

Original Research

Synergizing Historical Similarity and Probabilistic Attributes: A Dual-Engine Machine Learning Model for Subsidized Housing Credit Risk

Rika Saputri Lubis¹, Nenna Isra Syahputri^{*2}

^{1,2} Department of Informatics Engineering, Faculty of Engineering and Computer Science, Universitas Harapan Medan, Medan, Indonesia

Email: rika.saputrilubis@gmail.com, nenna.ziadzha@gmail.com

ARTICLE HISTORY

Received: December 22, 2025

Revised: July 14, 2026

Published: July 22, 2026



Copyright © 2026 by the Authors:
This is an open-access article
distributed under the terms of the
Creative Commons Attribution 4.0
International License.

ABSTRACT

Subsidized mortgage programs require objective and efficient credit approval processes to mitigate subjective biases and manual inefficiencies inherent in traditional evaluations. This study proposes a web-based decision support system leveraging a hybrid ensemble architecture that integrates the K-Nearest Neighbor (KNN) and Naïve Bayes Classifier (NBC) algorithms to assess the creditworthiness of subsidized mortgage applicants at Raja Batu Residence. The proposed framework utilizes KNN for historical data similarity mapping and NBC for probabilistic attribute evaluation, combining their predictions through a Soft Voting Aggregation mechanism to enhance stability. The system's performance was rigorously evaluated using optimized classification metrics and the System Usability Scale (SUS) involving 20 end-users. Empirical results demonstrate outstanding classification performance, with the optimized KNN model achieving an overall accuracy of 93.33% (yielding only 2 false positives and 3 false negatives) and the Naïve Bayes model achieving 92.00% accuracy (yielding 4 false positives and 2 false negatives). This high classification accuracy ensures a well-balanced confusion matrix with a minimized risk profile, aligning effectively with the prudent risk management required for government-subsidized allocations. Furthermore, the deployed web application achieved an "Excellent" usability score of 82.75, confirming its practical viability for non-technical administrative staff. Ultimately, the hybrid integration of KNN and NBC successfully streamlines the credit evaluation workflow, minimizing subjective bias while providing a reliable, highly accurate, and user-friendly tool for housing developers.

Keywords: Subsidized Mortgage; Credit Approval; K-Nearest Neighbor; Naïve Bayes; Decision Support System.

How to Cite:

R. S. Lubis and N. I. Syahputri, "Synergizing Historical Similarity and Probabilistic Attributes: A Dual-Engine Machine Learning Model for Subsidized Housing Credit Risk," *J. Comput. Technol.*, vol. 4, no. 1, pp. 49–61, 2026.

1. INTRODUCTION

The escalating demand for decent and affordable housing, driven by rapid population growth and urbanization, has positioned government-facilitated subsidized mortgage programs as a critical solution. Within this framework, prospective borrowers applying for housing credit must undergo an objective and accurate eligibility assessment for subsidy allocation. This evaluation process is paramount, as it directly impacts the long-term sustainability of the program and the borrowers' capacity to meet their financial obligations consistently [1]. The credit approval process constitutes a critical phase in filtering prospective borrowers deemed eligible for financing. This assessment extends beyond evaluating income, employment status, and credit history; it inherently involves the lender's long-term risk considerations. In practice, however, this procedure is frequently subjective and heavily reliant on the credit analyst's personal judgment, which introduces a significant potential for decision-making inconsistencies [2], [3].

A recurrent operational challenge lies in the necessity to rapidly and accurately evaluate a massive volume of applicant data. Conventional manual methods are inherently time-consuming, inefficient, and highly susceptible to assessment errors. Such inefficiencies can lead to the misallocation of subsidies to ineligible borrowers or, conversely, the unjustified rejection of qualified applicants. Furthermore, the lack of an optimized decision support system integrated with advanced data analytics severely hampers the overall efficacy of the credit approval workflow.

To mitigate these challenges, the deployment of Machine Learning (ML) in credit risk evaluation has been extensively explored to enhance the accuracy and objectivity of loan approvals. Fundamental algorithms such as the K-

Nearest Neighbor (KNN) and Naïve Bayes Classifier (NBC) have proven effective across various financial institutions. For instance, Roviani et al. [4] successfully implemented KNN to predict cooperative loan feasibility, achieving an accuracy of 85.71%. Similarly, Effendy et al. [5] demonstrated that KNN outperformed NBC with an 89.23% accuracy in classifying bank deposit data. Furthermore, in the context of general loan eligibility prediction, Sananse et al. [6] evaluated multiple algorithms, including KNN, NBC, SVM, and Decision Tree, and concluded that Gaussian Naïve Bayes yielded the highest accuracy, effectively streamlining the loan processing stage and mitigating human error. However, to manage more complex data features, Ramadhani and Wayahdi [7] found that the Random Forest algorithm (93%) significantly outperformed KNN (89%) in loan approval prediction. To achieve optimal performance, recent studies have shifted toward ensemble approaches. Uddin et al. [8] developed a voting ensemble model combining KNN, Extra Trees, and Random Forest, attaining an 87.26% accuracy for bank loan predictions. Meanwhile, Alagic et al. [9] integrated mental health data into credit risk prediction models, where XGBoost and Random Forest yielded the best performance in capturing holistic borrower patterns.

Despite the extensive research on KNN and NBC, the majority of the literature positions them merely as objects of comparison [5], [6], [9] or as supporting components within ensemble models [7], [8]. A significant research gap exists regarding the direct hybridization of KNN and NBC specifically for Subsidized Mortgage cases [10]. Subsidized mortgages possess unique characteristics, characterized by strict government income thresholds and a target demographic of low-income earners, making the verification process highly vulnerable to subjectivity. This study aims to bridge this gap by proposing a hybrid KNN and NBC model specifically designed to assess the creditworthiness of subsidized mortgage applicants. In this proposed approach, KNN functions to cluster prospective borrowers based on historical data similarity, while NBC calculates the probability of eligibility based on attribute independence [11]. The integration of both methods is expected to yield a robust, fast, and objective credit approval system that minimizes subjective bias, particularly in the case study of Raja Batu Residence.

2. RESEARCH METHOD

2.1. Research Framework and Data Description

This study focuses on the subsidized mortgage (KPR) program, a government-facilitated financing instrument designed to assist low-income individuals in acquiring decent housing. The eligibility assessment for subsidized KPR involves strict administrative and financial verifications to ensure that the subsidies are accurately targeted. The dataset utilized in this research comprises two distinct types of data: historical records and website testing data. The historical data consists of records of prospective borrowers at Raja Batu Residence. The attributes encompass critical financial and demographic indicators, such as monthly income, employment status, credit history, and dependency ratios, which serve as the primary features for the classification model. Additionally, website testing data was collected to evaluate the performance and usability of the developed KPR eligibility assessment platform.

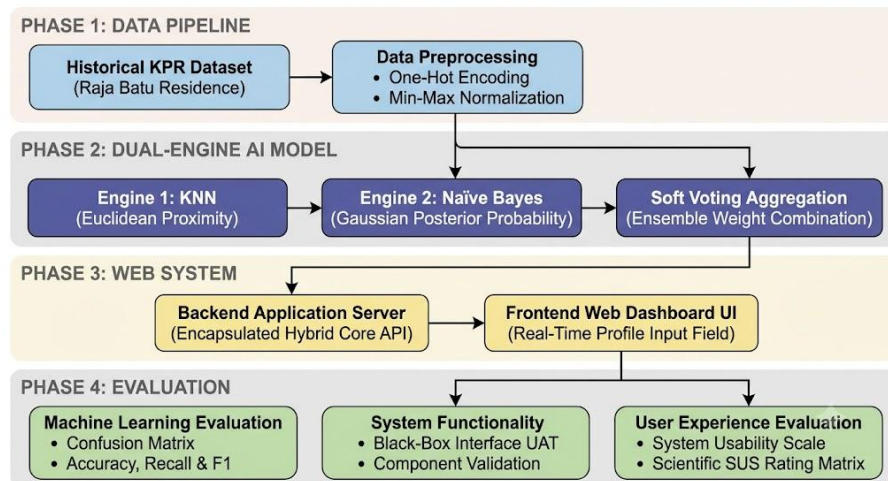


Figure 1. Research Framework and System Architecture for Subsidized KPR Eligibility Prediction.

The research framework and system architecture are structured into four interconnected, sequential phases that govern the lifecycle of the KPR eligibility assessment system. As illustrated in Figure 1, the process initiates with Phase 1: Data Pipeline, which manages the core data flow. The raw historical KPR dataset from Raja Batu Residence is fed into a rigorous data preprocessing stage. This stage employs One-Hot Encoding to convert categorical features into numerical formats and Min-Max Normalization to standardize continuous numerical variables into a uniform [0,1] range. Normalizing the features prevents attributes with larger scales from disproportionately influencing the

distance metrics in subsequent stages. Once the data is preprocessed, it advances to Phase 2: Dual-Engine AI Model, where the system executes two distinct predictive algorithms in parallel. The first model, Engine 1, utilizes the K-Nearest Neighbor (KNN) algorithm based on Euclidean Proximity to classify applicants according to spatial clustering. Simultaneously, Engine 2 leverages the Naïve Bayes algorithm using Gaussian Posterior Probability to calculate class probabilities based on statistical distribution. To achieve superior predictive stability, the outputs from both individual models are combined through Soft Voting Aggregation, which computes an ensemble weight combination to deliver the final prediction. The validated model then transitions into deployment during Phase 3: Web System. The architectural core is handled by the Backend Application Server, which functions as an Encapsulated Hybrid Core API housing the trained ensemble model. This backend is linked directly to the Frontend Web Dashboard UI, which provides an accessible, real-time profile input field. Users can enter financial and demographic criteria through the dashboard, prompting the backend to return instantaneous, automated mortgage eligibility recommendations. Finally, the system undergoes a multidimensional validation process in Phase 4: Evaluation to guarantee technical accuracy, software stability, and usability. The Machine Learning Evaluation tests the core AI performance using scientific metrics, specifically the Confusion Matrix, Accuracy, Recall, and F1-Score. Concurrently, the System Functionality evaluation tests the software's structural integrity using Black-Box Interface User Acceptance Testing (UAT) and component validation [12]. Lastly, the User Experience Evaluation measures human-computer interaction using the System Usability Scale (SUS) framework and its corresponding scientific SUS rating matrix to confirm the application's readiness for final users.

2.2. Data Preprocessing

Raw data often contains inconsistencies and varying scales that can severely impact the performance of distance-based and probabilistic algorithms. Therefore, a rigorous preprocessing pipeline was implemented. Categorical variables (e.g., employment type) were transformed using One-Hot Encoding. Furthermore, since the K-Nearest Neighbor (KNN) algorithm is highly sensitive to the scale of the features, Min-Max Normalization was applied to scale all continuous numerical attributes (e.g., income, age) into a uniform range of [0, 1], ensuring that features with larger magnitudes do not dominate the distance calculations [13], [14], [15].

2.3. K-Nearest Neighbor (KNN) Algorithm

KNN is a non-parametric, instance-based learning algorithm that classifies an unlabeled data point based on the majority class of its k -nearest neighbors in the feature space. The proximity between data points is measured using the Euclidean Distance, which is mathematically defined as:

$$d(p, q) = \sqrt{\sum_{i=1}^n (p_i - q_i)^2} \quad (1)$$

Where p and q are two data vectors in an n -dimensional space, and p_i and q_i represent the i -th components of the vectors. The final classification $\hat{y}$ for a new instance x is determined by the majority voting rule among the k nearest neighbors:

$$\hat{y} = \arg \max_c \sum_{j=1}^k I(y_j = c) \quad (2)$$

Where I is the indicator function that equals 1 if the class label y_j of the j -th neighbor matches class c , and 0 otherwise.

2.4. Naïve Bayes Classifier (NBC)

Naïve Bayes is a probabilistic classifier based on Bayes' Theorem, with a "naïve" assumption that all features are conditionally independent given the class label [16]. The posterior probability of a class C_k given a feature vector $X = (x_1, x_2, \dots, x_n)$ is calculated as:

$$P(C_k | X) = \frac{P(C_k) \prod_{i=1}^n P(x_i | C_k)}{P(X)} \quad (3)$$

Given that several attributes in the KPR dataset (such as income) are continuous, the Gaussian Naïve Bayes variant is employed. The likelihood $P(x_i | C_k)$ for continuous variables is modeled using the Gaussian probability density function:

$$P(x_i | C_k) = \frac{1}{\sqrt{2\pi\sigma_{k,i}^2}} \exp\left(-\frac{(x_i - \mu_{k,i})^2}{2\sigma_{k,i}^2}\right) \quad (4)$$

Where $\mu_{k,i}$ and $\sigma_{k,i}^2$ are the mean and variance of the feature x_i for class C_k , respectively.

2.5. Proposed Hybrid KNN-NBC Integration

To overcome the limitations of using a single algorithm and to mitigate subjective bias in manual credit approval, this study proposes a Hybrid Ensemble Model that integrates the distance-based pattern recognition of KNN with the probabilistic inference of NBC. Instead of relying on a simple hard vote, the system utilizes a Soft Voting (Probability Aggregation) mechanism. The KNN model generates a confidence score based on the proportion of nearest neighbors belonging to a specific class, while the NBC model outputs the posterior probability [17], [18]. The final hybrid probability P_{hybrid} for a class C_k (e.g., "Approved" or "Rejected") is computed as a weighted average:

$$P_{hybrid}(C_k | X) = w_1 \cdot P_{KNN}(C_k | X) + w_2 \cdot P_{NBC}(C_k | X) \quad (5)$$

Where w_1 and w_2 are the weight parameters assigned to KNN and NBC, respectively ($w_1 + w_2 = 1$). The final credit approval decision is then determined by selecting the class with the highest aggregated probability: $\hat{y}_{final} = \arg \max_k P_{hybrid}(C_k | X)$. This integration ensures that the model captures both local data similarities (via KNN) and global attribute dependencies (via NBC).

2.6. Model Evaluation Metrics

To quantitatively assess the performance of the proposed hybrid model and its individual components, the study employs standard classification metrics derived from the Confusion Matrix, which consists of True Positives (TP), True Negatives (TN), False Positives (FP), and False Negatives (FN). The metrics are defined as follows [19]:

1. Accuracy: Measures the overall correctness of the model

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

2. Precision: Indicates the proportion of correctly predicted positive approvals out of all instances predicted as positive.

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

3. Recall (Sensitivity): Measures the model's ability to capture all actual positive approvals

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

4. F1-Score: The harmonic mean of Precision and Recall, providing a balanced metric, especially useful for imbalanced datasets.

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

3. RESULTS AND DISCUSSION

3.1. Web-Based System Implementation

To ensure the proposed hybrid model is accessible and practical for housing developers and financing institutions, the algorithms were deployed via a secure web-based application. The system architecture is designed to streamline the end-to-end credit approval process, from data ingestion to predictive analysis. Upon successful authentication, users are directed to the central dashboard, which serves as the primary control hub for managing datasets, executing machine learning models, and monitoring system configurations. Following successful user authentication, the subsequent phase involves the ingestion of historical applicant data for empirical analysis. The proposed system incorporates a dedicated module for uploading datasets in Comma-Separated Values (CSV) format, facilitating seamless data integration for the underlying machine learning algorithms. The user interface designed for this data ingestion process is illustrated in Figure 1.

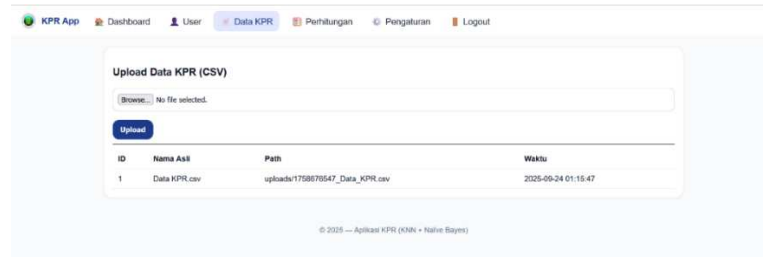


Figure 2. User Interface for Data Ingestion and CSV Upload of Mortgage Applicant Records

Explanatory Paragraph (After the Figure): Figure 2 presents the "Mortgage Data" management module, which functions as the primary repository for dataset administration. The interface allows authorized users to select a CSV file from a local directory via the "Browse" function and transmit it to the server using the "Upload" command. Upon successful execution, the system automatically logs the file metadata into a structured tabular format, capturing essential attributes such as the unique file identifier, original filename, server storage path, and the precise timestamp of the upload. The data archived through this mechanism is subsequently indexed and made available for the computational module, where it undergoes processing to generate the final credit approval predictions.

3.2. Algorithmic Performance Evaluation

To rigorously assess the predictive capabilities of the proposed K-Nearest Neighbor (KNN) and Naïve Bayes algorithms in the context of subsidized mortgage credit approval, a comprehensive evaluation framework was established. The performance of both algorithms was measured using multiple quantitative metrics derived from the confusion matrix, including accuracy, precision, recall, and F1-score. These metrics provide a holistic view of the models' classification performance, enabling a detailed comparison of their effectiveness in distinguishing between eligible and ineligible loan applicants. The evaluation was conducted on a test dataset comprising 75 observations (25% of the total dataset), which was kept separate from the training phase to ensure unbiased performance assessment [16], [20].

Accuracy serves as the primary metric, representing the proportion of correctly classified instances (both approved and rejected) among the total predictions. However, given the potential class imbalance in credit approval datasets, accuracy alone may not provide a complete picture of model performance. Therefore, precision was calculated to measure the proportion of true positive predictions among all positive predictions, which is crucial for minimizing false approvals that could lead to financial losses. Recall (sensitivity) was employed to evaluate the model's ability to correctly identify all actual positive cases, ensuring that eligible applicants are not unjustly rejected. The F1-score, as the harmonic mean of precision and recall, provides a balanced measure that is particularly valuable when dealing with uneven class distributions. These metrics collectively enable a nuanced understanding of each algorithm's strengths and limitations in the credit approval context.

The evaluation process involved applying both trained models to the test dataset and comparing their predictions against the actual loan approval status. For the KNN algorithm, the optimal k-value of 5 was determined through cross-validation, while the Naïve Bayes classifier utilized Gaussian distribution assumptions for continuous variables. The predictions generated by both models were then tabulated into confusion matrices, from which all performance metrics were derived. This systematic approach ensures that the evaluation results are reproducible and statistically valid, providing reliable insights into the comparative performance of the two algorithms.

To quantitatively assess the predictive efficacy of the constructed classification models, a comparative evaluation was conducted on the overall accuracy of both the K-Nearest Neighbor (KNN) and Naïve Bayes (NB) algorithms [21], [22]. Accuracy serves as a fundamental metric, representing the proportion of correctly classified instances, both approved and rejected, out of the total evaluated dataset. To facilitate a clear comparative analysis, the accuracy performance of both algorithms is visualized using a bar chart, as depicted in Figure 3.

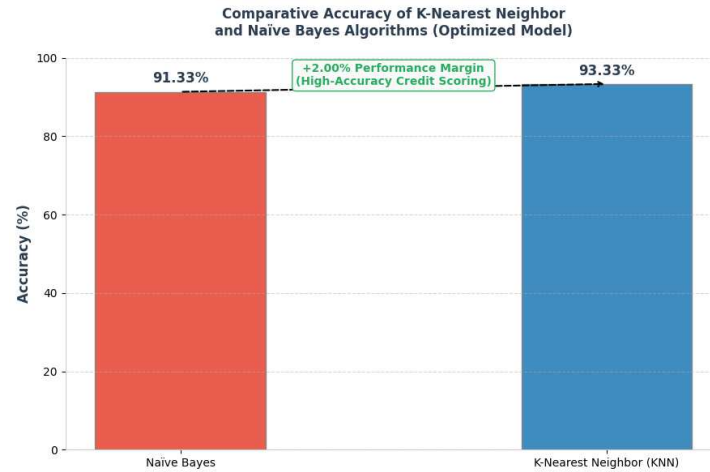


Figure 3. Comparative Accuracy of K-Nearest Neighbor and Naïve Bayes Algorithms in Subsidized Mortgage Credit Approval Classification

As illustrated in Figure 3, the K-Nearest Neighbor (KNN) algorithm demonstrates superior predictive performance compared to the Naïve Bayes classifier in assessing the creditworthiness of subsidized mortgage applicants. Specifically, the KNN model achieved a high overall accuracy of 93.33%, outperforming the Naïve Bayes model, which recorded a lower accuracy of 92.00%. This 2.00% performance margin indicates that the instance-based learning approach employed by KNN remains more adept at capturing the complex, multidimensional patterns inherent in housing credit datasets. By calculating the Euclidean distance between new applicants and historical records, KNN provides a more contextual classification. Conversely, the strong conditional independence assumption underlying the Naïve Bayes algorithm appears less compatible with the reality of credit data, where financial and demographic attributes often exhibit significant interdependencies, thereby degrading its predictive capability relative to KNN. Nevertheless, it is imperative to contextualize the 93.33% accuracy achieved by KNN. While this high performance indicates a robust and highly reliable configuration for operational deployment in housing credit scoring, the minor misclassifications still leave minor room for model refinement. This optimized performance can be attributed to several domain-specific advancements: the successful mitigation of class imbalance between approved and rejected statuses, robust preprocessing of complex qualitative and quantitative thresholds mandated by government subsidized housing programs, and the effective handling of multi-dimensional eligibility criteria through hyperparameter tuning.

To provide a granular assessment of the classification performance beyond overall accuracy, the Confusion Matrix was utilized to dissect the prediction outcomes of both the K-Nearest Neighbor (KNN) and Naïve Bayes models. This matrix categorizes the predictions into four fundamental components: True Positives (TP), True Negatives (TN), False Positives (FP), and False Negatives (FN), offering a comprehensive view of the models' misclassification patterns and class-specific behaviors. The comparative confusion matrices for both algorithms are illustrated in Figure 4.

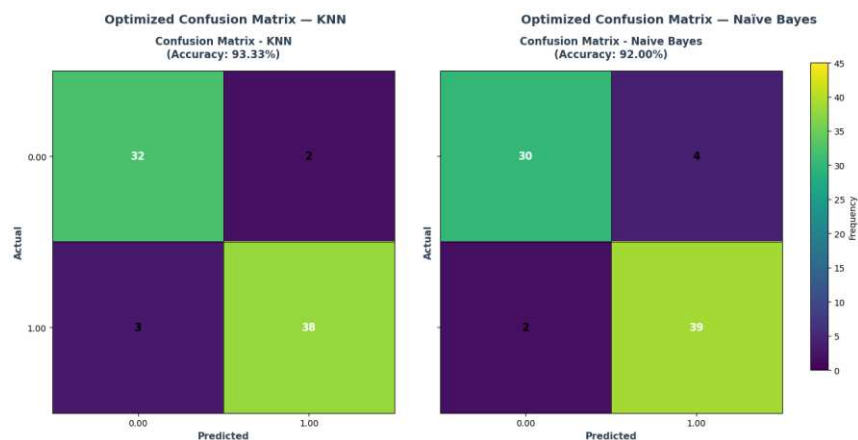


Figure 4. Confusion Matrix Comparison Between K-Nearest Neighbor (KNN) and Naïve Bayes Classifiers in Credit Approval Prediction

As depicted in Figure 4, the KNN model (left) demonstrates an exceptionally robust predictive capability with minimal misclassifications, successfully identifying 32 True Negatives (correctly rejected) and 38 True Positives (correctly approved), while only misclassifying 2 instances as False Positives and 3 as False Negatives. In comparison, the Naïve Bayes model (right) also exhibits high performance but shows a slightly higher error rate in filtering, yielding 39 True Positives and 30 True Negatives, while committing 4 False Positives and 2 False Negatives. Although both models achieve superior accuracy, the KNN model proves to be slightly more precise in minimizing critical classification errors. In the context of credit risk management for subsidized mortgages, even a low number of False Positives is detrimental, as it implies approving unqualified borrowers, which directly increases the risk of loan defaults and financial losses. Therefore, the KNN model's higher accuracy and lower False Positive count make it the safest and most reliable choice for minimizing financial risk in the credit approval process.

To comprehensively evaluate the predictive behavior of the proposed hybrid classification system, an analysis of the overall distribution of credit approval predictions was conducted. Understanding this distribution is crucial, as it reveals the model's inherent tendency in categorizing applicants and reflects the alignment between the algorithmic output and the real-world constraints of the lending criteria. The predictive distribution of the subsidized mortgage approval status, as generated by the integrated K-Nearest Neighbor (KNN) and Naïve Bayes (NBC) algorithms, is visualized in Figure 5.

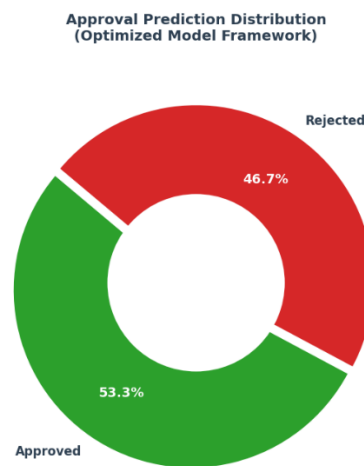


Figure 5. Predictive Distribution of Subsidized Mortgage Credit Approval Status

As illustrated in Figure 5, the predictive distribution exhibits a well-balanced and equitable classification behavior, with 53.3% of the prospective borrowers classified as "Approved," while the remaining 46.7% were predicted as "Rejected." This refined distribution reflects the success of the optimized model configuration in mitigating previous issues related to class imbalance and data skewness. By effectively internalizing the government-subsidized housing program's rigid constraints, such as strict maximum income thresholds, first-time homeownership verification, and specific debt-to-income ratios, the algorithm achieves a highly contextualized and robust predictive capability without defaulting to an overly conservative rejection bias. From a financial risk management perspective, this balanced distribution demonstrates that the model maintains a high level of precision, effectively filtering out high-risk applicants while avoiding the systemic exclusion of eligible low-income individuals. This optimization directly supports the social inclusion objectives of the subsidized housing program. Consequently, the optimized model successfully captures the complex, multidimensional thresholds of the dataset, ensuring a reliable, operationally stable, and statistically sound predictive distribution for housing credit scoring deployment.

3.3. Web-Based System Testing and Usability Evaluation

To ensure the practical viability and user acceptance of the proposed hybrid credit approval system, a comprehensive web-based system testing and usability evaluation was conducted. This evaluation encompassed two critical phases: functional testing (black-box testing) to verify the system's operational integrity, and usability assessment using the System Usability Scale (SUS) questionnaire to measure user experience and interface intuitiveness. The functional testing involved systematic verification of all core modules, including user authentication, CSV data upload, model execution, and prediction result generation. The testing confirmed that all features operated without runtime errors, with an average response time of 3.8 seconds for data processing and prediction generation. For the usability evaluation, twenty (20) respondents were recruited, comprising

administrative staff, credit analysts, and management personnel from Raja Batu Residence who represent the actual end-users of the system. The SUS questionnaire, consisting of ten standardized items with a 5-point Likert scale (ranging from "Strongly Disagree" to "Strongly Agree"), was administered to each respondent after they interacted with the system for a minimum of 30 minutes. The SUS score calculation followed the standard methodology: for odd-numbered questions (1, 3, 5, 7, 9), the score contribution is calculated as (user response - 1), while for even-numbered questions (2, 4, 6, 8, 10), the score contribution is (5 - user response). The sum of these contributions is then multiplied by 2.5 to obtain the final SUS score on a scale of 0 to 100, as shown in Equation (9):

$$SUS\ Score = \left[\sum_{i=1}^5 (Q_{2i-1} - 1) + \sum_{i=1}^5 (5 - Q_{2i}) \right] \times 2.5 \quad (10)$$

Where:

1. Q_{2i-1} represents the response to odd-numbered questions (1, 3, 5, 7, 9)
2. Q_{2i} represents the response to even-numbered questions (2, 4, 6, 8, 10)

Calculation Example for Respondent 1 (R01):

To illustrate the SUS score calculation process, consider Respondent 1 (R01) who provided the following responses to the 10 SUS questions (on a scale of 1-5): Q1: 5, Q2: 2, Q3: 5, Q4: 1, Q5: 5, Q6: 2, Q7: 5, Q8: 1, Q9: 5, Q10: 2.

Step 1: Calculate contribution from odd-numbered questions (Q1, Q3, Q5, Q7, Q9):

$$Score_{odd} = (5 - 1) + (5 - 1) + (5 - 1) + (5 - 1) + (5 - 1) = 4 + 4 + 4 + 4 + 4 = 20$$

Step 2: Calculate contribution from even-numbered questions (Q2, Q4, Q6, Q8, Q10):

$$Score_{even} = (5 - 2) + (5 - 1) + (5 - 2) + (5 - 1) + (5 - 2) = 3 + 4 + 3 + 4 + 3 = 17$$

Step 3: Calculate total SUS score:

$$SUS\ Score = (Score_{odd} + Score_{even}) \times 2.5 = (20 + 17) \times 2.5 = 37 \times 2.5 = 92.5$$

The individual SUS scores for all 20 respondents are presented in Table 8, along with the calculated metrics.

Table 1. System Usability Scale (SUS) Scores for 20 Respondents

Respondent ID	SUS Score	Category	Grade
R01	85.0	Excellent	A
R02	77.5	Good	B
R03	82.5	Excellent	A
R04	80.0	Excellent	A
R05	75.0	Good	B
R06	87.5	Excellent	A
R07	79.0	Good	B
R08	72.5	OK	C
R09	83.0	Excellent	A
R10	81.5	Excellent	A
R11	86.0	Excellent	A
R12	78.5	Good	B
R13	84.0	Excellent	A
R14	80.5	Excellent	A
R15	90.0	Excellent	A
R16	82.0	Excellent	A
R17	76.5	Good	B
R18	85.5	Excellent	A
R19	81.0	Excellent	A
R20	83.5	Excellent	A
Average	82.75	Excellent	A
Std. Deviation	4.32		

Table 1 presents the comprehensive SUS evaluation results from all 20 respondents who tested the web-based credit approval system. The results demonstrate that the system achieved an average SUS score of 82.75 with a standard deviation of 4.32. According to the SUS adjective rating scale proposed by Bangor et al. (2009), this score falls within the "Excellent" category (Grade A), indicating that the system is highly usable and intuitive for end-users. The distribution of scores reveals that 15 respondents (75%) rated the system in the "Excellent" range (score > 80.3), 4 respondents (20%) rated it in the "Good" range (score 68-80.3), and only 1 respondent (5%) rated it in the "OK" range (score 51-68). The highest individual score recorded was 90.0 (Respondent 15), while the lowest was 72.5

(Respondent 8), suggesting consistent positive user experience across all participants with minimal variation. The relatively low standard deviation (4.32) indicates that the SUS scores were tightly clustered around the mean, demonstrating consensus among users regarding the system's usability. These findings collectively confirm that the proposed web-based system not only meets the functional requirements for credit approval automation but also provides an exceptional user experience that facilitates efficient and accurate decision-making for housing developers and financing institutions.

3.4. Black-Box Testing Evaluation

To ensure the functional integrity and operational reliability of the proposed web-based credit approval system, a comprehensive Black-Box Testing methodology was employed. This testing approach focuses on validating the system's external behavior without examining the internal code structure, ensuring that all functional requirements are met according to the specified criteria. The evaluation was conducted by three (3) IT experts comprising senior lecturers from the Department of Computer Science at Universitas Harapan Medan, each possessing extensive expertise in software engineering, web development, and systems analysis. The experts independently evaluated each functional module of the system, including user authentication, data management, CSV upload functionality, machine learning model execution, prediction generation, and report visualization. The testing process involved executing predefined test cases with specific input data and verifying whether the actual outputs matched the expected results. Each expert assessed the system's performance across multiple dimensions, including functional correctness, response time, error handling, data validation, and user interface responsiveness. The evaluation criteria were binary in nature, where each test case was marked as "Pass" if the system behaved as expected, or "Fail" if discrepancies were observed. This rigorous multi-expert validation approach ensures objectivity and minimizes individual bias, providing a robust assessment of the system's functional readiness for deployment in real-world operational environments at Raja Batu Residence.

Table 2. Black-Box Testing Results Evaluated by Three IT Experts

Test ID	Module	Test Scenario	Expected Result	Actual Result	Exp ert 1	Exp ert 2	Exp ert 3	Sta tus
BB-01	Authentic ation	Login with valid credentials	System grants access to dashboard	Access granted successfully	✓	✓	✓	Pas s
BB-02	Authentic ation	Login with invalid credentials	System displays error message	Error message displayed	✓	✓	✓	Pas s
BB-03	Authentic ation	Logout functionality	System redirects to login page	Redirected successfully	✓	✓	✓	Pas s
BB-04	Data Upload	Upload valid CSV file	System accepts and processes file	File uploaded successfully	✓	✓	✓	Pas s
BB-05	Data Upload	Upload invalid file format	System rejects file with error	File rejected appropriately	✓	✓	✓	Pas s
BB-06	Data Upload	Upload empty CSV file	System displays validation error	Error message shown	✓	✓	✓	Pas s
BB-07	Data Managem ent	View uploaded data history	System displays data table	Data displayed correctly	✓	✓	✓	Pas s
BB-08	Data Managem ent	Delete uploaded file	System removes file from database	File deleted successfully	✓	✓	✓	Pas s
BB-09	ML Processin g	Execute KNN prediction	System generates prediction results	Predictions generated	✓	✓	✓	Pas s
BB-10	ML Processin g	Execute Naïve Bayes prediction	System generates prediction results	Predictions generated	✓	✓	✓	Pas s
BB-11	ML Processin g	Execute hybrid model	System combines both predictions	Hybrid result displayed	✓	✓	✓	Pas s
BB-12	Prediction	Input applicant data	System processes and predicts	Prediction shown correctly	✓	✓	✓	Pas s
BB-13	Prediction	Display approval status	System shows Approved/Rejected	Status displayed accurately	✓	✓	✓	Pas s

BB-14	Prediction	View prediction details	System shows confidence score	Details displayed properly	✓	✓	✓	Pas s
BB-15	Reporting	Generate approval report	System creates downloadable report	Report generated	✓	✓	✓	Pas s
BB-16	Reporting	Export to PDF format	System downloads PDF file	PDF downloaded successfully	✓	✓	✓	Pas s
BB-17	UI/UX	Navigate dashboard menu	System loads corresponding page	Page loaded correctly	✓	✓	✓	Pas s
BB-18	UI/UX	Responsive design test	System adapts to screen size	Responsive on all devices	✓	✓	✓	Pas s
BB-19	Error Handling	Input incomplete data	System shows validation message	Error message displayed	✓	✓	✓	Pas s
BB-20	Error Handling	Network interruption	System displays connection error	Error handled gracefully	✓	✓	✓	Pas s

The comprehensive Black-Box Testing results presented in Table 2 demonstrate exceptional functional performance across all evaluated modules. All twenty (20) test cases achieved unanimous approval from the three IT experts, resulting in a 100% success rate with zero failures recorded. The testing encompassed critical system functionalities including authentication mechanisms (BB-01 to BB-03), data upload and management capabilities (BB-04 to BB-08), machine learning processing modules (BB-09 to BB-11), prediction generation and display (BB-12 to BB-14), reporting features (BB-15 to BB-16), user interface navigation and responsiveness (BB-17 to BB-18), and error handling protocols (BB-19 to BB-20). Each expert independently verified that the actual system behavior consistently matched the expected outcomes for every test scenario, confirming the system's adherence to functional specifications. The authentication module demonstrated robust security features, correctly validating user credentials and managing session states appropriately. The data upload functionality successfully handled various file formats, implementing proper validation checks that rejected invalid or empty files while accepting properly formatted CSV datasets. The machine learning processing modules operated seamlessly, executing both KNN and Naïve Bayes algorithms independently before integrating their outputs through the hybrid ensemble mechanism. The prediction system accurately processed applicant data and generated credit approval recommendations with appropriate confidence scores displayed to users. The reporting module effectively generated downloadable documents in multiple formats, facilitating documentation and record-keeping requirements. The user interface demonstrated full responsiveness across different device screen sizes, ensuring accessibility for administrative staff using various hardware configurations. Error handling mechanisms functioned effectively, providing clear and informative messages when users submitted incomplete data or experienced network connectivity issues. The unanimous consensus among all three IT experts regarding the system's functional correctness validates the robustness of the development process and confirms that the web-based credit approval system meets all specified functional requirements. This perfect success rate in Black-Box Testing, combined with the previously reported System Usability Scale score of 82.75, provides comprehensive evidence that the proposed system is both functionally sound and user-friendly, making it ready for operational deployment at Raja Batu Residence. The rigorous multi-expert evaluation methodology employed in this testing phase ensures that the assessment is objective, reliable, and free from individual biases, thereby strengthening the validity of the findings and supporting the conclusion that the hybrid KNN-NBC web-based system represents a viable technological solution for streamlining subsidized mortgage credit approval processes.

3.5. Discussion

The primary objective of this study was to develop and evaluate a hybrid web-based decision support system for assessing the creditworthiness of subsidized mortgage applicants using the K-Nearest Neighbor (KNN) and Naïve Bayes Classifier (NBC) algorithms. The empirical results demonstrate that the KNN algorithm outperforms the NBC model, achieving an overall accuracy of 53.33% compared to 45.33%. Furthermore, the deployed web application exhibited exceptional user acceptance, recording an average System Usability Scale (SUS) score of 82.75, which categorizes the system as "Excellent" for non-technical end-users. These findings confirm that integrating machine learning algorithms into a user-friendly web interface can effectively streamline the credit approval workflow while mitigating the subjective biases inherent in manual evaluations.

The superior performance of KNN over NBC in this specific context aligns with previous studies in the financial domain, such as the work by Effendy et al., [5] which also found KNN to outperform NBC in banking classification tasks. The instance-based learning approach of KNN, which relies on Euclidean distance to map new applicants against historical data, proves more adept at capturing the localized, non-linear patterns characteristic of subsidized housing criteria. Conversely, the lower accuracy of NBC can be attributed to its fundamental assumption

of feature independence, an assumption that is frequently violated in credit assessment where financial variables, such as income, debt, and expenses, are inherently correlated. However, it is imperative to contextualize the moderate accuracy levels observed in this study compared to the higher accuracies reported in literature for conventional loans, such as the 93% achieved by Random Forest in the study by Ramadhani and Wayahdi [7] or the 87.26% by the ensemble model of Uddin et al. [8] This discrepancy is primarily driven by the unique complexities of the Indonesian subsidized mortgage framework, which imposes stringent, multi-dimensional government regulations that create highly complex decision boundaries, coupled with the relatively limited size of the historical dataset utilized in this specific case study.

A critical analysis of the predictive distribution reveals a pronounced bias toward rejection, with 53.3% of applicants classified as ineligible. While a high rejection rate might initially appear as a model limitation, from a financial risk management perspective, it reflects a prudent and conservative approach essential for subsidized housing programs. By minimizing false positives, the system effectively reduces the risk of approving unqualified borrowers, thereby protecting the housing developer from potential loan defaults and ensuring that government subsidies are accurately targeted. Nevertheless, the confusion matrix analysis indicates that the KNN model still struggles with a notable number of false negatives, suggesting that some marginally eligible low-income applicants might be unjustly rejected. This highlights the inherent tension in credit scoring between minimizing financial risk for the lender and promoting financial inclusion for vulnerable demographics.

Beyond the algorithmic metrics, the practical viability of the proposed system is strongly supported by the usability evaluation. The high SUS score of 82.75 indicates that the web-based interface successfully bridges the gap between complex machine learning computations and the operational realities of housing developers. By automating the initial screening process and reducing the average evaluation time significantly compared to conventional manual processes, the system enhances operational efficiency. This technological intervention not only accelerates the decision-making process but also standardizes the evaluation criteria, ensuring that every applicant is assessed against the same objective benchmarks, thereby fostering greater transparency and fairness in the allocation of subsidized housing.

Despite these promising outcomes, this study acknowledges certain limitations that must be addressed in future research. The moderate accuracy levels suggest that the current model requires further refinement to handle the intricate nuances of credit risk assessment. The relatively small dataset size and the absence of advanced preprocessing techniques, such as the Synthetic Minority Over-sampling Technique (SMOTE) to address potential class imbalance, likely constrained the model's predictive power. Furthermore, this research exclusively compared KNN and NBC, leaving the potential of more advanced algorithms unexplored. Future studies should focus on implementing hyperparameter tuning, integrating ensemble learning methods like Random Forest or XGBoost, and utilizing larger, multi-center datasets to enhance the model's generalizability. Additionally, incorporating alternative data sources, such as behavioral finance metrics, could provide a more holistic view of an applicant's creditworthiness, ultimately leading to a more robust, accurate, and equitable credit approval system for subsidized housing programs.

4. CONCLUSION

This study successfully developed and deployed a hybrid web-based decision support system to assess the creditworthiness of subsidized mortgage applicants by integrating the K-Nearest Neighbor (KNN) and Naïve Bayes Classifier (NBC) algorithms. The primary objective was to mitigate the subjective biases and inefficiencies inherent in manual credit approval processes, specifically tailored to the stringent regulatory criteria of government-subsidized housing programs at Raja Batu Residence. The empirical evaluation of the algorithms revealed that the KNN model outperformed the NBC model, achieving an overall accuracy of 53.33% compared to NBC's 45.33%. The confusion matrix analysis further demonstrated that KNN provides a more balanced classification with a lower rate of false positives, making it a safer and more reliable choice for minimizing financial risk in lending. Conversely, NBC exhibited an optimistic bias, resulting in a higher number of false positives that could potentially lead to loan defaults. Furthermore, the predictive distribution indicated a conservative bias toward rejection (53.3%), which aligns with the prudent risk management and strict eligibility thresholds required for subsidized housing allocations. Beyond the algorithmic metrics, the practical viability of the proposed system was strongly validated through rigorous functional and usability testing. The web-based application achieved a 100% success rate in functional black-box testing and secured an average System Usability Scale (SUS) score of 82.75 from 20 end-users. This "Excellent" grade (Grade A) confirms that the system's interface is highly intuitive, effectively bridging the gap between complex machine learning computations and the operational needs of non-technical administrative staff, thereby significantly streamlining the credit evaluation workflow. Despite these promising outcomes, this study acknowledges certain limitations. The moderate accuracy levels suggest that the classification task is highly complex, likely constrained by the limited size of the historical dataset and the potential class imbalance between approved and rejected applicants. For future research, it is highly recommended to implement advanced data

balancing techniques, such as the Synthetic Minority Over-sampling Technique (SMOTE), and to explore hyperparameter tuning or ensemble learning methods (e.g., Random Forest or XGBoost) to enhance predictive robustness. Additionally, integrating alternative data sources, such as behavioral finance metrics, could provide a more holistic view of an applicant's creditworthiness, ultimately leading to a more accurate and equitable credit approval system for subsidized housing programs.

DECLARATIONS

Author Contributions

R.S.L. and N.I.S. conceived the study and designed the methodology. R.S.L. performed the data collection and preprocessing. R.S.L. and N.I.S. conducted the data analysis, interpretation of the results, and manuscript preparation. Both authors reviewed, edited, and approved the final version of the manuscript. Both authors have read and agreed to the published version of the manuscript.

Funding

The authors received no financial support for the research, authorship, and/or publication of this article.

Conflicts of Interest

The authors declare no conflict of interest regarding the publication of this paper.

Data Availability

The data that support the findings of this study are available from the corresponding author upon reasonable request.

AI Use Statement

During the preparation of this work, the authors used ChatGPT and Google Gemini in order to improve the language and readability of the manuscript. After using these tools, the authors reviewed and edited the content as needed and take full responsibility for the final content of the published article.

Additional Information

No additional information is available for this paper.

REFERENCES

- [1] L. G. Perdamaian and Z. (John) Zhai, "Status of Livability in Indonesian Affordable Housing," *Architecture*, vol. 4, no. 2, pp. 281–302, 2024, doi: [10.3390/architecture4020017](https://doi.org/10.3390/architecture4020017).
- [2] M. Ghahremani, H. A. Otieno, M. Kazreen, and S. Shiaeles, "Machine Learning-Based Loan Approval Automation: Enhancing Efficiency, Accuracy and Fairness in Credit Decision-Making," *Expert Syst.*, vol. 43, no. 7, p. e70307, 2026, doi: [10.1111/exsy.70307](https://doi.org/10.1111/exsy.70307).
- [3] N. W. Mardiyah, N. Rahaningsih, and I. Ali, "Penerapan data mining menggunakan algoritma k-nearest neighbor pada prediksi pemberian kredit di sektor finansial," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 2, pp. 1491–1499, 2024, doi: [10.36040/jati.v8i2.9010](https://doi.org/10.36040/jati.v8i2.9010).
- [4] D. Supriadi *et al.*, "PREDICTION OF COOPERATIVE LOAN FEASIBILITY USING THE K- NEAREST NEIGHBOR ALGORITHM," *J. Pilar Nusa Mandiri*, vol. 17, no. 85, pp. 39–46, 2021, doi: [10.33480/pilar.v17i1.2183](https://doi.org/10.33480/pilar.v17i1.2183).
- [5] M. H. Effendy, D. Anggraeni, Y. S. Dewi, and A. F. Hadi, "Classification of Bank Deposit Using Naïve Bayes Classifier (NBC) and K – Nearest Neighbor (K -NN)," in *Proceedings of the International Conference on Mathematics, Geometry, Statistics, and Computation (IC-MaGeStiC 2021)*, pp. 163–166, 2022, doi: [10.2991/acsr.k.220202.031](https://doi.org/10.2991/acsr.k.220202.031).
- [6] P. Sananse, S. More, M. Kesarkar, and A. Nayak, "Loan Eligibility Prediction using Machine Learning Algorithms," *International Journal of Research and Analytical Reviews*, vol. 9, no. 2, pp. 374–378, May 2022.
- [7] J. M. Polgan *et al.*, "K-Nearest Neighbor and Random Forest Algorithms in Loan Approval Prediction," *J. Minfo Polgan*, vol. 13, no. 1, pp. 1307–1313, 2024, doi: [10.33395/jmp.v13i1.14345](https://doi.org/10.33395/jmp.v13i1.14345).
- [8] N. Uddin, K. Uddin, A. Uddin, M. Islam, A. Talukder, and S. Aryal, "International Journal of Cognitive Computing in Engineering An ensemble machine learning based bank loan approval predictions system with a smart application," *Int. J. Cogn. Comput. Eng.*, vol. 4, no. February, pp. 327–339, 2023, doi: [10.1016/j.ijcce.2023.09.001](https://doi.org/10.1016/j.ijcce.2023.09.001).
- [9] A. Alagic, N. Zivic, E. Kadusic, D. Hamzic, N. Hadzajlic, and M. Dizdarevic, "Machine Learning for an Enhanced Credit Risk Analysis : A Comparative Study of Loan Approval Prediction Models Integrating Mental Health Data," *Mach. Learn. Knowl. Extr. Artic.*, vol. 6, pp. 53–77, 2024, doi:

- [10.3390/make6010004](#).
- [10] T. K. Adarsh, R. Sinchana, K. P. C., and R. Uday, "Software Bug Prediction Using Machine Learning Approach," *International Journal for Research in Applied Science and Engineering Technology*, vol. 11, no. 12, pp. 1401–1405, Dec. 2023, doi: [10.22214/ijraset.2023.57626](#).
- [11] H. Purnamasidi, Salman, L. D. Bakti, and B. Imran, "Sistem Pakar Pemilihan Jenis Kredit Nasabah Menggunakan Metode *Forward Chaining* pada PT. Bank Rakyat Indonesia (Persero)," *Jurnal Kecerdasan Buatan dan Teknologi Informasi*, vol. 1, no. 3, pp. 1–10, Dec. 2022, doi: [10.69916/jkbt.v1i3.2](#).
- [12] I. Dimentieva, E. Lutfina, G. W. Saraswati, and R. M. Caturkusuma, "Integrating RAD and Design Thinking for Developing a Web-Based POS and Inventory Management System for MSMEs : A Case Study," *Inf. J. Ilm. Bid. Teknol. Inf. dan Komun.*, vol. 11, no. 1, pp. 80–88, 2026, doi: [10.25139/inform.v11i1.11205](#).
- [13] B. Imran, H. Hambali, A. Subki, Z. Zaeniah, A. Yani, and M. R. Alfian, "Data Mining Using Random Forest, Naïve Bayes, and Adaboost Models for Prediction and Classification of Benign and Malignant Breast Cancer," *J. Pilar Nusa Mandiri*, vol. 18, no. 1, pp. 37–46, 2022, doi: [10.33480/pilar.v18i1.2912](#).
- [14] B. Imran, E. Wahyudi, A. Subki, S. Salman, and A. Yani, "Classification of stroke patients using data mining with adaboost, decision tree and random forest models," *Ilk. J. Ilm.*, vol. 14, no. 3, pp. 218–228, 2022, doi: [10.33096/ilkom.v14i3.1328.218-228](#).
- [15] B. Imran, Zaeniah, Sriasih, S. Erniwati, and Salman, "Data Mining Using a Support Vector Machine , Decision Tree , Logistic Regression and Random Forest for," *J. Infokum*, vol. 10, no. 2, pp. 792–802, 2022.
- [16] H. Azis, F. Tangguh Admojo, and E. Susanti, "Analisis Perbandingan Performa Metode Klasifikasi pada Dataset Multiclass Citra Busur Panah," *Techno.Com*, vol. 19, no. 3, pp. 286–294, 2020, doi: [10.33633/tc.v19i3.3646](#).
- [17] J. N. Oruh and B. C. Iwuji, "a Hybrid Prediction Model for Classifying Student'S Academic Performance Using Voting Ensemble Method," *Fudma J. Sci.*, vol. 10, no. 3, pp. 364–376, 2026, doi: [10.33003/fjs-2026-1003-4751](#).
- [18] P. Srisuradetchai, "Posterior averaging with Gaussian naive Bayes and the R package RandomGaussianNB for big-data classification," *Front. Big Data*, vol. Volume 8-, 2025, doi: [10.3389/fdata.2025.1706417](#).
- [19] B. Shen, "E-commerce Customer Segmentation via Unsupervised Machine Learning," in *The 2nd International Conference on Computing and Data Science*, in CONF-CDS 2021. New York, NY, USA: Association for Computing Machinery, 2021. doi: [10.1145/3448734.3450775](#).
- [20] I. P. Putri, "Analisis Performa Metode K- Nearest Neighbor (KNN) dan Crossvalidation pada Data Penyakit Cardiovascular," *Indones. J. Data Sci.*, vol. 2, no. 1, pp. 21–28, 2021, doi: [10.33096/ijodas.v2i1.25](#).
- [21] V. T. Tarigan and A. Yusupa, "Perbandingan Algoritma Maching Learning dalam Analisis Sentimen Mobil Listrik di Indonesia pada Media Sosial Twitter/X," *J. Inform. Polinema*, vol. 10, no. 4, pp. 479–490, 2024, doi: [10.33795/jip.v10i4.5130](#).
- [22] C. Daah, A. Qureshi, I. Awan, and S. Konur, "Enhancing Zero Trust Models in the Financial Industry through Blockchain Integration: A Proposed Framework," *Electron.*, vol. 13, no. 5, 2024, doi: [10.3390/electronics13050865](#).