

Global Food Waste Prediction (2018-2024): Trend Analysis and Random Forest Regression Model Development

Kemal Pramayuda^{1,*}, Fidi Supriadi¹, David Setiadi¹

¹Department of Informatics, Faculty of Information Technology, Sebelas April University, Sumedang, Indonesia

*Author to whom any correspondence should be addressed.

E-mail: kemalpramayuda10@gmail.com

Received: November 12, 2025

Accepted for publication: Juni 22, 2026

Published xxxxxx

ABSTRACT

Food waste remains a major global sustainability challenge due to its environmental, economic, and social impacts. Understanding food waste patterns and their contributing factors is essential for supporting effective mitigation strategies. This study aims to (1) investigate food waste trends and patterns through Exploratory Data Analysis (EDA) and (2) develop a predictive model for estimating total food waste using Random Forest Regression. The study utilizes a publicly available dataset containing 5,000 records from 20 countries, covering eight food categories over the period 2018–2024. The dataset includes variables such as food category, economic loss, population, average waste per capita, and household waste percentage. Exploratory analysis reveals variations in waste generation across food categories and countries, with fruits and vegetables contributing a substantial share of total waste. A Random Forest Regression model was developed and evaluated, achieving a coefficient of determination (R^2) of 0.9582. In addition to predictive performance, the study highlights the importance of examining key contributing variables to better understand food waste patterns. The findings demonstrate the potential of machine learning techniques as decision-support tools for food waste analysis and management, while also acknowledging the limitations associated with secondary public datasets.

Keywords: food waste, random forest regression, machine learning, predictive analytics, sustainability

I. Introduction

Food loss and waste represent one of the most pressing sustainability challenges worldwide. It refers to food originally intended for human consumption that is subsequently lost or discarded throughout the food supply chain, from production to final consumption [1]. Globally, approximately 25–30% of all food produced annually is lost or wasted, resulting in substantial economic losses, environmental degradation, and persistent food insecurity. Furthermore, food waste contributes approximately 8–10% of global greenhouse gas emissions, making it a significant driver of climate change and a critical issue in achieving sustainable development goals [2, 3].

Understanding food waste generation and its associated factors is essential for developing effective mitigation strategies. Conventional statistical approaches have been widely employed to analyze food waste; however, they often face limitations when modeling complex and non-linear relationships among demographic, economic, and behavioral variables. In recent years, Machine Learning (ML) techniques have demonstrated strong potential for extracting patterns from large and heterogeneous datasets, enabling more accurate predictions and supporting data-driven decision making [4].

Previous studies have applied Machine Learning methods to food waste prediction in specific domains, including households, hospitality services, retail operations, and regional waste management systems [5], [6], [7]. Although these studies have reported promising predictive performance, most remain focused on particular sectors or geographic regions. In addition, many studies emphasize model accuracy while providing limited discussion regarding food waste patterns and the relative contribution of different influencing factors. Consequently, comprehensive analyses that combine exploratory investigation and predictive modeling using multi-country food waste datasets remain relatively limited [8], [9], [10], [11].

To address this gap, this study utilizes a publicly available food waste dataset consisting of 5,000 records covering 20 countries, eight food categories, and the period from 2018 to 2024. The objectives of this research are threefold: first, to analyze food waste trends and patterns through Exploratory Data Analysis (EDA); second, to develop and evaluate a Random Forest Regression model for predicting total food waste generation based on demographic, economic, and food-related variables; and third, to examine the contribution of key variables associated with food waste prediction through feature importance analysis. By integrating exploratory analysis with predictive modeling using the Random Forest algorithm [12], this study seeks to provide both predictive performance evaluation and analytical insights that may support future food waste management and sustainability initiatives.

II. Related work

Recent studies have increasingly employed Machine Learning (ML) techniques to support food waste prediction and management. Traditional statistical approaches, including regression and time-series forecasting models, have been widely used to estimate waste generation; however, their ability to capture complex and non-linear relationships is often limited. Consequently, ML-based approaches have gained attention due to their flexibility and predictive capability when handling heterogeneous datasets. The increasing availability of large-scale food-related datasets and advances in big data analytics have further enhanced the application of machine learning techniques for food system optimization and waste management [4], [13].

Several studies have explored ML applications in food waste prediction within specific sectors and regions. Lee and Shin [5] applied machine learning techniques to predict waste generation at the regional level and reported satisfactory predictive performance. Similarly, Turker [6] utilized data-driven forecasting methods in campus dining services to support food waste reduction strategies. Rodrigues et al. [7] investigated machine learning models for short-term demand forecasting in food catering services and demonstrated their effectiveness in minimizing food waste. Other studies have explored forecasting approaches based on time-series analysis and integrated machine learning frameworks for waste management applications [8], [10].

Tree-based ensemble algorithms have also been widely adopted due to their robustness and ability to model non-linear relationships. Popescu et al. [9] demonstrated the effectiveness of machine learning techniques for food loss and waste prediction. In addition, studies by Lakshmi et al. [11], Weerasooriya et al. [14], and Uganya et al. [15] further highlighted the potential of machine learning approaches in waste prediction and intelligent waste management systems. These methods are particularly advantageous when datasets contain both numerical and categorical variables, as they can capture complex interactions while reducing the risk of overfitting. The Random Forest algorithm, originally introduced by Breiman [12], remains one of the most widely adopted ensemble learning methods for predictive modeling tasks.

Despite these advances, existing studies predominantly focus on specific regions, sectors, or operational environments, such as households, retail systems, catering services, or municipal waste management. Furthermore, many studies emphasize predictive performance while providing limited discussion regarding food waste characteristics, category-level patterns, and the relative importance of influencing variables. Therefore, there remains an opportunity to combine exploratory analysis and predictive modeling to obtain both predictive and analytical insights from multi-country food waste datasets.

Based on this research gap, the present study investigates food waste patterns using a publicly available dataset consisting of 5,000 records covering 20 countries and eight food categories from 2018 to 2024. The study integrates Exploratory Data Analysis (EDA) and Random Forest Regression to analyze food waste characteristics, evaluate predictive performance, and examine the contribution of key variables associated with food waste generation.

III. Material and Methods

The dataset used in this study was obtained from a publicly available Kaggle repository entitled *Global Food Wastage Dataset (2018–2024)* and was utilized through the associated Kaggle machine learning project *Global Food Waste Prediction – ML Project*. The dataset contains 5,000 observations collected across 20

countries and eight food categories covering the period from 2018 to 2024. The dataset includes Country, Year, Food Category, Total Waste (Tons), Economic Loss (Million \$), Average Waste per Capita (Kg), Population (Million), and Household Waste (%). No missing values were identified during the preprocessing stage.

The dataset is publicly accessible through [Kaggle](#). Since the original data collection and generation procedures were not fully documented by the dataset provider, the dataset was treated as a secondary data source for exploratory analysis and predictive modeling. Therefore, the findings of this study should be interpreted within the context and limitations of a publicly available secondary dataset.

The dataset combines observations from multiple countries to identify general food waste patterns at a global scale. While country-specific characteristics may differ, the objective of this study is not to model individual national food waste behavior but rather to develop a generalized predictive framework capable of capturing broad relationships among demographic, economic, and food-related variables across diverse geographical contexts. More detail about the characteristics of the dataset is shown in Table 1.

Table 1. Dataset Characteristics.

Attribute	Description
Source	Kaggle
Records	5000
Countries	20
Food Categories	8
Period	2018-2024
Numerical Variables	5
Categorical Variables	2
Missing Values	0

A. Data Pre-processing

This crucial stage ensures that the data quality and format are compatible with the requirements of the ML algorithm. **Data Loading and Library Import:** The initial step involves importing essential Python libraries were imported for specific functions, including Pandas and NumPy for robust data manipulation, Matplotlib and Seaborn for data visualization, and Scikit-learn for core Machine Learning functionalities. More detail about the category of the sample dataset is shown in Table 2.

Table 2. Sample of Initial Input Data Used for Food Waste Predictive Modeling.

Category	Year	Food Category	Total Waste (Tons)	Economic Loss (Million \$)	Avg Waste per Capita (Kg)	Population (Million)	Household Waste (%)
Australia	2019	Fruits & Vegetables	19268.63	18686.68	72.69	87.59	53.64
Indonesia	2019	Prepared Food	3916.97	4394.48	192.52	1153.99	30.61
Germany	2022	Dairy Products	9700.16	8909.16	166.94	1006.11	48.08
France	2023	Fruits & Vegetables	46299.69	40551.22	120.19	953.05	31.91
France	2023	Beverages	33096.57	6980.82	104.74	1105.47	36.06

Handling Categorical Variables and Data Splitting: A crucial step in preparing the data for the Random Forest Regressor, a model requiring exclusively numerical inputs, involved the transformation of nominal categorical features. Specifically, the 'Country' and 'Food Category' variables were converted using the One-Hot Encoding technique. This process, executed via the `pd.get_dummies` function, generates new binary indicator columns (0 or 1) for each unique category within the original features. This step ensures that the entire feature set (X) is represented in a numerical format, suitable for subsequent model training.

Feature Standardization: Following categorical encoding, numerical variables were standardized using the `StandardScaler` method to ensure that all features were represented on a comparable scale. Standardization was applied to Year, Economic Loss (Million \$), Population (Million), Average Waste per Capita (Kg), Household Waste (%), and Total Waste (Tons). This process transforms each variable to have a mean of zero and a standard deviation of one, thereby improving numerical consistency during model training and evaluation.

Feature-Target Separation and Data Validation Partition: Following the encoding of categorical variables, the dataset was formally structured by separating the independent variables (features, X) from the dependent variable (target, y), which is defined as 'Total Waste (Tons)'.

To objectively assess the model's predictive performance and generalization capability, the dataset was subsequently partitioned into distinct subsets: a training set and a testing set. A standard division ratio was employed, allocating 80% of the data for training and 20% for testing ($\text{test_size}=0.2$). This rigorous partitioning is essential to validate the model's efficacy on previously unseen data, thereby mitigating the risk of *overfitting* and providing a robust estimate of its real-world predictive accuracy.

B. Machine Learning Modeling

To effectively model and predict the target variable ('Total Waste (Tons)'), this study employed the Random Forest Regressor algorithm. This tree-based *ensemble* method was selected due to its inherent robustness in handling non-linear data structures and its demonstrated efficacy in mitigating the risk of *overfitting* often associated with single decision trees.

The prediction generated by the Random Forest Regressor can be expressed in Equation 1, where $\hat{y}$ represents the predicted value, $T_b(x)$ denotes the prediction produced by the b -th decision tree, B represents the total number of decision trees in the forest.

$$\hat{y}(x) = \frac{1}{B} \sum_{b=1}^B T_b(x) \quad (1)$$

This ensemble mechanism combines predictions from multiple decision trees to improve predictive accuracy and reduce variance compared with a single decision tree. Table 3 below shown the configuration of Random Forest Hyperparameter.

Table 3: Random Forest Hyperparameter Configuration	
Hyperparameter	Value
n_estimators	100
max_depth	10
Min_samples_split	2
random_state	42
Cross-validation folds	5

A 5-fold cross-validation procedure was employed during hyperparameter optimization using GridSearchCV. The hyperparameter search evaluated two candidate values for the maximum tree depth ($\text{max_depth} = \text{None}$ and $\text{max_depth} = 10$), while n_estimators and min_samples_split were fixed at 100 and 2, respectively. Based on the GridSearchCV results, the optimal hyperparameter configuration was $\text{n_estimators} = 100$, $\text{max_depth} = 10$, and $\text{min_samples_split} = 2$. The default value of min_samples_leaf provided by Scikit-learn was retained during model development.

The Random Forest Regressor was subsequently trained using the pre-processed training dataset (X_{train} and y_{train}). A fixed random seed ($\text{random_state} = 42$) was applied to ensure reproducibility of the experimental results. The optimized model obtained from the hyperparameter tuning process was then used to generate predictions (y_{pred}) on the unseen testing dataset (X_{test}).

Upon completion of the prediction process, the generated outputs were evaluated using the R^2 Score, Mean Absolute Error (MAE), and Root Mean Squared Error (RMSE) metrics to assess the model's predictive performance and generalization capability.

C. Evaluation Metrics

The performance of the proposed model was evaluated using three commonly used regression metrics, namely the Coefficient of Determination (R^2), Mean Absolute Error (MAE), and Root Mean Squared Error (RMSE).

1) R^2 Score (Coefficient of Determination)

The coefficient of determination measures the proportion of variance in the target variable that can be explained by the predictive model. It is calculated as shown in Equation 2, where y_i represents the actual value, $\hat{y}_i$ represents the predicted value, $\bar{y}$ represents the mean of actual values, and n represents the number of observations. A value closer to 1 indicates better predictive performance.

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (2)$$

2) Mean Absolute Error (MAE)

The Mean Absolute Error represents the average absolute difference between actual and predicted values, as shown in Equation 3, where y_i represents the actual value, $\hat{y}_i$ represents the predicted value, and n represents the number of observations. Lower MAE values indicate better prediction accuracy.

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad (3)$$

3) Root Mean Squared Error (RMSE)

The Root Mean Squared Error measures the square root of the average squared prediction error, as formulated in Equation 4, where y_i represents the actual value, $\hat{y}_i$ represents the predicted value, and n represents the number of observations. RMSE penalizes larger errors more heavily than MAE and is therefore sensitive to outliers.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (4)$$

IV. Results and Discussion

A. Exploratory Data Analysis (EDA) Results

1) Temporal Food Waste Trends.

The exploratory analysis indicates variations in total food waste across the observation period from 2018 to 2024. As illustrated in Figure 1, annual waste volumes fluctuate over time, suggesting that food waste generation is influenced by multiple factors, including consumption patterns, economic conditions, and food category characteristics. Understanding these temporal patterns is important for supporting future waste reduction strategies and predictive modeling efforts.

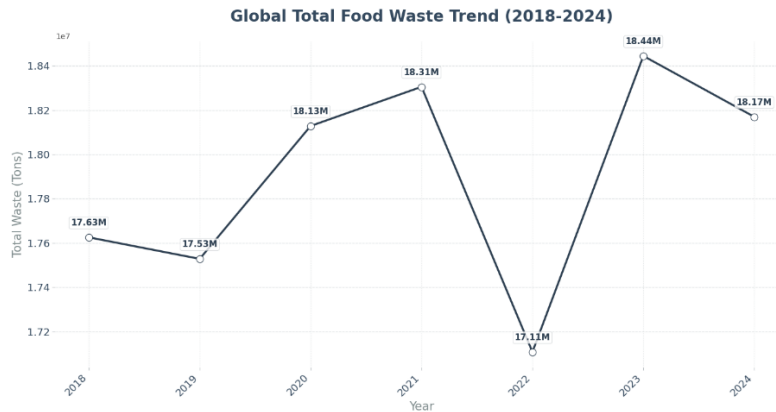


Figure 1. Global Total Food Waste Trend per Year (2018-2024)

2) Food Waste Distribution by Category

Figure 2 presents the distribution of food waste across different food categories. The results indicate that fruits and vegetables contribute a substantial proportion of total waste compared with several other categories. This finding is consistent with the highly perishable nature of fresh produce, which is more susceptible to spoilage during storage, transportation, and household consumption.

The observed differences among food categories suggest that food waste mitigation strategies should not be generalized across all food groups. Instead, category-specific interventions may be required, particularly for highly perishable products.

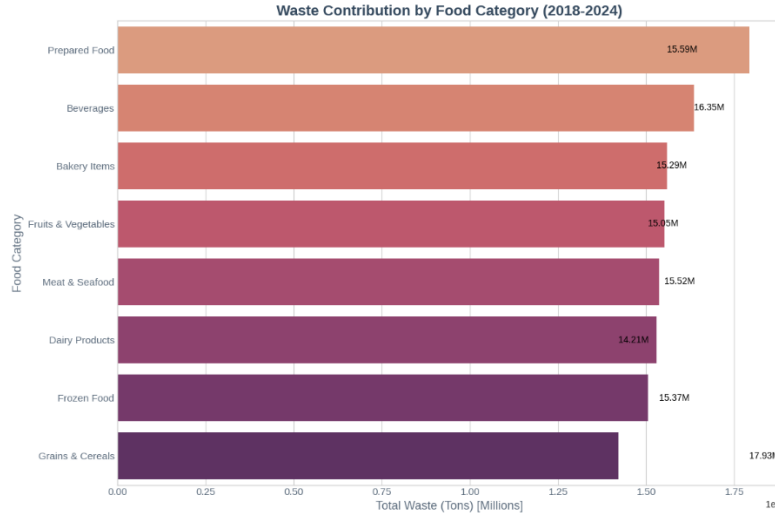


Figure 2. Food Waste Contribution by Food Category

3) Correlation Analysis

Figure 3 illustrates the correlation matrix among numerical variables included in the dataset. The analysis indicates that the relationships between Total Waste (Tons) and several explanatory variables are generally weak. In particular, Population (Million) exhibits only a weak linear association with Total Waste (Tons), suggesting that population size alone may not adequately explain food waste generation within the dataset. Similarly, Economic Loss (Million \$) demonstrates a relatively limited linear relationship with total waste levels.

These findings imply that food waste generation is likely influenced by multiple interacting factors rather than a single demographic or economic indicator. Consequently, the use of machine learning methods capable of capturing complex and non-linear relationships is justified for this prediction task.

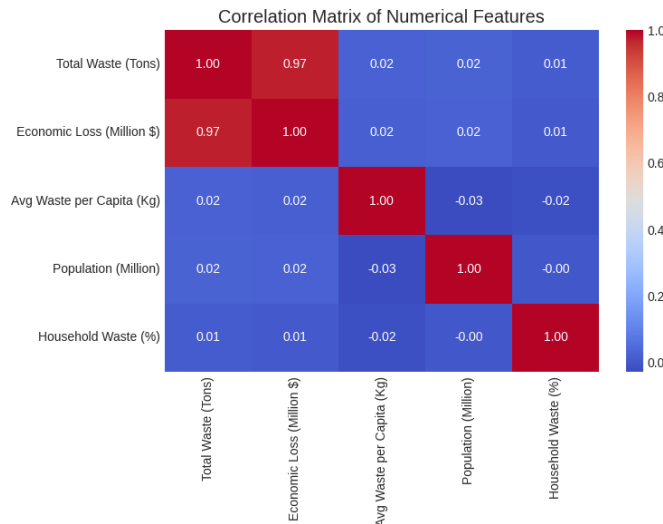


Figure 3. Heatmap of Correlation Among Numerical Variables

B. Random Forest Regression Results

The Random Forest Regressor model was evaluated using the testing dataset to assess its predictive performance. The results are presented in Table 4. Since the target variable was standardized during preprocessing using StandardScaler, the reported MAE and RMSE values are expressed in standardized units rather than the original tonnage scale. Therefore, these metrics should be interpreted as relative prediction errors within the normalized feature space.

Table 4. Random Forest Regression Performance

Evaluation Metric	Value
R ² Score	0.9582
Mean Absolute Error (MAE)	0.16
Root Mean Squared Error (RMSE)	0.21

The model achieved an R² score of 0.9582, indicating that approximately 95.82% of the variance in food waste generation can be explained by the selected predictor variables. This result demonstrates strong predictive capability and suggests that the Random Forest algorithm is effective in modeling complex relationships within the dataset.

The obtained MAE and RMSE values further indicate that prediction errors remain relatively limited compared with the overall range of food waste values contained in the dataset. These results suggest that the model is capable of generating stable predictions while maintaining a high level of explanatory power.

The strong performance observed in this study can be attributed to the ensemble nature of Random Forest, which combines multiple decision trees to reduce variance and improve generalization performance. By aggregating predictions from numerous trees, the model is able to capture complex interactions among demographic, economic, and categorical variables more effectively than a single decision tree model.

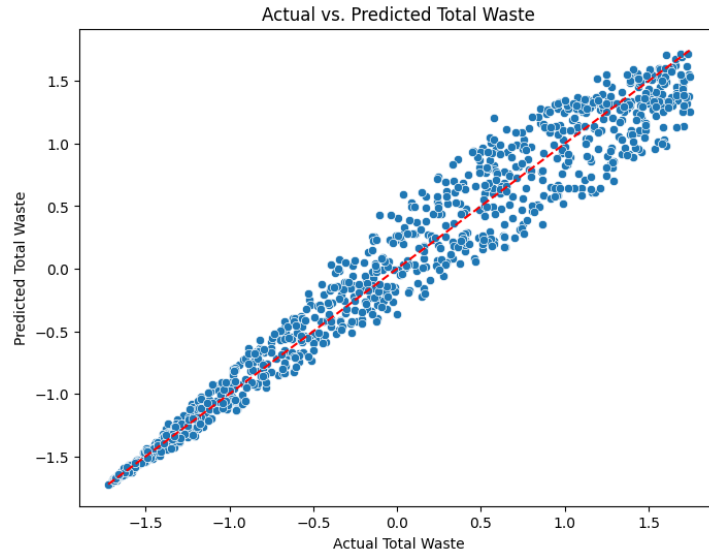


Figure 4. Actual versus Predicted Values Generated by the Random Forest Model

As shown in Figure 4, most observations are distributed close to the diagonal reference line, indicating a high degree of agreement between actual and predicted values. This pattern further confirms the model's ability to accurately estimate food waste levels across the dataset.

C. Feature Importance Analysis

To obtain additional analytical insights beyond predictive performance, feature importance analysis was conducted using the Random Forest model. Feature importance measures the relative contribution of each predictor variable to the model's prediction process. As shown in Figure 5, Economic Loss (Million \$) emerged as the most influential predictor of food waste generation, followed by Population (Million) and Household Waste (%). In contrast, most country-specific and food-category variables exhibited relatively lower importance scores.

These findings suggest that economic and demographic factors contribute more substantially to food waste estimation than individual country or category indicators within the dataset. The dominance of Economic Loss further indicates a close relationship between food waste generation and its associated economic consequences. Therefore, food waste mitigation strategies may benefit from simultaneously addressing both economic inefficiencies and household-level consumption behavior.

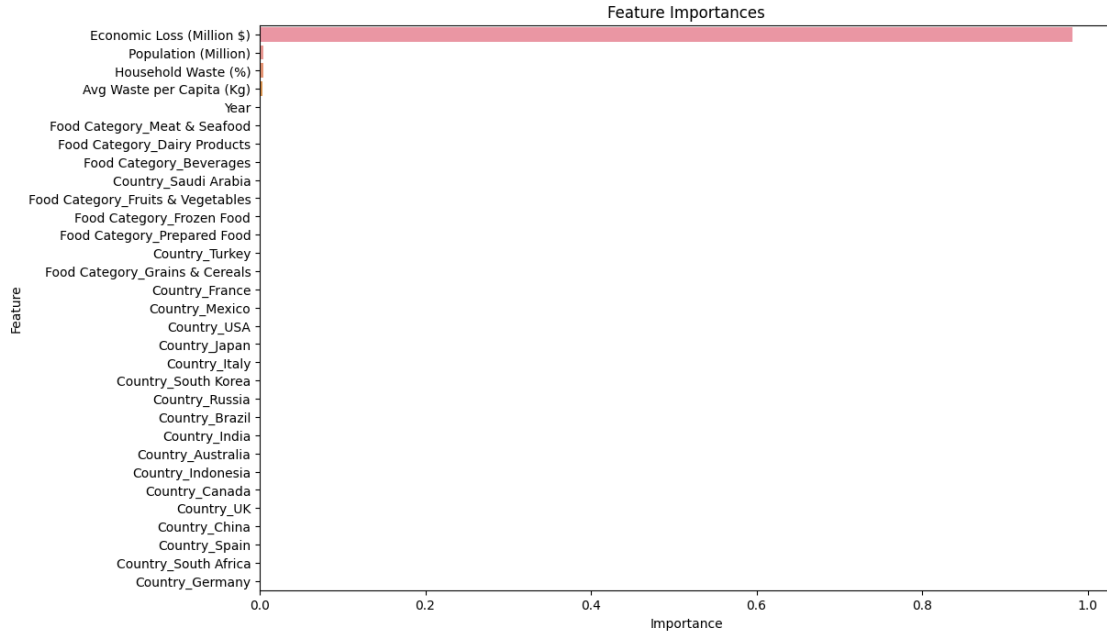


Figure 5. Feature Importance Generated by the Random Forest Model

D. Study Limitations

This study has several limitations that should be considered when interpreting the results.

First, the dataset was obtained from a publicly available Kaggle repository, and detailed information regarding the original data collection and generation procedures was not fully documented by the dataset provider. Consequently, the findings should be interpreted within the context of a secondary public dataset.

Second, the predictive model only considered variables available in the dataset, such as population, economic loss, food category, and household waste. Other potentially relevant factors, including government policies, climate conditions, logistics infrastructure, and socio-cultural behaviors, were not incorporated into the analysis.

Third, this study focused exclusively on the Random Forest Regression algorithm and did not perform comparative evaluations against alternative machine learning or deep learning approaches. Therefore, future research should conduct comparative experiments using multiple predictive models to identify the most suitable approach for global food waste prediction and to further assess model robustness, predictive performance, and interpretability.

Fourth, statistical significance testing, such as paired t-tests or Wilcoxon signed-rank tests, was not conducted because the study focused on the evaluation of a single predictive model rather than a comparative assessment of multiple algorithms. Future studies may incorporate statistical comparison methods when evaluating the performance of several machine learning models.

V. Conclusion

This study investigated food waste patterns using a publicly available dataset covering 5,000 records from 20 countries and eight food categories between 2018 and 2024. Through Exploratory Data Analysis (EDA), several important patterns were identified, including variations in food waste generation across food categories and years. The results indicate that fruits and vegetables contribute a substantial proportion of total food waste, highlighting the importance of targeted waste reduction strategies for highly perishable food products. The Random Forest Regression model demonstrated strong predictive capability, achieving an R^2 score of 0.9582. This result suggests that machine learning techniques can effectively model complex relationships among demographic, economic, and food-related variables for food waste prediction. In addition to predictive performance, feature importance analysis provided further insight into the variables that contribute to food waste estimation, supporting a more comprehensive understanding of food waste patterns. Nevertheless, this study has several limitations. The dataset was obtained from a secondary public repository, and detailed information regarding the original data collection process was not fully available. Furthermore, external factors such as climate conditions, government policies, logistics infrastructure, and

socio-cultural behaviors were not included in the predictive model. Future research may incorporate additional explanatory variables, compare multiple machine learning and deep learning approaches, and utilize more extensively documented datasets to improve both predictive performance and analytical interpretation of food waste generation.

Conflicts of Interest

The authors affirm that there are no actual or potential conflicts of interest that could inappropriately influence the interpretation of the results presented in this paper. This rigorous study, data analysis, and preparation of this manuscript were conducted independently without influence from external financial or personal interests.

Author Contributions Statement

Kemal Pramayuda is the sole author of this manuscript and is responsible for all aspects of the research. His responsibilities include, but are not limited to, conceptualization, methodology design, data preparation, exploratory data analysis (EDA), Machine Learning model development (utilizing the Random Forest Regressor), analysis of the findings (results), drafting the initial version of the manuscript, and final editorial work. The author has thoroughly reviewed and approved the final version of the manuscript for publication.

Acknowledgment

The author would like to express gratitude to the Information Technology Department, Universitas Sebelas April Sumedang, for providing the necessary academic environment and resources required to complete this research. Special appreciation is extended to the open-source community for making the Python libraries (*Pandas*, *Scikit-learn*, *Matplotlib*) and the Kaggle dataset publicly available, which were fundamental to this study.

References

- [1] Jenny. Gustavsson, *Global food losses and food waste : extent, causes and prevention : study conducted for the International Congress 'Save Food!' at Interpack 2011 Düsseldorf, Germany*. Food and Agriculture Organization of the United Nations, 2011.
- [2] Hamish Forbes, Tom Qusted, and Clementine O'Connor, *Food Waste Index Report 2021*. 2021.
- [3] S. Corrado *et al.*, 'Food waste accounting methodologies: Challenges, opportunities, and further advancements', *Glob. Food Sec.*, vol. 20, pp. 93–100, Mar. 2019, doi: 10.1016/j.gfs.2019.01.002.
- [4] X. Kong *et al.*, 'Deep learning for time series forecasting: a survey', *International Journal of Machine Learning and Cybernetics*, vol. 16, no. 7–8, pp. 5079–5112, Aug. 2025, doi: 10.1007/s13042-025-02560-w.
- [5] J.-S. Lee and D.-C. Shin, 'Prediction of Waste Generation Using Machine Learning: A Regional Study in Korea', *Urban Science*, vol. 9, no. 8, p. 297, Jul. 2025, doi: 10.3390/urbansci9080297.
- [6] G. F. Turker, 'Reducing Food Waste in Campus Dining: A Data-Driven Approach to Demand Prediction and Sustainability', *Sustainability*, vol. 17, no. 2, p. 379, Jan. 2025, doi: 10.3390/su17020379.
- [7] M. Rodrigues, V. Miguéis, S. Freitas, and T. Machado, 'Machine learning models for short-term demand forecasting in food catering services: A solution to reduce food waste', *J. Clean. Prod.*, vol. 435, p. 140265, Jan. 2024, doi: 10.1016/j.jclepro.2023.140265.
- [8] F. S. Farvin, S. A. Musthafa, and Ad. Carolina Rani, 'Time Series for Managing Food Waste Using Machine Learning', 2022.
- [9] D. Nijloveanu *et al.*, 'The Development of a Prediction Model Related to Food Loss and Waste in Consumer Segments of Agrifood Chain Using Machine Learning Methods', *Agriculture*, vol. 14, no. 10, p. 1837, Oct. 2024, doi: 10.3390/agriculture14101837.
- [10] Carlos A. Villanueva, 'Integration of Statistical and Machine Learning Models for Time Series Forecasting in Optimizing Decision Making for Smart Waste Management', *Journal of Information Systems Engineering and Management*, vol. 10, no. 51s, pp. 396–405, May 2025, doi: 10.52783/jisem.v10i51s.10399.
- [11] G. H. S. -, G. R. P. V. -, P. M. A. -, C. R. -, and M. A. L. -, 'Food Waste Management And Prediction Using Machine Learning : A Comprehensive Analysis', *International Journal on Science and Technology*, vol. 16, no. 2, Apr. 2025, doi: 10.71097/IJSAT.v16.i2.3275.
- [12] L. Breiman, 'Random Forests', *Mach. Learn.*, vol. 45, no. 1, pp. 5–32, Oct. 2001, doi: 10.1023/A:1010933404324.

- [13] A. Siddique, A. Gupta, J. T. Sawyer, T.-S. Huang, and A. Morey, 'Big data analytics in food industry: a state-of-the-art literature review', *NPJ Sci. Food*, vol. 9, no. 1, p. 36, Mar. 2025, doi: 10.1038/s41538-025-00394-y.
- [14] T. Weerasooriya and K. Kumar, 'Foodlignce – Predicting Expiry Date of Perishable Foods to Reduce Loss and Waste', *Computer Science & Engineering: An International Journal*, vol. 12, no. 6, pp. 155–164, Dec. 2022, doi: 10.5121/cseij.2022.12615.
- [15] G. Uganya, D. Rajalakshmi, Y. Teekaraman, R. Kuppusamy, and A. Radhakrishnan, 'A Novel Strategy for Waste Prediction Using Machine Learning Algorithm with IoT Based Intelligent Waste Management System', *Wirel. Commun. Mob. Comput.*, vol. 2022, no. 1, Jan. 2022, doi: 10.1155/2022/2063372.