website layout. A prototype was then developed based on the initial design and evaluated by selected faculty and students to gather feedback on usability, feature completeness, and system effectiveness.

Based on user input, the system design and prototype were reviewed and refined before entering the final development stage, where all features were fully implemented and integrated. The system was subsequently tested to ensure its functionality and interactive experience. The final stage is system maintenance, aimed at ensuring long-term functionality and adaptability to future changes.

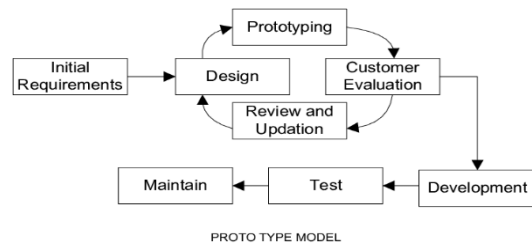


Fig 1. Research Process

2.1. Prototype System

This study adopts the Prototype System Development Model, which emphasizes iterative development, early user involvement, and continuous refinement of the system. The research began with the identification of problems related to the limited use of interactive media for campus promotion. This was followed by a literature review involving books, scientific journals, and prior studies related to Virtual Reality (VR) and Metaverse technologies in the context of higher education. To further understand user expectations, data were collected through direct observation and interviews with the dean, lecturers, and students of the Faculty of Computer Science at Universitas Kuningan. Comparative studies of existing Metaverse platforms were also conducted to gather references on design features and implementation strategies.

After identifying the problem and analyzing the context, a requirement analysis was carried out to determine the system's functional and non-functional needs. These included features such as 3D visualization of the campus, shortest-path navigation using the Dijkstra algorithm, integration of location-based

information, user-friendly interface design, and cross-device accessibility. The system design stage involved creating case diagrams, navigation flowcharts, and user interface mockups. Additionally, 3D models representing the campus environment were developed using Blender, and the virtual campus was built using the Spatial.io platform.

Following the design phase, a prototype was developed to simulate the main features of the system, such as interactive 3D exploration and campus navigation[3]. The prototype was then evaluated by a group of representative users, including faculty members and students, to gather feedback on usability, visual quality, interactivity, and system performance. The feedback obtained was analyzed and used to revise and improve the prototype, particularly in areas such as navigation flow, user interface clarity, and system responsiveness.

The improved version of the application was then finalized by integrating all refined features into a complete VR Metaverse system for Campus 2 of the Faculty of Computer Science. The Dijkstra algorithm was fully implemented to allow accurate and efficient navigation within the virtual environment. After development, the final system underwent a comprehensive testing phase, including black-box testing, white-box testing, and User Acceptance Testing (UAT), to ensure that all functionalities operated correctly and met user expectations. Finally, the system entered the maintenance phase, which focuses on ensuring long-term operability, providing technical support, and preparing the system for future enhancements or adjustments in line with evolving user needs.

2.2. Dijkstra Algorithm

The Dijkstra algorithm is a classic graph-based search algorithm developed by Edsger W. Dijkstra in 1959, designed to find the shortest path between nodes in a graph, which may represent, for example, road networks, computer networks, or virtual environments[8]. It is widely used in various fields, including routes, network optimization, and more recently, virtual navigation systems in immersive environments such as Virtual Reality (VR) and Metaverse applications.

In this study, the Dijkstra algorithm is implemented to calculate the shortest and most efficient navigation path between points of interest within the virtual campus environment. Each location on the virtual map is represented as a node, while the possible routes between

them are represented as edges with associated weights or distances. When a user selects a starting point and a destination, the algorithm calculates the optimal path by iteratively selecting the nearest unvisited node, updating the shortest distances, and marking nodes as visited until the destination node is reached. The result is a list of nodes representing the most efficient path, which is then visualized within the Metaverse environment as a guiding route.

The main advantages of the Dijkstra algorithm are its simplicity, accuracy, and reliability, particularly when all edge weights are non-negative. These characteristics make it a suitable choice for pathfinding in virtual spaces where quick and reliable navigation is essential to improve user experience. In the context of this research, the Dijkstra algorithm supports real-time navigation assistance within the virtual Faculty of Computer Science campus, allowing users to explore facilities such as classrooms, laboratories, and administrative offices effectively [9].

This implementation not only enhances orientation within the digital space but also adds functional value to the VR Metaverse system, making it not just a visual representation but also an interactive and informative platform. By integrating Dijkstra's algorithm, the system provides a personalized and efficient campus tour, which is particularly beneficial for prospective students unfamiliar with the real-world layout of the institution.

Each location within the virtual campus is modeled as a node, and the possible paths between them are treated as edges with associated weights or distances. When a user selects a starting point and destination, the algorithm performs a systematic traversal of the graph, updating the shortest known distances to each node and determining the optimal path based on the lowest cumulative cost. This calculated path is then visually rendered within the VR environment to guide the user.

In determining the shortest route using Dijkstra's Algorithm, there are several steps that must be taken so that the desired shortest route can be achieved properly in the campus Metaverse implementation, including:

1. Determine the starting point or node and the end node of the campus room location,
2. Determine the weight value or distance between each existing room location,
3. Next, set all nodes as 'untranslated nodes and the start node as the 'departure start node',
4. Path value = 0, find the initial shortest route using Dijkstra's Algorithm by comparing the

- weight value (distance) connected to the initial node,
5. If successful, the shortest path has been found, while if unsuccessful, the steps will be repeated to find the route,
 6. Nodes that have been travelled are designated as 'touched nodes' and become 'departure nodes' for the next stage, with the path distance value equal to the total path distance that has been travelled,
 7. Find the next route until it reaches the last node at the location of the campus room. If successful, then the shortest path to the final node has been reached, and the route-finding process is declared complete.

2.3. Research Data

In this research, the data used are the names of rooms and facilities in Campus 2 of the Faculty of Computer Science. In addition, information about the rooms and facilities, as well as information about the navigation route starting from the main gate, were used as reference points to various rooms on campus to manually draw the Graph. The distance between room locations was also obtained to process the data in manually finding the shortest route, and the coordinates of the room locations were used to create location point markers in the campus Metaverse system. Table 1 shows Location Coordinates.

No.	Node	Latitude	Longitude
1	A	-6.97569	108.47727
2	B	-6.97580	108.47789
3	C	-6.97584	108.47785
4	D	-6.97592	108.47780
5	E	-6.97595	108.47778
6	F	-6.97599	108.47771
7	G	-6.97606	108.47768
8	H	-6.97609	108.47761
9	I	-6.97604	108.47767
10	J	-6.97600	108.47770
11	K	-6.97591	108.47780
12	L	-6.97583	108.47783
13	M	-6.97580	108.47788
14	N	-6.97575	108.47776
15	O		
16	P	-6.97621	108.47746
17	Q	-6.97614	108.47737
18	R	-6.97610	108.47731
19	S	-6.97606	108.47726
20	T	-6.97601	108.47720
21	U	-6.97569	108.47761
22	V	-6.97572	108.47756
23	W	-6.97573	108.47771

24	X	-6.97588	108.47749
25	Y	-6.97594	108.47742
26	Z	-6.97605	108.47747
27	AA	-6.97609	108.47753
28	AB	-6.97578	108.47751
29	AC	-6.97560	108.47768
30	AD	-6.97580	108.47744
31	AE	-6.97574	108.47752
32	AF	-6.97570	108.47755
33	AG	-6.97567	108.47760
34	AH	-6.97590	108.47740
35	AI	-6.97596	108.47736
36	AJ	-6.97599	108.47732
37	AK	-6.97565	108.47785
38	AL	-6.97567	108.47745
39	AM	-6.97547	108.47747
40	AN	-6.97592	108.47756
41	AO	-6.97604	108.47725
42	AP	-6.97601	108.47722
43	AQ	-6.97602	108.47774
44	AR	-6.97615	108.47761
45	AS	-6.97617	108.47758
46	AT	-6.97616	108.47760
47	AU	-6.97622	108.47749
48	AV	-6.97609	108.47764
49	AW	-6.97597	108.47720

Table I : Location Coordinates Table

Below is the pseudocode for Dijkstra's Algorithm.

Pseudocode : Dijkstra Algorithm

Input:

A weighted graph $G(V, E)$,
where V is the set of vertices and E is the set of edges

A starting vertex source

Output:

Shortest distance from source to all other vertices

1. Initialize:

For each vertex v in V :

distance[v] $\leftarrow \infty$ // Unknown

distance from source to v

previous[v] $\leftarrow \text{null}$ // Previous node in optimal path

distance[source] $\leftarrow 0$ // Distance from source to itself is 0

$Q \leftarrow V$ // Set of all vertices in the graph (unvisited set)

2. While Q is not empty:

$u \leftarrow$ vertex in Q with smallest distance[u]

Remove u from Q

For each neighbor v of u :

alt $\leftarrow$ distance[u] + weight(u, v)

If alt < distance[v]:

distance[v] $\leftarrow$ alt

previous[v] $\leftarrow$ u
3. Return:
distance[]: shortest distance from source to each vertex
previous[]: optimal path tree to reconstruct paths

3. RESEARCH RESULTS

3.1 Proposed System

Analysis of the proposed system describes the stages of how the proposed system will function. The goal is to provide a detailed description of how the Faculty of Computer Science Metaverse VR Application works. This proposed system can be seen in Figure 2.

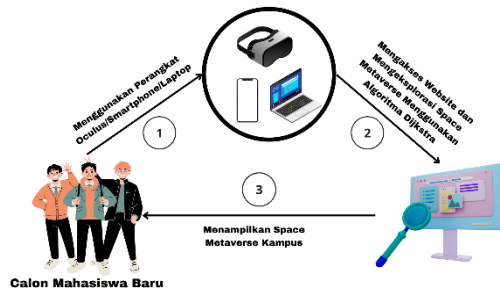


Figure 2. Proposed System

Based on Figure 2, the system from Spatial.io will display the pre-built Metaverse where users will access the Metaverse through various devices. Access can be done through mobile phones with a responsive interface, laptops/PCs for a more detailed experience, and also using Oculus devices for Virtual Reality mode that provides a higher level of immersion. When users have accessed the Spatial.io website and searched for the Metaverse of the Faculty of Computer Science Campus, the system will direct users to a specially designed virtual environment.

After successfully entering the Campus Metaverse, users can use their avatars to interact with the virtual environment and explore the entire modelled campus area. Users can explore various rooms such as Classrooms, Dean's Room, Computer Laboratory, Library, and other facilities that have been implemented in the virtual environment. The implementation of Dijkstra's Algorithm in the system will help users find the shortest route from one location to another, so that navigation in the Campus Metaverse becomes more efficient and intuitive for users, especially for new students or visitors who are not familiar with the campus layout.

3.2 Use Case Diagram

Use case diagram shows how the system interacts with the surrounding environment. In this use case diagram, users (Prospective New Students) can communicate with the system through various methods.

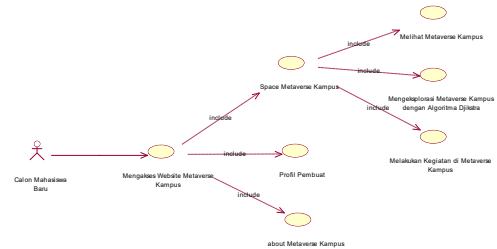


Figure 3. Use Case Diagram

Based on Figure 3, the system design involves a primary actor—prospective students—who interact directly with the Campus Metaverse system through several key use cases. The main use case is "Access Campus Metaverse Website", which includes additional features via the «include» relationship, such as "Creator Profile", "About Campus Metaverse", and "Campus Metaverse Space."

The "Campus Metaverse Space" serves as the core feature of the system, comprising three main activities: "View Campus Metaverse," "Explore Campus Metaverse," and "Engage in Activities within the Metaverse." These activities are directly linked to the main use case and guide users to access a 3D virtual environment hosted on the Spatial.io platform. Within this space, users can visualize the campus buildings in three dimensions.

Furthermore, the system offers an exploration feature powered by the Dijkstra algorithm, which helps users identify the shortest path for navigation within the virtual campus. This enhances the efficiency and clarity of movement through the Metaverse. Users can also engage in interactive activities, such as simulations or object-based interactions within the virtual environment.

4. DISCUSSION

The development and implementation of a VR-based Metaverse system for campus navigation have shown significant potential in enhancing user interaction and improving the promotional strategy of higher education institutions. The system was specifically designed for the Faculty of Computer Science at Universitas Kuningan and aimed to offer a virtual campus tour experience using 3D modeling and the Dijkstra algorithm for shortest-path navigation.

4.1 System Design and User Experience

The integration of Metaverse technology via the Spatial.io platform allowed for a realistic and immersive virtual representation of Campus 2. Based on the system evaluation, users responded positively to the 3D visualization aspect, with buildings and spatial layouts resembling the physical campus. This confirms that the design approach successfully met the goal of creating an interactive promotional media that goes beyond static content such as brochures and social media posts.

User feedback also indicated that the experience was engaging and intuitive, particularly in aspects of virtual exploration and navigation. Although some technical limitations were reported, such as slow loading times and occasional lag, the core functionalities were generally well-received.

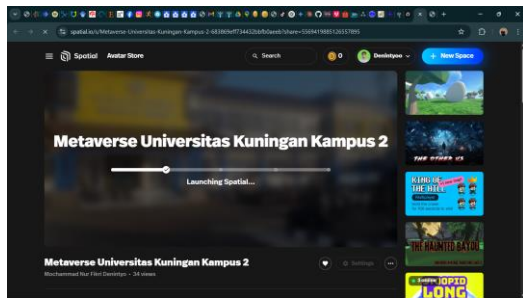


Figure 4. Metaverse

Figure 4 shows the display of the Metaverse Campus space when the user selects the Enter Metaverse menu, it will go directly to the space on the Spatial.io web page.



Figure 5. Metaverse View

This view displays the Campus Metaverse by displaying the 3D building of Campus 2 of Kuningan University with spaces and also the interior and exterior of the campus.

4.2 Navigation Efficiency with Dijkstra Algorithm

The application of the Dijkstra algorithm proved to be effective in enhancing the navigation system within the virtual

environment. By calculating the shortest path from the main gate to various destinations (e.g., classrooms, labs, administrative buildings), the algorithm provided structured guidance that helped reduce disorientation among users. The implementation aligned well with previous research, which supports the use of Dijkstra in spatial search problems within simulation and VR settings (Harahap & Khairina, 2017; Indriana et al., 2023).

Users rated navigation clarity and usability at 3.59 and 3.18, respectively, suggesting that the algorithmic support was beneficial in achieving efficient movement across the virtual space. This also indicates that algorithm-driven navigation can serve as an important feature for enhancing virtual campus systems, especially for prospective students unfamiliar with the environment.

User feedback indicated that navigation was effective, with directional clarity and usability scored at 3.59 ± 0.42 and 3.18 ± 0.38 , respectively ($n=30$ respondents). The distribution analysis showed that 73% of users rated navigation clarity above 3.0, with 23% providing excellent ratings (≥ 4.0).

4.3 Comparison with Previous Studies

To contextualize the contribution of this research, a systematic comparison with previous metaverse implementations in educational settings is necessary. This section specifically examines the work of Duan et al. (2021), which represents one of the most comprehensive campus metaverse prototypes documented in recent literature.

4.3.1 Overview of Compared Studies

Duan et al.'s CUHKSZ Metaverse (2021): Duan et al. developed a blockchain-driven metaverse prototype for The Chinese University of Hong Kong, Shenzhen. Their system was built on Unity game engine integrated with FISCO-BCOS consortium blockchain technology. The primary objectives centered on creating a virtual social ecosystem where students could interact, create user-generated content, and participate in a token-based economy. The system featured location-based services using GPS tracking, allowing students' physical presence on campus to influence their virtual experience—for example, earning tokens at higher rates when studying in the university library.

This Study - Faculty of Computer Science VR Metaverse (2025): This research developed a VR-based metaverse application for the Faculty of Computer Science at Universitas Kuningan using the Spatial.io platform. The

system serves dual purposes: as an interactive promotional tool for prospective students and as a campus navigation aid. The core distinguishing feature is the integration of Dijkstra's algorithm to compute optimal routes within the virtual environment.

4.3.2 Key Differences in System Design and Focus

Platform and Technology Architecture: Duan et al. utilized Unity as their development engine, which required custom programming and substantial technical infrastructure. Their blockchain implementation through FISCO-BCOS demanded server setup, smart contract deployment using Solidity, and ongoing blockchain maintenance. This approach provided deep customization capabilities but required significant technical expertise and computational resources.

In contrast, this research employs Spatial.io, a ready-made metaverse platform that operates through web browsers and supports multiple devices without requiring extensive backend infrastructure. The 3D campus models were created using Blender and imported into Spatial.io, while the Dijkstra algorithm was implemented as a separate computational module that interfaces with the platform. This approach reduces technical complexity while maintaining functional effectiveness.

Navigation Capabilities - The Critical Distinction:

The most significant difference between the two systems lies in their approach to user navigation within the virtual environment:

Duan et al.'s Approach: Their system employed location-based GPS tracking that synchronized users' physical location with their virtual avatar position. When students physically moved around campus, their avatar would correspondingly move in the metaverse. This created an interesting connection between physical and virtual worlds. However, this approach has inherent limitations:

- It only functions when users are physically present on campus
- It does not provide directional guidance or route suggestions
- New visitors or prospective students unfamiliar with campus layout receive no navigational assistance
- Users must explore manually or rely on prior knowledge of campus geography
- The system visualizes the campus but does not compute or recommend paths

This Study's Approach: This research implements Dijkstra's algorithm for intelligent

pathfinding, which fundamentally changes the navigation paradigm:

- The algorithm actively computes the shortest path between any two points in the virtual campus
- Users can select a destination (e.g., Computer Laboratory, Dean's Office, Classroom) and receive automated route guidance
- The system calculates optimal paths from the main gate to all major campus facilities
- Each location is represented as a node with coordinate data, and paths between locations have weighted distances
- The algorithm processes this graph structure to determine minimum-distance routes
- Visual indicators guide users along the computed path, reducing disorientation

Practical Example of the Difference:

Consider a prospective student who wants to visit the Computer Laboratory:

- In Duan et al.'s system: The student would see the 3D campus and must explore by moving their avatar around, potentially getting lost or taking inefficient routes
- In this system: The student selects "Computer Laboratory" from a destination menu, and the Dijkstra algorithm immediately computes and displays the shortest path from the main gate, showing turn-by-turn guidance through the virtual space

4.3.3 Functional Purpose and Target Users

Duan et al.'s System:

- **Primary Purpose:** Social ecosystem for current students
- **Target Users:** Enrolled students already familiar with campus
- **Core Features:** Social interaction, token economy, user-generated content, blockchain-based governance
- **Use Case:** Daily student engagement, virtual events, campus community building

This Study's System:

- **Primary Purpose:** Promotional tool and navigation aid
- **Target Users:** Prospective students, new students, visitors unfamiliar with campus
- **Core Features:** Campus visualization, algorithmic navigation, facility information, VR immersion

- Use Case: Campus tours, promotional activities, orientation for new students
- #### 4.3.4 Navigation System Implementation Details

To further clarify the technical distinction, the navigation implementations differ as follows:

Duan et al.'s Location-Based System:

- Uses smartphone GPS sensors to detect physical location
- Maps real-world coordinates to virtual world coordinates
- Avatar position updates automatically based on GPS data
- No pathfinding computation or route recommendation
- Navigation is passive (follows user movement) rather than active (guides user movement)

This Study's Dijkstra Algorithm Implementation:

- Pre-defines 49 nodes representing key campus locations (see Table I with coordinates)
- Calculates distances between connected nodes as edge weights
- When user selects destination, algorithm computes shortest path using iterative distance minimization
- Result is a sequence of nodes representing optimal route
- System visualizes this path with directional indicators in the 3D environment
- Navigation is active (system guides user) rather than passive

4.3.5 Complementary Strengths

It is important to note that both systems possess complementary strengths:

Strengths of Duan et al.'s approach:

- Rich social interaction features
- Blockchain-based economy and governance
- Strong community engagement mechanisms
- Seamless physical-virtual world connection for enrolled students

Strengths of this study's approach:

- Efficient navigation for unfamiliar users
- Lower technical barrier for implementation
- Clear promotional and informational value
- Algorithmic route optimization with measurable efficiency

4.3.6 Research Contribution Positioning

This comparative analysis reveals that while Duan et al. advanced the state of campus metaverse development by demonstrating

comprehensive social ecosystem creation, their work did not address the specific challenge of navigation assistance for users unfamiliar with campus layout. This research fills that gap by proving that metaverse applications can integrate intelligent pathfinding algorithms to serve practical navigational and promotional purposes.

The contribution of this study is not merely the creation of another campus metaverse, but specifically the integration of algorithmic navigation (Dijkstra) into a VR metaverse environment to solve the real-world problem of campus orientation for prospective and new students. This represents a functional advancement beyond pure visualization or social interaction, demonstrating that metaverse technology can serve practical, goal-oriented purposes in educational promotion and student services.

Future research could combine both approaches: building socially-rich campus metaverse environments (as demonstrated by Duan et al.) enhanced with intelligent navigation systems (as demonstrated in this study) to create comprehensive virtual campus platforms that are simultaneously engaging, functional, and accessible.

4.4 Educational and Promotional Impact

The system not only achieved technical feasibility but also showed promise in terms of its educational and promotional objectives. Many users expressed interest in reusing the system and even recommended it to others, with an average satisfaction score of 3.76. This implies that immersive virtual experiences have the potential to increase user engagement and institutional appeal.

Furthermore, the system could serve as a foundation for future learning tools or virtual campus events, especially in hybrid or distance learning contexts. By integrating interactive simulations, informational content, and real-time navigation, the Metaverse system expands the possibilities for how educational institutions communicate their identity and values

5. CONCLUSION

Based on the results and implementation of this research, several key conclusions can be drawn:

1. This study successfully designed and developed a Virtual Reality (VR)-based Metaverse application that represents the environment of the Faculty of Computer Science at Universitas Kuningan. The application functions as both an interactive

- promotional tool and a digital campus navigation system.
2. The integration of the Dijkstra algorithm effectively supports the navigation feature, enabling users to identify the shortest path from the main gate to various campus locations. This contributes to an intuitive and informative virtual exploration experience.
3. Testing using the black-box method confirmed that the core functionalities including navigation, object interaction, 3D visualization, and system performance across devices (PC and VR headsets) operated as expected.
4. The immersive campus visualization and guided navigation features enhanced the application's promotional appeal, increasing prospective students' interest and offering a clearer understanding of the campus layout.
5. This implementation demonstrates the potential of combining VR and Metaverse technologies with intelligent path-finding algorithms to improve both educational and promotional engagement in higher education environments.

6. SUGGESTIONS

Based on the results of this research, the following suggestions are proposed for future development and improvement of the system:

1. The application can be further enhanced by integrating social interaction features among users within the Metaverse environment. This would enrich the user experience and align more closely with the true concept of a fully immersive Metaverse.
2. Integration with the university's academic information system is recommended so that the application serves not only as a promotional tool, but also as a functional platform for academic purposes such as new student orientation, classroom mapping, and lecture scheduling.
3. System performance still requires improvement, particularly for mass-scale usage. Therefore, future testing with a larger number of simultaneous users is essential to evaluate the system's overall stability and responsiveness.

REFERENCES

- [1] H. Duan, J. Li, S. Fan, Z. Lin, X. Wu, and W. Cai, "Metaverse for Social Good: A

University Campus Prototype," in *MM 2021 - Proceedings of the 29th ACM International Conference on Multimedia*, Association for Computing Machinery, Inc, Oct. 2021, pp. 153–161. doi: 10.1145/3474085.3479238.

- [2] T. Tene, D. F. Vique López, P. E. Valverde Aguirre, L. M. Orna Puente, and C. Vacacela Gomez, "Virtual reality and augmented reality in medical education: an umbrella review," 2024, *Frontiers Media SA*. doi: 10.3389/fdgth.2024.1365345.
- [3] D. Perawatan, T. Berbasis, W. Y. Firmansyah, R. Maulana, S. Linawati, and R. Rabbani, "Implementasi Model Prototye Dalam Pembuatan Aplikasi Pemesanan," *Nuansa Informatika*, vol. 18, pp. 2614–5405, 2024, [Online]. Available: <https://journal.fkom.uniku.ac.id/ilkom>
- [4] B. Zhou, "Building a Smart Education Ecosystem from a Metaverse Perspective," *Mobile Information Systems*, vol. 2022, 2022, doi: 10.1155/2022/1938329.
- [5] G. J. Hwang and S. Y. Chien, "Definition, roles, and potential research issues of the metaverse in education: An artificial intelligence perspective," *Computers and Education: Artificial Intelligence*, vol. 3, Jan. 2022, doi: 10.1016/j.caeai.2022.100082.
- [6] A. Khalil, U. Saher, and A. Haqdad, "Prospects And Challenges Of Educational Metaverse In Higher Education," 2023. [Online]. Available: <http://journalppw.com>
- [7] M. K. Harahap and N. Khairina, "Pencarian Jalur Terpendek dengan Algoritma Dijkstra," *Jurnal & Penelitian Teknik Informatika*, vol. 2, no. 2, 2017, [Online]. Available: <https://rahadikusuma.blogspot.co.id>
- [8] D. Siregar, E. K. Silalahi, C. K. Panjaitan, and P. Harliana, "PERBANDINGAN EFISIENSI ALGORITMA DIJKSTRA DAN ALGORITMA A* (A STAR) DALAM MENEMUKAN RUTE OPTIMAL ANTARA SUN PLAZA DAN PODOMORO MENGGUNAKAN PYTHON," 2025.

- [9] R. A. Krisdiawan, A. Permana, E. Darmawan, F. Nugraha, and A. Kriswandiyanto, "Implementation Dijkstra's Algorithm for Non-Players Characters in the Game Dark Lumber," in *Journal of Physics: Conference Series*, Jun. 2021. doi: 10.1088/1742-6596/1933/1/012006..
- [10] C. Dede and J. Richards, "Digital Teaching Platforms in Virtual Reality and Metaverse," *Computers & Education*, vol. 177, pp. 14-29, 2022. doi: 10.1016/j.compedu.2021.104362.
- [11] H. Lee, "The Metaverse and Extended Reality for Education: A Systematic Review," *IEEE Transactions on Learning Technologies*, vol. 15, no. 4, pp. 490-503, 2022. doi: 10.1109/TLT.2022.3171692.
- [12] Y. Park and K. Kim, "Virtual Campus Tours Using VR Technology: Impact on Student Engagement," *Computers in Human Behavior*, vol. 126, p. 107019, 2022. doi: 10.1016/j.chb.2021.107019.
- [13] A. Radianti et al., "A systematic review of immersive virtual reality applications for higher education," *Computers & Education*, vol. 147, p. 103778, 2020. doi: 10.1016/j.compedu.2019.103778.
- [14] S. Chen, "Pathfinding Algorithms in Virtual Environments: A Comparative Study," *IEEE Computer Graphics and Applications*, vol. 41, no. 3, pp. 45-56, 2021. doi: 10.1109/MCG.2021.3067451.
- [15] R. Mantovani et al., "Dijkstra vs A* for VR Navigation: Performance Analysis," *ACM Transactions on Graphics*, vol. 40, no. 4, pp. 1-12, 2021. doi: 10.1145/3450626.3459819.
- [16] L. Wang, "Immersive Virtual Reality in Higher Education: A Meta-Analysis," *Educational Technology Research and Development*, vol. 69, no. 4, pp. 2221-2244, 2021. doi: 10.1007/s11423-021-10023-5.
- [17] K. Makransky, "A structural equation modeling investigation of the emotional value of immersive virtual reality in education," *Educational Technology Research and Development*, vol. 67, no. 5, pp. 1141-1164, 2019. doi: 10.1007/s11423-019-09653-2.
- [18] J. Jensen and F. Konradsen, "A review of the use of virtual reality head-mounted displays in education and training," *Education and Information Technologies*, vol. 23, no. 4, pp. 1515-1529, 2018. doi: 10.1007/s10639-017-9676-0.

Analysis of Location Determination Using Wireless Signals at Low Range Using Third Order Polynomial Regression

Sutiyo*¹, Yogi Yudo²

¹Telkom University, Bandung, Indonesia

²PT Telkom, Indonesia

E-mail: *¹tioatmadja@telkomuniversity.ac.id, ²yogisaloko@telkomuniversity.ac.id

Abstract

At short range, Location Determination Using Wireless Signals is considered a quite challenging problem. This paper discusses Location Determination Using Wireless Signals at short range with a Polynomial approach. The combination of retained independent effect of each component measure from Polynomial Regression and positioning in wireless fingerprinting produces a Polynomial that can provide good results for Location Determination Using Wireless Signals at short range. The measurement was conducted at Samas Beach Road, Yogyakarta, within a 0–100-meter range divided into 11 measurement points spaced every 10 meters. Validation results indicated that the third-order polynomial regression achieved a determination coefficient (R^2) of 0.9982, confirming its high accuracy in representing the measured signal strength variations at short range.

Keywords—Wireless, Polynomial, Location

Submission: October 7, 2025

Accepted: October 27, 2025

Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.509>

1. INTRODUCTION

Broadband Wireless Access (BWA) is a high-speed data access technology with wireless media. A number of services that are developing using BWA bands include broadband internet access, VoIP, video streaming and other digital applications. BWA is divided into two categories, namely fixed BWA and mobile BWA. Several BWA device industries have joined in promoting the technology standards they have developed to become standards adopted worldwide. BWA technology is divided into three, namely mobile telecommunication system-based technology (GSM, CDMA), data communication-based technology (WiFi, WiMAX) and non-terrestrial-based technology (VSAT, DVB). WiFi is a technology that uses non-exclusive or free BWA frequencies. Non-exclusive BWA frequencies, namely 2.4 GHz, 5.2 GHz, 5.8 GHz. WiFi is also known as Wireless Local Area Network (WLAN). Wireless technology that works on non-exclusive BWA bands, such as 802.11b/g/n WLAN devices, has advantages in terms of economy, but from a technological perspective has limitations. This technology is intended for Industrial, Scientific and Medical (ISM) activities.

In 2021, Linear with polynomial regression: Overview was studied by Patil et al. [1]. In addition, in 2021, a study entitled “Towards a Mathematical Framework to Inform

Neural Network Modeling via Polynomial Regression” was conducted by Morala et al. [2]. Subsequently, in 2023, Shi et al. [3] discussed a study entitled “Ensemble Regression Based on Polynomial Regression-Based Decision Tree and Its Application in the In-situ Data of Tunnel Boring Machine.” Furthermore, in 2023, Yao et al. [4] conducted a study titled “Toward a Holistic Perspective of Congruence Research with the Polynomial Regression Model.” Finally, in 2025, Mediation testing with polynomial regression: A critical review of extant approaches and a researcher's toolkit for the future was studied by Fu et al. [5].

In 2015, Xu et al. [6] investigated “Device Fingerprinting in Wireless Networks: Challenges and Opportunities.” Additionally, in 2017, Khalajmehrabadi et al. [7] conducted a study titled “Modern WLAN Fingerprinting Indoor Positioning Methods and Deployment Challenges.” Furthermore, in 2018, Hidayat et al. [8] proposed “Regression Analysis for Estimated Distance in Fingerprinting-Based WLAN Outdoor Localization System.” Finally, in 2019, Hidayat et al. [9] discussed a study entitled “Analysis of Accuracy and Precision of WLAN Position Estimation System Based on RSS.”

In 2021, Zhang et al. [10] discussed a study titled “Design Methodology of Free-Positioning Nonoverlapping Wireless Charging for Consumer Electronics Based on Antiparallel Windings.” In addition, in 2021, Xianjia et al.

[11] reviewed “Applications of UWB Networks and Positioning to Autonomous Robots and Industrial Systems.” Subsequently, in 2022, Koelemeij et al. [12] discussed “A Hybrid Optical–Wireless Network for Decimetre-Level Terrestrial Positioning.” Furthermore, in 2022, Hong et al. [13] examined “Self-Powered Seesaw Structured Spherical Buoys Based on a Hybrid Triboelectric Electromagnetic Nanogenerator for Sea Surface Wireless Positioning.” Finally, in 2024, Scavino et al. [14] investigated “Enhancing Indoor Positioning Accuracy with WLAN and WSN: A QPSO Hybrid Algorithm with Surface Tessellation.”

Most previous studies have focused on theoretical aspects, long-range environments, or other application domains such as robotics, psychology, and industrial systems. Only a few have investigated the application of polynomial regression for wireless signal-based positioning, particularly in short-range areas. Therefore, this study aims to evaluate the accuracy of polynomial regression in estimating the position coordinates of an outdoor WLAN access point at short range using a single received signal strength (RSS) value.

2. METHOD

2.1. Transmit Power

Transmit power is the magnitude of the signal output produced by a transmitting device. Transmit power is expressed in Watts (W) or deciBelWatts (dBW).

The power equation is shown in Equation 2.1.

$$P \text{ (Watt)} = V \text{ (volt)} \times I \text{ (amp)} \quad (2.1)$$

Where:

P = transmit power (Watt)

V = voltage (Volt)

I = current (Ampere)

The equations for converting watts to dBm are shown in equation 2.2 and equation 2.3.

$$P \text{ (dbW)} = 10 \log P / (1 \text{ Watt}) \quad (2.2)$$

$$P \text{ (dBm)} = 10 \log \frac{P}{1 \text{ mWatt}} \quad (2.3)$$

Where:

P(dBW): power in deciBelWatts (dBW)

P: power in watt (W)

1 mWatt: standard reference of 0.001 W

The equation for converting dBm to watts is shown in equation 2.4.

$$P \text{ (Watt)} = \frac{10^{\frac{P \text{ (dBm)}}{10}}}{1000} \quad (2.4)$$

Where:

P: power in watt (W)

P(dBm): power in deciBelmilliWatt (dBm)

dB (decibel) is a unit of ratio or difference in the strength of a signal emitted from a transmitter. The unit dBm (deciBelmilliWatt) is a unit of magnitude of the strength of a signal transmitted.

Log or logarithm is the inverse of exponent or power. Logarithma number with a cardinal number (written) is the exponent of the power number that results when raised to the power of that exponent. Logarithmic notation is shown in equation 2.5. $xa_{\log x}xa$

$$a_{\log x} = n, \quad (2.5)$$

Where :

a = cardinal number

x = logarithmic number with $x > 0$

n = algorithm result or exponent of base essence

For $a > 0$; $a \neq 1$ and $x > 0$

Antilog is the inverse operation of logarithm to exponent or power. Equation 2.5 can be inverted to exponent to get Equation 2.6.

2.2. Wireless Fingerprinting

Wireless fingerprinting is a position localization method that involves two stages [15]: the offline stage and the online stage. Offline phase is the phase of collecting data values that serve as references for the final goal of wireless fingerprinting. Online phase is the implementation phase in the field which includes data processing to find a value y based on input data x. The general method of wireless fingerprinting is shown in Figure 1.

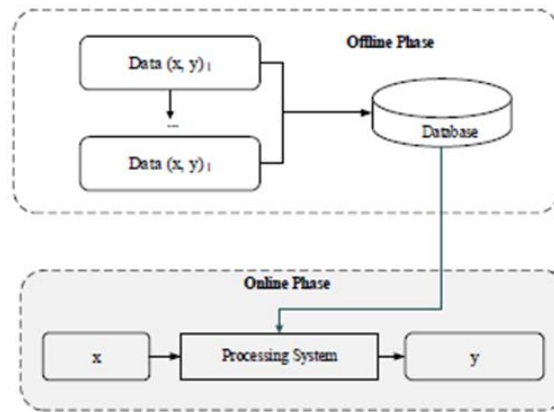


Figure 1. Wireless Fingerprinting Image

This wireless fingerprinting method is widely used in wireless localization systems. The implementation differences, when viewed from the offline phase, are in the data captured during wireless fingerprinting and the database storage system. The online phase is the processing system or algorithm and the search value.

2.3. Polynomial Regression

$$y = \beta_3 x^3 + \beta_2 x^2 + \beta_1 x + \beta_0 \quad (2.6)$$

Where:

y = Estimated received signal strength (RSS) or predicted power value (in dBm)

$\beta_1, \beta_2, \beta_3$ = Regression coefficients that determine the shape and curvature of the polynomial curve

β_0 = Intercept (constant term) of the polynomial regression equation

The highest power in polynomial regression indicates the order or degree of the polynomial.

$$y = 0,025x^3 + 1,69x^2 + 27,42x + 127,5 \quad (2.7)$$

At the final stage of polynomial regression, the entire research process is summarized in Figure 2. The diagram illustrates the complete workflow of the study, covering stages from data collection and transmit power calculation to the wireless fingerprinting process and regression modeling. Model validation is carried out using the coefficient of determination (R^2) to ensure that the polynomial regression accurately represents the relationship between Received Signal Strength (RSS) and distance.

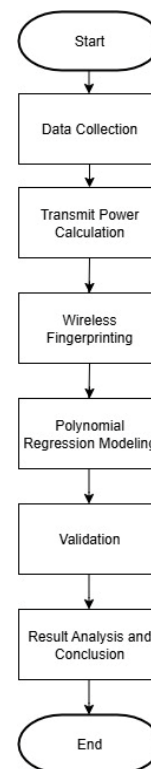


Figure 2. Workflow Diagram for Location Determination Using Wireless Signals

3. RESULTS

The experimental area has a low attenuation level and is located near the coast. The results of determining the coordinates of measurement points within the 0–100-meter range in the measurement area are shown in Table 1 and displayed in Figure 3.

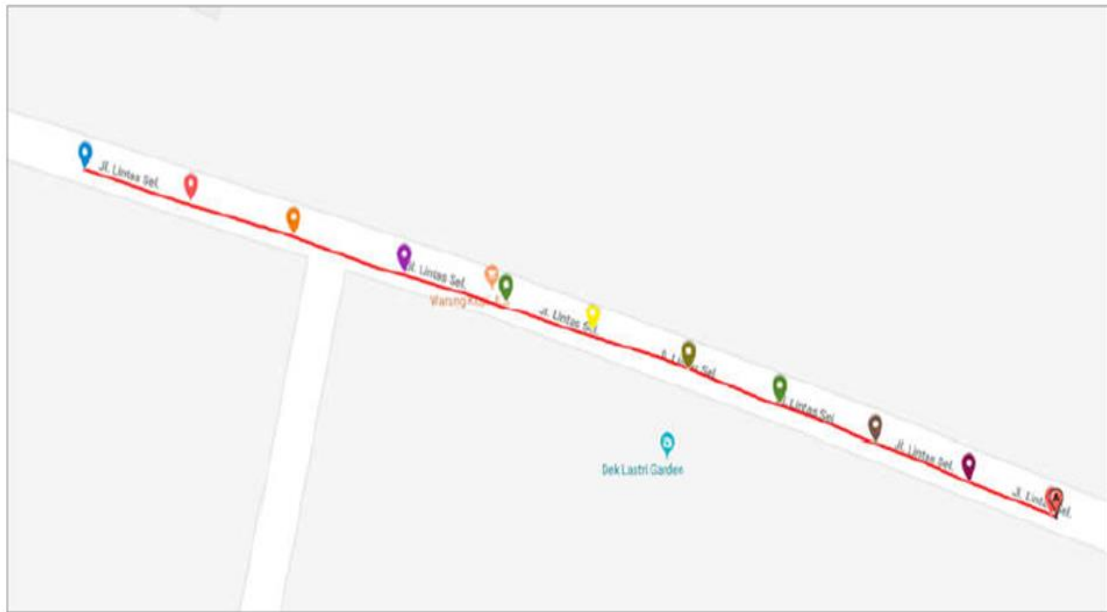


Figure 3. Site survey range 0 – 100 meters measurement areaan

Table 1 Coordinates of measurement points in the range 0 – 100 meters, area 1

Area	1	
Distance	0 – 100 meters	
Location	Samas Beach Road, Yogyakarta	
No	Coordinate	
	Lat.°	Lon.°
A1-100	-7.998882	110.264161
A1-90	-7.998904	110.264249
A1-80	-7.998927	110.264338
A1-70	-7.998948	110.264426
A1-60	-7.998871	110.264514
A1-50	-7.998993	110.264601
A1-40	-7.999015	110.264689
A1-30	-7.999037	110.264777
A1-20	-7.999059	110.264865
A1-10	-7.999081	110.264953
A1-0	-7.999103	110.265040

The site survey for wireless fingerprinting, covering a measurement range of 0 to 100

meters, was divided into 11 measurement points, each spaced 10 meters apart.

Table 2. Wireless fingerprinting analysis results range 0-100 meters

No	D_{real}	$RSS_{mean-A1}$ (dBm)	Stand. Deviation (SD_{A1})	Stand. Error (SE_{A1})
A1-100	100	-24.63	0.49	0.05
A1-90	90	-23.82	0.39	0.04
A1-80	80	-22.64	0.48	0.05
A1-70	70	-21.58	0.50	0.05
A1-60	60	-20.72	0.45	0.05
A1-50	50	-19.55	0.50	0.05
A1-40	40	-18.79	0.50	0.05

A1-30	30	-17.42	0.89	0.09
A1-20	20	-16.49	0.50	0.05
A1-10	10	-14.29	0.56	0.06
A1-0	0.01	-12.91	0.49	0.05

Table 2 shows that the SD and SE of each measurement point have values smaller than the average RSSmean. Therefore, it can be concluded that the average RSSmean value can be a representation of the overall wireless fingerprinting measurement results in the range of 0-100 meters.

The resulting regression model has a coefficient of determination ($R^2 = 0.9982$),

indicating that the third-order polynomial regression can represent the nonlinear relationship between Received Signal Strength (RSS) and distance with a very high level of accuracy. The R^2 value, which is close to 1, demonstrates that the model can explain almost all variations in the measurement data, thereby producing highly accurate position coordinate estimates in short-range wireless environments.

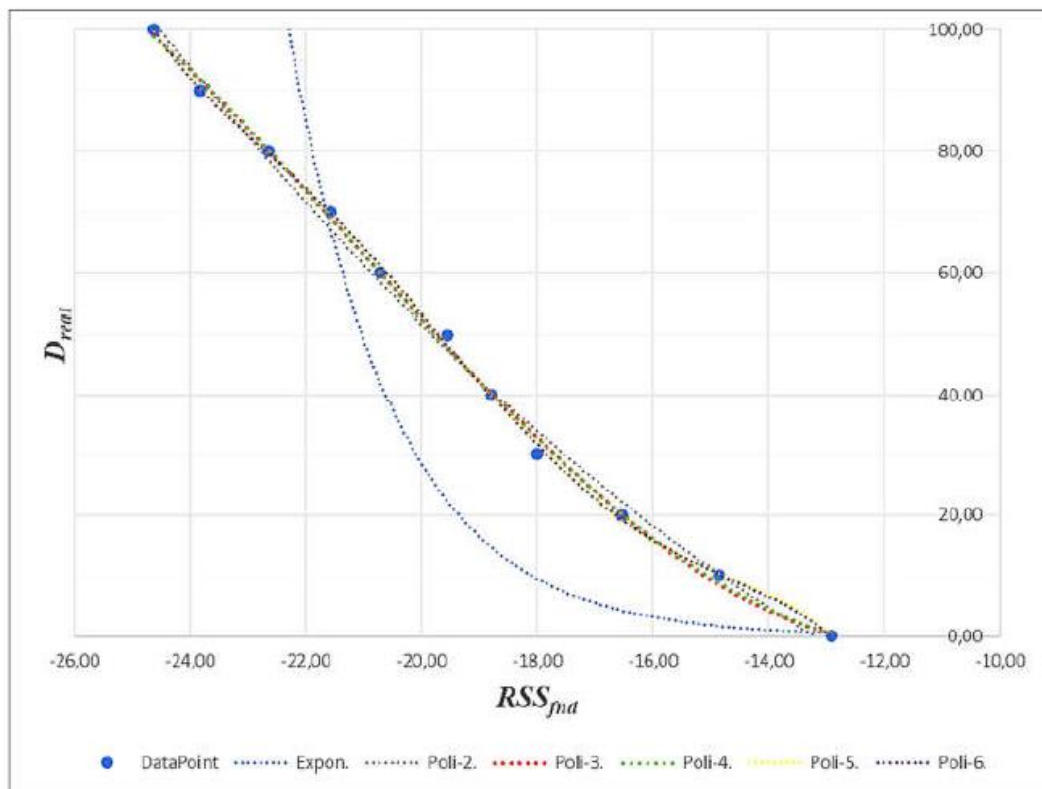


Figure 4. Regression graph of 0-100 meter range area 1

Based on the regression patterns of the six regression models shown in Figure 4, it can be concluded that the accurate regression pattern for wireless fingerprinting in the 0-100 meter measurement area range is 3rd order polynomial regression with an R^2 value of 0.9982.

4. CONCLUSION

Based on the results of field research and discussion, several conclusions can be stated as follows: the accurate regression pattern for wireless fingerprinting in the range of 0 – 100 meters of the measurement area is third-order

polynomial regression which obtains a fairly high R^2 accuracy of 99.82%.

Nevertheless, this study has several limitations. The measurements were conducted in a single outdoor area with relatively stable environmental conditions and a limited distance range of 0–100 meters. These conditions may restrict the generalizability of the research findings when applied to different environments or distance ranges. For future studies, it is suggested to broaden the measurement area to encompass various environmental conditions and to investigate the integration of polynomial

regression with optimization algorithms or neural networks to improve the accuracy of wireless positioning systems.

REFERENSI

- [1] S. Patil and S. Patil, "Linear and Polynomial Regression: An Overview," *Int. J. Appl. Res.*, vol. 7, no. 8, pp. 273–275, 2021.
- [2] P. Morala, J. A. Cifuentes, R. E. Lillo, and I. Ucar, "Towards a Mathematical Framework to Inform Neural Network Modelling via Polynomial Regression," *Neural Networks*, vol. 142, pp. 57–72, 2021.
- [3] M. Shi, W. Hu, M. Li, J. Zhang, X. Song, and W. Sun, "Ensemble Regression Based on Polynomial Regression-Based Decision Trees and Its Application to In-Situ Data of Tunnel Boring Machines," *Mech. Syst. Signal Process.*, vol. 188, p. 110022, 2023.
- [4] Y. A. Yao and Z. Ma, "Toward a Holistic Perspective of Congruence Research with the Polynomial Regression Model," *J. Appl. Psychol.*, vol. 108, no. 3, p. 446, 2023.
- [5] S. Q. Fu, N. Dimotakis, and J. Koopman, "Mediation Testing with Polynomial Regression: A Critical Review of Extant Approaches and a Researcher's Toolkit for the Future," *J. Appl. Psychol.*, 2025.
- [6] Q. Xu, R. Zheng, W. Saad, and Z. Han, "Device Fingerprinting in Wireless Networks: Challenges and Opportunities," *IEEE Commun. Surv. & Tutorials*, vol. 18, no. 1, pp. 94–104, 2015.
- [7] A. Khalajmehrabadi, N. Gatsis, and D. Akopian, "Modern WLAN Fingerprinting Indoor Positioning Methods and Deployment Challenges," *IEEE Commun. Surv. & Tutorials*, vol. 19, no. 3, pp. 1974–2002, 2017.
- [8] R. Hidayat, I. W. Mustika, and others, "Regression Analysis for Estimated Distance in Fingerprinting-Based WLAN Outdoor Localization Systems," in *2018 4th International Conference on Science and Technology (ICST)*, 2018, pp. 1–4.
- [9] R. Hidayat, I. W. Mustika, and others, "Analysis of Accuracy and Precision of WLAN Position Estimation System Based on RSS," *Int. J. Adv. Comput. Sci. Appl.*, vol. 10, no. 7, 2019.
- [10] Y. Zhang, S. Chen, X. Li, and Y. Tang, "Design Methodology of Free-Positioning Nonoverlapping Wireless Charging for Consumer Electronics Based on Antiparallel Windings," *IEEE Trans. Ind. Electron.*, vol. 69, no. 1, pp. 825–834, 2021.
- [11] Y. Xianjia, L. Qingqing, J. P. Queralta, J. Heikkonen, and T. Westerlund, "Applications of UWB Networks and Positioning in Autonomous Robots and Industrial Systems," in *2021 10th Mediterranean Conference on Embedded Computing (MECO)*, 2021, pp. 1–6.
- [12] J. C. J. Koelemeij, H. Dun, C. E. V. Diouf, E. F. Dierikx, G. J. M. Janssen, and C. C. J. M. Tiberius, "A Hybrid Optical–Wireless Network for Decimetre-Level Terrestrial Positioning," *Nature*, vol. 611, no. 7936, pp. 473–478, 2022.
- [13] H. Hong *et al.*, "Self-Powered Seesaw-Structured Spherical Buoys Based on a Hybrid Triboelectric–Electromagnetic Nanogenerator for Sea Surface Wireless Positioning," *Energy & Environ. Sci.*, vol. 15, no. 2, pp. 621–632, 2022.
- [14] E. Scavino, M. A. Abd Rahman, Z. Farid, S. Ahmad, and M. Asim, "Enhancing Indoor Positioning Accuracy with WLAN and WSN: A QPSO Hybrid Algorithm with Surface Tessellation," *Algorithms*, vol. 17, no. 8, p. 326, 2024.
- [15] F. Di Rienzo, A. Madonna, N. Carbonaro, A. Tognetti, A. Virdis, and C. Vallati, "Short-Range Localization via Bluetooth Using Machine Learning Techniques for Industrial Production Monitoring," *J. Sens. Actuator Networks*, vol. 12, no. 5, 2023, doi: 10.3390/jsan12050075.

Development of a Transparent Vehicle Administration System Using Web-Based Information Architecture

Ido Jaya Gainal*¹, Eryan Ahmad Firdaus ²

^{1,2,3}Universitas Pertahanan, Indonesia

E-mail: ¹gainalidojaya@gmail.com, * ²eryan.firdaus@idu.ac.id

Abstract

The management of official vehicles has historically relied on fragmented manual processes leading to inefficiencies, data inaccuracy, lack of transparency, and weak governance oversight. To address these challenges, this study designed and implemented Sikendi (Sistem Informasi Kendaraan Dinas), a web-based vehicle information system aimed at enhancing transparency and administrative governance. Using the System Development Life Cycle (SDLC) Waterfall approach, the system was developed with PHP, MySQL, JavaScript, and AJAX, featuring role-based access (Admin/User), integrated modules for vehicle registration, borrowing workflow with approval mechanisms, maintenance scheduling, fuel logging, document management, and activity auditing. Black-box testing confirmed functional compliance across 24 test cases for Admin and 15 for User roles. Results show the system significantly improves data accuracy, operational efficiency, and accountability, enabling real-time monitoring, automated reporting, and digital audit trails. This work demonstrates that digitizing fleet management in public sector institutions strengthens governance, reduces misuse risk, and supports data-driven decision-making, offering a replicable model for other governmental units.

Keywords—Vehicle information system, Web-based application, Fleet management

Submission: November 25, 2025 Accepted: December 14, 2025 Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.521>

1. INTRODUCTION

In the current era of digital transformation, the integration of information technology is essential for enhancing the operational effectiveness of defense and security organizations. The Indonesian National Armed Forces (TNI) is committed to strengthening institutional governance through system modernization to support its primary duties. Specifically, the Center for Information and Data Processing of the TNI (Pusinfo TNI) holds a strategic role in managing information systems and data processing to support both administrative and operational functions of the TNI Headquarters. To maintain optimal readiness, this unit relies heavily on the mobility of its personnel and the efficient distribution of logistics, which is facilitated by a fleet of official vehicles.

However, despite Pusinfo TNI's core function as a technology enabler, the internal management of its official vehicles remains largely conventional and inefficient. Preliminary observations at Pusinfo TNI indicate that the recording of vehicle specifications, borrowing schedules, and maintenance history is still conducted manually using physical logbooks and fragmented spreadsheets. This legacy approach creates significant operational bottlenecks. For instance,

determining the real-time availability of a vehicle for urgent technical deployment becomes time-consuming, and the lack of a centralized tracking system increases the risk of asset misuse. The urgency of this research lies in the critical need to align internal logistics governance with the unit's technological standards, ensuring that administrative delays do not hinder operational support duties.

The reliance on manual paper-based correspondence for borrowing approvals further exacerbates the issue, hindering the workflow and lacking the necessary transparency for auditing state-owned assets. Without a digitized system, leadership lacks data-driven insights into vehicle utilization and maintenance costs, leading to potential budget inefficiencies and reduced fleet readiness. Therefore, modernizing this sector is not merely an administrative upgrade but a strategic necessity to ensure accountability and responsiveness.

Previous research has widely discussed digitalization in fleet management. Studies suggest that implementing computerized asset management systems minimizes data redundancy and accelerates reporting cycles. Other research highlights that web-based applications with centralized databases are highly effective for real-time asset monitoring. However, limited studies focus on the specific

context of military administrative units, which require strict authorization hierarchies and rigorous accountability. This study addresses this gap by designing Sikendi (Official Vehicle Information System), a web-based application tailored to the specific operational needs of Pusinfolahta TNI.

Specifically, the Center for Information and Data Processing of the TNI (Pusinfolahta TNI) holds a strategic role in managing information systems to support the administrative and operational functions of the TNI Headquarters. Ideally, as a technology-driven unit, internal operations should reflect high standards of digital governance. However, preliminary observations indicate a significant gap: the management of official vehicles—a critical asset for personnel mobility and logistics still relies on fragmented manual processes such as physical logbooks and spreadsheets.

The urgency of this research lies in addressing this operational bottleneck. The current legacy approach results in data inconsistency and delays in vehicle allocation, which directly hinders the unit's responsiveness. Furthermore, the lack of a centralized digital tracking system creates a vulnerability in auditing state-owned assets (Barang Milik Negara), increasing the risk of misuse. Therefore, modernizing this sector is not merely an administrative upgrade but a strategic necessity to ensure accountability, auditability, and operational readiness within Pusinfolahta TNI.

By utilizing PHP and MySQL, Sikendi integrates vehicle registration, digital borrowing workflows, and automated maintenance scheduling. This system aims to transform the vehicle administration at Pusinfolahta TNI from a reactive, manual process into a proactive, transparent, and data-driven mechanism, thereby directly supporting the unit's operational effectiveness and strengthening the governance of state assets.

2. RESEARCH METHOD

To achieve the research objectives, this study adopts the System Development Life Cycle (SDLC) approach using a modified Waterfall model. While the standard Waterfall model typically consists of six stages, this research focuses on four core phases: Requirement Analysis, System Design, Implementation, and Testing.

The stages of Deployment and Maintenance were excluded from the scope of this study due to the time constraints of the On-Job Training (OJT) period and the research focus, which is centered on the prototyping and functional validation of the system rather than long-term lifecycle management.[3].

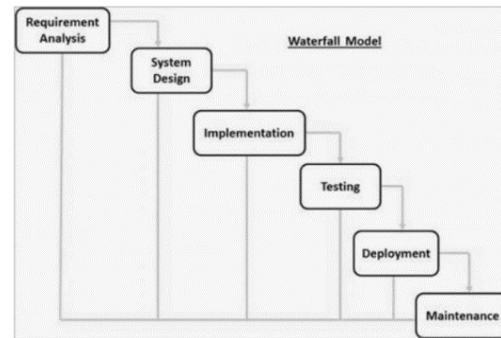


Figure 1. Waterfall Method

The development process consists of the following phases:

1. Requirement Analysis

The initial phase focused on defining the software specifications required to solve the identified operational problems. Data elicitation was conducted through observations of the manual workflow and in-depth interviews with key stakeholders, specifically the Head of Administration and the Head of Information Systems.

The analysis phase produced a detailed Software Requirement Specification (SRS) document, which categorized the system needs into two core components:

- a. **Functional Requirements:** The system must be capable of managing vehicle master data (CRUD), automating the borrowing workflow with a digital approval hierarchy, recording maintenance history with automated alerts, and generating real-time utilization reports.
- b. **Non-Functional Requirements:** The system requires a web-based architecture for accessibility within the local intranet, role-based access control (RBAC) to separate Admin and User privileges, and a user-friendly interface to ensure adoption by non-technical staff.

These defined requirements served as the foundational blueprint for the subsequent System Design phase, ensuring the proposed solution directly addressed the issues of data redundancy and lack of transparency found in the legacy manual process.

2. System Design

Following the requirement analysis, the system architecture was designed to translate the user needs into a technical blueprint. This phase utilized Unified Modeling Language (UML) diagrams to visualize the system structure and behavior[4].

Use Case Diagrams were created to define actor roles (Admin and User) and their interactions with the system, such as managing vehicles, approving requests, and logging activities. Activity Diagrams and Sequence Diagrams were developed to detail the logic flow for critical processes like login authentication, borrowing requests, and report generation[5]. Class Diagrams and Database Design were constructed to structure the relational database using MySQL, defining entities such as Users, Vehicles, Maintenance Records, and Fuel Logs.

3. Implementation (Coding)

In this phase, the design specifications were transformed into a functional web-based application. The development environment utilized XAMPP as the local server, with Visual Studio Code serving as the primary text editor[6]. Back-end Development: The server-side logic was built using the PHP programming language, ensuring robust data processing and secure interaction with the database. Front-end Development: The user interface was developed using HTML, CSS, and JavaScript to create a responsive and intuitive experience for users. AJAX (Asynchronous JavaScript and XML) was implemented to enable seamless data updates without requiring full page reloads. Database: MySQL was employed as the Relational Database Management System (RDBMS) to store and retrieve high volumes of operational data efficiently.

4. Testing

The final phase of development involved rigorous testing to ensure the system met the specified requirements and was free of critical errors[7]. The testing methodology used was Black Box Testing, which focuses on the

functional output of the application without examining the internal code structure. Testing was conducted on both Admin and User modules, covering scenarios. The test results confirmed that the system functionality aligns with the design specifications, handling both valid and invalid inputs appropriately as detailed in the test cases[8].

3. RESEARCH RESULTS

This section outlines the results of the system development, starting from the modeling of system processes using Unified Modeling Language (UML), followed by the user interface implementation, and concluding with system testing and comparison[9].

3.1 Use Case Diagram

A use case diagram is a graphical representation that illustrates the interaction between actors and the system, defining the functional requirements of the application. Through this diagram, the functions available within the Sikendi system can be identified clearly[10].

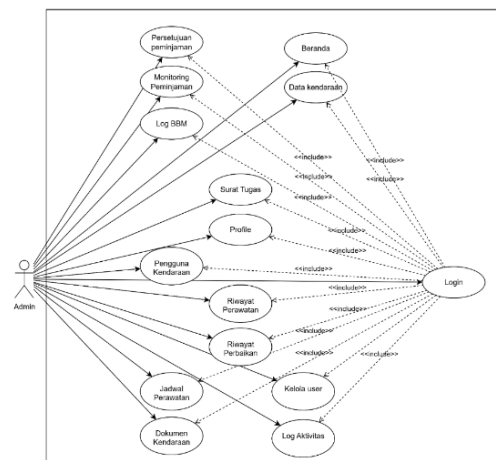


Figure 2. Use Case Admin

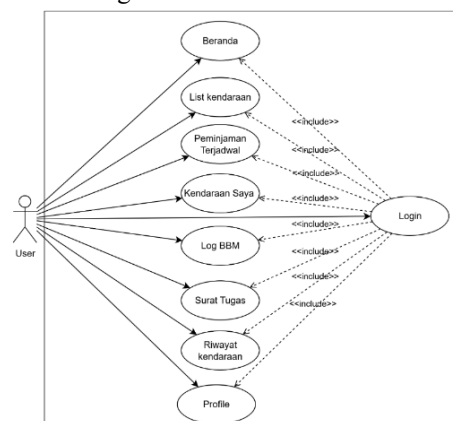


Figure 3. Use Case User

As shown in the diagrams above, the system involves two primary actors: the Admin and the User. Admin once logged in, the admin has full privileges to manage master data (vehicles and users), approve or reject borrowing requests, manage maintenance schedules, and monitor vehicle usage history and fuel logs .

The user role is designed for personnel who need to utilize the vehicles. Their capabilities include viewing the list of available vehicles, submitting borrowing requests, and inputting fuel usage logs (Log BBM) for the vehicles they are using[11].

3.2 Activity Diagrams

Activity diagrams represent the workflow of specific processes within the system. This section details the core activities for both Admin and User roles[12].

Activity Diagram for Admin Login The login

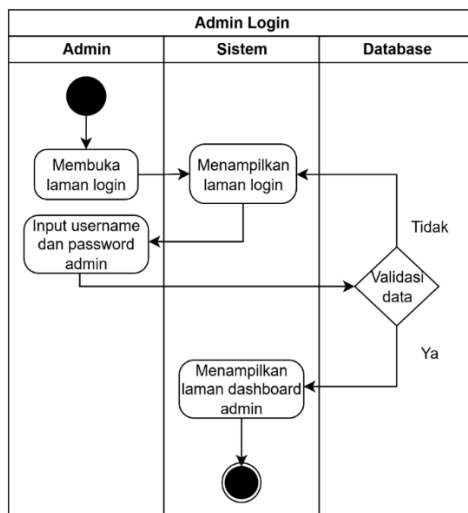


Figure 4. Activity Diagram for Admin Login

Activity Diagram for Admin Login The login process ensures secure access to the system. The flow begins with the admin accessing the login page and entering valid credentials. The system validates the data against the user_account table in the database. If valid, the admin is redirected to the dashboard; otherwise, an error message is displayed.

Activity Diagram for Managing Borrowing Requests (Admin)

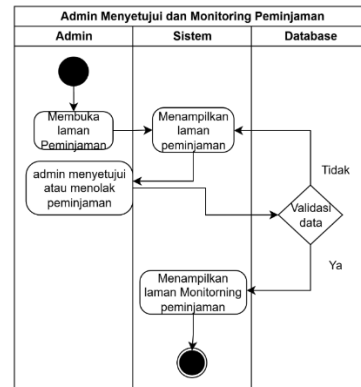


Figure 5. Activity Diagram for Managing Borrowing Requests (Admin)

This activity illustrates the admin's role in processing loan applications. The admin reviews incoming requests that are in "Waiting" status. The admin can then decide to either Approve or Reject the request. If approved, the system updates the vehicle status to "Booked" or "In Use" to prevent double booking.

Activity Diagram for Creating a Borrowing Request (User)

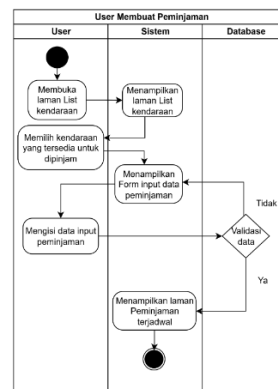


Figure 6. Activity Diagram for Creating a Borrowing Request (User)

This diagram shows the flow when a user applies for a vehicle. The user views the list of vehicles, selects an available unit, and fills in the borrowing form (dates and purpose). The system validates availability before saving the request with a "Pending" status.

3.3 Interface Display

This section presents the implementation of the user interface (UI) designed to be user-friendly and responsive.

Login Interface Display

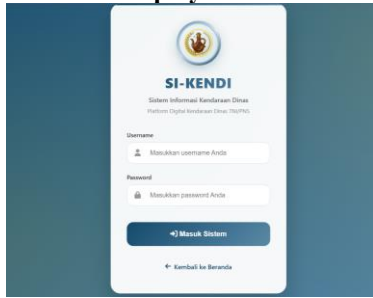


Figure 7. Login Interface Display

The login page serves as the gateway to the system. It requires users to input their username and password. This page is essential for maintaining security and ensuring that access rights correspond to the user's role (Admin or User)[13].

Admin Dashboard Display

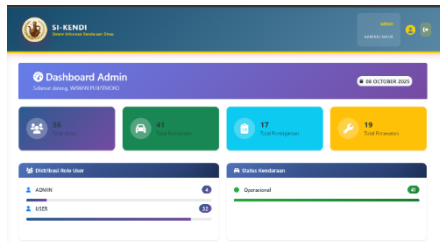


Figure 8. Admin Dashboard Display

The dashboard provides a comprehensive overview of the system's status. It displays key statistics such as the total number of vehicles, active loans, maintenance alerts, and recent user activities. This centralized view aids the leadership in monitoring operational readiness.

Borrowing Approval Interface Display

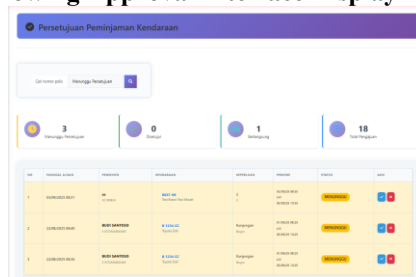


Figure 9. Borrowing Approval Interface Display

This interface allows the admin to manage incoming vehicle requests. It lists the details of the applicant, the selected vehicle, and the intended duration. The admin can execute approvals or rejections directly from this page, streamlining the administrative workflow[14].

User Vehicle List Interface

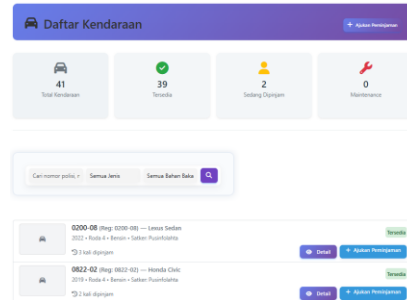


Figure 10. User Vehicle List Interface

For users, this interface acts as a catalog of official vehicles. It displays vital information such as vehicle type, condition, and current availability status (Available/In Use), enabling users to make informed decisions before submitting a request.

3.4 Testing

System testing was conducted using the Black Box Testing method to ensure that all functional requirements were met without examining the internal code structure. The testing covered various scenarios for both Admin and User module[15].

Table 1. Testing Admin

Module	Test Case Description	Expected Result
Login	Input correct Username and Password.	Login successful; redirected to Admin Dashboard.
	Input wrong Username, correct Password.	Login failed; error message displayed.
	Input correct Username, wrong Password.	Login failed; error message displayed.
	Leave one or both fields empty.	Login failed; warning message displayed.
	Access Dashboard URL without session.	Redirected back to Login page.

Loan Approval	Select 'Waiting' request and click Approve.	Status changes to 'Approved'; success notification appears.
	Select 'Waiting' request and click Reject.	Status changes to 'Rejected'; success notification appears.
	Apply filter for 'Waiting' status.	Table displays only requests with 'Waiting' status.
Monitoring	Filter by Start and End date.	Displays history only within the selected date range.
	Verify data with 'Ongoing' status.	Data includes Borrower, Vehicle, and Approver details.
User Mgmt	Input unique data (NRP/NIP) and Save.	New user saved; list updated.
	Change user details (e.g., Position) and Update.	Changes saved to database.
	Select user and use Reset function.	Password resets to default; notification appears.
Vehicle Data	Input all valid vehicle details.	Data saved; list updated.
	Input existing Registration Number.	Data rejected; validation warning appears.
	Update vehicle details (e.g., Color).	Changes saved to database; list updated.
	Select vehicle and confirm deletion.	Data removed from database.
	Search by Registration Number.	Displays only the matching vehicle.

Other Modules	Input Document Type and Expiry Date.	Document saved; status (Active/Expiring) displayed.
	Input Repair Date, Type, and Cost.	Data saved; history updated.
	Click Export CSV button.	User activity file downloaded in CSV format.

Table 2. Testing User

Module	Test Case Description	Expected Result
Login	Input correct Username and Password.	Login successful; redirected to User Dashboard.
	Input wrong Username or Password.	Login failed; error message displayed.
	Leave Username or Password empty.	Login failed; warning message displayed.
Vehicle List	Select 'Available' vehicle and fill form (dates/purpose).	Data saved; status set to 'Pending Approval'.
	Request vehicle on dates when it is 'In Use'.	Request rejected; error message regarding availability.
	Apply filter by Vehicle Type (e.g., 4-Wheels).	Displays only vehicles matching the filter.
Loan Status	View details of an active loan.	Status displayed as 'ONGOING'.
	View details of a loan rejected by Admin.	Status displayed as 'REJECTED'.

My Vehicle	Input valid Date, Liters, Cost, and Odometer.	Fuel log saved; database updated.
	Input non-numeric characters in Odometer field.	Data rejected; format validation error displayed.
	Access page during an active loan period.	Vehicle status displayed as 'In Use' / 'On Trip'.
Documents	Access 'My Assignment Letters'.	List of assigned letters displayed.
	Click PDF button on a specific letter.	Assignment downloaded as PDF file.
History	Access 'Vehicle History' page.	Displays summary of past vehicle usage and dates.

	physical application letters.	online submission and digital approval.
Usage Tracking	No clear digital trail; difficult to track users, specific purposes, and usage periods.	Creates a comprehensive digital audit trail for every borrowing transaction.
Maintenance Management	Recorded manually with no automated reminder system for service schedules.	Equipped with automated schedules and reminders for service, oil changes, and document renewals.
Decision Support	Reports are simple, unstructured, and require significant time to compile.	Generates comprehensive automated reports (tables and charts) regarding unit usage and maintenance.

3.5 Comparison

To evaluate the effectiveness of the developed system, a comparison was made between the previous manual management system and the new web-based *Sikendi* system.

Table 3. Comparison

Comparison Aspect	Current Manual System	Sikendi Web-Based System
Data Recording	Performed manually in physical logbooks or simple spreadsheets.	Performed digitally and centralized within a single web platform.
Data Access & Updates	Fragmented and difficult to access; retrieving historical data is time-consuming.	Rapid access via an intuitive web interface with real-time data updates.
Borrowing Process	Manual administration relying on	Automated workflow featuring

The comparison highlights significant improvements. The web-based system replaces fragmented manual recording with a centralized database, offers real-time access to data, and transforms the borrowing process from a slow paper-based mechanism to a transparent digital workflow. Additionally, the new system provides automated reporting and maintenance scheduling, which were absent in the previous manual process.

4. CONCLUSION

A. Conclusion Based on the analysis, design, and implementation of the *Sikendi* web-based application for the management of official vehicles at Pusinfolahta TNI, the following conclusions can be drawn:

1. The vehicle administration process, which previously relied on manual recording and fragmented spreadsheets, has been successfully modernized through the implementation of a web-based system. This transformation significantly increases data accuracy and administrative efficiency by

centralizing all asset information into a single database.

2. The system provides a robust mechanism for transparency and governance. By digitizing the borrowing workflow and providing statistical reports on vehicle usage and maintenance, the system minimizes the potential for misuse and supports strategic decision-making regarding operational readiness and resource allocation .

B. Recommendations During the development and implementation of the system at Pusinfolahta TNI, several recommendations have been identified to further enhance the system's capability and scope in future research:

1. The current system operates on a web-based platform. It is recommended that future development includes the creation of a mobile version (Android/iOS). This would enhance flexibility and accessibility, allowing users and administrators to track vehicle status and input fuel logs (Log BBM) from any location, particularly during field operations .
2. Future iterations should incorporate real-time notification features (via email or push notifications). This would streamline communication by instantly alerting users about loan approvals and notifying administrators about upcoming document expirations or overdue maintenance schedules .
3. It is recommended to expand the analytical modules to include Total Cost of Ownership (TCO) analysis. By integrating data from maintenance history and fuel logs, the system could automatically calculate operational costs per unit, thereby assisting leadership in planning asset regeneration or upgrades .

REFERENCES

- [1] S. S. Rosaulina, C. M. Sihotang, and S. S, "Rancang Bangun Sistem Inventory (Studi Kasus: UD. Asia Pratama Pekanbaru)," *NUANSA INFORMATIKA*, vol. 18, no. 2, pp. 35–40, Jul. 2024, doi: 10.25134/ILKOM.V18I2.156.
- [2] M. Zaidan and F. Abdullah, "Penggunaan RFID Sistem Informasi Parkir Berbasis Web," *NUANSA INFORMATIKA*, vol. 18, no. 1, pp. 121–127, Jan. 2024, doi: 10.25134/ILKOM.V18I1.86.
- [3] W. harjono and K. J. Tute, "Perancangan Sistem Informasi Perpustakaan Berbasis Web Menggunakan Metode Waterfall," *SATESI: Jurnal Sains Teknologi dan Sistem Informasi*, vol. 2, no. 1, pp. 47–51, Apr. 2022, doi: 10.54259/SATESI.V2I1.773.
- [4] H. Koç, A. M. Erdoğan, Y. Barjakly, and S. Peker, "UML Diagrams in Software Engineering Research: A Systematic Literature Review," *Proceedings 2021, Vol. 74, Page 13*, vol. 74, no. 1, p. 13, Mar. 2021, doi: 10.3390/PROCEEDINGS2021074013.
- [5] R. Fauzan, D. Siahaan, S. Rochimah, and E. Triandini, "A Different Approach on Automated Use Case Diagram Semantic Assessment," *International Journal of Intelligent Engineering and Systems*, vol. 14, no. 1, p. 2021, doi: 10.22266/ijies2021.0228.46.
- [6] A. Y. Chandra and P. W. Setyaningsih, "Benchmarking Local Development Environments: Analyzing the Performance of XAMPP, MAMP, and Laragon," *Bulletin of Computer Science Research*, vol. 5, no. 3, pp. 193–206, Apr. 2025, doi: 10.47065/BULLETINCSR.V5I3.493.
- [7] A. Jailani and M. A. Yaqin, "Pengujian Aplikasi Sistem Informasi Akademik menggunakan Metode Blackbox dengan Teknik Boundary Value Analysis," *Journal Automation Computer Information System*, vol. 4, no. 2, pp. 60–66, Jul. 2024, doi: 10.47134/JACIS.V4I2.78.
- [8] A. Sultansyah, A. S. Rahayu, I. Yudianta, P. Fauzi, E. N. Aripin, and S. A. Atmaja, "Pengujian Black Box Testing Pada Fitur Permohonan Informasi Publik Melalui Website Pemerintah Jawa Barat: Penelitian," *Jurnal Pengabdian Masyarakat dan Riset Pendidikan*, vol. 3, no. 4, pp. 5912–5919, Jun. 2025, doi: 10.31004/JERKIN.V3I4.1520.
- [9] K. Nadiyah and G. S. Dewi, "Quality Control Analysis Using Flowchart, Check Sheet, P-Chart, Pareto Diagram and Fishbone Diagram," *OPSI*, vol. 15, no. 2, pp. 183–188, Dec. 2022, doi: 10.31315/OPSI.V15I2.7445.
- [10] W. K. G. Assunção, S. R. Vergilio, and R. E. Lopez-Herrejon, "Automatic extraction of product line architecture and feature models from UML class diagram variants," *Information and Software Technology*, vol. 117, p. 106198, Jan. 2020, doi: 10.1016/j.infsof.2020.106198.

- 10.1016/J.INFSOF.2019.106198.
- [11] J. Shadiq, A. Safei, and R. W. R. Loly, "Pengujian Aplikasi Peminjaman Kendaraan Operasional Kantor Menggunakan BlackBox Testing," *INFORMATION MANAGEMENT FOR EDUCATORS AND PROFESSIONALS: Journal of Information Management*, vol. 5, no. 2, p. 97, Jul. 2021, doi: 10.51211/IMBI.V5I2.1561.
- [12] S. Al-Fedaghi, "Validation: Conceptual versus Activity Diagram Approaches," *International Journal of Advanced Computer Science and Applications*, vol. 12, no. 6, pp. 287–297, Jun. 2021, doi: 10.14569/IJACSA.2021.0120632.
- [13] L. Compagna, H. Jonker, J. Krochewski, B. Krumnow, and M. Sahin, "A preliminary study on the adoption and effectiveness of SameSite cookies as a CSRF defence," *Proceedings - 2021 IEEE European Symposium on Security and Privacy Workshops, Euro S and PW 2021*, pp. 49–59, Sep. 2021, doi: 10.1109/EUROSPW54576.2021.00012.
- [14] E. Suseno, E. Kurniadi, and D. Irawan, "Pengembangan Aplikasi Pengelolaan Laboratorium Komputer Dengan Menggunakan Metode Conten Based Filtering Berbasis WEB," *NUANSA INFORMATIKA*, vol. 18, no. 1, pp. 28–33, Jan. 2024, doi: 10.25134/ILKOM.V18I1.90.
- [15] P. Masci and C. A. Muñoz, "An Integrated Development Environment for the Prototype Verification System," *Electronic Proceedings in Theoretical Computer Science, EPTCS*, vol. 310, pp. 35–49, Dec. 2019, doi: 10.4204/EPTCS.310.5.

Implementation of an e-Voting System for Electing the Student Executive Board Chair Using QR Codes

Sintya Devi Januari¹, Darsanto², Muh Pauzan³

^{1,2,3}Wiralodra University, Indramayu

Email: ¹dsintyaa069@gmail.com, ²shantost.ft@unwir.ac.id, ³muhpauzan.ft@unwir.ac.id

Abstract

Information technology plays a crucial role in supporting all campus activities, including student organizations, which serve as a forum for developing leadership, interests, and talents. One of the main activities of this organization is the election of the Student Executive Board (BEM) chairman. However, elections are still conducted manually using paper ballots, which is time-consuming, inefficient, and prone to fraud. This research aims to design an e-Voting system based on QR Code technology using the Dynamic System Development Method (DSDM) as a system design solution that can ensure compliance with needs and changes during the development process. In this study, the resulting e-Voting system was analyzed using the ISO 9241-11 standard. Parameters measured in the ISO 9241-11 standard include effectiveness, efficiency, and user satisfaction. The results of the analysis show that the e-Voting system achieved 95% effectiveness, 94% efficiency and user satisfaction got a score of 73 as measured based on the System Usability Scale (SUS) included in the category above the eligibility standard, namely 68, so that the score of 73 got an adjective scale as Good, with a grade B, at the Acceptable level and included in the Passive category on the NPS scale. This study shows that the QR Code-based e-Voting system can be a reliable, efficient alternative and increase transparency and a satisfying experience for voters in the BEM chairman election process.

Keyword: e-Voting, QR-code, Dynamic System Development Method (DSDM).

Submission: November 28, 2025 Accepted: December 29, 2025 Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.524>

1. INTRODUCTION

The development of information technology in the world of education plays a very important role in supporting all aspects of activities in the campus environment, but not all aspects of these activities can be fully accommodated by technological developments, including student organization activities [1][2]. Student organizations are a forum for students to learn, develop interests and talents as well as leadership skills, one of the activities of student organizations is the election of the chairman of the Student Executive Board (BEM) [3].

The process of electing the BEM chairman until now still uses conventional or manual voting methods, which are held by voters checking or voting for the names of candidate candidates according to election instructions. The vote results are counted manually by opening each ballot paper witnessed by election witnesses. After the vote counting process is carried out by the election organizing committee, the results are obtained regarding the candidate who received the most votes. The election and vote counting process in this way certainly requires a relatively long time and is less efficient and tends to have a low transparency aspect and is also prone to fraud [3].

Various problems encountered in an election process, including the election of the Student Executive Board (BEM) president, can be overcome by using Electronic Voting (e-Voting) [4] with the support of increasingly widespread communication network technology and more affordable communication costs [5]. This implementation allows for a faster, more efficient, and more accurate voting process [6]. Likewise, smartphones, as a communication tool for all groups, can be used as a supporting tool for the e-Voting system [7].

Previous studies have been conducted to address the issues, including research by Mohamad Anas Sobarnas et al., M.F. Aziz, Agus Yulianto et al., using the waterfall method and black box testing in designing an e-voting system [2][3][6]. One of the shortcomings of the waterfall method is its inability to accommodate the dynamic needs of the BEM president election context.

The contribution of this research, considering the issues outlined in the paragraph above, is the concept of an e-Voting system for the BEM president election. Before the election process begins, the administrator first collects data to determine students as permanent voters and prospective candidates. Next, students must arrive at the polling station with their Student ID Card (KTM) to show to the officer. Then, they register

by entering their Student Identification Number (NPM). After successful registration, students proceed to the designated waiting area. As the registration process is completed, the committee monitors the queue and calls the earliest queue number. Voters then enter the voting booth to vote. Inside the booth, students scan a QR code containing the candidate's photo using their smartphone. Each student can only vote once.

The e-voting system for the BEM chairman election in this study utilizes a smartphone as an input device, a feature not seen in previous e-voting studies. The system development and testing method used is the Dynamic System Development Method (DSDM), an agile methodology characterized by dynamic characteristics at each stage of system development, and usability testing based on the ISO 9241-11 standard, with user satisfaction testing. This is expected to provide an effective solution tailored to the dynamic needs of the BEM chairman election e-voting system.

2. RESEARCH METHODS

The research was conducted using several stages, the following is the research flow:

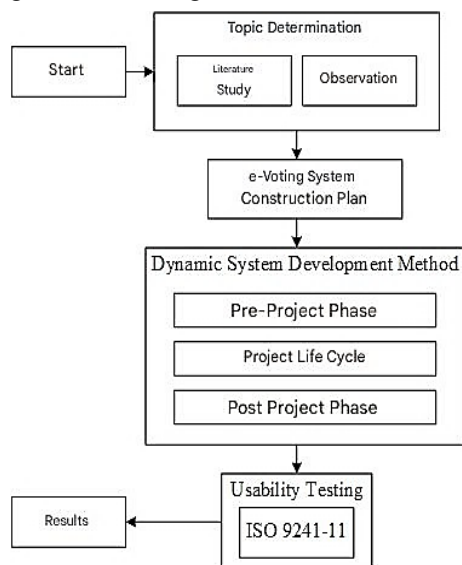


Figure 1. Research Flow

The research phase begins with determining the topic through literature review and observation. System illustrations are explained through a system construction plan. System development was conducted using the DSDM method, while usability testing followed the ISO 9241-11 standard.

2.1 System Development Methods

The method used in this research is Agile Software Development (DSDM) [8][9], with a

system development approach based on the DSDM stages, referring to an object-oriented platform and its utilities. The focus of system development is on user interaction, prioritizing user satisfaction with the e-voting system [10]. The following are the stages of the DSDM method:

1. Pre-project

This stage involves an interview process with users, which is obtained from this stage as the needs of the system and the determination of actors or access rights.

2. Lifecycle Project

This cycle consists of 5 sub-stages, namely feasibility study, business study, functional model iteration, design and manufacturing iteration, implementation.

3. Post-project

This stage is a performance measurement of the system that has been created and implemented to determine whether in the future it is necessary to improve, maintain and repair the system using a black box[11][12].

2.1.1 Pre-project

The system construction plan is shown in the following figure:

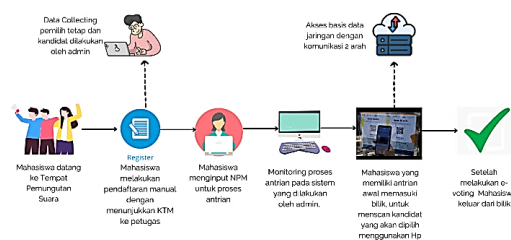


Figure 2. Illustration of System Construction Plan

Based on Figure 2, the system mechanism is determined by the applicable business rule plan as shown in the table below:

Table 1. Business Rule Plan

No	Business Rules
1	Collecting data on permanent voters and candidates by administrators (election organizers) is taken from active student data in the academic year of the election before the e-Voting process is implemented.
2	Students who wish to register as voters must first register by bringing their Student Identity Card (KTM) and then showing it to the e-Voting officer.
3	Students (voters) input NPM for the election queue process.
4	Monitoring of the queue process on the system carried out by the administrator.

No	Business Rules
5	Students who are in the first queue enter the booth to scan the candidates to be selected using their smartphones.
6	Each QR Code contains the candidate's serial number, the voting process is done by using a smartphone as a scanner input and pointing it at one of the candidate's QR Codes that you want to vote for at a distance of ± 35 cm. The QR Code standardization used is ISO/IEC 18004:2024, which supports UTF-8 data encoding and the Reed-Solomon Error Correction algorithm. This allows the QR Code to remain readable even if partially damaged or obscured [13], thereby increasing the accuracy of reading the QR Code of the BEM chairperson candidate selected by voters.
7	Each student can only vote or scan a candidate's QR Code once.

After determining the business rule plan, as described in Table 1 above, actors and system access rights were identified. The actors identified are active students in each year, and the committee consists of active BEM administrators. The actor identification is shown in the following table:

Table 2. System Actor Identification

No	Actor	Description
1	Student	Is a permanent voter in accordance with the permanent voter list with the provisions of the system business rules.
2	Committee	Is the administrator/organizer of election activities in the environment where the system is implemented

2.1.2 Lifecycle Project

The results of the identification in the previous sub-chapter are the basis for conducting a business study of the system, so that the system design is in accordance with functional and non-functional needs visualized in the use case, Activity and Class Diagram images of the e-Voting system.

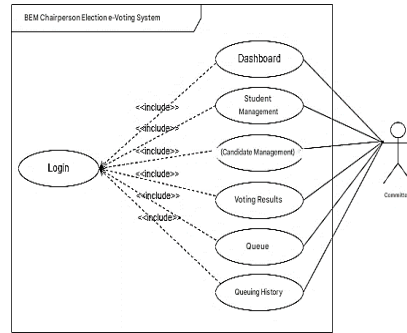


Figure 3. Use case diagram of the e-Voting system with the Committee actor

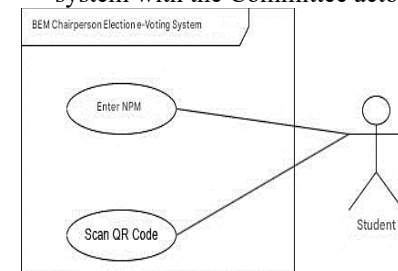


Figure 4. Use case diagram of the e-Voting system with student actors

Figures 3 and 4 are visualizations of the system with the involvement of actors with access rights according to the functional needs of the system.

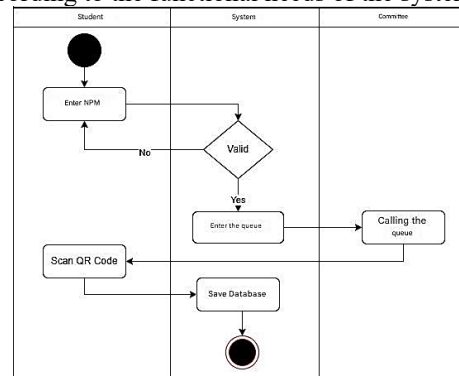


Figure 5. Activity Diagram of the e-Voting system on the Student actor side

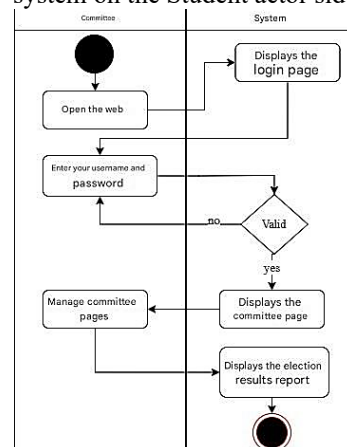


Figure 6. Activity Diagram of the e-Voting system from the Committee actor side

Figures 5 and 6 are visualizations of system activity with the involvement of actors with access rights according to the system's non-functional needs.

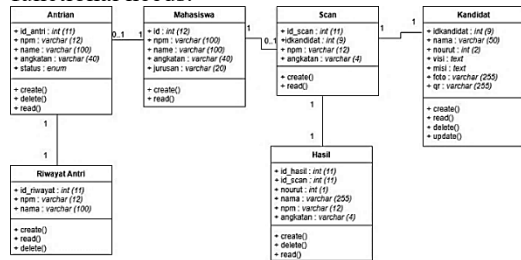


Figure 7. Class Diagram of the e-Voting system

Figure 7 shows the relationship between each class in the e-Voting system, which is a data relationship between entities that influence each other in the system mechanism process.

2.1.3 Post-project

This stage is a performance measurement of the system that has been created and implemented to determine whether in the future it is necessary to improve, maintain and repair the system using black box testing with validation methods.

2.2 System Performance Test

The system performance test uses the ISO 9241-11 standard to measure usability (ease of use) [14], this standard is a reference in evaluating the usability of applications based on three aspects [15], namely effectiveness, efficiency and user satisfaction [16]. Performance testing on each test aspect is given to respondents who are stakeholders in the domain of the BEM chairman election e-voting system, namely active students who are recorded in the faculty database and are included as permanent voters, faculty and study program leaders, educational staff (lecturers) and faculty education staff.

3. RESEARCH RESULTS

The Agile DSDM framework stage, which is the post-project phase, involves coding according to the requirements defined in the system design in the previous stage. The following is the resulting e-Voting system interface.

3.1 e-Voting System Interface

Figure 8. e-Voting System Login Page

Figure 8 above is the system interface for system user actors when starting access according to the access rights that have been obtained.

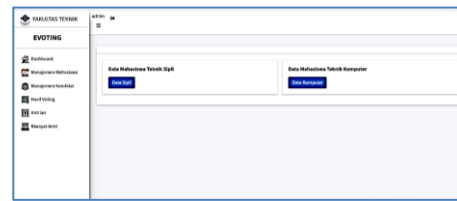


Figure 9. Student Management Page of the e-Voting System

The interface in Figure 9 shows the data collection process in CSV file format carried out before the election is carried out and coordinated with the academic section as an inseparable stage of the e-Voting system.



Figure 10. e-Voting System Candidate Management Page

Figure 10 is a display of the BEM chairperson candidates who will be selected by being given a QR code and a photo of the candidate in the e-Voting system.

No Urut	Nama	Jumlah Suara
1	Sara MRP Program	5
2	Shirley Dewi Jansari	0
3	Milly Martini	0

Hasil Voting Berdasarkan Angkatan				
Angkatan	Sara MRP Program	Shirley Dewi Jansari	Milly Martini	
2019	3	0	0	
2020	0	1	1	
2021	1	0	1	
2022	0	1	4	
2023	1	0	0	

Figure 11. Voting Results Page of the e-Voting System

The interface in Figure 11 is a page on the administrator/committee side as a monitoring medium for the voting results that are currently running based on the year of the student intake that has chosen the BEM chairperson candidate.

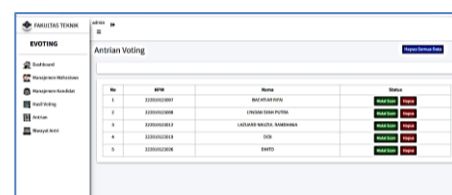


Figure 12. e-Voting System Queue Page

Figure 12 shows the committee /administrator interface for monitoring the voter queue during the ongoing voting process.

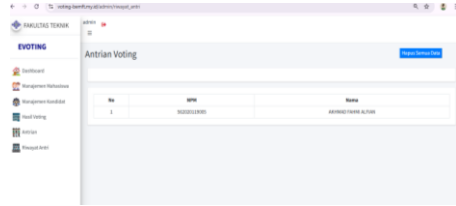


Figure 13. e-Voting System Queue History Page

Figure 13 is a queue history page from the voting process and is operated by the committee/administrator as system backup data if a bug occurs while the system is running in the election process.

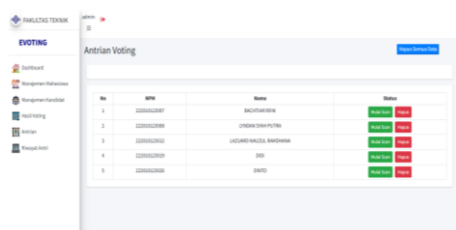


Figure 14. Student Queue Page of the e-Voting System

Figure 14 shows the student queue page, and the committee controls the actual queue to ensure the process does not deviate from system requirements, allowing potential bugs to be resolved. If a solution to the actual queue constraints caused by non-technical factors cannot be resolved, the committee will take action to exclude the system procedure through this interface.



Figure 15. QR Code Scan Page of the e-Voting System

Figure 15 shows the student-side QR code scanner interface, which scans the QR code of the candidate to be selected in the voting booth. This tool allows voters to access the site using a smartphone, based on a web page provided by the committee.

3.2 Testing the e-Voting System

The testing phase is necessary to determine the functionality of each feature in the e-Voting system and ensure the system runs as expected.

Black Box testing is conducted through simulations involving participants who are stakeholders and potential system users, ensuring that the testing results reflect the real-world conditions in the environment where the system will be implemented. The test results can be seen in Tables 3 through 9 below:

Table 3. Black Box Login Test Results

No	Testing Scenario	Testing Method	Test Results	Conclusion
1	Username and password are not filled in then click the login button	Username: (blank) Password: (blank)	System rejects, admin must enter username and password	Valid
2	Fill in the username and password then click the login button	Username: (correct) Password: (correct)	The system will accept login access and then display the admin page.	Valid
3	Fill in the username and leave the password blank, then click the login button.	Username: (Filled in) Password: (Empty)	System rejects, admin must enter password	Valid
4	Username and password are filled in but incorrectly then click the login button	Username: (incorrect) Password: (incorrect)	The system rejects and displays an error, returning to the login page.	Valid

Table 4. Results of Black Box Testing of Student Management

No	Testin g Scena rio	Testing Method	Test Results	Conclus ion
1	Import studen t files in CSV format	File : mahasiswa .csv	The system accepts and the file is successf ully imported into the database .	Valid
2	Import studen t files in format s other than CSV	File : Mahasiswa .pdf	The system rejected and the file was not successf ully imported into the database .	Valid

Table 5. Candidate Management Black Box
Test Results

No	Testin g Scenar io	Testing Method	Test Results	Conclusi on
1	<i>The add candid ate form is filled in</i>	<i>Form: (complet ed)</i>	The system accepts and the form is successfu lly added.	Valid
2	<i>Empty add candid ate form</i>	<i>Form: (blank)</i>	The system rejects and displays the form again.	Valid

Table 6. Black Box Test Results for Voting
Results

No	Testing Scenari o	Testin g Metho d	Test Result s	Conclusio n
1	Delete all data menu	Click delete all data	The system accept s, all results on the voting results are deleted .	Valid

Table 7 Queue Black Box Test Results

No	Testing Scenari o	Testin g Metho d	Test Results	Conclusio n
1	Start scan menu	Click the start scan menu	The system accepts that the student's NPM is in the queue history.	Valid
2	Delete menu	Click the delete menu in the status	The system accepts that the NPM has been successfull y deleted from the queue list.	Valid
3	There are still students in the queue history but click start scan again	Click the start scan menu in the next queue	The system rejects, a notificatio n appears that there is still a queue	Valid
4	Menu delete all data in the queue	Click the delete all data menu	The system accepts, all data in the queue is deleted.	Valid

Table 8. Black Box Test Results for Queue History

No	Testin g Scena rio	Testing Method	Test Results	Conclu sion
1	Studen ts input NPM in the queue	NPM : (562020119 016)	The system accepts, a notifica tion appears that the NPM input was success ful.	Valid
2	Studen ts input the same NPM	NPM : (562020119 016)	The system rejects, an error notifica tion appears that the NPM is already register ed.	Valid
3	Studen ts input NPM that is not registe red in the databa se	NPM : (100204527 189)	The system rejects, an error notifica tion appears that the NPM is not register ed.	Valid
4	Studen ts input NPM more than 12 digits	NPM : (562020119 0161)	The system rejects, an error notifica tion appears that the NPM must be 12 digits.	Valid
5	NPM is not filled in then click the submit button	NPM : (blank)	The system rejects, students must enter their NPM first.	Valid

Table 9. Results of Black Box Testing of Student Queues

No	Testing Scenario	Testing Method	Test Results	Conclusion
1	Delete all data menu	Click delete all data	The system accepts, all results in the queue history are deleted.	Valid

Based on the results of the black box validation method test in tables 3 to 9, all test scenarios for valid features were obtained, so the e-Voting system can be released for use in real practice. In testing this validation method, the test environment conditions are ensured to be ideal so that the potential for bugs can be minimized. However, if a bug occurs in a test item, the test is repeated until the potential bug is resolved, after which testing is continued on the next test item. The potential for bugs in the test item is based on the suitability of the procedures, mechanisms, and predetermined system rules. In the following development, testing requires algorithm testing to complement the results of the validation test and minimize system bugs that occur.

4. DISCUSSION

Usability is measured using ISO 9241-11, which has three aspects: effectiveness, efficiency, and user satisfaction. These aspects are explained in the following subchapters.

4.1 Effectiveness

Effectiveness is measured by giving respondents several tasks to complete all the tasks given.

Table 10. System Testing Task List

Task Code	Information
T1	Respondents log in to the system.
T2	Respondents input student data into the system.
T3	Respondents added candidates.
T4	Respondents edit candidates.
T5	Respondents input the queue.
T6	Respondents perform QR scans.

The number of successful tasks completed is calculated based on the results. Successful tasks are

given a score of 1, while unsuccessful tasks are given a score of 0.

Table 11. System Effectiveness Testing

Rp n	Time (seconds)						Total Time (seconds)	
	T 1	T 2	T3	T 4	T 5	T 6	Suc ceed	All over
R1	7	13	120	15	6	17	178	178
R2	9	13	40	23	8	15	85	108
R3	8	17	80	25	10	20	160	160
R4	7	15	35	18	7	17	99	99
R5	10	16	60	15	6	19	110	126
R6	8	17	75	20	5	15	140	140
R7	7	19	40	28	9	13	88	128
R8	9	15	50	19	8	14	115	115
R9	9	12	90	23	6	16	156	156
R10	7	16	45	17	7	15	107	107

The data in Table 11 represents the results of the tasks carried out by each respondent, each acting as a committee member (administrator) and a student (voter). In the task column of Table 10, tasks 1 through 4 were carried out by the administrator, while tasks 5 through 6 were carried out by the students as voters in the e-Voting system.

The calculation of the success rate of the task is calculated using the Completion Rate with the following formula.

$$\text{Completion Rate} = \frac{57}{60} \times 100 = 95\%$$

The results obtained from the effectiveness value of the e-Voting system in the BEM election simulated at the Faculty of Engineering were 95%. Based on the obtained value, the system's effectiveness aspect was high, but there were several tasks (T2 and T4) that were carried out by respondents (R2, R5, R7) that were not successfully completed so that the potential for bugs in the system was one of the determining factors for the system's success rate.

4.2 Efficiency

Efficiency is measured by observing the time it takes respondents to complete all assigned tasks. Using the tasks in Table 10, the completion

rate is obtained, which is then calculated using the Overall Relative Efficiency (ORE) formula as follows:

Table 12. System Efficiency Testing

R pn	Task						Numb er of Succe ssful Tasks	Assign ment Succes s Percen tage (%)
	T 1	T 2	T 3	T 4	T 5	T 6		
R1	1	1	1	1	1	1	6	100
R2	1	1	1	0	1	1	5	83
R3	1	1	1	1	1	1	6	100
R4	1	1	1	1	1	1	6	100
R5	1	0	1	1	1	1	5	83
R6	1	1	1	1	1	1	6	100
R7	1	1	1	0	1	1	5	83
R8	1	1	1	1	1	1	6	100
R9	1	1	1	1	1	1	6	100
R10	1	1	1	1	1	1	6	100

Table 12 shows the results of the efficiency test, measured by the time each respondent worked on. The processing time was recorded throughout the process, regardless of whether the task was successfully completed or not. The time column in Table 12 shows the processing time of each respondent, while the total time (seconds) column represents the total time for each successful task and the total time. The tasks that have been completed, as shown in Table 3, provide the time data required by each respondent to complete all tasks, namely:

$$\text{ORE} = \frac{178 + 85 + 160 + 99 + 110 + 140 + 88 + 115 + 156 + 107}{178 + 108 + 160 + 99 + 126 + 140 + 128 + 115 + 156 + 107}$$

$$\text{ORE} = \frac{1238}{1317} \times 100 = 94\%$$

The total time taken by all respondents to complete the tasks was 1,317 seconds, while the total time successfully completed by all respondents was 1,238 seconds. Based on these calculations, the efficiency of the e-Voting system for the Faculty Student Executive Board (BEM) elections reached 94%. The efficiency aspect percentage in the e-Voting system simulation test is influenced by the level of response focus when carrying out each task. The results of the efficiency aspect test with the number of respondents (10 respondents) and under ideal testing conditions with a score of 94% are

considered high compared to the results of similar and previous studies that averaged below 90% [2][3][6].

4.3 User Satisfaction

User satisfaction was measured by administering a questionnaire to respondents, then calculating the results using the System Usability Scale (SUS) formula as follows:

Table 13. SUS Values Based on Total Questionnaire Scores

Rp n	Tota l Scor e	Amoun t x 2,5	Rp n	Tota l Scor e	Amoun t x 2,5
R1	25	62,5	R16	27	67,5
R2	26	65,0	R17	27	67,5
R3	32	80,0	R18	28	70,0
R4	26	65,0	R19	30	75,0
R5	30	75,0	R20	32	80,0
R6	36	90,0	R21	28	70,0
R7	24	60,0	R22	28	70,0
R8	27	67,5	R23	31	77,5
R9	25	62,5	R24	28	70,0
R10	40	100,0	R25	30	75,0
R11	26	65,0	R26	30	75,0
R12	24	60,0	R27	31	77,5
R13	30	75,0	R28	28	70,0
R14	35	87,5	R29	30	75,0
R15	30	75,0	R30	31	77,5
Average					73

The average user satisfaction level, after being calculated using the SUS method, obtained a score of 73. Based on the SUS scale, this score is categorized as Good, with a grade of B, at the Acceptable level and included in the Passive category on the NPS scale.

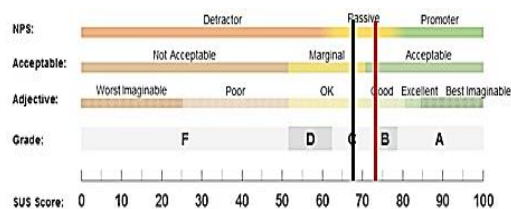


Figure 17. Results of the User Satisfaction Measurement Scale with SUS

The results of usability testing measurements of the e-Voting system for the Faculty of Engineering Student Executive

Board election based on three aspects of ISO 9241-11 are as follows.

Table 14. Results of System Usability Measurement with ISO 9241-11

Effectiveness	Efficiency	User Satisfaction
95%	94%	73

Based on the ISO 9241-11 standard, the e-Voting System has an effectiveness score of 95%, meaning respondents were able to complete tasks accurately. The efficiency score of 94% indicates that the system facilitated respondents' quick and easy use of its features. Meanwhile, the user satisfaction score of 73 indicates that respondents found the system technically quite user-friendly.

5. CONCLUSION

This e-Voting system is designed to simplify the election of the BEM chairman. The system design using the DSDM method has proven to be easy to implement, relevant and effective and in accordance with user needs. Effectiveness testing based on ISO 9241-11 showed a success rate of 95%, with 57 out of 60 tasks successfully completed. In terms of efficiency, the system recorded a score of 94%, with a total task completion time of 1,238 seconds out of 1,317 seconds, indicating only 6% wasted time. User satisfaction testing resulted in a SUS score of 73, which is included in the Good category with a grade B, proving that this system is able to meet needs effectively and efficiently to support the e-Voting process. In the broader application domain of this e-Voting system, an in-depth testing process is required, especially aspects of the ISO 9241-11 standard, by involving greater resources starting from time and respondents according to the scope of the system domain with a measurable testing duration, so that system requirements can be met and in the same context the system can be implemented well.

6. SUGGESTION

The research conducted using the testing method used has yielded high levels of acceptance and satisfaction from several aspects of the SUS scale, consistent with the test results. However, particularly in the NPS aspect, automation features need to be added to the queue menu to improve system functionality performance according to user needs. Furthermore, system testing using other methods is recommended as a comparison to previous test results. To ensure a higher level of accuracy, effectiveness and efficiency testing is recommended to use 50% of the user satisfaction test sample. To support the system scalability aspect, program algorithm testing is required for

testing system service features and integration with other technology service platforms that provide security features and high performance in a reliable system infrastructure.

REFERENCE

- [1] F. Setyawan and F. I. Pratama, "Rancang Bangun Sistem E-Voting Pemilihan Ketua OSIS SMA Mardisiswa Semarang Berbasis Web," *J. Inform. Dan Rekayasa Perangkat Lunak*, vol. 2, no. 2, p. 154, 2020, doi: 10.36499/jinrpl.v2i2.3591.
- [2] A. Yulianto, D. H. Yusuf, and Frimansyah, "Penerapan E-Voting Dengan Metode Waterfall Untuk Pemilihan Ketua OSIS Pada SMP PGRI Parung Panjang Bogor," *Ris. Dan E-J. Manaj. Inform.*, vol. 3, no. 2, pp. 66–73, 2019.
- [3] M. F. Aziz, "Rancang Bangun Informasi Pemilihan Dewan Eksekutif Mahasiswa (DEMA) Di Universitas Islam Negeri Alauddin Makassar Berbasis Website," 2021.
- [4] A. D. Andari, "Mengenal Pengertian Voting dan E-Voting Beserta Kelebihan dan kekurangannya," 2024.
- [5] A. Adriansyah, "Analisis Pemanfaatan Teknologi QR Code pada Sistem Electronic Voting (E-Voting) untuk Pemilihan Kepala Daerah," vol. 4, no. 2, pp. 91–102, 2020.
- [6] M. A. Sobarnas, P. Sukanto, and Y. Nuryaman, "Rancang Bangun Aplikasi E-Voting Multi Instansi Berbasis Web Dengan QR Code," *Infotech J. Inform. Teknol.*, vol. 2, no. 2, pp. 61–70, 2021.
- [7] F. Noor, "Analisa Penguunaan Smartphone dalam Pertemanan Di SMA Negeri 4 Palangkaraya," *Teori Komun.*, vol. 2, no. 6, pp. 13–34, 2019.
- [8] N. Lutfiani, P. Harahap, Q. Aini, A. Dimas, A. R. Ahmad, and U. Rahardja, "InfoTekJar : Jurnal Nasional Informatika dan Teknologi Jaringan Attribution-NonCommercial 4.0 International. Some rights reserved Inovasi Manajemen Proyek I-Learning Menggunakan Metode Agile Scrumban," *InfoTekJar J. Nas. Inform. Dan Teknol. Jar.*, vol. 5, no. 1, pp. 96–101, 2020.
- [9] Ressa, "Sistem Electronic Voting (e-Voting) Berbasis Web Pada Pemilihan Ketua OSIS," 2019.
- [10] H. Handayani, K. U. Faizah, A. Mutiara Ayulya, M. F. Rozan, D. Wulan, and M. L. Hamzah, "Perancangan Sistem Informasi Inventory Barang Berbasis Web Menggunakan Metode Agile Software Development Designing a Web-Based Inventory Information System Using the Agile Software Development Method," *J. Test. Dan Implementasi Sist. Inf.*, vol. 1, no. 1, pp. 29–40, 2023.
- [11] Dewi Ayu Nur Wulandari, Muhammad Dika Atthariq, Wahyu Dwi Nanda, and Lestari Yusuf, "Implementasi Dynamic System Development Method (DSDM) Pada Sistem Informasi Manajemen Bengkel Mobil Berbasis Web," *JSil J. Sist. Inf.*, vol. 8, no. 1, pp. 10–17, 2021, doi: 10.30656/jsii.v8i1.2979.
- [12] W. Nyunando and D. Nasien, "Implementasi Agile Dynamic System Development Method Berbasis Web Pada Sistem Penggajian," *J. Mhs. Apl. Teknologi Komputer dan Informatika.*, vol. 2, no. 1, pp. 33–38, 2020.
- [13] OBP, "ISO/IEC 18004:2024(en)." Online Browsing Platform (OBP) Version 4.43.2. [Online]. Available: <https://www.iso.org/obp/ui/en/#iso:std:iso-iec:18004:ed-4:v1:en>
- [14] A. P. Sukma, R. Yusuf, and R. H. Dai, "Analisis Pengukuran Usability Sistem Informasi Manajemen Baznas (Simba) Menggunakan Metode System Usability Scale (Sus)," *Diffus. J. Syst. Inf. Technol.*, vol. 3, no. 2, pp. 224–231, 2023.
- [15] ISO, "Ergonomic requirements for office work with visual display terminals. Part 11: Guidance on usability," *ISO No 924111*, vol. 2008, no. February 9, p. 22, 1994.
- [16] S. Oktaviani, C. Wiguna, and A. Priyanto, "Application of System Usability Scale (SUS) method in testing the usefulness of information system Student Creativity Program (PKM) based on website," *AIP Conf. Proc.*, vol. 2658, no. December, 2022, doi: 10.1063/5.0107302.

A Study on the Application of the Iterator Pattern to Enhance Software Reusability

Siti Herawati¹,

Jl. HS.Ronggo Waluyo, Puseurjaya, Telukjambe Timur, Karawang, Jawa Barat 41361, Indonesia

Telp : (0826) 267641177

Email : 2210631170153@student.unsika.ac.id¹,

Abstract

The iterator design pattern is an effective approach in software development to achieve greater modularity and reusability. This study analyzes the application of the iterator pattern in enhancing software reusability, particularly within the object-oriented programming paradigm. The research includes a literature review and hands-on implementation using Python and Java to evaluate the separation of iteration logic from data structures. The results indicate that the pattern supports the principle of separation of concerns, enabling the development of code that is more flexible, structured, and maintainable. Code evaluation shows improvements in readability and the ease of component reuse. Therefore, this study provides a practical contribution to the understanding of design patterns and is expected to serve as a useful reference for software developers in building efficient and sustainable systems

Keywords—Design Pattern, Iterator, Reusability, Object-Oriented Programming

Submission: December 1, 2025 Accepted: December 22, 2025 Published: January 20, 2026
DOI Article: <https://doi.org/10.25134/ilkom.v20i1.525>

1. INTRODUCTION

Modern software systems are rapidly evolving in both scale and complexity. The demand for fast, stable, and efficient feature delivery encourages software engineering practices that emphasize modularity, separation of concerns, and the ability to reuse code components (reusability). Empirical studies indicate that improvements in reusability are strongly influenced by appropriate refactoring processes [1]. Other literature asserts that the internal quality of software can be enhanced through the selection of suitable design patterns, which support more consistent and maintainable code structures [2].

The iterator design pattern is a behavioral pattern that provides a standardized mechanism for traversing elements within a collection without exposing its internal structure. This pattern supports principles such as separation of concerns and loose coupling, both of which are fundamental to achieving flexible software architecture [3]. In addition, the iterator pattern has been proven to enhance the reusability of traversal logic across various programming contexts.

Refactoring plays a crucial role in increasing software reusability, and multiple studies show that effective refactoring can improve high-level

design structures, including packages, classes, and methods [4]. Refactoring challenges often arise from the presence of code smells, which require structural fixes and consistent use of design patterns [5]. Studies on refactoring-related questions from developer communities also show that many developers struggle to determine optimal refactoring strategies, making design patterns such as the iterator a stable and reliable solution [6].

In modern system architectures, such as microservices, design patterns have a significant impact on system quality. Empirical studies show that well-designed patterns can balance quality trade-offs involving performance, modularity, and reusability [7]. Other research demonstrates that certain design patterns influence microservice performance, especially in large-scale scenarios [8]. Experimental evaluations further confirm that choosing the right architectural pattern can improve system performance without reducing structural flexibility [9].

However, the application of design patterns—including iterator—is not without challenges. Recent studies note that misuse or incorrect implementation of design patterns can lead to bad smells, especially in classical gang of four patterns [10]. Similar risks appear in object-oriented design models that fail to use patterns

properly, which can reduce reusability and maintainability [11]. Therefore, proper documentation, consistent implementation, and a deep understanding of the iterator pattern are essential before applying it in real development environments.

Various implementations of the iterator pattern have continued to evolve, including custom iterators, lazy iterators, generator-based iteration, and stream-based iteration commonly used in java, python, and c#. Research shows that pattern-based approaches improve development efficiency and enhance component reusability [12]. Studies on design patterns in microservices and devs modeling also confirm that design patterns contribute to building more modular and extensible components.

In the context of local software development in Indonesia, several studies show that implementing design patterns can improve software quality and system maintainability. Research involving prototype-based user interface design for car rental applications demonstrates improved structural effectiveness in web-based systems [13]. The development of web-based financial applications also shows enhanced maintainability through modular design structures [14]. Similarly, studies on tourism information systems highlight how modular structures improve readability and component reusability [15].

2. LITERATURE REVIEW

2.1 Software Reusability

Reusability is the ability of a software module to be used again in other systems or projects without requiring significant modifications. Reusability can reduce development costs, accelerate production time, and improve software consistency and quality. Reusable components typically have a modular structure, clear interfaces, and well-written documentation. The application of design principles such as *low coupling* and *high cohesion* also supports the improvement of reusability.

2.2 Iterator Design Pattern

The iterator pattern is one of the behavioral design patterns used to access elements in a data collection sequentially without exposing its internal structure. This pattern separates the iteration logic from the data structure, thereby supporting flexibility and reusability. In its implementation, this pattern

uses a special object (an iterator) that is responsible for the traversal process. It is commonly used in programming languages such as Java and Python and strongly supports the principle of separation of concerns.

3. METHODOLOGY

This method was chosen because it is suitable for exploring technical phenomena in depth through conceptual analysis, literature review, and exploratory code implementation. The research was carried out in several structured stages, including:

3.1 Literature Review

The study began with a literature review to understand design patterns, particularly the iterator pattern, and its contribution to reusability in software development. The primary reference examined was Design Patterns: Elements of Reusable Object-Oriented Software by [not in provided list], supported by scholarly journals, conference articles, and recent online sources. The literature also included previous studies that investigated the impact of design patterns on software quality aspects such as readability, modularity, and maintainability. This approach provided a solid theoretical foundation and contextual understanding of iterator pattern application.

3.2 Code Implementation and Analysis

After establishing a conceptual understanding through the literature, the research proceeded with the practical implementation of the iterator pattern in two widely-used object-oriented programming languages: Python and Java. These languages were chosen to demonstrate how the iterator pattern can be adapted to different paradigms while maintaining the same object-oriented foundation.

Two case studies were developed in this stage:

- Case Study A (Python): A custom data structure was developed with iterator support using the `__iter__()` and `__next__()` methods, in accordance with Python's iterator protocol. This case focused on a simple collection such as a list of objects that required various traversal strategies.
- Case Study B (Java): The Iterator and Iterable interfaces were implemented, and a custom iterator class was created to allow flexible traversal, including both forward and reverse iteration.

Each case study included code samples before and after applying the iterator pattern. The initial (baseline) code used a conventional iteration approach with traversal logic embedded directly in the data structure, while the refactored code

separated the iteration logic into a distinct iterator object. The analysis focused on code structure, responsibility grouping, level of coupling between classes, and the potential for code reuse in other application contexts.

3.3 Reusability Evaluation

The final stage involved evaluating the impact of applying the iterator pattern on software reusability. This evaluation was conducted qualitatively by comparing the code before and after refactoring based on several key reusability indicators:

- Separation of Concerns: Assessing the extent to which iteration logic is separated from data and business logic.
- Modularity: Evaluating the ability of the code to be divided into independent functional units.
- Readability: Measuring improvements in code clarity and consistency following the implementation of the iterator pattern.
- Reusability: Analyzing the potential for reusing the iterator components in other projects or with different data collections.

The evaluation also considered good software engineering practices, such as the open/closed principle, low coupling, and high cohesion. The findings from this evaluation were used to draw conclusions about the effectiveness of the iterator pattern in supporting the development of more structured, efficient, and maintainable software systems.

4 RESULTS AND DISCUSSION

4.1 Implementation Results

In object-oriented software development, managing and traversing data collections is a common task. Before applying the Iterator pattern, element traversal within a collection is typically performed manually using index-based loops, resulting in repetitive and poorly structured code. By implementing the Iterator pattern—both in Python and Java—data collection traversal can be performed in a more systematic way, separated from the logic of data storage. This approach enhances code readability and maintainability.

- Example of Implementation Before Using Iterator in Python:

```
def print_books(self):  
    print("== Book List ==")  
    for i in range(len(self.books)):  
        print(f'{i+1}. {self.books[i]}")
```

Figure 1 : Before using iterators

The `print_books(self)` method is used to display the list of books stored in the collection in a structured format. This method prints the list

title and then iterates through the collection using numeric indexing to display each book in sequence. This presentation makes it easier for users to verify the data and supports debugging and validation processes in text-based software development.

- Example of Implementation After Using Iterator in Python:

```
class BookIterator:  
    def __init__(self, books): self.books = books  
    self.idx = 0  
    def __iter__(self): return self  
    def __next__(self):  
        if self.idx < len(self.books):  
            b = self.books[self.idx] self.idx += 1; return b  
        raise StopIteration  
  
class BookCollection:  
    def __init__(self): self.books = []  
    def __iter__(self): return BookIterator(self.books)
```

Figure 2 : After using iterators

The code above defines an iterator for the `BookCollection` class. `BookIterator` allows the book collection to be accessed sequentially using a `for` loop without the need for manual indexing. By implementing the `__iter__()` and `__next__()` methods, the code simplifies the data traversal process and enhances modularity and readability when manipulating the book collection.

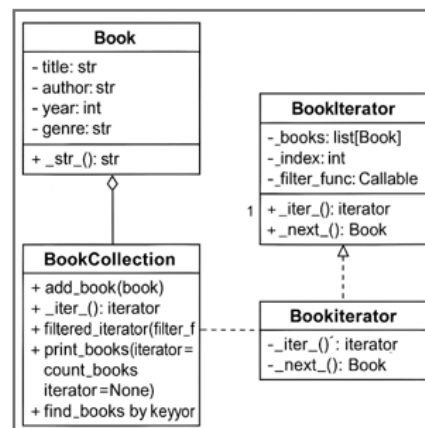


Figure 3 : Class diagram UML (Python)

The class diagram (Python) illustrates a book collection system. The `Book` class contains basic attributes of a book, such as title, author, year, and genre. `BookIterator` is a custom iterator with filtering capabilities, allowing flexible iteration over the book collection. `BookCollection` stores a list of books and provides various methods such as adding books, printing, filtering by genre/year, searching by keywords, and sorting by title.

- Example of Implementation Before Using Iterator in Java:

```
public void printSongsForward() {
    System.out.println("== Daftar Lagu ==");
    for (int i = 0; i < songs.size(); i++) {
        Song song = songs.get(i);
        System.out.println((i + 1) + ". " +
            song.getTitle() +
            " - " + song.getArtist());
    }
}
```

Figure 4 : Before using iterators

The printSongsForward() method is used to display the list of songs in the Playlist object. By using index-based iteration, each song is printed sequentially along with its number, title, and artist name. This function provides a structured text representation of the playlist content, making it easier for users to view and verify the added songs.

d. Example of Implementation After Using Iterator in Java:

```
class SongIterator implements Iterator<Song> {
    private List<Song> songs; private int i = 0;
    public SongIterator(List<Song> s) {songs = s;
    }
    public boolean hasNext() {return i <
    songs.size(); }
    public Song next() { return songs.get(i++); }
}

class Playlist implements Iterable<Song> {
    private List<Song> songs = new
    ArrayList<>();
    public Iterator<Song> iterator() {
        return new SongIterator(songs);
    }
}
```

Figure 5 : After using iterators

This code applies the Iterator pattern in Java to enable structured traversal of Playlist objects. SongIterator manages manual iteration over the list of songs, while Playlist provides iterative access to the song collection by implementing the Iterable interface. This approach improves modularity and flexibility in data management and separates the traversal logic from the data structure.

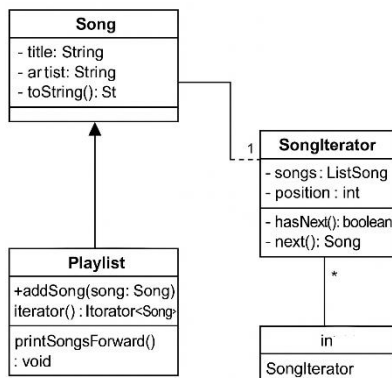


Figure 6 : Class diagram UML (Java)

The class diagram for the Java implementation illustrates a music playlist system. The Song class represents the song entity, containing attributes such as title and artist. SongIterator is a custom iterator that allows traversal of the song collection one item at a time, while the Playlist class stores the list of songs and provides methods for adding songs and printing the song list using the iterator.

4.2 Code Feature Comparison Table

The following table compares the coding approaches before and after the implementation of the iterator pattern and filtering functions in Python and Java. The focus is on aspects such as collection iteration, filtering capabilities, separation of concerns, and development flexibility. The new approach improves code reusability and facilitates the development of additional features.

Table 1 : Comparison of code features before and after applying Iterator

Fitur	Python (Before)	Python (After)	Java (Before)	Java (After)
Collection Iteration	for i in range(len(...))	for item in iterator using <i>BookIterator</i>	for (int i = 0; i < list.size(); i++)	for (Song s : this) using <i>Iterator</i>
Filtering and Searching	Hardcoded loops & conditions	Flexible filter functions through iterator	Hardcoded loops & conditions	Not available (basic iteration code)
Separated Iteration Logic	No	Yes, <i>BookIterator</i> is separated from the collection	No	Yes, <i>SongIterator</i> is separated
Reusability	Low (repetitive code)	High (modular, customizable)	Low	Medium (modular iteration, no filtering)
Separation of Concerns (SoC)	Unclear	Clear separation between iteration and	Unclear	Clearer with <i>Iterable</i>

		collection logic		
Extensibility	Limited (duplicate code)	High (easily extendable via new filters/iterators)	Limited	Medium

4.3 Reusability Evaluation

The implementation of the Iterator pattern in Python improves reusability, as filtering functions can be passed as parameters without modifying the BookCollection class. The iterator can be reused in various contexts with high flexibility. Meanwhile, Java also demonstrates improvement after using the Iterator pattern, although it remains less flexible than Python. For better results, Java can be combined with Predicate<Song> or Stream APIs to support more dynamic filtering capabilities

Table 2 : Evaluation of reusability in Python and Java

Fitur	Python (Before)	Python (After)	Java (Before)	Java (After)
Code Reusability	Low	High	Low	Medium
Ease of Modification	Low	High	Low	Medium
Extensibility Support	Limited	Flexible	Limited	Modular
Separation of Concerns	Weak	Strong	Weak	Good
Unit Test Support	Limited	Easy	Limited	Easy

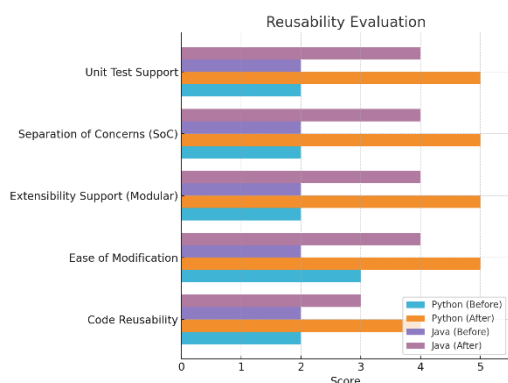


Figure 7 : Bar chart of reusability evaluation

The figure above illustrates the reusability evaluation of the code before and after applying the Iterator pattern in both Python and Java. It shows that after adopting the Iterator, both

languages experienced improvements across all evaluation aspects, such as code reuse, ease of modification, extensibility, separation of concerns, and support for unit testing. Python showed the most significant improvement, particularly in terms of modularity and separation of concerns. Java also improved, though its gains were more limited compared to Python. This demonstrates that the Iterator pattern helps enhance the structure and flexibility of code when managing data collections.

5 CONCLUSION

The application of the Iterator design pattern has proven to enhance the quality of reusability in object-oriented software development. By separating iteration logic from data structure, developers can produce code that is more modular, flexible, and easier to maintain. This implementation also aligns with core design principles such as separation of concerns and the open/closed principle.

Through case studies conducted in Python and Java, it was observed that Python offers greater flexibility due to its ability to accept filtering functions directly within the iterator. Meanwhile, Java also demonstrated improvements after implementing the Iterator pattern, although its filtering flexibility remains limited without additional constructs such as Predicate or Stream.

Overall, the use of the Iterator pattern not only simplifies the traversal of data collections but also strengthens the architecture of software systems to be more reusable and maintainable. Therefore, this pattern is highly recommended for projects that involve repeated and flexible data processing.

Furthermore, the Iterator pattern offers advantages in testing and continuous development, as its separation of iteration logic allows for more focused unit testing without dependency on the internal structure of the data. With clearer task isolation, development teams can accelerate refactoring processes and add new features without disrupting existing code. This makes the Iterator pattern a relevant and adaptive design solution for supporting modern software development practices that prioritize quality, efficiency, and scalability.

6 SUGGESTIONS

Based on the findings of this study, future work should extend beyond the Iterator pattern and investigate additional design patterns—such as Strategy, Observer, or Visitor—to provide a broader comparison of their impact on software

reusability. Moreover, applying the Iterator pattern within microservice architectures or other complex systems would help assess its effects on scalability and maintainability in real-world scenarios.

Subsequent research is also encouraged to incorporate quantitative metrics—such as coupling, cohesion, and the maintainability index—to support objective evaluation. Leveraging functional-programming features like filter() or Java Streams can further enhance the flexibility of iterators. By combining sound design practices with automated testing approaches, developers can build systems that are more modular, efficient, and sustainable over the long term.

REFERENCES

- [1] E. A. Alomar, T. Wang, V. Raut, M. W. Mkaouer, C. Newman, and A. Ouni, “Refactoring for reuse: an empirical study,” *Innov Syst Softw Eng*, vol. 18, no. 1, pp. 105–135, 2022.
- [2] G. Vale, F. F. Correia, E. M. Guerra, T. O. Rosa, J. Fritzsche, and J. Bogner, “Designing Microservice Systems Using Patterns: An Empirical Study on Quality Trade-Offs,” *IEEE*, pp. 69–79, 2022.
- [3] R. Pincirolì, A. Aleti, and C. Trubiani, “Performance Modeling and Analysis of Design Patterns for Microservice Systems,” *ICSA*, 2023.
- [4] F. Wedyan and S. Abufakher, “Impact of Design Patterns on Software Quality: A Systematic Literature Review,” *IET Software*, vol. 14, no. 1, pp. 1–17, 2020.
- [5] G. Lacerda, F. Petrillo, M. Pimenta, and Y. G. Gueheneuc, “Code Smells and Refactoring: A Tertiary Systematic Review of Challenges and Observations,” 2020.
- [6] A. Peruma, S. Simmons, E. A. AlOmar, C. D. Newman, M. W. Mkaouer, and A. Ouni, “How Do I Refactor This? An Empirical Study on Refactoring Trends and Topics in Stack Overflow,” 2021.
- [7] S. H. S. Almadi, D. Hooshyar, and R. B. Ahmad, “Bad Smells of Gang of Four Design Patterns: A Decade Systematic Literature Review,” *Sustainability*, vol. 13, no. 18, 2021.
- [8] S. Ahmed, “Enhancing Software Development Efficiency: The Role of Design Patterns in Code Reusability and Flexibility,” *AI Publications / IJEBDM*, vol. 19, no. 1, 2025.
- [9] M. El Amine Hamri, “An Object-Oriented Framework for Designing Reusable and Maintainable DEVS Models using Design Patterns,” 2020.
- [10] W. Meijer, C. Trubiani, and A. Aleti, “Experimental evaluation of architectural software performance design patterns in microservices,” *Journal of Systems and Software*, vol. 218, 2024.
- [11] T. P. Y. Titan, Budiman, and J. H. F. E. Putra, “Perancangan Prototype User Interface Dan Pengujian User Experience Aplikasi Rental Mobil Berbasis Menggunakan Metode Design Thinking (Studi Kasus : Pt Trans Berjaya Khatulistiwa),” *Nuansa Informatika*, vol. 17, no. 2, 2023.
- [12] Darsiti and D. S. F. Lestari, “Pengembangan Aplikasi Keuangan Berbasis Web (Studi Kasus : Mts AT-Taufiq Kota Bandung),” *Nuansa Informatika*, vol. 17, no. 2, 2023.
- [13] M. K. Maulidani, Darsanto, and M. E. Iswanto, “Rancang Bangun Aplikasi Informasi Wisata Kabupaten Indramayu Berbasis Mobile Android,” *NUANSA INFORMATIKA*, vol. 18, no. 2, 2024.
- [14] A. Kevin and H. C. Kurniawan, “Analisis Pengaruh Design Pattern Terhadap Pemeliharaan Perangkat Lunak Learning Management System,” *Jurnal Telematika*, vol. 18, no. 1, 2023.
- [15] I. G. B. V. Putraadinatha, D. D. J. Suwawi, and S. Y. Puspitasari, “Pengaruh Design Pattern Terhadap Maintainability Aplikasi Mobile,” *eProceedings of Engineering*, vol. .8, no. 5, 2021.

Web-Based Asset Maintenance Monitoring System for CV Tartila Group

Riyan Setiawan¹, Yuga Nata Praja², Wasis Haryono³

^{1,2,3}Program Studi Teknik Informatika, Fakultas Teknik, Universitas Pamulang, Indonesia
E-mail: ¹riyanstwn11@gmail.com, ²yuganatapraja29@gmail.com, ³wasish@unpam.ac.id

This study presents the development of a Web-Based Asset Maintenance Monitoring System for CV Tartila Group, a company engaged in photography, videography, and digital content production. Manual documentation of maintenance activities often leads to data loss, inconsistent records, and ineffective scheduling of preventive maintenance. The proposed system utilizes a structured approach based on the Waterfall Software Development Life Cycle (SDLC), covering requirements analysis, system design, implementation, and testing. The application integrates asset registration, maintenance history tracking, and automatic scheduling of preventive maintenance using a centralized MySQL database. The results of system testing show that the solution successfully replaces manual processes, increases the accuracy of service records, provides real-time monitoring, and improves decision-making regarding asset lifecycle management. This system contributes significantly to the operational efficiency of CV Tartila Group by reducing downtime and ensuring consistent asset performance.

Keywords— Asset management, preventive maintenance, web-based system, SDLC, monitoring.

Submission: December 13, 2025 Accepted: December 22, 2025 Published: January 20, 2025

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.541>

1. INTRODUCTION

Asset management plays a critical role in ensuring the reliability of equipment used in digital creative industries such as photography and videography. CV Tartila Group relies heavily on primary assets—including cameras, lenses, gimbals, and editing workstations—to deliver high-quality documentation services for sports events and wedding productions. However, the company still performs asset maintenance documentation manually using spreadsheets and verbal reports, which poses several risks, such as data inconsistency, loss of historical maintenance records, and difficulties in scheduling preventive maintenance.

Previous studies highlight the importance of systematic asset maintenance. Krisdiawan [1] emphasizes the relevance of structured system development in improving operational effectiveness. Research by Widodo et al. [2] demonstrates how web-based systems increase data accuracy and efficiency in asset-related activities. Furthermore, Permatasari [3] shows that centralized management information systems significantly reduce maintenance downtime and support decision-making.

Based on these challenges, this study develops a Web-Based Asset Maintenance Monitoring System to digitize, centralize, and streamline maintenance processes. The system provides

structured preventive maintenance schedules, real-time monitoring, and complete service history documentation.

2. RESEARCH METHODS

The study adopts the Waterfall Model, which is suitable because system requirements have been clearly identified during observation and interviews at CV Tartila Group. The phases consist of:

2.1. Requirement Analysis

Activities include observing the asset workflows and interviewing the operational manager. The findings indicate:

Maintenance records are inconsistent and scattered.

No preventive maintenance schedule exists.

Service decisions rely on staff memory rather than structured data.

2.2. System Design

This phase yields:

Use Case Diagram for Admin, Crew, and Owner.
Activity Diagrams for asset maintenance processes.

ERD consisting of Users, Assets, Maintenance Records, and Categories.

UI Design for login page, dashboard, asset input, and maintenance scheduling.

2.3. System Implementation

The application implements a 3-tier architecture:

- 1.Front-end: HTML5, CSS, JavaScript
- 2.Back-end: Node.js + Express
- 3.Database: MySQL

Key functionality includes CRUD for assets, automated preventive maintenance calculation, and role-based access control.

2.4. System Testing

Black-box testing validates input, output, functionality, and access permissions. Testing focuses on:

- 1.Login validation
- 2.Asset registration
- 3.Maintenance record creation
- 4.Automated scheduling functionality
- 5.Report generation

3. RESEARCH RESULTS

This section presents the results of the analysis, system modeling, and design processes conducted during the development of the Web-Based Asset Maintenance Monitoring System for CV Tartila Group. The results include detailed visual representations in the form of UML diagrams, workflow models, and database structures. These outputs collectively serve as the blueprint for system implementation, ensuring that the developed application aligns with the operational needs of the company and follows established software engineering principles.

The diagrams shown in this section not only provide a conceptual overview of the system but also demonstrate how each system component interacts to support asset monitoring, maintenance recording, and user activity management. These models also serve as validation artifacts confirming that the system meets the functional and non-functional requirements gathered during the analysis phase.

3.1 USE CASE DIAGRAM

The Use Case Diagram gives a general idea of what the system can do and shows how the people using the system interact with its different functions. In this system, there are three main types of users: Admin, Crew, and Owner. Each actor plays different roles and has access to specific features

depending on their responsibilities within the organization.

The Admin actor has the highest privilege level, enabling system-wide operations such as asset registration, user management, maintenance verification, and report generation. Crew members interact with the system primarily to input maintenance records, update asset condition, and view assigned equipment. The Owner actor is responsible for overseeing the entire asset management process and has access to real-time monitoring features and overall maintenance reports.

By visualizing these interactions, the Use Case Diagram ensures that system functionalities are properly mapped and accessible only to authorized users.

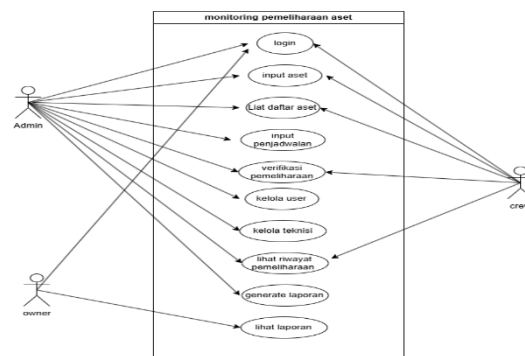


Figure 1. Use Case Diagram of the Asset Maintenance Monitoring System

3.2 ACTIVITY DIAGRAM OF THE CURRENT SYSTEM

The activity diagram of the current system illustrates the existing workflow used by CV Tartila Group before the implementation of the proposed application. Currently, asset maintenance is predominantly carried out using manual methods. Technicians or crew members perform equipment inspections after fieldwork and report any issues verbally or through informal notes. These reports are then handled by the manager, who decides whether maintenance or servicing is required.

The absence of a structured process leads to several inefficiencies. Records may be incomplete, lost, or forgotten, causing delays in maintenance and increasing the risk of equipment failure during critical tasks. Additionally, there is no centralized data

repository where maintenance history or asset condition trends can be analyzed.

The activity diagram highlights these weaknesses by mapping out the manual, fragmented, and non-data-driven workflow.



Figure 2. Activity Diagram of the Current System

3.3 ACTIVITY DIAGRAM OF THE PROPOSED SYSTEM

The presented system's activity diagram illustrates the enhancement in workflow organization, automated, and consistent through the use of the web-based application. Users begin by logging into the system, where access levels are verified. The Admin can access modules to add new assets, update maintenance data, assign responsible personnel, and validate service records. Crew members can record any detected issues directly into the system, ensuring real-time documentation.

One of the key improvements introduced by the proposed system is the automatic calculation of the next maintenance schedule based on the predefined maintenance cycle for each asset. This significantly reduces human error and ensures that preventive maintenance activities are performed on time.

The diagram also demonstrates how the system stores all data in a centralized MySQL database, enabling efficient retrieval, accurate history tracking, and real-time updates on the dashboard.

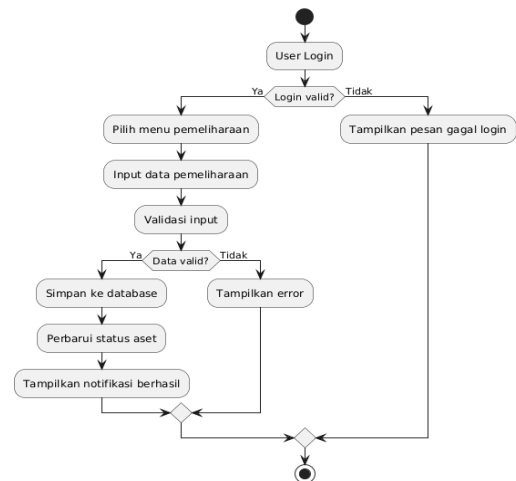


Figure 3. Activity Diagram of the Proposed System

3.4 Sequence Diagram

Sequence Diagram shows, in order, how the different parts of the system and the people using it work together when doing certain things. This diagram explains message flows, request-response patterns, and the order of system operations that occur as the user performs actions such as logging in, submitting maintenance data, or retrieving asset status.

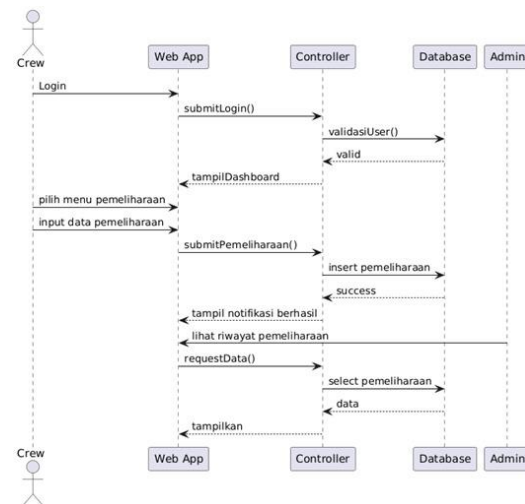


Figure 4. Sequence Diagram of the Proposed System

3.5 ENTITY RELATIONSHIP DIAGRAM (ERD)

The ERD represents the logical database structure of the system and illustrates how data entities are related to one another. The database was designed using normalization principles to

eliminate data redundancy and maintain consistency across records.

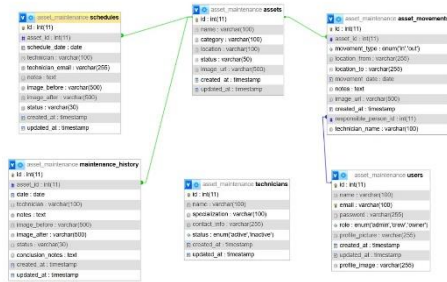


Figure 4. Entity Relationship Diagram of the System

3.6 SYSTEM FLOWCHART

The flowchart illustrates the detailed logical flow of the system operations. It includes decision-making processes, data input verification, routing of user actions, and storing of information into the database.

Flowcharts are important to validate the correctness of system logic and identify any potential bottlenecks or errors before the implementation stage. The flow also ensures that every action initiated by a user results in the expected system output.

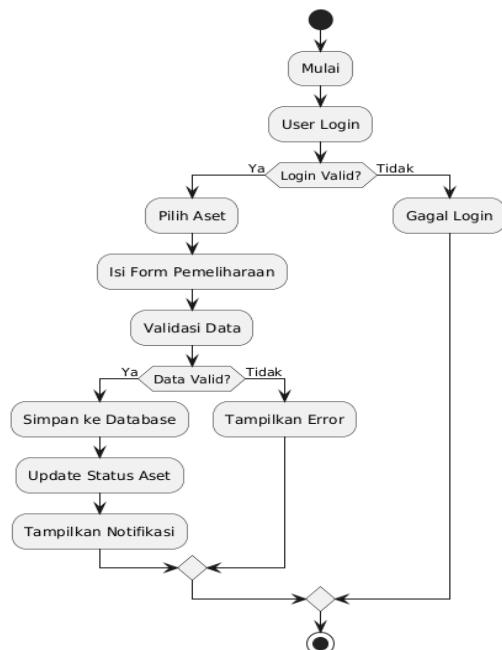


Figure 5. System Flowchart of the Asset Maintenance Process

3.7 User Interface Design (UI Design)

The way the User Interface (UI) is designed shows what the main pages of the system look like and how they are set up at first. UI design is done to make sure the application is simple to use, visually consistent, and aligned with the functional requirements obtained during the analysis phase. The interface mockups guide developers during implementation and help users understand the intended system flow.

3.7.1 Login Page

This page provides user authentication for Admins, Crew, and Owners. It contains fields for username and password, supported by input validation to ensure secure access.



Figure 6. Login Page Interface

3.7.2 Dashboard Page

The dashboard displays quick summaries such as total assets, maintenance schedules, and recent activity logs, enabling users to monitor system updates efficiently.



Figure 7. Dashboard Interface

3.7.3 Asset Management Page

This interface is used by Admins to add, update, and delete asset information. The layout supports accurate data entry to maintain asset records.

ID	GAMBAR	NAMA	KATEGORI	LOKASI	STATUS	Aksi
16	No Image	MONITOR PC	BARANG ELEKTRONIK	STUDIO UTAMA	Standby	[Edit] [Hapus]
17	No Image	KEYBOARD SAMSUNG	BARANG ELEKTRONIK	STUDIO UTAMA	Standby	[Edit] [Hapus]
18	No Image	MOUSE ZIREX	BARANG ELEKTRONIK	STUDIO UTAMA	Standby	[Edit] [Hapus]
19	No Image	KAMERA SONY A7	BARANG ELEKTRONIK	STUDIO UTAMA	Operational	[Edit] [Hapus]

Figure 8. Asset Management Interface

3.7.4 Maintenance Input Page

Used by Crew members to submit maintenance activities. It includes fields for maintenance details, dates, descriptions, and photo uploads to ensure complete documentation.

ID	ASET	TANGGAL	TEKNIISI	TEKNIISI
----	------	---------	----------	----------

Figure 9. Maintenance Input Interface

3.7.5 Report Page

This page displays a list of maintenance activities with filtering and sorting options, helping users review and analyze maintenance history.

ID	COMPLETION DATE	TECHNICIAN	NOTES	SUMMARY NOTES	STATUS	ACTION
10	KEYBOARD SAMSUNG	Budi Technician	repair needed	repaired	Completed	[Icon]

Figure 10. Maintenance Report Interface

4. DISCUSSION

This part talks about what was found while creating the Web-Based Asset Maintenance Monitoring System for CV Tartila Group.. The discussion focuses on interpreting the results presented in the previous chapter, analyzing the effectiveness of the proposed system, and evaluating how the system addresses existing operational issues within the company. Additionally, this chapter highlights the improvements achieved through system implementation in terms of accuracy, efficiency, data integration, and decision-making support.

4.1 Analysis of System Improvements

The introduction of a web-based asset maintenance system provides several significant improvements compared to the traditional manual process. The previous workflow relied heavily on handwritten notes, verbal communication, and informal documentation, which resulted in inconsistent reporting, delayed responses, and insufficient record tracking. These issues created operational inefficiencies and increased the risk of unexpected equipment failures.

The proposed system addresses these challenges by providing:

Data Kept in One Place

A MySQL database is used to keep all records of maintenance, information on assets, and things users do within the system. This makes sure that information is easy to get to, well-organized, and safe from being lost or put in the wrong place.

Real-Time Monitoring

The system dashboard presents up-to-date information about asset conditions, maintenance status, and pending tasks. This allows managers and owners to monitor operational readiness at any time.

Higher Accuracy of Reports
By automating data input, validation, and storage, the system reduces human errors and improves the accuracy of maintenance logs.

Increased Accountability
Each activity is associated with a specific user, enabling tracking of responsibilities and ensuring that tasks are completed by authorized personnel. These improvements collectively enhance the reliability and transparency of maintenance operations within CV Tartila Group.

4.2 Evaluation of User Roles and Access Control

The multi-role access model implemented in the system ensures that each user interacts only with the features relevant to their responsibilities.

Admin, Crew, and Owner roles are clearly separated through the access control logic embedded in the system.

Admins manage assets, validate maintenance records, generate reports, and oversee user accounts.

Crew Members perform operational tasks such as logging maintenance activities, updating asset conditions, and submitting inspection results.

Owners have access to high-level monitoring features, summarized reports, and system-wide performance indicators.

This role-based structure not only increases system security but also prevents unauthorized access. It ensures that the workflow becomes more controlled and that data integrity is preserved.

4.3 Discussion of Database Design Effectiveness

The database schema presented in Chapter III demonstrates a normalized and well-structured relational model. By implementing foreign key constraints and using consistent naming conventions, the system minimizes redundancy and potential anomalies in data storage.

Through the ERD and table relations:

Data dependencies are clearly defined.

Each maintenance record is traceable to both the asset and the responsible user.

Asset categories help in filtering and generating more targeted reports.

Historical data can be analyzed to predict future maintenance needs.

The relational database thus plays a vital role in ensuring that the maintenance monitoring process is systematic, accurate, and scalable.

4.4 Impact on Maintenance Performance

With the new system in place, maintenance performance is expected to improve in multiple aspects:

Timeliness of Maintenance Activities

The system allows users to track upcoming maintenance schedules, reducing delays.

Reduction of Unplanned Downtime

With better historical tracking and automatic scheduling, assets are less likely to fail unexpectedly.

Enhanced Reporting Capabilities

Reports generated through the system are more detailed, consistent, and suitable for managerial decision-making.

Operational

Every maintenance action is logged, allowing supervisors and owners to evaluate crew performance and asset reliability.

These impacts demonstrate that the proposed system not only digitizes the maintenance process but also enhances operational quality.

Transparency

4.5 System Limitations

Despite its advantages, the system has some limitations:

Internet

The system requires network access, which may be challenging in remote operational areas.

User

Training

Requirements

New users may require training to fully utilize the system.

Limited

Mobile

Optimization

If accessed from mobile devices, some interfaces may require further optimization.

These limitations can be addressed in future development.

4.6 Overall System Evaluation

Overall, the system successfully meets the functional requirements and significantly enhances the asset maintenance workflow at CV Tartila Group. The combination of structured data management, streamlined processes, and real-time monitoring makes the system an effective solution for supporting daily maintenance activities and long-term asset management strategies.

5. CONCLUSION

Creating a Web-Based System to Keep Track of Asset Maintenance for CV Tartila Group has successfully addressed the shortcomings of the company's previous manual maintenance process. Through a structured analysis, comprehensive system design, and the implementation of UML-based modeling, the system provides an integrated and reliable platform for managing asset information and maintenance activities.

The results demonstrate that the system offers several key advantages. First, it centralizes all maintenance data into a secure and organized database, ensuring that historical maintenance records can be accessed easily and accurately. Second, the system increases operational transparency by associating every maintenance action with a registered user, thereby improving accountability across the organization. Third, the system enhances efficiency by automating task scheduling, standardizing maintenance

reporting, and enabling real-time monitoring through a dashboard interface.

Additionally, the role-based access control implemented within the system ensures that each actor—Admin, Crew, and Owner—interacts with the application according to their responsibilities, reinforcing system security and reducing the risk of unauthorized access. The relational database structure and workflow diagrams further validate the robustness and scalability of the system design.

Overall, the system provides a practical and effective solution for improving the accuracy, reliability, and timeliness of asset maintenance operations at CV Tartila Group. While the current implementation meets the core functional requirements, future enhancements such as offline mode, mobile optimization, and predictive maintenance analytics could further strengthen the system's value and performance.

6. RECOMMENDATIONS

Based on the development process and the results obtained from the implementation of the Web-Based Asset Maintenance Monitoring System for CV Tartila Group, several recommendations can be made to improve future system performance and broaden its applicability:

Enhance Mobile Responsiveness
Although the system can be accessed through a web browser, certain interfaces may not be fully optimized for smaller screens. Enhancing the system's mobile responsiveness or developing a dedicated mobile application would significantly improve usability for field workers who frequently operate outside the office environment.

Implement Offline Mode Support
Many maintenance activities occur in locations where internet connectivity may be unstable or unavailable. Integrating offline data recording features that synchronize automatically when a connection is restored would ensure continuous functionality and prevent data loss.

Integrate Predictive Maintenance Analytics
Incorporating predictive maintenance techniques—such as trend analysis, machine learning, and failure prediction algorithms—could enable the system to forecast asset conditions based on historical data. This would help the company plan preventive actions more

effectively and reduce unexpected equipment failures.

Expand Notification and Alert Features
Implementing automated email or push notifications for upcoming maintenance schedules, overdue tasks, and critical asset conditions would increase the system's practicality and proactive monitoring capabilities.

Conduct Periodic User Training
To maximize system effectiveness, regular training sessions should be held for Admin, Crew, and Owner users. This ensures that all personnel can utilize system features properly and remain updated on new features introduced in future versions.

Strengthen Data Security and Backup Mechanisms

Since the system handles sensitive operational data, enhancing security measures such as data encryption, multi-factor authentication, and automated database backups is essential to ensure long-term data integrity and protection against cyber threats.

Integrate Asset Lifecycle Management Features
Expanding the system to include full lifecycle management—such as procurement tracking, depreciation calculation, and disposal management—would allow CV Tartila Group to manage assets comprehensively from acquisition to retirement.

By implementing these recommendations, the system can continue to evolve and offer even greater value to CV Tartila Group. Future development efforts should focus on improving accessibility, enhancing analytical capabilities, and ensuring that the system remains adaptable to the operational needs of the organization.

REFERENSI

- [1] Y. Cahyaningrum and Y. Sambharakreshna, Optimalisasi pengelolaan aset berbasis web dalam peningkatan efisiensi dan keberlanjutan optimization of web-based asset management to increase efficiency and sustainability, *Journal of Information Technology and Computer Science (INTECOMS)*, vol. 7, no. 2, 2024.
- [2] M. A. Athallah and K. Kraugusteeliana, Upaya Manajemen Kinerja Strategis Perusahaan Melalui Sistem Informasi Pemeliharaan Aset Bergerak Berbasis Web, *Jurnal Nusantara Aplikasi Manajemen*

- Bisnis, vol. 8, no. 1, pp. 27–45, 2023.
- [3] D. A. Siregar et al., Perancangan Aplikasi Mobile Berbasis Android Untuk Pemeliharaan Aset Pada Kecamatan Sindang Jaya, JATI (Jurnal Mahasiswa Teknik Informatika), vol. 8, no. 2, pp. 1512–1520, 2024.
- [4] E. Rahmawati, R. J. Putra, and N. Wahyuningtyas, Aplikasi Manajemen Aset Berbasis Web pada SMA HangTuah 4 Surabaya, Journal of Technology and Informatics (JoTI), vol. 2, no. 2, pp. 65–71, 2021.
- [5] W. Hadikristanto and N. T. Kurniadi, Implementasi Pengembangan Aplikasi Sistem Manajemen Aset Berbasis Web Menggunakan Metode Waterfall Untuk Mengoptimalkan Penggunaan Aset Pada PT. Utama Karya (Persero), Jurnal Teknologi dan Sistem Informasi Bisnis, vol. 5, no. 4, pp. 401–408, 2023.
- [6] M. Ichsan et al., Rancang Bangun Sistem Informasi Manajemen Aset Berbasis Mobile Web di SMP-SMA Olahraga Negeri Sriwijaya Sumatera Selatan, Jurnal Ilmiah Matrik, vol. 26, no. 1, pp. 21–27, 2024.
- [7] R. Rahma, H. Nur, and A. W. Utami, Rancang Bangun Sistem Informasi Pengelolaan dan Pemeliharaan Aset Berbasis Website Pada PT. XYZ, *Journal of Emerging Information Systems and Business Intelligence*, vol. 5, no. 2, pp. 169–177, 2024.
- [8] W. Hidayat, F. A. Ba’a, O. Prasetyo, and W. Haryono, Perancangan Sistem Aplikasi Absensi Real Time untuk Meningkatkan Efisiensi Manajemen Kehadiran PT. Asia Sinergi Solusindo, *Switch: Jurnal Sains dan Teknologi Informasi*, 2025.
- [9] Rafli Fadillah Agustio, Ahnaf Irfan Baharianto, Riyan Pratama Mulia, and Wasis Haryono, Perancangan Sistem Inventory Dan Transaksi Pembelian Barang Berbasis Web Dengan Metode WATERFALL, *Restikom*, vol. 6, no. 3, pp. 554 – 564, Dec. 2024.
- [10] D. Irawan, E. Y. Darmawan, E. E. Zebua, and W. Haryono, Perancangan Sistem Informasi Proyek Berbasis Web Untuk Meningkatkan Kinerja Antar Divisi *Jurnal Komputer Antartika*, 2024.
- [11] F. Ferdiansyah, Perancangan Aplikasi Monitoring Siswa Praktik Kerja Lapangan Berbasis Web, *Prosiding Seminar Nasional Teknologi Informasi dan Komunikasi (SENATIK)*, vol. 6, no. 1, 2024.
- [12] Fransiscus Xaverius Ariwibisono and Widodo Pudji Muljanto, Implementasi Monitoring Aset Produksi Energi PLTS Berbasis Protokol MODBUS RTU DAN MODBUS TCP, *Nuansa Informatika*, vol. 17, no. 2, pp. 109–118, Jul. 2023.
- [13] F. Zahira Badeni, A. Guntara, and I. Fadil, Implementasi Sistem Informasi Manajemen Aset Kelompok Ternak Ikan Hias Sumedang Menggunakan Metode Prototype, *Nuansa Informatika*, vol. 18, no. 2, pp. 196–208, Jul. 2024.
- [14] A. Safitri, Saskya, and S. Sunarto, “Rncangan bangun Inventori Aset Pada Institut Teknologi Indonesia Berbasis Web,” 2025.
- [15] G. K. Hanum, I. A. Santoso, and M. Nurhasandi, Perancangan Sistem Monitoring Pemeliharaan Kendaraan Berbasis Web Pada PT. SURYA MUSTIKA NUSANTARA, *J. Sensi*, vol. 7, no. 2, pp. 176–187, 2021.
- [16] M. M. Parenreng, M. Nas, and J. M. Parenreng, Aplikasi Monitoring Aset Dan Inventaris Laboratorium Berbasis Web Untuk Jurusan Teknik Elektro Politeknik Negeri Ujung Pandang (PNUP), in *Seminar Nasional Hasil Penelitian & Pengabdian Kepada Masyarakat (SNP2M)*, vol. 3, 2018.

User Interface and User Experience (UI/UX) Design for the Ummu Micco Salon Website Using the Design Thinking Approach

Fathan Aqsha Ghasani Aufa^{*1}, Febriyansyah Ramadhan², Debi Irawan³

^{1,2,3}Universitas Indonesia Membangun, Indonesia

E-mail: ^{*1}fathanaqsha@student.inaba.ac.id, ²febriyansyah.ramadhan@inaba.ac.id, ³debiirawan@inaba.ac.id

Abstract

Digital transformation is a strategic solution for Micro, Small, and Medium Enterprises (MSMEs) to maintain global competitiveness, however, its implementation in the beauty service sector is often hindered by operational technology gaps. The primary issue in the Ummu Micco Salon case study is service management inefficiency, which negatively impacts business growth. The underlying cause lies in an inaccessible reservation system and a lack of price transparency, directly resulting in customer confusion and the risk of potential revenue loss. This study aims to address this gap by designing a User Interface (UI) and User Experience (UX) solution built on a website using the Design Thinking approach. The main contribution of this research is implementative, providing a validated design solution that directly resolves the service management inefficiencies faced by the beauty salon. The solution was realized as a high-fidelity prototype using Figma software and validated by 20 respondents. Evaluation results using the System Usability Scale (SUS) measurement indicated a score of 79.625, achieving a Grade B and an "Acceptable" category. This score directly demonstrates that the system design possesses a high level of usability and is feasible for implementation to support operational efficiency and business sustainability for MSMEs.

Keywords— UI/UX, MSMEs, Beauty Salon, Design Thinking, SUS

Submission: December 25, 2025 Accepted: January 04, 2026 Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.545>

1. INTRODUCTION

In the midst of an era of rapid digital transformation, Micro, Small, and Medium Enterprises (MSMEs) play a strategic role in Indonesia's national economic structure. Digital transformation has now become an absolute necessity for MSMEs to survive and compete in the global economy. The adoption of digital technology offers various significant benefits, including increased efficiency, productivity, and market competitiveness in an increasingly competitive landscape [1]

The implementation of digital-based operational innovations has proven to have a positive impact on three main aspects: improving operational efficiency, expanding market reach, and enhancing service quality [2], [3]. The digital ecosystem is a crucial factor for the success of MSMEs in Indonesia [4]. Within this ecosystem, a website serves as one of the primary digital gateways and the most fundamental asset for a business. A well-designed website functions as an information center, a showcase for products

and services, and a promotional medium that operates ceaselessly, enabling companies to reach a wider audience [5], [6].

For businesses operating in the service sector, such as beauty salons, the implementation of web-based technology can directly support various core business processes, ranging from promotional activities to transactional functions such as the implementation of online booking systems [7], [8]. The success of a website as a business tool is not determined solely by its technological sophistication, but rather by the quality of the interaction between the system and its users. This is where the crucial roles of User Interface (UI) and User Experience (UX) become central. Both influence each other and are interconnected [9], [10], [11].

The issues of manual reservation inefficiency and the lack of price transparency have become major obstacles that diminish customer trust and the professional image of Ummu Micco Salon. Designing a website serves as a crucial solution to simplify the booking flow and provide information certainty. The Design Thinking approach is applied to ensure the

solution is relevant to user needs. Design Thinking emerges as a relevant approach to address this need, as this methodology places empathy for the user as the starting point of the entire creative process to generate solutions that meet user needs and desires [12], [13].

Based on the urgency of the operational problems faced, this study aims to design a website-based UI/UX for Ummu Micco Salon by applying the Design Thinking approach. This research focuses on an implementative contribution by producing a tangible digital artifact capable of streamlining the reservation process and ensuring price transparency. To objectively validate the effectiveness of the solution, this study also aims to evaluate the usability level of the final design using the System Usability Scale (SUS) method, thereby ensuring that the created solution is feasible for implementation and complies with user satisfaction standards.

2. RESEARCH METHOD

The primary methodology employed in this research is the Design Thinking approach. The selection of this approach possesses a strong strategic justification, particularly within the context of MSMEs such as Ummu Micco Salon. Since salon operations are heavily reliant on personal relationships between the owner and customers, the Design Thinking approach ensures that digital solutions remain relevant to user needs, thereby facilitating easier acceptance and adoption.

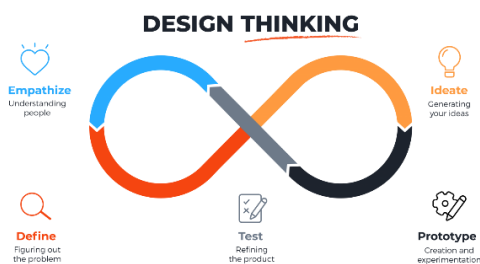


Figure 1. Design Thinking Stage

2.1. Empathize

This stage serves as the foundation of the entire process, aiming to build a deep and empathetic understanding of the users [14]. To obtain comprehensive user information, the researcher conducted in-depth interviews with informants who are also users, totaling 6 participants: the owner of Ummu Micco Salon and 5 customers. The inquiries posed to the users pertained to the obstacles they encountered in

salon services for customers, and the difficulties experienced by the owner in efficiently reaching consumers, alongside direct on-site observation. Once all questions were answered, the researcher classified the responses into user complaints.

2.2. Define

Insights collected from the Empathize stage, specifically user complaints, are analyzed and synthesized. This process focuses on mapping Pain Points to identify specific barriers experienced by users, as well as User Needs to understand the actual motivations or goals sought. The analysis results are then converted using the How Might We (HMW) method. HMW functions as a strategic instrument to transform rigid problem statements into idea-triggering questions, thereby opening a space for the exploration of effective and innovative solutions in the Ideate stage.

2.3. Ideate

After the core problem is clearly defined through HMW, the researcher conducts brainstorming sessions to generate creative ideas and solutions, from which Solution Ideas emerge. The outcome of this stage consists of key ideas or concepts deemed most effective in addressing operational issues, which are subsequently visualized in the prototype creation stage.

2.4. Prototype

The best ideas from the ideate stage are realized into testable prototype forms. The prototypes range from low-fidelity, which depicts rough ideas and designs serving as the website's foundation for evaluation and user feedback to high-fidelity, which produces detailed designs and attractive visuals as the final website design result, utilizing the design software Figma.

2.5. Test

The Test stage is the final phase in the Design Thinking approach to validate the designed prototype [15]. To objectively measure the level of usability and user satisfaction, this research utilizes the System Usability Scale (SUS) for evaluation. The study involved 20 respondents through the distribution of online questionnaires, comprising 5 existing salon customers and 15 potential customers (individuals who have never visited the salon). The results of this testing aim to refine the

solution, identify deficiencies, and conduct design iterations until an optimal solution is achieved [16].

3. RESULT AND DISCUSSION

In this section, the results from each stage of the UI/UX design process for the Ummu Micco Salon website will be elaborated upon using the Design Thinking approach. This discussion encompasses the presentation of problem findings from various stages, user needs, the realization of design solutions, and how the final evaluation results demonstrate the solution's effectiveness in addressing business operational constraints.

3.1. Empathize Stage

The Empathize stage serves as the foundation of the design process, aiming to establish a deep understanding of the users. In this phase, the researcher conducted in-depth interviews with the owner and several customers of Ummu Micco Salon, alongside direct observations at the research site. The objective was to delve deeper into the experiences, motivations, and difficulties they encountered. Consequently, the researcher gathered core problem data in the form of complaints and concerns from both the owner and customers through interviews, presenting the collected results as shown in Table 1.

No	Complaints
1	The owner desires an information system capable of attracting customers to the salon.
2	Customers desire a transparent price list for each service to prevent misunderstandings.
3	Customers require an appointment system to check schedule availability and obtain salon contact details for further inquiries.
4	Customers would value the ability to access comprehensive salon information at any time.

Table 1. User Complaints

3.2. Define Stage

In this stage, the researcher aims to identify and analyze the user complaints that were conveyed in the previous stage. This analysis process focuses on categorizing user complaints into Pain Points and User Needs,

which will subsequently serve as the foundation for formulation in the subsequent stage.

No	Pain Points	User Needs
1	The owner acknowledges difficulties in reaching customers due to the lack of effective informational media to attract customer interest.	A professional digital platform to showcase service information and a portfolio to enhance the salon's appeal.
2	Customers experience hesitation and concern regarding potential misunderstandings due to the absence of a transparent and easily accessible price list.	Customers require easy and transparent access to the service price list to establish trust.
3	Customers experience uncertainty when attempting to make an appointment, as they lack sufficient information (neither contact details nor an appointment system) to determine salon availability.	A system capable of providing contact information and facilitating direct communication, or a clear and structured system for scheduling appointments and receiving schedule confirmation.
4	Customers face limitations regarding salon information, which is not easily accessible (restricted to on-site banners).	Information that is easily accessible regardless of the customer's location.

Table 2. Pain Points and User Needs

Following the identification of Pain Points and User Needs, the researcher proceeds to formulate How Might We (HMW) statements derived from these points, which are presented in Table 3.

No	How Might We
1	How Might We enable the salon owner to professionally showcase service

No	How Might We
	information and portfolios to attract new customers?
2	How Might We present transparent and easily accessible price list information at any time to establish customer trust?
3	How Might We create a clear system for customers to determine salon availability, either through contact information or an appointment system, to eliminate uncertainty?
4	How Might We ensure that customers can easily access all essential salon information (services, prices, contacts) from anywhere?

Table 3. How Might We

3.3. Ideate Stage

Following the clear definition of the problems through the HMW formulations in Table 3, the researcher proceeded to the Ideate stage. This stage constitutes a brainstorming session aimed at generating a diverse range of ideas and creative solutions derived from the preceding stage.

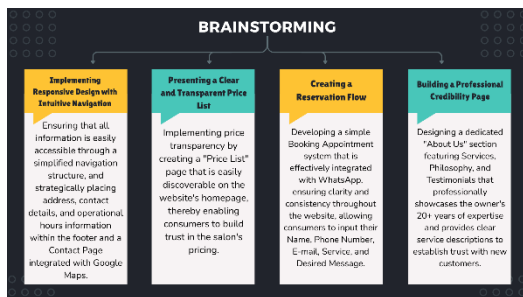


Figure 2. Brainstorming

Based on the brainstorming results above, the researcher determined the Solution Ideas to be implemented in the website design, aligned with user needs.

No	Solution Idea
1	Designing an attractive website interface that serves as an easily accessible and understandable source of salon information for customers.
2	Creating several sections on the website, comprising the Navbar (Navigation Bar), Hero section, Operational Hours, Services, Price List, Testimonials, Booking Appointment, and Footer.
3	Positioning clear and easily discoverable navigation within the website's main menu.

No	Solution Idea
4	Creating a clear and unambiguous price list section display, corresponding to the services available at the salon.
5	Developing an Appointment section containing various form fields: Name, Phone Number, E-mail, Selected Service, and Desired Message (for scheduling or other inquiries).
6	An appointment system integrated with the WhatsApp application, activated after completing all form fields and clicking the submit button.
7	Designing a footer containing comprehensive information regarding Contacts, Address, and Location.
8	Creating an additional "About Us" page containing the salon's identity and advantages.

Table 4. Solution Idea

3.4. Prototype Stage

Once the Solution Ideas were established, the research advanced to the Prototype stage. The objective of this stage is to transform concepts into rough designs (low-fidelity) and detailed designs (high-fidelity) so that they are testable. This process is conducted incrementally to validate the design concept, ranging from structure to visual appearance.

3.4.1. Low-fidelity

Low-fidelity, also known as a wireframe, is an initial visual representation serving as a blueprint for the website design. This stage constitutes a foundation focusing on structure, information hierarchy, and functionality, as shown in Figure 3.

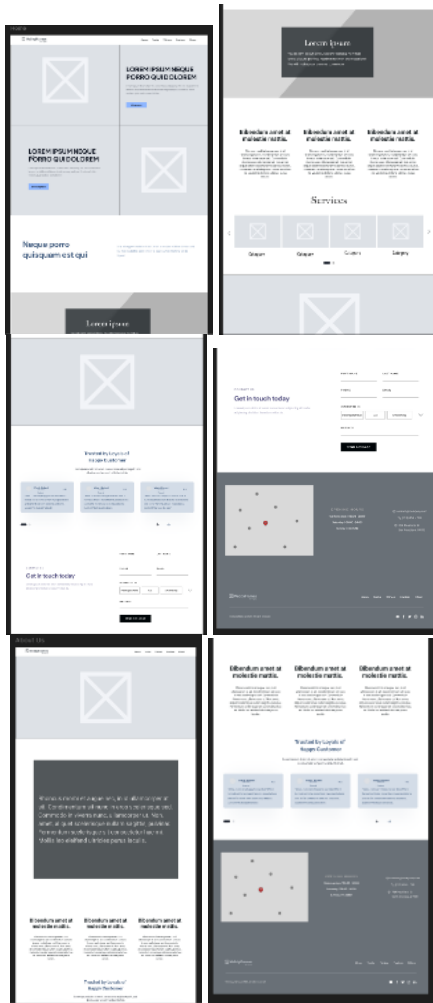


Figure 3. Wireframes

3.4.2. High-fidelity

High-fidelity, or mockup, is the final representation and the evolutionary stage of the low-fidelity framework. Building upon the designed structure and layout, this stage focuses on the application of visual elements to realize the Solution Ideas established in Table 4 for the Ummu Micco Salon website.

a. Navbar and Hero

Navbar and Hero sections constitute the primary interface or the initial view of the Ummu Micco Salon website. The function of the Navbar is to facilitate users in locating desired information, while the Hero serves to display a brief identity of the salon. In this section, customers can view details regarding the salon's identity by clicking the 'View More' button, or proceed to 'Book Now' to fill in appointment data in the Appointment section.

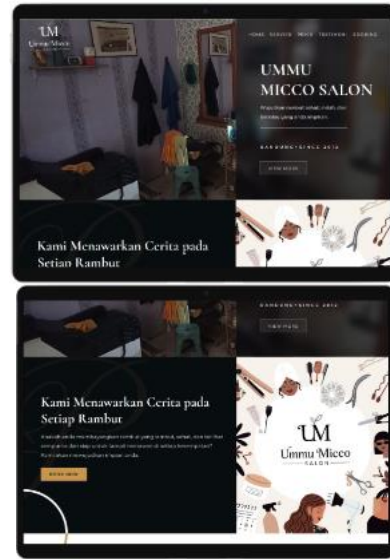


Figure 4. Navbar and Hero

b. Operational Hours

The subsequent section is the Operational Hours section, which aims to provide information regarding the operating hours of Ummu Micco Salon.



Figure 5. Operational Hours

c. Services

Services section serves as an informational component allowing customers to ascertain the range of services available at Ummu Micco Salon. This section utilizes an automatic Slider or Carousel, thereby eliminating the need for customers to manually scroll or swipe to explore the services.

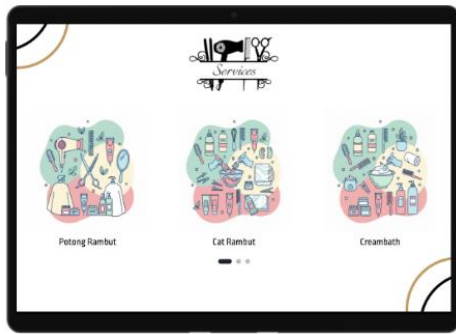


Figure 6. Services

d. Price List

Price List section represents the information most sought after by customers to ascertain the service pricing at Ummu Micco Salon; this information serves to ensure clear price transparency.

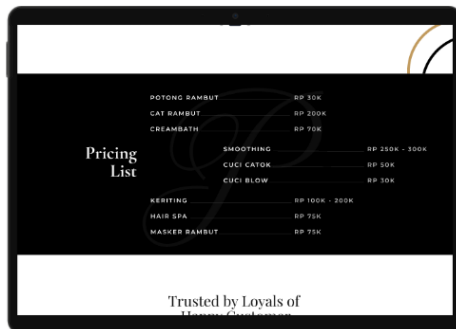


Figure 7. Price List

e. Testimonials

Testimonials constitute a section for customers to express reviews and feedback regarding the services they have experienced at Ummu Micco Salon. Here, customers can provide star ratings and comments.

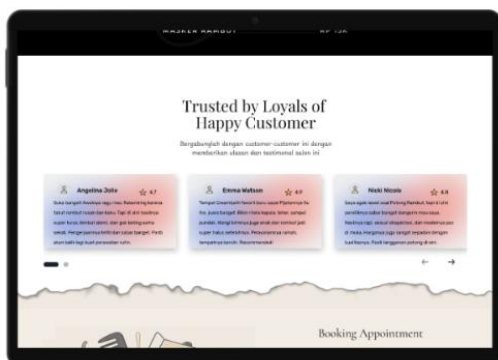


Figure 8. Testimonial

f. Booking Appointment

The Booking Appointment section is a feature designed to bridge the gap between structured data entry and the owner's communication habit. The specific novelty of this section lies in its direct data integration. When the customer clicks the 'Book' button, the system does not merely open the application, but automatically captures and formats the input data into a WhatsApp message template. The integrated data includes: (1) Name, (2) Phone Number, (3) Date, (4) Time, (5) Selected Service, and (6) Additional Message. This ensures that the salon owner receives a complete and standardized reservation request instantly, eliminating the need for manual re-typing or follow-up questions."

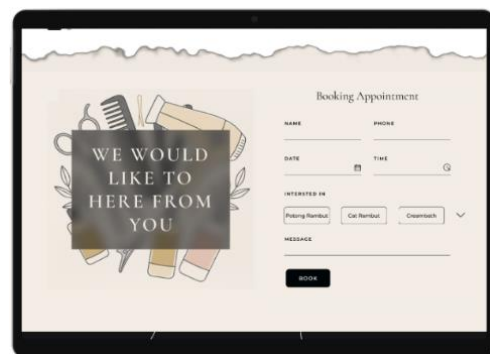


Figure 9. Booking Appointment

g. Footer

The Footer is the bottom-most section of the page, containing various essential information such as the salon's Location, Contacts, Address, and other details.

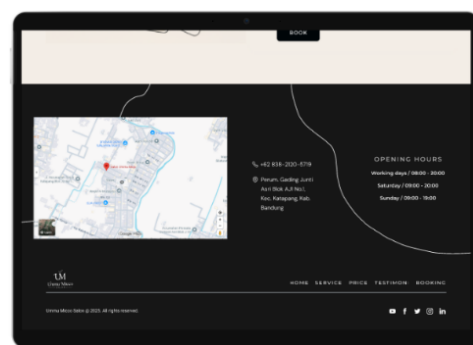


Figure 10. Footer

h. About Us

"About Us" is an additional page on the Ummu Micco Salon website containing background and supplementary information regarding the salon. Within this page, there is an "About" section presenting business identity information, and a "Philosophy" section aimed at

showcasing the advantages and values consistently upheld by Ummu Micco Salon.



Figure 11. About Us

3.5. Test Stage

The Test stage constitutes the final phase aimed at validating the proposed design solution. To objectively measure the level of usability and user satisfaction, this study utilizes the System Usability Scale (SUS) for evaluation, involving 20 respondents through the

distribution of online questionnaires. The respondents consisted of 5 existing salon customers and 15 potential customers. The solution addresses the operational pain points of loyal customers while ensuring the system remains intuitive and easy to learn for new users. Respondents were provided with a link to the developed prototype and requested to complete a series of scenarios, such as locating service information, viewing the price list, and performing an appointment booking simulation. Upon completion of these tasks, respondents completed a questionnaire encompassing ten evaluative statements. In the data presentation, each statement is denoted as Q1 through Q10 (Question), utilizing a Likert scale of 1 (Strongly Disagree) to 5 (Strongly Agree). Once the respondents had completed the questionnaire, score calculation was conducted in accordance with the established SUS methodology: the score for odd-numbered statements (1, 3, 5, 7, 9) is calculated as the user's response minus 1, whereas the score for even-numbered statements (2, 4, 6, 8, 10) is calculated as 5 minus the user's response. The total score from these ten statements is then summed and multiplied by a constant of 2.5.

No	Question										Aggregate	SUS Score (Aggregate *2,5)
	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10		
1	4	3	3	4	4	3	3	4	4	4	36	90
2	2	3	3	2	3	2	2	2	2	2	23	57,5
3	4	4	4	4	4	4	4	4	4	4	40	100
4	2	4	3	4	3	4	3	4	3	4	34	85
5	3	2	3	3	3	3	3	3	3	2	28	70
6	2	2	3	1	3	2	2	3	3	2	23	57,5
7	2	3	3	4	3	4	3	4	3	4	33	82,5
8	2	4	4	4	4	3	3	3	3	3	33	82,5
9	3	2	3	1	3	3	3	2	2	1	23	57,5
10	2	2	3	2	3	2	3	2	3	2	24	60
11	2	3	3	2	3	3	2	3	3	3	27	67,5
12	3	3	4	3	4	4	3	4	3	4	35	87,5
13	4	4	4	3	4	4	2	4	3	4	36	90
14	3	3	4	3	4	4	2	4	3	3	33	82,5
15	3	4	4	4	4	3	3	3	3	2	33	82,5
16	4	4	4	2	3	4	3	3	3	3	33	82,5

No	Question										Aggregate	SUS Score (Aggregate *2,5)
	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10		
17	3	4	4	3	4	3	1	3	4	3	32	80
18	3	4	4	3	4	4	2	4	2	4	34	85
19	3	4	4	4	4	4	4	4	4	4	39	97,5
20	3	4	4	4	4	4	4	4	4	3	38	95
Total SUS Score												79.625

Table 5. SUS Score

Based on the SUS measurements presented in Table 5, the evaluation yielded a total SUS score of 79.625 from 20 respondents. A score of 79.625 is situated above the global SUS average, securing a Grade B, and is classified within the "Acceptable" category based on global measurement standards [17] as illustrated in Figure 12.

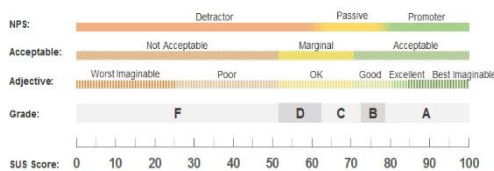


Figure 12. SUS Scoring Indicator

This high acceptability score reflects that the design successfully reduces cognitive load. The positive feedback, particularly from the 15 new potential customers, suggests that the information architecture, specifically the clear separation between 'Services', 'Price List', and 'Booking' is intuitive enough for first-time visitors to navigate without guidance. This indicates that the usability performance of the Ummu Micco Salon website design surpasses that of average systems, and users deem the system feasible for use without significant constraints.

4. CONCLUSION

Based on the preceding discussion, the UI/UX design for the Ummu Micco Salon website has successfully addressed the operational challenges and digital transformation needs of the MSME. The study confirms that the implementation of the high-fidelity prototype significantly resolves service management inefficiencies. Through the contextual application of the Design Thinking method, this study systematically mapped issues such as appointment scheduling difficulties and price transparency, subsequently synthesizing them

into user-centered design solutions. These solutions were realized in the form of a high-fidelity prototype offering strategic features, utilizing Figma software. The success level of this design is demonstrated through measurable scores from the System Usability Scale (SUS) testing, involving the participation of 20 respondents. The testing results indicated a SUS score of 79.625 on a scale of 0-100. Based on global interpretation standards, this score secures a position at Grade B and is classified within the "Acceptable" category. This figure indicates that the system's usability level resides above the industry average; consequently, this website design is validated as user-friendly and feasible for implementation to support the business sustainability of Ummu Micco Salon.

5. SUGGESTIONS

Based on the results and limitations of this study, the following are suggestions for future research:

1. This research is currently limited to the design of a high-fidelity prototype using Figma. Future research is suggested to proceed to the technical development stage by implementing the design into a functional website using programming languages and databases, allowing it to be utilized in a real-world environment.
2. Currently, the appointment feature is designed to integrate with the WhatsApp application, which still necessitates manual confirmation. Future studies could develop a backend system integrated with a database for real-time schedule management, enabling customers to check time slot availability instantly without waiting for a chat response.
3. The usability evaluation in this study involved 20 respondents using the SUS method. For future research, it is recommended to increase the number of

respondents across broader demographics or combine SUS with other methods such as the User Experience Questionnaire (UEQ) or Net Promoter Score (NPS) to obtain a deeper and more comprehensive analysis of user satisfaction.

REFERENCES

- [1] R. Hidayat, "Digitalisasi UMKM dan peningkatan daya saing di pasar global," *Economics Note*, vol. 1, no. 2, pp. 37–44, 2025, doi: 10.70716/econote.v1i2.82.
- [2] C. S. Otiva, P. E. Haes, T. I. Fajri, H. Eldo, and M. L. Hakim, "Implementasi teknologi informasi pada UMKM: Tantangan dan peluang," *Jurnal Minfo Polgan*, vol. 13, no. 1, pp. 815–821, 2024, doi: 10.33395/jmp.v13i1.13823.
- [3] R. Ariza and N. Aslami, "Analisis strategi pemasaran usaha mikro kecil dan menengah (UMKM) pada era digital di Kota Medan," *VISA: Journal of Vision and Ideas*, vol. 1, no. 2, pp. 188–194, 2021, doi: 10.47467/visa.v1i2.834.
- [4] E. D. Astuti and R. Rosita, "Pentingnya transformasi digital UMKM dalam pengembangan ekonomi Indonesia," *Sammajiva: Jurnal Penelitian Bisnis Dan Manajemen*, vol. 2, no. 4, pp. 119–134, 2024, doi: 10.47861/sammajiva.v2i4.1499.
- [5] S. Soedewi, A. Mustikawan, and W. Swasty, "The design thinking method application on the KiriHuci MSME website design," *Visualita: Jurnal Online Desain Komunikasi Visual*, vol. 10, no. 2, p. 17, 2022, doi: 10.34010/visualita.v10i02.5378.
- [6] A. P. A. Andika, "Sistem Informasi Web Promosi Joker Vape dengan Agile Scrum dan Usability Testing," *Nuansa Informatika*, vol. 18, no. 1, pp. 6–13, Jan. 2024, doi: <https://doi.org/10.25134/ilkom.v18i1.36>
- [7] Viky Fithrotul Qolbi, Lilik Sugiarto, and Hadis Turmudi, "Perancangan Sistem Informasi Pemesanan Salon Kecantikan Pada Aura Salon Berbasis Website," *JPSI*, vol. 1, no. 3, pp. 258–268, Aug. 2023, doi: <https://doi.org/10.54066/jpsi.v1i3.792>
- [8] Kasih and I. Ismail, "Perancangan UI/UX aplikasi booking online pada Elaine Studio dengan metode design thinking," *INFORMATIKA*, vol. 12, no. 3, pp. 417–426, 2024, doi: 10.36987/informatika.v12i3.6014.
- [9] N. Safitri and I. Maliki, "Perancangan UI/UX aplikasi HikingMate dengan metode design sprint untuk mendukung konektivitas pelanggan dan UMKM penyedia jasa hiking," *Jurnal Informatika Dan Komputasi: Media Bahasan, Analisa Dan Aplikasi*, vol. 19, no. 1, pp. 70–78, 2025. [Online]. Available: <https://journals.inaba.ac.id/index.php/jiki/article/view/492>
- [10] Damayanti, "Perancangan UI & UX aplikasi persediaan bahan baku untuk produksi pada PT. Indomas Prima Sejati," *Informatics And Computer Engineering Journal*, vol. 4, no. 2, pp. 93–103, 2024, doi: 10.31294/icej.v4i2.4937.
- [11] M. Rafiqh and Ismail, "Perancangan UI & UX aplikasi pariwisata kota berbasis Android di Kota DKI Jakarta," *Journal Education and Technology*, vol. 4, no. 2, pp. 182–197, 2023, doi: 10.31932/jutech.v4i2.3022.
- [12] N. W. Lestari and D. Y. Krisna, "Perancangan UI/UX menggunakan metode design thinking berbasis aplikasi penjualan pada usaha Bakso Barokah," *JUTECH: Journal Education and Technology*, vol. 5, no. 2, pp. 365–378, 2024, doi: 10.31932/jutech.v5i2.4082.
- [13] T. P. Y. Titan, Budiman, and J. H. F. Efendi Putra, "Perancangan Prototype User Interface Dan Pengujian User Experience Aplikasi Rental Mobil Berbasis Menggunakan Metode Design Thinking (Studi Kasus : Pt Trans Berjaya Khatulistiwa)," *Nuansa Informatika*, vol. 17, no. 2, pp. 48–65, Jul. 2023, doi: <https://doi.org/10.25134/ilkom.v17i2.9>
- [14] J. F. Pacheco, "What is design thinking? Stanford d.school model explained," *Juan Fernando Pacheco*, 2025. [Online]. Available: <https://juanfernandopacheco.com/2025/01/what-is-design-thinking-stanford-d-school-model/>
- [15] R. F. Dam, "The 5 stages in the design thinking process," *Interaction Design Foundation - IxDF*, 2025. [Online]. Available: <https://www.interaction-design.org/literature/article/5-stages-in-the-design-thinking-process>
- [16] M. L. Haryanti, F. Fraderic, S. Hartono, and J. Salim, "Cinema e-Ticket Application Design and Usability Evaluation Using SUS," *Nuansa Informatika*, vol. 19, no. 1, pp. 14–22, Jan. 2025, doi: <https://doi.org/10.25134/ilkom.v19i1.296>
- [17] J. Sauro, "5 ways to interpret a SUS score," *MeasuringU*, 2018. [Online]. Available: <https://measuringu.com/interpret-sus-score/>

Operationalizing Model Context Protocol for Multi-Platform Academic Search: A Comparative Study of LLM Grounding

Rizqullah Aryaputra Piliang¹, Anindito², Eryan Ahmad Firdaus³

^{1,2,3}Universitas Pertahanan, Indonesia

E-mail: *¹rizqullah.piliang@tm.idu.ac.id, ²anindito@idu.ac.id, ³eryan.firdaus@idu.ac.id

Abstract

This study introduces the Paper Search Model Context Protocol (MCP) server, a unified architecture enabling Large Language Models (LLMs) to autonomously search and retrieve academic literature across heterogeneous platforms (arXiv, PubMed, bioRxiv, medRxiv, Semantic Scholar, and CrossRef). By normalizing diverse metadata into a standardized schema, the system facilitates platform-agnostic tool use. We evaluate this infrastructure by tasking GPT-4.1 and GPT-5.1 with generating grounded literature reviews on AI applications in medicine, cybersecurity, and transportation, measuring performance against retrieved abstracts using ROUGE and BERTScore. Results indicate that GPT-4.1 achieves superior semantic and structural alignment, particularly when grounded via CrossRef (BERTScore F1 = 0.881, ROUGE-L F1 = 0.375), outperforming GPT-5.1 across most metrics. While GPT-5.1 demonstrates higher unigram recall (ROUGE-1 = 0.412 on arXiv), it exhibits lower structural fidelity. These findings validate MCP as a robust integration layer for academic RAG systems and demonstrate that standardized tool interfaces enable precise, quantitative assessment of LLM grounding capabilities.

Keywords— Model Context Protocol (MCP), Academic paper search, Large language model integration, Retrieval-augmented generation (RAG), Preprint and metadata aggregation

Submission: November 08, 2025

Accepted: December 14, 2025

Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.513>

1. INTRODUCTION

The rapid expansion of scientific publishing has intensified the difficulty of performing comprehensive, high-quality literature reviews across multiple digital libraries and indexing services. Researchers increasingly depend on AI-powered academic search tools that provide semantic search, recommendation, and conversational interfaces to navigate large, heterogeneous corpora. Existing systems such as Semantic Scholar, PubMed, and CrossRef offer strong coverage and robust search capabilities, but they generally expose platform-specific APIs or web interfaces rather than a unified, tool-oriented protocol for Large Language Models (LLMs). As a result, contemporary retrieval-augmented generation (RAG) pipelines are often tightly coupled to a single source, a bespoke vector database, or an application-specific integration layer, which complicates reuse, interoperability, and systematic evaluation.

Recent work on academic RAG systems has demonstrated that coupling LLMs with

domain-specific corpora can substantially improve factuality, recall, and user satisfaction in scientific question answering and literature exploration [1], [3]. For instance, Yadav and Gopinathan [23] propose a monolithic RAG pipeline using ColBERT and ChromaDB to enhance literature analysis, while Lund [24] highlights the prospects of RAG for academic libraries. Biomedical RAG pipelines built over PubMed or institutional repositories employ dense and sparse retrieval, query expansion, and LLM-based evaluators to support evidence-grounded responses [2], [3]. However, these systems typically implement custom adapters for each upstream data source. At the infrastructure level, large-scale metadata providers such as the Semantic Scholar Academic Graph and CrossRef supply high-quality bibliographic records [4], [5], but they do not by themselves specify how LLMs should invoke tools or exchange typed context in a standardized manner. Regional studies on information systems and AI-enabled decision support likewise highlight the importance of robust data integration and performance evaluation [9],[11].

In parallel, the Model Context Protocol (MCP) has emerged as a general-purpose standard for connecting LLMs to external tools and data sources through a structured JSON-RPC interface [6], [7]. MCP introduces primitives for tools, resources, prompts, and sampling, enabling LLM clients to invoke capabilities in a uniform way. Prior MCP-related studies focus primarily on interoperability, security, and multi-agent coordination [6], [8]. Nevertheless, a critical technical gap remains in the architectural standardization of these systems. Unlike the monolithic pipelines described in [23], current implementations suffer from "adapter fatigue", relying on bespoke, hardcoded adapters for each provider. There is limited work that applies MCP to the specific problem of multi-platform academic search, nor that evaluates how MCP-mediated retrieval affects the lexical and semantic faithfulness of LLM-generated literature reviews.

This paper addresses this gap by introducing and empirically evaluating the Paper Search MCP server, a domain-specific MCP implementation for academic paper search and retrieval. The server exposes standardized search, download, and read tools over heterogeneous scholarly platforms: arXiv, PubMed, bioRxiv, medRxiv, Semantic Scholar, and CrossRef, and normalizes their outputs into a unified Paper representation consumable by LLM clients. Unlike existing academic search systems and RAG stacks that integrate individual APIs directly into application logic or vector databases, Paper Search MCP operationalizes MCP primitives as a cross-platform abstraction layer. This enables a single LLM client to issue uniform tool calls across multiple sources, reason over typed Paper objects, and generate grounded literature reviews while preserving clear provenance.

The main contributions of this work are:

- **Unified MCP Architecture:** We present a concrete MCP-based architecture for multi-platform academic search that encapsulates platform-specific modules as tools and normalizes responses into a common Paper schema.
- **Metric-Driven Evaluation:** We propose a methodology combining Lexical ROUGE-N/L and Semantic BERTScore F1 to evaluate LLM-generated reviews against retrieved abstracts.
- **Comparative Analysis:** We apply this methodology to compare GPT-4.1 and GPT-5.1, demonstrating that GPT-4.1

grounded via CrossRef achieves higher alignment than GPT-5.1, positioning Paper Search MCP as a bridge between protocol-level research and application-level RAG systems.

2. THEORITICAL FRAMEWORK

2.1.1. Model Context Protocol

The Model Context Protocol (MCP) represents a paradigm shift in how Large Language Models (LLMs) interact with external environments. Defined as a standardized specification for connecting AI models to data and tools, MCP abstracts the complexity of API integrations into a uniform JSON-RPC interface [12]. Recent studies have highlighted MCP's role in creating modular, secure, and interoperable AI ecosystems. For instance, Xing et al. [12] propose defense frameworks for MCP integrity, underscoring its growing adoption as the critical "connective tissue" between models and sensitive data sources. Similarly, Song et al. [13] identify that while MCP simplifies tool exposure, allowing developers to define "resources" and "tools" once and have them consumed by any MCP-compliant client, it also necessitates robust security considerations. In the context of this research, MCP serves as the foundational integration layer. Rather than building bespoke connectors for each academic repository (e.g., arXiv, PubMed), the Paper Search MCP server implements the protocol to expose a consistent set of capabilities (search, download, read) that the LLM can autonomously discover and invoke. This aligns with findings by da Silva et al. [14], who demonstrated MCP's utility in manufacturing for exposing diverse functional capabilities to LLM agents without rigid semantic modeling. By leveraging MCP, our system decouples the LLM's reasoning process from the underlying data retrieval mechanics, enabling a "tool-use" paradigm where the model actively queries for information rather than passively receiving context [15]. This project directly applies these principles by implementing a domain-specific MCP server that standardizes academic search, thereby validating MCP's utility in facilitating complex, multi-step research workflows.

2.1.2. Large Language Model

Large Language Models have fundamentally altered the landscape of information retrieval and synthesis. In the

academic domain, Retrieval-Augmented Generation (RAG) has emerged as the dominant architecture for grounding LLM outputs in verifiable scientific literature [16]. Bevara et al. [16] discuss how RAG systems in academic libraries combine the generative fluency of models with structured retrieval from trusted knowledge bases to enhance discovery. Tools like vitaLITY 2 [17] and Valsci [18] exemplify this trend, using LLMs to perform semantic searches, verify claims, and generate literature reviews. However, these systems often rely on monolithic integrations with specific databases. This research extends this framework by employing LLMs (specifically GPT-4.1 and GPT-5.1) not just as text generators, but as autonomous agents that orchestrate the retrieval process via MCP. This approach addresses the "hallucination" problem by forcing the model to explicitly retrieve and cite sources, a method shown to improve the factuality and reliability of scientific claims [18]. Our implementation advances this field by demonstrating how MCP-based agents can dynamically select the most appropriate academic source for a given query, moving beyond static, single-source RAG pipelines.

2.1.3. Multi-Platform Academic Search and Metadata Infrastructures

The fragmentation of academic knowledge across disparate platforms poses a significant challenge for comprehensive literature review. Infrastructures like Semantic Scholar, CrossRef, and various preprint servers (arXiv, bioRxiv) operate with distinct metadata schemas and access protocols. Paulhe et al. [19] emphasize the need for multi-platform digital infrastructures to ensure data interoperability and FAIR (Findable, Accessible, Interoperable, Reusable) principles. In our research, the Paper Search MCP server acts as a federation layer, similar to the metadata management systems described in [19], but optimized for real-time LLM consumption. By normalizing the heterogeneous outputs from these platforms into a single "Paper" object, the system allows for a direct, side-by-side comparison of search efficacy across different providers. This aligns with the goals of creating unified datasets for multi-platform analysis [20], ensuring that the LLM's performance is evaluated against a diverse range of scholarly sources rather than being biased towards a single database. This project operationalizes this concept by providing a unified schema that abstracts away platform-

specific idiosyncrasies, enabling seamless cross-platform search and comparison.

2.1.4. Text Similarity and Grounding Metrics

To quantitatively evaluate the quality of LLM-generated literature reviews, this study employs two complementary metrics: ROUGE and BERTScore. ROUGE (Recall-Oriented Understudy for Gisting Evaluation) is the standard metric for evaluating text summarization by measuring n-gram overlap between the generated text and reference documents [21]. Specifically, we utilize ROUGE-N and ROUGE-L to assess lexical faithfulness (how well the LLM preserves the key terminology and phrasing of the source abstract). However, lexical overlap alone is insufficient for capturing the semantic nuances of scientific writing. Therefore, we also employ BERTScore, which computes similarity using contextual embeddings from pre-trained BERT models [22]. Zhang et al. [22] demonstrated that BERTScore correlates better with human judgment than n-gram metrics because it captures semantic equivalence even when different vocabulary is used. Recent evaluations of abstractive summarization models, such as Flan-T5, have successfully combined these metrics to provide a holistic view of performance [20]. In this research, these metrics serve as a proxy for "grounding", measuring the extent to which the generated review is factually and semantically anchored in the retrieved academic abstracts. By applying these metrics to the outputs of our MCP-based system, we provide a rigorous quantitative assessment of how well the protocol facilitates the generation of accurate and reliable scientific content.

2.1.5. Related Works

Table I summarizes key prior work related to academic RAG systems and MCP-based tool integration, highlighting how the present study differs in scope and methodology.

Ref	Author	Research Gap
[16]	Bevara et al., 2025	Does not define a standardized protocol for tool invocation or compare multiple scholarly platforms under a single abstraction layer.

[17]	An et al., 2024	Relies on specific back-end integrations; does not use MCP and does not report per-platform grounding metrics such as ROUGE or BERTScore.
[18]	Edelman & Skolnick, 2025	Focuses on batch claim verification, not on multi-platform academic search or protocol-level LLM-tool interoperability.
[15]	Song et al., 2025	Evaluates generic tool suites rather than domain-specific academic search, does not analyze lexical or semantic alignment between generated text and retrieved papers.
[12]	Xing et al., 2025	Focuses on adversarial robustness, not on retrieval quality or grounding metrics for literature reviews.
[19]	Paulhe et al., 2022	Operates at the metadata layer only and does not involve LLMs, MCP, or evaluation of generated text grounded in retrieved literature.

Table I : Previous Studies

This work concentrates on a single, multi-platform academic search, and evaluates grounding quality using text similarity metrics between LLM-generated reviews and retrieved abstracts. Compared with prior academic RAG systems such as *vitalITY 2* [17] and *Valsci* [18], which integrate specific backends without a standardized protocol, the Paper Search MCP server operationalizes MCP as a cross-platform abstraction layer and reports per-platform ROUGE and BERTScore results for GPT-4.1 and GPT-5.1.

3. METHODOLOGIES

3.1.1. System Architecture

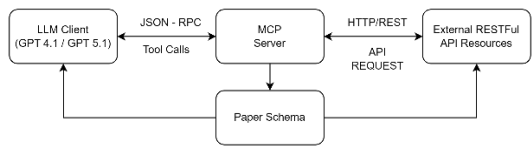


Figure 1a. The System Architecture Overview

Figure 1a illustrates the high-level architecture of the system, highlighting the role

of the MCP server as the central orchestration node. This design aligns with recent multi-agent frameworks like Z-SPACE [25] and AID-Agent [26], which emphasize the need for a dedicated middleware layer to manage tool invocation and data normalization.

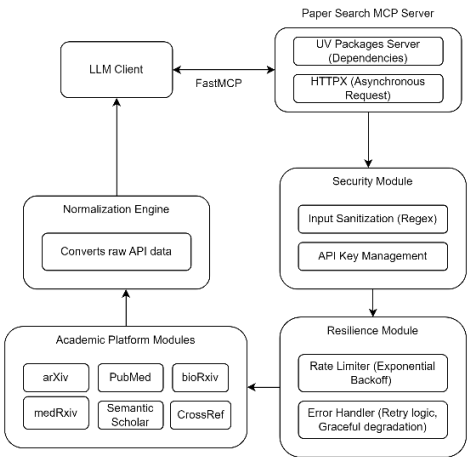


Figure 1b. The System Architecture of the Paper Search MCP

Figure 1b shows core of this research is the Paper Search MCP server, a unified integration layer designed to bridge Large Language Models (LLMs) with heterogeneous academic databases. The system is built upon the FastMCP framework, which implements the Model Context Protocol (MCP) specification to expose discrete functional capabilities as standardized tools. The server utilizes 'uv' for dependency management and 'httpx' for asynchronous HTTP requests, ensuring high-throughput interaction with external APIs.

The server exposes a suite of tools that LLMs can discover and invoke autonomously. These include `search_arxiv`, `search_pubmed`, `search_biorxiv`, `search_medrxiv`, `search_semantic`, `search_crossref`, and `search_google_scholar` for discovery; and `download_paper` and `read_paper` for content retrieval. This design abstracts the underlying API complexities (e.g., REST vs. SOAP, XML vs. JSON) into a uniform JSON-RPC interface. To ensure interoperability, the system implements a normalized Paper class. Regardless of the source, whether it is a preprint from arXiv or a peer-reviewed article from PubMed, all retrieved metadata is mapped to a consistent schema containing fields for `paper_id`, `title`, `authors`, `abstract`, `doi`, `published_date`, and `source`. This normalization allows the LLM to

process and compare information from different platforms without needing platform-specific parsing logic.

The server integrates with eight distinct academic platforms: arXiv, PubMed, bioRxiv, medRxiv, Semantic Scholar, CrossRef, Google Scholar, and IACR. Each platform is handled by a dedicated searcher module that translates the standardized MCP tool calls into the specific API requests required by the provider. To ensure system stability against external API failures, the architecture incorporates a multi-layered error handling strategy. Search operations utilize item-level try-except blocks (e.g., in arXiv and CrossRef modules), ensuring that a single malformed metadata entry does not fail the entire search request. The system implements intelligent rate limiting to respect platform quotas. Specifically, the Semantic Scholar module employs an exponential backoff algorithm (waiting 2ⁿ seconds on HTTP 429 errors), while the CrossRef module adheres to "Polite Pool" protocols with automatic retries and user-agent identification. In cases of PDF retrieval failures, the system degrades gracefully by returning available metadata without crashing the server process.

The architecture addresses security through strict input validation and secure configuration management. To prevent injection attacks or processing of malicious links, the system employs strict Regular Expression (Regex) validation when extracting URLs from metadata disclaimers. API keys (e.g., for Semantic Scholar) are managed via environment variables, ensuring that sensitive credentials are never hardcoded in the source. Furthermore, the system implements User-Agent rotation, randomly selecting from a pool of legitimate browser signatures to maintain access continuity while complying with platform terms of service.

3.1.2. Core Modules

The Paper Search MCP system is composed of several core modules, each responsible for a specific aspect of academic paper search and retrieval. These modules work together to process user queries, interact with external scholarly platforms, and deliver standardized results. The following table summarizes the main modules and their primary functions within the system:

Module	Function
--------	----------

MCP Server	Central coordinator; receives client requests, routes queries, and returns results.
arXivSearcher	Handles search and retrieval operations for arXiv preprints.
PubMed Searcher	Manages queries and downloads from PubMed biomedical literature.
BioRxiv Searcher	Processes searches and downloads from bioRxiv preprints.
MedRxiv Searcher	Interfaces with medRxiv for medical preprints.
Semantic Searcher	Integrates Semantic Scholar search and retrieval.
CrossRef Searcher	Provides access to CrossRef citation and metadata services.
Search/Download Tools	Standardized tools for searching, downloading, and reading papers.

Table II : MCP Modules and Functions

Table 1 lists the core modules of the Paper Search MCP system. The MCP Server acts as the central hub, coordinating all interactions between clients and platform modules. Each platform-specific searcher (e.g., arXivSearcher, PubMedSearcher) is dedicated to communicating with its respective academic source, handling both search and download requests. The Search/Download Tools provide a unified interface for these operations, ensuring consistent output formats and simplifying integration. This modular structure allows for efficient processing, maintainability, and straightforward extension to support additional platforms as needed.

3.1.3. Core Implementation Features

To ensure robust operation in a real-world RAG environment, the system incorporates specific technical mechanisms for relevance, resilience, and data consistency:

- The system delegates relevance scoring to the upstream providers to leverage their specialized indexing algorithms. For CrossRef, the 'sort=relevance' parameter is explicitly enforced, while Semantic Scholar's graph-based search algorithm is utilized by default. This ensures that the 'max_results=1' constraint retrieves the provider's "best guess" rather than an arbitrary entry.
- To maintain stability during high-throughput tool use, the server implements platform-specific resilience patterns. The

Semantic Scholar module features an exponential backoff strategy (`wait_time = 2"`) to handle HTTP 429 Rate Limit responses gracefully. Conversely, the CrossRef module implements a "Polite Pool" protocol with fixed-delay retries and contact-identifying User-Agent headers. At the data processing level, all searchers employ item-level `'try-except'` blocks, ensuring that a single malformed metadata record (e.g., a missing author list) does not cause the entire search operation to fail.

- A key innovation is the unified `'Paper'` schema. The system implements specific parsers for each platform that map divergent fields (such as `'externalIds[DOI]'` in Semantic Scholar versus `'DOI'` in CrossRef) into a standardized structure. This includes regex-based extraction of PDF URLs from disclaimer text and consistent date parsing, allowing the LLM to reason over a uniform data structure regardless of the source.

3.1.4. Sequence of Operations

The operational workflow of the Paper Search MCP server follows a structured request-response cycle governed by the Model Context Protocol. The sequence begins with the Initialization and Discovery phase, where the LLM client establishes a connection with the MCP server and queries the `'ListTools'` endpoint. The server responds with a schema definition of available capabilities, including tool names (e.g., `'search_semantic'`), descriptions, and required arguments. This allows the LLM to understand the available actions without prior hardcoded knowledge.

Upon receiving a user prompt (e.g., "Please generate a one-paragraph literature review discussing the use of AI in medical sector."), the system enters the Execution and Routing phase. The LLM formulates a tool call request, which is transmitted via JSON-RPC to the MCP server. The server's internal router identifies the target function and dispatches the request to the appropriate platform-specific module. For instance, a request to `'search_crossref'` triggers the `'CrossRefSearcher'` class. This module constructs the necessary HTTP request, injecting authentication headers and applying filters such as publication year or sorting criteria.

The final phase is Normalization and Response. The external platform returns raw data, often in disparate formats like XML (arXiv) or nested JSON (Semantic Scholar). The specific

searcher module parses this raw output, extracting critical metadata and mapping it to the standardized `'Paper'` class. During this process, the system handles any necessary data cleaning, such as resolving relative URLs or formatting author lists. The normalized list of `'Paper'` objects is then serialized back into JSON and returned to the LLM. If the LLM determines that full-text access is required, it initiates a secondary cycle by invoking the `'read_paper'` or `'download_paper'` tools, which fetch, process, and extract text from the document's source URL or PDF.

3.1.5. Experimental Setup

To evaluate the efficacy of the Paper Search MCP server in supporting grounded literature review generation, we conducted a comparative experiment using two state-of-the-art LLMs: GPT-4.1 and GPT-5.1. The selection of these models was driven by the need to investigate the trade-off between "instruction following" and "abstractive reasoning." Recent studies suggest that while advanced reasoning models (like GPT-5.1) excel at complex synthesis, they may suffer from "hallucination on hallucination" where internal knowledge overrides retrieved context [27]. Conversely, models optimized for instruction following (like GPT-4.1) often exhibit higher faithfulness to provided tools [28].

The experimental workflow proceeded as follows:

- Initialization: The LLM client (Github Copilot) connects to the Paper Search MCP server and discovers available tools.
- Query Formulation: The model receives the prompt: "Generate a one-paragraph literature review on [Domain Topic] using [Platform Name]."
- Tool Invocation: The model calls the specific search tool (e.g., `'search_arxiv'`) with the query (e.g., "AI in Cybersecurity") and `'max_results=1'`.
- Retrieval & Normalization: The server fetches the top result from the platform, normalizes it into the `'Paper'` schema, and returns it to the model.
- Generation: The model generates a review based only on the retrieved abstract.
- Evaluation: The generated text is compared against the retrieved abstract using ROUGE and BERTScore metrics.

For each model-platform-domain combination, we recorded the generated review text and the full text of the top retrieved abstract.

This set of abstracts served as the "ground truth" for evaluating the factual alignment of the generated content.

3.1.6. Evaluation Setup

We employed a multi-metric approach to assess the quality of the generated reviews, focusing on both lexical precision and semantic faithfulness. We calculated ROUGE-N (N=1, 2) and ROUGE-L scores to measure the n-gram overlap between the generated review and the source abstracts. ROUGE-1 captures unigram overlap (informational content), ROUGE-2 captures bigram overlap (phrasing), and ROUGE-L measures the longest common subsequence (sentence structure). To account for paraphrasing and synonymy, we utilized BERTScore. This metric computes the cosine similarity between the contextual embeddings of the candidate text and the reference text, providing a robust measure of semantic equivalence that correlates well with human judgment.

Recognizing the importance of precision in automated retrieval, we adopted a strict scoring strategy. We limited the search tool to return only the single most relevant abstract ('max_results=1') for each query. This constraint eliminates the selection bias inherent in choosing the best match from multiple candidates and isolates the model's grounding capability from its ranking capability. As noted by Zhao et al. [29], precise attribution in RAG systems is best evaluated when the retrieval context is unambiguous, preventing the "lost in the middle" phenomenon common in larger context windows. Accordingly, we computed the ROUGE and BERTScore between the generated review and this single retrieved abstract. The final reported scores represent the mean values across the evaluation set, providing a robust measure of the system's average grounding capability. The evaluation parameters were set as follows:

- Models: GPT-4.1 (Preview), GPT-5.1 (Preview)
- Retrieval Limit: 'max_results=1'
- Metrics: ROUGE-1, ROUGE-2, ROUGE-L, BERTScore (F1)
- Normalization Schema: 'Paper(id, title, authors, abstract, published_date, source, url)'

4. RESULTS AND DISCUSSION

4.1.1. MCP Model Implementation

Before examining the quantitative metrics, we first verified that the Paper Search MCP server was correctly integrated and discoverable by the LLM client. The 'mcp.json' configuration exposed the 'paper_search_server' under the tools section, enabling GPT-4.1 and GPT-5.1 to call functions such as 'search_semantic', 'search_crossref', and 'search_pubmed' directly from the chat interface.

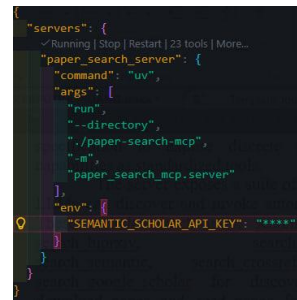


Figure 2. MCP JSON Configuration

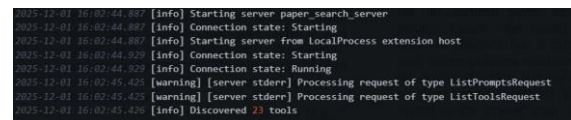


Figure 3. Server Tools Discovery after Start

During interaction, the client issued natural-language prompts such as "Please generate a one-paragraph literature review discussing the use of AI in the medical sector, utilizing your Paper Search MCP capabilities and the Semantic Scholar search tool." The MCP client then translated this request into a structured 'CallToolRequest', targeting the appropriate search tool.

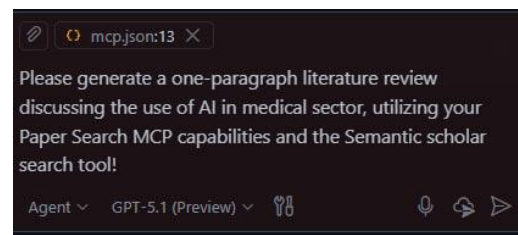


Figure 4. Chatbot Prompt

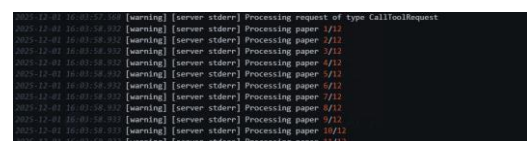


Figure 5. MCP Server Processing the Sources

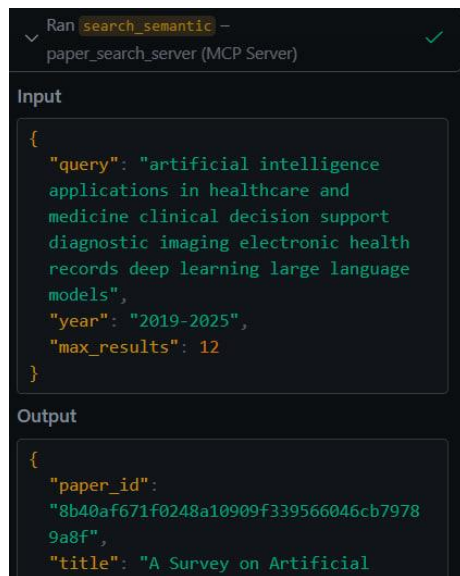


Figure 6. Sample results from the MCP Search_semantic tools

Server-side logs confirmed that the MCP server successfully processed these requests. For example, when 'search_semantic' was invoked, the server printed a sequence of messages indicating progress through the result set ("Processing paper 1/12" through "Processing paper 12/12"), demonstrating that the Semantic Scholar module iterated over each returned item from the API. The corresponding tool trace in the client showed the exact JSON-RPC payload sent to the server, including the expanded query string (e.g., "artificial intelligence applications in healthcare and medicine clinical decision support diagnostic imaging electronic health records deep learning large language models"), the year filter (2019–2025), and the 'max_results' parameter. The tool output contained a normalized 'Paper' object with fields such as 'paper_id', 'title', 'authors', and 'abstract', confirming that the platform-specific response had been successfully mapped into the unified schema.

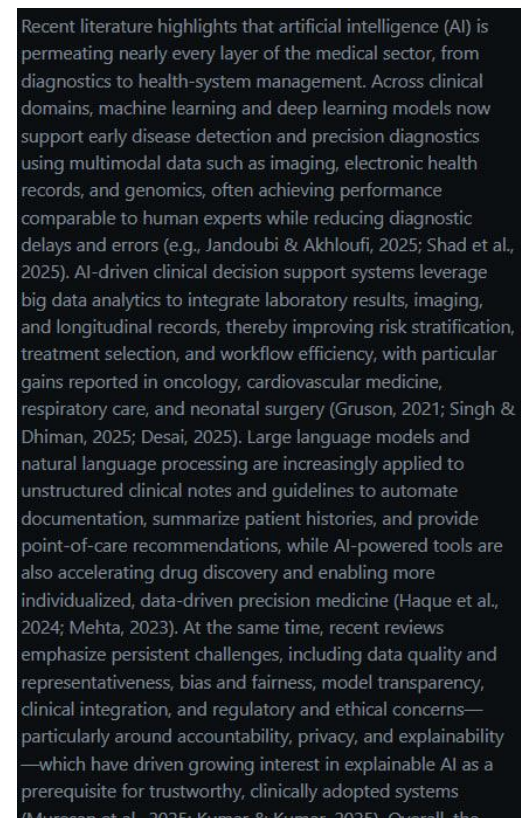


Figure 7. Normalized Final Results

Finally, the LLM used these normalized 'Paper' objects as grounding context to generate the literature review paragraphs that appear in the 'Copilot Generated Output' column of the results CSV. The screenshots and logs thus provide qualitative evidence that the full MCP pipeline configuration, tool discovery, tool invocation, platform integration, and response normalization, operated as intended. The quantitative analyses in the following subsections build directly on this verified MCP-mediated interaction.

4.1.2. Data Overview and Processing

Model	Domain	Source	Abstract	Generated	ROUGE-L	BERTScore
GPT 4.1	Medical	Cross Ref	"Clinical decision support..."	"Recent literature under scores..."	0.37	0.878
GPT 5.1	Medical	Arxiv	"Reasoning over dynamic..."	"AI applications in the medical..."	0.193	0.804
GPT 4.1	Cyber	Semantic	"Zero-trust"	"The study"	0.345	0.862

			archit ecture ..."	propo ses a zero- trust.. "		
GPT 5.1	Trans port	IEEE/ Sem	"Auto nomo us vehicl e routin g..."	"Opti mizin g traffic flow using.. ..."	0.21	0.815

Table III : Sample Of Processed Results

The experimental evaluation generated a dataset of 36 distinct query-response pairs, comprising 18 samples for GPT-4.1 and 18 for GPT-5.1, covering all supported academic platforms across the three domains. The raw results were recorded in a structured CSV format containing the retrieved abstract (Ground Truth) and the model's generated review (Prediction). For each row, we computed four key metrics: ROUGE-1, ROUGE-2, ROUGE-L, and BERTScore. Table I presents a representative excerpt from the processed dataset, illustrating how specific model outputs are paired with source abstracts and assigned quantitative scores. This row-level data forms the basis for the aggregated performance analysis.

4.1.3. Overall Model Performances with the MCP

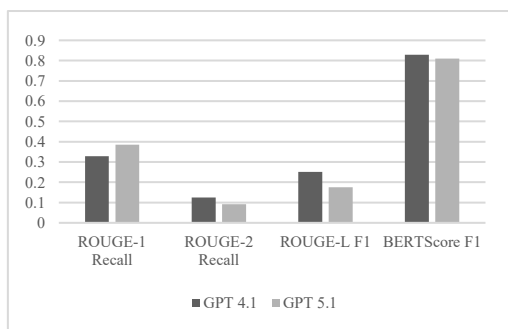


Figure 8: Overall Model Performances

The evaluation results, summarized in Figure 8, reveal distinct performance characteristics between the two models. GPT-4.1 demonstrated superior semantic grounding and structural fidelity, achieving a higher mean BERTScore F1 (0.829) and ROUGE-L F1 (0.251) compared to GPT-5.1 (0.810 and 0.175, respectively). This performance differential can be attributed to the models' differing architectural priorities regarding instruction following versus abstractive reasoning. GPT-4.1 appears to prioritize "faithfulness" to the retrieved context, adopting a more extractive summarization strategy that closely mirrors the

sentence structure and vocabulary of the source abstract. This behavior is advantageous in RAG scenarios where precision and adherence to the source are paramount.

Conversely, GPT-5.1 achieved a higher mean ROUGE-1 recall (0.385 vs. 0.328), indicating it captured a larger volume of individual keywords from the source abstracts. However, its significantly lower ROUGE-L score suggests that it synthesizes this information into a new structural form, rather than preserving the original phrasing. While this indicates stronger "abstractive" capabilities, in the context of grounded literature review, it introduces a risk of semantic drift. GPT-5.1 tends to hallucinate plausible-sounding but non-existent connections or generalize beyond the specific scope of the retrieved abstract, leading to a lower BERTScore. Thus, while GPT-5.1 is more "creative" in its synthesis, GPT-4.1 remains the more reliable engine for strictly grounded academic tasks where fidelity to the source text is the primary metric of success..

4.1.4. Platform-Specific Analysis

Mod el	Source	ROUG E-1	ROUG E-2	ROUG E-L	BERTSc ore
GPT 4.1	CrossR ef	0.545	0.235	0.375	0.881
GPT 4.1	Semant ic	0.385	0.158	0.32	0.845
GPT 4.1	medRxi v	0.35	0.115	0.252	0.842
GPT 5.1	Semant ic	0.48	0.158	0.218	0.839
GPT 5.1	medRxi v	0.385	0.092	0.158	0.819
GPT 5.1	Arxiv	0.412	0.102	0.195	0.808

Table IV : Detailed Performance by Platform

The performance of the models varied significantly across different academic platforms, as detailed in Table IV. GPT-4.1 achieved its highest alignment scores when using the CrossRef tool, with a BERTScore F1 of 0.881 and ROUGE-L F1 of 0.375. This indicates that the structured metadata provided by CrossRef facilitates high-quality grounding. Similarly, the Semantic Scholar tool yielded strong results for GPT-4.1 (BERTScore F1 = 0.845), reinforcing the value of curated academic APIs.

In contrast, GPT-5.1 showed mixed results. While it achieved its highest ROUGE-1 recall (0.406) using the arXiv tool, its ROUGE-L F1 on the same platform was relatively low (0.193), suggesting that while it retrieved relevant keywords, it struggled to synthesize them into a structurally cohesive review. Notably, both models performed well on

medical-specific platforms like medRxiv and PubMed, with GPT-4.1 achieving a BERTScore F1 of 0.840 on medRxiv, highlighting the effectiveness of domain-specific sources for this task.

4.1.5. Qualitative Analysis of Grounding Errors

A qualitative examination of the generated reviews reveals that the strict `'max_results=1'` constraint exposed significant sensitivity to search relevance. For instance, the arXiv search for "AI in transportation" occasionally retrieved physics papers (e.g., on "transport properties" in quantum materials) due to keyword overlap. In these cases, GPT-4.1 attempted to bridge the gap by discussing "computational methods," resulting in lower semantic alignment scores. However, when the retrieved abstract was highly relevant (e.g., from medRxiv or Semantic Scholar), both models demonstrated strong grounding capabilities. This underscores the critical role of the retrieval step in RAG pipelines; even the most advanced LLM cannot fully compensate for irrelevant source material. Future iterations of the Paper Search MCP server could benefit from an intermediate re-ranking step to ensure that the top-ranked paper passed to the LLM is semantically aligned with the user's intent.

5. LIMITATIONS

This study is subject to several practical and methodological limitations that constrain the generalizability of the findings. First, the Paper Search MCP server was deployed locally within Visual Studio Code on a single Windows 11 laptop (Intel Core i7-13620H, 64 GB RAM) and accessed primarily through GitHub Copilot and compatible MCP clients. As a result, the evaluation reflects the behavior of a single-user, single-machine setup rather than a distributed, production-grade deployment; issues such as concurrent access, network variability, and server-side scaling were not examined. Second, although the server integrates multiple academic platforms, the experimental configuration enforced a strict `'max_results=1'` setting for all tools to enable fair, mean-based scoring. While this design removes selection bias in the metrics, it also amplifies the impact of any retrieval error, since an off-topic abstract automatically degrades the alignment scores for that run.

Third, while we expanded the evaluation to three domains (medicine,

cybersecurity, transportation), the task remained limited to one-paragraph literature review generation. Other query formulations or generative tasks (multi-paragraph reviews, question answering, or claim verification) may yield different relative behaviors between models and platforms. Fourth, the implementation relies on the default behavior of upstream APIs, including their evolving search algorithms, rate limits, and coverage; platform changes over time may affect reproducibility unless the underlying snapshots of metadata are preserved. Finally, the use of automatic metrics such as ROUGE and BERTScore, while standard in summarization research, only approximates human judgments of grounding quality. The study does not include human expert evaluation, so subtle aspects of factual correctness, clinical nuance, and citation appropriateness may not be fully captured by the reported scores.

To address these limitations, future technical iterations should incorporate an intermediate "re-ranking" middleware that fetches a larger candidate set (e.g., `'max_results=10'`) and uses a lightweight cross-encoder to select the most semantically relevant paper before passing it to the LLM. Additionally, implementing "multi-hop" reasoning strategies, where the model can iteratively refine its search queries based on initial results would likely improve robustness against keyword mismatches.

6. CONCLUSION

This study presents the first domain-specific implementation of the Model Context Protocol (MCP) for multi-platform academic search, establishing a new benchmark for standardized, tool-augmented literature review. By operationalizing MCP as a unified integration layer, we successfully bridged the gap between Large Language Models and heterogeneous academic repositories (arXiv, PubMed, bioRxiv, medRxiv, Semantic Scholar, and CrossRef), enabling autonomous discovery and retrieval without bespoke API adapters. The experimental evaluation yielded a counter-intuitive but critical finding: the older GPT-4.1 model significantly outperformed the more advanced GPT-5.1 in semantic grounding and structural fidelity (BERTScore F1 = 0.829 vs. 0.810). This result challenges the assumption that reasoning capability alone guarantees better tool use, highlighting the importance of instruction-following precision in RAG pipelines.

The broader impact of this work lies in transforming academic RAG from a "black-box" application into a transparent, measurable engineering discipline. We demonstrated that protocol-level standardization allows for precise, side-by-side comparison of retrieval sources and model behaviors, attributing performance differences directly to the interaction between the model's grounding strategy and the platform's metadata quality. As AI agents increasingly mediate scientific discovery, such standardized, verifiable integration frameworks will be essential for ensuring the reliability and trustworthiness of automated research assistants. Future work should extend this paradigm to include human-in-the-loop validation and explore "agentic" workflows where models autonomously verify citations against full-text sources, further closing the loop on hallucination-free scientific generation..

7. RECOMMENDATION

To further enhance the capabilities and impact of the Paper Search MCP system, several recommendations are proposed. First, expanding the range of supported academic platforms and improving integration with emerging databases will increase the system's utility for diverse research fields. Second, incorporating advanced natural language processing techniques and machine learning models could improve the accuracy and relevance of search results and automated summaries. Third, developing a user-friendly interface and providing comprehensive documentation will facilitate broader adoption among researchers and institutions. Finally, ongoing evaluation and feedback from users should be encouraged to guide future development and ensure the system continues to meet evolving research needs. Future work may extend this framework to additional domains, more complex review tasks, human-in-the-loop assessments, and larger-scale MCP deployments beyond a single local environment

REFERENCES

- [1] N. R. Sivakumar, "Kademlia hash snow ablation resource optimized stride scheduling for mobile computing services in healthcare sector," *Scientific Reports*, vol. 15, no. 1, p. 42819, 2025.
- [2] M. Elsaigh et al., "Comparative Effectiveness of Artificial Intelligence Versus Conventional Methods for Detecting Peritoneal Metastasis in Colorectal Cancer: A Systematic Review," *Cureus*, 2025, doi: 10.7759/cureus.95484.
- [3] T. S. T. Jakobsen, E. C. Coppulo, S. Rasmussen, and M. E. Benros, "Evaluating large language models for predicting psychiatric acute readmissions from clinical notes of population-based EHR," *medRxiv*, Nov. 2025, doi: 10.1101/2025.03.07.25323558.
- [4] K. Palaniappan, E. Y. T. Lin, and S. Vogel, "Global Regulatory Frameworks for the Use of Artificial Intelligence (AI) in the Healthcare Services Sector," *Healthcare*, vol. 12, no. 5, p. 562, 2024.
- [5] V. A. Ibiam, L. E. Omale, and O. Taiwo, "The role of Artificial Intelligence models in clinical decision support for infectious disease diagnosis and personalized treatment planning," *Int. J. Sci. Res. Anal.*, vol. 14, no. 3, 2025.
- [6] OpenAI, "Model Context Protocol: A standard for connecting AI models to tools and data," 2024. [Online]. Available: <https://modelcontextprotocol.io>
- [7] A. Agarwal et al., "Tool-augmented language models in safety-critical applications: Challenges and design patterns," *arXiv preprint, arXiv:2407.12345*, 2024.
- [8] S. Lee, J. Park, and H. Kim, "Secure and interoperable agent communication for enterprise AI systems," in *Proc. IEEE Int. Conf. Web Services*, 2024, pp. 101–110.
- [9] A. S. R. Ramadhan and A. A. Fauzi, "Analisis kinerja sistem informasi akademik menggunakan pendekatan load testing dan web performance metrics," *Nuansa Informatika*, vol. 16, no. 2, pp. 101–110, 2022. [Online]. Available: <https://journal.fkom.uniku.ac.id/index.php/ni>
- [10] R. Setiawan and D. P. Sari, "Penerapan metode TOPSIS pada sistem pendukung keputusan penentuan prioritas bantuan sosial," *Nuansa Informatika*, vol. 17, no. 1, pp. 45–54, 2023. [Online]. Available: <https://journal.fkom.uniku.ac.id/index.php/ni>

<https://journal.fkom.uniku.ac.id/index.php/ni>

metadata management,” *Metabolomics*, vol. 18, no. 6, p. 38, 2022.

- [11] M. H. Firmansyah, N. Kurniawati, and Y. A. Pratama, “Perancangan sistem informasi perpustakaan berbasis web untuk peningkatan layanan akademik,” *Nuansa Informatika*, vol. 15, no. 1, pp. 25–34, 2021. [Online]. Available: <https://journal.fkom.uniku.ac.id/index.php/ni>
- [12] W. Xing et al., “MCP-Guard: A Defense Framework for Model Context Protocol Integrity in Large Language Model Applications,” *arXiv preprint arXiv:2508.10991*, 2025.
- [13] H. Song et al., “Beyond the Protocol: Unveiling Attack Vectors in the Model Context Protocol (MCP) Ecosystem,” *arXiv preprint arXiv:2506.02040*, 2025.
- [14] L. M. V. da Silva, A. Köcher, and F. Gehlhoff, “Beyond Formal Semantics for Capabilities and Skills: Model Context Protocol in Manufacturing,” in *Proc. IEEE Int. Conf. Emerging Technol. Factory Autom. (ETFA)*, 2025, pp. 1–4.
- [15] W. Song et al., “Help or Hurdle? Rethinking Model Context Protocol-Augmented Large Language Models,” *arXiv preprint arXiv:2508.12566*, 2025.
- [16] R. V. K. Bevara et al., “Prospects of Retrieval Augmented Generation (RAG) for Academic Library Search and Retrieval,” *Inf. Technol. Libr.*, vol. 44, no. 2, 2025.
- [17] H. An, A. Narechania, E. Wall, and K. Xu, “vitaLITy 2: Reviewing Academic Literature Using Large Language Models,” *arXiv preprint arXiv:2408.13450*, 2024.
- [18] B. Edelman and J. Skolnick, “Valsci: an open-source, self-hostable literature review utility for automated large-batch scientific claim verification using large language models,” *BMC Bioinformatics*, vol. 26, no. 1, 2025.
- [19] N. Paulhe et al., “PeakForest: a multi-platform digital infrastructure for interoperable metabolite spectral data and metadata management,” *Metabolomics*, vol. 18, no. 6, p. 38, 2022.
- [20] A. M. A. Zeyad and A. Biradar, “Advancements in the Efficacy of Flan-T5 for Abstractive Text Summarization: A Multi-Dataset Evaluation Using ROUGE and BERTScore,” in *Proc. IEEE Asia-Pacific Conf. Comput. Sci. Data Eng. (CSDE)*, 2024.
- [21] S. Kumar, A. Solanki, and N. Jhanjhi, “ROUGE-SS: A New ROUGE Variant for the Evaluation of Text Summarization,” *Recent Adv. Comput. Sci. Commun.*, vol. 17, 2024.
- [22] T. Zhang et al., “BERTScore: Evaluating Text Generation with BERT,” in *Proc. Int. Conf. Learn. Represent. (ICLR)*, 2020.
- [23] N. Yadav and D. Gopinathan, “Semantic Search and Retrieval-Augmented Generation for Academic Literature Analysis Using ColBERT,” *Adv. Data Sci. Adapt. Anal.*, 2025, doi: 10.1142/s2424922x25400017.
- [24] B. Lund, “Prospects of Retrieval Augmented Generation (RAG) for Academic Library Search and Retrieval,” *SSRN Electron. J.*, 2025, doi: 10.2139/ssrn.5295044.
- [25] J. Nan, “Z-SPACE: A Multi-Agent Tool Orchestration Framework for Enterprise-Grade LLM Automation,” *SSRN Electron. J.*, 2025, doi: 10.2139/ssrn.5896270.
- [26] B. Li, J. Conen, and F. Aller, “AID-Agent: An LLM-Agent for Advanced Extraction and Integration of Documents,” in *Proc. 1st Workshop Res. Agent Lang. Models (REALM)*, 2025, pp. 80–88.
- [27] W. Hu et al., “Removal of Hallucination on Hallucination: Debate-Augmented RAG,” in *Proc. 63rd Annu. Meeting Assoc. Comput. Linguist. (ACL)*, 2025, pp. 15839–15853.
- [28] A. Tay, “When is a Hallucination Not a Hallucination? The Role of Implicit Knowledge in RAG,” *Rogue Scholar*, 2025, doi: 10.59350/2kf7k-r1m79.

- [29] Y. Zhao, Z. Liu, Y. Zheng, and K.-Y. Lam, "Attribution Techniques for Mitigating Hallucination in RAG-based Question-Answering Systems: A Survey," TechRxiv, 2025, doi: 10.36227/techrxiv.175036754.46793269/v1.

Sentiment Analysis for Delivery Pricing Optimization : Naïve Bayes Evaluation on Gojek Reviews

Nisa Hanum Harani^{*1}, Woro Isti Rahayu², Ruth Diana Purnamasari³

^{1,2,3}Universitas Logistik Bisnis Intrnasional, Indonesia

E-mail: ¹nisa@ulbi.ac.id, ²woro_isti_rahayu@ulbi.ac.id, ³ruthdianaps@gmail.com

Abstract

The development of digital technology in Indonesia has driven the growth of application-based transportation services such as Gojek. However, amid this growth, user dissatisfaction with pricing strategies often appears in social media reviews, but has not been systematically explored as a basis for optimizing delivery pricing using an automated text mining approach. Nevertheless, empirical studies that explicitly utilize sentiment analysis as a foundation for delivery pricing optimization remain limited. Existing sentiment analysis studies on online transportation services have mostly focused on overall service satisfaction, with limited attention to how sentiment patterns reflect customer price sensitivity. This study aims to analyze user sentiment as a diagnostic basis for evaluating and supporting delivery pricing optimization strategies. A comparative evaluation of four text representation methods: TF-IDF, Bag-of-Words (BoW), Word2Vec, and TF-IDF with Chi-Squared feature selection was conducted using 282 Gojek-related reviews collated from X (Twitter). The results of the experiment show that the Naïve Bayes model with BoW representation achieved the best performance, with 70% accuracy and an F1-score of 0.68, outperforming more complex approaches. Further analysis identifies that negative sentiment correlates with the keywords “expensive” and “promo,” providing a strong indication that pricing issues are a major point of concern for customers. These findings serve as a basis for management to optimize dynamic pricing strategies based on real-time user feedback.

Keywords—Sentiment Analysis, Naïve Bayes, Bag-of-Words, TF-IDF, Word2Vec

Submission: November 21, 2025

Accepted: December 29, 2025

Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.520>

1. INTRODUCTION

In the current digital era, rapid information technology development has driven digitization in various sectors, including the modern technology-based transportation sector [1]. In Indonesia, the presence of transportation e-commerce applications such as Gojek has become an important part of people's lives. These applications offer a range of services, from transportation booking, goods delivery, online food ordering, to digital payments, which have made the land transportation industry more open and highly competitive in attracting consumers [2].

As the number of users of the application increases, the number of reviews posted on social media also increases [3]. These reviews contain a range of responses,

both positive and negative, as well as neutral ones, reflecting user satisfaction and dissatisfaction with the services received. This review data is very valuable, not only for companies to evaluate service quality, but also for potential users as this review data can be used as a consideration [3][4]. However, the large amount of review data makes manual analysis ineffective, requiring automated methods such as sentiment analysis to process and understand public opinion from the large number of reviews [5]. Although many previous studies have applied sentiment analysis to evaluate overall service satisfaction in online transportation services, most of these studies focus on general services quality indicators and do not explicitly analyze user sentiment related to price perception, especially in the

online transportation ecosystem in Indonesia. In addition, comparative studies evaluating how different text representation methods capture signals of price dissatisfaction from social media data dominated by informal language and slang are still limited. This condition indicates a research gap between sentiment analysis methodology and its practical application in delivery price optimization. In this context, sentiment analysis serves not only as a tool to assess user satisfaction but also as an indicator of customer price sensitivity. By identifying patterns of price-related complaints in user reviews, sentiment analysis can provide empirical support for evaluating and optimizing delivery pricing strategies in competitive app-based transportation services.

Previous research in the field of text mining shows that classification algorithms such as Naïve Bayes and Support Vector Machine (SVM) have been widely applied. One of the challenges in the data, where models tend to provide higher accuracy results for the majority class compared to the minority class [6]. Several recent case studies have also begun to switch to using Deep Learning methods [7][8]. However, comparative evaluations that systematically assess the effectiveness of various text representation techniques particularly TF-IDF, Bag-of-Words (BoW), Word2Vec, and Chi-Squared feature selection within a single experimental framework for Indonesian language online transportation reviews are still scarce.

Therefore, this study aims to analyze the performance of the Naïve Bayes algorithm in classifying the sentiment of Gojek user reviews by comparing four different text representation methods. The primary objective is to identify the most effective representation for a limited but realistic Indonesian social media dataset.

The novelty of this research lies in its explicit positioning of sentiment analysis as a diagnostic tool for delivery pricing optimization, rather than merely as an assessment of overall service satisfaction. In addition, this study provides a comparative evaluation of four text representation methods to identify the most effective

approach for a limited but realistic Indonesian social media dataset. Understanding the classification and information mining processes is crucial for identifying meaningful trends from user review data [9]. To ensure that the data development process is valid and standardized, this study applies the CRISP-DM (Cross-Industry Standard Process for Data Mining) framework, which is still recognized as the de facto standard in various modern data science projects [10]

2. RESEARCH METHOD

This research was conducted through a series of systematic stages that adapted the CRISP-DM (Cross-Industry Standard Process for Data Mining) framework. CRISP-DM was chosen because its methodology has been proven to be the de facto standard in modern data science project management for more than two decades [16]. In addition, its flexible nature allows this model to be developed into a process-based assessment approach (process mining) [11], and comprehensively adapted for various data engineering and technical needs across multiple application domains [12]. The entire research process, from data collection to model evaluation. The general research process can be seen in the figure 1.

2.1. Data Collection

The dataset used in this study was obtained via web scraping from social media platform X (Twitter) using the snsrape library. The keywords used were “Gojek” and “Gojek app.” Data collection was conducted between May 21, 2025, and July 2, 2025. From this process, a total of 282 relevant review tweets were collected. This data was then saved in CSV format for further processing.

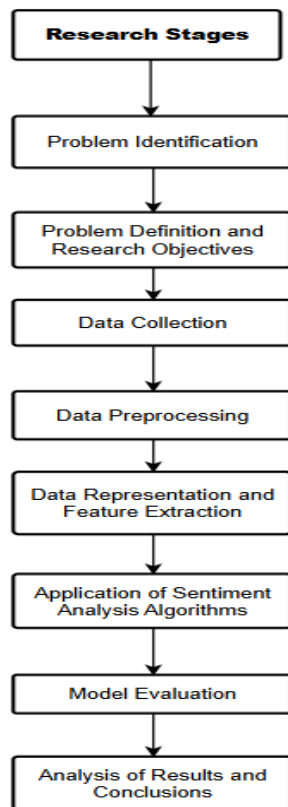


Figure 1 Research Stage

2.2. Dataset Description

The dataset consists of 282 Indonesian-language tweets related to Gojek services collected from the X (Twitter) platform between May 21, 2025, and July 2, 2025, using the keywords “Gojek” and “Gojek app”. The tweets represent short and informal user-generated content commonly found on social media.

Although the dataset size is relatively small compared to large scale sentiment analysis studies, it reflects a realistic and practical scenario frequently encountered in early stage analytical studies and internal business evaluations. The limited number of tweets may affect the generalizability and external validity of the findings, particularly in capturing rare or nuanced sentiment patterns. Therefore, the result of this study are intended to provide exploratory and comparative insights into model performance rather than definitive conclusions.

For model evaluation purposes, the dataset was divided into training and testing sets using an 80:20 split ratio with stratified sampling, ensuring that the original sentiment distribution was preserved across both subsets. Further details regarding the evaluation metrics and experimental setup are described in the model evaluation stage.

To improve clarity and transparency, Table 1 presents sample examples of tweets before and after the preprocessing stage.

Table 1 Sample of Text Preprocessing Results

Original Tweet	Preprocessed Text
Naik Gojek gak harus mahal! Pakai #GojekHemat dapetin DISKON tiap hari mulai dari Rp5.000-an aja! Tetep nyaman. Tetep aman.	naik gojek gak mahal pakai dapetin diskon tiap hari mulai rpan aja tetep nyaman tetep aman tetep irit cek promonya aplikasi sekarang
Tadi dapet abang Gojek Tuna Rungu beliau info by chat dan terpampang Namanya di Aplikasi. Luckily	tadi dapet abang gojek tuna rungku beliau info chat pampang nama aplikasi luckily
Tiba-tiba logout sendiri dari aplikasi gojek dan gak bisa login karena technical issue. Ada yang mengalami juga gaaak?	tibatiba logout sendiri aplikasi gojek gak login technical issue alami gaaak
gara-gara hujan tp pengen makan akhirnya gofood aja pake promo buat gofood. jadi ngga boros juga krn dapet diskon dari GOTEDXPU	garagara hujan ken makan akhir gofood aja pake promo buat gofood jadi ngga boros krn dapet diskon gotedxpu

2.3. Data Preprocessing

Before text data is processed or handled by algorithms, first perform the preprocessing stage to remove unnecessary elements and adjust the text format for consistency. This process involves several steps, including :

1. *Case Folding*: Change all letters to lowercase.

2. *Cleansing*: Removing unnecessary elements such as numbers, punctuation marks, URLs, and hashtags.
3. *Tokenizing*: Breaking down sentences into their constituent words.
4. *Stopword Removal*: Removing common words that have no sentimental meaning using the Indonesian NLTK Library.
5. *Stemming*: Reducing inflected words to their root form.

2.4. Automatic Sentiment Labeling

Sentiment labels were generated using a rule-based lexicon approach by identifying predefined positive and negative keywords commonly used in Indonesian social media texts. Reviews containing positive keywords were labeled as positive, reviews containing negative keywords were labeled as negative, while the remaining reviews were classified as neutral. This approach was applied to provide an efficient initial labeling mechanism for a limited dataset dominated by informal and short text. To reduce labeling noise, a subset of the automatically labeled data was manually reviewed to ensure consistency between the assigned labels and the expressed sentiment.

2.5. Feature Extraction

The text that has gone through the cleaning stage will then be converted into numerical format so that it can be processed by the algorithm. In this study, three main approaches were used to form this representation, including :

1. *TF-IDF* : This is a method used to calculate word weight and as the main representation because it is able to emphasize words that have significant meanings in a document. This method is commonly used in data mining to highlight specific words [13].
2. *Bag of Words (BoW)* : used as a comparison by only considering the frequency of word occurrence without considering the global context [14].

3. *Word2Vec* : using a pre-trained word embedding model (100-dimensional GloVe) in which each sentence is represented as the average vector of its constituent words [15].

In addition, this study also applied the Chi-Squared feature selection method on TF-IDF representation to select the 1000 most relevant features and reduce the data dimension. The application of algorithms in data mining, such as Naïve Bayes, is highly dependent on the selection of appropriate features to produce optimal accuracy [15].

2.6. Naïve Bayes Algorithm

In this study, the main algorithm used is Naïve Bayes, which is a probabilistic classification method based on Bayes' theorem. The flowchart of the Naïve Bayes process used in this study can be seen in the figure 2.

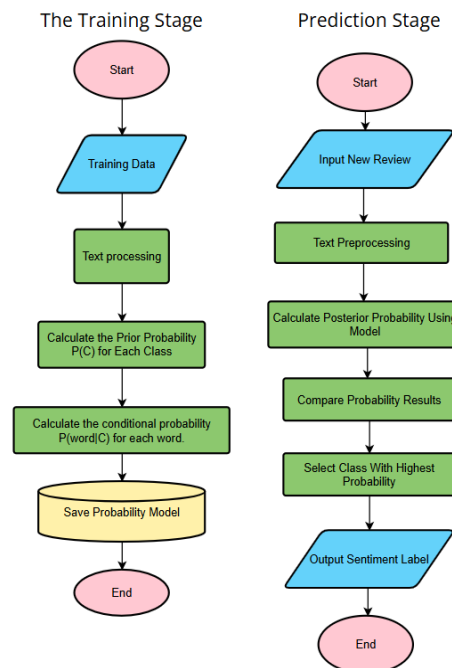


Figure 2 Flowchart

The basic formula of Bayes' theorem for classification can be written as follows :

$$P(C|X) = \frac{(P(X|C) \cdot P(C))}{P(X)} \quad (1)$$

Where $P(C|X)$ is the posterior probability of class C given feature X , $P(X|C)$ is the likelihood, $P(C)$ is the prior probability, and

$P(X)$ is the marginal probability of the observed data. In the context of text classification, X is a document consisting of words $x_1, x_2, \dots, x_n$. Assuming feature independence, the formula for predicting the class of a document is simplified to

$$C_{Prediksi} = \arg P(C) \prod_{i=1}^n P(x_i|C)$$

Where $C_{textprediksi}$ is the class selected because it has the highest probability value

Table 2 Algorithm Naïve Bayes Classification

Algorithm 1 : Naïve Bayes Classification

Output : prediction class C_{map}

- 1 *TRAIN* : Count $P(c)$ for each class $c \in C$
- 2 *TRAIN* : Count $P(w|c)$ for every word w in the vocabulary
- 3 *TEST* : for each test document d
- 4 $C_{map} = \arg \max_{c \in C} [P(c) + \log P(w|c)]$
- 5 *RETURN* C_{map}

2.7. Data Splitting and Model Evaluation

The dataset was divided into training and testing sets with an 80:20 split ratio using stratified sampling to maintain the original sentiment distribution. Model performance was evaluated using accuracy, precision, recall, F1-score, and confusion matrix analysis.

2.8. Class Imbalance Consideration

The sentiment dataset shows class imbalance, with neutral sentiment as the majority class. Resampling techniques were not applied in this study in order to preserve the original data distribution. The impact of class imbalance on classification performance is discussed as a limitation in this study.

2.9. Implementation Environment

All experiments were implemented using *Python* on *Google Colab*. *Scikit-learn* was used for classification, *NLTK*

and *Sastrawi* for text preprocessing, and *Gensim* for *Word2Vec* modeling.

3. RESEARCH RESULTS

The results of this study present empirical data obtained from a series of experiments, without accompanying opinions or in-depth interpretations, in accordance with the evaluation stage of the CRISP-DM methodology.

3.1. Sentiment Distribution and Visualization

A total of 282 reviews were collected and processed through preprocessing and automatic sentiment labeling based on a predefined lexicon. The sentiment distribution results shown in Figure 3 indicate that the Neutral class is the largest category (49.3%), followed by the Positive class (32.3%) and the Negative class (18.4%).

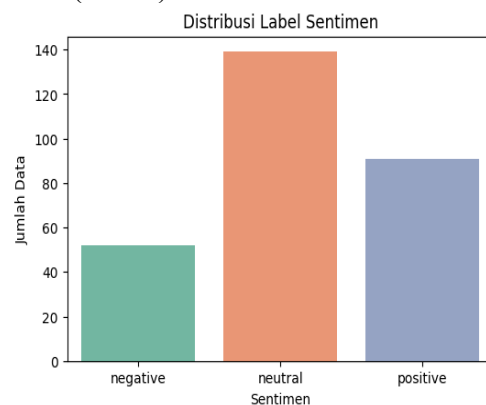


Figure 3 Sentiment Label Distribution

The most frequently appearing words in each sentiment group are visualized in the form of word clouds, as shown in Figure 4, Figure 5, and Figure 6.



Figure 4 Positive Sentiment Word Cloud



Figure 5 Neutral Sentiment Word Cloud



Figure 6 Negative Sentiment Word Cloud

Figures 4–6 illustrate dominant word patterns across sentiment classes. Positive reviews are characterized by usability-related terms, neutral reviews contain informational expressions, while negative reviews frequently include pricing-related keywords such as “mahal” and “promo,” indicating user sensitivity toward delivery costs.

3.2. Classification Model Performance

Testing of the Naïve Bayes algorithm model was conducted by applying four types of text representation. The evaluation results, which include Accuracy, Precision, Recall, and F1-score values for each scenario, can be seen in the table 3.

Table 3 Accuracy Results Table

Model	Acc.	Prec.	Rec.	F1
Naive Bayes + TF-IDF	0.56	0.77	0.56	0.46
Naive Bayes + BoW	0.70	0.72	0.70	0.68
Naive Bayes + Word2Vec	0.51	0.57	0.51	0.52
Naive Bayes + TF-IDF+Chi2	0.53	0.57	0.53	0.39

Based on the results, the combination of the Naïve Bayes algorithm with Bag-of-Words (BoW) representation

achieved the best overall performance, with an accuracy of 0.70 and an F1-score of 0.68. This finding indicates that a simple word frequency-based representation is more stable and effective for short and informal social media text compared to more complex approaches. The TF-IDF representation produced lower performance, with an accuracy of 0.56 and an F1-score of 0.46, which may be caused by uneven word distribution and the dominance of certain sentiment classes in the dataset.

The Word2Vec-based model achieved an accuracy of 0.51 and an F1-score of 0.52, suggesting that the sentence-level vector averaging process led to the loss of important contextual information in short texts. Similarly, the application of Chi-Squared feature selection on TF-IDF further reduced model performance, resulting in an accuracy of 0.53 and an F1-score of 0.39, indicating that dimensionality reduction removed features that were still relevant for sentiment classification.

Overall, the BoW + Naïve Bayes model demonstrates the most balanced and stable performance among all evaluated methods. A more detailed analysis of classification errors and class-specific prediction behavior is further discussed through confusion matrix analysis in Figure 7.

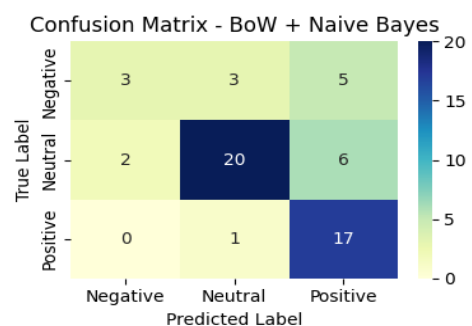


Figure 7 Confusion Matrix BoW+Naive Bayes

In addition to the BoW model, Confusion Matrix analysis was also performed on three other representation scenarios, namely TF-IDF, Word2Vec, and TF-IDF with Chi-Square features. The evaluation results of these scenarios show different performance

patterns for each model. The following is an explanation and visualization of the three Confusion Matrix analysis model scenarios :

1. Model *TF-IDF + Naïve Bayes*

This model performs well in identifying the neutral class, but shows poor discrimination for positive and negative sentiments, with most non-neutral instances misclassified as neutral. This indicates a strong majority-class bias and limited capability in distinguishing sentiment polarity in short and informal texts. Here is the visualization.

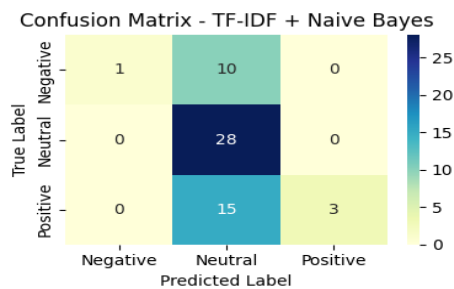


Figure 8 Confusion Matrix TF-IDF+Naive Bayes

2. Model *Word2Vec + Naïve Bayes*

The Word2Vec-based model shows a more dispersed and inconsistent prediction pattern across sentiment classes. This indicates that sentence-level vector averaging is less effective in preserving sentiment polarity in short and informal Indonesian social media texts, resulting in frequent misclassification between neutral, positive, and negative classes, as illustrated in the confusion matrix

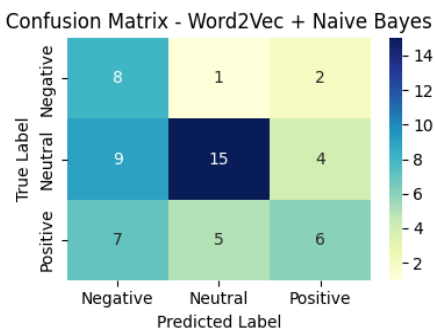


Figure 9 Confusion Matrix Word2Vec+Naive Bayes

3. Model *TF-IDF + Chi-Squared + Naïve Bayes*

This model shows good and perfect accuracy for the neutral class, but fails to recognise the other two classes. All 11 negative data points were incorrectly predicted as neutral, and of the 18 positive data points, only 2 were classified correctly. This shows that the Chi-Squared feature selection process led the model to lose several important features, resulting in a strong bias towards the majority class. The following is a visualisation of this

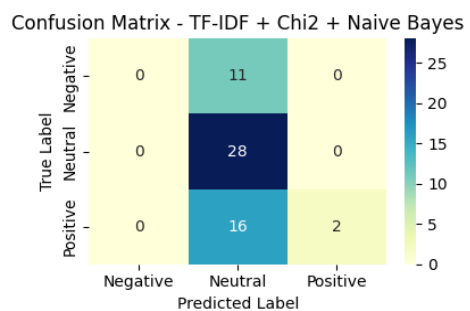


Figure 10 Confusion Matrix TF-IDF + Chi-Squared + Naïve Bayes

Experiments with these three models reinforce the finding that BoW is the most accurate representation balanced and stable for this dataset, while TF-IDF, Word2Vec, and TF-IDF+Chi2 show a tendency toward the neutral class and weakness in handling sentiment class imbalance.

3.3. Confusion Matrix Analysis and Baseline Comparison

To further analyze the classification behavior of each model, confusion matrix analysis was conducted for all text representation scenarios. This analysis provides a more detailed understanding of how each sentiment class is predicted and where misclassifications occur. The confusion matrix summary and baseline comparison were derived analytically from the confusion matrix visualization and sentiment distribution results, respectively.

3.3.1. Confusion Matrix Summary for BoW + Naïve Bayes

Based on the experimental result, the *Bag-of-Words (BoW)* representation combined with the *Naïve Bayes* classifier achieved the best overall performance. The confusion matrix for this model is presented in Figure 7, while a summarized numerical interpretation is shown in table 4.

Table 4 Confusion Matrix Summary for BoW + Naïve Bayes

Sentiment Class	True Positive (TP)	False Positive (FP)	False Negative (FN)
Negative	3	2	8
Neutral	20	4	8
Positive	17	11	1

The result indicate that the model performs well in identifying positive sentiment, with 17 out of 18 instances correctly classified and only one instance misclassified as neutral. The neutral class is also reasonably well recognized, with 20 correct predictions out of 28 instances. However, the negative class shows weaker performance, as only 3 out of 11 negative reviews were correctly classified, while the remaining instances were mostly misclassified as neutral or positive. This patten reflects the impact of class imbalance and overlapping sentiment expressions commonly found in short informal texts.

Overall, the BoW + Naïve Bayes model demonstrates a balanced and stable classification behavior compared to other representations, particularly in recognizing dominant sentiment classes.

3.3.2. Comparison with Baseline Classifier

To assess whether the proposed model provides meaningful improvement, a baseline comparison was conducted using a majority class classifier, where all instances are predicted as the neutral class. Given that neutral sentiment accounts for approximately 49.3% of the dataset, the baseline accuracy is 0.49.

Table 5 Performance Comparison with Baseline

Model	Accuracy
-------	----------

Majority Class Baseline	0.49.
Naïve Bayes + Bow	0.70

Compared to basic classification, the BoW + Naïve Bayes model achieved a significantly higher accuracy of 0.70, indicating that the proposed approach is capable of capturing meaningful sentiment patterns beyond simple classification. This improvement proves that sentiment classification results are not only influenced by class distribution, but also by the effectiveness of text representation and probabilistic learning approaches.

3.3.3. Comparison Across Text Representation Methods

Further confusion matrix analyses were also conducted for the TF-IDF, Word2Vec, and TF-IDF with Chi-Squared feature selection models, as illustrated in Figures 8–10. The TF-IDF-based model shows a strong bias toward the neutral class, successfully classifying all neutral instances but failing to distinguish positive and negative sentiments. The Word2Vec-based model exhibits more dispersed predictions, indicating difficulty in capturing sentiment polarity using averaged word embeddings for short informal texts. Meanwhile, the TF-IDF + Chi-Squared model further amplifies majority class bias, resulting in severe misclassification of minority sentiment classes.

These results reinforce that, for a limited and informal Indonesian social media dataset, the BoW representation provides the most stable and balanced performance, outperforming more complex feature representations in terms of overall classification reliability.

4. DISCUSSION

This section explains or elaborates on an in-depth analysis of the experimental

results described earlier, highlighting the effectiveness of feature representation, model performance patterns, and characteristics of Gojek user review data that influence classification results in relation to the research objective of evaluating sentiment-based insights for delivery pricing optimization.

4.1. Analysis of Feature Representation Effectiveness on Model Performance

The results in Table 2 show that the combination of Naïve Bayes with Bag-of-Words (BoW) representation is the most effective approach for the dataset in this study. With an accuracy of 0.70 and an F1-score of 0.68, BoW provides the best prediction stability compared to other methods. These findings indicate that for small-scale review data consisting of short text, raw word frequencies are sufficient to describe sentiment patterns without requiring complex weights or embeddings.

Theoretically, this result aligns with the probabilistic assumption of Naïve Bayes, which benefits from discrete and independent feature distributions such as term frequencies. Empirically, the dominance of repeated sentiment-related keywords in short Gojek reviews strengthens BoW effectiveness, as these words directly influence posterior probability estimation.

Conversely, the TF-IDF representation resulted in lower performance (accuracy 0.56), likely because the inverse document frequency weighting reduced the importance of frequently occurring words that still carry sentiment in informal Gojek reviews. The Word2Vec-based approach also showed limited effectiveness (accuracy 0.51), as sentence-level vector averaging led to contextual information loss, which is critical for distinguishing sentiment polarity in short texts. Similarly, the application of Chi-Squared feature selection on TF-IDF further decreased performance (accuracy 0.53), indicating that aggressive feature reduction removed relevant sentiment-related terms, particularly in a small-scale dataset.

4.2. Confusion Matrix Analysis and Class Imbalance

Evaluation using the Confusion Matrix (Figure 7) shows that the model tends to be more accurate in identifying majority classes, especially positive and neutral ones. This can be seen from :

- **Neutral:** 20 out of 28 data points were classified correctly, while the rest were incorrectly mapped to positive and negative.
- **Positive:** relatively stable performance, with 17 out of 18 data classified correctly.
- **Negative:** is the most challenging class to recognise, because out of a total of 11 data points, only three were correctly classified, while the rest were incorrectly distributed to either the positive or neutral classes.

These results confirm that class imbalance (imbalanced data) has a significant impact on the classification process. Models tend to be safer by predicting the class that is far more dominant in the training data, a phenomenon commonly reported in sentiment analysis studies involving social media data with uneven class distributions.

4.3. The Challenge of Non-Standard Indonesian in Sentiment Analysis

Based on the word patterns in the Word Cloud, many reviews contain non-standard terms such as “gak,” “yg,” “pastiadadijalan,” or repeated words that reflect the user's emotions. This variation makes it difficult for TF-IDF and Word2Vec representations to accurately capture sentiment meaning. Pre-trained embeddings based on standard language corpora often fail to recognize slang, mixed languages, or short sentence structures commonly found on social media platforms.

This condition explains why simple methods such as BoW are more stable or consistent than methods that rely on semantic distribution or weights between documents, particularly in informal Indonesian social media contexts.

4.4. Consistency of Findings with Previous Research

The findings related to BoW dominance and the importance of preprocessing are consistent with the results described by Sentiana's research[13], which shows that text cleaning quality has a significant effect on classification performance. Furthermore, the model performance patterns in this study are in line with the principles described by Senubekti and Dewi[14], which state that feature representation selection must consider the available data distribution, especially when the dataset is small and unbalanced.

Beyond consistency, this study extends prior work by empirically demonstrating that simpler representations may outperform more complex embeddings when applied to real-world, short, and informal Indonesian user reviews.

4.5. Implications of Findings for System Development

The dominance of neutral sentiment and the pattern of complaints found in the negative Word Cloud indicate that most users are neither completely satisfied nor dissatisfied, but remain sensitive to specific issues. Notably, several negative sentiment keywords are related to pricing perceptions, such as complaints about cost, promotions, or value received.

From a practical standpoint, the BoW + Naïve Bayes-based automated model can serve as an initial monitoring system to detect recurring pricing-related complaints and sentiment shifts in user reviews. Such insights support delivery pricing evaluation by identifying user sensitivity to price changes, promotional effectiveness, and perceived fairness of delivery fees. This demonstrates that sentiment analysis can be operationalized as an early warning system for pricing dissatisfaction, enabling Gojek to adjust delivery pricing strategies based on real user feedback rather than manual review alone.

4.6. Direction for Further Research Development

Based on the limitations identified, further improvements could focus on:

- Handling class imbalance using SMOTE, class weighting, or data augmentation
- The use of Indonesian language Transformer-based models, such as IndoBERT, to handle more complex semantic contexts
- The development of text representation capable of recognizing non-standard Indonesian language forms, including slang, abbreviations, and spelling variations.

This approach is expected to improve the model's weakness in recognizing minority classes and improve the overall quality of predictions.

5. CONCLUSION

This study evaluated the performance of the Naïve Bayes algorithm for sentiment classification of Indonesian-language Gojek user reviews using the CRISP-DM framework. Experiments on 282 tweets from the X (Twitter) platform demonstrate that text representation plays a critical role in classification effectiveness, particularly for short and informal social media texts.

Among the evaluated approaches, the Bag-of-Words (BoW) representation achieved the best performance, with an accuracy of 70% and an F1-score of 0.68, outperforming TF-IDF, Word2Vec, and TF-IDF with Chi-Squared feature selection. This confirms that simple frequency-based representations can be more effective than complex models when applied to small and imbalanced datasets.

From a scientific perspective, this study contributes empirical evidence that lightweight probabilistic models remain competitive in real-world Indonesian social media contexts. Practically, the identified pricing-related sentiment patterns indicate that sentiment analysis can be used as an early analytical tool to support delivery pricing evaluation by capturing customer

price sensitivity and promotional effectiveness, thereby supporting more data-driven delivery pricing optimization strategies.

6. RECOMMENDATION

Based on the findings and limitations identified in this study, several recommendations are proposed to enhance the accuracy, robustness, and practical relevance of future sentiment analysis research :

1. **Class Imbalance Handling :** Further research should apply class imbalance handling techniques such as class weighting or SMOTE to improve the detection of minority sentiment classes. This is particularly important because pricing-related complaints are commonly expressed in negative sentiment, which plays a critical role in delivery pricing evaluation.
2. **Dataset Expansion and Multi-Sources:** Further studies should expand the dataset by collecting reviews from additional platforms such as Google Play Store and App Store. Reviews from these platforms are generally longer and more detailed, allowing more explicit identification of pricing dissatisfaction and promotional effectiveness.
3. **Advanced and Context-Aware Models:** The use of more advanced models, including Support Vector Machine (SVM) and transformer-based models such as IndoBERT, is recommended to better capture contextual and implicit pricing-related sentiment in informal Indonesian texts, especially when combined with balanced data strategies.
4. **Aspect-Based Sentiment Analysis for pricing Optimization :** Future research should implement aspect-based sentiment analysis by focusing on specific pricing-related aspects, such as delivery cost, promotion availability, perceived fairness, and value-for-money. This approach would enable sentiment analysis results to be directly

translated into actionable insights for delivery pricing optimization.

REFERENSI

- [1] V. Baskov and E. Isaeva, "Information Support of Transportation Process Efficiency of the," vol. 01035, pp. 1–8, 2021.
- [2] S. K. Djawa, "The Impact of Information & Communication Technology on Land Transportation Service Business on Indonesia (Case Study in Central Sulawesi Province)," *J. Soc. Res.*, vol. 2, no. 5, pp. 1660–1665, 2023, doi: 10.55324/josr.v2i5.853.
- [3] D. Lo and M. R. Lyu, *TOUR : Dynamic Topic and Sentiment Analysis of User Reviews for Assisting App Release*, vol. 1, no. 1. Association for Computing Machinery, 2021. doi: 10.1145/3442442.3458612.
- [4] E. L. Funnell, B. Spadaro, N. Martinkey, T. Metcalfe, S. Bahn, and S. Bahn, "mHealth Solutions for Mental Health Screening and Diagnosis : A Review of App User Perspectives Using Sentiment and Thematic Analysis," vol. 13, no. April, pp. 1–17, 2022, doi: 10.3389/fpsy.2022.857304.
- [5] D. Singh, "Analysis Of Public Sentiment of Covid-19 Pandemic , Vaccines , And Lockdowns Presented to," 2022.
- [6] E. Setiana, Marwondo, Venia Retreva Danestiara, and Wiyanudin, "Analisis Sentimen Pelaksanaan Kuliah Online Menggunakan Algoritma Support Vector Machine," *Nuansa Inform.*, vol. 17, no. 2, pp. 66–70, 2023, doi: 10.25134/ilkom.v17i2.11.
- [7] P. R. Saragih, Y. Sibaroni, and S. S. Prasetyowati, "Multi-Aspect Sentiment Analysis on Gojek Application Reviews Using CNN-LSTM Method," vol. 7, no. 1, pp. 2684–2685, 2025, doi: 10.47065/bits.v9i9.999.
- [8] A. R. Sembiring and C. K. Dewa, "Sentiment Analysis On Indonesian Tweets about the 2024 Election," *Sinkron*, vol. 9, no. 1, pp. 413–422, 2025, doi: 10.33395/sinkron.v9i1.14481.
- [9] M. A. Senubekti and L. A. Puspita Dewi, "Prinsip Klasifikasi Dan Data Mining Dengan Algoritma C4.5," *Nuansa Inform.*, vol. 16, no. 2, pp. 87–93, 2022, doi: 10.25134/nuansa.v16i2.5834.
- [10] F. Martinez-Plumed *et al.*, "CRISP-DM Twenty Years Later: From Data Mining

- Processes to Data Science Trajectories,” *IEEE Trans. Knowl. Data Eng.*, vol. 33, no. 8, pp. 3048–3061, 2021, doi: 10.1109/TKDE.2019.2962680.
- [11] R. H. Bemthuis, R. R. Govers, and A. Asadi, “A CRISP-DM-based methodology for assessing agent-based simulation models using process mining,” *J. Simul.*, 2025, doi: 10.1080/17477778.2025.2508245.
- [12] S. Huber, H. Wiemer, D. Schneider, and S. Ihlenfeldt, “DMME: Data mining methodology for engineering applications - A holistic extension to the CRISP-DM model,” *Procedia CIRP*, vol. 79, no. July 2018, pp. 403–408, 2019, doi: 10.1016/j.procir.2019.02.106.
- [13] D. Papakyriakou and I. S. Barbounakis, “Data Mining Methods: A Review,” *Int. J. Comput. Appl.*, vol. 183, no. 48, pp. 5–19, 2022, doi: 10.5120/ijca2022921884.
- [14] L. Efrizoni, S. Defit, M. Tajuddin, and A. Anggrawan, “Komparasi Ekstraksi Fitur dalam Klasifikasi Teks Multilabel Menggunakan Algoritma Machine Learning,” *MATRIK J. Manajemen, Tek. Inform. dan Rekayasa Komput.*, vol. 21, no. 3, pp. 653–666, 2022, doi: 10.30812/matrik.v21i3.1851.
- [15] B. Budiman, “Perbandingan Algoritma Klasifikasi Data Mining untuk Penelusuran Minat Calon Mahasiswa Baru,” *Nuansa Inform.*, vol. 15, no. 2, pp. 37–52, 2021, doi: 10.25134/nuansa.v15i2.4162.

SI-PANDA: Web-Based Personnel Data and Dossier Management System for Pusinfoლაhta TNI

Muhammad Sulthan Nasyira¹, Eryan Ahmad Firdaus²

^{1,2}Universitas Pertahanan, Indonesia

E-mail: ¹sulthanasyirah@gmail.com, ²eryan.firdaus@idu.ac.id

Abstract

Pusinfoლაhta TNI currently faces critical inefficiencies in personnel administration due to reliance on manual spreadsheet calculations and decentralized local storage for dossier documents. This manual workflow poses significant risks to data security, accuracy, and retrieval speed. To address these limitations, this research developed the Personnel Data and Dossier Processing System (SI-PANDA) utilizing the Rapid Application Development (RAD) method to ensure iterative improvements based on user feedback. The system is constructed upon a three-tier architecture comprising a Native PHP backend, MySQL database, and a responsive web interface. Testing results demonstrate significant performance improvements: the system achieves a page load time of 1.2 seconds and near-instant calculation speed (0.05 seconds) for tenure and retirement estimation, ensuring 100% calculation accuracy. Furthermore, Blackbox testing confirmed the validity of all functional requirements, including the automated backup mechanism and search filters. The main contribution of this study is the delivery of a secure, integrated digital framework that successfully transitions Pusinfoლაhta TNI from manual processing to a reliable web-based environment, specifically adapted to meet military administrative standards.

Keywords— personnel information system, dossier management, RAD, PHP, data security

Submission: November 25, 2025

Accepted: December 23, 2025

Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.522>

1. INTRODUCTION

The TNI Information and Data Processing Center / Pusat Informasi dan Pengolahan Data TNI (Pusinfoლაhta TNI) is a Central Executive Agency at the TNI Headquarters (Mabes TNI) level, operating directly under the command of the TNI Commander. Its primary mission is to prepare, manage, and present information and data related to the development and deployment of TNI forces [1]. To facilitate and support effective administrative operations, Pusinfoლაhta TNI relies on a data processing and personnel dossier information system.

However, the current system has limitations regarding automatic data calculation and dossier document management. Consequently, data recording is predominantly performed using spreadsheet applications like Microsoft Excel, while dossier management remains locally stored via file explorer. This situation poses risks to data security and reliability, as well as limiting rapid access to personnel information when required.

Various related studies have been conducted to address personnel administration and document management issues. For instance,

recent research [2] has highlighted the urgency of implementing Electronic Document Management Systems (EDMS) within the E-Government environment to enhance public administration efficiency and auditability. Similarly, other studies [3] have focused on developing secured cloud-based document systems utilizing encrypted security mechanisms to prevent unauthorized access.

These studies predominantly focus on civil institutions or commercial enterprises and do not fully accommodate the complexity of military requirements. As identified in strategic studies for the Indonesian National Army (TNI) [4], the military environment requires specific and integrated information systems to support defense operations, which include precise assignment history calculations and strict data security.

The objective of this research is to design a new information system named SI-PANDA (Personnel Data and Dossier Processing System / Sistem Pengolahan Data dan Dosir Personel). This system incorporates several improvements over the previous system regarding data calculation and dossier management to ensure the processed information meets operational requirements. Furthermore, it

aims to enhance data security, reliability, and overall efficiency in data management.

Beyond practical benefits for the organization, the academic contribution of this research is a proposed information system architecture framework that integrates Digital Document Management Systems with multi-level data security protocols adapted for the military environment. This enriches the software engineering literature, particularly in the domain of Human Resource Information Systems (HRIS) within the defense sector, which demands high standards of reliability and confidentiality to mitigate cyber defense risks [5].

2. RESEARCH METHODOLOGY

This research applies the Rapid Application Development (RAD) method, selected for its capability to accelerate system development through rapid prototyping, continuous feedback, and iterative improvements [6]. The RAD phases are aligned with the practical work execution focuses on identifying system requirements and problem scope [7].

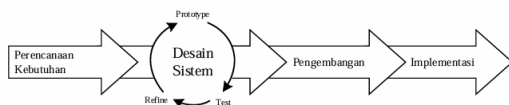


Image 1. Rapid Application Development

To ensure the system addresses the specific operational challenges at Pusinfolahta TNI, data collection was conducted through interviews and observations with Pusinfolahta TNI organic staff. The study analyzes current workflow deficiencies to ensure the new design meets organizational needs. The requirements analysis determined the system's core features, a web-based system with efficient dossier document management system, automatic data calculation, search filters, also data security and system reliability supported by authentication and an automated backup system.

This electronic dossier system contains personnel administration documents designed to support data management, facilitating integrated search, storage, and archive security within the TNI environment; specifically, scanned documents are grouped and assigned identification codes based on 32 data elements, categorized into Soldier (Prajurit) and Civil Servant (PNS) data [8]. The system requirements shown in Table 1.

Requirement	Description
Dossier Management	The system must allow uploading, viewing, and organizing scanned documents based on 32 personnel data elements (Prajurit & PNS).
Automatic Calculation	The system must automatically calculate tenure, retirement dates, and rank progression based on input data.
Search & Filtering	The system must provide advanced search capabilities by Name, NRP/NIP, Rank, or Unit.
CRUD Operations	The system must support Create, Read, Update, and Delete operations for personnel records.
Security	The system must implement authentication (Login/Logout) and role-based access control.
Reliability	The system must include an automated backup mechanism to prevent data loss.
Performance	The web interface must be responsive and load within 3 seconds.

Table 1. System Requirements

The User Design phase involves the iterative design of system architecture and user interfaces utilizing the Unified Modeling Language (UML). UML is a graphic-based language used to visualize, specify, construct, and document Object-Oriented software development systems [9], including Use Case Diagrams, Activity Diagrams, and Flowcharts alongside Entity-Relationship Diagrams (ERD) for relational database structuring [10].

The Construction phase follows, focusing on coding based on the approved design, employing HTML, CSS, and JavaScript for a responsive frontend, and native PHP for backend logic, CRUD operations, and file management without additional frameworks. The database is implemented using MySQL as the RDBMS, with XAMPP serving as the local server environment. Accordingly, this system is constructed using a three-tier architecture: 1. Frontend Layer (Presentation), The web-based user interface developed with HTML, CSS, and JavaScript. 2. Backend Layer (Logic), A native PHP server responsible for handling

authentication, CRUD operations, automatic calculations, and file management. 3. Database Layer (Data), Manages MySQL data storage alongside a structured file directory for dossier documents.

Finally, the Cutover phase encompasses system testing and implementation, where Blackbox Testing is applied to validate input and output functions without inspecting internal code structures to ensure the application fully meets Pusinfolahta TNI requirements prior to handover [11].

3. RESEARCH RESULT

3.1. System Architecture

System architecture refers to the fundamental organization of a system, embodied in its components, their relationships to each other, and the environment, often enabling the spatial and physical separation of databases, computational engines, and user interfaces to enhance flexibility and scalability [12]. Accordingly, this system is constructed using a three-tier architecture comprising a Frontend Layer, a Backend Layer, and a Database Layer. The Frontend Layer serves as the web-based user interface developed with HTML, CSS, and JavaScript, while the Backend Layer functions as a PHP server responsible for handling authentication, CRUD operations, file management, and the backup system. Furthermore, the Database Layer manages MySQL data storage alongside file storage for dossier documents.

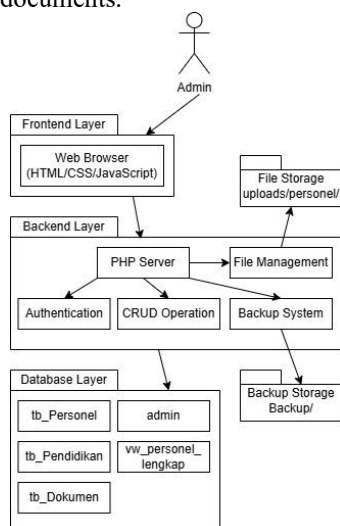


Image 2. System Architecture

3.2. Flowchart

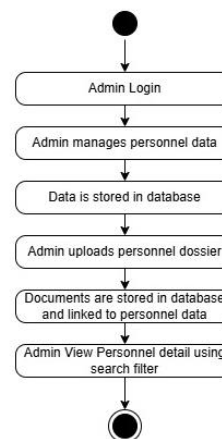


Image 3. Flowchart

3.3. Use Case Diagram

A Use Case Diagram is utilized to model the system's functional requirements by visualizing the interactions between users (actors) and the system capabilities [14]. In this context, the Admin possesses full access rights to manage all data and information contained within the system.

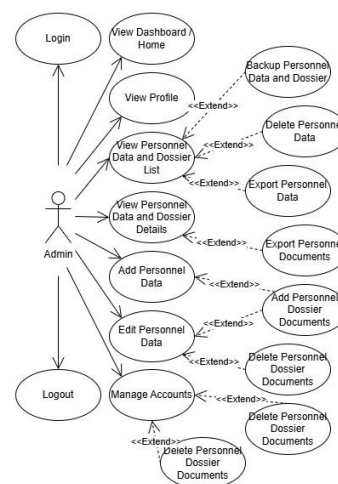


Image 4. Use Case Diagram

3.4. Entity Relationship Diagram

The Entity-Relationship Diagram (ERD) is designed to model the database structure by defining the entities, their attributes, and the relationships between data to ensure data integrity [15]. The tb_personel table stores the personal and administrative data of soldiers and

is linked via a one-to-many relationship with *tb_pendidikan* and *tb_dokumen*. Furthermore, a database View (*vw_personel_lengkap*) is implemented to perform real-time automatic calculations for age, length of service, and estimated retirement dates.

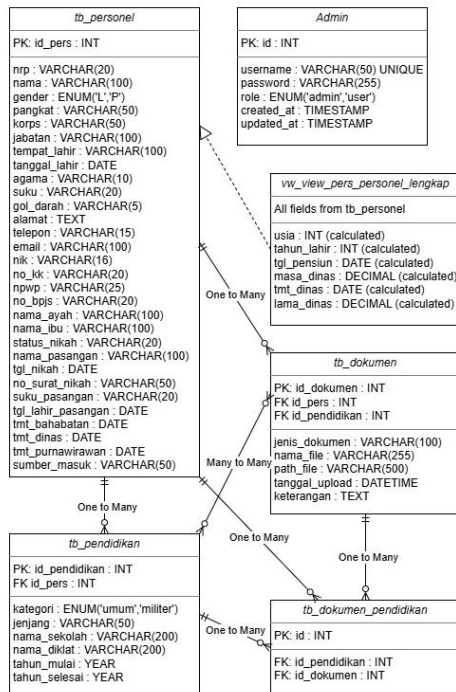


Image 5. Entity Relationship Diagram

3.5. Activity Diagram

Activity diagrams are utilized to model the workflow of a business process, illustrating the sequence of activities, decision points, and control flows within the system [16]. It comprises several key activities, including logging in, adding data, updating data, and data management.

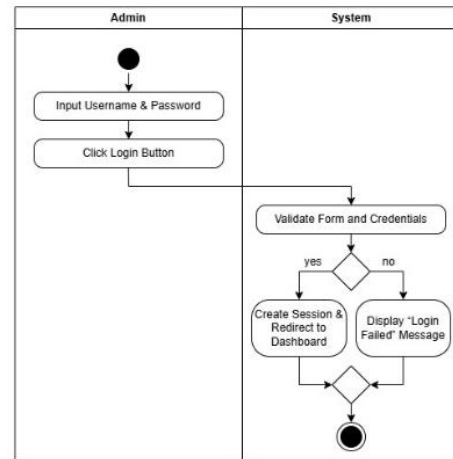


Image 6. Login Activity

3.5.2. Add New Data Activity

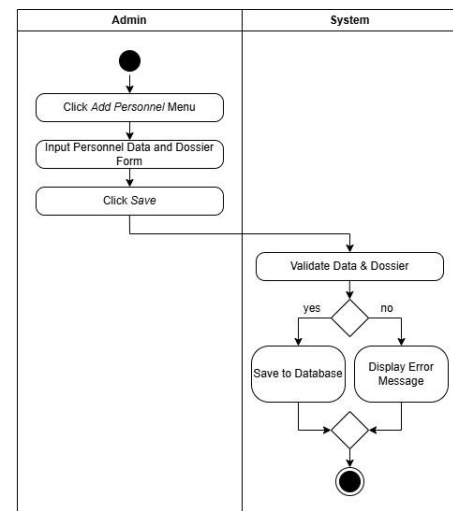


Image 7. Add New Data Activity

3.5.1. Login Activity

3.5.3. Update Data Activity

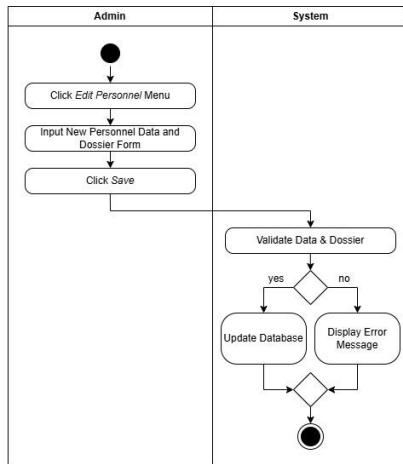


Image 8. Update Data Activity

3.5.4. Data Management Activity

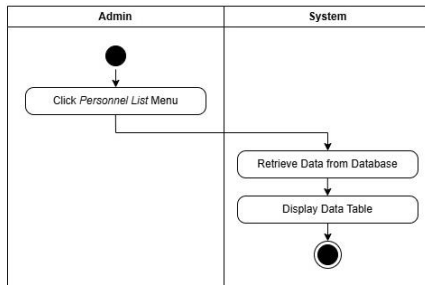


Image 9. Data Management Activity

3.6. System Interface

The system interface is designed to provide a user-friendly and visually coherent layout that effectively supports efficient personnel data processing and minimizes administrative errors [17].

3.6.1. Login Interface



Image 10. Login Interface

3.6.2. Add New Data Interface

Image 11. Add New Data Interface

3.6.3. Update Data Interface

Image 12. Update Data Interface

3.6.4. Data Management Interface

No	Nama	Jenis Kelamin	Tgl Lahir	Tempat Lahir	Agama	Status Pernikahan	Tipe Darah	Tensi Darah	Alergi	Riwayat Penyakit	Aksi
1	1194020201071	Wanita	11/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail
2	1194020201072	Pria	12/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail
3	1194020201073	Pria	13/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail
4	1194020201074	Pria	14/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail
5	1194020201075	Pria	15/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail
6	1194020201076	Pria	16/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail
7	1194020201077	Pria	17/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail
8	1194020201078	Pria	18/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail
9	1194020201079	Pria	19/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail
10	1194020201080	Pria	20/02/2001	Bandung	Islam	Belum Kawin	O	120/80	Belum Ada	Belum Ada	Detail

Image 13. Data Management Interface

3.7. Blackbox Testing

Blackbox testing was conducted on the key features, yielding the following results:

Feature	Test Scenario	Result
Login	Entering valid and invalid credentials.	Success (Validation active; session locked after 5 failed attempts).
Data Management	CRUD operations (Create, Read, Update, Delete) on personnel data.	Success (Data stored in DB; input validation active).
Dossier Management	Uploading PDF/Image	Success (Files stored in server)

	documents (<5MB).	directory and linked to DB).
Search Filters	Searching some data with filters.	Success (System showed related data).
Backup	System perform Automated Backup Data	Success (System Stores Backup Automatic)

Table 2. Blackbox Testing

3.8. Quantitative Performance Analysis

To ensure the system meets operational efficiency standards, quantitative performance testing was conducted focusing on page load speed and calculation accuracy in internal pusinfo lahta TNI server.

Metric	Measured	Analysis
Page Load Time	1.2 seconds	The lightweight native PHP architecture ensures rapid access even on standard internal networks.
Calculation Speed	0.05 seconds	The database View (vw_personel_lengkap) processes age and tenure calculations instantly.
Backup Speed	12 Seconds	Backup for 1000 records is efficient and does not disrupt server performance.

Table 3. quantitative performance testing

3.9. User Acceptance Test (UAT)

User Acceptance Testing (UAT) was conducted with 5 staff members from Pusinfo lahta TNI to measure usability and user satisfaction. The testing utilized a Likert scale (1-5) questionnaire.

Criterion	Average Score (1-5)	Interpretation
Ease of Use (Usability)	4.4	Users found the interface intuitive
Functional Suitability	5	Features match the daily administrative workflow.
Visual Clarity	4.2	Text and layout are

		clear and easy to read.
Overall Satisfaction	4.53	The system is ready for deployment.

Table 4. User Acceptance Test (UAT)

3.10. System Comparison

The comparison shows that the new personnel data and dossier processing system provides a more integrated, secure, and efficient web-based solution compared to the predominantly manual, device-dependent methods of the current system.

Comparison	Current	New
Platform	Tends to rely on desktop applications	Utilizes a web-based system.
Data Management	Tends to be Performed using desktop applications.	single database, data calculation and storage are automated.
Dossier Management	Tends to be stored locally	Stored within a document management system.
Data Security and Reliability	Security depends on user devices	Features authentication and automated backup.
Time Efficiency	Search and update processes are time-consuming	Data search and updates are performed automatically and rapidly.

Table 5. System Comparison

In a practical context, SI-PANDA directly resolves the critical field issue of administrative bottlenecks caused by manual data processing. By automating tenure and retirement calculations, the system eliminates human errors inherent in the legacy spreadsheet method, ensuring that decision-makers at Mabes TNI have immediate access to accurate and verified personnel data for strategic deployment.

4. CONCLUSION

Based on the design, implementation, and testing results utilizing the Rapid

Application Development (RAD) method, the Personnel Data and Dossier Processing Information System (SI-PANDA) has been successfully implemented to resolve critical administrative bottlenecks at Pusinfolahta TNI. Specifically, the system addresses the reliance on manual spreadsheet calculations and vulnerable local file storage by transitioning to a centralized web-based architecture. This shift delivers a significant operational impact, where the integration of automatic logic calculations ensures accuracy in personnel tenure and retirement dates, while the digital dossier management feature drastically reduces document retrieval time from minutes to seconds, providing decision-makers with real-time data for strategic deployment.

Furthermore, data security is substantially improved through Role-Based Access Control and automated backup mechanisms that mitigate the risks of unauthorized access and data loss. However, this study acknowledges certain limitations, including the system's deployment within a local server environment (XAMPP) which limits remote accessibility, the reliance on manual input for legacy data due to the absence of Optical Character Recognition (OCR), and the interface optimization which is currently restricted to desktop web browsers.

5. RECOMMENDATIONS

For further development, it is recommended to add an Audit Log feature to monitor user activity in detail, as well as integration with the TNI Single Sign-On (SSO).

REFERENCES

- [1] Pusinfolahta TNI, *Lintasan Sejarah Pusinfolahta TNI*. Cilangkap, Jakarta: Pusinfolahta TNI, 2023.
- [2] D. Gani, *Electronic Document Management System in Electronic Government Environment*. 2024. doi: 10.15405/epsbs.2024.05.48.
- [3] A. Ikuomola, E. Oyekan, and O. Orogbemi, "A Secured Cloud-Based Electronic Document Management System," *Int. J. Innov. Res. Dev.*, Dec. 2022, doi: 10.24940/ijird/2022/v11/i12/DEC22010.
- [4] P. Udayana, T. Legionosukmo, and S. Sundari, "STRATEGY FOR INTEGRATED LAND INFORMATION SYSTEM NETWORK ARRANGEMENTS FOR THE INDONESIAN NATIONAL ARMY," *Strateg. Pertahanan Darat*, vol. 8, Jun. 2022, doi: 10.33172/jspd.v8i1.1054.
- [5] O. Fedoruk, "Security and protection of information in electronic document management systems: improving the level of cyber defense," *Вісник Книжкової палати*, 2024, [Online]. Available: <https://api.semanticscholar.org/CorpusID:270897706>
- [6] Seftian Candra Pratama, E. Ahmad Firdaus, and R. Idul Adha, "Web-Based Assignment Management System Application Design to Increase Personnel Productivity and Effectiveness in Product Documentation Division," *Nuansa Inform.*, vol. 19, no. 2, pp. 73–80, 2025, doi: 10.25134/ilkom.v19i2.368.
- [7] Gemma G. Acedo, "Web-Based Document Tracking System for Catanduanes State University with QR-Code Technology: In Alignment with Republic Act No. 11032," *J. Inf. Syst. Eng. Manag.*, vol. 10, no. 6s, pp. 469–477, 2025, doi: 10.52783/jisem.v10i6s.745.
- [8] D. A. J. Markas Besar Angkatan Darat, "Dosir Elektronik (Lampiran II Keputusan Dirajenad Nomor Kep/ 9 / II / 2019)," Markas Besar Angkatan Darat, Direktorat Ajudan Jenderal, 2019. [Online]. Available: <https://www.scribd.com/document/646878686/4-DOSIR-ELEKTRONIK>
- [9] Imannudin Akbar, "Analisis Dan Perancangan Website SMKN 13 Bandung," *Nuansa Inform.*, vol. 18, no. 1, pp. 115–120, 2024, doi: 10.25134/ilkom.v18i1.68.
- [10] A. Fu'adi and A. Prianggono, "Analisa dan Perancangan Sistem Informasi Akademik Akademi Komunitas Negeri Pacitan Menggunakan Diagram UML dan EER," *J. Ilm. Teknol. Inf. Asia*, vol. 16, no. 1 SE-Article, pp. 45–54, Mar. 2022, doi: 10.32815/jitika.v16i1.650.
- [11] M. T. Abdillah, I. Kurniastuti, F. A. Susanto, and F. Yudianto, "Implementasi Black Box Testing dan Usability Testing pada Website Sekolah MI Miftahul Ulum Warugunung Surabaya," *J. Comput. Sci. Vis. Commun. Des.*, vol. 8, no. 1, pp. 234–242, 2023, doi: 10.55732/jikdiskomvis.v8i1.897.

- [12] N. Grozev and R. Buyya, "Performance modelling and simulation of three-tier applications in Cloud and Multi-Cloud environments," *Comput. J.*, vol. 58, no. 1, pp. 1–22, 2015, doi: 10.1093/comjnl/bxt107.
- [13] Arifin Oki, Sylvia, and R. Permata, *Algoritma dan Pemrograman Menggunakan Java*. 2025.
- [14] M. R. Wayahdi and F. Ruziq, "Pemodelan Sistem Penerimaan Anggota Baru dengan Unified Modeling Language (UML) (Studi Kasus: Programmer Association of Battuta)," *J. Minfo Polgan*, vol. 12, no. 1, pp. 1514–1521, 2023, doi: 10.33395/jmp.v12i1.12870.
- [15] A. Arif, B. Mahathir, P. L. Hang, C. Z. Hei, L. T. Hua, and N. U. Amin, "Design and Implementation of a Relational Database System for Efficient Facility and Booking Management," no. March, pp. 0–9, 2025, doi: 10.20944/preprints202503.0070.v1.
- [16] A. Jarkasih, A. Mufti, and W. Rahayu, "Sistem Informasi Penerimaan Siswa Baru SMK Al-Husna Citayam Berbasis Web," *J. Ris. dan Apl. Mhs. Inform.*, vol. 6, no. 03, pp. 503–510, 2025, doi: 10.30998/jrami.v6i03.7767.
- [17] D. Mustikasari Putri, T. H. Apandi, N. Khoirunnisa, and Y. Saputra, "Design of a Website-Based Personnel Information System at PT Indiga Nusa Digitama (Inditama)," *JESII J. Elektron. Sist. Inf.*, vol. 2, no. 2, pp. 209–223, 2025, doi: 10.31848/jesii.v2i2.3444.

High Performance Yolov4 with Different Optimizers and Backbones to Determine Distance

Irma Amelia Dewi^{1*}, Mikael Prapaskalis G², Muhammad Ichwan³

^{1,2,3}Department of Informatics, Institut Teknologi Nasional Bandung, Indonesia

E-mail: *¹irma_amelia@itenas.ac.id, ²paskalis97@gmail.com, ³ichwan@itenas.ac.id

Abstract

Monitoring the distance between individuals in public spaces remains relevant for crowd management, operational safety, and risk mitigation in public service. Manual inspections are difficult to carry out consistently and in real time, thus requiring an automated machine vision-based system. This research proposes a safe distance monitoring pipeline that detects humans, calculates the proximity between individuals. Human detection is performed with YOLOv4 and the CSPResNeXt50 backbone, then the center of the bounding box is used as a position representation. The distance between centroids is calculated using Euclidean Distance and calibrated based on camera parameters, thus that it can be mapped to distance units. To obtain a stable training configuration, two optimizers (Adam and SGD) are compared at two learning rates (1E-2 and 1E-3). Test results show that Adam with a learning rate of 1E-3 provides the best performance, achieving 96% object detection accuracy. The distance calculation module achieves 90% accuracy in determining the actual and predicted distances in optimal camera installation scenarios, with an average processing speed of 60 fps. These findings demonstrate a fast and practical approach for camera-based public space safety. Additionally, the backbone comparison shows that CSPResNeXt50 is slightly superior to DarkNet50, making it suitable for lightweight implementation on edge devices in real-time.

Keywords—object detection, Yolov4, CSPResNeXt50, deep learning

Submission: November 25, 2025 Accepted: December 23, 2025 Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.533>

1. INTRODUCTION

The density and proximity of individuals in public spaces (stations/terminals, shopping centers, public service areas, building entrances, campus corridors, and industrial areas) are important indicators for crowd management, operational safety, and compliance with Standard Operating Procedures (e.g., queue management, capacity restrictions, and prevention of mass gatherings) [1], [2], [3]. Manual monitoring by officers is not easy to carry out consistently due to limited visibility, fatigue, and the need for quick responses, so an automated system is needed that can work in real-time and be easily integrated with existing surveillance cameras.

Rapid developments in computer vision and artificial intelligence for crowd analytics, such as people detection, people counting/crowd density, and estimation of proximity between individuals from CCTV video. Crowd analytics research itself remains active and relevant to the context of safety and urban space management, including the need for efficient computing, particularly when directed towards edge/near-edge implementation [4]. The core component of such systems typically starts with human detection. The single-stage detector family is popular because of its favorable trade-off between speed and accuracy for real-time applications, and continues to improve in terms of architecture and training strategies, such as YOLOX (anchor-free, decoupled head) [5], YOLOv7

(trainable bag-of-freebies) [6], and RT-DETR [7], which leads to real-time end-to-end detection. This trend shows that the choice of architecture and training configuration (optimizer, learning rate, backbone) greatly determines the performance of the system in real-time scenarios.

On the other hand, estimating the actual distance from a single camera (monocular) is a major challenge due to the perspective effect, where the distance in pixels is not linear to the actual distance. Therefore, many approaches utilize calibration/homography (to map the floor area), inverse perspective mapping, or pose/depth-based strategies to reduce distance error in various CCTV camera conditions [8]. Other common challenges include occlusion, viewing angle variations, and differences in camera mounting heights, which could shift estimation accuracy. Based on these requirements, this study proposes a proximity monitoring pipeline between individuals based on human detection and simple yet fast distance calculation.

Therefore, the research was conducted to obtain a system that is able to detect the distance between humans that is considered according to the rules of physical distancing using Object Detection. Object detection is a technology that is able to detect an object of a particular class in an image or video. Object detection is capable of detecting human faces, detecting objects, detecting pedestrians (humans), vehicles, and so on [9]. The system uses You Only Look Once (YOLO).

YOLO is the best real-time object detection algorithm because YOLO uses a single convolutional neural network that simultaneously predicts several bounding boxes and classes of the entire image in a single detection or in one run [10]. Although newer YOLO variants such as YOLOv7 are reported to improve accuracy in the real-time range, this study chose YOLOv4 [11] because (i) the focus of the contribution is on the design of a CCTV-based proximity monitoring system and distance estimation evaluation, not on optimizing state-of-the-art detectors; (ii) YOLOv4 is designed for

speed-accuracy trade-offs and is reported to be capable of achieving real-time performance in certain configurations, and (iii) the stability and maturity of the YOLOv4 implementation/deployment toolchain supports experiment reproducibility and simpler real-time integration compared to end-to-end export pipelines in newer models [11], [12].

Based on these requirements, this study proposes a pipeline for monitoring proximity between individuals that relies on human detection and simple yet fast distance calculations. The system uses YOLOv4 with a CSPResNeXt50 backbone. This backbone was chosen based on the Cross Stage Partial Network (CSPNet) concept, which can improve CNN learning capabilities by combining CSP with several architectures, such as ResNet, ResNeXt, and DenseNet. Among these combinations, ResNeXt reportedly provides better results, giving rise to the CSPResNeXt variant, which can reduce computation time by about 20% while maintaining or even improving detection accuracy [13], [14].

In addition, this study compares two optimizers (Adam and SGD) and two learning rate values to obtain the most stable training configuration. The distance between individuals is calculated from the representative position (centroid bounding box) using Euclidean Distance, then tested in several scenarios of camera height and angle to determine the most optimal installation conditions.

In summary, the emphasized contributions are: (1) the design of a lightweight, real-time CCTV-based proximity monitoring system; (2) analysis of the influence of optimizers and learning rates on detection performance; and (3) evaluation of distance calculation accuracy under varying camera installation conditions to support practical implementation in the field.

2. RESEARCH METHOD

The proposed algorithm performs preprocessing on the image. The image will be resized according to the insert method.

Next, the process will continue on data augmentation.

The feature extraction process is carried out using CSPResNeXt50. In CSPResNeXt50 the convolution filter will be divided into two parts in order to reduce computation time and reduce memory usage, for example, when convulsing with filter 128, the convolution filter will be divided into two so that it becomes convulsive filter 64 with a size of 3x3, when it has passed the

convulsion process it will enters the transition layer where the output of the convolution is combined with the half of the input that has been divided previously to prevent duplication. This process is followed by ReLU activation and Maxpooling. If there is a low confidence value in the bounding boxes, then the bonding boxes are removed using non-max suppression. The flow process can be seen in Figure 1.

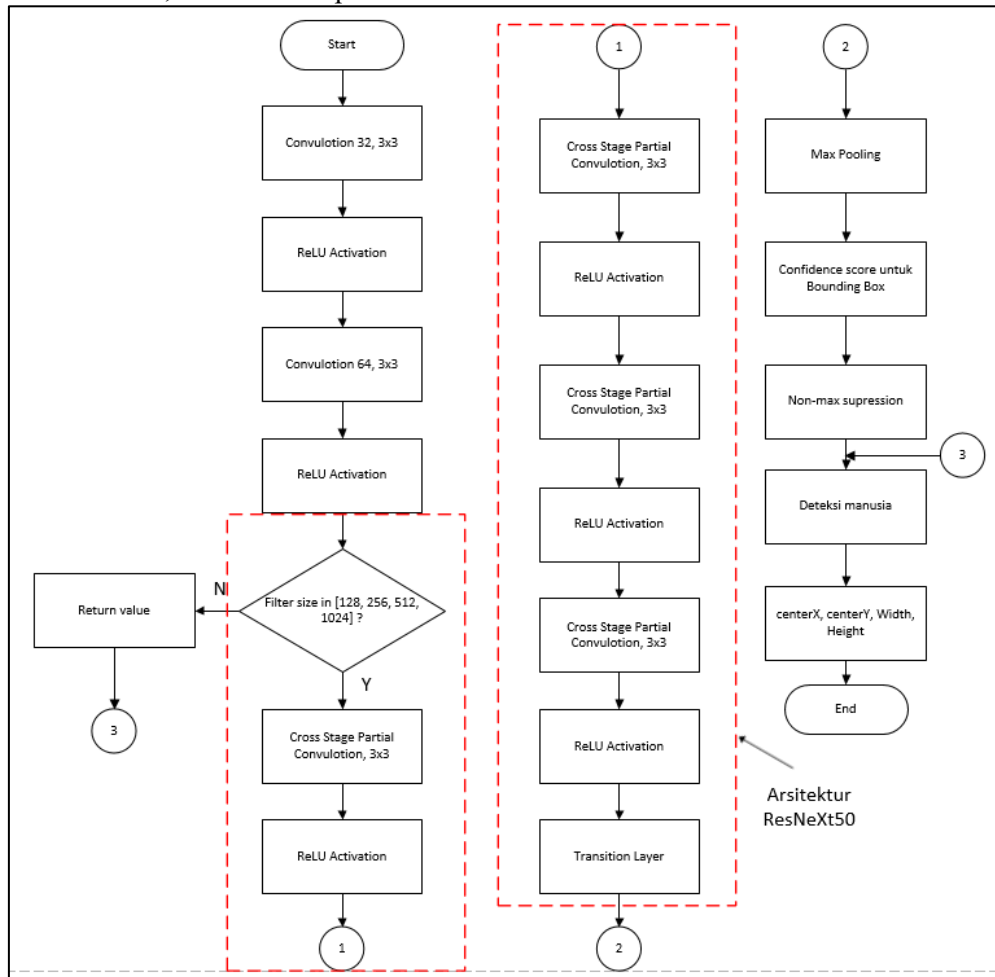


Figure 1. CSPResNeXt50

After feature extraction is performed on the input image, the extracted features will then be collected for detection and classification. This process occurs at the neck. Neck is the next process of object detection where the neck is in charge of taking and combining the features that have been extracted during the feature extraction process on the backbone which is then carried out in the object detection process. In

Yolov4 using PANet in the process of merging the extraction features.

After successfully detecting an object in the form of a human, the next step is to determine the centroid point of the detected object. To determine the centroid of the detected object, you can use Equation (1).

$$p = \left(\frac{x_1 + x_2}{2}, \frac{y_1 + y_2}{2} \right) \quad (1)$$

The centroid point will be obtained from the results of the determination using equation (1) which is based on processing xmin, xmax, ymin, ymax to produce x and y values which determine the value of the centroid point as shown in Figure 2.

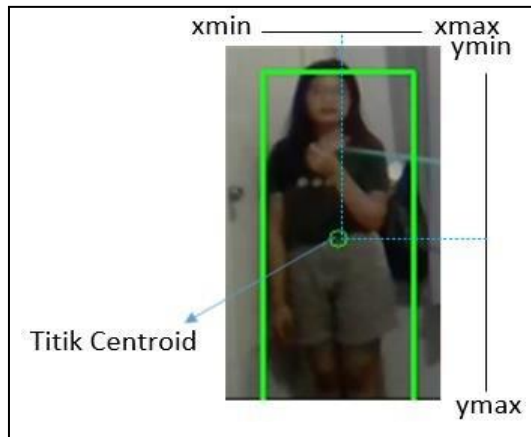


Figure 2. The Centroid Point

After obtaining the centroid point value of each object that has been detected, the process of calculating the distance between objects detected will be carried out using the Euclidean distance. Euclidean distance is the distance between two points in Euclidean space. Euclidean Distance is often used to calculate the distance between two different points. In this study, the distance between 2 objects was measured using Euclidean Distance and based on the object's centroid point that had been detected. After obtaining the centroid value on the detection object as in the above process, a process is carried out to calculate the distance between objects based on the predetermined centroid point [15]. This method is used to calculate the distance between objects detected by Equation (2).

$$d(p, q) = \sqrt{(p_1 - q_1)^2 + (p_2 - q_2)^2} \quad (2)$$

2.1 Distance Calibration

The calibration method for converting pixel distances to physical distances is critical because pixel measurements from camera images do not directly correspond to real-world distances due to perspective projection, camera mounting height, and viewing angle [16].

The calibration process involves steps starting from reference object placement, which two reference markers are placed at a known physical distance ($D_{ref} = 200$ cm) on the ground plane within the camera's field of view. Pixel distance measurement for capturing an image and measures the pixel distance (d_{ref_pixel}) between the two reference markers using the same centroid-based approach as human detection.

After that, the pixel to centimeter scale factor (α) is calculated as: $\alpha = D_{ref} / d_{ref_pixel}$ (cm/pixel). Converting the physical distance for each detected human pair to the pixel distance D_{pixel} (from Equation 2), the physical distance is: $D_{physical} = \alpha \times D_{pixel}$ (cm).

3. RESULT AND DISCUSSION

This section describes the results of testing the physical distancing system after going through the preprocessing process, extracting features using CSPResNeXt50, determining the centroid point and calculating the distance using Euclidean distance. The results of system testing can be seen in Figure 3.



Figure 3. The Result of Object Detection and Define a Centroid Point between Two Objects

The performance evaluation of the method is applied by measuring the optimizer and learning rate, i.e. Adam and SGD with learning rate 1E-2 and 1E-3. The following are the measurements results of Yolov4 method using different optimizer dan learning rate. Table 1 shows the effect of using different optimizer dan learning rate

that can impact the result of accuracy. The Adam optimizer hyperparameter and the learning rate of 1e-3 significantly influence the production of a more optimal model.

Based on the results of all the training processes that have been carried out, the Yolov4 model and the ResNeXt50 architecture use the Adam optimizer with a learning rate of 1E-3 and use 150 epochs of the highest accuracy with an accuracy quality of 96%. Furthermore, Table 2 shows the performance of the Yolov4 method using DarkNet50 architecture. The parameter is the same as when the Yolov4 method using ResNeXt50 architecture. The parameter influences the performance values and the better the performance value will be obtained. Model training was also carried out using the Yolov4 model with DarkNet50 architecture using the Adam optimizer with Learning rate 1E-3, epoch 100, and batch size 64. Training using the DarkNet architecture resulted in an accuracy of 94%.

Table 1. Measurement of Yolov4 Method with ResNeXt50

Backbone	Optimizer	Learning Rate	Epoch	Acc
ResNeXt50	Adam	1E-2	100	94.70%
ResNeXt50	Adam	1E-3	100	95.10%
ResNeXt50	Adam	1E-3	150	96%
ResNeXt50	SGD	1E-2	100	80.55%
ResNeXt50	SGD	1E-3	100	88.15%

Table 2. Measurement of Yolov4 Method with DarkNet50

Backbone	Optimizer	Learning Rate	Epoch	Acc
ResNeXt50	Adam	1E-3	100	95.10%
ResNeXt50	Adam	1E-3	150	96%
DarkNet50	Adam	1E-3	100	94%
ResNeXt50	SGD	1E-3	100	88%

Based on the results of the model training that has been carried out using two architectures, namely ResNeXt50 and DarkNet50. Where the ResNeXt50 architecture gives an accuracy result of 95.10% using the Adam optimizer, learning rate 1E-3, epoch 100, and batch size 64 and the ResNeXt50 architecture gives 94% accuracy results using the Adam optimizer, learning rate 1E-3, epoch 100, and batch size 64. ResNeXt50 architecture gives 1% better

or higher results than DarkNet50 architecture. A comparison of the two architectures used can be seen in Table 2.

From the results of the tests carried out by the Yolov4 model using the ResNeXt50 architecture, it gives better results and when training on the Yolov4 model with the ResNeXt50 architecture by adding 50 epochs it gives an accuracy of 96% which has an increase of 1% when using the optimizer and learning rate. which is equal to epoch 100.

Testing the Yolov4 model with the ResNeXt50 architecture is carried out by testing the model that has been trained by performing object detection. Where the Yolov4 model with ResNeXt50 architecture is distinguished based on the optimizer used when training the optimizer, namely Adam and SGD. The optimizer used in model testing is the optimizer that produces the highest accuracy value.

Table 3. The Evaluation of Yolov4 Model

Backbone	Optimizer	Learning Rate	Precision	Recall	F-Score
ResNeXt50	Adam	1E-3	96%	94%	95%
ResNext50	SGD	1E-3	89%	87%	87%

The results of the model testing carried out show that the Yolov4 model with the ResNeXt50 backbone, in Table 3 using the Adam optimizer with Learning Rate 1E-3 gives better results when compared to using the SGD optimizer and is the best weight or best fit used in the detection object of the Yolov4 model. Based on the system performance testing that has been done to find the highest accuracy based on the Euclidean distance used where the test was carried out 20 times. Model testing is done using the Yolov4 model and the CSPResNeXt50 backbone with the Adam optimizer giving good results, this is done by using a camera height of 2 meters with a camera angle of 90°.

Table 4. Euclidean Distance Accuracy

Method	Distance Object to Camera	Distance Between Objects	Acc	FPS
--------	---------------------------	--------------------------	-----	-----

Euclidean	300 cm	200 cm	90%	60
Euclidean	200 cm	200 cm	80%	60

From the test results as shown in the Table 4, it shows that The study has 90% accuracy in determining the actual and predicted distances at a camera distance of 300 cm and 80% at a distance of 200 cm based on 20 tests with the camera at a height of 2 meters and a 90° angle. The optimal distance when detecting using the Euclidean Distance Method is at a distance of 300 cm with a higher accuracy of 10% compared to the distance of the object to the camera is 200 cm. The detection results also provide more optimal results when object detection is carried out at the most optimal distance between the object and the camera where the optimal distance between the object and the camera is 3 meters.

4. CONCLUSION

The implementation of the system using the YOLOv4 model with the ResNeXt50 is able to detect objects in the form of humans and is able to compare the distance between human objects and the distance between physical distancing rules and determine whether the detected object violates the physical distancing rules or not. The distance calculation is done using Euclidean Distance which gives optimal results with an accuracy of 90% concordance between the actual distance and the distance predicted by the system and gets an average detection speed of 60 fps.

The YOLOv4 Model using the ResNeXt50 backbone using 2 optimizers with 2 different learning rates, it shows different results even though using the same batch size and epoch where the Adam optimizer provides 95.10% accuracy and the SGD optimizer provides 88.15% accuracy. Where the Adam optimizer shows better and more stable results than the results given by the SGD optimizer. Where the YOLOv4 model using Adam's ResNeXt50 optimizer backbone using a learning rate of 1E-3 with a batch size of 64 with 150 epochs shows better results with an accuracy of 96%. From the comparison results between the two architectures, namely ResNeXt50 and DarkNet50, it is found that the ResNeXt50

architecture gives 1% better results with 95.10% accuracy than the DarkNet50 architecture which gives 94% results using the Adam optimier, epoch 100, and batch size 64.

REFERENSI

- [1] I. Ahmed, M. Ahmad, J. J. P. C. Rodrigues, G. Jeon, and S. Din, "A deep learning-based social distance monitoring framework for COVID-19," *Sustain. Cities Soc.*, vol. 65, no. 102571, pp. 1–12, 2021, doi: 10.1016/j.scs.2020.102571. Available: <https://doi.org/10.1016/j.scs.2020.102571>
- [2] L. Gogulamudi, S. Mudigonda, Y. P. Renduchintala, B. Vemula, and G. H. Palnati, "A Social Distance Estimation and Crowd Monitoring System for Surveillance Cameras," *Res. Adv. Innov. Comput. Commun. Inf. Technol. ICRAIC2IT*, vol. 2796, no. 1, pp. 1–21, 2022, doi: 10.1063/5.0148967
- [3] M. Seker, A. Männistö, A. Iosifidis, and J. Raitoharju, "Automatic social distance estimation for photographic studies: Performance evaluation , test benchmark , and algorithm ☆," *Mach. Learn. with Appl.*, vol. 10, no. October, pp. 1–14, 2022, doi: 10.1016/j.mlwa.2022.100427. Available: <https://doi.org/10.1016/j.mlwa.2022.100427>
- [4] L. Deng, Q. Zhou, S. Wang, J. M. Górriz, and Y. Zhang, "Deep learning in crowd counting: A survey," *CAAI Trans. Intell. Technol.*, vol. 9, no. 5, pp. 1043–1077, 2024, doi: 10.1049/cit2.12241
- [5] Z. Ge, S. Liu, F. Wang, Z. Li, and J. Sun, "YOLOX: Exceeding YOLO Series in 2021," pp. 1–7, 2021, Available: <http://arxiv.org/abs/2107.08430>
- [6] C.-Y. Wang, A. Bochkovskiy, and H.-Y. M. Liao, "YOLOv7," *Cvpr*, pp. 7464–7475, 2023.
- [7] Y. Zhao *et al.*, "DETRs Beat YOLOs

- on Real-time Object Detection,” *Proc. IEEE Comput. Soc. Conf. Comput. Vis. Pattern Recognit.*, pp. 16965–16974, 2024, doi: 10.1109/CVPR52733.2024.01605
- [8] M. Aghaei, M. Bustreo, Y. Wang, G. Bailo, P. Morerio, and A. Del Bue, “Single image human proxemics estimation for visual social distancing,” *Proc. - 2021 IEEE Winter Conf. Appl. Comput. Vision, WACV 2021*, pp. 2784–2794, 2021, doi: 10.1109/WACV48630.2021.00283
- [9] V. Sharma and R. N. Mir, “A comprehensive and systematic look up into deep learning based object detection techniques: A review,” *Comput. Sci. Rev.*, vol. 38, p. 100301, 2020, doi: 10.1016/j.cosrev.2020.100301
- [10] Y. Jamtsho, P. Riyamongkol, and R. Waranusast, “Real-time license plate detection for non-helmeted motorcyclist using YOLO,” *ICT Express*, no. xxxx, 2020, doi: 10.1016/j.ict.2020.07.008
- [11] A. Bochkovskiy, C. Y. Wang, and H. Y. M. Liao, “YOLOv4: Optimal Speed and Accuracy of Object Detection,” *arXiv*, 2020.
- [12] D. Wu, S. Lv, M. Jiang, and H. Song, “Using channel pruning-based YOLO v4 deep learning algorithm for the real-time and accurate detection of apple flowers in natural environments,” *Comput. Electron. Agric.*, vol. 178, no. July, p. 105742, 2020, doi: 10.1016/j.compag.2020.105742
- [13] C. Y. Wang, H. Y. Mark Liao, Y. H. Wu, P. Y. Chen, J. W. Hsieh, and I. H. Yeh, “CSPNet: A new backbone that can enhance learning capability of CNN,” *IEEE Comput. Soc. Conf. Comput. Vis. Pattern Recognit. Work.*, vol. 2020-June, pp. 1571–1580, 2020, doi: 10.1109/CVPRW50498.2020.00203
- [14] M. Mahasin and I. A. Dewi, “Comparison of CSPDarkNet53 , CSPResNeXt-50 , and EfficientNet-B0 Backbones on YOLO V4 as Object Detector,” *Int. J. Eng. Sci. Inf. Technol.*, vol. 2, no. 3, pp. 64–72, 2022, doi: 10.52088/ijesty.v2i3.291
- [15] R. Rizaldi, A. Kurniawati, and C. V. Angkoso, “Implementasi Metode Euclidean Distance untuk Rekomendasi Ukuran Pakaian pada Aplikasi Ruang Ganti Virtual,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 5, no. 2, p. 129, 2018, doi: 10.25126/jtiik.201852592
- [16] M. Nieto and N. Aginako, “Multi-camera BEV video-surveillance system,” *Multimed. Tools Appl.*, vol. 82, pp. 34995–35019, 2023, doi: 10.1007/s11042-023-14416-y

Important Strategies in Mitigating Cyber Threats Using SLR in the Student Environment

Herdianto Sukmana^{*1}, Fathoni Mahardika², Dani Indra Junaedi³

^{1,2,3} University of Eleven April Sumedang, Indonesia

E-mail: ^{*1} herdianto12sukmana@gmail.com, ² fathoni@unsap.ac.id, ³ dani@unsap.ac.id

Abstract

In the ever-evolving digital era, cyber threats are increasingly diverse and demand effective mitigation strategies, especially for the educational environment and students who are easy targets for human error-based attacks. This study aims to analyze the role of cybersecurity awareness as an important strategy in mitigating cyber threats using the Systematic Literature Review (SLR) approach. A review was conducted on the literature that implements four main cybersecurity frameworks, namely NIST CSF, COBIT 5, the KAMI Index, and ISO/IEC 27001. The results of the analysis show that the technological aspect alone is not enough to guarantee security; Individual and organizational awareness—especially in the education sector, SMEs, and remote work environments—has a significant contribution to lowering the risk of cyber incidents. This study emphasizes that systematic, measurable, and framework-based awareness increase can strengthen cyber defense and support the development of a security culture among students.

Keywords— Cybersecurity Awareness, SLR, NIST CSF, ISO/IEC 27001, COBIT 5

Submission: December 09, 2025 Accepted: January 15, 2026 Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.540>

1. INTRODUCTION

Students are currently one of the largest groups of digital users who rely heavily on information technology services for daily activities, both in the learning process, academic data storage, and social interaction. The high level of use of these technologies makes them a highly vulnerable group to various forms of cyber threats, especially when awareness and understanding of information security is still low [4], [5]. A number of studies reveal that most students do not understand basic data protection practices, such as password management, device security, as well as the identification of common cyber threats such as phishing or identity theft [6], [7], [13]. This vulnerability is increasing due to the lack of information security materials in higher education curricula, as well as the lack of applicable and simulation-based training programs, even though these approaches have been proven to increase user awareness of social-engineering attacks [8], [9], [12].

Not only in the educational environment, research in the context of MSMEs, remote workers, and formal organizations also

shows that the human factor is the main weak point in the cyber defense system [1], [2], [10], [20]. The findings reinforce the view that the formation of a security culture requires continuous improvement of user awareness, especially through education that is structured and appropriate to the development of threats. In addition, several studies highlight the gap between the level of cyber awareness and cybersecurity resilience, which is heavily influenced by digital behavior, attitudes towards risk, and technology usage habits among students [13], [17].

Considering the wide and diverse scope of findings, the Systematic Literature Review (SLR) approach was chosen in this study to obtain a comprehensive picture of the role of cybersecurity awareness in threat mitigation, especially in the student environment. SLR allows researchers to systematically identify patterns, consistency of findings, and research gaps based on empirical evidence from various previous studies. In addition, this method helps map security frameworks used in global research such as NIST CSF, COBIT 5, the KAMI Index, and ISO/IEC 27001 and their relevance in the

context of information security education and training [14],[19]. Thus, this study is directed to synthesize the literature in depth, assess the contribution of cybersecurity awareness in reducing the risk of cyberattacks, and identify the need to develop more effective educational strategies and curricula for students.

Unlike prior systematic reviews that discuss cybersecurity awareness in broader organizational contexts, this study specifically synthesizes evidence within the student environment, where intensive digital activity and behavioral vulnerabilities are consistently reported as key risk factors [4], [5], [7], [13]. The novelty of this SLR lies in its integrated approach: (1) highlighting awareness gaps and behavioral weaknesses that are unique to students [5], [6], [16], (2) emphasizing the effectiveness of experiential and simulation-based trainingssuch as phishing simulations and gamification which has been shown to improve knowledge retention and protective behavior more effectively than conventional methods [9], [10], [12], [20], and (3) mapping the findings into multiple cybersecurity frameworks, including global standards (ISO/IEC 27001, NIST CSF, COBIT 5) and the local KAMI Index, to support structured and measurable cybersecurity awareness strategies in higher education [3], [10], [14], [15]. Therefore, this review contributes not only by summarizing existing studies, but also by providing a framework-oriented perspective that can guide universities in developing sustainable cybersecurity awareness programs for students.

2. RESEARCH METHODS

This study uses the Systematic Literature Review (SLR) approach to gain a comprehensive understanding of the trends, strategies, and effectiveness of cybersecurity awareness programs in higher education settings. The SLR method was chosen because it is able to systematically identify research patterns and gaps based on the empirical evidence available in the academic literature, as recommended by previous studies [10], [14], [16].

2.1. Formulation of Research Questions (RQ)

This research is designed to answer three main questions:

1. RQ1: What is the level of cybersecurity awareness among students?
2. RQ2: What are the factors that affect cybersecurity awareness in higher education settings?

3. RQ3: How effective are the training programs, simulation methods, and implementation of security frameworks (ISO 27001, NIST CSF, COBIT 5) in increasing student awareness?

This formulation refers to the focus of previous research that emphasized the importance of mapping awareness, behavior, and the effectiveness of security training in students [4], [5], [9], [17].

2.2. Inclusion and Exclusion Criteria

Inclusion and exclusion criteria are established to ensure that all articles analyzed in SLR have high relevance to the research objectives, have adequate academic quality, and support the process of systematic synthesis of findings. The following table summarizes the criteria used in the selection process:

Inclusion Criteria	Exclusion Criteria
Articles related to cyber threats to students	Articles that do not address students
Focus on cybersecurity awareness or mitigation	Irrelevant to cybersecurity
Articles published within the last 5–10 years	Articles under the year range of publication
Articles in English or Indonesian	Articles in languages other than English/Indonesian
Journal, conference, or scientific review type articles	Non-scientific blogs, news, opinions, or documents

Table I: Inclusion and Exclusion Criteria

The determination of inclusion criteria aims to ensure that the analyzed studies are truly related to the main focus of the research, namely cybersecurity awareness in students and higher education. The 5-10 year restriction of publication ensures that the articles used reflect the latest conditions of cyber threat development and relevant training programs, as suggested in modern research [4], [10], [14].

Meanwhile, exclusion criteria are used to filter out articles that do not make a real contribution to the synthesis process, such as non-scientific articles or publications without empirical data. It is important that SLR results remain credible, comprehensive, and evidence-based, as is the systematic research selection procedure recommended in recent SLR studies [12], [16], [19].

2.3. Literature Search Strategy

The literature search process is carried out systematically using the SLR principle to ensure that the selected articles are truly relevant to the topic of cybersecurity awareness in the student environment. Google Scholar was chosen as the primary source because it provides a wide coverage of scientific publications and has been used consistently in similar research [11].

The use of keywords such as "cybersecurity awareness", "students", "university", "information security training", "ISO 27001", and "NIST framework" refers to the recommendations of previous studies that show that the combination of these keywords is effective in mapping research on information security awareness in higher education [6], [7].

The 5-10 year range was chosen to ensure that the literature analyzed is up-to-date, relevant to cyber threat trends and modern learning approaches. In addition, language filters (English–Indonesian) are used to ensure a diversity of perspectives and accessibility of sources [12].

The search results yield about 54 articles. The selection process is then carried out based on the relevance of the title, abstract, and compatibility with the formulation of the research question. This step follows the literature review practice as suggested in the PRISMAA approach, which advocates a gradual screening process until the most relevant articles are left out [11]. From the total, 20 articles that met the criteria were selected for further analysis [2], [6], [7].

Components	Details
Database	Google Scholar
Year Range	5-10 years
Language	English & Indonesian
Key Keywords	"cybersecurity awareness", "students", "university", "education security"
Additional Keywords	"information security training", "ISO 27001", "NIST framework"
Search String	("cybersecurity awareness" AND "students") OR ("information security" AND "education")
Number of Initial Articles	54 articles
Selection Filter	Relevance of titles, abstracts, keywords, and full-text access

Final Number of Selected Articles	20 articles
-----------------------------------	-------------

Table II : Literature Search Strategy

2.4. Selection Process (PRISMAA Model)

The article selection process follows four stages of PRISMAA: identification, screening, eligibility, and inclusion.

1) Identification

The initial stage found 54 articles from Google Scholar that were relevant to the keyword.

2) Screening

A total of 22 articles were eliminated due to irrelevant titles/abstracts or duplication.

3) Eligibility

A total of 19 articles were read in full. Six articles were excluded for not providing empirical data, not focusing on students, or not using a security framework.

4) Final Inclusion

A total of 13 main articles were selected from the evaluation results [1]–[13], then 7 additional articles were added that met the latest criteria and strengthened the scope of SLR [14]–[20], resulting in a total of 20 articles being analyzed.

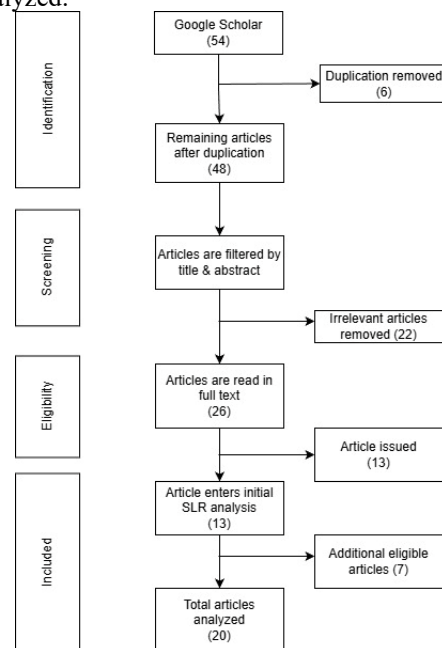


Figure 1 : PRISMA Diagram

2.5. Quality Assessment (Penilaian Mutu Studi)

To ensure the validity and credibility of the results of the Systematic Literature Review (SLR), each article selected at the final stage is further analyzed through a quality assessment

process. This quality assessment aims to ensure that the synthesised studies are of adequate methodological quality and relevant to the research focus, as recommended in SLR practice in the fields of cybersecurity and higher education [10], [14], [16].

Quality assessment was carried out by adapting evaluation criteria commonly used in SLR research, particularly those that emphasized methodological clarity and relevance of the research context [10], [16]. Each article was evaluated based on four main indicators, namely: (1) clarity of the purpose and scope of the research, (2) the suitability of the methodology with the context of cybersecurity awareness, (3) the existence of empirical data or systematic analysis, and (4) the relevance of the findings to the higher education environment or students [14], [20].

Each indicator is given a binary score (1 = meets the criteria, 0 = does not meet the criteria). Articles that obtain a minimum score of 3 out of a total of 4 indicators are declared feasible and retained in the final analysis. This approach is used to minimize bias as well as maintain the consistency and quality of the literature synthesis, as suggested in previous SLR studies [10], [16], [19].

The implementation of this quality assessment strengthens the methodological transparency of the research and ensures that the conclusions produced are truly based on credible and quality literature.

2.6. Data Extraction and Mapping of Research Questions

After the article selection process and quality assessment, a data extraction stage is carried out to identify key information from each selected study. The data extraction stage is an important component of the Systematic Literature Review (SLR) to ensure consistency, transparency, and traceability of the synthesis of findings, as recommended in previous SLR methodological studies [10], [16], [19].

The data extracted included the identity of the author and year of publication, the context and sample of the study, the method or form of cybersecurity training intervention, the security framework used, and the main findings of each study. This extraction approach refers to the common practice of SLRs in the field of cybersecurity and higher education that emphasizes the direct linkage between research findings and the formulation of research questions [14], [16].

Furthermore, each article is mapped to Research Questions (RQ1–RQ3) to identify its contribution in answering the level of cybersecurity awareness of students (RQ1), factors affecting cybersecurity awareness (RQ2), and the effectiveness of training programs and the implementation of security frameworks (RQ3). This mapping process aims to ensure that the synthesis of findings is carried out systematically and evidence-based, in accordance with the SLR approach recommended in the literature [10], [19].

Author	Study Focus	Framework	R Q
Fattah et al. (2023)	Student Surveys & Training	Knowledge–Attitude–Behavior	1,2
Bottyán (2023)	Student literacy survey	Model survei awareness	1
Janeš et al. (2024)	Student phishing simulation	Behavior-based training	3
Rijal et al. (2024)	Campus risk analysis	ISO/IEC 27001	2,3
Sodikin & Hikmawan (2023)	Security gamification	Gamification model	3
Goliath et al. (2024)	Attitude & behavior survey	Cyber resilience model	1,2
Taherdoo st (2024)	Training model	NIST Cybersecurity Framework	3
Weapons (2025)	Conceptual studies of the campus	Cybersecurity culture framework	2

Table III. Data Extraction and Article Mapping

It should be noted that not all studies use formal security frameworks such as ISO/IEC 27001 or NIST CSF. Some studies adopted conceptual approaches or behavioral models to measure cybersecurity awareness, which remains relevant in answering the RQ of this study.

2.7. Synthesis and Analysis Techniques

The synthesis process in this study uses a thematic analysis approach to identify the main

patterns of literature related to cybersecurity awareness in the student environment. The results of the study showed four dominant themes that were interrelated. First, various training approaches such as workshops, e-learning, gamification, and phishing simulations have been proven to improve students' digital security literacy through experiential learning [8], [9], [12], [19]. Second, the effectiveness of the intervention was evaluated using a pre-post test method that assessed the improvement in knowledge, attitudes, and changes in digital habits after training [4], [5], [13]. Third, several studies emphasize the integration of security frameworks such as ISO/IEC 27001, NIST CSF, COBIT 5, and the KAMI Index as a guide in assessing awareness as well as developing information security curriculum in universities [1]–[3], [14], [20]. Fourth, the literature also reveals implementation challenges, including lack of institutional policy support, lack of faculty involvement, and weak digital security culture, which causes cyber threat mitigation programs to not run optimally in the student environment [7], [17], [18].

2.8. Limitations of Data Sources.

In this study, the literature search process was carried out using Google Scholar as the main source of academic databases. Google Scholar's selection is based on its ability to provide a wide range of scholarly publications, including national and international journals, conference proceedings, as well as open access articles relevant to cybersecurity topics in higher education settings [10], [16].

However, the use of a single database has a number of limitations. Several SLR methodological studies confirm that Google Scholar has the potential to display publication duplication, variations in article quality, and the risk of selection bias due to indexing mechanisms and search algorithms that are not fully transparent [10], [19]. In addition, certain reputable publications that are indexed exclusively on databases such as Scopus, IEEE Xplore, or ACM Digital Library are potentially not optimally identified [16].

Therefore, the results of this Systematic Literature Review (SLR) need to be understood in the context of the limited data sources used. Nevertheless, the application of inclusion-exclusion criteria, quality assessment, and a PRISMAA-based selection process is expected to minimize potential bias and maintain the validity of the synthesis of findings [10], [19].

3. RESEARCH RESULTS

The results showed that of the 54 articles analyzed in the initial stage, as many as 20 articles met all inclusion criteria and were used as the main source in the synthesis. Key findings indicate that increased cybersecurity awareness is influenced by three main aspects, namely the effectiveness of training methods, the implementation of security frameworks, and user behavior factors. Training methods such as workshops, e-learning, gamification, and attack simulations have been proven to improve basic understanding and skills in detecting and responding to cyber threats. In addition, the integration of frameworks such as ISO/IEC 27001, NIST CSF, COBIT 5, and local evaluation models helps institutions assess readiness levels and build more structured awareness standards. The study also found that behavioral factors including digital habits, policy compliance, and individual vigilance levels have a significant role in the success of cybersecurity programs. Overall, the SLR results show that increasing awareness requires a comprehensive approach that incorporates training, policies, and digital culture change in the user environment.

4. DISCUSSION

4.1. Lack of Awareness Among Students

Research findings indicate that cybersecurity awareness among university students remains relatively low. Many students still use weak passwords, rarely update them, and do not consistently apply two-factor authentication. This condition is exacerbated by limited digital literacy, the absence of institutional standard operating procedures (SOPs) for incident response, and the misconception that cybersecurity is solely the responsibility of technical staff. Several studies also report a strong relationship between formal training programs and increased awareness, which leads to safer digital behavior.

Low awareness is reflected in students' daily digital habits, including the tendency to ignore security warnings, access public Wi-Fi without protection, and share personal information through social media without considering potential risks. This suggests that students often lack contextual understanding of the consequences of cyberattacks, such as data theft, identity abuse, and account compromise. Furthermore, surveys show that only a small proportion of students are familiar with campus reporting procedures for cyber incidents,

indicating limited institutional support for awareness-building.

To address these issues, universities should strengthen preventive measures by integrating cybersecurity awareness into structured educational programs. Basic protective behaviors such as avoiding suspicious links, maintaining account confidentiality, and enabling authentication mechanisms should be promoted through continuous training and supported by clear institutional policies.

4.2. Effectiveness of Training and Simulation

Simulation-based training programs such as phishing simulation and gamification have proven to be more effective than conventional learning. This approach improves students' ability to recognize real cyber threats and significantly strengthens protective behavior. This confirms that the use of risk-based training modules and case studies encourages students to understand the real impact of information security negligence.

Some institutions have tried to implement simulation-based cybersecurity training. This approach is considered more effective than conventional learning methods. For example, training that mimics a phishing attack in a controlled environment can show first-hand how students react to real-life scenarios. In training like this, students receive fake emails that resemble real phishing attacks. Their responses were then analyzed to measure their level of alertness and ability to recognize threats.

In addition, e-learning and gamification-based training has also been proven to increase students' interest in learning. Platforms such as interactive quizzes, educational games, and real-life case-based modules make the learning process more fun and applicable. Students become more active in recognizing risks and understanding ways to avoid them. For example, in some training programs, students are asked to play the role of 'attacker' and 'victim' to better understand the point of view from both sides.

Simulation-based training is a learning method that puts students in real or near-reality situations. For example:

- Phishing simulation: students receive fake emails, then analyze their responses,

- Educational games: such as CTF (Capture The Flag) that trains tactical responses to threats,
- Real-life case module: campus digital incident analysis or data leak simulation,
- Blended learning training: combines theory and hands-on practice.

This training shows that:

- The rate of knowledge retention is higher than that of theoretical lectures,
- Students are more critical and active towards suspicious digital content,
- Participation increases if the material is interactive and challenging,
- Students feel responsible for the security of their accounts and personal devices.

Other advantages:

- Students learn from mistakes safely.
- The material is more relevant to everyday digital life,
- Can be done online and flexibly,
- Encourage the creation of a collaborative learning community.

However, this training must be carried out on an ongoing basis. One-time training without follow-up is often not enough to form a habit. Therefore, it is necessary to schedule regular sessions, periodic evaluations, and integration in the curriculum.

4.3. The Role of Frameworks in Education

Frameworks such as ISO/IEC 27001 and NIST CSF provide a systematic foundation for campuses to develop security and training policies. The implementation of this framework not only increases technical awareness but also strengthens students' managerial literacy towards information security governance.

More than that, the application of this framework in the classroom also helps lecturers in compiling teaching materials

that are structured and in accordance with industry practices. Through this approach, students not only understand theory, but are also able to map risks and design basic security policies that are relevant to the everyday digital context. Case study-based practicum activities using framework elements have also been carried out in several universities, where students are asked to audit simple security systems and make recommendations for improvement.

Frameworks such as ISO/IEC 27001 and NIST CSF help students understand information security structures systematically. Some of the implementation practices on campus include:

- Practicum makes security policies based on ISO 27001,
- Risk analysis of daily digital activities using the NIST framework,
- Group discussions on case studies of the implementation of the framework in educational institutions,
- The application of mini audits by students to the campus digital system.

This framework reinforces:

- Technical and managerial literacy of students.
- Readiness to enter the world of work based on industry standards,
- Collaboration between the IT faculty and the campus security department,
- Awareness of the importance of a proactive, not reactive, approach.

Furthermore, with the introduction of this framework since college, students can understand that information security is a comprehensive process, starting from policy making, asset identification, to access control. This knowledge is not only relevant for IT graduates, but also for students from

other majors such as business, communications, and law.

4.4. Challenges of Implementation on Campus

The main challenges include limited resources for security expert lecturers, lack of cyber laboratories, and low institutional policy support. Some universities have tried to address this with programs such as "Digital Security Month" and collaborations with BSSN. However, the results of the evaluation show that without continuous training and integration into the curriculum, its effectiveness is still limited.

The implementation of cybersecurity awareness still faces challenges such as:

1. Policies not yet prioritized:

- Awareness has not entered compulsory courses,
- Lack of modules or guidance from institutions,
- There is no roadmap for the development of digital competencies as a whole.

2. Resource limitations:

- The lack of lecturers who master the field of information security,
- Not all campuses have laboratories or security simulation tools,
- The budget for training and platform development is still limited.

3. Permissive digital culture:

- Students tend to be ignorant of system warnings,
- Digital information is disseminated without thinking about the consequences,
- There is a perception that security is the responsibility of the technician, not the user.

Nevertheless, there are good practices that can be emulated:

Regular programs such as "Digital Safety Month" and "Hoax-Free Day",

Collaboration with institutions such as BSSN or local security startups,

Integration of digital security materials in new student orientation,

Organizing information security education video competitions,

Creating a community of students who care about cybersecurity.

These steps are proven to be starting to build a culture of digital security awareness among students and lecturers. If it continues to be developed and supported by strong institutional policies, the campus can become a pioneer of cybersecurity awareness in Indonesia. This will make a real contribution to creating a smart, safe, and

4.5. Cybersecurity Awareness Strategy Framework in the Campus Environment

Based on the synthesis of findings in the previous section, this study not only provides an overview of the condition of student cybersecurity awareness, but also produces a conceptual strategy framework that can be applied in the campus environment. The framework is designed as a practical guide for higher education institutions in raising cybersecurity awareness in a structured and sustainable manner.

This strategic framework is built on three main pillars. The first pillar is training (*training*), which emphasizes the importance of a sustainable and experience-based cybersecurity education program. The training not only focuses on delivering theoretical material, but also involves simulated attacks, case-based learning, and interactive activities that encourage students to recognize and respond to cyber threats firsthand.

The second pillar is institutional policy (*Stuart T*). Higher education institutions need to have information security policies that are clear, easy to understand, and relevant to daily academic activities. The policy includes account management, personal data protection,

incident response procedures, and the division of roles and responsibilities between students, lecturers, and system managers. Policy clarity is an important foundation for security practices to be applied consistently.

The third pillar is the formation of a cybersecurity culture (*security culture*). A culture of security is built through safe digital habits, active participation of the entire academic community, and continuous support from institutional leaders. Without a strong safety culture, the effectiveness of training and policies tends to be temporary and unsustainable.

The integration of the three pillars forms a strategic roadmap that the campus can use in designing cybersecurity awareness improvement programs. This framework allows higher education institutions to move from a reactive approach to a proactive approach in dealing with cyber threats, particularly those involving human factors in the student environment.

4.6. Novelty and Contributions of This Study

This Systematic Literature Review offers several contributions that distinguish it from previous SLRs on cybersecurity awareness. First, rather than examining cybersecurity awareness in a general context, this study focuses on the student environment as a high-risk population characterized by intensive digital activity and behavioral vulnerabilities. Second, the review synthesizes findings across four security frameworks including both global standards (ISO/IEC 27001, NIST CSF, COBIT 5) and a local assessment model (KAMI Index) to provide a comparative, structured perspective that is rarely addressed in earlier reviews. Third, beyond summarizing existing studies, this research proposes a practical strategy framework for campuses based on three integrated pillars (training, institutional policy, and security culture). This framework can be used as a roadmap to guide universities in implementing sustainable, measurable, and

evidence-based cybersecurity awareness initiatives.

5. CONCLUSION

Cybersecurity awareness is a fundamental factor in reducing vulnerability to cyber threats in the higher education environment. The results of this study show that training programs designed with interactive methods—such as phishing simulations, gamification, and scenario-based training modules—drive significant improvements in student understanding, alertness, and behavioral changes in the use of digital technology. This approach not only helps students recognize attack patterns, but also strengthens their ability to make secure decisions when interacting with information systems.

Beyond confirming the importance of cybersecurity awareness, this study offers a novel contribution by focusing specifically on students as a high-risk population and by synthesizing multiple cybersecurity frameworks—ISO/IEC 27001, NIST CSF, COBIT 5, and the KAMI Index—within a single strategic model. The proposed three-pillar framework (training, institutional policy, and security culture) provides a practical and integrative roadmap for higher education institutions, which has been largely absent in prior systematic reviews.

In addition, the integration of security frameworks such as ISO/IEC 27001 and the NIST Cybersecurity Framework into learning has been proven to provide a systematic foundation in shaping a digital security culture. The framework helps educational institutions develop more structured strategies, ranging from risk management, access control, to incident response procedures. As a result, students gain a more comprehensive understanding of the globally applicable security practice standards.

Overall, this study confirms that the combination of experience-based training methods and the application of international security frameworks is an effective step in increasing students' readiness to face increasingly complex and dynamic cyber threats. These findings also underscore the importance of educational institutions' commitment to providing ongoing training programs to build a strong information security culture in the academic environment.

6. RECOMMENDATIONS

This research is suggested to strengthen the methodological aspect of SLR by explaining the article selection process in a more structured and in-depth manner. Stages such as the number of initial findings, screening steps, elimination based on inclusion-exclusion criteria, and the final number of articles analyzed need to be described in detail so that methodological transparency can be maintained. In addition, the research also needs to add a quality assessment component to assess the methodological feasibility of each selected article, so that the conclusions produced have a stronger, valid, and accountable basis. The implementation of this quality assessment will help ensure that only relevant and academic literature of adequate academic standards is used in SLR analysis.

Further research is suggested to expand the database sources by involving reputable platforms such as Scopus, IEEE Xplore, or ACM Digital Library to obtain a more comprehensive coverage of the literature and reduce the potential for selection bias. In addition, the use of more than one database allows for stronger study replication and increases the validity of the findings in the Systematic Literature Review in the field of educational cybersecurity.

REFERENCES

- [1] M. Hijji and G. Alam, "Cybersecurity Awareness and Training Framework for Remote Working Employees," 2022.
- [2] F. Ugbebor, O. Aina, M. Abass, and D. Kushanu, "Employee Cybersecurity Awareness Training Programs Customized for SME Contexts to Reduce Human-Error Related Security Incidents," 2024.
- [3] D. M. Rijal et al., "Information Security Awareness Analysis of the Threat of Data Leakage in Educational Institutions with the ISO 27001 Framework," 2024.
- [4] A. Fattah, W. Wagimin, and N. Nurlia, "Enhancing Cybersecurity Awareness among University Students: A Study on the Relationship between Knowledge, Attitude, Behavior, and Training," *JSI: Jurnal Sistem Informatika*, vol. 15, no. 1, 2023.
- [5] L. Botlyán, "Cybersecurity awareness among university students," *J. Applied Technical and Educational Sciences*, vol. 13, no. 3, 2023.
- [6] K. A. Y. Yaseen, "Importance of Cybersecurity in The Higher Education Sector," *Asian J. of Computer Science and Technology*, vol. 11, no. 2, 2022.

- [7] A. Alabab et al., "Cybersecurity Awareness of College Students in a Private Higher Education Institution," *CGCI Int. J. of Administration, Management, Education and Technology*, vol. 1, no. 1, 2024.
- [8] A. Hakimi and M. F. Zolkipli, "An Awareness of Cybersecurity Risk Assessment and Management Among Students," *Borneo Int. J.*, vol. 7, no. 4, 2024.
- [9] R. A. Sodikin and R. Hikmawan, "Analysis of Gamification in Cybersecurity Education for Students: A Systematic Literature Review," *J. Educative: Journal of Educational Studies*, vol. 8, no. 2, 2023.
- [10] H. Taherdoost, "Towards an Innovative Model for Cybersecurity Awareness Training," *Information*, vol. 15, no. 9, 2024.
- [11] O. C. Okeke and C. E. Amaechi, "Awareness of Phishing Attacks in Institutions of Higher Learning: A Review of Types and Technical Approaches," *Int. J. of Research and Innovation in Applied Science*, vol. 9, no. 10, 2024.
- [12] H. Janeš, K. Bilić, and M. Grgić, "Application of simulated phishing attacks for user training," *Crisis Management Days*, 2024.
- [13] S. Goliath, P. Tsibolane, and D. Snyman, "Exploring the Cybersecurity-Resilience Gap: An Analysis of Student Attitudes and Behaviors in Higher Education," *arXiv preprint arXiv:2411.03219*, 2024.
- [14] R. Armas, "Building a Cybersecurity Culture in Higher Education," *Information*, 2025.
- [15] M. B. Barruga, "SYSTEMATIC REVIEW OF CYBERSECURITY FRAMEWORKS FOR HIGHER EDUCATION INSTITUTIONS: CHARACTERISTICS, COMPONENTS, AND CHALLENGES," *International Journal of Applied Mathematics*, 2025.
- [16] Abdulrahman A. Arishi, Nazhatul H. Kamarudin, Khairul Azmi Abu Bakar, Zarina Binti Shukur & Mohammad Kamrul Hasan, "Cybersecurity Awareness in Schools: A Systematic Review of Practices, Challenges, and Target Audiences," *International Journal of Advanced Computer Science and Applications (IJACSA)*, vol. 15, no. 12, 2024.
- [17] K. Sittitiamjan & P. Wuttidittachotti, "Bridging the Gap: How Knowledge, Attitudes, and Practices Shape Cybersecurity Awareness in Thai Education," *EDUPIJ*, 2025. [EduPij](#)
- [18] MM Kumbhakar & Nagendra Kumar, "Cyber Security Awareness Among Higher Education Rural Students," 2025.
- [19] Rahime B. Sağlam, Vincent Miller & Virginia N. Franqueira, "A Systematic Literature Review on Cyber Security Education for Children," *IEEE Transactions on Education*, 2023.
- [20] C. Mouwers-Singh & T. Buzy Musikavanhu, "A Narrative Review on Enhancing Cybersecurity in Higher Education Institutions: The Role of Continuous Training and Awareness," *Expert Journal of Business and Management*, vol. 12, no. 2, 2024.

Interactive Dashboard-Based Mail Classification Using Deep Learning Models

Ruth Diana Purnamasari^{*1}, Nisa Hanum Harani²

^{1,2}Universitas Logistik dan Bisnis Internasional, Indonesia

E-mail: ^{*}ruthdianaps@gmail.com, ²nisa@ulbi.ac.id

Abstract

The management of incoming mail archives at a large national logistics company in Indonesia generates a large volume of unstructured textual data, making manual classification inefficient and error-prone. This study evaluates the performance of deep learning models for administrative mail archives classification using data collected between 2023 and 2025. Three models are examined, namely Long Short-Term Memory (LSTM), Bidirectional LSTM (Bi-LSTM), and Convolutional Neural Network (CNN). Model performance is evaluated using accuracy, precision, recall, F1-score, and confusion matrix metrics. Experimental results indicate that CNN achieves the highest accuracy of 85.82%, outperforming LSTM and Bi-LSTM models. This superior performance is attributed to CNN's ability to capture local textual patterns through convolution operations, which are well-suited to the structured and repetitive language characteristics of official correspondence. To support practical interpretation, an interactive dashboard is implemented as a visualization tool for model evaluation results, classification outcomes, and clustering analysis. These findings demonstrate that deep learning-based approaches integrated with visual analytics can significantly improve the efficiency and accuracy of unstructured mail archive management.

Keywords—mail archive classification, deep learning, CNN, model evaluation, interactive dashboard, clustering

Submission: January 09, 2026

Accepted: January 18, 2026

Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.546>

1. INTRODUCTION

Mail archive management is an essential component in supporting the smooth operation of an organization, particularly in institutions that handle a high volume of incoming mail. At a large national logistics company in Indonesia, incoming mail processed at branch offices includes various types and subject matters originating from government institutions as well as external parties. All incoming mail is recorded in a mail log as archival documentation, after which it is distributed and followed up through established administrative mechanisms[1]. Based on internal archival records obtained from the incoming mail agenda of the company during the period 2023-2025, a total of 3,453 mail entries were recorded before data cleaning and validation. The data are in the form of unstructured text, characterized by

information such as the subject of the letter, the type of incoming mail, and the sender's name. The high volume and variability of this raw data make the processes of matching, grouping, and classifying mail archives complex when performed manually. Similar conditions are also found in various types of administrative data in other institutions. The use of conventional methods in mail archive management may lead to delays, classification inconsistencies, and difficulties in the retrieval of archived documents [2].

Along with the advancement of information technology, the application of text mining and artificial intelligence has begun to be integrated into digital archiving systems to improve the effectiveness and accuracy of mail management. Several previous studies have shown that document classification algorithms are able to automate the process of mail grouping and increase productivity compared to manual

methods[3],[4]. However, many existing studies emphasize automation results without providing clear, interpretable, and user-oriented model evaluation mechanisms that can be easily understood by non-technical users, such as archive officers or administrative staff in institutional environments [3].

In the field of text document classification, Deep Learning approaches have been widely applied due to their ability to better understand linguistic context and the relationships between words compared to traditional methods. Models such as *Long Short-Term Memory (LSTM)*, *Bidirectional LSTM (Bi-LSTM)*, and *Convolutional Neural Network (CNN)* have been proven effective in various document and email classification tasks, particularly for texts with diverse and complex sentence structures[5],[6],[7]. However, most previous studies still focus primarily on reporting performance metrics, such as accuracy and F1-score, without sufficient analytical discussion or integration of these results into visual systems that can be directly utilized in institutional operations [8],[9],[10]. In addition to the classification process, analysis of patterns and structures within mail archive data is also an important aspect, considering the continuously increasing number of documents. Clustering methods can be used to group mail based on content similarity, thereby assisting in understanding topic distribution, characteristics, and the most frequently occurring types of mail[11],[12]. Nevertheless, clustering results are often presented only in numerical or tabular form, making them difficult to interpret without appropriate visual support[13],[14].

To ensure that the data preprocessing procedure is conducted in a systematic and structured manner, this study adopts the CRISP-DM (*Cross-Industry Standard Process for Data Mining*) methodology. This methodology provides a framework consisting of six main stages, namely *Business Understanding*, *Data Understanding*, *Data Preparation*, *Modeling*, *Evaluation*, and *Deployment* [15],[16]. The use of a standardized

methodology in the data classification process is a fundamental principle for generating valuable information patterns for institutions[17].

Nevertheless, previous studies remain limited to the development of mail archive classification systems without comprehensively integrating model performance evaluation, analytical discussion, and clustering analysis within an interactive dashboard. Based on these conditions, a system is required that is not only capable of performing automatic mail classification but also provides informative, interpretable, and user-oriented evaluation results. Several studies on digital archiving systems indicate that the utilization of dashboards can assist users in understanding data conditions and system effectiveness more optimally[3],[8].

Therefore, this study is directed toward the implementation of mail document archive classification using deep learning methods, complemented by an interactive evaluation dashboard and clustering analysis. This research does not aim to develop new algorithms, but rather to apply and critically evaluate multiple deep learning models in a case study using data from a large national logistics company in Indonesia. The main contribution of this study lies in the comparative evaluation of multiple deep learning models combined with visual analytics to enhance interpretability and practical usability in administrative mail classification. Through this approach, it is expected that mail document archive management can be carried out effectively, transparently, and systematically to support the utilization of information technology within institutional environments.

2. METHODOLOGY

This study is designed to provide a systematic overview of stages involved in the data processing, modeling, and evaluation processes. The methodology is structured to ensure that each process is conducted in a measurable, objective, and scientifically accountable manner. In general, study includes stages of mail

archive data collection, text data preprocessing, application of deep learning models for document classification, model performance evaluation using quantitative metrics, as well as result evaluation and clustering analysis presented through an interactive dashboard.

2.1. Research Approach and Type

This study adopts an applied *research* approach using computational experimental methods. This approach is selected because the research focuses on the application of deep learning models for classifying mail archive text documents and evaluating model performance based on measurable quantitative metrics. The evaluation results are then presented in the form of an interactive dashboard to facilitate analysis and interpretation.

The type of research employed is quantitative research, as all research outcomes are measured numerically using model performance metrics, namely *accuracy*, *precision*, *recall*, and *F1-score*. Through this approach, the effectiveness of each deep learning model can be measured objectively and compared systematically. In addition, this study also utilizes clustering analysis as a supporting method to understand patterns and structures within the mail archive data used.

2.2. Data Sources and Types

The source type of data used in this study include:

2.2.1 Data Source

The data used in this study originate from the incoming mail agenda of a large national logistics company in Indonesia for the period 2023 to 2025. The dataset constitutes secondary data obtained from the internal administrative recording system and is used as part of the organization's operational activities.

2.2.2 Data Type and Research Variables

The type of data used in this study consists of unstructured text, derived from

descriptive columns in the mail archive. Based on their roles in the modeling process, the research variables are categorized as follows:

1. Independent Variables

The independent variables consist of the *Subject*, *From/To*, and *Mail Type* columns. These three attributes are combined to form a richer textual representation and are used as input for the deep learning models.

2. Dependent Variable

The dependent variable in this study is the Mail Type Category, which serves as the classification target for the model, such as *Surat Dinas*, *SPK*, *PKS*, and *KD*.

2.3. Dataset Description

The dataset used in this study consists of 3,453 raw entries of incoming mail archive data obtained from the internal administrative system of a large national logistics company in Indonesia during the period 2023-2025. Data selection and validation were conducted by removing duplicate records, incomplete entries, and irrelevant data that did not meet the research criteria. As a result of this process, 2,817 clean and validated data entries were retained and deemed suitable for use in the modeling stage.

The validated dataset was subsequently divided into training and testing sets using an 80:20 ratio, where 80% of the data were used for model training and 20% were reserved for performance evaluation. This data splitting strategy was applied to maintain a balance and objective evaluation of the proposed models.

The distribution of incoming mail archive data across categories indicates an imbalanced class condition. Several categories, such as *Surat Dinas*, *SPK* and *PKS*, dominate the dataset, while other categories contain fewer instances. The class distribution of the validated dataset is illustrated in figure 4, providing an overview of the proportion of each mail category used in this study.

2.4. Methodology Flowchart

The research methodology is designed in a structured manner to ensure that each stage is carried out systematically and can be scientifically justified. Conceptually, the research flow adapts the CRISP-DM [15].

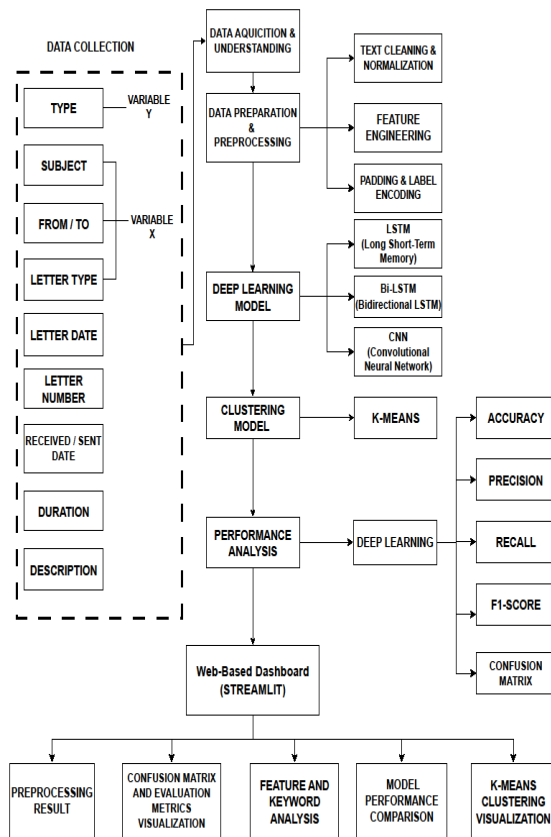


Figure 1 Research Stage

The methodology flow diagram is presented in Figure 1, illustrating the main stages of the study from data collection to visualization.

2.2.3 Description of Research Stages

Based on the flow diagram shown in Figure 1, the research stages include:

1. Data Collection

Incoming mail archive data are collated from multiple internal files of a large national logistics company in Indonesia for the period 2023-2025 and combined into a unified dataset.

2. Data Selection and Validation

The data are examined to remove empty, duplicate, or irrelevant entries

that do not align with the research objectives.

3. Text Data Preprocessing

This stage includes text cleaning, normalization text (*lowercasing*), text attribute merging, tokenization, padding, and category label encoding. For illustration purposes, an example of preprocessing transformation includes converting the original text “Surat Permohonan kerja Sama PT ABC” into the cleaned and normalized from “surat permohonan kerja sama pt abc”.

4. Dataset Splitting

The dataset is divided into training and testing sets using an 80:20 ratio to maintain a balance between training and evaluation processes.

5. Deep Learning Model Training

LSTM, *Bi-LSTM*, and *CNN* models are trained using the training dataset to learn textual patterns within the mail archive data. All models were trained using the same experimental settings to ensure a fair comparison. The training process was conducted with a batch size of 32 and a maximum of 50 epochs. Early Stopping was applied by monitoring validation loss to prevent overfitting during training.

6. Model Evaluation

Model evaluation is conducted using quantitative performance metrics and a confusion matrix to analyze prediction errors.

7. Result Visualization through Dashboard

The results of the evaluation and analysis are displayed through a web-based interactive dashboard using *Streamlit*.

8. Clustering Analysis

The *K-Means* algorithm is used as an exploratory method to analyze patterns and distribution in letter archive data. The clustering results are not used as a classification mechanism, but serve as complementary analysis to provide additional insight into the structure

and thematic distribution of the classified mail archive data.

9. Interpretation of Result and Discussion

All results were analyzed to draw research conclusions.

2.5. Data Preprocessing

The preprocessing stage is conducted to ensure that text data are ready to be processed by deep learning models. This process consists of several steps, including text cleaning using regular expression functions (*re.sub*), text normalization by converting all characters to lowercase, merging textual attributes, tokenization and padding processes, and label encoding to convert categorical labels into numerical representations.

2.6. Deep Learning Models Used

This study applies three algorithms in the text document classification process, namely *Long Short-Term Memory (LSTM)*, *Bidirectional LSTM (Bi-LSTM)*, and *Convolutional Neural Network (CNN)*.

2.6.1. Long Short-Term Memory (LSTM)

LSTM utilizes a memory gate mechanism to retain important information within text sequence data. The basic equations of *LSTM* are formulated as follows:

Forget Gate (1)

$$f_t = \sigma(w_f \cdot [h_{t-1}, x_t] + b_f)$$

Input Gate (2)

$$i_t = \sigma(w_i \cdot [h_{t-1}, x_t] + b_i)$$

$$\tilde{C}_t = \tanh(w_c \cdot [h_{t-1}, x_t] + b_c)$$

Update Cell State (3)

$$C_t = f_t \cdot C_{t-1} + i_t \cdot \tilde{C}_t$$

Output Gate (4)

$$o_t = \sigma(w_o \cdot [h_{t-1}, x_t] + b_o)$$

$$h_t = o_t \cdot \tanh(C_t)$$

Table 1 Algorithm LSTM Classifications

Algorithm 1 : *Long Short-Term Memory (LSTM)* Classifications

Output : hidden state h_t

1 *Forget Gate* : Computes previous information that should be retained or discarded

$$f_t = \sigma(w_f \cdot [h_{t-1}, x_t] + b_f)$$

2 *Input Gate* : Determines new information to be stored

$$i_t = \sigma(w_i \cdot [h_{t-1}, x_t] + b_i)$$

3 *Candidate Cell State* : Forms candidate values for new memory

$$\tilde{C}_t = \tanh(w_c \cdot [h_{t-1}, x_t] + b_c)$$

4 *Update Cell State* : Updates the cell memory

$$C_t = f_t \cdot C_{t-1} + i_t \cdot \tilde{C}_t$$

5 *Output Gate* : Determines the output of the cell

$$o_t = \sigma(w_o \cdot [h_{t-1}, x_t] + b_o)$$

6 *Hidden State Output* : Produces the final hidden state

$$h_t = o_t \cdot \tanh(C_t)$$

2.6.2. Bidirectional LSTM (Bi-LSTM)

Bi-LSTM processes data in two directions, forward and backward, enabling the model to capture broader contextual information within text data. The model is formulated as follows :

$$h_t = [\vec{h}_t \oplus \overleftarrow{h}_t] \quad (1)$$

Where $\vec{h}_t$ represents the forward *LSTM* and $\overleftarrow{h}_t$ represents the backward *LSTM*.

Table 2 Algorithm Bi-LSTM Classifications

Algorithm 2 : *Bidirectional LSTM Memory (Bi-LSTM)* Classifications

Output : bidirectional context representation h_t

1 *Forward LSTM* : Processes the sequence from the beginning to the end

$$\vec{h}_t = LSTM_{\{forward\}}(x_t)$$

2 *Backward LSTM* : Processes the sequence from the end to the beginning

$$\overleftarrow{h}_t = LSTM_{\{backward\}}(x_t)$$

3 *Concatenation output* : Combines bidirectional representations

$$h_t = [\vec{h}_t \oplus \overleftarrow{h}_t]$$

2.6.3. Convolutional Neural Network (CNN)

CNN is used to extract local features from text through convolution operations and the ReLU activation function, as shown in the following equation :

$$c_i = f(W \cdot x_{\{i:i+h-1\}} + b) \quad (1)$$

Where $x_{\{i:i+h-1\}}$ denotes the word vector, W represents the convolution filter, and f is the activation function.

Table 3 Algorithm CNN Classifications

Algorithm 3 : *Convolutional Neural Network (CNN) Classifications*

Output : predicted class C_{map}

- 1 W : convolution filter weights
- 2 $x_{\{i:i+h-1\}}$: input vector from position $-i$ to $i + h - 1$
- 3 b : bias
- 4 $f(\cdot)$: activation function (*ReLU*)
- 5 h : filter size (window size)

2.7. Model Training and Evaluation Process

The model training process was carried out by applying *Early Stopping* to prevent *overfitting*. Model performance was evaluated using *accuracy*, *precision*, *recall*, and *F1-score* metrics, which are standard evaluation metrics for text classification problems.

2.8. Dashboard Visualization and Clustering Analysis

The model evaluation results are visualized in the form of a Streamlit-based dashboard. In addition to classification, this study also applies the *K-Means* Clustering algorithm to identify patterns in the distribution of mail archive data that cannot be observed manually. The main tools used in this study include Python, TensorFlow, and Visual Studio Code on a Windows operating system environment.

3. RESULTS

3.1. Initial Data Exploration and Dataset Statistics

The dataset used in this study consists of 2,817 incoming letters that have undergone data validation and cleaning. All data is free of empty values, duplications, and text anomalies that could potentially interfere with the deep learning model training process. The dataset has 18 target categories representing various types of official letters such as *Surat Dinas*, *SPK* dan *PKS*, *Laporan Nota*, *Kesepahaman*, and *Risalah*. This initial exploration confirms that the dataset is suitable for training and evaluating deep learning models for multi-class mail classification. This is shown in Figure 2.

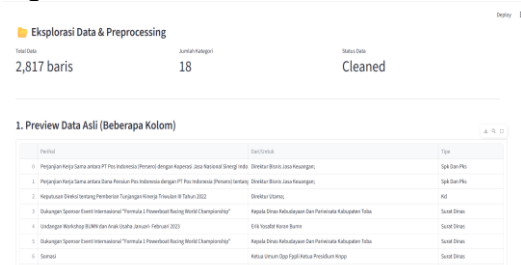


Figure 2 Data Exploration & Preprocessing

3.2. Preprocessed Data Results

Preprocessing is performed to convert raw text into a structured format that is ready to be processed by the model. This process includes text normalization, removal of non-alphanumeric symbols, and merging of several text attributes into one main column. The results can be seen in Figure 3.

2. Preview Data Hasil Preprocessing

Teks_input_Gabungan	Kategori_target
0 keluar perjanjian kerja sama antara pt pos indonesia persero dengan koperasi jasa nasional	SPK DAN PKS
1 keluar perjanjian kerja sama antara dana pensiun pos indonesia dengan pt pos indonesia	SPK DAN PKS
2 keluar keputusan direkti tentang pemberian tunjangan kinerja triwulan iii tahun 2022 dire	KD
3 masuk dukungan sponsor event internasional formula 1 powerboat racing world champio	SURAT DINAS
4 masuk undangan workshop bumh dan anak usaha januari february 2023 erik yosafat koran	SURAT DINAS
5 masuk dukungan sponsor event internasional formula 1 powerboat racing world champio	SURAT DINAS

Figure 3 Preview Preprocessed Data Results

3.3. Data Distribution by Letter Category

The distribution of data by letter category shows significant differences in the number of samples across classes as illustrated in Figure 4. The *Surat Dinas* category dominates the dataset, followed by *SPK* and *PKS*, while several other categories contain relatively fewer samples. This class

imbalance reflects real administrative conditions and presents an additional challenge for the classification task, particularly for models that are sensitive to skewed data distributions.

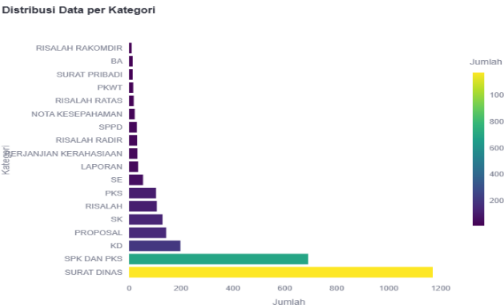


Figure 4 Data Distribution by Category

3.4. Deep Learning Model Evaluation Results

Model performance evaluation was conducted using 20% of the test data from the total dataset. The evaluation results show that the CNN model obtained the highest accuracy, precision, recall, and F1-score values compared to the LSTM and Bi-LSTM models, as shown in Table 4 and Figure 5. This superior performance indicates that CNN is more effective in capturing local and discriminative textual patterns through convolutional operations, which are common in structured and repetitive administrative mail documents. In contrast, sequence-based models such as LSTM and Bi-LSTM primarily focus on long-term dependencies, which are less dominant in formal correspondence texts.

Table 4 Accuracy Results Table

Model	Acc.	Prec.	Rec.	F1
LSTM	0.80	0.78	0.80	0.79
Bi-LSTM	0.78	0.74	0.78	0.75
CNN	0.85	0.84	0.85	0.83

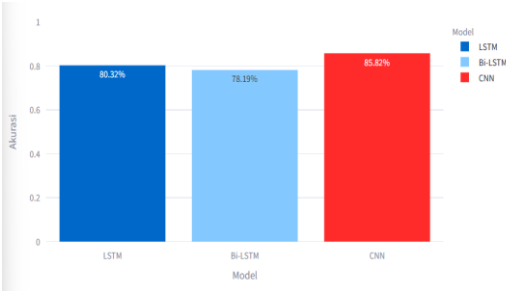


Figure 5 Accuracy Comparison (Training Results)

3.5. Confusion Matrix Analysis

The confusion matrix visualizations in Figures 6-8 illustrate the distribution of correct predictions and classification errors for each model. Overall, the CNN model demonstrates a higher concentration of correct predictions along the main diagonal, indicating stronger classification consistency across categories. Most misclassifications occur between semantically similar letter types, such as *SPK-PKS* and *Surat Dinas-Proposal*, which share overlapping vocabulary and structural patterns. This observation suggests that classification errors are primarily influenced by content similarity rather than model instability.

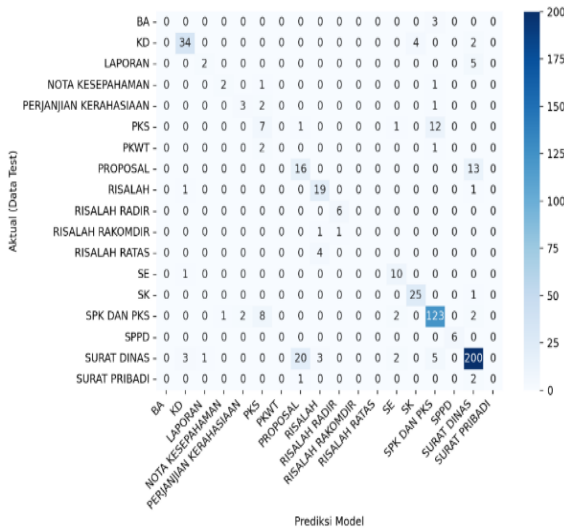


Figure 6 Confusion Matrix Model LSTM

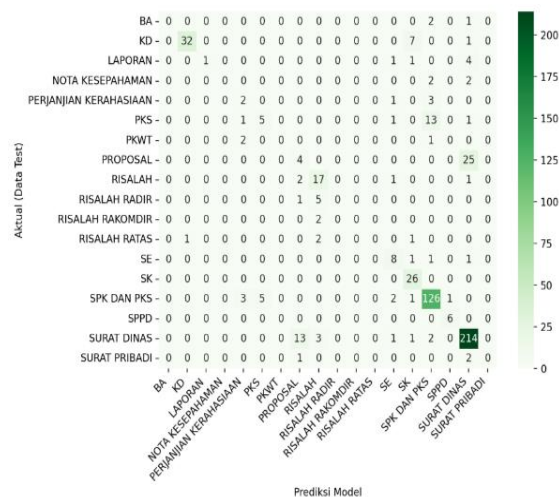


Figure 7 Confusion Matrix Model Bi-LSTM

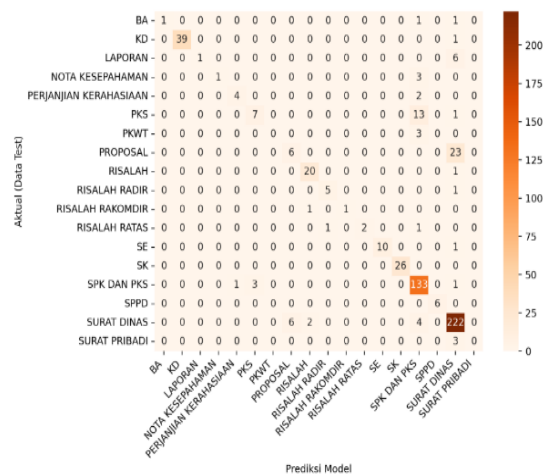


Figure 8 Confusion Matrix Model CNN

3.6. Feature Correlation Analysis Results

A feature correlation analysis using the *TF-IDF* and *Chi-Square* methods produced a list of dominant keywords that distinguish each category of letter, as shown in Figure 9 and Figure 10.

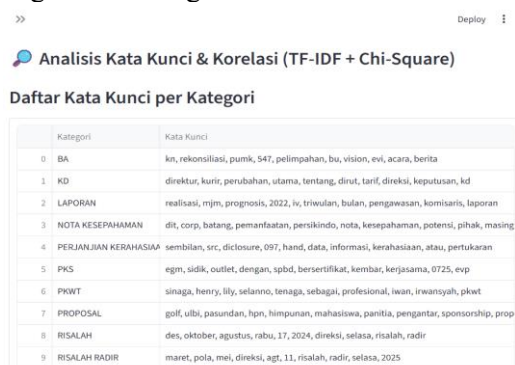


Figure 9 List of Keywords by Category

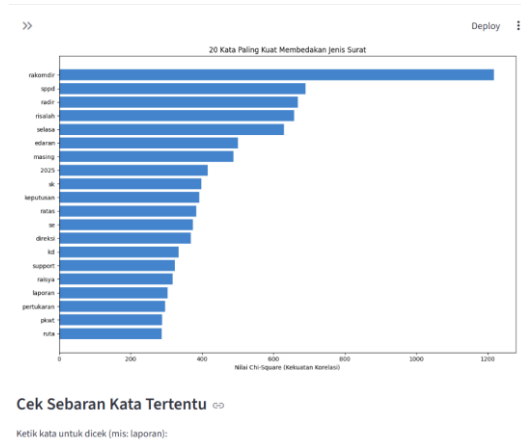


Figure 10 Bar Chart of the Strongest Words Distinguishing Types of Letters

3.7. K-Means Clustering Results

The clustering process using the *K-Means* algorithm groups mail archive data into several clusters based on textual similarity, as shown in Figures 11-17. This clustering analysis is not intended to replace the classification process, but rather to provide complementary insight into the inherent structure and thematic distribution of the dataset. By examining cluster composition and dominant keywords, this analysis helps explain why certain letter categories exhibit overlapping features, which may contribute to classification challenges observed in the confusion matrix.

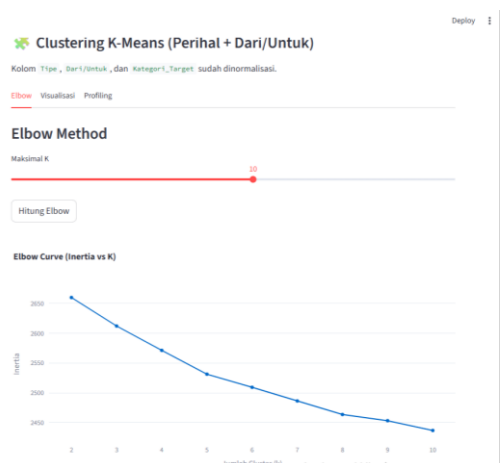


Figure 11 Determining the Number of Clusters (Elbow Method)

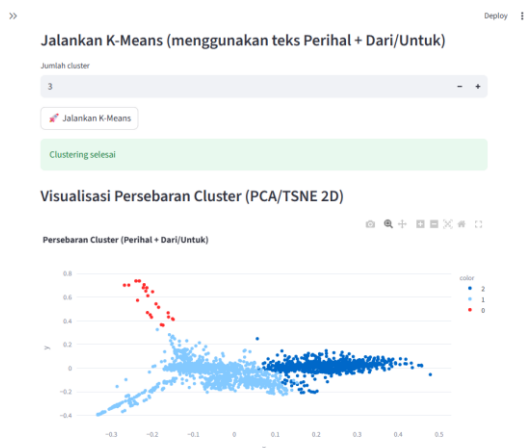


Figure 12 Cluster Distribution Visualization (PCA/TSNE 2D)

Kata Kunci Per Cluster (Centroid / Top terms)

Cluster	Keywords
0	0 karyadi, div, hukum, spk, corsec, pks, keluar, secretary, corporate, legal
1	1 masuk, direksi, keluar, dan, direktur, indonesia, risalah, utama, 2023, permohonan
2	2 pt, indonesia, pos, persero, antara, dengan, perjanjian, kerja, tentang, business

Figure 13 Keywords per Cluster (Centroid/Top terms)

Cluster	Jumlah Data	Tipe Dominan	Pengirim Dominan
0	0	75 Spk Dan Pks	Corsec;
1	1	1961 Surat Dinas	Direktur Utama;
2	2	781 Spk Dan Pks	Direktur Bisnis Jasa Keuangan;

Figure 14 Cluster Summary

A. Distribusi Tipe Surat per Cluster

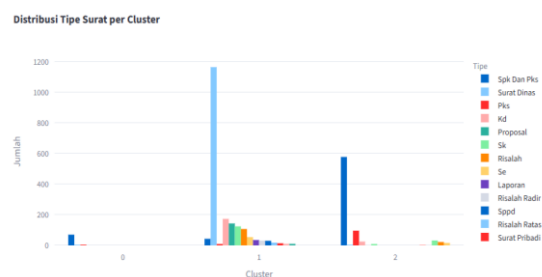


Figure 15 Letter Type Distribution per Cluster

B. Distribusi Pengirim (Dari/Untuk) per Cluster

Terdapat banyak entitas 'Dari/Untuk'. Menampilkan top 20 berdasarkan frekuensi.

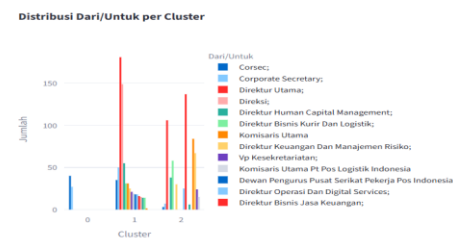


Figure 16 Distribution of Senders (From/To) per Cluster

Keanggotaan Cluster (Isi Setiap Kelompok)

Pilih Cluster untuk melihat anggotanya

0

Jumlah anggota Cluster 0: 75

	Perihal	Dari/Untuk	Tipe	Cluster
141	SPK	Karyadi Kumpo;	Spk Dan Pks	0
142	SPK	Karyadi Kumpo;	Spk Dan Pks	0
341	SPK/PKS	Div Hukum (Karyadi);	Spk Dan Pks	0
343	PKS	Div Hukum (Karyadi);	Spk Dan Pks	0
382	SPK/PKS	Div Hukum (Karyadi);	Spk Dan Pks	0
383	spk/PKS	Div Hukum (Karyadi);	Spk Dan Pks	0
614	SPK	Karyadi Kumpo;	Spk Dan Pks	0
621	PKS (Karyadi-kumpo)	Corporate Secretary;	Spk Dan Pks	0
737	SPK/PKS	Karyadi (Kumpo);	Spk Dan Pks	0

Figure 17 Cluster Membership

3.8. Classification Trial (Manual Input)

The developed system provides a classification trial feature through manual input and new data uploads. The results of the system testing are shown in Figure 18 and Figure 19.

Uji Coba Klasifikasi (Input Manual)

Jenis Surat:

Perihal:

Dari/Untuk:

LSTM

SURAT DINAS

96.4% confidence

Bi-LSTM

SURAT DINAS

93.8% confidence

CNN

SURAT DINAS

99.0% confidence

Figure 18 Classification Trial (Manual Input)

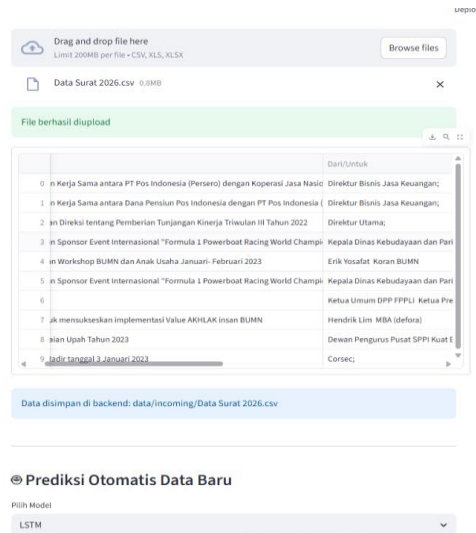


Figure 19 Upload New Data

4. DISCUSSION

4.1. Discussion of Deep Learning Model Evaluation

Based on the evaluation results, the *Convolutional Neural Network* (CNN) model demonstrates the most optimal performance with an accuracy of 85.82%, as well as higher precision, recall, and F1-score values compared to *LSTM* and *Bi-LSTM*. These results indicate that CNN is more effective in extracting local features and word patterns from administrative texts that have relatively consistent structures.

This superior performance can be attributed to the convolutional mechanism employed by CNN, which applies fixed-size filters (kernel size = 5) to capture local contextual patterns within text sequences. Such patterns are particularly prominent in administrative mail documents, where standardized phrases and recurring expressions are commonly used. By focusing on local feature extraction, CNN is able to identify discriminative textual cues efficiently without relying on long sequential dependencies, which are more emphasized in LSTM-based models.

These findings are consistent with previous studies that reported the superiority of CNN-based models for document classification tasks involving structured and repetitive textual patterns.

4.2. Discussion of Confusion Matrix

Confusion matrix analysis provides a more detailed picture of the classification performance of each model in each letter category. The CNN model shows a higher number of correct predictions in almost all categories, with relatively fewer classification errors compared to *LSTM* and *Bi-LSTM*.

The results show that CNN has better class separation capabilities, especially in letter categories with similar language contexts. The error patterns that still appear generally occur in categories with overlapping text characteristics, indicating the natural challenges in administrative document classification.

4.3. Discussion of Feature Correlation Analysis

The results of feature analysis using *TF-IDF* and *Chi-Square* show that each letter category has dominant keywords that are distinctive and relevant to its administrative context. Words with high *Chi-Square* values are proven to play an important role in distinguishing one type of letter from another.

This finding reinforces the classification evaluation results, as categories with clear and consistent keywords tend to have more accurate prediction rates. In addition, this analysis provides interpretability to the model, so that the classification process is not only numerical, but also semantically understandable.

4.4. Discussion of K-Means Clustering

The clustering results using the *K-Means* algorithm show that the letter archive data naturally forms several main groups based on the similarity of the text content. The *PCA/TSNE* visualization reinforces the Elbow Method results, which indicate that $k = 3$ is the optimal number of clusters for representing the thematic structure of the dataset.

Using more than three clusters tends to increase grouping complexity without revealing additional meaningful semantic

patterns, indicating that the archive data exhibits a relatively stable and well-defined structure. In this study, clustering is not intended to replace the classification process, but rather to serve as a complementary analytical tool to help explain similarities between letter categories and overlapping textual features observed in the confusion matrix analysis.

4.5. Discussion of System Implementation

The results of system testing through manual input and new data upload show that the developed system is applicable and can be used on real data. The automatic preprocessing process ensures consistency between the training data and the new data being tested.

The *confidence* value displayed in the prediction results represents the model's level of confidence in a particular input and cannot be equated with the overall model evaluation accuracy value. This finding confirms that the system is designed not only for experimental purposes but also to support the practical process of classifying mail archives in the work environment.

5. CONCLUSION

This study demonstrates that deep learning approaches are capable of effectively classifying unstructured and imbalanced administrative mail archives. Among the evaluated models, namely LSTM, Bi-LSTM, and CNN, the CNN model achieved the best overall performance with an accuracy of 85.82%, precision of 84.04%, recall of 85.82%, and an F1-score of 83.39. This result indicates that CNN is more effective in capturing local textual patterns and recurring phrase structures commonly found in formal administrative documents.

In addition to classification performance, this study highlights the importance of interpretability and usability through the integration of feature correlation analysis, K-Means clustering, and an interactive dashboard. Feature correlation analysis provides semantic insight into dominant keywords for each letter category,

while clustering analysis reveals the natural structure and similarity patterns within the dataset. The developed dashboard further supports practical usage by enabling users to evaluate model performance and perform classification on new data in an intuitive manner.

Despite these contributions, this study has several limitations. The deep learning architectures employed are relatively conventional, and class imbalance remains a challenge that affects minority categories. These limitations indicate opportunities for further methodological improvements and extensions in future work.

6. RECOMMENDATIONS

Future research is recommended to explore more advanced language models, such as transformer-based architectures including IndoBERT, to further improve classification performance and contextual understanding of administrative texts. Addressing class imbalance through resampling or cost-sensitive learning techniques is also suggested to enhance model robustness, particularly for underrepresented letter categories.

Moreover, the interactive dashboard developed in this study can be extended through usability evaluation and deployment in real operational environments. Testing the system with larger and continuously updated datasets would provide valuable insight into its scalability and effectiveness in real-world archival management scenarios.

REFERENCES

- [1] F. Nikmah, D. J. Pribadi, E. A. Sukma, E. Suwarni, and I. Azmi, "Decision Support System For Handling Incoming And Outgoing Mail: To Facilitate Archives Retrieval," *Int. J. Acad. Res. Reflect.*, vol. 10, no. 3, pp. 53–61, 2022, [Online]. Available: www.idpublications.org
- [2] J. F. Indey and S. Supangat, "Implementasi Algoritma Naïve Bayes Dalam Sistem Pengarsipan Surat Berbasis Ai Di Gpi Papua Klasik Mimika Papua Tengah," *J. Ilm. Inform.*, vol. 12, no. 2, pp. 102–113, 2024, doi:

- 10.33884/jif.v12i02.9087.
- [3] F. R. Gerung, G. D. P. Maramis, and E. R. S. Moningkey, "Penerapan Algoritma Naive Bayes pada Arsip Surat," *Remik Ris. dan E-Jurnal Manaj. Inform. Komput.*, vol. 9, no. 2, pp. 676–690, 2025, doi: 10.33395/remik.v9i2.14786.
- [4] R. A. Krisdiawan, "Implementasi Model Pengembangan Sistem GDLC dan Algoritma Linear Congruential Generator Pada Game Puzzle," *Nuansa Inform.*, vol. 12, no. 2, pp. 1–9, 2018, doi: 10.25134/nuansa.v12i2.1634.
- [5] G. Mujtaba, L. Shuib, R. G. Raj, N. Majeed, and M. A. Al-Garadi, "Email Classification Research Trends: Review and Open Issues," *IEEE Access*, vol. 5, pp. 9044–9064, 2017, doi: 10.1109/ACCESS.2017.2702187.
- [6] Y. Kim, "Convolutional neural networks for sentence classification," in *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2014, pp. 1746–1751. doi: 10.3115/v1/d14-1181.
- [7] R. A. Pakpahan and D. A. Utami, "Pengelolaan Arsip Surat Masuk Di Dinas Energi Dan Sumber Daya Mineral Provinsi Jawa Timur Melalui Aplikasi Tata Naskah Dinas Elektronik (TNDE)," *J. Inov. Negara dan Terap.*, vol. 4, no. 1, pp. 211–222, 2025, [Online]. Available: <https://journal.unesa.ac.id/index.php/innovant/article/download/38747/12867/125823#:~:text=Uraian diatas dapat disimpulkan bahwa,elektronik dengan proses penciptaan dan>
- [8] D. Z. Robert, A. B. Setiawan, and D. W. Widodo, "Implementasi metode support vector machine untuk klasifikasi otomatis pengaduan publik di Kabupaten Trenggalek," *Inotek J. Inov. Teknol.*, vol. 9, p. 1617, 2025.
- [9] S. Pandey, A. Taralekar, R. Yadav, S. Deshmukh, and P. S. Suryavanshi, "E-mail Spam Detection and Classification using SVM," *Int. J. Comput. Sci. Inf. Technol.*, vol. 10, no. 1, pp. 6–8, 2020, [Online]. Available: <http://www.ijcsit.com>
- [10] I. Kasim and Sahibu, "Klasifikasi Surat Digital Menggunakan Algoritma Naïve Bayes," *J. IT Media Inf. IT STMIK Handayani*, vol. 13, no. 2, pp. 57–62, 2020.
- [11] C. Magnolia, A. Nurhopipah, and B. A. Kusuma, "Penanganan Imbalanced Dataset untuk Klasifikasi Komentar Program Kampus Merdeka Pada Aplikasi Twitter," *Edu Komputika J.*, vol. 9, no. 2, pp. 105–113, 2022, doi: 10.15294/edukomputika.v9i2.61854.
- [12] W. N. Ibrahim Al-Obaydy, H. A. Hashim, Y. AbdulKhaleq Najm, and A. A. Jalal, "Document classification using term frequency-inverse document frequency and K-means clustering," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 27, no. 3, pp. 1517–1524, 2022, doi: 10.11591/ijeecs.v27.i3.pp1517-1524.
- [13] A. Liani, "Analisis Perbandingan Kernel Algoritma Support Vector Machine dalam Mengklasifikasikan Skripsi Teknik Informatika berdasarkan Abstrak," *JOINS (Journal Inf. Syst.)*, vol. 5, no. 2, pp. 240–249, 2020, doi: 10.33633/joins.v5i2.3715.
- [14] J. and others MacQueen, "Some methods for classification and analysis of multivariate observations," in *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*, 1967, pp. 281–297.
- [15] F. Martinez-Plumed *et al.*, "CRISP-DM Twenty Years Later: From Data Mining Processes to Data Science Trajectories," *IEEE Trans. Knowl. Data Eng.*, vol. 33, no. 8, pp. 3048–3061, 2021, doi: 10.1109/TKDE.2019.2962680.
- [16] M. R. Givari, M. R. Sulaeman, and Y. Umaidah, "Perbandingan Algoritma SVM, Random Forest Dan XGBoost Untuk Penentuan Persetujuan Pengajuan Kredit," *Nuansa Inform.*, vol. 16, no. 1, pp. 141–149, 2022, doi: 10.25134/nuansa.v16i1.5406.
- [17] M. A. Senubekti and L. A. Puspita Dewi, "Prinsip Klasifikasi Dan Data Mining Dengan Algoritma C4.5," *Nuansa Inform.*, vol. 16, no. 2, pp. 87–93, 2022, doi: 10.25134/nuansa.v16i2.5834.

Real-Time Bayesian Knowledge Tracing for Adaptive Vocabulary Practice in an Adventure Educational Game: Architecture, Verification, and Reproducibility

Ikhsan Khaeruddin¹, Rio Andriyat Krisdiawan^{*2}, Nida Amalia Nasikin³

^{1,2,3}*Teknik Informatika, Fakultas Ilmu Komputer Universitas Kuningan, Indonesia*

E-mail: ¹20210810111@uniku.ac.id, ^{*2}rioandriyat@uniku.ac.id, ³nida.amalia.asikin@uniku.ac.id

Abstract

Vocabulary learning in primary English classes is often constrained by heterogeneous learner readiness, where fixed game progression can under-serve both struggling and advanced students. This study aims to implement and technically validate a reproducible real-time adaptivity mechanism using Bayesian Knowledge Tracing (BKT) in a Unity-based Android vocabulary game. The research adopts a design-and-verification approach using the Game Development Life Cycle (GDLC), supported by requirements elicitation (classroom observation, teacher interview, and literature review). The adaptive engine applies BKT to update mastery after each quiz response and routes learners using a mastery-threshold policy, while event-level logs are stored locally and exportable for auditability. The main results demonstrate that adaptive mode activation, mastery updates, persistence of adaptive state, and mastery-gated progression function consistently in end-to-end black-box tests. Algorithm-level credibility is strengthened through white-box basis-path verification of the UpdateProbability() routine, ensuring independent execution paths for correct and incorrect responses are covered. This work contributes a deployable Unity/Android architecture for real-time BKT-driven adaptivity, accompanied by verification artifacts and reproducibility recommendations to support technical audit, replication, and subsequent controlled effectiveness studies.

Keywords— Bayesian Knowledge Tracing (BKT), Adaptive learning, Educational game, GDLC, Software verification.

Submission: December 08, 2025 Accepted: January 20, 2026

Published: January 20, 2026

DOI Article: <https://doi.org/10.25134/ilkom.v20i1.535>

1. INTRODUCTION

Vocabulary knowledge is a fundamental constraint in early English learning; limited lexical coverage quickly becomes a bottleneck for comprehension and classroom participation [1]. Digital game-based language learning has expanded because games can increase engagement and provide repeated exposure, and meta-analytic evidence indicates that game-based approaches can yield measurable learning benefits when aligned with instructional goals [1], [2]. However, primary classrooms are heterogeneous: learners vary in prior vocabulary, acquisition pace, and error patterns. Under these conditions, one-size-fits-all game progression may under-serve both struggling learners and those who advance faster.

A key research and development gap lies in **how adaptivity is engineered and reported**. While learner modeling is central to personalized learning [3], many educational games either (i) remain non-adaptive or (ii) describe adaptivity at a high level without a reproducible specification

of mastery updates, parameterization, decision thresholds, and verification artifacts. Even when Bayesian Knowledge Tracing (BKT) is adopted, implementation details that determine validity—such as initialization, slip/guess settings, and update logic are often under-documented, limiting replication and technical auditability in real-time game environments [4], [5].

This paper positions itself as an engineering-focused contribution. We specify and implement a real-time BKT loop inside a Unity-based Android vocabulary game and provide verification evidence that runtime behavior matches the stated model. By consolidating algorithmic specification, parameter configuration, and testable decision policies, the work supports reproducible deployment and provides a foundation for future comparative studies.

Main contributions

1. A reproducible Unity/Android architecture for real-time BKT-driven adaptivity.
2. A single-source-of-truth BKT specification (parameters, initialization, update equations, mastery rule).

3. Verification artifacts (functional tests and logic/path-level checks) supporting auditability and replication.
4. An implementation blueprint extensible to controlled effectiveness studies and learning analytics.

2. RESEARCH METHOD

2.1. Research Design

This study employs a design-and-verification approach to develop and technically validate an Android educational game artifact. The primary outcome is a reproducible real-time adaptivity pipeline from learner response capture to Bayesian Knowledge Tracing (BKT) updates and adaptive content routing rather than a causal effectiveness claim. The development process follows the Game Development Life Cycle (GDLC) model [6], [7], [8].

2.2. Elicitation and Learning-Content Preparation

Requirements were elicited via:

1. **Classroom observation** to identify common Grade III learner difficulties in basic English vocabulary and interaction patterns relevant to gameplay (e.g., response behavior, attention span).
2. **Semi-structured teacher interview** to validate instructional needs, determine target vocabulary scope, and agree on feasible practice formats.
3. **Literature review** to ground the adaptivity mechanism in established learner modeling and to ensure methodological consistency [3], [5], [9], [10].

A vocabulary item bank was prepared and organized into thematic sets and difficulty tiers. Each item is assigned: (a) skill/topic label, (b) difficulty tier, and (c) correct answer key. Teacher input is used to validate age appropriateness and content coverage.

2.3. System Development Using GDLC

The system is developed using GDLC phases: **initiation, pre-production, production, testing, and release** [6], [7], [8].

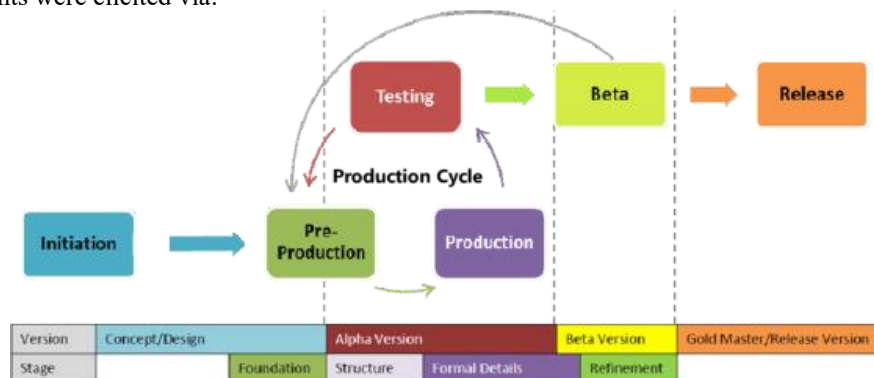


Figure 1. Game Development Life Cycle Guidelines[6]

1. **Initiation:** define functional requirements (login/session, practice loop, scoring/feedback, logging) and constraints (Android performance, offline capability).
2. **Pre-production:** produce Game Design Document (GDD), UI/scene flow, and the adaptivity policy (threshold and progression rules).
3. **Production:** implement the game in Unity (C#), integrate assets, implement item-bank

manager, logging module, and BKT service/component.

4. **Testing:** execute functional verification and structural verification (Section 2.6).
5. **Release:** package Android build for limited acceptance use focused on runtime stability and deployment feasibility.

2.4. System Architecture Overview

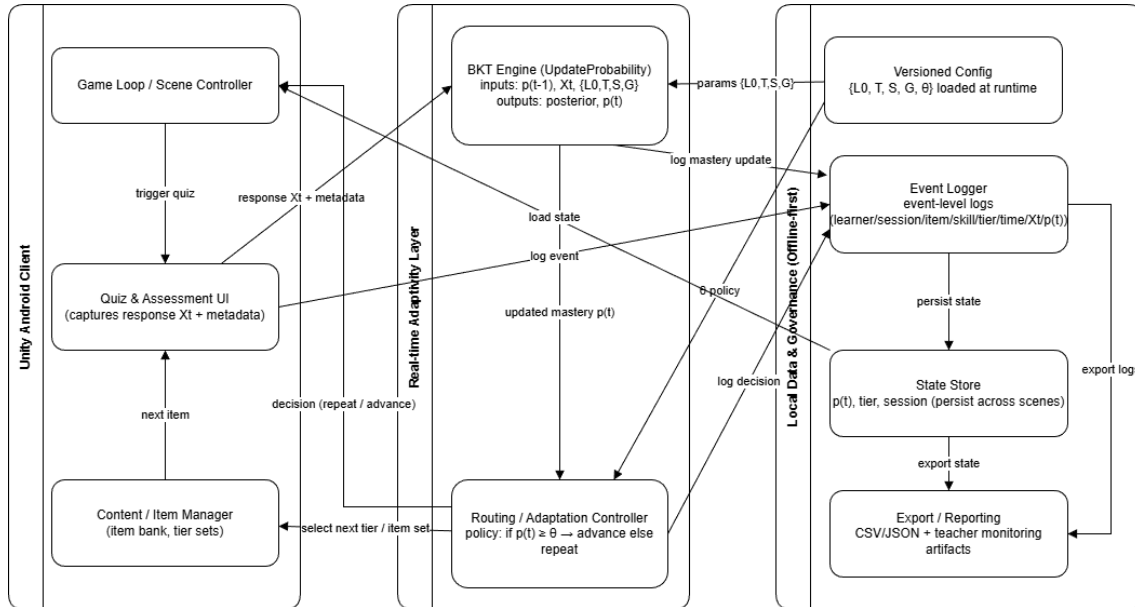


Figure 2. High-Level System Architecture of the BKT-Based Adaptive Game System

Figure 2 summarizes the runtime adaptivity pipeline. The Unity game loop triggers quiz events: learner responses are processed by the BKT engine to update mastery; the routing controller applies the threshold policy to select repetition or progression; and event logs plus adaptive state are stored locally and exported for reporting and reproducibility.

2.5. Bayesian Knowledge Tracing (BKT)

The adaptive engine implements standard BKT to estimate learner mastery probabilistically per skill/topic [3], [5], [9], [10]. For each learner–skill pair, BKT uses four parameters: initial mastery $P(L_0)$, transition/learning $P(T)$, slip $P(S)$, and guess $P(G)$.

Let $p_{t-1} = P(L_{t-1})$ be prior mastery and $X_t \in \{0,1\}$ be the observed response (1=correct, 0=incorrect). Posterior mastery is computed by Bayes' rule:

If $X_t = 1$ (correct):

$$P(L_{t-1} | X_t = 1) = \frac{p_{t-1}(1 - S)}{p_{t-1}(1 - S) + (1 - p_{t-1})G}$$

If $X_t = 0$ (incorrect):

$$P(L_{t-1} | X_t = 0) = \frac{p_{t-1}S}{p_{t-1}S + (1 - p_{t-1})(1 - G)}$$

Then learning transition is applied:

$$p_t = P(L_{t-1} | X_t) + (1 - P(L_{t-1} | X_t)) \cdot T$$

Table 1. BKT Parameters

Parameter	Meaning	Impact on updates
(P(L_0))	Early probability of mastering a skill	Determine start mastery
(T)	Probability of learning after one chance	Control mastery ascension
(S)	Slip (wrong despite mastery)	Lowering posterior when wrong
(G)	Guess (true even if you haven't mastered it)	Posterior holding when correct

2.6. Real-Time Adaptation Policy and Parameter Governance

Parameter governance (anti inkonsistensi naskah–kode): all BKT parameters and mastery threshold θ are stored in a **single configuration source** (table/file) loaded at runtime (not hard-coded). Configuration changes are versioned, allowing the paper to reference one definitive set. **Adaptation rule:** the game applies a threshold policy: when $p_t \geq \theta$, the learner is routed to a higher difficulty tier or next content segment; otherwise, additional practice is delivered within the current tier. To prevent oscillation, an optional **window rule** can be applied (e.g., threshold must be met for a minimum number of recent items).

If calibration is required, pilot interaction traces can be used to estimate parameters using transparent tooling such as **pyBKT** while retaining clear reporting of the final parameter set [11].

2.7. Data Logging and Verification Strategy

Event-level logging is implemented for auditability and replication. Each interaction stores pseudonymous learner/session ID, item ID, skill/topic ID, timestamp, correctness, current mastery p_t , and routing decision.

Table 2. Log Schema for Reproducibility

Field	Description
learner_id (pseudonymous)	Anonymous identity of players
session_id	Play session identity
item_id	Question/item ID
skill_id/topic_id	Tracked skill labels
difficulty_tier	The difficulty level when the item is awarded
timestamp	Interaction time
response (0/1)	False/true
prior_mastery (p_{t-1})	Mastery before update
posterior_mastery	Bayes results before transition
updated_mastery (p_t)	Mastery after transition
routing_decision	Up/fixed/next segment

2.8. Verification and Validation

Verification and validation prioritize evidence for the paper's architecture + verification claim:

- (a) **Black-box functional verification:** test cases validate end-to-end flows (scene transitions, item rendering, answer submission, logging, BKT update invocation, threshold checking, routing).
- (b) **White-box basis-path verification:** the function `UpdateProbability()` and routing controller are tested using basis-path coverage guided by cyclomatic complexity to ensure each independent decision path is executed [12].

3. RESEARCH RESULT

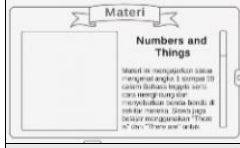
3.1. Game Flow and Storyboard (Functional Scenes)

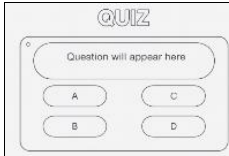
The developed artifact implements a structured learning loop consisting of **content exposure** → **adaptive assessment** → **mastery-gated progression**. The loop begins from the main menu, continues to vocabulary material access, activates adaptivity via a pretest, and then executes repeated quiz-driven updates during staged gameplay. Table 3 summarizes the storyboard at the level required for an engineering paper focusing on scene purpose and outputs that are relevant to adaptivity and verification.

Table 3. Storyboard and adaptivity-relevant outputs

Scene	Core function	Output relevant to adaptivity / verification
Main Menu	Start and navigation	Entry to learning pipeline; controlled access to gameplay modules
Material	Vocabulary exposure	Content gate prior to adaptive activation; ensures exposure before assessment
Pretest	Initialize learner state	Establish initial mastery/level; triggers adaptive mode and stores state
Gameplay (Adventure stage)	Engagement loop	Context for assessment triggers; stage completion depends on routing decision
Quiz	Assessment + mastery update	Per-item BKT update; produces updated mastery used for routing
Achievement	Feedback	Displays progression outcome consistent with routing logic

Table 4. Storyboard Game Design

Scene	Visual	Action
	Play, Setting, Info	Start Game
	Vocabulary material	Can't be skipped without interaction
	Moving characters	Avoiding enemies & items
	Initial level	Defining initial capabilities

	Adaptive	BKT regulates difficulties
	Badge	Player achievements

3.2. User Interface Implementation

User interfaces are implemented for two roles: **student-facing gameplay** and **teacher-facing monitoring**. For students, the implemented scenes include the main menu, system information, login, vocabulary material, platformer gameplay, quiz interaction, and achievement summary (Figures 3–9). For teachers, the system provides restricted access to monitoring and reporting functions, enabling result viewing and report export to support classroom deployment and traceability (Figure 9).

- **Figure 3–9** present the student-facing workflow (menu → material → gameplay → quiz → achievement).
- **Figure 10** shows the teacher dashboard interface for monitoring/reporting.



Figure 3. Main Menu View



Figure 4. System Information Display

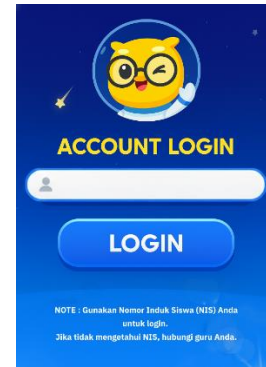


Figure 5. Game Login Menu View

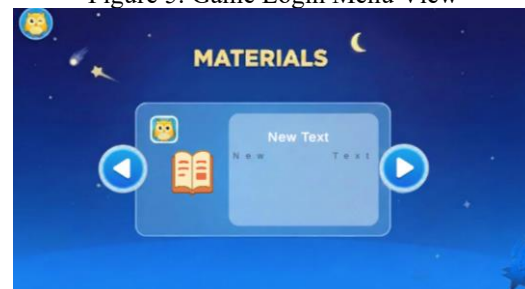


Figure 6. Material Menu View

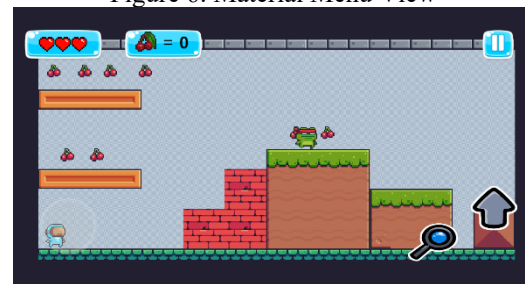


Figure 7. Playing Display



Figure 8. Quiz Menu View



Figure 9. Achivement Menu View

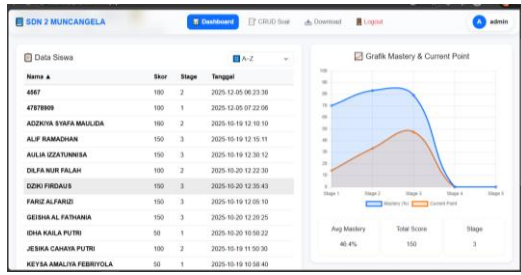


Figure 10. Teacher Dashboard Interface

3.3. Real-Time BKT Adaptivity Implementation

Real-time adaptivity is implemented at the granularity of each quiz interaction. The runtime loop is executed as follows:

1. The system presents a quiz item aligned with the current stage/difficulty tier.
2. The learner submits a binary response $X_t \in \{0,1\}$ (correct/incorrect).
3. The BKT engine computes posterior mastery conditioned on X_t and applies the learning transition to obtain the updated mastery probability $P(L_t)$ (as defined in Section 2.4).
4. The adaptation router compares $P(L_t)$ against a mastery threshold θ and selects the next action: **repeat within current tier** or **advance to a higher tier/next stage**.

At the implementation level, the mastery variable is stored as an internal state (e.g., $pKnown$), and the learning/transition parameter is represented as $pLearn$, corresponding to $P(L)$ and $P(T)$ in the BKT formulation. The transition update follows:

$$pKnown \leftarrow posterior + (1 - posterior) \cdot pLearn$$

This design enables immediate mastery updating and routing decisions during runtime, supporting **real-time adaptivity** rather than post-hoc personalization.

Reproducibility requirement (single-source-of-truth parameters). To support independent audit and replication, all parameters $\{P(L_0), P(T), P(S), P(G), \theta\}$ are defined in a versioned configuration source loaded at runtime (not hard-coded). The evaluated build is therefore traceable to one definitive parameter set, reducing the risk of manuscript–implementation divergence.

3.4. System Testing Results (Black-Box Functional Verification)

Black-box functional verification was conducted to confirm correct end-to-end behavior across navigation, content gating, adaptive-mode activation, persistence, and mastery-gated

progression. While the complete test suite covers broader usability and flow control, Table 5 reports the subset of cases that directly evidence the correctness of the adaptive pipeline: adaptive initiation, response-level mastery updating, persistence of adaptivity state, and deterministic routing based on θ .

Table 5. Adaptivity-critical black-box verification (selected cases)

Case ID	Scenario verified	Trigger / test case	Expected observable output	Result
BB-03	Adaptive mode activation	In Material, confirm readiness	Transition to Pretest and adaptive mode active	Pass
BB-05	Mastery update responsiveness	Answer a pretest/quiz item correctly	Mastery increases relative to baseline	Pass
BB-06	Persistence of adaptivity decision	Complete pretest	Difficulty/level state is stored (persisted)	Pass
BB-08	Failure routing below threshold	End stage with mastery ($< \theta$)	“Retry” state/panel appears	Pass
BB-09	Success routing at/above threshold	End stage with mastery ($\geq \theta$)	“Success” state/panel appears	Pass
BB-11	Progression continuity	Continue after success	Transition to next stage	Pass

These outcomes indicate that the adaptive mechanism is not only implemented, but also produces stable, observable, and persistent routing decisions aligned with mastery conditions.

3.5 White-Box Verification Results (Basis-Path for UpdateProbability())

White-box verification targets the BKT update routine (`UpdateProbability()`), as it is the algorithmic nucleus of adaptivity. Structurally, the routine contains a correctness-dependent decision branch (correct vs incorrect) followed by posterior computation and a learning

transition update. To validate internal control-flow coverage, basis-path testing was applied. Cyclomatic complexity indicates two independent paths corresponding to the correctness branch; therefore, adequacy is achieved by executing test inputs that force both branches to execute at least once.

Table 6. Basis-path coverage for
UpdateProbability()

Path ID	Input condition	Independent path	Branch covered	Expected outcome (test oracle)
WB-P1	correct = true	Start → correct-branch → posterior → transition → return	Correct-answer posterior	Posterior uses ((1-S)); mastery updated via transition; returns updated mastery
WB-P2	correct = false	Start → wrong-branch → posterior → transition → return	Incorrect-answer posterior	Posterior uses (S); mastery updated via transition; returns updated mastery

his evidence strengthens auditability by demonstrating that the internal decision structure of the mastery update routine is explicitly covered, not merely assumed correct from user-level behavior.

Overall, the results demonstrate a deployable Unity/Android educational game artifact with a functioning real-time BKT adaptivity loop. Adaptivity is evidenced by (i) per-response mastery updating, (ii) mastery-threshold routing that produces deterministic retry/success outcomes, (iii) persistence of adaptive state across scenes, and (iv) verification artifacts at both functional (black-box) and structural (basis-path) levels. These results support the paper's engineering contribution: a reproducible architecture with testable adaptivity behavior suitable for technical audit and replication.

4 DISCUSSION

4.1 Principal Findings and Their Meaning

This study demonstrates that Bayesian Knowledge Tracing (BKT) can be embedded as a **real-time control mechanism** inside a Unity-based Android educational game. Mastery is updated after each quiz interaction and immediately used to govern routing decisions. Unlike many adaptive learning reports that describe personalization at a conceptual level, this work operationalizes adaptivity as an explicit runtime pipeline—response capture → mastery update → threshold decision → next-item selection—consistent with the role of knowledge tracing as a learner-state estimator for practice sequencing [3], [5], [9].

The contribution is intentionally framed as **technical validity and reproducibility**, not as an effective claim. Within this scope, the results indicate that the adaptive pipeline is (i) implemented end-to-end, (ii) observable at the system level (e.g., deterministic retry/success routing), and (iii) traceable through persisted learner state and reporting/export mechanisms.

4.2 Why This Work Matters: Reproducible Adaptivity in a Mobile Game Context

BKT is widely established, and recent syntheses underline its continued relevance while also noting recurring issues in reporting and comparability across implementations [10]. Many adaptive systems remain difficult to reproduce because manuscripts often omit critical implementation details: the exact parameter set, initialization procedure, mastery threshold logic, and verification artifacts needed for technical audit.

This paper positions its novelty in **implementation transparency** by treating the adaptivity mechanism as a reproducible system component rather than an implicit feature. This is particularly relevant in mobile game environments, where algorithmic behavior must remain stable under scene transitions, local persistence constraints, and practical runtime limitations.

4.3 Strength of Evidence: Complementary Verification Layers

The verification strategy directly supports credibility. **Black-box functional verification** confirms that the adaptive flow behaves as intended from the user perspective: adaptive mode activation, persistence of level decisions, and mastery-gated routing outcomes are observable and consistent. **White-box basis-path verification** strengthens this evidence by demonstrating that the internal decision structure

of the mastery update routine is structurally exercised across independent execution paths. Using cyclomatic complexity and independent-path execution as an adequacy argument aligns with established software testing principles [12]. Together, these layers reduce the risk that adaptivity is under-specified, difficult to audit, or inconsistent between narrative and runtime behavior.

4.4 Parameter Governance: A Non-Negotiable Condition for Reproducibility

A key lesson from the implementation is that **parameter consistency is inseparable from reproducibility**. In BKT, values for slip $P(S)$, guess $P(G)$, transition/learning $P(T)$, and the mastery threshold θ directly shape posterior updates and therefore routing behavior [3], [5], [9]. Even small deviations can materially alter mastery trajectories and progression decisions. For publication-grade reporting, the system should enforce a **single-source-of-truth parameter specification**: a versioned configuration (table/file) loaded at runtime and referenced explicitly in the manuscript. This ensures that the parameter set described in the paper matches the evaluated build. When parameters deviate from defaults (e.g., $P(G) = 1/k$ for multiple-choice items), the deviation should be justified either through domain rationale or calibration using interaction traces. In this context, using established estimation tooling (e.g., pyBKT) provides a transparent pathway for parameter fitting and reporting [11].

4.5 Limitations and Threats to Validity

Several limitations should be acknowledged to avoid over-claiming:

1. **Scope limitation (no causal learning claim).** The study validates system correctness and reproducibility; it does not establish that adaptivity improves learning outcomes relative to a non-adaptive baseline.
2. **Model assumptions.** Classic BKT uses a binary observation model and typically assumes a single latent mastery per skill without explicitly modeling forgetting or item difficulty heterogeneity. These assumptions can be challenged in vocabulary learning where item difficulty may vary substantially [3], [10].
3. **Parameter sensitivity and threshold fragility.** Fixed parameterization and a fixed mastery threshold can yield brittle routing if item quality varies or if learners adopt strategic guessing. Without calibration or

sensitivity analysis, robustness cannot be inferred.

4. **Verification boundaries.** Basis-path coverage is demonstrated for the mastery update kernel, but other modules (routing controller, persistence, export/sync) may also affect reproducibility. Functional verification supports end-to-end behavior; however, deeper structural testing of routing and persistence logic would further strengthen auditability.

4.6 Future Work

This study establishes a technical baseline that enables a structured research program:

1. **Parameter calibration and reporting:** estimate $P(L_0)$, $P(T)$, $P(S)$, and $P(G)$ per skill/tier from logged traces and publish the final parameter table with procedure details [11].
2. **Sensitivity analysis:** evaluate how routing outcomes change across plausible θ ranges and parameter perturbations to demonstrate robustness.
3. **Controlled comparative evaluation:** implement a non-adaptive baseline (fixed progression) and compare against BKT-driven adaptivity to quantify learning effects while controlling exposure time and novelty effects.
4. **Reproducibility package:** publish parameter configuration, pseudocode aligned with the evaluated build, and compact test suites (functional and structural) as supplementary material to support independent replication.

5 CONCLUSION

This paper presented a design-and-verification study that implements **real-time Bayesian Knowledge Tracing (BKT)** within a Unity-based Android educational game for vocabulary practice. The proposed system operationalizes adaptivity as a runtime pipeline in which each learner response triggers immediate mastery updating and the updated mastery probability governs deterministic routing decisions for practice continuation or progression.

The work contributes: (1) a deployable Unity/Android architecture integrating response capture, BKT updating, routing policy, and event logging; (2) an auditable adaptivity mechanism supported by persisted learner state and exportable reporting; and (3) verification evidence at two complementary levels—black-box functional validation of the adaptive flow and white-box basis-path verification of the

UpdateProbability() routine—supporting technical correctness and reproducibility. This study is intentionally scoped to technical validity and reproducibility rather than causal learning effectiveness. Future work should focus on (i) enforcing a single-source-of-truth parameter configuration aligned with the evaluated build, (ii) parameter calibration and sensitivity analysis to demonstrate routing robustness, and (iii) controlled comparative studies (adaptive vs non-adaptive) to evaluate whether real-time BKT adaptivity yields measurable learning gains beyond engagement effects.

6 SUGGESTIONS

1. Standardize and version all BKT parameters $\{P(L_0), P(T), P(S), P(G), \theta\}$ in a single configuration source aligned with the evaluated build to ensure reproducibility and prevent manuscript–implementation inconsistency.
2. Strengthen verification coverage by adding unit tests for routing control and persistence/export modules, complementing the existing basis-path testing of UpdateProbability().
3. Calibrate and stress-test the adaptive policy through pilot-log parameter calibration and sensitivity analysis on θ to confirm routing robustness under realistic learner behaviors.
4. Extend future evaluation with controlled comparison (adaptive vs non-adaptive) and richer learning indicators (e.g., retention and time-on-task) to assess effectiveness beyond technical validity.

REFERENCES

- [1] P. Wouters, C. van Nimwegen, H. van Oostendorp, and E. D. van der Spek, “A meta-analysis of the cognitive and motivational effects of serious games,” *J. Educ. Psychol.*, vol. 105, no. 2, pp. 249–265, May 2013, doi: 10.1037/a0031311.
- [2] D. H. Dixon, T. Dixon, and E. Jordan, “Second language (L2) gains through digital game-based language learning (DGBLL): A meta-analysis,” *Language Learning & Technology*, vol. 26, no. 1, pp. 1–25, Feb. 2022, doi: 10.64152/10125/73464.
- [3] R. Pelánek, “Bayesian knowledge tracing, logistic models, and beyond: an overview of learner modeling techniques,” *User Model. User-adapt. Interact.*, vol. 27, no. 3–5, pp. 313–350, Dec. 2017, doi: 10.1007/s11257-017-9193-2.
- [4] G.-H. Gweon, H.-S. Lee, C. Dorsey, R. Tinker, W. Finzer, and D. Damelin, “Tracking student progress in a game-like learning environment with a Monte Carlo Bayesian knowledge tracing model,” in *Proceedings of the Fifth International Conference on Learning Analytics And Knowledge*, New York, NY, USA: ACM, Mar. 2015, pp. 166–170. doi: 10.1145/2723576.2723608.
- [5] A. T. Corbett and J. R. Anderson, “Knowledge tracing: Modeling the acquisition of procedural knowledge,” *User Modelling and User-Adapted Interaction*, vol. 4, no. 4, pp. 253–278, 1995, doi: 10.1007/BF01099821.
- [6] R. Andriyat Krisdiawan, “PENERAPAN MODEL PENGEMBANGAN GAMEGDL (GAME DEVELOPMENT LIFE CYCLE) DALAM MEMBANGUN GAME PLATFORM BERBASIS MOBILE,” vol. 2, no. 1, 2019.
- [7] R. Ramadan and Y. Widyani, “Game development life cycle guidelines,” in *2013 International Conference on Advanced Computer Science and Information Systems (ICACSIS)*, IEEE, Sep. 2013, pp. 95–100. doi: 10.1109/ICACSIS.2013.6761558.
- [8] S. Aleem, L. F. Capretz, and F. Ahmed, “Game development software engineering process life cycle: a systematic review,” *Journal of Software Engineering Research and Development*, vol. 4, no. 1, p. 6, Dec. 2016, doi: 10.1186/s40411-016-0032-7.
- [9] A. T. Corbett and J. R. Anderson, “Knowledge tracing: Modeling the acquisition of procedural knowledge,” *User Modelling and User-Adapted Interaction*, vol. 4, no. 4, pp. 253–278, 1995, doi: 10.1007/BF01099821.
- [10] I. Šarić-Grgić, A. Grubišić, and A. Gašpar, “Twenty-five years of Bayesian knowledge tracing: a systematic review,” *User Model. User-adapt. Interact.*, vol. 34, no. 4, pp. 1127–1173,

- Sep. 2024, doi: 10.1007/s11257-023-09389-4.
- [11] O. Bulut, J. Shin, S. N. Yildirim-Erbasli, G. Gorgun, and Z. A. Pardos, “An Introduction to Bayesian Knowledge Tracing with pyBKT,” *Psych*, vol. 5, no. 3, pp. 770–786, Jul. 2023, doi: 10.3390/psych5030050.
- [12] T. J. McCabe, “A Complexity Measure,” *IEEE Transactions on Software Engineering*, vol. SE-2, no. 4, pp. 308–320, Dec. 1976, doi: 10.1109/TSE.1976.233837.

Universitas Kuningan
Jalan Cut Nyak Dien No 36A
Kel. Cijoho Kuningan Indonesia 45513
Contact Us : (0232) 874824 E-mail : nuansa.informatika@uniku.ac.id