

# Integrating Potentiostat Measurements and Ensemble Learning for Water Pollution Estimation

Rizqy Ahsana Putri  
Department of Informatics Engineering  
Institut Teknologi Sepuluh Nopember  
Surabaya 60111, Indonesia  
[7025242013@student.its.ac.id](mailto:7025242013@student.its.ac.id)

Riyanarto Sarno\*  
Departement of Informatics Engineering  
Institut Teknologi Sepuluh Nopember  
Surabaya, Indonesia  
[riyanarto@its.ac.id](mailto:riyanarto@its.ac.id)

Wahyu Prasetyo Utomo  
Departement of Chemistry  
Institut Teknologi Sepuluh Nopember  
Surabaya, Indonesia  
[wp.utomo@its.ac.id](mailto:wp.utomo@its.ac.id)

Dwi Sunaryono  
Departement of Informatics Engineering  
Institut Teknologi Sepuluh Nopember  
Surabaya, Indonesia  
[dwi@its.ac.id](mailto:dwi@its.ac.id)

Taufiq Choirul Amri  
Department of Informatics Engineering  
Institut Teknologi Sepuluh Nopember  
Surabaya 60111, Indonesia  
[7025222001@student.its.ac.id](mailto:7025222001@student.its.ac.id)

**Abstract**—Microplastic pollution in water has emerged as a serious environmental concern due to its persistence and impact on aquatic ecosystems. Detecting microplastics in aqueous environments remains a complex task, particularly across different concentration levels and due to the lack of observable visual indicators. This study presents a machine learning approach for estimating microplastic concentrations using current-voltage signals generated by a potentiostat. These signals were processed through a feature extraction stage that identified six numerical descriptors, including peak current, voltage, and area from upper and lower signal regions. The resulting feature set was used as input for several machine learning algorithms, including Random Forest, K-Nearest Neighbors, Support Vector Machine, Logistic Regression, and an ensemble learning. Model evaluation was conducted using stratified 5-fold cross-validation to ensure balanced data partitioning. Performance was further enhanced through hyperparameter optimization using GridSearch. Among all tested models, the ensemble learning achieved the best results, with an accuracy of 92.33% after optimization, outperforming individual models in all evaluation metrics. These findings support the potential of ensemble learning strategies in improving the reliability of microplastic estimation based on potentiostat signals and offer a foundation for more scalable monitoring tools in future water quality assessment systems.

**Keywords**—ensemble learning, microplastic, potentiostat, voltammetry, water pollutant.

## I. INTRODUCTION

Water pollution remains one of the most pressing environmental concerns globally, threatening ecosystems, public health, and water resources. Among various types of pollutants, microplastics have recently gained increased attention due to their pervasive presence and potential risks [1]. Microplastics are plastic particles less than 5 mm in size, originating from the degradation of larger plastics or released directly through consumer products [2]. Their small size allows them to easily pass through filtration systems and enter aquatic environments, where they accumulate and persist. The detection of microplastics in water is particularly challenging due to their microscopic nature and diverse physical-chemical properties. Moreover, their ability to adsorb toxic chemicals and enter the food chain makes them a serious environmental and health hazard, highlighting the urgency of developing reliable detection and monitoring methods [3].

Recent advancements have focused on electrochemical sensing methods to address this challenge, particularly those employing potentiostats and techniques like voltammetry. A potentiostat is an electronic device which regulates the voltage between a working electrode and a reference electrode, simultaneously measuring the ensuing current flowing through a counter electrode [4]. This setup allows for the capture of the unique redox behavior of compounds within a solution, encompassing substances associated with microplastics. Cyclic voltammetry has gained popularity due to its high sensitivity and capacity to identify faint electrochemical signals. Distinct electrical signals are generated by the interaction between the microplastic-associated compounds and the electrode surface, enabling real-time monitoring.

However, the output of a potentiostat, commonly referred to as a voltammogram, is often complex and difficult to interpret, particularly for non-experts. These signals typically consist of nonlinear current-voltage relationships that vary depending on the nature and concentration of the analytes. This complexity necessitates the application of advanced data-driven approaches. Machine learning (ML) has emerged as a promising solution for this task [5, 6], offering tools to learn patterns and make predictions directly from raw voltammetric signals. Previous studies have widely employed ML to automate and enhance the detection of pollutants, leveraging its ability to handle high-dimensional, noisy, and nonlinear data.

Several previous works have explored the integration of voltammetry and machine learning for pollutant detection. For instance, one study proposed a voltammetric electronic tongue system built on COOH-functionalized multi-walled carbon nanotubes and screen-printed carbon electrodes [7]. This sensor array was able to discriminate ions such as  $K^+$ ,  $Ca^{2+}$ , and  $Mg^{2+}$  using cyclic voltammetry, achieving excellent recognition accuracy through PCA and SVM. The method showed promising long-term stability and high reproducibility over ten months. Another study by Leon-Medina et al. [8] applied support vector machines (SVM) to classify heavy metal ions such as arsenic (As), lead (Pb), and cadmium (Cd), achieving a classification accuracy of 98.31%. This confirmed the strength of SVM in handling voltammetric data for environmental applications.

Ye et al. [9] conducted a study examining classifications of voltammetric signals for six types of heavy metal ions through the application of various machine learning models such as SVM, KNN, decision trees, and ensemble methods. The results revealed a high level of classification performance, with accuracy rates reaching as high as 90%. The study primarily focused on comparing standard classifiers, without examining more sophisticated or adaptive ensemble methods. The findings suggested potential, yet highlighted the need for further methodological research to enhance performance and practical application.

A common limitation in previous research is the reliance on standard, off-the-shelf models without further structural enhancements or optimization. This research aims to address that gap by exploring the use of ensemble learning, specifically soft voting, to combine multiple classifiers for improved performance. Ensemble learning refers to the strategy of combining several models to produce a more robust and accurate prediction compared to individual models [10]. Soft voting, in particular, leverages the predicted probabilities from each model rather than their discrete class outputs, making it more flexible and informative than hard voting [11–13]. In addition to proposing soft voting as a primary classification strategy, this study also incorporates GridSearch for hyperparameter optimization, ensuring that each model operates under its most effective configuration.

This paper is organized as follows. Chapter 1 introduces the research background and objectives. Chapter 2 details the methodology, including data acquisition, signal preprocessing, and model development. Chapter 3 presents experimental results and a comparative analysis of the proposed models. Finally, Chapter 4 concludes the findings and outlines recommendations for future work.

## II. METHODOLOGY

This section outlines the methodological framework adopted in this study, which is summarized in Fig. 1. The process begins with data preparation and feature extraction. Following this, two classification strategies are employed: first, using individual machine learning models; and second, applying an ensemble learning approach to combine the strengths of multiple classifiers. Lastly, model performance is further optimized through hyperparameter tuning using GridSearch. Each of these stages is detailed in the subsections below.

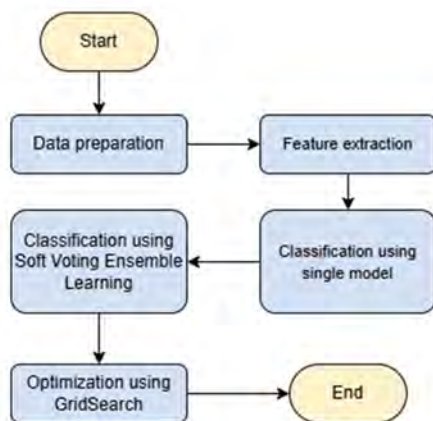


Fig. 1. Research stages

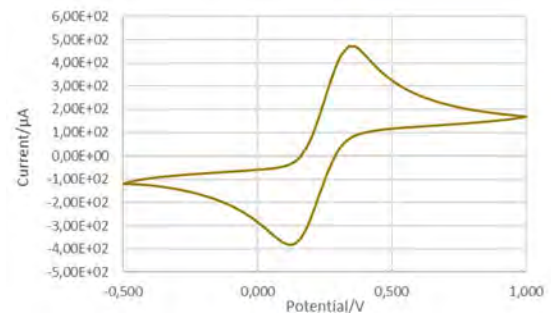


Fig. 2. Example of a voltammogram

### A. Data Preparation

The data utilised in this study were gathered through a series of controlled experiments that included microplastic-contaminated water samples. The microplastics were manually added to a base solution at different concentrations to simulate various levels of contamination. The prepared solutions were then evaluated using a potentiostat connected to a standard three-electrode setup comprising a working electrode, a reference electrode, and a counter electrode [4]. These electrodes were immersed in the prepared solution and connected to a potentiostat device, which was in turn interfaced with a personal computer.

During each measurement, the potentiostat applied a predefined voltage sweep and captured the resulting current responses, generating electrochemical profiles known as voltammograms. These voltammograms served as the primary raw data for subsequent analysis. The experiment was conducted across six distinct concentration levels of microplastics: 0 ppm, 200 ppm, 400 ppm, 600 ppm, 800 ppm, and 1000 ppm. For each concentration level, 50 samples were recorded, resulting in a total of 300 voltammogram signals. These signals were then compiled and labeled according to their respective concentration levels, forming a well-balanced dataset suitable for classification modeling.

### B. Feature Extraction

Once the voltammogram signals were collected, the next step was to extract key features that could serve as input for the classification models. A voltammogram reflects the relationship between applied voltage and measured current, and it typically contains peaks that indicate the electrochemical activity of compounds present in the solution. These peaks provide valuable information about the presence and behavior of analytes in this case, microplastics. Fig. 2 shows an example of a voltammogram recorded during the experiment.

To make the data suitable for machine learning, specific characteristics of the voltammogram were identified and extracted. From each sample, six features were derived: the peak current and peak voltage from both the upper and lower regions of the voltammogram, along with the area under the curve in each of those regions. These six numerical values capture the essential signal patterns and form a compact yet informative representation of each voltammogram, allowing the classifier to distinguish between different levels of microplastic contamination effectively.

### C. Classification Algorithm

This study investigates various classification methods to forecast microplastic concentration levels using extracted voltammetric characteristics. Initially, four conventional machine learning models were utilised individually as

classifiers: Random Forest (RF), K Nearest Neighbors (KNN), Support Vector Machine (SVM), and Logistic Regression (LR). A soft voting ensemble method is proposed in addition to these standalone models, to combine the predictions of multiple classifiers. Soft voting is chosen due to its capacity to take into account the predicted probabilities from each model, thereby enabling a more balanced and informed ultimate decision. The performance of the soft voting ensemble is compared with two deep learning models, specifically the Deep Neural Network (DNN) and the Gated Recurrent Unit (GRU). Four performance metrics were used to evaluate all classification models: accuracy, precision, recall, and F1 score.

#### 1) Random Forest (RF)

Random Forest (RF) is a potent ensemble learning method created to enhance classification accuracy by integrating the outputs of numerous independently trained decision trees [14]. During the training phase, the RF produces a diverse collection of decision trees by introducing randomness in both feature selection and data sampling. The model aggregates the predictions of all trees for each input instance, then assigns the class label with the most votes. A majority-voting approach helps to enhance the model's stability and predictive capability. One of the key benefits of RF is its capacity to generalise well, thereby significantly reducing the likelihood of overfitting that frequently occurs in single-tree models. RF also exhibits a high level of resistance to data noise and outliers.

#### 2) K-Nearest Neighbors (KNN)

The K-Nearest Neighbors (KNN) algorithm is a non-parametric classification technique that determines the class of a given sample based on the majority label among its  $k$  closest training instances in the feature space. This method operates based on a simple yet effective principle that samples located near each other in the input space tend to belong to the same class [15]. Unlike model-based approaches, KNN does not rely on prior assumptions about the underlying data distribution, and is therefore particularly suitable for situations involving unknown or highly irregular data structures.

A critical component of KNN is the choice of distance metric used to quantify similarity between data points. Among the available distance measures, Euclidean and Manhattan distances are the most frequently applied. Euclidean distance calculates the straight-line distance between two points in an  $n$ -dimensional space, while Manhattan distance measures the total absolute difference across dimensions, resembling movement along orthogonal grid paths.

The Euclidean distance is defined in Eq. (1), and the Manhattan distance is presented in Eq. (2).

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (1)$$

$$d(x, y) = \sum_{i=1}^n |x_i - y_i| \quad (2)$$

Where  $d(x, y)$  denotes the distance between two data points  $x$  and  $y$ . The terms  $x_i$  and  $y_i$  refer to the values of the  $i$ -th feature or dimension of  $x$  and  $y$ , respectively, and  $n$  indicates the total number of dimensions considered in the feature space.

#### 3) Support Vector Machine (SVM)

The Support Vector Machine is a supervised learning algorithm that performs classification by finding the optimal hyperplane which maximally separates data points from two distinct classes [16]. The SVM uses a kernel function to map input data into a higher-dimensional feature space, thereby handling non-linearly separable data. This approach has demonstrated strong performance, especially when working with unstructured or high-dimensional datasets. The decision function of an SVM classifier is mathematically expressed in Eq. (3).

$$f(x) = \text{sign}(w \cdot x + b) \quad (3)$$

In Eq. (3),  $f(x)$  denotes the predicted class label for a given input vector  $x$ ,  $w$  is the weight vector defining the orientation of the hyperplane, and  $b$  is the bias term that shifts the hyperplane from the origin.

#### 4) Logistic Regression (LR)

Logistic Regression (LR) is a statistical classification method used to model the relationship between input features and a categorical outcome [17]. By applying the logistic (sigmoid) function, LR estimates the probability of an event occurring based on the given input variables. The logistic regression model is expressed in Eq. (4).

$$P(y = 1 | x) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x_1 + \dots + \beta_n x_n)}} \quad (4)$$

Where  $P(y = 1 | x)$  represents the probability of the positive class,  $\beta_0$  is the intercept, and  $\beta_1, \dots, \beta_n$ , are the coefficients associated with each input feature  $x_1, \dots, x_n$ .

#### 5) Ensemble Learning

Ensemble learning is a powerful machine learning strategy that combines the predictions of multiple base models to improve overall performance, particularly in classification tasks. One of the most common ensemble approaches is voting, where multiple models "vote" on the final prediction. In hard voting, each model casts a single vote for a class label,

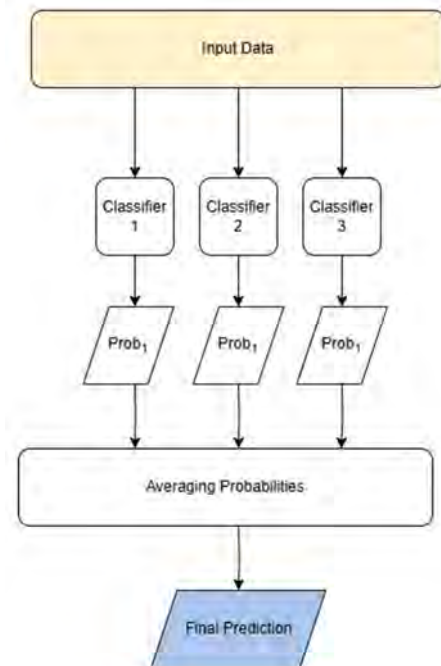


Fig. 3. Conceptual diagram of soft voting ensemble

and the majority vote determines the outcome. However, hard voting ignores the confidence level or predicted probability associated with each model's output, which can lead to suboptimal results, especially when classifiers vary in performance. In contrast, soft voting aggregates the predicted probabilities from each model rather than their final class decisions [11, 13]. The class with the highest average probability across all models is chosen as the final output. This allows models with higher confidence to have a greater influence, making soft voting generally more robust and accurate, especially when combining diverse classifiers. The overall concept of soft voting as used in this study is illustrated in Fig. 3.

#### 6) Deep Neural Network (DNN)

Artificial neural networks of the Deep Neural Network (DNN) class are characterised by multiple hidden layers between input and output layers [18, 19]. The network is able to learn complex patterns through the extraction of hierarchical representations from the input data. In deep neural networks, each neuron within a layer is connected to every neuron in the subsequent layer, with each connection having an associated weight that is adjusted during training through backpropagation. The activation functions, such as ReLU or sigmoid, introduce non-linearity into the model, allowing it to capture complex relationships in the data.

#### 7) Gated Recurrent Unit (GRU)

The Gated Recurrent Unit (GRU) is a type of recurrent neural network (RNN) designed to model sequential data more efficiently than traditional RNNs [20]. GRU addresses the vanishing gradient problem by using gating mechanisms that regulate the flow of information through the network. GRU combines the forget and input gates of an LSTM into a single update gate, making the architecture simpler and faster to train while still capturing long-term dependencies in the data.

#### D. Optimization using Gridsearch

After the initial classification was performed using default parameters, this study proceeded with hyperparameter tuning to improve the overall model performance. Hyperparameters play a crucial role in how machine learning models learn from data, and selecting the right combination can significantly enhance accuracy and generalization. To systematically

TABLE I. HYPERPARAMETER SPACE FOR EACH MODEL

| Model                  | Hyperparameter        | Values Tested              |
|------------------------|-----------------------|----------------------------|
| Random Forest          | n_estimators          | [50, 100]                  |
|                        | max_depth             | [10, 20]                   |
|                        | min_samples_split     | [2, 5, 10]                 |
| K-Nearest              | n_neighbors           | [3, 5, 7]                  |
| Neighbors              | weights               | ['uniform', 'distance']    |
|                        | metric                | ['euclidean', 'manhattan'] |
| Logistic Regression    | C                     | [0.01, 0.1, 0.5, 1, 5, 10] |
|                        | solver                | ['lbfgs', 'liblinear']     |
| Support Vector Machine | C                     | [0.1, 1, 10]               |
|                        | gamma                 | [0.1, 1, 10]               |
| Deep Neural Network    | num_dense_layers_list | [2, 4, 6]                  |
|                        | hidden_units_list     | [64, 128]                  |
|                        | dropout_rate_list     | [0.2, 0.3]                 |
| Gated Recurrent Unit   | num_gru_layers_list   | [2, 4, 6]                  |
|                        | hidden_units_list     | [64, 128]                  |
|                        | dropout_rate_list     | [0.2, 0.3]                 |

TABLE II. FEATURE EXTRACTION RESULT

| ID  | Upper Peak Area | Upper Peak Voltage (V) | Upper Peak Current (μA) | Lower Peak Area | Lower Peak Voltage (V) | Lower Peak Current (μA) | Label    |
|-----|-----------------|------------------------|-------------------------|-----------------|------------------------|-------------------------|----------|
| 1   | 173.15          | 0.3407                 | 571.34                  | 128.11          | 0.1405                 | -451.33                 | 0 ppm    |
| 2   | 168.95          | 0.3307                 | 558.07                  | 129.61          | 0.1405                 | -453.51                 | 0 ppm    |
| 3   | 168.13          | 0.3307                 | 554.85                  | 129.95          | 0.1405                 | -453.89                 | 0 ppm    |
| 4   | 167.96          | 0.3307                 | 554.18                  | 130.06          | 0.1405                 | -454.17                 | 0 ppm    |
| 5   | 167.92          | 0.3307                 | 553.90                  | 130.06          | 0.1405                 | -454.55                 | 0 ppm    |
| ... | ...             | ...                    | ...                     | ...             | ...                    | ...                     | ...      |
| 296 | 145.52          | 0.3407                 | 512.38                  | 123.97          | 0.1405                 | -457.87                 | 1000 ppm |
| 297 | 145.50          | 0.3407                 | 512.85                  | 123.93          | 0.1405                 | -458.06                 | 1000 ppm |
| 298 | 145.50          | 0.3407                 | 513.23                  | 123.92          | 0.1405                 | -458.25                 | 1000 ppm |
| 299 | 145.67          | 0.3407                 | 513.80                  | 123.92          | 0.1405                 | -458.44                 | 1000 ppm |
| 300 | 145.66          | 0.3407                 | 513.80                  | 123.82          | 0.1405                 | -458.63                 | 1000 ppm |
| 296 | 145.52          | 0.3407                 | 512.38                  | 123.97          | 0.1405                 | -457.87                 | 1000 ppm |

explore the best configuration, we employed GridSearch, a widely used optimization method in machine learning.

GridSearch operates by exhaustively evaluating all possible combinations of predefined hyperparameter values [21, 22]. Each combination is assessed using cross-validation to ensure consistent performance across different data splits. The configuration that yields the best evaluation score is selected as the optimal one. While this method can be computationally demanding, it offers a reliable way to fine-tune models beyond their default settings. In the context of this study, GridSearch allows each classifier to adapt more precisely to the characteristics of the potentiostat-derived signal data. The following table presents the full parameter space explored for each model in this study. Table I summarizes the hyperparameter search space defined for each classification model evaluated in this work.

### III. RESULT AND DISCUSSION

#### A. Feature Extraction Result

The raw voltammetric signals produced by the potentiostat were initially in the form of continuous current-voltage curves, which are not readily usable as direct input for classification algorithms. Therefore, a feature extraction process was carried out to convert these signals into structured data. As described in the methodology section, this involved identifying the upper and lower peaks in each voltammogram and calculating key values associated with these regions.

Each voltammogram was summarized into six numerical features: the area under the upper peak, the voltage at which the upper peak occurs, the corresponding peak current, and the same three values extracted from the lower peak. This transformation resulted in a tabular dataset where each row represents a single voltammogram sample with its extracted features and label. Table II presents a preview of the structured dataset obtained after feature extraction.

#### B. Classification and Optimization Result

After the features were successfully extracted from each voltammogram, the next stage focused on evaluating the ability of various classification models to distinguish between different levels of microplastic concentration. The evaluation included both conventional machine learning models and deep learning architectures, with comparisons made before and after parameter optimization. In addition, an ensemble learning was examined to assess whether combining multiple classifiers could offer improvements over individual models.

The dataset was divided using Stratified K-Fold Cross Validation with k set to 5 to guarantee a balanced and impartial assessment. The method splits the data into five subsets, ensuring that each subset retains the original class distribution. One fold was used for testing at a time, with the other four employed for model training. Unlike simple train-test splits, the stratified K-fold method guarantees that each class is evenly represented across every fold, especially in multi-class classification scenarios like this one. This approach minimizes the risk of performance bias and offers a more stable and reliable evaluation of model generalizability.

The initial evaluation was performed using default parameters for all models, as shown in Table III. Among the machine learning classifiers, the ensemble learning achieved the highest overall performance with an accuracy of 89.33%, precision of 90.96%, recall of 89.33%, and F1 score of 89.07%. Individual models such as Random Forest, K-Nearest Neighbors, Support Vector Machine, and Logistic Regression followed closely behind with accuracy ranging from 87% to 88.33%. Deep learning models were also included for comparison, where Deep Neural Network reached an accuracy of 87% and Gated Recurrent Unit showed significantly lower performance with only 74%.

Ensemble learning demonstrated clear advantages over individual machine learning models in terms of all evaluation metrics. By averaging the predicted probabilities from multiple classifiers, the ensemble learning was able to leverage the strengths of each base model, resulting in more balanced and accurate predictions. Although Deep Neural Network approached the performance of ensemble learning, the ensemble still maintained a slight edge. Gated Recurrent Unit, on the other hand, underperformed due to the nature of the input data, which consisted of extracted features rather than sequential time-series. The fixed-length feature vectors did not fully exploit the temporal modeling capabilities of GRU, which contributed to its lower classification accuracy.

To further enhance model performance, hyperparameter optimization was performed using GridSearch, and the results are presented in Table IV. For each model, the best-performing parameter configuration was identified based on cross-validation results. After optimization, almost all models showed improvements in accuracy and other metrics. Logistic Regression achieved an accuracy of 91%, SVM improved to

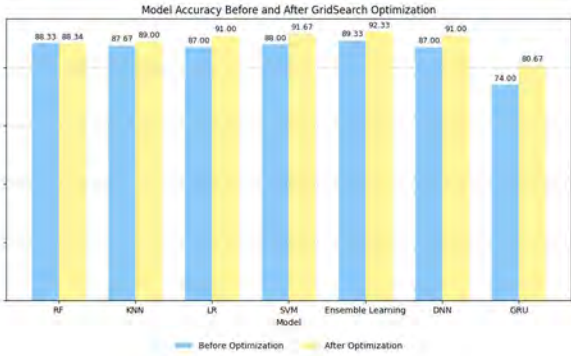


Fig. 4. Accuracy comparison before and after GridSearch optimization

91.67%, and K-Nearest Neighbors increased to 89%. The optimized ensemble learning reached an accuracy of 92.33%, along with a precision of 93.91%, recall of 92.33%, and F1 score of 92.19%, confirming its position as the most effective model.

The post-optimization results confirmed that ensemble learning consistently outperformed all individual model, even when each model was operating under its best configuration. Deep Neural Network also improved to an accuracy of 91%, closely following ensemble learning. Nevertheless, the ensemble approach remained slightly superior, likely due to its ability to integrate diverse decision boundaries from multiple algorithms. Gated Recurrent Unit also showed an increase in performance, reaching 80.67% accuracy, but remained the least effective among all evaluated models.

To better visualize the impact of optimization, a comparison of model performance before and after GridSearch is presented in Fig. 4. The bar chart highlights noticeable improvements across nearly all models. Most notably, Logistic Regression, SVM, and Deep Neural Network experienced significant gains in both accuracy and F1 score. The enhancements observed in ensemble learning were also consistent, reinforcing the robustness of ensemble-based methods.

IV. CONCLUSION

This study demonstrated that machine learning models can effectively classify microplastic concentration levels in water based on features extracted from voltammetric signals generated by a potentiostat. Among all models tested, the ensemble learning consistently achieved the best performance, with an accuracy of 89.33% before optimization and improving to 92.33% after hyperparameter tuning. The transformation of raw electrochemical signals into structured numerical features enabled a clearer interpretation by the models, and further improvements were observed after optimization using GridSearch. Although the optimization process required additional computational time, the performance gains confirmed its value. For future work, other ensemble approaches may be explored to further improve classification results, and expanding the range of hyperparameters could lead to better configurations while maintaining a balance between accuracy and efficiency.

ACKNOWLEDGMENT

This study was funded by the Directorate General of Higher Education, Research and Technology of the Republic of Indonesia through the Master Program of Education Leading to a Doctoral Degree for Excellent Graduates (PMDSU) Batch VII, under Grant Number

TABLE III. MODEL PERFORMANCE BEFORE HYPERPARAMETER OPTIMIZATION

| Model               | Accuracy | Precision | Recall | F1 Score |
|---------------------|----------|-----------|--------|----------|
| Random Forest       | 88.33    | 89.76     | 88.33  | 88.20    |
| KNN                 | 87.67    | 89.55     | 87.67  | 87.50    |
| Logistic Regression | 87.00    | 89.62     | 87.00  | 87.11    |
| SVM                 | 88.00    | 90.53     | 88.00  | 87.92    |
| Ensemble Learning   | 89.33    | 90.96     | 89.33  | 89.07    |
| DNN                 | 87.00    | 89.62     | 87.00  | 87.07    |
| GRU                 | 74.00    | 76.91     | 74.00  | 73.09    |

TABLE IV. MODEL PERFORMANCE AFTER HYPERPARAMETER OPTIMIZATION

| Model               | Accuracy | Precision | Recall | F1 Score |
|---------------------|----------|-----------|--------|----------|
| Random Forest       | 88.34    | 89.66     | 87.97  | 87.76    |
| KNN                 | 89.00    | 89.31     | 87.78  | 87.55    |
| Logistic Regression | 91.00    | 89.85     | 87.44  | 87.51    |
| SVM                 | 91.67    | 88.74     | 86.33  | 86.36    |
| Ensemble Learning   | 92.33    | 93.91     | 92.33  | 92.19    |
| DNN                 | 91.00    | 92.73     | 91.00  | 90.93    |
| GRU                 | 80.67    | 84.84     | 80.67  | 80.95    |

\*all results in percentage(%)

017/C3/DT.05.00/PL/2025 and Local Grant Number 1179/PKS/ITS/2025. Additional support was provided by Institut Teknologi Sepuluh Nopember (ITS) through Riset Kolaborasi Indonesia, Penelitian Kemitraan A, and Penelitian Departemen. This research also received funding from the Ethereum Foundation under Academic Grants Round 2025.

## REFERENCES

- [1] L. Borriello, M. Scivico, N. A. Cacciola, F. Esposito, L. Severino, and T. Cirillo, "Microplastics, a Global Issue: Human Exposure through Environmental and Dietary Sources," *Foods*, vol. 12, no. 18, p. 3396, Sep. 2023, doi: 10.3390/foods12183396.
- [2] A. Baruah, A. Sharma, S. Sharma, and R. Nagraik, "An insight into different microplastic detection methods," *International Journal of Environmental Science and Technology*, vol. 19, no. 6, pp. 5721–5730, Jun. 2022, doi: 10.1007/s13762-021-03384-1.
- [3] E. C. Emenike, C. J. Okorie, T. Ojeyemi, A. Egbemhenge, K. O. Iwuozor, O. D. Saliu, H. K. Okoro, and A. G. Adeniyi, "From oceans to dinner plates: The impact of microplastics on human health," *Heliyon*, vol. 9, no. 10, p. e20440, Oct. 2023, doi: 10.1016/j.heliyon.2023.e20440.
- [4] A. W. Colburn, K. J. Levey, D. O'Hare, and J. V. Macpherson, "Lifting the lid on the potentiostat: a beginner's guide to understanding electrochemical circuitry and practical operation," *Physical Chemistry Chemical Physics*, vol. 23, no. 14, pp. 8100–8117, 2021, doi: 10.1039/D1CP00661D.
- [5] X. Liu, D. Lu, A. Zhang, Q. Liu, and G. Jiang, "Data-driven machine learning in environmental pollution: gains and problems," *Environ Sci Technol*, vol. 56, no. 4, pp. 2124–2133, 2022.
- [6] R. Sarno and D. R. Wijaya, "Recent development in electronic nose data processing for beef quality assessment," *Telkomnika (Telecommunication Computing Electronics and Control)*, vol. 17, no. 1, pp. 337–348, Feb. 2019, doi: 10.12928/TELKOMNIKA.v17i1.10565.
- [7] T. Lu, A. Al-Hamry, A. Adiraju, M. Yan, R. Azhati, M. Talbi, Z. Wu, G. Shi, and O. Kanoun, "Electronic Tongue Based on Composites of Metal Phthalocyanine and Carbon Nanotubes and Electrochemically Deposited Metal Nanoparticles for Metal Ions Detection Enhanced by Machine Learning," In: *2023 International Workshop on Impedance Spectroscopy (IWIS)*, Sep. 2023, pp. 148–153. doi: 10.1109/IWIS61214.2023.10302765.
- [8] J. X. Leon-Medina, D. A. Tibaduiza, J. C. Burgos, M. Cuenca, and D. Vasquez, "Classification of As, Pb and Cd heavy metal ions using square wave voltammetry, dimensionality reduction and machine learning," *IEEE Access*, vol. 10, pp. 7684–7694, 2022.
- [9] J.-J. Ye, C.-H. Lin, and X.-J. Huang, "Analyzing the anodic stripping square wave voltammetry of heavy metal ions via machine learning: Information beyond a single voltammetric peak," *Journal of Electroanalytical Chemistry*, vol. 872, p. 113934, 2020.
- [10] M. Malikah, R. Sarno, and S. Sabilla, "Ensemble Learning for Optimizing Classification of Pork Adulteration in Beef Based on Electronic Nose Dataset," *International Journal of Intelligent Engineering and Systems*, vol. 14, no. 4, pp. 44–55, Aug. 2021, doi: 10.22266/ijies2021.0831.05.
- [11] N. Rai, N. Kaushik, D. Kumar, C. Raj, and A. Ali, "Mortality prediction of COVID-19 patients using soft voting classifier," *International Journal of Cognitive Computing in Engineering*, vol. 3, pp. 172–179, Jun. 2022, doi: 10.1016/j.ijcce.2022.09.001.
- [12] S. Kumari, D. Kumar, and M. Mittal, "An ensemble approach for classification and prediction of diabetes mellitus using soft voting classifier," *International Journal of Cognitive Computing in Engineering*, vol. 2, pp. 40–46, Jun. 2021, doi: 10.1016/j.ijcce.2021.01.001.
- [13] M. A. Khan, N. Iqbal, Imran, H. Jamil, and D.-H. Kim, "An optimized ensemble prediction model using AutoML based on soft voting classifier for network intrusion detection," *Journal of Network and Computer Applications*, vol. 212, p. 103560, Mar. 2023, doi: 10.1016/j.jnca.2022.103560.
- [14] R. A. Putri, S. I. Sabilla, and R. Sarno, "Optimal Feature Selection Algorithm (FSA) for Electronic Nose Signal," In: *2023 1st International Conference on Advanced Engineering and Technologies (ICONNIC)*, 2023, pp. 310–315.
- [15] I. A. Sabilla, R. A. Irawan, R. Sarno, A. Al Fauzi, D. R. Wijaya, and R. Gunawan, "Classification of Male and Female Sweat Odor in the Morning Using Electronic Nose," In: *2021 3rd East Indonesia Conference on Computer and Information Technology (EIConCIT)*, Apr. 2021, pp. 320–324. doi: 10.1109/EIConCIT50028.2021.9431909.
- [16] U. I. Shabrina, R. Sarno, R. N. E. Anggraini, A. T. Haryono, and A. F. Septiyanto, "Sentiment Analysis of Presidential Candidate Debates from YouTube Videos," In: *2024 IEEE International Conference on Artificial Intelligence and Mechatronics Systems (AIMS)*, 2024, pp. 1–6.
- [17] A. Cemiloglu, L. Zhu, A. B. Mohammednour, M. Azarafza, and Y. A. Nanehkaran, "Landslide susceptibility assessment for Maragheh County, Iran, using the logistic regression algorithm," *Land (Basel)*, vol. 12, no. 7, p. 1397, 2023.
- [18] R. Sarno, S. I. Sabilla, and D. R. Wijaya, "Electronic Nose for Detecting Multilevel Diabetes using Optimized Deep Neural Network," *Engineering Letters*, vol. 28, no. 1, 2020.
- [19] A. Kamilah, R. Sarno, K. R. Sungkono, A. F. Septiyanto, R. A. Putri, T. C. Amri, F. Taufany, D. Sunaryono, and S. I. Sabilla, "Deep Neural Network (DNN) for Outlier Detection in Heavy Metal Voltammetric Data," In: *2024 Beyond Technology Summit on Informatics International Conference (BTS-I2C)*, Dec. 2024, pp. 415–420. doi: 10.1109/BTS-I2C63534.2024.10942294.
- [20] A. Bhavani, A. V. Ramana, and A. S. N. Chakravarthy, "Comparative Analysis between LSTM and GRU in Stock Price Prediction," In: *2022 International Conference on Edge Computing and Applications (ICECAA)*, Oct. 2022, pp. 532–537. doi: 10.1109/ICECAA55415.2022.9936434.
- [21] D. P. Mishra, H. K. Gupta, G. Saajith, and R. Bag, "Optimizing Heart Disease Prediction Model with GridsearchCV for Hyperparameter Tuning," In: *2024 1st International Conference on Cognitive, Green and Ubiquitous Computing (IC-CGU)*, Mar. 2024, pp. 1–6. doi: 10.1109/IC-CGU58078.2024.10530772.
- [22] D. D. C. Maceda and J. C. D. Cruz, "Rainfall Classification Model for the Philippines using Optimized K-nearest Neighbor Algorithm with GridSearchCV Hyperparameter Tuning," In: *2023 IEEE 13th International Conference on Control System, Computing and Engineering (ICCSCE)*, Aug. 2023, pp. 51–55. doi: 10.1109/ICCSCE58721.2023.10237156.