

Original Research

Random Forest-Based Performance Prediction for Teachers and Staff at Yayasan Al Istiqomah Al-Islamiyah Prabumulih

Hilmy Putra Dwi Anwarsyah ^{a*}, Rendra Gustriansyah ^b, Evi Purnamasari ^c

^{a,b,c} Department of Informatics Engineering, Faculty of Computer Science and Science, Universitas Indo Global Mandiri, Palembang, Indonesia

*Corresponding Author: 12021110066@students.uigm.ac.id

ABSTRACT

Teacher and staff performance evaluation is important for supporting managerial decision-making and maintaining the quality of educational services in educational institutions. However, the manual evaluation process at the Al-Istiqomah Al-Islamiyah Foundation in Prabumulih is time-consuming and less effective in supporting decisions regarding the appointment of Permanent Foundation Teachers (GTY) and Permanent Foundation Employees (PTY). This study aims to develop and evaluate a Random Forest-based model for predicting teacher and staff performance. The study used a quantitative approach based on 100 historical performance evaluation records collected from January 2022 to January 2024, with ten assessment variables. The research stages included data preprocessing, training-testing data splitting using five ratios (50:50, 60:40, 70:30, 80:20, and 90:10), Random Forest modeling, and evaluation using accuracy, precision, recall, F1-score, confusion matrix, feature importance, ROC-AUC, and k-fold cross-validation. The 70:30 split produced 96% training accuracy and 97% testing accuracy. The testing confusion matrix correctly classified 24 of 25 GTY/PTY-appointed cases and all 5 not-appointed cases. The model also achieved an AUC of 0.992, indicating strong classification capability. These results show that the Random Forest model can provide a systematic, consistent, and objective approach to supporting teacher and staff performance evaluation and appointment decisions.

Copyright © 2026 by Author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution 4.0 International License (CC BY 4.0).

How to Cite:

H. P. D. Anwarsyah, R. Gustriansyah, and E. Purnamasari, "Random Forest-Based Performance Prediction for Teachers and Staff at Yayasan Al Istiqomah Al-Islamiyah Prabumulih," *J. Kecerdasan Buatan dan Teknol. Inf.*, vol. 5, no. 3, pp. 717–726, 2026.
Doi: <https://doi.org/10.69916/jkbt.v5i3.593>

KEYWORDS

Machine Learning
Performance Evaluation
Performance Prediction
Random Forest
Teachers and Employees

ARTICLE INFO

Received: July 20, 2026

Accepted: August 28, 2026

Available Online: September 18, 2026

1 INTRODUCTION

Human resources (HR) are a strategic component of organizational performance, including in educational institutions where teacher and staff quality directly affects the effectiveness of educational services [1,2]. Human resource analytics supports the use of employee-related data to improve evidence-based decision-making and connect HR information with organizational outcomes [3].

In educational institutions, human resource practices such as performance appraisal, training, and development are closely associated with organizational performance [4,5]. Empirical evidence from educational institutes indicates that performance appraisal is an important HR practice related to organizational performance, reinforcing the need for systematic and objective evaluation mechanisms [6].

At the Al-Istiqomah Al-Islamiyah Foundation, teacher and staff performance evaluation is still conducted manually by school principals and evaluation teams [7,8]. This process can require substantial time and may lead to

differences among evaluators, while decisions concerning the appointment of Permanent Foundation Teachers (GTY) and Permanent Foundation Employees (PTY) may be delayed [9,10]. A data-driven approach can help transform historical evaluation records into information that supports more consistent managerial decisions [11]. Recent research demonstrates that machine learning can be applied to employee performance prediction using demographic, job-related, and historical performance variables [12,13]. In particular, Random Forest has been evaluated alongside other machine learning algorithms for employee performance classification and has shown strong predictive performance while supporting data-driven HR decision-making [14,15]. The Random Forest algorithm itself is an ensemble method that combines multiple decision trees and determines predictions through aggregation, making it suitable for tabular classification problems [16].

Machine learning has also been widely applied in educational data mining [17]. Educational data mining provides computational methods for analyzing educational data and identifying patterns that can support educational decision-making [18,19]. Studies have applied Random Forest to teacher teaching-ability evaluation and to student performance prediction, demonstrating the relevance of Random Forest-based approaches to educational evaluation and prediction tasks [20].

2 RESEARCH METHODS

2.1 Research Design

This study was conducted through a series of systematic research stages. Each stage serves as a framework that describes the entire research process from the initial phase to the final stage. The overall research workflow used in this study is illustrated in Figure 1.

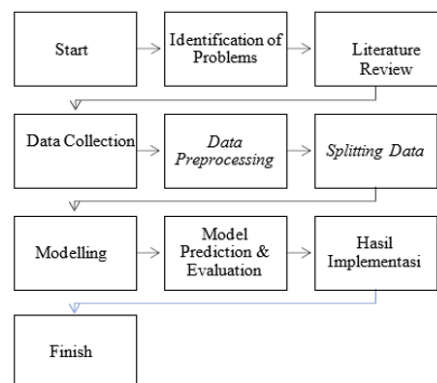


Figure 1. Research Stages

2.2 Problem Identification

The educational operations at the Al-Istiqomah Al-Islamiyah Foundation, Prabumulih, face challenges in collecting, processing, and utilizing teacher and staff performance evaluation data effectively. This issue arises from the absence of a data-driven system capable of accurately predicting employee performance. Consequently, performance monitoring and decision-making processes, particularly regarding the appointment of permanent teachers and staff, are often delayed and rely on subjective evaluations.

To improve human resource quality and ensure compliance with the institution's Standard Operating Procedures (SOPs), a systematic and accurate approach to processing historical performance evaluation data is required. Historical evaluation records provide valuable information that can be utilized to develop a performance prediction model using machine learning techniques.

However, implementing a Random Forest-based prediction model requires careful evaluation of factors that influence model performance, particularly the data-splitting process. For relatively small datasets, different training and testing data split ratios may affect prediction accuracy and the model's generalization capability. Therefore, the primary problem addressed in this study is to analyze how different data-splitting ratios influence the performance of the Random Forest algorithm in developing an accurate, objective, and reliable prediction model for teacher and staff performance based on historical evaluation data.

2.3 Literature Review

This study reviewed peer-reviewed literature on human resource analytics, educational data mining, employee performance prediction, teacher evaluation, and Random Forest classification. The literature was selected based on its direct relevance to the research problem and methodological approach.

HR analytics provides a foundation for using employee-related data to support evidence-based HR decisions. In educational settings, performance appraisal is an important HR practice associated with organizational performance. Recent machine learning research has also demonstrated the feasibility of predicting employee performance from HR-related variables, including the use of Random Forest classifiers.

Random Forest was introduced as an ensemble learning method that constructs multiple decision trees using randomized samples and feature selection, with the final prediction obtained through aggregation. This characteristic makes it appropriate for classification problems involving multiple performance indicators.

2.4 Data Collection

This study employed the documentation method to collect secondary data from historical teacher and staff performance evaluation records at the Al-Istiqomah Al-Islamiyah Foundation, Prabumulih. The records covered January 2022 to January 2024 and contained ten performance-assessment indicators for each employee.

A total of 100 records were available for analysis. This small dataset is an important methodological limitation because performance metrics obtained from small testing sets can have high sampling variability. Therefore, the reported results are interpreted as an exploratory evaluation of the available institutional data rather than as a population-level estimate of model performance.

The target variable is the employee's appointment status. The two target classes are (1) GTY/PTY Appointed, representing employees recorded as appointed as Permanent Foundation Teachers (GTY) or Permanent Foundation Employees (PTY), and (2) Not Appointed, representing employees without a GTY/PTY appointment outcome in the historical records. The dataset contains 82 Appointed records and 18 Not Appointed records, corresponding to an 82:18 class distribution.

Table 1. Teacher and Staff Data Sample

Employee	V1	V2	V3	V4	V5	V6	V7	V8	V9	V10
P1	45	85	70	65	95	65	75	50	50	70
P2	24	37	35	45	45	50	38	45	30	45
P3	70	85	75	60	75	55	80	80	50	85
P4	75	70	70	75	75	85	78	78	80	85
P5	70	75	70	82	85	85	70	70	70	70
P6	25	25	25	20	25	30	25	25	20	25
P7	85	85	73	95	95	75	75	74	85	70
P8	85	75	70	85	80	85	85	90	90	85
P9	80	75	75	89	70	70	75	70	80	75
...	...	...	...	...	...	...	...	...	...	...
P100	20	25	25	25	20	22	20	20	22	22

The predictor variables are retained in the source dataset as V1V10. The available manuscript documentation identifies performance dimensions including discipline, communication skills, work efficiency, and teamwork, but it does not provide the exact mapping between these dimensions and V1V10. To avoid inventing information, the exact V-number-to-indicator mapping must be verified against the original institutional evaluation form before publication.

2.5 Data Preprocessing

Data preprocessing was performed to ensure that the dataset was clean, consistent, and suitable for Random Forest classification. The process included checking missing, incomplete, and invalid records and ensuring that predictor values were represented in a consistent numeric format.

The target distribution is imbalanced (82 Appointed and 18 Not Appointed). The original experiment did not apply SMOTE, ADASYN, or class weighting. This is therefore reported explicitly as a limitation, and accuracy is not interpreted alone; class-specific precision, recall, F1-score, and the confusion matrix are considered together. A class-weighted Random Forest or a resampling strategy should be evaluated in a future rerun before corrected performance values are claimed.

2.6 Data Splitting

After preprocessing, the dataset was divided into training and testing sets. The training set was used to build the Random Forest model, while the testing set was used to evaluate predictive performance on unseen records. Five split ratios were examined: 50:50, 60:40, 70:30, 80:20, and 90:10. With 100 total records, these scenarios produce test sets of 50, 40, 30, 20, and 10 observations, respectively. The 90:10 scenario is therefore particularly sensitive to sampling variation and is treated as an exploratory comparison rather than a statistically stable estimate of generalization performance.

The data were split while maintaining class proportions as closely as possible through stratification. K-fold

cross-validation was also used as a supplementary assessment of model stability. Nevertheless, cross-validation cannot fully overcome the fundamental limitation imposed by the small number of observations.

2.7 Modeling

After preprocessing and data splitting, the Random Forest algorithm was used to classify the target variable into GTY/PTY Appointed and Not Appointed. The model was trained separately for each data-splitting scenario. Random Forest constructs multiple decision trees from randomized subsets of observations and predictors, and the final class is determined through aggregation. The original experiment used 100 trees, a maximum depth of 2, and `random_state = 42`. These settings were kept consistent across split scenarios to make the comparison reproducible.

2.8 Model Evaluation

The performance of the Random Forest models was evaluated using accuracy, precision, recall, and F1-score for each data-splitting ratio. Because the target classes are imbalanced, class-specific evaluation metrics and the confusion matrix were emphasized in addition to overall accuracy. The evaluation metrics were calculated based on the number of true positives (TP), true negatives (TN), false positives (FP), and false negatives (FN). Accuracy measures the proportion of correctly classified instances among all observations and is calculated as:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

Precision measures the proportion of correctly predicted positive instances among all instances predicted as positive:

$$\text{Precision} = \frac{TP}{TP + FP} \quad (2)$$

Recall, also referred to as sensitivity or true positive rate, measures the proportion of actual positive instances that are correctly identified by the model:

$$\text{Recall} = \frac{TP}{TP + FN} \quad (3)$$

To provide a balance between precision and recall, the F1-score was calculated as the harmonic mean of precision and recall:

$$\text{F1-score} = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (4)$$

These class-specific metrics are particularly important in the presence of imbalanced target classes because accuracy alone may provide an overly optimistic assessment of model performance.

The confusion matrix was also examined to identify the distribution of correct and incorrect predictions across the target classes. It provides the values of TP , TN , FP , and FN , allowing the classification errors of the Random Forest model to be analyzed in greater detail.

In addition, Receiver Operating Characteristic (ROC) analysis and the Area Under the Curve (AUC) were used as supplementary measures of the model's discrimination capability. The ROC curve illustrates the relationship between the true positive rate (TPR) and false positive rate (FPR) across different classification thresholds. The true positive rate and false positive rate are defined as:

$$TPR = \frac{TP}{TP + FN} \quad (5)$$

$$FPR = \frac{FP}{FP + TN} \quad (6)$$

The AUC represents the area under the ROC curve and summarizes the model's ability to distinguish between the target classes across classification thresholds. An AUC value closer to 1 indicates better discrimination, whereas a value around 0.5 indicates discrimination that is close to random classification.

To assess the stability of the Random Forest models across different data partitions, k -fold cross-validation was also employed. In k -fold cross-validation, the dataset is divided into k approximately equal subsets, with each subset used once as the validation set while the remaining $(k - 1)$ subsets are used for model training. The

cross-validation score can be expressed as the mean performance across all folds:

$$CV_{\text{mean}} = \frac{1}{k} \sum_{i=1}^k M_i \quad (7)$$

where M_i represents the evaluation metric obtained from the i -th fold and k denotes the number of folds. The variability across folds may additionally be assessed using the standard deviation:

$$CV_{\text{std}} = \sqrt{\frac{1}{k-1} \sum_{i=1}^k (M_i - CV_{\text{mean}})^2} \quad (8)$$

Given the small sample size, the evaluation metrics were interpreted cautiously because both hold-out and cross-validation estimates may vary substantially depending on the sampled records and data partitions.

Furthermore, hyperparameter tuning and feature-importance analysis were performed to examine the behavior and contribution of the predictors in the Random Forest models. Feature importance was used to identify the variables that contributed most strongly to the model's decision-making process.

The experimental results reported in this manuscript reflect the original model configuration. No additional results involving unreported resampling or class-balancing techniques, such as SMOTE, ADASYN, or class-weighting, were introduced.

3 RESULTS AND DISCUSSION

3.1 Model Performance Evaluation

a. Confusion Matrix Results

This study employed five data-splitting ratios (50:50, 60:40, 70:30, 80:20, and 90:10) to evaluate the performance of the Random Forest model. Among these, the 70:30 ratio was selected as the primary experimental configuration because it provides a balanced proportion of training and testing data, enabling a more reliable assessment of the model's generalization capability.

Across the five data-splitting scenarios, the 70:30 configuration provided the most balanced testing performance. It achieved 97% testing accuracy and maintained strong precision, recall, and F1-score values for both classes. The comparison in Table 2 therefore supports 70:30 as the primary configuration for interpreting the model's predictive performance.

Table 2. Confusion Matrix Results across Data-Splitting Scenarios

	50:50		60:40		70:30		80:20		90:10	
	Train	Test	Train	Test	Train	Test	Train	Test	Train	Test
Accuracy	94%	98%	95%	97%	96%	97%	100%	95%	96%	100%
Precision	100%	& 100%	& 100%	& 100%	& 100%	& 100%	& 100%	& 100%	& 100%	& 100% (1)
	79%	& 88%	& 80%	& 86%	& 81%	& 83%	& 100%	& 80%	& 82%	
Recall	92%	& 98%	& 94%	& 97%	& 95%	& 96%	& 94%	& 100%	& 94%	& 100% (1)
	100%	& 100%	& 100%	& 100%	& 100%	& 100%	& 100%	& 100%	& 100%	
F1-Score	96%	& 99%	& 97%	& 99%	& 97%	& 98%	& 97%	& 100%	& 97%	& 100% (1)
	88%	& 93%	& 89%	& 92%	& 90%	& 91%	& 89%	& 100%	& 90%	

b. Training Data Performance

Based on the evaluation of the Random Forest model on the training dataset, the model achieved an accuracy of 96%, indicating excellent classification performance. The confusion matrix shows that 54 of 57 teachers and staff appointed as GTY/PTY were correctly classified, while 3 instances were misclassified. In addition, all 13 non-appointed cases were correctly classified, demonstrating the model's strong ability to identify this class.

The classification report indicates that the GTY/PTY Appointed class achieved a precision of 1.00, recall of 0.95, and an F1-score of 0.97. For the Not Appointed class, the model obtained a precision of 0.81, recall of 1.00, and an F1-score of 0.90. Overall, the weighted average F1-score was 0.96, while the macro average F1-score reached 0.93, indicating stable and balanced performance despite class imbalance.

These results demonstrate that the Random Forest model effectively learned the underlying patterns in the training data and was well prepared for evaluation on the testing dataset.

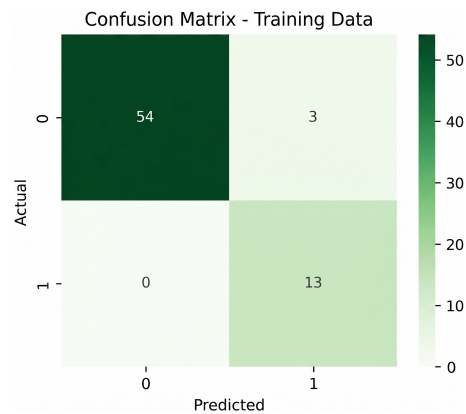


Figure 2. Training Data Confusion Matrix

c. Testing Data Performance

Based on the evaluation of the Random Forest model on the testing dataset, the model achieved an accuracy of 97%, demonstrating excellent generalization performance. The confusion matrix indicates that 24 of 25 teachers and staff appointed as GTY/PTY were correctly classified, with only one misclassification. In addition, all five non-appointed cases were correctly predicted, confirming the model’s strong ability to identify both classes on unseen data.

According to the classification report, the GTY/PTY Appointed class achieved a precision of 1.00, recall of 0.96, and an F1-score of 0.98. The Not Appointed class obtained a precision of 0.83, recall of 1.00, and an F1-score of 0.91. Overall, the model achieved a weighted average F1-score of 0.97, indicating accurate, consistent, and reliable classification performance on previously unseen data.



Figure 3. Testing Data Confusion Matrix

3.2 Feature Importance and Interpretation

The feature-importance results show that the performance indicators do not contribute equally to the prediction of GTY/PTY appointment eligibility. This is important for the foundation because it indicates which aspects of the existing evaluation process are most closely associated with the historical appointment outcomes. Discipline was the strongest predictor in the model, corresponding to the first indicator listed in the dataset description (V1). Work Efficiency was also among the strongest predictors and corresponds to V3 based on the order in which the indicators are introduced in the manuscript. A second high-ranking indicator (V5) also contributed substantially, but its substantive label is not specified in the current dataset description; therefore, its organizational meaning should be documented in the final data dictionary before the result is interpreted. The remaining indicators, particularly those ranked lower in the feature-importance plot, contributed less to the model’s predictions. These differences suggest that the foundation should pay particular attention to discipline and work efficiency when reviewing performance evidence related to GTY/PTY appointment decisions, while avoiding the assumption that lower-importance indicators are unimportant for employee development. Overall, the feature importance analysis demonstrates that the Random Forest model successfully identified the most influential performance indicators. These findings provide valuable insights for improving performance evaluation criteria and support future development of data-driven decision-making models.

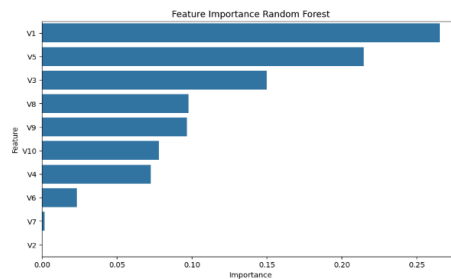


Figure 4. Feature Importance Ranking

3.3 ROC Curve and AUC

Based on the ROC curve, the Random Forest model performed well above the diagonal reference line, indicating excellent discrimination between the GTY/PTY Appointed and Not Appointed classes. The model achieved an Area Under the Curve (AUC) of 0.992, demonstrating outstanding classification performance.

An AUC of 0.992 indicates a 99.2% probability that the model assigns a higher prediction score to an employee who is actually eligible for appointment than to one who is not. The ROC curve's close proximity to the upper-left corner further confirms a high True Positive Rate and a low False Positive Rate, indicating accurate predictions with minimal misclassification.

The ROCAUC results support the model's strong discrimination between the GTY/PTY Appointed and Not Appointed classes. The AUC of 0.992 indicates excellent separation in this dataset. Nevertheless, a high AUC does not by itself establish that the model is appropriate as an autonomous personnel-decision mechanism; operational use should remain subject to human review, SOP compliance, and periodic validation on new data.

3.4 Discussion and Practical Implications

The findings should be interpreted not only as evidence of predictive accuracy but also as evidence about how performance information is currently associated with GTY/PTY appointment outcomes at the Al-Istiqomah Al-Islamiyah Foundation. The model achieved 97% testing accuracy and an AUC of 0.992, indicating strong discrimination in the available dataset. These results suggest that historical evaluation records contain sufficiently consistent patterns to support a data-driven decision-support approach. However, the model should support, rather than replace, the professional judgment of school principals and the evaluation team.

The feature-importance results provide a more actionable interpretation for educational human resource management. Discipline and Work Efficiency emerged as the most influential named indicators in the available description of the dataset. In the foundation's context, this suggests that consistent attendance, adherence to work responsibilities, and efficient completion of assigned duties may be closely associated with historical appointment outcomes. The result can therefore inform more focused performance monitoring and feedback. At the same time, feature importance should not be interpreted as proof that these indicators cause appointment decisions; it only shows their contribution to the model's predictions.

For SOP improvement, the foundation can use the model as an early screening and evidence-summarization tool. Before an appointment recommendation is finalized, the predicted class should be reviewed together with the employee's underlying evaluation records, documented achievements, disciplinary history, and applicable institutional criteria. A practical workflow would be: the model generates a prediction, the HR or evaluation team checks the supporting indicators, discrepancies are documented, and the final appointment decision remains with authorized decision-makers. This approach reduces repetitive manual screening while preserving accountability.

Ethical and governance safeguards are also necessary before deployment. Employee data should be access-controlled and used only for legitimate HR purposes. Predictions should be explainable to decision-makers, and employees should not be disadvantaged solely because of an automated prediction. The model should be periodically re-evaluated for changes in employee characteristics, evaluation practices, and class distribution. Because the dataset contains only 100 historical records, the reported performance should be treated as promising rather than definitive, and the model should be validated on larger and more recent data before being used for high-stakes appointment decisions.

Overall, the main contribution of the model is not simply its accuracy but its potential to make the foundation's evaluation process more consistent, evidence-based, and auditable. The strongest indicators can guide attention, while the prediction itself can function as one source of evidence within a transparent human-led HR process.

4 CONCLUSION

This study successfully developed a Random Forest-based prediction model for evaluating teacher and staff performance using historical performance evaluation data from the Al-Istiqomah Al-Islamiyah Foundation, Prabumulih. The proposed model achieved 96% training accuracy and 97% testing accuracy, demonstrating strong predictive performance in classifying GTY/PTY appointment eligibility. Among the five evaluated data-splitting ratios (50:50, 60:40, 70:30, 80:20, and 90:10), the 70:30 ratio produced the most balanced overall performance, with high accuracy, precision, recall, and F1-score, as well as a low number of misclassifications. The model also achieved an AUC of 0.992, indicating strong discriminatory capability between the GTY/PTY Appointed and Not Appointed classes.

From a practical perspective, the findings indicate that Random Forest can support a more objective, consistent, and data-driven performance evaluation process at the foundation. The model can be used as a decision-support tool to assist HR managers and school administrators in identifying performance patterns relevant to GTY/PTY appointment decisions. However, the model should not be used as the sole basis for employee appointment decisions. Its predictions should be considered together with institutional SOPs, professional judgment, and other relevant administrative and performance considerations.

A key limitation of this study is the relatively small dataset, which consists of only 100 performance evaluation records collected from January 2022 to January 2024. Although the model achieved high predictive performance on the available data, the limited sample size may constrain the generalizability of the findings to other periods, employee groups, or educational institutions. Therefore, the current performance results should be interpreted within the context of the available dataset and should not yet be considered universally applicable.

Future research should validate the proposed model using a substantially larger dataset collected over multiple years and, where possible, from different units or educational institutions. Such validation would provide a stronger assessment of the model's stability, robustness, and generalization capability across different employee characteristics and evaluation periods. Future studies may also incorporate additional relevant performance indicators and compare Random Forest with other machine learning algorithms to determine whether predictive performance and practical usefulness can be further improved. Furthermore, integrating the validated model into a decision-support application could facilitate its responsible implementation within the foundation's human resource management process.

ACKNOWLEDGMENTS

The authors would like to express their gratitude to Universitas Indo Global Mandiri and Yayasan Al-Istiqomah Al-Islamiyah Prabumulih for providing administrative support, facility access, and historical evaluation data that contributed to the completion of this research.

DECLARATIONS

Author Contributions:

Hilmy Putra Dwi Anwarsyah: Conceptualization, Methodology, Software, Writing – Original Draft.

Rendra Gustriansyah: Formal Analysis, Resources, Supervision.

Evi Purnamasari: Validation, Supervision, Writing – Review & Editing.

Funding:

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

Conflicts of Interest:

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

Data Availability:

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

AI Use Statement:

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

Additional Information:

No additional information is available for this paper.

REFERENCES

- [1] L. Breiman, "Random Forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [2] M. R. Adrian, M. P. Putra, M. H. Rafialdy, and N. A. Rakhmawati, "Perbandingan Metode Klasifikasi Random Forest dan SVM Pada Analisis Sentimen PSBB," *Jurnal Sains dan Informatika*, vol. 7, no. 1, pp. 36–40, 2021.
- [3] O. Pahlevi and Y. Handrianto, "Implementasi Algoritma Klasifikasi Random Forest Untuk Penilaian Kelayakan Kredit," *Jurnal Informatika dan Komputasi*, vol. 5, no. 1, pp. 71–76, 2023.
- [4] D. Jurinaldo and S. Arlis, "Teacher Performance Evaluation Analysis Using K-Means Clustering Algorithm and Random Forest Classification," *Sebatik*, vol. 29, no. 1, pp. 215–223, 2026, doi: 10.46984/sebatik.v29i1.2744.
- [5] G. M. Agung, R. A. Zuama, and E. S. Budi, "Analysis of Student Academic Performance Using Random Forest and Support Vector Machines," *Journal of Data Science and Analytics*, vol. 6, no. 1, pp. 57–65, 2026.
- [6] A. Rajamanickam and C. Kamalakannan, "Land Use and Land Cover Prediction in Tamilnadu of India, Using Random Forest Machine Learning Technique," *Journal of Applied Remote Sensing*, vol. 20, 2025.
- [7] N. Febriana, F. Fadzira, and M. A. Senubekti, "Analisis Klasifikasi Resign Karyawan dengan Random Forest," *Jurnal Teknologi Informasi*, vol. 18, no. 1, pp. 2580–2582, 2025.
- [8] E. A. Apriadi, R. Julianto, and M. Bisri, "Prediksi Kelayakan Pendidikan Sekolah Dasar di Indonesia Menggunakan Metode Random Forest Berbasis Python Tahun 2023–2024," *Jurnal Teknologi dan Informasi*, vol. 12, no. 1, pp. 45–52, 2026.
- [9] E. Y. Boateng, J. Otoo, and D. A. Abaye, "Basic Tenets of Classification Algorithms K-Nearest-Neighbor, Support Vector Machine, Random Forest and Neural Network: A Review," *Journal of Data Analysis and Information Processing*, vol. 8, no. 4, pp. 341–357, 2020, doi: 10.4236/jdaip.2020.84020.
- [10] S. Amaliah and M. Nusrang, "Penerapan Metode Random Forest Untuk Klasifikasi Varian Minuman Kopi Di Kedai Kopi Konijiwa Bantaeng," *Variansi*, vol. 4, no. 2, pp. 121–127, 2022, doi: 10.35580/variasiunm31.
- [11] A. Boukerche, L. Zheng, and O. Alfandi, "Outlier Detection: Methods, Models, and Classification," *ACM Computing Surveys*, vol. 53, no. 3, pp. 1–35, 2020.
- [12] A. Nugroho, "Analisa Splitting Criteria Pada Decision Tree dan Random Forest untuk Klasifikasi Evaluasi Kendaraan," *Jurnal Computer Science and Information Technology*, vol. 1, no. 1, pp. 41–49, 2022.
- [13] R. A. Nurhafiz, F. D. Marleny, and A. A. Ningrum, "Optimasi Algoritma Random Forest dalam Mengukur Kepuasan Peserta Pelatihan Guru pada Lembaga HAFECS," *Jurnal Media Informatika (JUMIN)*, vol. 7, no. 1, pp. 302–311, 2026.
- [14] T. Jurnal, P. Teknologi, N. Febriana, F. Fadzira, and M. A. Senubekti, "Prediksi Penilaian Kinerja Hakim Dengan Penerapan Machine Learning Menggunakan Tools Python," *Jurnal Teknologi Informasi*, vol. 2, no. 1, pp. 44–51, 2024.
- [15] R. Leonardo and J. Pratama, "Perbandingan Metode Random Forest Dan Naïve Bayes Dalam Prediksi Keberhasilan Klien Telemarketing," *Jurnal Sistem Informasi*, vol. 3, pp. 455–459, 2020.
- [16] C. E. Murwaningtyas, A. Kristiamita, A. Lintang, and A. Ika, "Analisis Pengaruh Media Sosial terhadap Produktivitas Akademik Mahasiswa menggunakan Metode Decision Tree dan Random Forest," *Jurnal Sains Data*, vol. 6, no. 2, pp. 499–509, 2024.
- [17] D. Pandey, K. Niwaria, and B. Chourasia, "Machine Learning Algorithms: A Review," *International Journal of Engineering and Advanced Technology*, vol. 8, no. 6, pp. 916–922, 2019.
- [18] S. Kurniawan and A. Nugroho, "Analisis Faktor yang Mempengaruhi Promosi Karyawan Menggunakan Random Forest pada Dataset Employee Promotion," *Jurnal Komputasi*, vol. 5, no. 2, pp. 177–187, 2025.

- [19] B. Said, A. Rashid, O. Asem, A. Azmi, and A. Abdullah, "Machine learning-based monitoring of renewable energy systems using random forest and LSTM," *E3S Web of Conferences*, vol. 500, p. 02004, 2026.
- [20] H. P. D. Anwarsyah, R. Gustriansyah, and E. Purnamasari, "Random Forest-Based Performance Prediction for Teachers and Staff at Yayasan Al Istiqomah Al-Islamiyah Prabumulih," *J. Kecerdasan Buatan dan Teknol. Inf.*, vol. 5, no. 3, pp. 709–723, 2026.