



Volume 20, Issue 1, 2026

Jurnal Ilmiah Teknologi Informasi Asia

Journal Homepage: <https://jurnal.asia.ac.id/index.php/jitika>



Article

Implementation of Feature Selection to Improve the Accuracy of Gender Classification Based on Voice Data with Random Forest

Suhardiyanto *, Fitroh Amaluddin, Aris Wijayanti

Universitas PGRI Ronggolawe, Tuban Regency, 62391, Indonesia

Abstract—Voice-based gender recognition has gained increasing importance in biometrics, security, forensics, and human-computer interaction. While humans can easily distinguish male and female voices, automatic classification remains challenging due to variability and high-dimensional acoustic data. This study investigates the role of feature selection in enhancing the performance and efficiency of Random Forest for gender classification. The dataset, obtained from Kaggle, consists of 3,168 balanced voice samples with 23 acoustic features. Using Pearson's correlation analysis, five features with the strongest associations to the target variable were selected. Random Forest classification was then conducted using both the full set of 22 features and the reduced set of 5 features. Results suggest that although the accuracy gain was marginal (98% to 99%), computation time decreased substantially from 0.3 to 0.1 seconds, representing a 66% efficiency improvement. These findings suggest that lightweight correlation-based feature selection can simplify models and enable faster real-time applications without compromising predictive performance. The study emphasizes efficiency rather than accuracy as the main contribution, providing a methodological insight for designing scalable and inclusive voice-based gender recognition systems.

Keywords—feature selection; gender classification; random forest; voice data.

1. Introduction

With the rapid development of intelligent voice assistants such as Alexa, Siri, and Google Assistant, gender recognition remains an open challenge for creating more adaptive and personalized interactions. Recent studies highlight that, although virtual assistants have advanced significantly, limitations still exist in tailoring responses to user demographics, including gender, in ways that are fair and context-sensitive (De Kloet & Yang, 2022; Taha et al., 2024). This indicates that the integration of lightweight gender recognition mechanisms remains relevant for improving personalization in human-computer interaction. If this capability were available, systems could provide more personalized services and better understand behavioral patterns. Pattern recognition systems play a crucial role in replicating human sensory abilities, especially hearing, by enabling computers to recognize logical patterns in audio signals (Kushwah et al., 2019). This study is thus motivated by the need to explore computational methods that allow machines to accurately recognize and classify voice patterns.

Human voices are inherently unique, influenced by physical and genetic characteristics, especially the anatomy of the vocal tract. Differences in vocal fold size and resonance cause male voices to typically have a lower pitch of around 120 Hz, while female voices average closer to 210 Hz (Shafhah et al., 2020). These differences, while biological, create measurable acoustic

* Corresponding author.

E-mail Address: suhardiyanto1680@gmail.com (Suhardiyanto)

Author E-mail(s): S (suhardiyanto1680@gmail.com), FA (amfitroh@gmail.com), AW (ariswijayanti5@gmail.com)

Digital Object Identifier 10.32815/jitika.v20i1.1204

Manuscript submitted 21 September 2025; revised 16 October 2025; accepted 18 October 2025.

ISSN: 2580-8397(O), 0852-730X(P).

features that can be leveraged for computational analysis. In the context of soft biometrics, gender recognition through voice is not only accessible and non-intrusive but also enhances other biometric systems such as facial and fingerprint recognition, thereby improving identification accuracy and reliability (Aithal et al., 2024; K. & Aithal, 2022; Dantcheva et al., 2016).

The importance of gender recognition through voice extends beyond biometrics to areas such as security, forensics, and human-computer interaction. Voice verification systems reduce the risks of fraud and identity misuse compared to password-based authentication (Bora et al., 2023; Štítilis et al., 2023). However, challenges remain, including algorithmic bias, as recognition accuracy often varies across demographic groups, potentially reinforcing social inequalities (Chen et al., 2022; Jansen et al., 2021). To address this, researchers emphasize the need for fairer algorithms and more diverse datasets to ensure inclusivity. In human-computer interaction, gender recognition enriches user experiences, enabling more adaptive and responsive systems (De Kloet & Yang, 2022; Taha et al., 2024).

At a technical level, gender classification relies on the extraction of acoustic features such as pitch, median frequency, formants, spectral entropy, and Mel-Frequency Cepstral Coefficients (MFCC), which are then processed through machine learning models (Jha et al., 2024; Sandhan et al., 2014). Yet, despite the richness of these features, the complexity of voice signals, such as frequency variations, background noise, and high dimensionality, poses challenges to building accurate classifiers. Without careful data preprocessing and feature management, machine learning models may fail to generalize effectively.

Feature selection has therefore become an essential step to enhance model performance in gender classification tasks. By reducing dimensionality and removing redundant variables, feature selection not only improves accuracy but also reduces computational costs and prevents overfitting (Ishak, 2022). Studies highlight that correlation-based feature selection and related approaches help focus the classification process on features with the most discriminative power, thereby yielding better accuracy and generalization (Farida & Mustopa, 2023; Swastika et al., 2023; Syukron et al., 2023; Yudiantoro & Sasongko, 2023). Moreover, normalization and class balancing techniques such as SMOTE have been proven effective in improving classifier robustness on imbalanced voice datasets (Suryanegara et al., 2021).

Random Forest has emerged as one of the most robust ensemble learning algorithms, capable of handling high-dimensional, non-linear data while maintaining resilience against noise and outliers (Wilson & Anwar, 2024). Built upon bootstrap aggregation of decision trees, Random Forest aggregates predictions through majority voting, thereby minimizing overfitting and improving accuracy. Its intrinsic ability to rank feature importance also makes it particularly suitable for classification tasks involving complex voice datasets (Maxwell et al., 2018; Wen & Hughes, 2020). Studies demonstrate its versatility in diverse applications, from medical diagnostics to environmental monitoring and speech analysis, underscoring its wide applicability (Islam et al., 2024; Ren et al., 2022; Xu et al., 2025; Zhou et al., 2022).

Despite numerous studies combining feature selection with ensemble methods such as Random Forest, most prior works have primarily emphasized accuracy improvement without

considering computational efficiency in real-time scenarios. In addition, while advanced feature selection techniques and Random Forest's internal importance measures (e.g., Gini or permutation importance) are frequently discussed, there is limited research highlighting the practicality of lightweight and easily interpretable methods such as Pearson correlation when applied to balanced voice datasets. This creates a research gap in exploring how simple feature selection can reduce model complexity while still achieving competitive results.

Therefore, the contribution of this study lies in two aspects: (1) demonstrating that even a straightforward correlation-based feature selection approach can meaningfully reduce computation time while maintaining high classification accuracy in Random Forest models, and (2) emphasizing computational efficiency as a critical dimension for real-time applications such as chatbots, virtual assistants, and digital authentication systems. By reframing efficiency rather than marginal accuracy gains as the core benefit, this study addresses an overlooked aspect in the literature and offers a methodological perspective for designing lightweight, scalable, and inclusive speech-based gender recognition systems.

2. Method

This study employs feature selection methods to optimize gender detection using the Random Forest algorithm in classifying human voices. The research process consists of several stages, beginning with feature selection to identify five features with correlation values closest to 1 with the target label, followed by splitting the dataset into training and testing sets, conducting gender detection using Random Forest, and finally comparing the performance of gender detection between using 22 features and 5 features to evaluate differences in accuracy and computation time. The overall research framework is as follows (Fig. 1):

- a. Dataset voice.csv (3,168 samples and 23 features)
- b. Feature Selection (Pearson Correlation), with 5 features selected
- c. Data Split (80% training and 20% testing)
- d. Model: Random Forest (n_estimators=100, max_depth=None, ...)
- e. Metrics Evaluation (Accuracy, Precision, Recall, F1-score, and Processing Time)

2.1. Human Voice Dataset

The dataset used in this research was obtained from Kaggle under the file name voice.csv. It consists of 3,168 labeled audio records, including 1,584 male voices and 1,584 female voices. Each record contains 23 acoustic features, with the target variable labeled as "label," representing the gender of each voice sample. This dataset is balanced and well-suited for gender classification tasks.

2.2. Feature Selection

Feature selection was conducted to identify and retain the most significant attributes contributing to gender classification. Using Pearson's correlation analysis, five features with correlation values approaching 1 with the target variable were

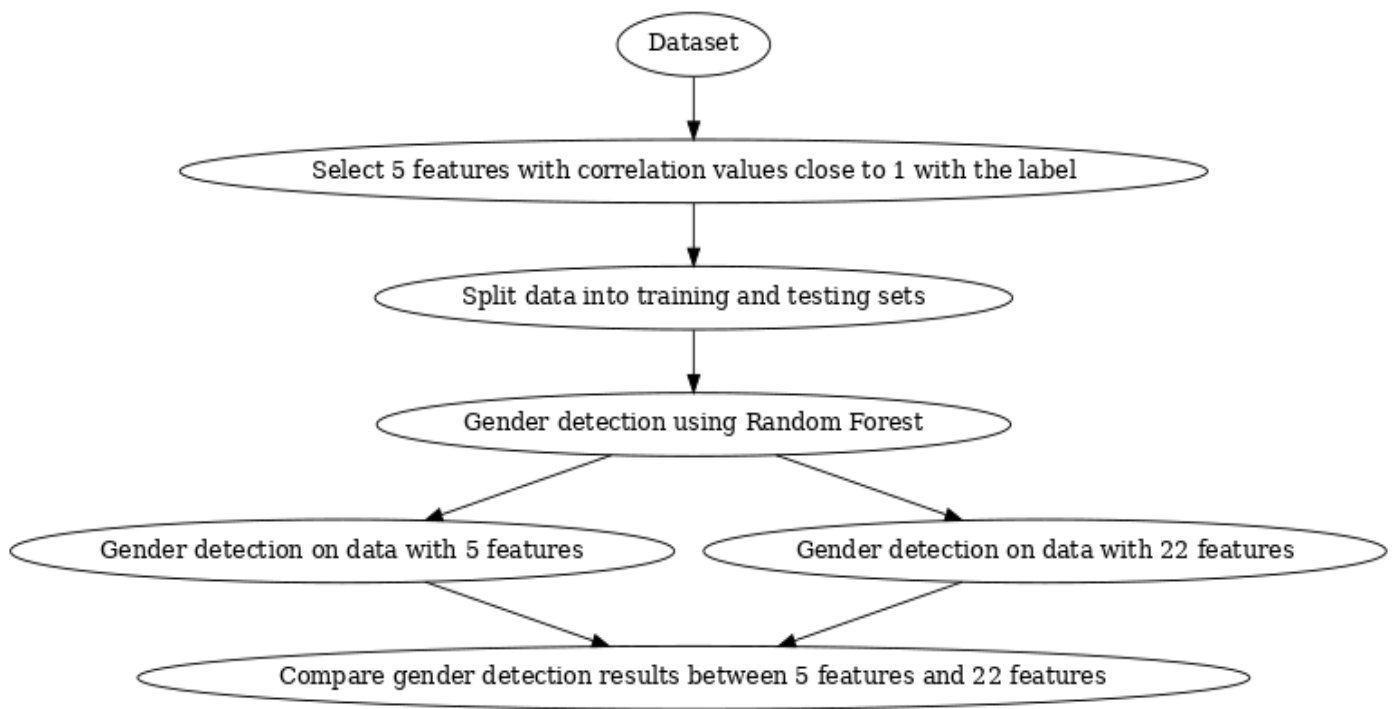


Fig. 1. Research block diagram

selected, indicating a very strong positive relationship. This step ensured that only the most relevant features were included in the model, thereby reducing complexity and enhancing efficiency. By focusing on these features, the study not only improved model performance but also provided more interpretable results, offering valuable insights for data-driven decision-making.

Pearson's correlation analysis was chosen as a lightweight and interpretable baseline method to identify features most associated with the target variable. Its simplicity makes it suitable for exploratory feature reduction in balanced datasets. However, it is acknowledged that acoustic data often exhibits non-linear relationships, and Random Forest itself provides built-in feature importance measures such as Gini Importance and Permutation Importance. While these alternatives may capture more complex dependencies, this study emphasizes correlation-based selection to evaluate whether a simple, computationally efficient approach can yield competitive results. Future work should incorporate a comparative analysis of these approaches to further validate robustness.

2.3. Data Splitting

After feature selection, the dataset was divided into training and testing subsets to evaluate model generalization. The split ratio was set at 80% for training and 20% for testing. Random splitting was applied to ensure that both subsets were representative of the entire dataset. This procedure allowed the model to learn patterns from the training data and subsequently be validated on unseen testing data. Such a method is critical to ensure that the evaluation reflects real-world performance and reduces the risk of overfitting.

For this study, a single 80/20 train-test split was employed. We explicitly acknowledge that this method is not as robust as k-fold cross-validation (e.g., $k = 5$ or 10) and that the results presented are subject to variance from this specific data partition. Therefore, the findings of this paper should not be interpreted as a definitive performance benchmark. Rather, this study serves as an illustrative proof-of-concept to investigate the potential impact of simple feature selection on computational efficiency. Given that the dataset is large and perfectly balanced, a single split is considered sufficient for this preliminary, illustrative purpose. We strongly emphasize that a follow-up study using rigorous k-fold cross-validation is essential to definitively validate these findings.

2.4. Random Forest Detection

The Random Forest classification was carried out twice: once using all 22 features and once using only the five selected features. This dual approach was designed to compare performance metrics, including accuracy, precision, recall, and F1-score, across both scenarios. By analyzing the results, the study assessed the effectiveness of feature selection in improving classification outcomes while also examining the robustness of Random Forest in processing acoustic signals. The comparative evaluation provides evidence of the significance of feature selection in machine learning pipelines and highlights Random Forest's potential as a reliable algorithm for voice-based gender recognition. Ultimately, the findings are expected to contribute to the broader field of speech recognition and open avenues for future research into more efficient and accurate classification algorithms.

The Random Forest model was implemented using scikit-

Table 2. Correlation values of selected features

Feature	Correlation Value
Meanfun (Mean Frequency)	0.83
IQR (Interquartile Range)	0.61
Q25 (Quartile 25)	0.51
sp.ent (Spectral Entropy)	0.49
sd (Standard Deviation)	0.47

learn with the following hyperparameters:

- `n_estimators = 100`,
- `max_depth = None`,
- `min_samples_split = 2`,
- `min_samples_leaf = 1`, and
- `random_state = 42`.

These default settings were chosen to emphasize the effect of feature selection rather than hyperparameter tuning. While no extensive optimization was performed, future work could apply grid search or Bayesian optimization to further improve performance.

3. Results

3.1. Feature Selection Results

From the initial 22 features, the five with the highest correlation values—Mean Frequency, Interquartile Range, Quartile 25, Spectral Entropy, and Standard Deviation—were retained. While these values ranged between 0.47 and 0.83, they still represent the strongest associations within the dataset. The choice of five features was made to balance interpretability and efficiency, serving as a proof of concept rather than claiming optimality. We acknowledge that a systematic ablation study (e.g., testing models with 3, 4, 6, or more features) would provide deeper insights into the optimal feature set, and recommend this direction for future research.

3.2. Comparison Between 5 Features and 22 Features

The comparison between models using all 22 features and those limited to 5 selected features highlights the crucial role of feature selection in enhancing model performance. Implementing feature selection led to notable improvements across evaluation metrics, demonstrating that Random Forest is particularly effective in managing high-dimensional data while delivering accurate predictions. By restricting the model to only five highly relevant features, the complexity of the model was reduced, resulting in faster processing time without compromising accuracy. In contrast, the model utilizing all 22 features yielded slightly lower performance metrics on the test set, where the model performed well on training data but struggled to generalize effectively on unseen data.

As shown in Table 2, the evaluation metrics—including accuracy, precision, recall, and F1-score—all improved slightly when using 5 features compared to 22. Furthermore, the processing time decreased significantly from 0.3 seconds to 0.1 seconds, underscoring the efficiency gained through feature

Table 1. Performance comparison between 22 features and 5 features

Evaluation Metric	22 Features	5 Features
Accuracy (%)	98	99
Precision (%)	98	99
Recall (%)	97	99
F1 Score (%)	98	99
Processing Time (s)	0.3	0.1

selection. These findings suggest that feature selection not only simplifies the model but also enhances its overall performance, making it a valuable approach in developing efficient and effective machine learning models for voice-based gender recognition.

4. Discussion

The experimental results indicate that feature selection improved model efficiency more substantially than accuracy. While the accuracy increased only marginally from 98% to 99%—a difference likely not statistically significant—the computational efficiency gain was substantial, with processing time reduced by 66% (0.3s to 0.1s). This suggests that the primary contribution of correlation-based feature selection lies in enabling faster and lighter models suitable for real-time deployment, rather than in marginal accuracy improvements (Aithal et al., 2024; K. & Aithal, 2022; Dantcheva et al., 2016). Voice-based biometrics offer the advantage of being unique to each individual and relatively easy to collect, making the improvements from feature selection particularly valuable for real-world applications.

The findings further highlight the critical role of voice-based gender recognition in security, forensics, and human–computer interaction. Higher accuracy obtained through feature selection suggests that systems using optimized models can improve identity verification, reducing risks of fraud and misuse that often arise with password-based systems (Bora et al., 2023; Štitilis et al., 2023). At the same time, the reduction in computational complexity without compromising accuracy is especially relevant in real-time applications such as voice assistants and surveillance systems. However, the ethical concerns regarding bias in biometric systems remain pressing. As Chen et al (2022) and Jansen et al (2021) argue, demographic biases in recognition systems can perpetuate inequality, underscoring the need for diverse datasets and fairness-oriented algorithms in future research.

The application of feature selection in this study directly addresses key challenges in voice data classification, including complexity, frequency variation, noise, and high dimensionality. The comparable performance obtained with fewer features supports the argument that correlation-based feature selection can help simplify models by removing potentially redundant features, thereby improving efficiency (Farida & Mustopa, 2023; Swastika et al., 2023; Yudiantoro & Sasongko, 2023). Furthermore, the efficiency gains shown in processing time reflect findings from previous research that emphasize the role of preprocessing, normalization, and data balancing techniques such as SMOTE in optimizing classification outcomes (Suryanegara et al., 2021; Syukron et al.,

2023). Thus, the present study reinforces the notion that effective feature management is as important as algorithm selection in machine learning pipelines.

The robustness of Random Forest as a classifier is also confirmed in this study. The algorithm's ensemble mechanism, which aggregates results from multiple decision trees, allowed it to effectively manage high-dimensional voice data while minimizing overfitting. Similar results have been documented across domains, including medical diagnostics and environmental monitoring, where Random Forest consistently outperforms other classifiers in handling noisy and non-linear data (Islam et al., 2024; Ren et al., 2022; Xu et al., 2025; Zhou et al., 2022). The ability of Random Forest to internally assess feature importance also complements the explicit feature selection process, further supporting its role as a reliable algorithm for speech-based applications (Maxwell et al., 2018; Wen & Hughes, 2020; Zhang et al., 2022).

Taken together, these findings suggest that combining feature selection with Random Forest may not significantly alter predictive accuracy but can substantially enhance efficiency, interpretability, and scalability in voice-based gender recognition. This study supports earlier research highlighting Random Forest's adaptability to complex and non-linear datasets (Han et al., 2017; Sen, 2024), while extending the discussion to the context of gender classification. At the same time, the results invite further exploration into integrating advanced feature extraction methods, such as deep learning-based representations (Kim & Yang, 2018) and multimodal approaches combining audio with visual or contextual data (De Kloet & Yang, 2022; Taha et al., 2024). Such directions will help ensure that future systems are not only accurate and efficient but also fair, inclusive, and ethically aligned with broader societal needs.

The selection of Mean Frequency as a key predictor aligns with the well-established pitch differences between male and female voices, typically centered around 120 Hz and 210 Hz, respectively. Interquartile Range (IQR) and Quartile 25 (Q25) reflect the distribution of spectral energy and provide insights into vocal tract resonance characteristics, which differ structurally between genders. Spectral Entropy measures the disorder of the frequency spectrum, with male voices often exhibiting lower entropy due to stronger harmonic structures, while female voices display greater spectral variability. Standard Deviation captures variability in frequency distribution, further distinguishing the dynamic range of male and female vocal patterns. Collectively, these features provide a strong acoustic rationale for why they are effective predictors of gender classification, beyond their statistical correlation values.

5. Conclusion

This study illustrated that applying correlation-based feature selection enhances the efficiency of Random Forest in classifying gender from voice data. By reducing the original 22 acoustic features to 5 key attributes, the model achieved a marginal improvement in accuracy from 98% to 99%. Furthermore, it substantially reduces computation time to approximately 66%. This suggests that efficiency, rather than accuracy alone, is the primary benefit of lightweight feature selection for real-time applications.

The outcomes of this research suggest that practical voice-based biometric systems can be designed to be both accurate

and computationally efficient, making them suitable for deployment in resource-constrained or real-time environments such as chatbots, digital authentication, and surveillance. Beyond technical performance, this study contributes by reframing efficiency as a critical dimension of machine learning pipelines and encourages future research to explore ablation studies for feature optimization as well as the integration of non-linear or fairness-oriented feature selection methods.

Data availability

All data produced or examined during this study are presented in this paper.

Declaration of 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.

Authors' contributions

All authors participated in the study design, writing, and manuscript revision. S drafted the initial manuscript, FA and AW revised it. All authors have reviewed and approved the final manuscript.

References

- Aithal, P. S., Prabhu, S., & Aithal, S. (2024). Future of Higher Education through Technology Prediction and Forecasting. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4901474>
- Amjad Hassan Khan M. K., & Aithal, P. S. (2022). Voice Biometric Systems for User Identification and Authentication – A Literature Review. *International Journal of Applied Engineering and Management Letters*, 198–209. <https://doi.org/10.47992/ijaeml.2581.7000.0131>
- Bora, B., Emanet, A. E., Elmaci, E., Kandaz, D., & Uçar, M. K. (2023). Hybrid AI-based Voice Authentication. *Turkish Journal of Forecasting*, 07(2), 17–22. <https://doi.org/10.34110/forecasting.1260073>
- Chen, X., Li, Z., Setlur, S., & Xu, W. (2022). Exploring Racial and Gender Disparities in Voice Biometrics. *Scientific Reports*, 12(1). <https://doi.org/10.1038/s41598-022-06673-y>
- Dantcheva, A., Elia, P., & Ross, A. (2016). What Else Does Your Biometric Data Reveal? A Survey on Soft Biometrics. *IEEE Transactions on Information Forensics and Security*, 11(3), 441–467. <https://doi.org/10.1109/tifs.2015.2480381>
- De Kloet, M., & Yang, S. (2022). The effects of anthropomorphism and multimodal biometric authentication on the user experience of voice intelligence. *Frontiers in Artificial Intelligence*, 5, 831046. <https://doi.org/10.3389/frai.2022.831046>
- Farida, F., & Mustopa, A. (2023). Comparison of Logistic Regression and Random Forest Using Correlation-Based Feature Selection for Phishing Website Detection. *Sistemasi*, 12(1), 13. <https://doi.org/10.32520/stmsi.v12i1.1832>
- Han, T., Jiang, D., Zhao, Q., Wang, L., & Yin, K. (2017). Comparison of Random Forest, Artificial Neural Networks and Support Vector Machine for Intelligent Diagnosis of Rotating Machinery. *Transactions of the Institute of Measurement and Control*, 40(8), 2681–2693. <https://doi.org/10.1177/0142331217708242>
- Ishak, R. (2022). Volume 4 Nomor 2 Juli 2022 Implementasi Seleksi Fitur Klasifikasi Waktu Kelulusan Mahasiswa Menggunakan Correlation Matrix With Heatmap. *Jambura Journal of Electrical and Electronics Engineering*, 169. <https://siakun.unisan.ac.id/>
- Islam, R., Imran, A., & Rabbi, Md. F. (2024). Prostate Cancer Detection

- From MRI Using Efficient Feature Extraction With Transfer Learning. *Prostate Cancer*, 2024, 1–28. <https://doi.org/10.1155/2024/1588891>
- Jansen, F., Sánchez-Monedero, J., & Dencik, L. (2021). Biometric Identity Systems in Law Enforcement and the Politics of (Voice) Recognition: The Case of SiP. *Big Data & Society*, 8(2). <https://doi.org/10.1177/20539517211063604>
- Jha, A. K., Singhal, N., & Chhabra, A. (2024). Voice Recognition Techniques: A Review Paper. *Educational Administration Theory and Practices*. <https://doi.org/10.53555/kuey.v30i3.5944>
- Kim, J.-H., & Yang, Y.-M. (2018). An Enhanced Classification Scheme With AdaBoost Concept in BCI. *Journal of Intelligent & Fuzzy Systems*, 35(1), 63–68. <https://doi.org/10.3233/jifs-169567>
- Kushwah, S., Singh, S., Vats, K., & Nemade, M. V. (2019). Gender Identification Via Voice Analysis. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 746–753. <https://doi.org/10.32628/cseit1952188>
- Maxwell, A. E., Warner, T. A., & Fang, F. (2018). Implementation of Machine-Learning Classification in Remote Sensing: An Applied Review. *International Journal of Remote Sensing*, 39(9), 2784–2817. <https://doi.org/10.1080/01431161.2018.1433343>
- Ren, Y., Wu, D., Singh, A. K., Kasson, E., Huang, M., & Cavazos-Rehg, P. (2022). Automated Detection of Vaping-Related Tweets on Twitter During the 2019 EVALI Outbreak Using Machine Learning Classification. *Frontiers in Big Data*, 5. <https://doi.org/10.3389/fdata.2022.770585>
- Sandhan, T., Sonowal, S., & Choi, J. Y. (2014). Audio Bank: A high-level acoustic signal representation for audio event recognition. *International Conference on Control, Automation and Systems*, 82–87. <https://doi.org/10.1109/ICCAS.2014.6987963>
- Sen, S. (2024). Comparison of Boosting and Random Forest Models in Forecasting Bank Failures: Revisiting the 2008 Financial Crisis From a Supervisory Perspective. *Economy & Finance*, 11(3), 258–281. <https://doi.org/10.3390/ef.2024.3.2>
- Shafhah, A. A., Adikara, P. P., & Adinugroho, S. (2020). *Klasifikasi Jenis Kelamin Berdasarkan Suara Menggunakan Metode Learning Vector Quantization*. <http://j-ptiik.ub.ac.id>
- Štitil, D., Laurinaitis, M., & Verenius, E. (2023). The Use of Biometric Technologies in Ensuring Critical Infrastructure Security: The Context of Protecting Personal Data. *Journal of Entrepreneurship and Sustainability Issues*, 10(3), 133–150. [https://doi.org/10.9770/jesi.2023.10.3\(10\)](https://doi.org/10.9770/jesi.2023.10.3(10))
- Suryanegara, G. A. B., Adiwijaya, A., & Purbolaksono, M. D. (2021). Peningkatan Hasil Klasifikasi Pada Algoritma Random Forest Untuk Deteksi Pasien Penderita Diabetes Menggunakan Metode Normalisasi. *Jurnal Resti (Rekayasa Sistem Dan Teknologi Informasi)*, 5(1), 114–122. <https://doi.org/10.29207/resti.v5i1.2880>
- Swastika, W., Widodo, R. B., & Oepojo, A. A. (2023). Perbandingan Akurasi Deteksi Emosi Pada Suara Menggunakan Multilayer Perceptron, Random Forest, Decision Tree Dan K-NN. *Journal of Intelligent System and Computation*, 5(1), 17–22. <https://doi.org/10.52985/insyst.v5i1.264>
- Syukron, A., Sardiarinto, S., Saputro, E., & Widodo, P. (2023). Penerapan Metode Smote Untuk Mengatasi Ketidakseimbangan Kelas Pada Prediksi Gagal Jantung. *Jurnal Teknologi Informasi Dan Terapan*, 10(1), 47–50. <https://doi.org/10.25047/jtit.v10i1.313>
- Taha, T. M., Messaoud, Z. B., & Frikha, M. (2024). Convolutional Neural Network Architectures for Gender, Emotional Detection From Speech and Speaker Diarization. *International Journal of Interactive Mobile Technologies (IJIM)*, 18(03), 88–103. <https://doi.org/10.3991/ijim.v18i03.43013>
- Wen, L., & Hughes, M. G. (2020). Coastal Wetland Mapping Using Ensemble Learning Algorithms: A Comparative Study of Bagging, Boosting and Stacking Techniques. *Remote Sensing*, 12(10), 1683. <https://doi.org/10.3390/rs12101683>
- Wilson, A., & Anwar, M. R. (2024). The Future of Adaptive Machine Learning Algorithms in High-Dimensional Data Processing. *International Transactions on Artificial Intelligence (ITALIC)*, 3(1), 97–107. <https://doi.org/10.33050/italic.v3i1.656>
- Xu, H., Xie, W., Pang, M., Li, Y., Jin, L., Huang, F., & Shao, X. (2025). Non-Invasive Detection of Parkinson's Disease Based on Speech Analysis and Interpretable Machine Learning. *Frontiers in Aging Neuroscience*, 17. <https://doi.org/10.3389/fnagi.2025.1586273>
- Yusdiantoro, S. Y., & Sasongko, T. B. (2023). Implementasi Algoritma MFCC Dan CNN Dalam Klasifikasi Makna Tangisan Bayi. *Indonesian Journal of Computer Science*, 12(4). <https://doi.org/10.33022/ijcs.v12i4.3243>
- Zhang, L., Wang, Y., Chen, J., & Chen, J. (2022). RFtest: A Robust and Flexible Community-Level Test for Microbiome Data Powerfully Detects Phylogenetically Clustered Signals. *Frontiers in Genetics*, 12. <https://doi.org/10.3389/fgene.2021.749573>
- Zhou, Y., Shen, H., & Zhang, M. (2022). A Distributed and Privacy-Preserving Random Forest Evaluation Scheme With Fine Grained Access Control. *Symmetry*, 14(2), 415. <https://doi.org/10.3390/sym14020415>



Suhardiyanto is currently a student at PGRI Ronggolawe University, Tuban, Indonesia. The focus of his studies is on the field of information technology and intelligent systems. He is actively developing research interests in the areas of machine learning, speech processing, and the application of artificial intelligence algorithms in various practical contexts. As a novice researcher, he is currently involved in several academic projects oriented towards data-based problem-solving and the development of innovative applications to support the needs of society and the world of education.



Fitroh Amaluddin is currently a lecturer at PGRI Ronggolawe University, Tuban, Indonesia. His areas of expertise and research include autonomous robots, image processing, and intelligence systems. His academic works have been published in various journals and proceedings, both national and international, with topics including vehicle classification using the Gaussian Mixture Model (GMM) and Fuzzy Cluster K-Means (FCM), microcontroller-based temperature and pressure sensor analysis, and the design of sales information systems using the Fast method. He is also active in applied research, such as the development of the E-Tracer Study system based on usability and user experience evaluation, coding introduction training for elementary school teachers using Scratch Jr, and the design of Arduino and QR Code-based systems.



Aris Wijayanti is currently a lecturer at the Informatics Engineering Study Program, PGRI Ronggolawe University, Tuban, Indonesia. Her areas of expertise include data mining, expert systems, and web-based and smart device application development. Her work has been published in various journals and national proceedings, with research topics that include the implementation of a priori algorithm for pharmacy data analysis, expert systems based on forward chaining and case-based reasoning, as well as facility improvement monitoring systems using Web GIS. She is also active in community service

activities, including coding introduction training for elementary school teachers with Scratch Jr. By 2025, her publications have obtained nearly 100 citations with an h-index of 5.