

Evaluation of Support Vector Machines and Adaptive Boosting in Classifying the Compliance Levels of Property and Building Taxpayers Using Receiver Operating Characteristic (ROC)

Saumina ^{1*}, Munirul Ula ^{2*}, Asrianda ^{3*}

*Department of Information Technology, Universitas Malikussaleh, Lhokseumawe, Indonesia
saumina2011adli@gmail.com ¹, munirul.ula@unimal.ac.id ², asrianda@unimal.ac.id ³

Article Info

Article history:

Received 2026-07-10

Revised 2026-08-01

Accepted 2026-08-11

Keyword:

*Adaptive Boosting,
Area Under Curve,
Land and Building Tax,
Receiver Operating
Characteristic,
Support Vector Machine.*

ABSTRACT

Taxpayer compliance is a critical factor in increasing Property Tax (PBB) revenue. A low level of compliance can reduce local government revenue, making accurate classification methods essential for identifying taxpayer compliance. This study aims to compare the performance of Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) in classifying property taxpayer compliance. The dataset consisted of 58,998 property tax records collected from Lhokseumawe City, covering the districts of Banda Sakti, Blang Mangat, Muara Dua, and Muara Satu. The research stages included data preprocessing, label encoding, Min-Max normalization, data splitting using 80:20 and 70:30 scenarios, model training, and performance evaluation using accuracy, precision, recall, F1-score, and Area Under the Curve (AUC). Under the 80:20 data split, SVM achieved an accuracy of 92.89%, precision of 93.12%, recall of 98.81%, F1-score of 95.87%, and AUC of 80.59%, while AdaBoost achieved an accuracy of 92.86%, precision of 93.11%, recall of 98.78%, F1-score of 95.86%, and AUC of 80.76%. Under the 70:30 data split, SVM achieved an accuracy of 93.02%, precision of 93.18%, recall of 98.90%, F1-score of 95.95%, and AUC of 80.32%, whereas AdaBoost achieved an accuracy of 92.99%, precision of 93.18%, recall of 98.87%, F1-score of 95.93%, and AUC of 80.97%. Overall, both methods demonstrated comparable classification performance, while AdaBoost exhibited slightly better discriminative capability based on the AUC values.



This is an open access article under the [CC-BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.

I. INTRODUCTION

Regional taxes constitute one of the primary sources of Regional Original Revenue (PAD) and play a vital role in supporting regional governance and development. One of the regional taxes that significantly contributes to PAD is the Rural and Urban Land and Building Tax (Pajak Bumi dan Bangunan Perdesaan dan Perkotaan, PBB-P2) [1]. The level of taxpayer compliance in fulfilling PBB-P2 payment obligations is a crucial factor in optimizing regional revenue. Low taxpayer compliance may prevent tax revenue from reaching its full potential, thereby affecting the capacity of local governments to finance regional development [2], [3].

In Lhokseumawe City, the administration of PBB-P2 is regulated under Lhokseumawe City Regional Regulation

(Qanun) Number 1 of 2024 concerning Regional Taxes and Regional Retributions [3]. Although PBB-P2 represents one of the most stable sources of local revenue, a considerable number of taxpayers still make payments after the due date or remain in arrears for multiple tax periods. Currently, taxpayer compliance assessment at the Regional Financial Management Agency (Badan Pengelolaan Keuangan Daerah, BPKD) of Lhokseumawe City is primarily conducted through administrative procedures and manual record-keeping, which limits the ability to effectively identify taxpayer behavioral patterns. Consequently, tax collection and monitoring strategies have not yet been fully developed based on taxpayer characteristics [4].

Recent advances in machine learning offer promising solutions to address these challenges through classification

techniques capable of identifying taxpayer compliance patterns using historical data. One of the most widely adopted algorithms is the Support Vector Machine (SVM), which has demonstrated excellent performance in handling high-dimensional datasets, constructing optimal hyperplanes, and achieving strong generalization capability [5]. Another effective approach is Adaptive Boosting (AdaBoost), an ensemble learning algorithm that improves classification performance by combining multiple weak classifiers into a strong classifier, enabling the model to iteratively correct previous classification errors [6]. The Support Vector Machine was chosen because it performs well in handling high-dimensional data and nonlinear relationships through the RBF kernel. Meanwhile, Adaptive Boosting was chosen because it can improve classification performance through an ensemble mechanism that iteratively corrects previous classification errors. Both methods were selected because they have distinct characteristics, making them interesting to compare in the context of taxpayer compliance classification. To evaluate classification performance more comprehensively, the Receiver Operating Characteristic (ROC) curve and the Area Under the Curve (AUC) are commonly employed, as they provide a more objective assessment of a model's ability to distinguish between positive and negative classes than accuracy alone [7] [8].

Several previous studies have demonstrated the effectiveness of machine learning techniques in improving taxpayer compliance classification. Binarto and Wibowo reported that the Support Vector Machine achieved an accuracy of 67.05% with a recall of 99.31% in taxpayer compliance classification [9]. Nafisa obtained an accuracy of 86% and a recall of 100% using the SVM algorithm for Land and Building Tax classification [10]. Meanwhile, H. Wibowo and Istianah found that Adaptive Boosting achieved an accuracy of 89.20% with an AUC value of 0.842, indicating excellent classification performance [11]. Previous studies have shown that both Support Vector Machines and Adaptive Boosting are capable of achieving good classification performance in various cases of taxpayer compliance. However, these studies used different dataset characteristics, attributes, and evaluation methods, so no conclusion has yet been reached regarding the most suitable method for classifying property tax compliance using real-world data from local governments. Furthermore, most studies still focus on accuracy scores without comparing the models' discriminative capabilities through ROC and AUC analyses. Therefore, this study was conducted to address this gap by directly comparing SVM and AdaBoost using the Lhokseumawe City property tax dataset.

The novelty of this study lies in the application and comparison of the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) methods to classify the compliance levels of Land and Building Tax (PBB) taxpayers using a real-world dataset of 58,998 taxpayers from four subdistricts in the city of Lhokseumawe. Unlike previous studies, which generally focused only on measuring accuracy

or used smaller-scale datasets, this study conducted a more comprehensive evaluation of model performance using the Receiver Operating Characteristic (ROC) curve and the Area Under the Curve (AUC). This approach provides a better understanding of each model's discriminatory ability to distinguish between compliant and non-compliant taxpayers, thereby generating more relevant information to support local government decision-making aimed at improving the effectiveness of tax administration and compliance.

This study aims to evaluate and compare the performance of the Support Vector Machine and Adaptive Boosting algorithms in classifying Land and Building Tax (PBB-P2) taxpayer compliance based on ROC analysis. The findings are expected to identify the classification method with the best predictive performance and provide practical recommendations for the Government of Lhokseumawe City in developing more effective, data-driven strategies for tax collection, monitoring, and improving taxpayer compliance.

II. METHODOLOGY

A. Research Stages

The research stages were conducted systematically and included problem identification, literature review, data collection, preprocessing, data segmentation, modeling, evaluation, and analysis of results. The flow of the research stages is shown in Figure 1.

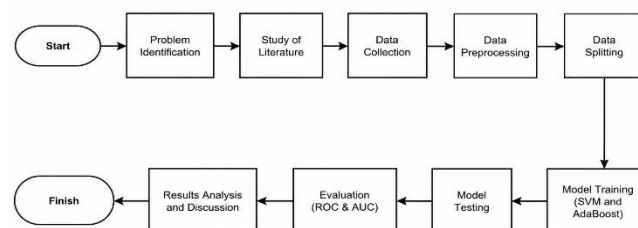


Figure 1. Research Stages

Based on Figure 1, the study began with problem identification to determine the research focus regarding the classification of Land and Building Tax (PBB) taxpayer compliance levels. Next, a literature review was conducted to establish a theoretical foundation regarding taxpayer compliance, the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) methods, as well as the Receiver Operating Characteristic (ROC) and Area Under the Curve (AUC) evaluation techniques. PBB data obtained from the Regional Financial Management Agency (BPKD) of Lhokseumawe City was then processed through a preprocessing stage that included data cleaning, variable selection, categorical data transformation using label encoding, and data normalization using the Min-Max Normalization method. These steps were carried out to produce structured data ready for use in the modeling process.

The dataset that had undergone the preprocessing stage was then divided into training and testing datasets using two

splitting scenarios: 80:20 and 70:30. The use of these two scenarios aimed to evaluate the model's consistency and generalization ability across different data proportions. The training data is used to build classification models using the SVM and AdaBoost methods, while the testing data is used to test the models' ability to classify taxpayer compliance levels into two classes: compliant and non-compliant.

The classification results were then evaluated using a confusion matrix to obtain accuracy, precision, recall, and F1-score values. Additionally, ROC curves and AUC values were used to measure each model's discriminative ability in distinguishing between compliant and non-compliant taxpayers. The evaluation results were subsequently analyzed to assess the performance of the SVM and AdaBoost methods and served as the basis for drawing research conclusions.

B. Research Dataset

This study uses 2025 Rural and Urban Land and Building Tax (PBB-P2) data for Lhokseumawe City obtained from the Lhokseumawe City Regional Financial Management Agency (BPKD). The dataset consists of 58,998 taxpayer records from the subdistricts of Banda Sakti, Blang Mangat, Muara Dua, and Muara Satu. The independent variables include the principal tax amount, the penalty amount, the total payment, the number of months in arrears, and the payment status, while the dependent variable is the taxpayer compliance rate, classified into two categories: compliant and non-compliant. The description of each variable is shown in Table I.

TABLE I
RESEARCH VARIABLES

Variable	Data Type	Role
Tax Principal Amount	Numerical	Input
Tax Penalty Amount	Numerical	Input
Total Payment	Numerical	Input
Payment Delay (Months)	Numerical	Input
Payment Status	Categorical	Input
Taxpayer Compliance Level (Compliant / Non-Compliant)	Categorical	Target

As shown in Table I, this study employs five input variables, namely tax principal amount, tax penalty amount, total payment, payment delay (months), and payment status. These variables represent taxpayers' payment characteristics and are used as input features for the classification models. The target variable is taxpayer compliance level, which is categorized into two classes: compliant and non-compliant. All input variables are utilized to train and evaluate the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) classification models. Next, the distribution of data counts across each compliance class is shown in Table II.

TABLE II
DATA DISTRIBUTION BASED ON COMPLIANCE LABELS

Compliance Label	Data Count
Compliant	6,526
Non-Compliant	52,472
Total	58,998

Based on Table II, the research dataset consists of 58,998 taxpayer records, comprising 52,472 non-compliant taxpayers (88.94%) and 6,526 compliant taxpayers (11.06%). This distribution indicates a class imbalance that reflects the actual compliance rate of property tax payers in Lhokseumawe City, with a predominance of taxpayers who have not fulfilled their payment obligations. Therefore, a preprocessing stage was conducted to prepare the data prior to developing the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) models. Model performance evaluation utilized not only accuracy but also the Confusion Matrix, Receiver Operating Characteristic (ROC), and Area Under the Curve (AUC) to provide a more comprehensive assessment of the models' ability to classify taxpayer compliance levels.

Based on the characteristics of the dataset, the number of months of payment delay, payment status, principal tax amount, total payment, and penalty amount were the primary attributes used by the models to distinguish taxpayer compliance levels. However, this study did not perform a quantitative feature importance analysis. Therefore, the individual contribution of each variable to the classification results remains an opportunity for future research.

C. Support Vector Machine (SVM)

Support Vector Machines (SVMs) are a type of supervised learning method used to solve classification and regression problems. In the classification process, SVMs construct an optimal decision boundary in the form of a hyperplane that can separate data into different classes with minimal error and good generalization ability [12], [13]. The main principle of SVM is to determine a hyperplane with the maximum distance from the nearest data points in each class. This distance is called the margin, while the data points closest to the hyperplane are called support vectors [14]. An illustration of how SVM works is shown in Figure 2.

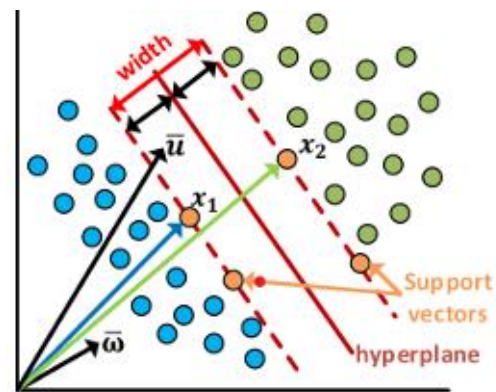


Figure 2. Illustration of the Support Vector Machine (SVM) Mechanism

As shown in Figure 2, the hyperplane is represented as a separating line between two data classes, while the lines on either side of the hyperplane indicate the margin boundaries. The data points closest to the separating line are the support vectors, which play a role in determining the position and orientation of the hyperplane. SVM aims to maximize the

margin so that the decision boundary is more optimal and has good performance in classifying new data. Mathematically, the SVM hyperplane function is formulated as follows [15]:

$$f(x) = w \cdot x + b = 0 \quad (1)$$

where w is the weight vector that determines the direction of the hyperplane, x is the input data vector, and b is the bias parameter. This function is used to determine the position of the data relative to the hyperplane. The value of the function serves as the basis for determining the class of each processed data point.

For data that cannot be linearly separated, SVM uses the concept of a soft margin, which allows for tolerance toward data points that fall within the margin or are misclassified. This approach enables the model to strike a balance between creating an optimal margin and controlling classification errors. In addition, SVM utilizes a kernel function to map data from the input space to a higher-dimensional feature space so that nonlinear data patterns can be separated more effectively [16].

This study uses the Radial Basis Function (RBF) kernel because the characteristics of taxpayer compliance data involve complex relationships between variables that cannot always be separated linearly. The RBF kernel has the ability to form decision boundaries that are flexible and adaptive to data patterns. The RBF kernel function is formulated as follows:

$$K(x_i, x_j) = \exp(-\gamma \|x_i - x_j\|^2) \quad (2)$$

where $K(x_i, x_j)$ represents the kernel value between the i -th and j -th data points, γ is the kernel parameter that controls the curvature of the decision boundary, and $\|x_i - x_j\|^2$ represents the squared Euclidean distance between the two data points. The value of γ determines the range of influence of each data point in forming the decision boundary. A larger γ value produces a more complex and localized decision boundary, whereas a smaller γ value generates a smoother and broader decision boundary.

The use of the RBF kernel in this study aims to identify nonlinear patterns in taxpayer compliance data based on the variables of tax principal amount, penalty amount, total payment, number of months of delinquency, and payment status [17]. By applying the principles of maximum margin, soft margin, and the RBF kernel, the SVM model is expected to accurately classify taxpayer compliance levels into compliant and non-compliant classes.

D. Adaptive Boosting (AdaBoost)

Adaptive Boosting (AdaBoost) is an ensemble learning algorithm in the supervised learning category designed to improve classification performance by combining multiple weak classifiers into a single strong classifier [18], [19]. The basic concept of AdaBoost was introduced by Freund and Schapire through a step-by-step adaptive learning mechanism. Each weak classifier is built sequentially, taking into account the classification errors produced by the model in the previous iteration. This mechanism allows the model to

focus its learning on data that is difficult to classify, thereby gradually reducing prediction errors [19], [20].

The AdaBoost learning process begins by assigning equal weights to all training data. In each iteration, a weak classifier is trained based on the distribution of these data weights to form an initial decision boundary. Data that are correctly classified receive smaller weights, while misclassified data are assigned larger weights in the next iteration. This increase in weight causes the subsequent weak classifier to focus more on data that were previously misclassified. Thus, each model formed seeks to correct the errors of the previous model [20], [21].

Each weak classifier makes a different contribution to the formation of the final model. The magnitude of this contribution is determined by the classification error rate it produces. A weak classifier with a low error rate will receive a larger weight, thereby exerting a more dominant influence on the final prediction result. Conversely, weak classifiers with high error rates are assigned smaller weights. The final classification result is obtained through a weighted combination of all weak classifiers, so the model's decision does not depend on a single classifier but is the result of aggregating several complementary models [21], [22]. The working mechanism of the AdaBoost method is shown in Figure 3.

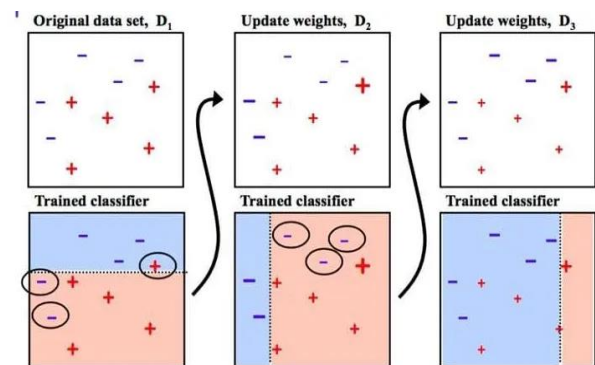


Figure 3. The Adaptive Boosting (AdaBoost) Process

Based on Figure 3, the AdaBoost process begins with the original dataset, where all training data are assigned the same initial weight. Assigning uniform weights indicates that, in the early stages, each data point has the same level of importance in the model-building process. The initial weight of each data point is formulated as follows [19]:

$$w_i = \frac{1}{N} \quad (3)$$

where w_i represents the initial weight of the i -th data point, N represents the total number of training data points, and i denotes the data index, where $i = 1, 2, \dots, N$. Based on this initial weight distribution, the first *weak classifier* is trained to establish an initial decision boundary. However, due to the limited classification capability of the *weak classifier*, some data points may still be misclassified.

The classification error is then calculated to determine the model error rate at each iteration. The error rate at the t -th iteration is formulated as follows [21]:

$$\varepsilon_t = \frac{\sum_{i=1}^N w_i I(h_t(x_i) \neq y_i)}{\sum_{i=1}^N w_i} \quad (4)$$

where ε_t represents the model error rate at the t -th iteration, w_i represents the weight of the i -th data point, $h_t(x_i)$ represents the prediction generated by the *weak classifier* at the t -th iteration, y_i represents the actual class label, and $I(\cdot)$ represents an indicator function that takes a value of 1 when the predicted class differs from the actual class and a value of 0 when the predicted class corresponds to the actual class.

The error rate is then used to determine the weight or contribution of each *weak classifier* to the construction of the final model. The weight of the *weak classifier* is calculated using the following equation:

$$\alpha_t = \frac{1}{2} \ln \left(\frac{1 - \varepsilon_t}{\varepsilon_t} \right) \quad (5)$$

where α_t represents the weight of the *weak classifier* at the t -th iteration, while ε_t represents the resulting error rate. The smaller the error rate ε_t , the greater the weight α_t . This indicates that a *weak classifier* with a lower error rate contributes more significantly to determining the final class.

The next stage involves updating the data weights to increase the model's attention to data points that are difficult to classify. Misclassified data points are assigned higher weights, thereby exerting a greater influence on the subsequent training process. Conversely, the weights of correctly classified data points are relatively reduced. The data weights are updated using the following equation:

$$w_{i,t+1} = w_{i,t} \cdot \exp(\alpha_t \cdot I(h_t(x_i) \neq y_i)) \quad (6)$$

where $w_{i,t+1}$ represents the weight of the i -th data point in the subsequent iteration, $w_{i,t}$ represents the weight of the i -th data point at the t -th iteration, α_t represents the weight of the *weak classifier*, and $I(h_t(x_i) \neq y_i)$ represents the classification error indicator. The weight-updating process enables misclassified data points to receive greater attention from the subsequent *weak classifier*, allowing the resulting decision boundary to be gradually improved.

The processes of training, error calculation, classifier weight determination, and data weight updating are performed iteratively until the predefined number of iterations is reached or the model achieves optimal performance. The final prediction is obtained by combining the predictions of all *weak classifiers* according to their respective weights. The final classification model is formulated as follows:

$$H(x) = \text{sign}(\sum_{t=1}^T \alpha_t \cdot h_t(x)) \quad (7)$$

where $H(x)$ represents the final prediction, T represents the total number of iterations, α_t represents the weight of the *weak classifier* at the t -th iteration, $h_t(x)$ represents the prediction generated by the *weak classifier* for data point x , and $\text{sign}(\cdot)$ represents the function used to determine the final class based on the weighted sum.

The application of AdaBoost in this study aims to classify taxpayer compliance levels into compliant and non-compliant

classes. The adaptive weight-updating mechanism allows the model to pay more attention to taxpayer data that is difficult to classify. However, AdaBoost is sensitive to outliers because data that are repeatedly misclassified will receive increasingly larger weights and potentially influence the learning process [23], [24]. Therefore, a preprocessing stage is conducted before model building to ensure that the data used are structured and suitable for the classification process.

E. Implementation of the SVM and AdaBoost Methods

The Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) methods were implemented to build a classification model for Land and Building Tax (PBB) compliance levels using data that had undergone preprocessing. The implementation process included data splitting, model training, model testing, and performance evaluation. The implementation flow for both methods is shown in Figure 4.

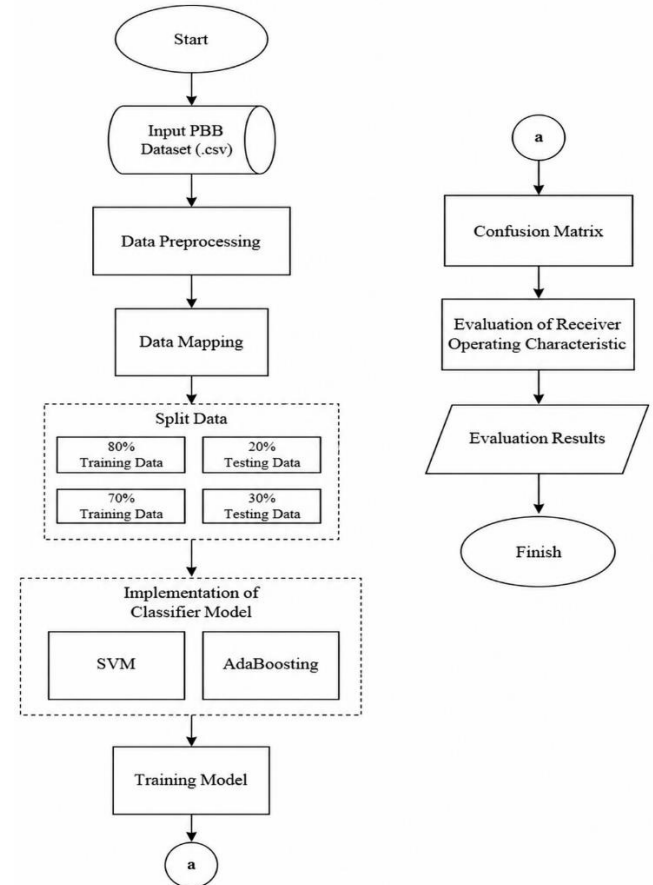


Figure 4. Implementation of the SVM and AdaBoost Methods

Based on Figure 4, the process begins by inputting the PBB dataset, which contains the variables: tax principal amount, penalty amount, total payment, number of months in arrears, and payment status. The data then undergoes a preprocessing stage that includes checking for missing values, removing duplicate data, transforming categorical data into

numerical form using label encoding, and normalizing the data to standardize the scale across variables.

The data that has undergone preprocessing is then split into training and testing sets using two splitting scenarios: 80:20 and 70:30. The training data is used to build classification models using the SVM and AdaBoost methods, while the testing data is used to evaluate the models' ability to classify taxpayer compliance levels into compliant and non-compliant classes. The use of two data splitting scenarios aims to evaluate the consistency of the model's performance and its generalization ability across different data proportions.

The built models were then tested using the test data to generate predictions of taxpayer compliance levels. The prediction results were evaluated using a Confusion Matrix to obtain accuracy, precision, recall, and F1-score values. Additionally, the Receiver Operating Characteristic (ROC) curve and the Area Under the Curve (AUC) value were used to measure the model's ability to distinguish between the compliant and non-compliant classes. The evaluation results were then analyzed to assess the performance of the SVM and AdaBoost methods and to determine which method has the best classification capability.

F. Receiver Operating Characteristic (ROC)

Receiver Operating Characteristic (ROC) is an evaluation method used to measure the ability of a classification model to distinguish between positive and negative classes [25]. The ROC curve illustrates the relationship between the *True Positive Rate* (TPR) and *False Positive Rate* (FPR) at various decision thresholds [26]. TPR represents the proportion of positive data points that are correctly classified, whereas FPR represents the proportion of negative data points that are incorrectly classified as positive [27]. The TPR and FPR values are formulated as follows:

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

$$FPR = \frac{FP}{FP+TN} \quad (9)$$

where TP (*True Positive*) represents the number of positive data points correctly classified, TN (*True Negative*) represents the number of negative data points correctly classified, FP (*False Positive*) represents the number of negative data points incorrectly classified as positive, and FN (*False Negative*) represents the number of positive data points incorrectly classified as negative.

The ROC curve is constructed by plotting the TPR values on the vertical axis against the FPR values on the horizontal axis at various decision thresholds. A model with good classification performance produces a curve that approaches the upper-left corner of the graph, indicating a high TPR and a low FPR. Conversely, a curve that approaches the diagonal line indicates that the model performs similarly to random classification.

The performance represented by the ROC curve is quantitatively measured using the *Area Under the Curve* (AUC). The AUC value represents the area under the ROC curve and indicates the model's overall ability to distinguish

between positive and negative classes without depending on a specific decision threshold. Mathematically, the AUC is formulated as follows:

$$AUC = \int_0^1 TPR(FPR) d(FPR) \quad (10)$$

where AUC represents the area under the ROC curve, TPR represents the proportion of positive data points correctly classified, FPR represents the proportion of negative data points incorrectly classified as positive, and the integration is performed over the FPR range from 0 to 1.

The AUC value ranges from 0 to 1. A value closer to 1 indicates a better discriminatory ability of the classification model. The interpretation of the AUC values used in this study is presented in Table III.

TABLE III
INTERPRETATION OF AUC VALUES

AUC Value	Model Performance Interpretation
AUC = 0.5	Equivalent to random classification
$0.6 \leq AUC < 0.7$	Poor
$0.7 \leq AUC < 0.8$	Fair
$0.8 \leq AUC < 0.9$	Good
$AUC \geq 0.9$	Excellent

Based on Table III, a model with a higher AUC value has better discriminatory ability in distinguishing between classes. The use of ROC and AUC in this study is relevant because the dataset has an imbalanced class distribution. These evaluation measures provide a more comprehensive assessment of model performance and determine the ability of the SVM and AdaBoost methods to distinguish between compliant and non-compliant taxpayers [28].

G. Confusion Matrix

A Confusion Matrix is an evaluation method used to measure the performance of a classification model by comparing the model's predicted classes with the actual classes [29]. The Confusion Matrix provides information regarding the number of data points that are correctly classified and those that are misclassified. Based on these results, model performance can be evaluated using several metrics, including accuracy, precision, recall, and F1-score [30].

A Confusion Matrix consists of four main components: True Positive (TP), True Negative (TN), False Positive (FP), and False Negative (FN). The general structure of a Confusion Matrix is presented in Table IV.

TABLE IV
CONFUSION MATRIX

Predicted Value	Actual Positive	Actual Negative
Positive	TP (<i>True Positive</i>)	FP (<i>False Positive</i>)
Negative	FN (<i>False Negative</i>)	TN (<i>True Negative</i>)

As presented in Table IV, True Positive (TP) represents the number of positive data points correctly predicted as

positive, whereas True Negative (TN) represents the number of negative data points correctly predicted as negative. False Positive (FP) represents the number of negative data points incorrectly predicted as positive, while False Negative (FN) represents the number of positive data points incorrectly predicted as negative [31].

Accuracy represents the proportion of all data points that are correctly classified by the model and is calculated using the following equation:

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (11)$$

Recall measures the model's ability to correctly identify all data points belonging to the positive class and is calculated as follows:

$$Recall = \frac{TP}{TP+FN} \quad (12)$$

Precision measures the proportion of correctly predicted positive data points among all data points predicted as positive and is calculated using the following equation:

$$Precision = \frac{TP}{TP+FP} \quad (13)$$

Furthermore, the *F1-score* is the harmonic mean of *precision* and *recall*, providing a balanced evaluation of both metrics. The *F1-score* is formulated as follows:

$$F1-Score = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (14)$$

where TP (True Positive) represents the number of positive data points correctly classified by the model, TN (True Negative) represents the number of negative data points correctly classified, FP (False Positive) represents the number of negative data points incorrectly classified as positive, and FN (False Negative) represents the number of positive data points incorrectly classified as negative.

The use of the Confusion Matrix in this study aims to provide a more detailed interpretation of model performance, particularly in identifying the types of classification errors. Therefore, the Confusion Matrix complements the ROC and AUC evaluations in comprehensively assessing the performance of the SVM and AdaBoost classification models.

III. RESULTS AND DISCUSSION

A. Data Preprocessing

The preprocessing stage was conducted to ensure that the Land and Building Tax (PBB) dataset was of high quality before being used in the classification process using the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) methods. This process aimed to improve data consistency, standardize attribute structures, and prepare the dataset for use in the training and testing phases of the classification model.

The data used consists of 2025 PBB taxpayer data obtained from the Regional Financial Management Agency (BPKD) of Lhokseumawe City, covering the subdistricts of Banda Sakti,

Blang Mangat, Muara Dua, and Muara Satu. During this stage, the dataset was examined to ensure there were no missing values or duplicate data, thereby maintaining data quality and consistency. The results of the inspection showed that all data were valid, so the number of data points before and after the preprocessing stage remained at 58,998, as presented in Table V.

TABLE V
PBB DATASET PREPROCESSING RESULTS

Dataset	Number of Data Before Preprocessing	Number of Data After Preprocessing
Banda Sakti	18,224	18,224
Blang Mangat	13,837	13,837
Muara Dua	18,898	18,898
Muara Satu	8,039	8,039
Total	58,998	58,998

Based on Table V, the number of data records in each district remained unchanged after the preprocessing stage. Banda Sakti District contained 18,224 records, Blang Mangat District contained 13,837 records, Muara Dua District contained 18,898 records, and Muara Satu District contained 8,039 records. Overall, the total number of records before and after preprocessing remained at 58,998. These results indicate that no data needed to be removed due to missing values, duplicate records, or inconsistencies in the data structure. Therefore, all data records were retained because they met the requirements for the analysis and classification processes.

The preprocessing stage also involved selecting relevant attributes for the classification process, namely Tax Base, Penalty, Total Payment, Months in Delinquency, and Payment Status, with Compliance Level as the target variable. Next, label encoding was performed on the categorical attributes so they could be processed by the classification algorithm. The Payment Status variable was converted to a value of 0 for "Unpaid" and 1 for "Paid," while the Compliance Level variable was converted to a value of 0 for "Compliant" and 1 for "Non-Compliant." After the mapping process was completed, all numerical attributes were normalized using the Min-Max Normalization method so that each variable had the same range of values and no single attribute dominated during the model training process.

This study did not apply class balancing (oversampling) techniques, such as the Synthetic Minority Oversampling Technique (SMOTE), because it aimed to preserve the original distribution of taxpayer data so that the classification results could accurately represent the actual compliance levels of taxpayers in Lhokseumawe City. The final output of the preprocessing stage is a dataset ready to be used as input for training and testing the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) models.

B. Data Mapping

The data mapping stage was conducted to transform categorical attributes into numerical values so that they could be processed by the *Support Vector Machine* (SVM) and *Adaptive Boosting* (AdaBoost) algorithms. The transformation was performed using the *label encoding* technique. In this process, the Payment Status categories “Unpaid” and “Paid” were assigned numerical values of 0 and 1, respectively. Meanwhile, the Compliance Label categories “Compliant” and “Non-Compliant” were represented by numerical values of 0 and 1, respectively.

The mapping process enabled categorical information to be represented in a numerical format without changing its original meaning. A sample of the dataset after the mapping process is presented in Table VI.

TABLE VI
SAMPLE DATA AFTER THE MAPPING PROCESS

Tax Principal (IDR)	Penalty (IDR)	Total Payment (IDR)	Payment Delay (Months)	Payment Status	Compliance Label
94,339	11,320	105,659	6	0	1
96,218	11,546	107,764	6	0	1
96,218	11,546	107,764	6	0	1
52,130	6,255	58,385	6	0	1
88,718	0	88,718	0	1	0
88,718	10,646	99,364	6	0	1
88,718	10,646	99,364	6	0	1
67,841	0	67,841	0	1	0
96,218	0	96,218	0	1	0
89,921	10,790	100,711	6	0	1

Based on Table VI, the Payment Status and Compliance Label attributes were successfully converted into numerical formats according to the requirements of the classification algorithms. This numerical representation enables the SVM and AdaBoost algorithms to process the data effectively. Therefore, the mapped dataset was ready for the subsequent stages of data splitting, model training, and model testing.

C. Confusion Matrix Evaluation

Model performance evaluation was conducted using a confusion matrix to measure the algorithm’s ability to classify taxpayer compliance levels based on predicted results and actual labels. The confusion matrix provides information on the number of correct and incorrect predictions, which can be used as a basis for calculating accuracy, precision, recall, and F1-score. The structure of the confusion matrix used in this study is presented in Table VII.

TABLE VII
CONFUSION MATRIX STRUCTURE

Actual	Predicted Compliant	Predicted Non-Compliant
Compliant	True Negative (TN)	False Positive (FP)
Non-Compliant	False Negative (FN)	True Positive (TP)

Notes:

1. True Positive (TP) is the number of non-compliant taxpayers successfully predicted as non-compliant.
2. True Negative (TN) is the number of compliant taxpayers who were correctly predicted to be compliant.

3. False Positive (FP) is the number of compliant taxpayers who were incorrectly predicted to be non-compliant.
4. False Negative (FN) is the number of non-compliant taxpayers who were incorrectly predicted to be compliant.

Based on the confusion matrix structure presented in Table VII, the classification results of the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) models were evaluated under two data split scenarios (80:20 and 70:30). The confusion matrices obtained from each model are presented in Tables VIII–IX and are used as the basis for calculating the accuracy, precision, recall, and F1-score of the proposed classification models.

1. Support Vector Machine (SVM) – 80:20 Data Split

TABLE VIII
CONFUSION MATRIX OF SVM (80:20)

Actual	Predicted Compliant	Predicted Non-Compliant
Compliant	1,199	721
Non-Compliant	118	9,762

Based on Table VIII, the Support Vector Machine (SVM) model correctly classified 9,762 non-compliant taxpayers (True Positive, TP) and 1,199 compliant taxpayers (True Negative, TN). However, 721 compliant taxpayers were incorrectly classified as non-compliant (False Positive, FP), while 118 non-compliant taxpayers were incorrectly classified as compliant (False Negative, FN). The model achieved an accuracy of 92.89%, indicating that the majority of taxpayer records were classified correctly.

2. Adaptive Boosting (AdaBoost) – 80:20 Data Split

TABLE IX
CONFUSION MATRIX OF ADABOOST (80:20)

Actual	Predicted Compliant	Predicted Non-Compliant
Compliant	1,198	722
Non-Compliant	121	9,759

Based on Table IX, the Adaptive Boosting (AdaBoost) model correctly classified 9,759 non-compliant taxpayers (TP) and 1,198 compliant taxpayers (TN). The classification errors consisted of 722 False Positives (FP) and 121 False Negatives (FN). The model achieved an accuracy of 92.86%, demonstrating a classification performance that was comparable to the SVM model under the 80:20 data split.

3. Support Vector Machine (SVM) – 70:30 Data Split

TABLE X
CONFUSION MATRIX OF SVM (70:30)

Actual	Predicted Compliant	Predicted Non-Compliant
Compliant	1,808	1,073
Non-Compliant	163	14,656

Under the 70:30 data split scenario, the SVM model correctly classified 14,656 non-compliant taxpayers (TP) and 1,808 compliant taxpayers (TN). Meanwhile, 1,073 compliant taxpayers were incorrectly classified as non-compliant (FP), and 163 non-compliant taxpayers were incorrectly classified as compliant (FN). The model achieved an accuracy of 93.02%, indicating excellent classification performance.

4. Adaptive Boosting (AdaBoost) – 70:30 Data Split

TABLE XI
CONFUSION MATRIX OF ADABOOST (70:30)

Actual	Predicted Compliant	Predicted Non-Compliant
Compliant	1,808	1,073
Non-Compliant	167	14,652

Based on Table XI, the AdaBoost model correctly classified 14,652 non-compliant taxpayers (TP) and 1,808 compliant taxpayers (TN). The classification errors consisted of 1,073 False Positives (FP) and 167 False Negatives (FN). The model achieved an accuracy of 92.99%, indicating a performance that was very close to that of the SVM model under the 70:30 data split scenario.

Based on the confusion matrix results, both Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) demonstrated strong classification performance by correctly identifying the majority of non-compliant taxpayers, as indicated by the high number of True Positives (TP) in both data split scenarios. Furthermore, the relatively low number of False Negatives (FN) indicates that only a small proportion of non-compliant taxpayers were incorrectly classified as compliant.

On the other hand, the number of False Positives (FP) was higher than that of False Negatives (FN) for both models. This indicates that the models were more likely to classify compliant taxpayers as non-compliant than vice versa. This behavior can be attributed to the class imbalance in the dataset, where the number of non-compliant taxpayer records substantially exceeded the number of compliant records.

To provide a quantitative assessment of the performance of both classification algorithms, the models were evaluated using accuracy, precision, recall, and F1-score metrics. The evaluation was conducted under two data split scenarios, namely 80:20 and 70:30, to compare the consistency of the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) models in classifying taxpayer compliance. The performance comparison of both methods is presented in Table XII.

TABLE XII
PERFORMANCE COMPARISON OF SVM AND ADABOOST

Method	Data Split	Accuracy (%)	Precision (%)	Recall (%)	F1-Score (%)
SVM	80:20	92.89	93.12	98.81	95.87
AdaBoost	80:20	92.86	93.11	98.78	95.86
SVM	70:30	93.02	93.18	98.90	95.95
AdaBoost	70:30	92.99	93.18	98.87	95.93

Based on Table XII, both algorithms demonstrated excellent classification performance under both data split scenarios. Under the 80:20 data split, the SVM model achieved an accuracy of 92.89%, which was slightly higher than the 92.86% achieved by AdaBoost. Similarly, under the 70:30 data split, SVM obtained the highest accuracy (93.02%), while AdaBoost achieved 92.99%. The precision, recall, and F1-score values of both methods were also very similar, indicating that the two algorithms exhibited comparable performance in classifying taxpayer compliance. Nevertheless, SVM consistently achieved slightly higher accuracy, recall, and F1-score values than AdaBoost across both evaluation scenarios, suggesting a marginally better overall classification performance.

The high accuracy values indicate that most instances were correctly classified by the models. However, the AUC values, which are around 0.80, indicate that the models' ability to distinguish between compliant and non-compliant taxpayers is classified as *Good*, but not yet perfect. This difference is influenced by the imbalanced class distribution, where the number of non-compliant taxpayers is substantially higher than the number of compliant taxpayers. Consequently, the accuracy metric tends to be dominated by the majority class, whereas the AUC provides a more comprehensive evaluation by measuring the models' ability to discriminate between the two classes regardless of class distribution.

D. Receiver Operating Characteristic (ROC) Evaluation

The Receiver Operating Characteristic (ROC) analysis was performed to evaluate the ability of the classification models to distinguish between compliant and non-compliant taxpayers. The evaluation was based on the True Positive Rate (TPR), False Positive Rate (FPR), and Area Under the Curve (AUC). The ROC evaluation results for the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) models under the two data split scenarios are presented in Table XIII.

TABLE XIII
ROC CURVE AND AUC EVALUATION RESULTS

Method	Data Split	TPR (%)	FPR (%)	AUC (%)	Classification Category
SVM	80:20	98.81	37.55	80.59	Good
AdaBoost	80:20	98.78	37.60	80.76	Good
SVM	70:30	98.90	37.24	80.32	Good
AdaBoost	70:30	98.87	37.24	80.97	Good

Based on Table XIII, both methods demonstrated good classification performance, with AUC values ranging from 80.32% to 80.97%. In both data split scenarios, AdaBoost achieved slightly higher AUC values than SVM, while both models obtained TPR values above 98%, indicating a strong ability to correctly identify non-compliant taxpayers. Overall, the performance difference between the two methods was relatively small; however, AdaBoost showed slightly better discriminative capability based on the AUC values.

To provide a visual representation of the models' ability to distinguish between compliant and non-compliant taxpayers, the ROC curves generated by SVM and AdaBoost using the 80:20 data-splitting scenario are presented in Figure 5.

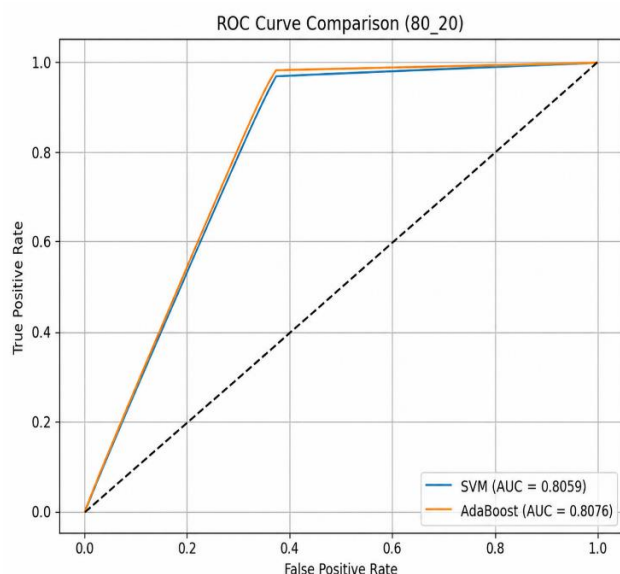


Figure 5. ROC Curves of SVM and AdaBoost Using the 80:20 Data Split

Figure 5 shows the ROC curves of the SVM and AdaBoost models under the 80:20 data split. Both models demonstrated good discriminative performance, with AUC values of 0.8059 and 0.8076, respectively. Although AdaBoost achieved a slightly higher AUC than SVM, the difference was minimal, indicating that both algorithms performed comparably in classifying taxpayer compliance.

The ROC curve evaluation was also conducted using the 70:30 data-splitting scenario to assess model performance with a larger proportion of testing data. The resulting ROC curves are presented in Figure 6.

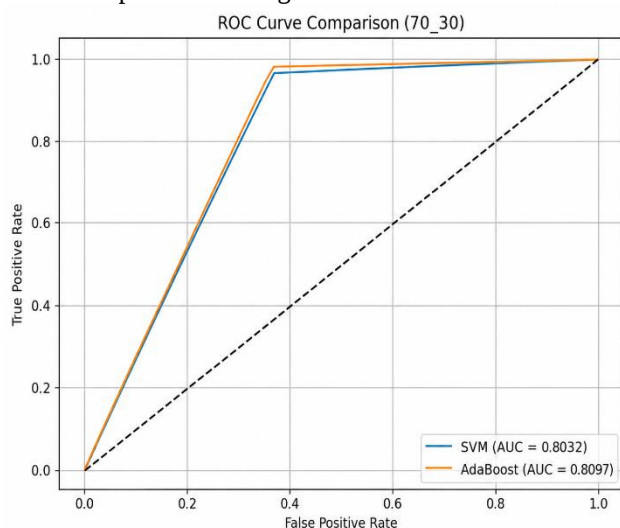


Figure 6. ROC Curves of SVM and AdaBoost Using the 70:30 Data Split

Figure 6 presents the Receiver Operating Characteristic (ROC) curves of the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) models using the 70:30 data split. Both models demonstrated good classification performance, with AUC values of 0.8032 for SVM and 0.8097 for AdaBoost. Although the difference in AUC values was relatively small, AdaBoost achieved a slightly higher AUC than SVM, indicating a marginally better ability to distinguish between compliant and non-compliant taxpayers. Overall, both models exhibited good and consistent classification performance under the 70:30 data split scenario.

Overall, the ROC and AUC evaluation results demonstrate that both SVM and AdaBoost have good discriminatory ability in classifying PBB taxpayer compliance. However, AdaBoost consistently achieved slightly higher AUC values than SVM in both data-splitting scenarios. Although the differences were relatively small, the results indicate that AdaBoost provided marginally better performance in distinguishing between compliant and non-compliant taxpayers across different classification thresholds.

Both algorithms produced nearly identical evaluation metrics because the dataset exhibited relatively clear class separation after the preprocessing and normalization stages. Furthermore, both methods utilized the same input attributes, resulting in highly similar classification decisions. A slight difference was observed in the AUC values, indicating that AdaBoost demonstrated marginally better discriminative performance than SVM in distinguishing between compliant and non-compliant taxpayers.

E. Classification Results of Land and Building Tax Compliance by District

The classification of Land and Building Tax (PBB) compliance by district was conducted to identify the distribution of compliant and non-compliant taxpayers across each district in Lhokseumawe City. This analysis provides an overview of taxpayer compliance levels in each region and can be used as supporting information for developing more targeted tax monitoring and collection strategies. The classification results using the 80:20 data-splitting scenario are presented in Table XIV and Table XV.

TABLE XIV
CLASSIFICATION RESULTS OF TAXPAYER COMPLIANCE BY DISTRICT USING SVM (80:20 DATA SPLIT)

District	Compliant	Non-Compliant	Total	Non-Compliant (%)	Compliant (%)
Banda Sakti	655	3,027	3,682	82.21	17.79
Blang Mangat	148	2,692	2,840	94.79	5.21
Muara Dua	399	3,235	3,634	89.02	10.98
Muara Satu	115	1,529	1,644	93.00	7.00

TABLE XV
CLASSIFICATION RESULTS OF TAXPAYER COMPLIANCE BY
DISTRICT USING ADABOOST (80:20 DATA SPLIT)

District	Compliant	Non-Compliant	Total	Non-Compliant (%)	Compliant (%)
Banda Sakti	655	3,027	3,682	82.21	17.79
Blang Mangat	151	2,689	2,840	94.68	5.32
Muara Dua	399	3,235	3,634	89.02	10.98
Muara Satu	114	1,530	1,644	93.07	6.93

Based on Tables XIV and XV, the classification results obtained using the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) models exhibited highly similar patterns across all districts. Blang Mangat recorded the highest percentage of non-compliant taxpayers, accounting for 94.79% using the SVM model and 94.68% using the AdaBoost model, indicating the lowest level of taxpayer compliance among the four districts. In contrast, Banda Sakti had the highest proportion of compliant taxpayers, reaching 17.79% under both classification models.

Meanwhile, Muara Dua and Muara Satu also showed consistent classification results, with only minor differences in the predicted numbers of compliant and non-compliant taxpayers between the two models. These findings indicate that SVM and AdaBoost exhibit comparable performance in identifying taxpayer compliance across districts. Therefore, the classification results can serve as valuable information for the Lhokseumawe City Government in identifying priority areas for taxpayer guidance, monitoring, and property tax collection efforts.

To examine the consistency of the classification results under the 70:30 data split, the prediction results of both models were further analyzed by district. This analysis presents the number and percentage of compliant and non-compliant taxpayers predicted by the Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) models across the districts of Lhokseumawe City. The classification results using the 70:30 data-splitting scenario are presented in Table XVI and Table XVII.

TABLE XVI
CLASSIFICATION RESULTS OF PROPERTY TAX (PBB)
COMPLIANCE BY DISTRICT USING THE SVM METHOD (70:30
DATA SPLIT)

District	Compliant	Non-Compliant	Total	Non-Compliant (%)	Compliant (%)
Banda Sakti	951	4,490	5,441	82.52	17.48
Blang Mangat	208	4,033	4,241	95.10	4.90
Muara Dua	627	4,957	5,584	88.77	11.23
Muara Satu	185	2,249	2,434	92.40	7.60

TABLE XVII
CLASSIFICATION RESULTS OF PROPERTY TAX (PBB)
COMPLIANCE BY DISTRICT USING THE ADABOOST METHOD
(70:30 DATA SPLIT)

District	Compliant	Non-Compliant	Total	Non-Compliant (%)	Compliant (%)
Banda Sakti	951	4,490	5,441	82.52	17.48
Blang Mangat	213	4,028	4,241	94.98	5.02
Muara Dua	627	4,957	5,584	88.77	11.23
Muara Satu	184	2,250	2,434	92.44	7.56

Based on Tables XIV and XV, the classification results under the 70:30 data split exhibited a pattern similar to that observed under the 80:20 data split. Blang Mangat recorded the highest percentage of non-compliant taxpayers, reaching 95.10% using the SVM model and 94.98% using the AdaBoost model, whereas Banda Sakti had the highest proportion of compliant taxpayers (17.48%) under both models. The classification results for Muara Dua and Muara Satu also showed only minor differences between the two algorithms. These findings indicate that both SVM and AdaBoost provide consistent classification results across all districts and can support the identification of priority areas for improving property tax (PBB) compliance in Lhokseumawe City.

Overall, the classification results indicate that both the SVM and AdaBoost models produced consistent predictions across all districts. Blang Mangat recorded the highest proportion of non-compliant taxpayers, whereas Banda Sakti had the highest proportion of compliant taxpayers. These findings suggest that both models can effectively support the identification of priority areas for improving property tax (PBB) compliance in Lhokseumawe City.

F. Classification Results of Land and Building Tax Compliance by District

The web-based system was developed to facilitate the analysis, visualization, and prediction of Land and Building Tax (PBB) taxpayer compliance using the SVM and AdaBoost methods. The system supports BPKD Lhokseumawe City in monitoring and evaluating taxpayer compliance more effectively and informatively. The system consists of three main interfaces: the main dashboard, the taxpayer compliance prediction page, and the compliance distribution map.

1. Main Dashboard

The main dashboard presents an overview of taxpayer compliance, including the total number of taxpayers, the number of compliant and non-compliant taxpayers, and the highest AUC value. It also displays taxpayer compliance distributions and model performance across districts.

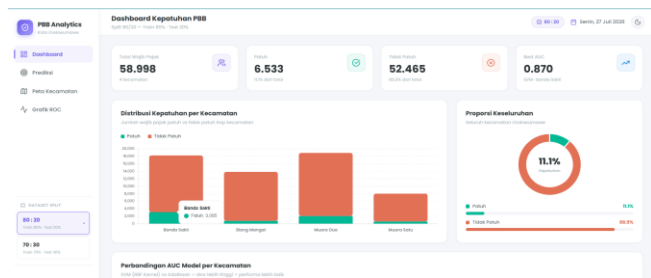


Figure 7. Main Dashboard Interface

The dashboard provides a comparison of SVM and AdaBoost performance based on the AUC values. The results indicate that AdaBoost achieved slightly higher AUC values than SVM in distinguishing between compliant and non-compliant taxpayers.

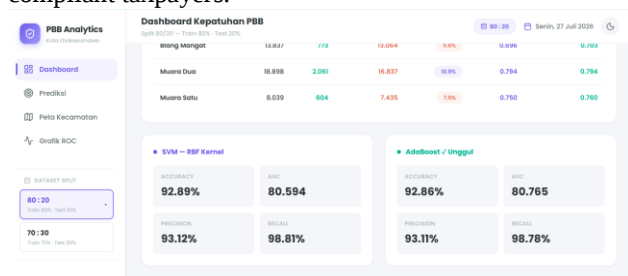


Figure 8. Model Performance Analysis on the Main Dashboard

2. Taxpayer Compliance Prediction Page

The prediction page classifies taxpayer compliance based on Tax Principal, Penalty, Total Payment, Payment Delay, and Payment Status. The system processes the input data using the trained classification model and displays the prediction as compliant or non-compliant.

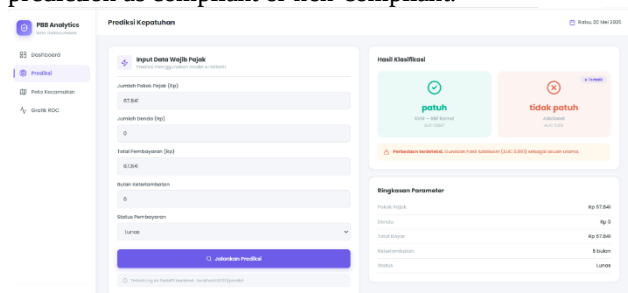


Figure 9. Taxpayer Compliance Prediction Interface

3. Taxpayer Compliance Distribution Map

The distribution map provides a geographical visualization of taxpayer compliance across the districts of Lhokseumawe City. It displays the number of compliant and non-compliant taxpayers, compliance rates, and the AUC values generated by SVM and AdaBoost.

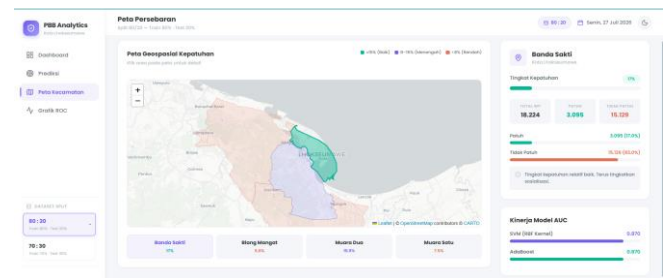


Figure 10. Taxpayer Compliance Distribution Map

Overall, the web-based system integrates machine learning models with interactive visualization to support taxpayer compliance analysis, prediction, and regional monitoring.

IV. CONCLUSION

Based on the implementation and evaluation results, both Support Vector Machine (SVM) and Adaptive Boosting (AdaBoost) demonstrated good performance in classifying Property Tax (PBB) compliance. Under the 80:20 data split, SVM achieved an accuracy of 92.89% with an AUC of 80.59%, while AdaBoost achieved an accuracy of 92.86% with an AUC of 80.76%. Under the 70:30 data split, SVM achieved an accuracy of 93.02% with an AUC of 80.32%, whereas AdaBoost achieved an accuracy of 92.99% with an AUC of 80.97%. Overall, both methods showed comparable classification performance, with AdaBoost exhibiting slightly better discriminative capability based on the AUC values.

The classification results indicate that Blang Mangat had the highest proportion of non-compliant taxpayers, whereas Banda Sakti had the highest proportion of compliant taxpayers. These findings can support the Lhokseumawe City Government in prioritizing tax collection, payment reminder, and taxpayer awareness programs.

This study has several limitations, including the use of a dataset collected only from Lhokseumawe City, the absence of k-fold cross-validation, and the lack of feature importance analysis and statistical significance testing. Future studies are recommended to use more diverse datasets, apply cross-validation techniques, and conduct more comprehensive analyses to improve the reliability and interpretability of the classification models.

REFERENCES

- [1] F. Hidayat, A. Frinaldi, and Asnil, "Tax Contribution to Regional Original Income (PAD) Tanah Datar District," *Sustain. Theory, Pract. Policy*, vol. 1, no. 1, pp. 40–52, 2023.
- [2] Pemerintah Republik Indonesia, "Undang-Undang Republik Indonesia Nomor 1 Tahun 2022 tentang Hubungan Keuangan antara Pemerintah Pusat dan Pemerintahan Daerah," 2022.
- [3] Pemerintah Kota Lhokseumawe, "Qanun Kota Lhokseumawe Nomor 1 Tahun 2024 tentang Pajak Daerah dan Retribusi Daerah," 2024. [Online]. Available: <https://lhokseumawekota.go.id/>
- [4] Badan Pengelolaan Keuangan Daerah Kota Lhokseumawe, "Data Pajak Bumi dan Bangunan Tahun 2025," 2025.
- [5] L. Chen, Y. Zhang, and H. Wang, "Support Vector Machine for Taxpayer Compliance Classification," *Int. J. Adv. Comput. Sci.*

- Appl.*, vol. 14, no. 5, pp. 215–223, 2023.
- [6] A. H. Salman and W. A. Al-Jawher, “Performance Comparison of Support Vector Machines, AdaBoost, and Random Forest for Classification,” *Int. J. Comput. Digit. Syst.*, vol. 15, no. 2, pp. 89–99, 2024.
- [7] J. Li, “Area under the ROC Curve has the most consistent evaluation for binary classification,” *PlosOne*, vol. 19, no. 12, pp. 1–28, 2024, doi: 10.1371/journal.pone.0316019.
- [8] A. M. Carrington, D. G. Manuel, P. W. Fieguth, and others, “Deep ROC Analysis and AUC as Balanced Average Accuracy, for Improved Classifier Selection, Audit and Explanation,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 45, no. 1, pp. 329–341, 2022, doi: 10.1109/TPAMI.2022.3145392.
- [9] R. Binarto and H. Wibowo, “Klasifikasi Kepatuhan Wajib Pajak PBB-P2 Menggunakan Support Vector Machine,” *J. Sist. Inf. Pajak*, vol. 10, no. 1, pp. 12–20, 2025.
- [10] N. Nafisa and others, “Compliance of Land and Building Taxpayers with Support Vector Machine,” in *Proceedings of the International Conference on Community Development*, 2024, pp. 210–217.
- [11] H. Wibowo and R. Istianah, “Evaluasi Kinerja Metode Ensemble Adaptive Boosting untuk Klasifikasi Kepatuhan PBB-P2,” *J. Data dan Pajak*, vol. 8, no. 1, pp. 33–42, 2025.
- [12] N. Aula, M. Ula, and L. Rosnita, “Analisis Sentimen Review Customer Terhadap Perusahaan Ekspedisi Jne , J & T Express Dan Pos Indonesia Menggunakan Metode Support Vector Machine (Svm) Analysis Of Customer Review Sentiment To Jne , J & T Express And Pos Indonesia Expedition Companies Using Svm Method,” vol. 9, no. 1, pp. 81–86, 2023.
- [13] Yohannes, A. R. M. Ezar, S. Devella, and Meiriyama, “Jurnal Teknologi Terpadu Klasifikasi Motif Songket Palembang Menggunakan Support Vector Machine Berdasarkan Histogram Of Oriented Gradients,” *J. Teknol. Terpadu Vol.*, vol. 9, no. 2, pp. 143–149, 2023.
- [14] I. K. Nti, O. Nyarko-boateng, F. A. Adekoya, and B. A. Weyori, “An empirical assessment of different kernel functions on the performance of support vector machines,” *Bull. Electr. Eng. Informatics*, vol. 10, no. 6, pp. 3403–3411, 2021, doi: 10.11591/eei.v10i6.3046.
- [15] Z. Abidin, W. Destian, and R. Umer, “Combining support vector machine with radial basis function kernel and information gain for sentiment analysis of movie reviews,” *J. Phys. Conf. Ser.*, pp. 1–5, 2020, doi: 10.1088/1742-6596/1918/4/042157.
- [16] I. A. Sapitri, M. Fikry, U. Islam, N. Sultan, and S. Kasim, “Pengklasifikasian sentimen ulasan aplikasi whatsapp pada google play store menggunakan support vector machine,” vol. 6, pp. 1–7, 2023, doi: 10.37600/tekinkom.v6i1.773.
- [17] D. Rahmawati and Y. S. Nugroho, “Penerapan Support Vector Machine untuk analisis kepatuhan pajak daerah,” *J. RESTI (Rekayasa Sist. dan Teknol. Informasi)*, vol. 6, no. 4, pp. 612–620, 2022.
- [18] R. R. Isnanto, A. Setiawan, and A. Nugroho, “Analisis Perbandingan Metode AdaBoost dan Support Vector Machine pada Klasifikasi Data Tidak Seimbang,” *J. Teknol. Inf. dan Ilmu Komput.*, vol. 10, no. 2, pp. 245–254, 2023.
- [19] Y. Freund and R. E. Schapire, “A Decision-Theoretic Generalization of On-Line Learning and an Application to Boosting,” *J. Comput. Syst. Sci.*, vol. 55, no. 1, pp. 119–139, 1997.
- [20] M. R. Fachrurrozi and E. Sutoyo, “Penerapan Adaptive Boosting untuk Meningkatkan Akurasi Klasifikasi Data Mining,” *J. RESTI (Rekayasa Sist. dan Teknol. Informasi)*, vol. 7, no. 1, pp. 45–53, 2023.
- [21] T. G. Dietterich, “Ensemble Methods in Machine Learning,” in *Lecture Notes in Computer Science*, Springer, 2000, pp. 1–15.
- [22] A. Wibowo and I. Kurniawan, “Optimasi Klasifikasi Menggunakan Algoritma AdaBoost pada Data Kompleks,” *J. Inform. Mulawarman*, vol. 18, no. 3, pp. 189–197, 2023.
- [23] H. Zhang, J. Zhang, and Q. Wang, “Robust AdaBoost Algorithms Against Noisy Data,” *Inf. Sci. (Ny.)*, vol. 512, pp. 379–390, 2020.
- [24] S. Sun, C. Zhang, and G. Yu, “A Bayesian Perspective on AdaBoost,” *Pattern Recognit.*, vol. 48, no. 12, pp. 3823–3833, 2019.
- [25] S. Kurnia, N. Nurdin, and A. Khaidar, “Perbandingan Metode Machine Learning Menggunakan Metode Support Vector Machine Dan Artificial Neural,” *J. Inform. Kaputama*, vol. 9, no. 2, pp. 87–94, 2025.
- [26] J. Li, “Area under the ROC Curve has the Most Consistent Evaluation for Binary Classification,” *Patterns*, 2024, doi: 10.1016/j.patter.2024.101147.
- [27] O. Rainio, “Evaluation metrics and statistical tests for machine learning,” *Sci. Rep.*, pp. 1–14, 2024, doi: 10.1038/s41598-024-56706-x.
- [28] N. S. Basuni and A. M. Siregar, “JURNAL RESTI Comparison of the Accuracy of Drug User Classification Models Using,” *J. RESTI (Rekayasa Sist. dan Teknol. Informasi)*, vol. 5, no. 158, pp. 1348–1353, 2026.
- [29] R. Nurhidayat and K. E. Dewi, “Penerapan Algoritma K-Nearest Neighbor Dan Fitur Ekstraksi N-Gram Dalam Analisis Sentimen Berbasis Aspek,” *KOMPUTA J. Ilm. Komput. dan Inform.*, vol. 12, no. 1, pp. 91–100, 2023.
- [30] Y. Afrillia, L. Rosnita, D. Siska, M. Rigayatsyah, and Nurqamarina, “Analisis Sentimen Ciutan Twitter Terkait Penerapan Permendikbudristek Nomor 30 Tahun 2021 Menggunakan TextBlob dan Support Vector Machine,” vol. 6, no. 2, pp. 387–394, 2022.
- [31] S. Sathyanarayanan and B. R. Tantri, “Confusion Matrix-Based Performance Evaluation Metrics,” vol. 27, no. 4, 2024.