

Analysis Combination of Machine Learning Classification with Feature Selection Technique for Lecturer Performance Measurement

Ni Luh Putri Srinadi¹, I Nyoman Suraja Antarajaya¹, Luh Putu Wiwien Widhyastuti¹, Dandy Pramana Hostiadi¹, Erma Sulistyo Rini¹, Pharan Chawaphan²

¹Institut Teknologi dan Bisnis STIKOM Bali, Denpasar, Indonesia

²Rajamangala University of Technology, Thanyaburi, Thailand

Article Info

Article history:

Received August 07, 2024

Revised February 22, 2025

Accepted March 06, 2025

Keywords:

Classification;

Feature selection;

Lecturer performance;

Machine learning;

Work performance analysis.

ABSTRACT

Machine learning-based classification techniques are widely utilized for accurate analysis in various fields. This study focuses on assessing lecturer performance in higher education to enhance teaching standards and produce high-quality learning outcomes. Previous studies have employed multi-parameter approaches, such as statistical correlation analysis, but these methods fail to achieve optimal accuracy and precision due to limited alignment with data characteristics. This research proposes a lecturer performance measurement model by evaluating three machine learning algorithms: k-Nearest Neighbors (k-NN), Decision Tree, and Naïve Bayes. The model integrates three feature selection techniques to improve classification performance: ANOVA, Information Gain, and Chi-Square. The study aims to enhance classification accuracy and assess the impact of feature selection techniques on performance metrics. A significant contribution of this research is introducing a dynamic feature selection approach tailored to data characteristics, which improves classification model performance. The methodology comprises three main stages: data loading and measurement of relevant parameters; data preprocessing, including filtering, cleaning, transformation, normalization, and feature selection; and performance evaluation using a machine learning-based classification approach. Experimental results demonstrate that the Decision Tree algorithm combined with Chi-Square feature selection achieved an accuracy of 0.887, precision of 0.903, recall of 0.887, and F1-score of 0.884. The proposed model provides a reliable framework for evaluating lecturer performance and can be utilized to recognize and reward high-performing lecturers effectively.

Copyright ©2025 The Authors.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Ni Luh Putri Srinadi, +62 812-3602-702,

Department of Information Systems,

Institut Teknologi dan Bisnis STIKOM Bali, Denpasar, Indonesia,

Email: putri@stikom-bali.ac.id.

How to Cite:

Author, "Analysis Combination of Machine Learning Classification with Feature Selection Technique for Lecturer Performance Measurement", *MATRIK: Jurnal Manajemen, Teknik Informatika, dan Rekayasa Komputer*, Vol.24, No.2, pp. xxx-xxx, March, 2025. This is an open access article under the CC BY-SA license (<https://creativecommons.org/licenses/by-sa/4.0/>)

1. INTRODUCTION

The pervasive use of artificial intelligence-based (AI) technology in data and information processing [1, 2], underscores its significance in modern data analysis. AI aids companies in illustrating business development trends, identifying development patterns, and providing decision-making information [3]. Various data analysis techniques, including classification approaches, clustering, association, and a combination of statistical data processing [2], can be employed with the help of AI. However, optimizing artificial analysis models such as classification requires the correct technique. One of them is choosing the proper feature selection method to produce an accurate analysis.

The data processing results in the classification model depend on the characteristic patterns of the data [4]. For example, measuring the performance of lecturers in a university has various variables sub-parameters [5, 6]. Previous research measured the quality of lecturer performance from several variables such as motivation [7], training [8, 9], organizational culture [7], work environment [10], leadership behaviour, organizational commitment [11] and satisfaction [12]. There are sub-parameters in each variable that can be used to measure the quality of lecturer performance, especially in educational institutions [13–16]. The performance variable has sub-parameters of responsibility, competency, motivation, coaching, work results, quality and quantity [17, 18]. The work motivation variable contains the sub-parameters of desires, goals, needs, effort, attitude, facilities, and teamwork [18, 19]. Meanwhile, the training variable consists of the sub-parameters ability, education, training, knowledge, skills, attitude and flexibility. Of all the sub-parameters, not all educational institutions have complex measurements and assessments, so they cannot be modelled using a data analytical approach [20]. The gap lies in the fragmented adoption and application of these sub-parameters across institutions, leading to inconsistencies in performance assessment models. Furthermore, existing research often focuses on isolated variables rather than a holistic integration of all relevant parameters into a unified analytical model. Addressing this gap requires a systematic framework that captures the complexity and variability of these parameters while ensuring scalability and adaptability across different educational environments.

In case studies of lecturer performance measurement, the analysis technique that is often used is statistical approaches, such as correlation and similarity [14]. Measurement analysis can produce connections between variables or sub-parameters represented in unidirectional, bidirectional correlation or causality. Research [21], shows that training positively influences the quality of lecturers' work, where only two measurement variables were used, namely organizational commitment and work motivation. Specifically, the analysis needs to explain how significant the influence is on the training variables separately. In [21], explains that training significantly influences the quality of lecturers' work. The research results show the influence of training on lecturers' performance, which states that $t_{count} > t_{table}$, with $2.452 > 1.971$, is in the H_0 rejection area, so H_a is accepted. Training has a partially significant effect on lecturers' performance. Previous research succeeded in analyzing the quality of lecturer performance using a statistical approach, but the analysis model requires accurate and optimal measurement evaluation based on data characteristics.

This paper proposes a model for measuring the quality of lecturer performance using machine learning classification methods and improving by implementing feature selection methods obtained from the best-comparing results. Previous research has yet to use analysis of the quality of lecturer performance involving classification. The research aims to analyze the influence of feature selection approaches and obtain the best selection approach that can increase the accuracy of the classification model. The novelty of the research is knowing a dynamic feature selection approach based on the performance produced by the classification method. This model analyzes three feature selection approaches: chi-square, ANOVA and information gain. Besides, the classification models used and analyzed are k-NN, Naïve Bayes and Decision Tree. The use of three classification models is because they can produce the best performance in learning data patterns, especially in prediction [2, 15, 16, 22]. The model proposed in this paper can be used as a basis for a university to assess lecturer performance and as a reference for determining rewards.

2. RESEARCH METHOD

This study aims to measure lecturer performance using machine learning classification improved using a feature selection approach. The model has three main stages: data load, which determines three data variables from questionnaires, including work motivation, work coaching, and work performance. The second stage is the data preprocessing stage, namely data preparation, which consists of filtering, cleaning, transformation, normalization, and feature selection. The last is the classification stage, namely measurement modeling using machine learning-based classification. The three classification models used are Naïve Bayes, decision tree, and k-NN. The three classification models used are Naïve Bayes, decision tree, and k-NN which are commonly used algorithms for evaluating lecturer performance due to their suitability for classification and predictive analysis tasks. Naïve Bayes excels in handling textual feedback and sentiment analysis, while Decision Trees provide interpretable results and highlight key performance factors from structured data. k-NN, on the other hand, effectively predicts lecturer performance based on similarity to other lecturers with known outcomes. These algorithms are chosen for their simplicity, adaptability to various data types, and ease of implementation, making them valuable tools in educational performance evaluation systems [3, 17]. The proposed model is shown in Figure 1.

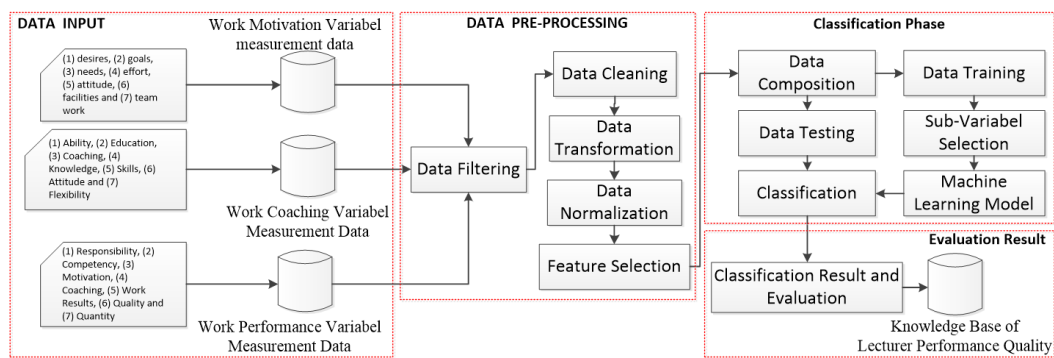


Figure 1. Proposed Model Overview

2.1. Problem Definition and Notation

The model contains three main variables for measuring the quality of lecturers' performance: coaching, work motivation, and performance results. The performance variables include the sub-parameter's responsibility, competency, motivation, coaching, work results, quality, and quantity. The work motivation variable contains the sub-parameters of desires, goals, needs, effort, attitude, facilities, and teamwork. Meanwhile, the coaching variable consists of the sub-parameters ability, education, coaching, knowledge, skills, attitude, and flexibility.

In this study, the data used is real data on a campus. Data collection was carried out by giving a questionnaire to each lecturer. The University leadership assessed the results of the questionnaire answers, which became the basis for labeling the quality of performance. The collected and labelled data can be unstructured data, so pre-processing techniques are needed for the form of data to be processed in a classification model. Choosing the right use of sub-parameters as features can influence the results of classification measurements, so it is necessary to determine the ideal number of features in the classification model for measuring the quality of lecturer performance. Classification models, such as decision trees, logistic regression and naïve Bayes, have different performance results and can be influenced by pre-processing techniques and feature selection. The examples questionnaire shows in Table 2.

Table 1. Example Performance Questionnaire

Variabels	Question	Value Answer
Responsibility	How often do you complete assigned tasks on time?	Never, Rarely, Sometimes, Often, always
	How do you ensure that your work aligns with institutional goals?	Poorly, Fairly, Well, Very Well, Exceptionally Well
	How do you handle accountability for your decisions and actions?	Not at all, Rarely, Sometimes, Often, Always
Competency	How confident are you in your ability to deliver lectures effectively?	Not confident, slightly confident, Neutral, Confident, Very confident
	How do you rate your subject knowledge compared to institutional standards?	Poor, Fair, Good, Very Good, Excellent
	How often do you seek opportunities to enhance your professional skills?	Never, Rarely, Sometimes, Often, Always
Desires and Goals	How motivated are you to achieve your professional goals?	Not motivated, slightly motivated, Neutral, Motivated, Highly motivated
Needs and effort	How much effort do you put into preparing your teaching materials?	None, Minimal, Moderate, High, Very High
Teamwork	How well do you collaborate with colleagues to achieve academic objectives?	Poorly, Fairly, Well, Very Well, Exceptionally Wel
Knowledge, Skills, Attitude, and Flexibility	I am adaptable to changes in teaching methods or technology. (Agreement)	1 = Never, 2 = Rarely, 3 = Sometimes, 4 = Often, 5 = Always
...	...	...

In this paper, notation is used to describe the proposed model. The performance variable P consists of seven sub-parameters denoted as p , namely responsibility (p_1), Competency (p_2), motivation (p_3), Coaching (p_4), work result (p_5), quality (p_6) dan quantity (p_7). Thus, it denoted as $p \in P, P = p_1, p_2, p_3, p_4, p_5, p_6, p_7$. The work motivation variable (M) consists of 7 sub-parameters denoted as m , namely desires (m_1), goals (m_2), needs (m_3), effort (m_4), attitude (m_5), facilities (m_6), teamwork (m_7). So, it is written as $m \in M, M = \{m_1, m_2, m_3, m_4, m_5, m_6, m_7\}$. The coaching variable (T) has seven sub-parameters (t) namely ability (t_1), education (t_2), training (t_3), knowledge (t_4), skills (t_5), attitude (t_6) dan flexibility (t_7). So, it is written as $t \in T, T = \{t_1, t_2, t_3, t_4, t_5, t_6, t_7\}$.

2.2. Data Filtering

Data filtering is an essential preprocessing step where raw questionnaire data is standardized into a structured tabular format. This process ensures consistency and prepares the data for analysis. Each questionnaire response is transformed into a specific data point, with values accurately mapped into table cells. Sub-parameter questions, often grouped under broader categories, are converted into distinct features or columns in the dataset. Different response types, such as numerical ratings, categorical choices, or text inputs, are uniformly formatted for easier interpretation. As a result, the filtered data becomes clean, organized, and suitable for statistical analysis, machine learning, or visualization.

2.3. Data Cleansing

Data cleansing is a critical preprocessing technique focused on handling incomplete or inconsistent data within a dataset to ensure accuracy and reliability in analysis. Specifically, this process involves systematically examining each column to identify missing or empty values that could compromise the quality of insights derived from the data. When null values are detected, the data record removes entire rows. For example, if a questionnaire response lacks critical answers across key sub-parameters, that respondent's data may be excluded to prevent skewed results. This approach helps maintain data integrity by ensuring that only complete and reliable entries are included in further analysis. Ultimately, data cleansing minimizes errors, reduces noise, and improves the overall quality and trustworthiness of the dataset for statistical modelling or machine learning processes.

2.4. Data Transformation

Data transformation is an essential step in preprocessing, where categorical data is converted into numerical data to ensure compatibility with statistical and machine-learning models. This process often uses the one-hot encoding technique, which generates binary features for each category in a variable. Each unique category is represented numerically without introducing any unintended ordinal relationships. The transformation increases the number of features based on the unique categories present in the original data. These newly created numerical features are essential for algorithms that require numerical input for accurate processing. Ultimately, this step prepares the data for effective analysis, modelling, and interpretation in subsequent stages.

2.5. Data Normalization

Data normalization is a crucial preprocessing step aimed at standardizing the scale of values across different columns in a dataset to ensure consistency and improve model performance. This process adjusts numerical data to fall within a specific range, typically between 0 and 1, without distorting the relationships between data points. By doing so, normalization eliminates biases caused by varying scales or units of measurement across features, ensuring that no single feature dominates the analysis or modelling process. The technique is essential for algorithms that rely on distance-based calculations, as unnormalized data can lead to inaccurate results. Normalization enhances the comparability of features, reduces computational complexity, and speeds up the convergence of machine-learning models during training. Ultimately, this step ensures that all numerical data contributes equally to the analysis, improving the accuracy and reliability of the results.

2.6. Feature Selection

Feature selection is a critical stage in data preprocessing where the most relevant and significant features are chosen from the set extracted during the data transformation phase. This process aims to reduce dimensionality, eliminate irrelevant or redundant features, and improve the efficiency and accuracy of machine-learning models. Various feature selection techniques are applied to evaluate the importance of each feature, and only the most impactful ones are retained for further analysis. This research compares three

widely used approaches, chi-square, ANOVA, and information gain, to determine their effectiveness in identifying the best subset of features. Each method uses different statistical or information-theoretic criteria to evaluate the relationship between features and the target variable. The selected features are crucial in enhancing model performance, reducing overfitting, and minimizing computational complexity in the subsequent analysis stages.

2.7. Classification

In this section, each sub-parameter of the three variables, motivation, coaching, and performance, will be modeled using a machine learning-based measurement model. The data used in modeling is training data, where there are naïve Bayes, decision trees and k-NN methods will be used as models for measuring the quality of lecturers' work. Then, the test data is used to test and assess the quality of the lecturer's work. This research shows the k-NN method in equation 1.

$$\hat{y} = \underset{b}{\operatorname{argmax}} \sum_{i=1}^{k=5} 1(y_i = b) \quad (1)$$

In this paper, the classification is carried out by determining k of 5. The variable $\hat{y}$ is the predicted class for a data point in the data distribution X. b iterates through all possible class labels. $1(y_i = b)$ is an indicator function that returns one if y_i equals c, otherwise. k is the number of nearest neighbors to consider. The second classification algorithm used is Naïve Bayes, with adopted the Bayesian concept. The naïve Bayes equation is shown in equation 2.

$$P(A|B) = \frac{P(B|A) \cdot P(A)}{P(B)} \quad (2)$$

The Equation 2 represents Bayes' Theorem, a fundamental concept in probability theory and statistics. It describes the probability of the event A occurring, given that the event B has already occurred. In this formula, $P(A|B)$ is the posterior probability, representing the updated probability of A given B. $P(B|A)$ is the likelihood, showing the probability of observing B given that A has occurred. $P(A)$ is the prior probability, representing the initial probability of A before considering B. Lastly, $P(B)$ is the marginal probability, which represents the total probability of B across all possible scenarios. The third algorithm used in this paper is the decision tree, which is shown in equation 3.

$$E(S) = \sum_i^N -P_i * \log_2(P_i) \quad (3)$$

The Equation 3 represents the entropy of a dataset, which measures the level of uncertainty or impurity within the data. Here, P_i refers to the probability or proportion of class C_i in the sample dataset, indicating how often a particular class appears. The term $\log_2(P_i)$ quantifies the amount of information needed to describe the class distribution. The sum across all possible classes gives the overall entropy, with higher values indicating more mixed or uncertain data. The lower the entropy, the more pure or homogenous the data is, suggesting less uncertainty in classification. This entropy calculation is commonly used in decision tree algorithms to determine how best to split the data, aiming to reduce uncertainty with each decision. It is shown in Equations 4 and 5.

$$S = x_1, x_2, \dots x_k \quad (4)$$

$$P_i = \frac{\sum x_k \in C_i}{S} \quad (5)$$

3. RESULT AND ANALYSIS

In the research, data processing was carried out using a computer with Intel Core i5 processor specifications, 8 GB RAM and 500 GB storage capacity. The tools used are orange data mining. The data used is lecturer assessment data with data contained in [14]. There are 21 sub-parameters from 3 variables, known as data features. Variable descriptions and data sub-parameters are shown in Table 2, and examples of lecturer performance assessment data are shown in Table 3.

Table 2. Variable and Parameters Description

Variabel	Sub-parameter / Features	Type	Values
Work Performance	<i>Responsibility</i> (p_1)	Categorical	Average, Excellent, Good, Poor
	<i>Competency</i> (p_2)	Categorical	Average, Very Good
	<i>Motivation</i> (p_3)	Numeric	0-100
	<i>Coaching</i> (p_4)	Numeric	0-100
	<i>Work Results</i> (p_5)	Numeric	0-100
	<i>Quality</i> (p_6)	Categorical	Average, Excellent
	<i>Quantity</i> (p_7)	Numeric	0-100
Work Motivation	<i>Desires</i> (m_1)	Numeric	0-100
	<i>Goals</i> (m_2)	Numeric	0-100
	<i>Needs</i> (m_3)	Categorical	Average, Very Good
	<i>Effort</i> (m_4)	Numeric	0-100
	<i>Attitude</i> (m_5)	Numeric	0-100
	<i>Facilities</i> (m_6)	Numeric	0-100
	<i>Teamwork</i> (m_7)	Numeric	0-100
Work Coaching	<i>Ability</i> (t_1)	Numeric	0-100
	<i>Education</i> (t_2)	Categorical	Master-Degree
	<i>Training</i> (t_3)	Numeric	0-100
	<i>Knowledge</i> (t_4)	Numeric	0-100
	<i>Skills</i> (t_5)	Numeric	0-100
	<i>Attitude</i> (t_6)	Numeric	0-100
	<i>Flexibility</i> (t_7)	Numeric	0-100
Class Label		Categorical	Excellent, Fair, Good

Table 3. Example of Questionnaire Data

Data	p_1	p_2	$p \dots$	p_7	m_1	m_2	$m \dots$	m_7	t_1	t_2	$t \dots$	t_7	Class-label
D1	Excelent	Average	...	71	77	79	...	81	84	Master Degree	...	71	Good
D2	Poor	Very Good	...	84	95	88	...	85	81	Master Degree	...	84	Fair
D3	Excelent	Average	...	91	79	71	...	73	78	Master Degree	...	81	Fair
D4	Average	Average	...	85	83	83	...	79	88	Master Degree	...	86	Good
D5	Excelent	Very Good	...	79	93	95	...	71	84	Master Degree	...	85	Excellent
D6	Poor	Average	...	71	84	86	...	82	75	Master Degree	...	87	Excellent
...	...	...	...	...	...	...	...	...	...	...	...	...	...

The next stage is data normalization, where the value of each numerical feature is converted into data in the 0-1 range. This research compares three feature selection approaches: Chis-square, ANOVA, and Information Gain. In feature selection, the model used 75% of the total features in the study and obtained 20 features used in modeling. The results show that the information gain approach obtained m_1 as the best feature, with a score of 0.125, and m_6 as the 20th feature, with a score of 0.017. The ANOVA method obtained the feature m_1 as the best feature with a score of 5.343, and m_6 is the 20th feature with a score of 0.321. The Chi-Square method obtained m_1 as the first best feature with a score of 6.403, and the $m_3 = \text{verygood}$) as the 20th feature with a score of 0.412. Of the three feature selection methods, the best feature based on data characteristics is the feature m_1 . The results of the comparative analysis of feature selection methods are shown in Figure 2.

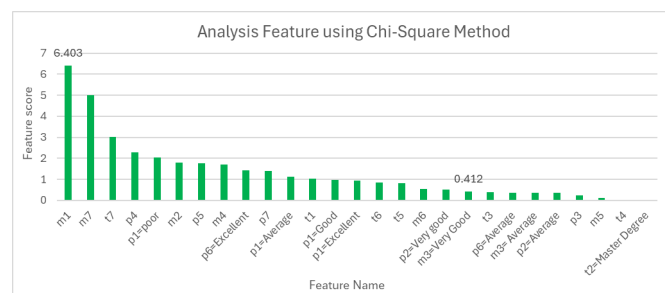
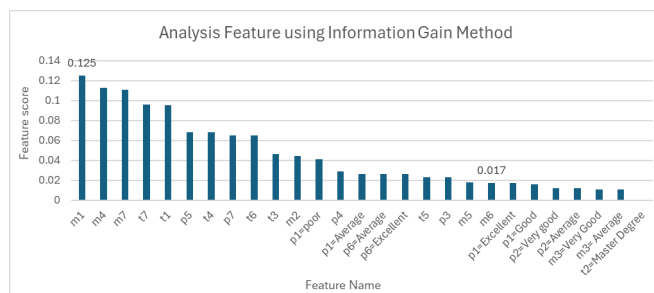


Figure 2 (continued)

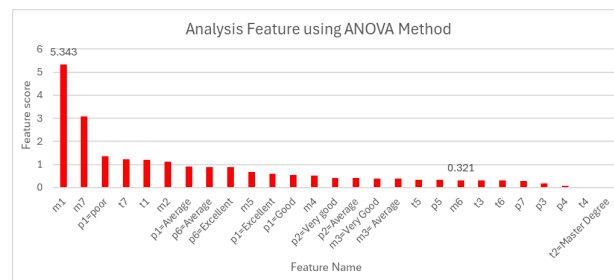


Figure 2. Comparison result of features analysis

Next is the data composition process, dividing the data composition into 70% training and 30% test data. The selection of data used in the composition data stage is random. The data distribution is shown in Table 4. After data composition stage, classification was carried out using three methods, namely k-NN, Decision Tree and Naïve Bayes. Each classification method is tested against the use of a feature selection approach. The results of testing the feature selection method on classification performance are shown in Table 5.

Table 4. Data Compositon for Classification

Total Data	Number of Data Training	Number of Data Testing
53	37	16

Table 5. Classification Comparison Result

Classification Model - Feature Selection	Evaluation			
	accuracy	precision	recall	f1-score
K-NN (Information Gain)	0.491	0.488	0.491	0.484
K-NN (ANOVA)	0.491	0.488	0.491	0.484
K-NN (Chi-Square)	0.434	0.422	0.434	0.418
Naïve Bayes (Information Gain)	0.774	0.775	0.774	0.771
Naïve Bayes (ANOVA)	0.774	0.775	0.774	0.771
Naïve Bayes (Chi-Square)	0.736	0.737	0.736	0.735
Decision Tree (Information Gain)	0.887	0.896	0.887	0.886
Decision Tree (ANOVA)	0.887	0.896	0.887	0.886
Decision Tree (Chi-Square)	0.887	0.903	0.887	0.884

From the classification results, the model found that the best classification method was shown as the Decision Tree classification method, with the use of the best feature selection method being Chi-Square with a classification accuracy of 0.887, precision of 0.903, recall of 0.887, and f1-score of 0.884. The second-best classification method is Naïve Bayes, which uses the Information Gain and ANOVA feature selection approach with classification accuracy of 0.774, precision of 0.775, recall of 0.774, and f1-score of 0.771. The k-NN method shows the lowest classification model performance using the Information Gain and ANOVA feature selection approach, which obtains a classification accuracy of 0.491, precision of 0.488, recall of 0.491 and f1-score of 0.484.

The comparison showed that the chi-square feature selection method had better results than the Chi-Square feature selection and worked more optimally in the Decision Tree classification model. This research conducted a comparative analysis between classification models using feature selection and those without feature selection. This comparative analysis is carried out to see how much the performance of the classification model has increased or decreased. The results of the comparative analysis between using and without using the feature selection method are shown in Table 6.

Table 6. Analysis of Feature Selection Effect in Classification Model

Classification – Feature Selection method	Evaluation	No Selection Feature	With Selection Feature	Increasing (%)	Decreasing (%)
Decision Tree (Chi-Square)	Accuracy	0.887	0.887	0.00	0.00
	Precision	0.896	0.903	0.78	0.00
	Recall	0.887	0.887	0.00	0.00
	F1-score	0.886	0.884	0.00	0.23

Table 6 (continued)

Classification – Feature Selection method	Evaluation	No Selection Feature	With Selection Feature	Increasing (%)	Decreasing (%)
Naïve Bayes (Information Gain and ANOVA)	Accuracy	0.774	0.774	0.00	0.00
	Precision	0.775	0.775	0.00	0.00
	Recall	0.774	0.774	0.00	0.00
	F1-score	0.771	0.771	0.00	0.00
K-NN (Information Gain and ANOVA)	Accuracy	0.453	0.491	8.39	0.00
	Precision	0.460	0.488	6.09	0.00
	Recall	0.453	0.491	8.39	0.00
	F1-score	0.460	0.484	5.22	0.00

In testing the Decision Tree model, it was found that the feature selection method was influenced by the performance increase in precision evaluation by 0.78%, from 0.896 to 0.903. The chi-square feature selection method does not influence increasing or decreasing the value of accuracy and recall. Meanwhile, in the f1-score evaluation value, a performance value decreased by 0.23% from 0.886 to 0.884. Analysis of the influence of the feature selection method on the Decision Tree model is shown in Figure 3, and the percentage effect of increasing or decreasing performance due to using Chi-Square feature selection is shown in Figure 4.

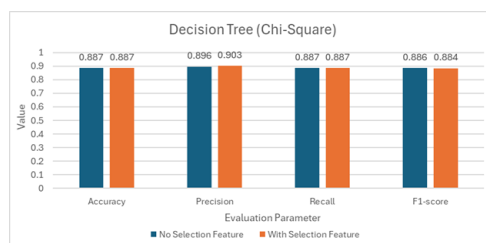


Figure 3. Comparison Result Between Selected and Unselected Feature Selection in Decision Tree Model

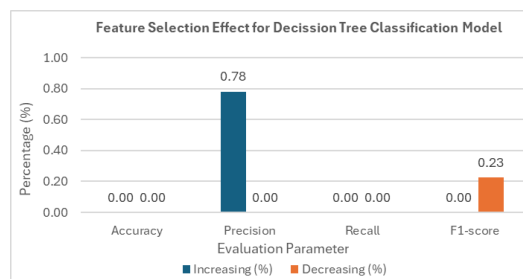


Figure 4. Chi-Square Feature Selection Effect in Decision Tree Model

In the Naïve Bayes classification model, no influence was found using the ANOVA or information gain selection methods on evaluation performance. The results of the influence feature selection method analysis on the Naïve Bayes method are shown in Figure 5, and the analysis of increasing influence is shown in Figure 6.

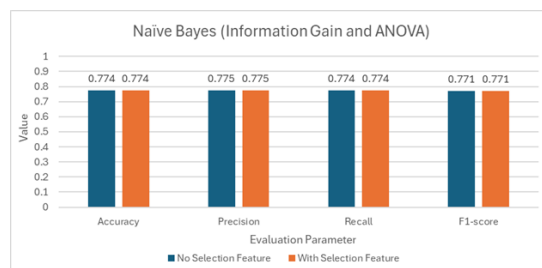


Figure 5. Comparison Result Between Selected and Unselected Feature Selection in Naïve Bayes

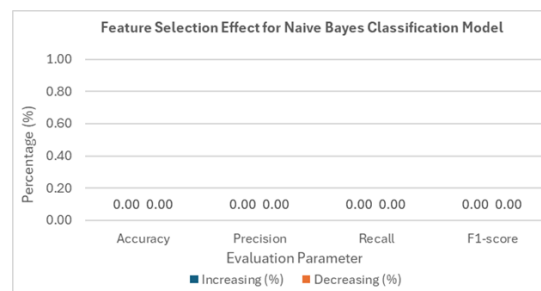


Figure 6. ANOVA and Information Gain Feature Selection Effect in Naïve Bayes Model

Meanwhile, in the k-NN model, using the Information Gain and ANOVA feature selection methods is influenced by model evaluation. Even though it obtained the lowest value, the feature selection method in the k-NN model can increase model performance in terms of accuracy by 8.39% from 0.453 to 0.491, precision by 6.09% from 0.46 to 0.488, recall by 8.39% from 0.453 to 0.491, and f1-score by 5.22% from 0.46 to 0.484. The results of the analysis of the influence of the feature selection method on the k-NN method are shown in Figure 7, and the analysis of the increase in influence is shown in Figure 8.

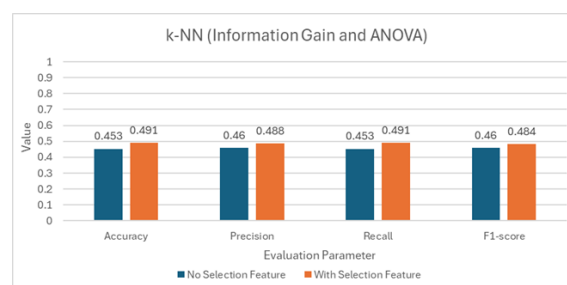


Figure 7. Comparison Result Between Selected an Unselected Feature Selection in k-NN

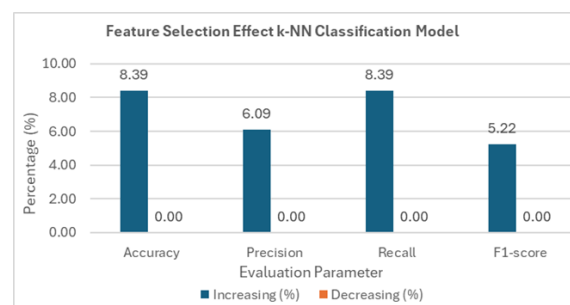


Figure 8. ANOVA and Information Gain Feature Selection Effect in k-NN Model

4. CONCLUSION

This research proposes an improving classification model with a feature selection approach to measure lecturer performance. The analysis approach was carried out on three feature selection approaches, ANOVA, Chi-Square and Information Gain, by implementing classification models such as k-NN, Naïve Bayes and Decision tree. Twenty-one basic features are extracted using the one-hot-encode technique, and 27 are obtained. The feature selection method reduces 75% of the number of features, and 20 different features are obtained by each feature selection method. The Information Gain feature selection technique obtained from the m_1 as the best feature with a score of 0.125, and the 20th feature is m_6 with a score of 0.017. The ANOVA method produces feature m_1 as the best feature with a score of 5.343, and the 20th feature is m_6 with a score of 0.321. The Chi-Square method obtained the m_1 as the first best feature with a score of 6.403, and the 20th feature is $m_3=verygood$ with a score of 0.412. The findings of this study showed the Decision Tree model is the best classification model using the Chi-Square feature selection method with a classification accuracy

of 0.887, precision of 0.903, recall of 0.887, and f1-score of 0.884. In addition, the Chi-square feature selection technique increased precision performance by 0.78%, from 0.896 to 0.903. However, on accuracy and recall, the chi-square feature selection method had no effect, and on the f1-score value, there was a decrease in performance value of 0.23% from 0.886 to 0.884. The classification model can measure lecturer performance by optimizing the feature selection approach. The feature selection method can reduce and increase the classification model's performance based on the data's characteristics used in testing the classification model.

In future research, we plan to analyze the number of ideal features that can be used in the classification model and compare our feature selection method with other variants. The goal is to find the optimal number of features and the best method for selecting features, which we believe will significantly enhance the classification model's performance in measuring the quality of lecturer performance. This potential for improvement should instill hope and optimism in our audience about the future of lecturer performance measurement.

5. ACKNOWLEDGEMENTS

Thank you to Institut Teknologi dan Bisnis STIKOM Bali for supporting this research to be completed and appropriately published and to improve our performance quality.

6. DECLARATIONS

AUTHOR CONTRIBUTION

Ni Luh Putri Srinadi was responsible for the conceptualization, methodology, supervision, writing of the original draft, review, and editing. Dandy Pramana Hostiadi contributed to the conceptualization and data curation. Luh Putu Wiwien Widhyastuti conducted formal analysis, investigation, and writing review and editing. I, Nyoman Suraja Antara, handled the software, validation, resources, visualization, and writing review and editing. Erma Sulistyo Rini was in charge of project administration, writing reviews, and editing.

FUNDING STATEMENT

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

COMPETING INTEREST

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

REFERENCES

- [1] V. V, "Comparison of Some Classification Algorithms for the Analysis of Students Academic Performance in Educational Data Mining Using Orange," *International Journal of Advanced Research in Science, Communication and Technology*, vol. 6, no. 1, pp. 318–324, 2021, <https://doi.org/10.48175/ijarsct-1394>.
- [2] I. Khan, A. R. Ahmad, N. Jabeur, and M. N. Mahdi, "An artificial intelligence approach to monitor student performance and devise preventive measures," *Smart Learning Environments*, vol. 8, no. 1, 2021, <https://doi.org/10.1186/s40561-021-00161-y>.
- [3] S. Sasmal, "Predictive Analytics in Data Engineering: An AI Approach," *International Research Journal of Engineering & Applied Sciences*, vol. 12, no. 1, pp. 13–18, 2024, <https://doi.org/10.55083/irjeas.2024.v12i01003>.
- [4] B. Mašić, M. Dželetović, and S. Nešić, "Big data analytics as a management tool: An overview, trends and challenges," *Anali Ekonomskog fakulteta u Subotici*, vol. 58, no. 48, pp. 101–118, 2022, <https://doi.org/10.5937/anebsub2248101m>.
- [5] R. Rosdiana and T. Husaen, "Analysis of Lecturer Characteristics on Lecturer Performance Through Learning Effectiveness Case Study of Nuku Tidore University," *Sosiohumaniora*, vol. 24, no. 3, p. 305, 2022, <https://doi.org/10.24198/sosiohumaniora.v24i3.36469>.
- [6] N. Nurmawati, A. Assagaf, and S. Sukiman, "the Influence of Leadership, Motivation, Organizational Culture, Lecturer Competence on Lecturer Performance With Lecturer Discipline As Intervening," *IJEED (International Journal of Entrepreneurship and Business Development)*, vol. 7, no. 1, pp. 198–208, 2024, <https://doi.org/10.29138/ijebed.v7i1.2640>.

- [7] Kasmiaty, Baharuddin, M. N. Fattah, H. Nasaruddin, Y. Yusriadi, M. I. Usman, and Suherman, "The influence of leadership and work motivation on work effectiveness through discipline," *Proceedings of the International Conference on Industrial Engineering and Operations Management*, vol. 565, no. INCoEPP, pp. 3648–3655, 2021, <https://doi.org/10.46254/SA02.20210847>.
- [8] E. KARAAHMETOĞLU, S. ERSÖZ, A. K. TÜRKER, V. ATEŞ, and A. F. İNAL, "Evaluation of Profession Predictions for Today and the Future with Machine Learning Methods : Emperical Evidence From Turkey," *Politeknik Dergisi*, vol. 26, no. 1, pp. 107–124, 2021, <https://doi.org/10.2339/politeknik.985534>.
- [9] S. Kim and E. Usui, "Employer learning, job changes, and wage dynamics," *Economic Inquiry*, vol. 59, no. 3, pp. 1286–1307, 2021, <https://doi.org/10.1111/ecin.12980>.
- [10] R. Restu and S. Sriadhi, "Lecturer Performance and its Determining Factors in a Blended Learning System During The COVID-19 Pandemic," *TEM Journal*, vol. 11, no. 4, pp. 1680–1686, 2022, <https://doi.org/10.18421/TEM114-32>.
- [11] I. R. Suwarma and S. Apriyani, "Explore Teachers' Skills in Developing Lesson Plan and Assessment That Oriented on Higher Order Thinking Skills (HOTS)," *Journal of Innovation in Educational and Cultural Research*, vol. 3, no. 2, pp. 106–113, 2022, <https://doi.org/10.46843/jiecr.v3i2.66>.
- [12] M. Hasim, N. F. Umar, A. Amiruddin, and A. Sinring, "Career Commitment Based on Career Identity Diffusion among Students in Vocational Higher Education," *Journal of Innovation in Educational and Cultural Research*, vol. 4, no. 2, pp. 220–228, 2023, <https://doi.org/10.46843/jiecr.v4i2.536>.
- [13] A. Tahniah, H. Fitria, and A. Wahidy, "The influence of organization culture on teacher performance of elementary school," *JPGI (Jurnal Penelitian Guru Indonesia)*, vol. 6, no. 2, p. 460, 2021, <https://doi.org/10.29210/021071jpgi0005>.
- [14] N. Luh, P. Srinadi, L. Putu, W. Widhyastuti, D. P. Hostiadi, and C. Ahmadi, "Lecturer Performance Measurement Based on Organizational Culture and Leadership Behavior Analysis Using Pearson Correlation," vol. 5, no. 1, pp. 131–139, 2024, <https://doi.org/10.46843/jiecr.v5i1.1163>.
- [15] I. K. Liem, A. E. Fatril, and F. A. Husna, "Satisfaction of lecturers and undergraduate students of medical faculties in Indonesia towards online anatomy learning during COVID-19 pandemic," *BMC Medical Education*, vol. 24, no. 1, pp. 1–10, 2024, <https://doi.org/10.1186/s12909-024-05620-x>.
- [16] A. Helmina, Dedy Irfan, Fahmi Rizal, and Kasmita, "Development of Teaching Performance Evaluation Application for Lecturers Using K-Nearest Neighbor Method with Manhattan Distance Approach," *JTP - Jurnal Teknologi Pendidikan*, vol. 26, no. 1, pp. 278–290, 2024, <https://doi.org/10.21009/jtp.v26i1.44443>.
- [17] Mahpud, S. Setyaningsih, and O. Sunardi, "Improving Lecturer Performance Through Strengthening Service Leadership, Empowerment, Achievement Motivation, And Trust in Muhammadiyah Tangerang University Lecturers," *International Journal of Business and Social Science Research*, vol. 5, no. 6, pp. 226–232, 2024, <https://doi.org/10.47742/ijbssr.v5n6p2>.
- [18] J. Hazzam and S. Wilkins, "The influences of lecturer charismatic leadership and technology use on student online engagement, learning performance, and satisfaction," *Computers and Education*, vol. 200, no. July, p. 104809, 2023, <https://doi.org/10.1016/j.compedu.2023.104809>.
- [19] H. Hermansyah, A. Moeins, and Y. Zain, "The Influence of Leadership, Learning Organization, Compensation, and Work Motivation on Lecturer Performance in the Legal Study Program at L13dikti Jakarta," *International Journal of Social Science and Human Research*, vol. 7, no. 04, pp. 2425–2435, 2024, <https://doi.org/10.47191/ijsshr/v7-i04-25>.
- [20] J. Turkiewicz, A. Bęczkowska, and R. Skroback, "The Importance of Building a Relationship Between Lecturers and Students for Student Satisfaction with Remote Learning," *Przegląd Badań Edukacyjnych Przegląd Badań Edukacyjnych Educational Studies Review*, vol. 1, no. 41, pp. 121–138, 2023, <https://doi.org/10.12775/PBE.2023.010>.
- [21] S. Faris, "Pengaruh kompetensi, pelatihan dan motivasi terhadap kinerja dosen tetap pada universitas prima indonesia 1," *Agriprimatech*, vol. 4, no. 1, pp. 16–24, 2020, <https://doi.org/10.34012/agriprimatech.v4i1.1317>.
- [22] M. Milkhatun, A. A. F. Rizal, N. W. W. Asthiningsih, and A. J. Latipah, "Performance Assessment of University Lecturers: A Data Mining Approach," *Khazanah Informatika : Jurnal Ilmu Komputer dan Informatika*, vol. 6, no. 2, pp. 73–81, 2020, <https://doi.org/10.23917/khif.v6i2.9069>.

[This page intentionally left blank.]