

Explainable Machine Learning for Food Calorie Prediction Using Tree-Based Ensemble Models and SHAP Analysis

Wilsen Grivin Mokodaser^{*1}, Tonny Irianto Soewignyo², Argha Orion Silitonga³, Regi Fernando Najohan⁴

Faculty of Computer Science, Universitas Klabat, Minahasa Utara

E-mail : wilsenm@unklab.ac.id^{*1}, tonnysoewignyo@unklab.ac.id², argha@unklab.ac.id³, reginajohan@unklab.ac.id⁴

^{*}Corresponding author

Received 25 June 2026; Revised 26 July 2026; Accepted 29 July 2026

Abstract - Accurate food calorie prediction is essential for nutritional assessment, dietary planning, and intelligent health applications. However, achieving high predictive accuracy while maintaining model interpretability remains a challenge for many machine learning approaches. This study proposes an explainable machine learning framework for food calorie prediction using tree-based ensemble models and SHAP (SHapley Additive exPlanations). A dataset containing 1,346 food samples with three nutritional attributes- protein, fat, and carbohydrate-was used. Data preprocessing included logarithmic transformation and an 80:20 train-test split. Three machine learning models, namely Linear Regression, Random Forest, and XGBoost, were developed and evaluated using Mean Absolute Error (MAE), Mean Squared Error (MSE), Root Mean Squared Error (RMSE), and the coefficient of determination (R^2). Hyperparameter optimization for XGBoost was performed using RandomizedSearchCV with five-fold cross-validation. Experimental results showed that Random Forest achieved the best predictive performance with an MAE of 0.1149, RMSE of 0.2869, and an R^2 score of 0.9008, outperforming both XGBoost and Linear Regression. Cross-validation further demonstrated the robustness of the selected model, yielding a mean R^2 of 0.9316 with a standard deviation of 0.0239. Residual analysis indicated prediction errors centered near zero without noticeable systematic bias. SHAP analysis provided both global and local model interpretability, identifying carbohydrate as the most influential feature, followed by fat and protein. The findings show that integrating tree-based ensemble learning with SHAP enables accurate, transparent calorie prediction, making the proposed approach suitable for nutritional decision support and explainable artificial intelligence applications.

Keywords - Food calorie prediction, Explainable artificial intelligence, SHAP, Random Forest, XGBoost, Machine learning, Nutritional analysis.

1. INTRODUCTION

Food calorie prediction has become increasingly important in nutritional assessment, dietary planning, and intelligent healthcare applications. Accurate estimation of calorie content enables healthcare professionals and individuals to evaluate dietary intake more efficiently while reducing reliance on manual calculations and nutritional tables. Since the relationship between macronutrients and calorie content is often nonlinear, machine learning techniques have recently attracted considerable attention because of their capability to learn complex patterns from nutritional data and provide more accurate predictions than conventional statistical approaches.

Recent advances in machine learning have introduced various predictive models for calorie estimation, including linear models, ensemble learning methods, and gradient boosting algorithms. Most existing studies primarily emphasize improving prediction accuracy by optimizing model performance. However, relatively limited attention has been devoted to explaining how individual nutritional components influence prediction outcomes. As a result,

many predictive models remain difficult to interpret, reducing user confidence and limiting their practical adoption in nutrition-related decision support systems where transparency is an important consideration.

Accurate determination of daily energy requirements is a fundamental aspect of supporting physical and academic activities, particularly among specific groups such as sports nutrition students who require energy intake ranging from 2480.7 to 3551.3 kcal/day [1]. In estimating these requirements, predictive equations such as the Harris–Benedict equation are still considered to provide acceptable accuracy, reaching approximately 62% precision for individuals with overweight or obesity conditions [2]. Understanding energy expenditure is crucial because human metabolic patterns continuously change throughout life; energy expenditure increases rapidly during infancy, remains relatively stable between the ages of 20 and 60, and gradually declines afterward [3]. However, challenges arise because the current Estimated Energy Requirements (EER) standards are considered insufficient in representing global ethnic and geographical diversity [4]. The inaccuracy of these standards contributes to the phenomenon of worldwide energy imbalance, where average population energy intake exceeds healthy requirements by approximately 80 kcal/day [5]. Technically, total daily energy expenditure is influenced by three major components, namely resting energy expenditure, the thermic effect of food, and physical activity [6]. Determining appropriate dietary reference values is further complicated by inconsistencies that frequently occur in estimation methods [7].

As a solution, the evaluation of energy adequacy has increasingly shifted toward monitoring body weight stability rather than relying solely on self-reported dietary intake [8]. Although predictive equations remain widely used in daily clinical practice, experts still recommend indirect calorimetry as a more reliable method whenever available [9]. In addition, biological factors such as pregnancy must also be carefully considered due to significant increases in energy expenditure variability during this period [10]. The implementation of the eXtreme Gradient Boosting (XGBoost) algorithm has become an important breakthrough in the development of transparent artificial intelligence systems, particularly in the medical field, where the model can serve as a communication bridge for clinicians through interpretable decision-tree structures [11]. The superior predictive performance of XGBoost compared to other supervised learning methods is further strengthened by the integration of SHapley Additive exPlanations (SHAP), which transforms black-box models into human-understandable systems by accurately revealing feature interdependencies and feature importance rankings [12]. Nevertheless, researchers must remain aware of potential interpretability bias, as the incremental addition of decision trees in XGBoost may inflate feature importance scores, while SHAP values themselves remain highly dependent on model structure and feature interactions [13].

Based on these considerations, a research gap exists because previous studies have predominantly focused on improving predictive accuracy, while comparatively little attention has been given to model interpretability in food calorie prediction. Although explainable artificial intelligence techniques have been increasingly adopted in various machine learning applications, their implementation for interpreting calorie prediction models remains limited. Therefore, this study proposes an explainable machine learning approach for food calorie prediction by comparing multiple predictive models and employing SHAP to interpret feature contributions. The proposed approach aims not only to identify an accurate predictive model but also to provide transparent explanations regarding how macronutrient attributes contribute to calorie prediction, thereby supporting more reliable decision-making in healthcare and nutritional analysis.

2. RESEARCH METHOD

In this study, several stages will be conducted as a guide to achieve better results. The stages used can be seen in the image below:

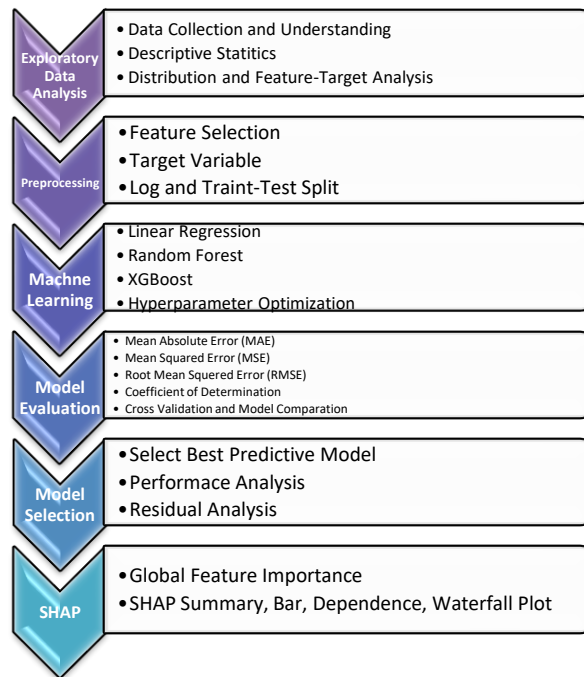


Figure 1. Research Workflow Stages

In Figure 1, it can be seen that this study will be conducted in 4 scientific steps to produce an effective comparison of each algorithm that will be used.

2.1 Exploratory Data Analysis (EDA)

The dataset used in this study consists of 1,345 food records containing nutritional information commonly consumed by the Indonesian population. The dataset includes several attributes, namely ID, calories, protein, fat, carbohydrate, food name, and image, where calories is defined as the target variable. Exploratory Data Analysis (EDA) is conducted to understand the initial characteristics of the dataset before model development. The analysis includes dataset structure inspection using the `info()` and `describe()` functions, visualization of feature distributions through histograms, correlation analysis using a heatmap, and exploration of feature-target relationships. This stage aims to identify data distribution patterns, detect potential anomalies, and obtain a comprehensive understanding of the relationships among nutritional attributes prior to the machine learning process.

2.2 Data Preprocessing

Data preprocessing is performed to prepare the dataset for machine learning modeling. Numerical attributes consisting of proteins, fat, and carbohydrate are selected as predictor variables, while calories is defined as the target variable. Non-predictive attributes, including ID, food name, and image, are excluded from the modeling process to improve computational efficiency and eliminate irrelevant information. Since the nutritional variables exhibit right-skewed distributions, a logarithmic transformation is applied using the `log1p` function to reduce skewness and stabilize feature distributions before training the models. Subsequently, the dataset is split into training data (80%) and testing data (20%) using the train-test split technique [14][15]. This step ensures that the data is properly prepared for the model training process.

2.3 Machine Learning Model Development

To evaluate the effectiveness of different predictive approaches, this study develops three regression models, namely Linear Regression, Random Forest, and eXtreme Gradient Boosting (XGBoost). Linear Regression is employed as a conventional baseline model, whereas Random Forest and XGBoost represent tree-based ensemble learning algorithms capable of capturing

nonlinear relationships within structured nutritional data [16]. To improve the predictive performance of XGBoost, hyperparameter optimization is performed using RandomizedSearchCV, where the optimal values of `n_estimators`, `learning_rate`, `max_depth`, `subsample`, `colsample_bytree`, `gamma`, and `min_child_weight` are determined automatically [17]. All models are trained using the same training dataset to ensure a fair performance comparison.

2.4 Model Performance Evaluation

The predictive performance of each regression model is evaluated using multiple quantitative metrics, including Mean Absolute Error (MAE), Mean Squared Error (MSE), Root Mean Squared Error (RMSE), and the Coefficient of Determination (R^2). MAE measures the average prediction error, MSE penalizes larger prediction errors more heavily [18], RMSE provides error values in the original measurement scale, and R^2 evaluates the proportion of variance explained by each model [19]. In addition, five-fold cross-validation is conducted to assess model stability and generalization capability. The evaluation is further supported by visualization techniques, including Actual versus Predicted plots and Residual Distribution plots, to analyze prediction accuracy and residual behavior.

2.5 Best Model Selection

Following model evaluation, the regression models are compared based on their predictive performance using the evaluation metrics and cross-validation results. The model demonstrating the best balance between predictive accuracy and generalization capability is selected as the representative explainable machine learning model for subsequent interpretation. This selection process ensures that the explainability analysis is performed on a model with reliable predictive performance rather than solely on a predefined algorithm.

2.6 Explainable AI (SHAP)

The final stage focuses on interpreting the selected prediction model using SHapley Additive exPlanations (SHAP). Explainability is performed through both global and local interpretation. Global interpretation is conducted using the SHAP Summary Plot and SHAP Bar Plot to quantify the overall contribution and importance ranking of each nutritional feature [20]. Local interpretation is performed using the SHAP Waterfall Plot to explain how individual features contribute to a specific prediction [21]. Furthermore, the SHAP Dependence Plot is employed to analyze the relationship between feature values and their corresponding SHAP values, providing additional insight into how nutritional attributes influence calorie prediction.

3. RESULTS AND DISCUSSION

3.1 Exploratory Data Analysis (EDA)

Before developing the machine learning models, an Exploratory Data Analysis (EDA) was conducted to examine the characteristics of the nutritional dataset and assess its suitability for predictive modeling. This stage provides an initial understanding of the relationships among the predictor variables and the target variable, identifies the distribution patterns of the nutritional attributes, and evaluates the overall quality of the data. The findings obtained from this analysis serve as the foundation for the subsequent preprocessing, model development, and explainable machine learning stages.

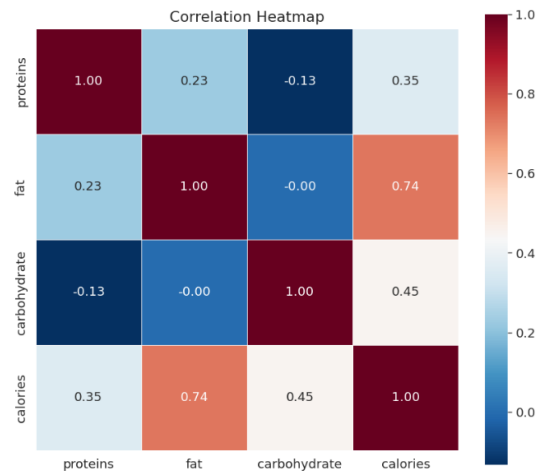


Figure 1. Pearson Correlation Matrix of Nutritional Features and Calories

Figure 1 presents the Pearson correlation matrix between the nutritional attributes and the target variable. The results indicate that fat has the strongest positive correlation with calories ($r = 0.74$), followed by carbohydrate ($r = 0.45$) and proteins ($r = 0.35$). In contrast, the correlations among the predictor variables remain relatively low, indicating that each feature provides complementary information for calorie prediction. These findings support the inclusion of all nutritional attributes in the subsequent machine learning modeling process

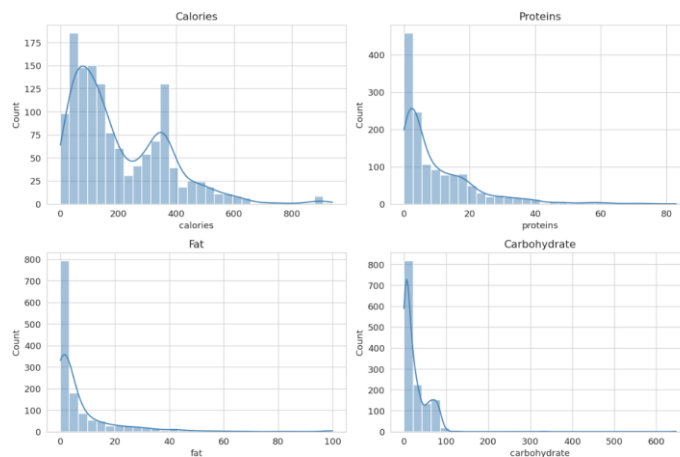


Figure 2. Distribution of Calories and Nutritional Features.

Figure 2 illustrates the distribution of the target variable (calories) and the three predictor variables (proteins, fat, and carbohydrate). Overall, the nutritional attributes exhibit positively skewed distributions, with most observations concentrated at lower values and a smaller number of observations extending toward higher values. A similar pattern is observed for calorie values, indicating that foods with lower to moderate calorie content are more common in the dataset than high-calorie foods. These distribution characteristics suggest that the data are not normally distributed and provide justification for applying data transformation during the preprocessing stage to improve model learning.

3.2 Data Preprocessing

During the preprocessing stage, a logarithmic transformation ($\log_{10}$) was applied to the numerical variables, including the predictor attributes and the target variable, to reduce the positive skewness observed during the exploratory data analysis. This transformation helps

stabilize the data distribution while preserving zero values, thereby improving the suitability of the data for machine learning algorithms. After the transformation, the dataset was partitioned into 1,076 training samples (80%) and 270 testing samples (20%). The training set was used to develop the prediction models, whereas the testing set was reserved for evaluating model performance on previously unseen data, ensuring an unbiased assessment of predictive capability.

3.3 Machine Learning Models

Table 1. Linear Regression (Baseline)

Evaluation Metric	Result
MAE	0.38510022724545445
MSE	0.25115368565235036
RMSE	0.5011523577240262
R2	0.6972296744653366

Table 1 presents the performance of the Linear Regression model, which serves as the baseline approach in this study. The model achieved a Mean Absolute Error (MAE) of 0.3851, a Mean Squared Error (MSE) of 0.2512, and a Root Mean Squared Error (RMSE) of 0.5012. In addition, the model obtained an R^2 score of 0.6972, indicating that approximately 69.7% of the variation in calorie values can be explained by the linear relationship between the nutritional features and the target variable. Although these results demonstrate that Linear Regression is capable of capturing the overall trend in the data, the prediction accuracy remains limited, suggesting that the relationship between macronutrients and calorie values is not entirely linear. Therefore, more advanced tree-based ensemble models are further evaluated to determine whether nonlinear learning algorithms can provide improved predictive performance.

Table 2. Random Forest Performance Metrics

Evaluation Metric	Result
MAE	0.11491407602296133
MSE	0.08232820597665451
RMSE	0.2869289214712496
R2	0.9007518537524467

Table 2 summarizes the predictive performance of the Random Forest model. The model achieved a Mean Absolute Error (MAE) of 0.1149, Mean Squared Error (MSE) of 0.0823, and Root Mean Squared Error (RMSE) of 0.2869, indicating substantially lower prediction errors than the baseline Linear Regression model. Furthermore, the Random Forest model obtained an R^2 score of 0.9008, demonstrating that approximately 90.1% of the variation in calorie values can be explained by the model. This improvement suggests that Random Forest is more effective in capturing the underlying relationship between nutritional attributes and calorie values. The superior performance is likely due to its ensemble learning mechanism, which combines multiple decision trees to model complex and nonlinear patterns that cannot be adequately represented by a simple linear model.

Table 3. Optimal Hyperparameter Configuration of the XGBoost Model

Parameter	Value
N_estimators	300
Max_depth	5
Learning_rate	0.1
Subsample	0.9
Consample_bytree	1.0
Gamma	0.1

Min_child_weight	3
Best Cross - Validation Score	0.9336

Table 3 presents the optimal hyperparameter configuration obtained through RandomizedSearchCV, where 30 candidate parameter combinations were evaluated using five-fold cross-validation, resulting in a total of 150 model training iterations. The optimization process selected an XGBoost model with 300 estimators, a maximum tree depth of 5, a learning rate of 0.1, subsample of 0.9, colsample_bytree of 1.0, gamma of 0.1, and min_child_weight of 3. This configuration achieved the highest cross-validation score ($R^2 = 0.9336$), indicating that the selected hyperparameters provide strong predictive capability while maintaining good generalization across different validation folds.

Table 4. Performance Evaluation of the Optimized XGBoost Model

Evaluation Metric	Result
MAE	0.1257
MSE	0.0836
RMSE	0.2891
R2	0.8992

Table 4 presents the predictive performance of the optimized XGBoost model using the testing dataset. The model achieved a Mean Absolute Error (MAE) of 0.1257, Mean Squared Error (MSE) of 0.0836, and Root Mean Squared Error (RMSE) of 0.2891, indicating a low level of prediction error. Furthermore, the model obtained an R^2 score of 0.8992, demonstrating that approximately 89.9% of the variation in calorie values can be explained by the nutritional attributes included in the model. Compared with the baseline Linear Regression model, XGBoost provides a substantial improvement in predictive performance, highlighting its ability to capture nonlinear relationships between macronutrients and calorie content. Although its predictive performance is comparable to that of Random Forest, a comprehensive comparison among all evaluated models is presented in the subsequent section.

3.4 Model Selection

Table 5. Comparative Performance of the Evaluated Machine Learning Models

Model	MAE	MSE	RMSE	R2 Score
MAE	0.114914	0.082328	0.286929	0.900752
MSE	0.125701	0.083584	0.289109	0.899238
R2	0.385100	0.251154	0.501152	0.697230

Table 5 compares the predictive performance of the three evaluated machine learning models. Among them, Random Forest achieved the best overall performance, producing the lowest prediction errors (MAE = 0.1149, MSE = 0.0823, and RMSE = 0.2869) together with the highest coefficient of determination ($R^2 = 0.9008$). The optimized XGBoost model demonstrated comparable performance, with only a slight reduction in predictive accuracy ($R^2 = 0.8992$), while Linear Regression produced substantially higher prediction errors and a lower R^2 score (0.6972). These results indicate that tree-based ensemble models are more effective than the linear model in capturing the relationship between nutritional attributes and calorie values. Based on the overall evaluation, Random Forest was selected as the best-performing model and was subsequently used for the explainable machine learning analysis.

Table 6. Five-Fold Cross-Validation Results of the Selected Random Forest Model

Fold	R ² Score
1	0.901651
2	0.926389
3	0.966789
4	0.950189
5	0.912881
Mean R ²	0.931580
Standard Deviation	0.023909

Table 6 presents the five-fold cross-validation results for the selected Random Forest model. The R² scores ranged from 0.9017 to 0.9668, with a mean R² of 0.9316 and a standard deviation of 0.0239. The consistently high R² values across all validation folds, together with the low standard deviation, indicate that the model maintains stable predictive performance and generalizes well to different data partitions. These results further support the selection of Random Forest as the final model for the subsequent SHAP-based explainable machine learning analysis.

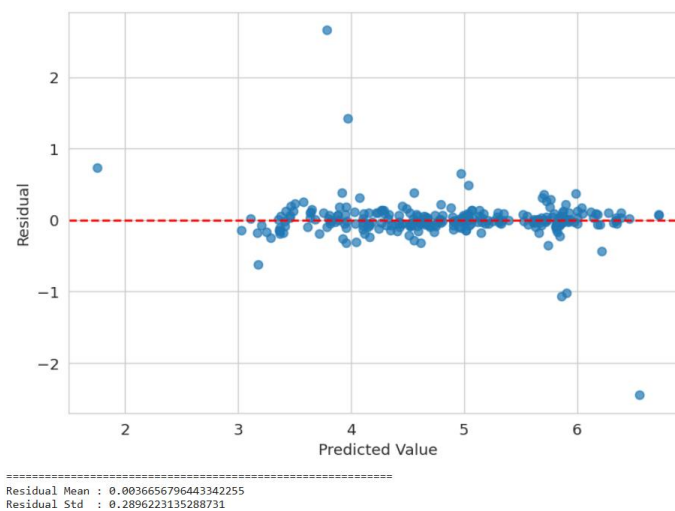


Figure 3. Residual Distribution of the Random Forest Predictions

Figure 3 illustrates the residual distribution of the selected Random Forest model by plotting the prediction errors against the predicted calorie values. Most residuals are randomly scattered around the zero reference line without forming a clear systematic pattern, indicating that the model does not exhibit noticeable bias across the prediction range. The mean residual of 0.0037, which is very close to zero, further suggests that the model neither consistently overestimates nor underestimates calorie values. In addition, the residual standard deviation of 0.2896 indicates that the prediction errors remain relatively small and concentrated around the zero line. Although a few observations exhibit relatively large positive or negative residuals, they appear to be isolated cases rather than evidence of systematic model errors. Overall, the residual analysis supports the reliability of the selected Random Forest model for subsequent SHAP-based interpretation.

3.5 Feature Importance Analysis

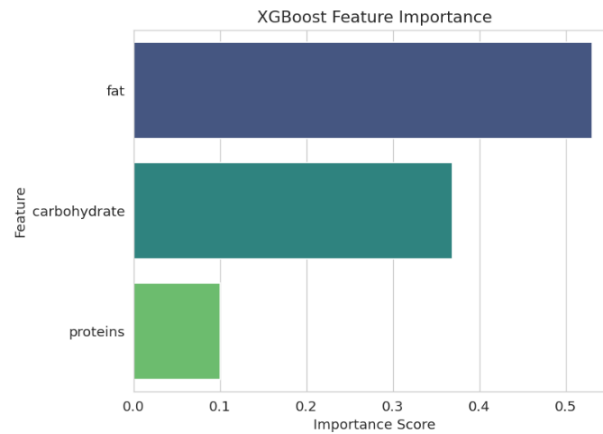


Figure 4. Global Feature Importance Derived from the XGBoost Model

Figure 4 presents the feature importance scores generated by the XGBoost model. Among the three nutritional attributes, fat has the highest importance score (approximately 0.53), indicating that it contributes the most to calorie prediction. Carbohydrate is the second most influential feature with an importance score of approximately 0.37, while proteins exhibit the lowest contribution, with a score of approximately 0.10. These results are consistent with the nutritional composition of foods, where fat generally provides more energy per gram than carbohydrates and proteins. The feature importance analysis offers an initial global interpretation of the model by identifying which variables contribute most to the prediction process. However, it only reflects the relative importance of each feature and does not explain how individual feature values influence specific predictions. Therefore, SHAP analysis is subsequently employed to provide a more comprehensive interpretation at both the global and local levels.

3.6 Explainable AI (SHAP)

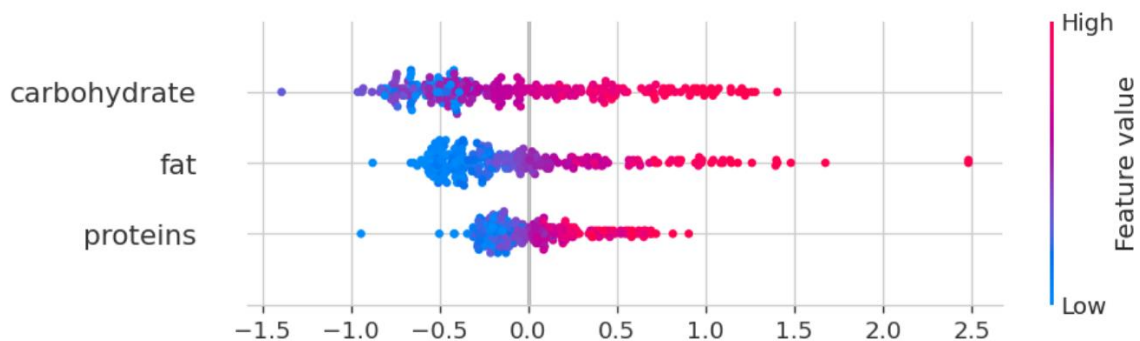


Figure 5. SHAP Summary Plot of Feature Contributions in the Selected Random Forest Model

Figure 5 presents the SHAP summary plot, which provides a global interpretation of how each nutritional feature contributes to calorie prediction. Each point represents an individual food sample, where the horizontal position indicates the SHAP value (feature contribution), and the color represents the original feature value, ranging from low (blue) to high (red). Features are ordered according to their overall impact on the model predictions.

The results indicate that carbohydrate has the greatest overall influence on the model predictions, followed by fat and proteins. For all three features, higher feature values (red points) are generally associated with positive SHAP values, indicating that foods with higher carbohydrate, fat, or protein content tend to increase the predicted calorie value. Conversely,

lower feature values (blue points) are concentrated on the negative side of the SHAP axis, suggesting that lower nutrient levels contribute to lower predicted calorie values. In addition, the wider spread of SHAP values for carbohydrate and fat indicates greater variability in their contributions across different food samples compared with proteins.

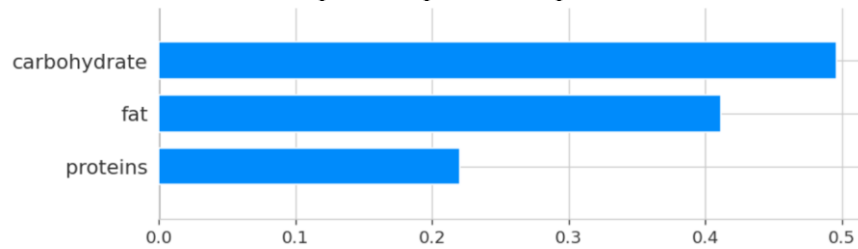


Figure 6. Global Feature Importance Based on Mean Absolute SHAP Values

Figure 6 presents the global feature importance based on the mean absolute SHAP values, which quantify the average magnitude of each feature's contribution to the model predictions. Unlike conventional feature importance, this approach measures how much each feature influences the prediction regardless of whether its effect is positive or negative. The results show that carbohydrate has the highest mean absolute SHAP value (0.4962), indicating that it is the most influential feature in predicting calorie content. Fat ranks second with a mean absolute SHAP value of 0.4112, while proteins have the lowest overall contribution (0.2203). These findings suggest that variations in carbohydrate content produce the largest average impact on calorie predictions, followed by fat, whereas protein contributes comparatively less to the model's decision-making process. The consistency between the SHAP summary plot and the mean absolute SHAP ranking further confirms that carbohydrate and fat are the dominant nutritional factors influencing the predictive model.

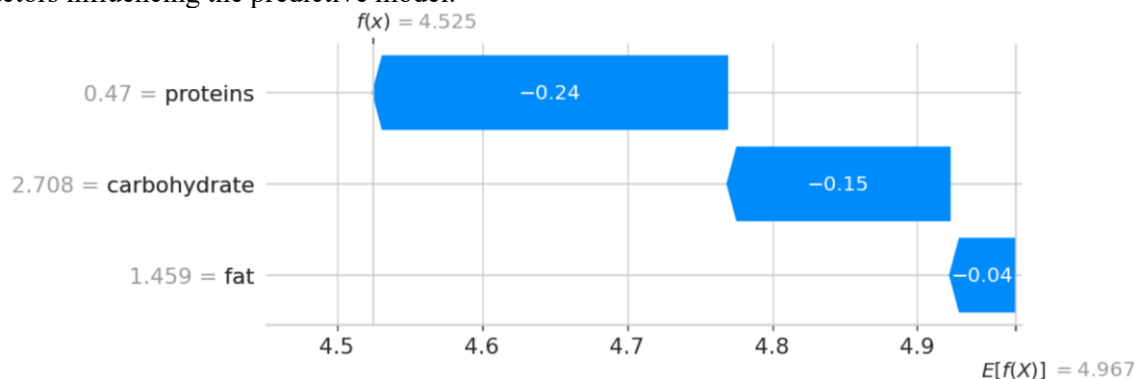


Figure 7. SHAP Waterfall Plot for an Individual Prediction

Figure 7 presents a SHAP waterfall plot that explains the prediction for an individual food sample. The model starts from the baseline prediction ($E[f(X)] = 4.967$), which represents the average prediction across all samples in the dataset. Each feature then contributes either positively or negatively to the final prediction according to its SHAP value. In this example, proteins have the largest negative contribution (-0.24), followed by carbohydrate (-0.15) and fat (-0.04). The combined effect of these feature contributions reduces the prediction from the baseline value to a final predicted value of $f(x) = 4.525$. This local explanation demonstrates how the SHAP framework decomposes an individual prediction into feature-level contributions, providing a transparent explanation of the model's decision for a specific food sample.

4. CONCLUSION

This study evaluated three machine learning algorithms, namely Linear Regression, Random Forest, and XGBoost, for food calorie prediction using nutritional attributes consisting

of proteins, fat, and carbohydrate. The experimental results demonstrated that tree-based ensemble methods significantly outperformed the baseline Linear Regression model. Among the evaluated models, Random Forest achieved the best predictive performance, with an MAE of 0.1149, RMSE of 0.2869, and an R^2 score of 0.9008, indicating its ability to explain approximately 90% of the variation in calorie values. The robustness of the selected model was further confirmed through five-fold cross-validation, which produced a mean R^2 score of 0.9316 with a low standard deviation (0.0239), while residual analysis showed prediction errors centered near zero, suggesting that the model does not exhibit significant systematic bias.

To improve model transparency, SHAP (SHapley Additive exPlanations) was employed to interpret both global and local model behavior. The SHAP analysis consistently identified carbohydrate as the most influential feature, followed by fat and proteins, indicating that carbohydrate contributes the largest average effect on calorie prediction across the dataset. Furthermore, the SHAP waterfall plot demonstrated how individual nutritional attributes contribute to specific predictions, providing an interpretable explanation of the model's decision-making process.

Overall, the results indicate that combining a high-performing tree-based ensemble model with SHAP-based explainability provides an accurate and transparent approach for food calorie prediction. This framework not only improves predictive performance but also enhances the interpretability of machine learning models, thereby increasing their potential for practical applications in nutritional assessment and decision support systems.

Future work may involve expanding the dataset with additional nutritional attributes, such as dietary fiber, sugar, sodium, and saturated fat, to further improve prediction accuracy. In addition, evaluating deep learning models and other explainable AI techniques, such as LIME or Integrated Gradients, could provide further insights into model interpretability and predictive performance across more diverse food datasets.

REFERENCES

- [1] Y. Tarigan and J. J. Tampi, "Analysis of daily calorie requirements," *Klasikal Journal*, 2025.
- [2] M. L. Macena, D. T. C. Paula, A. E. da Silva Junior, et al., "Estimates of resting energy expenditure and total energy expenditure using predictive equations in adults with overweight and obesity: a systematic review with meta-analysis," *Nutrition Reviews*, vol. 80, 2022.
- [3] H. Pontzer, Y. Yamada, H. Sagayama, P. N. Ainslie, et al., "Daily energy expenditure through the human life course," *Science*, vol. 373, no. 6556, 2021.
- [4] L. M. Neufeld and C. U. Loechl, "Dietary energy in balance: time for an update of energy requirements," *The Lancet*, 2025.
- [5] M. Springmann, "Estimates of energy intake, requirements and imbalances based on anthropometric measurements at global, regional and national levels and for sociodemographic groups," *BMJ Public Health*, 2025.
- [6] R. Fernández-Verdejo, et al., "Energy expenditure in humans: principles, methods, and changes throughout the life course," *Annual Review of Nutrition*, vol. 44, 2024.
- [7] L. Cloetens and L. Ellegård, "Energy—a scoping review for the Nordic Nutrition Recommendations 2023 project," *Food & Nutrition Research*, 2023.
- [8] National Academies of Sciences, Engineering, and Medicine, "Applications of the Dietary Reference Intakes for Energy," in *Dietary Reference Intakes for Energy*, Washington, DC: The National Academies Press, 2023.
- [9] I. Bendavid, D. N. Lobo, R. Barazzoni, T. Cederholm, et al., "The centenary of the Harris–Benedict equations: How to assess energy requirements best? Recommendations from the ESPEN expert group," *Clinical Nutrition*, vol. 40, no. 3, 2021.

- [10] National Academies of Sciences, Engineering, and Medicine, "Factors affecting energy expenditure and requirements," in *Dietary Reference Intakes for Energy*, Washington, DC: The National Academies Press, 2023.
- [11] J. Price, T. Yamazaki, K. Fujihara, et al., "XGBoost: interpretable machine learning approach in medicine," in *2022 5th World Conference on Computing and Communication Technologies (WCCCT)*, IEEE, 2022.
- [12] A. Homafar, H. Nasiri, and S. C. Chelgani, "Modeling coking coal indexes by SHAP-XGBoost: Explainable artificial intelligence method," *Fuel Communications*, vol. 11, 2022.
- [13] Y. Takefuji, "Beyond XGBoost and SHAP: unveiling true feature importance," *Journal of Hazardous Materials*, vol. 481, 2025.
- [14] Q. H. Nguyen *et al.*, "Influence of data splitting on performance of machine learning models in prediction of shear strength of soil," *Math. Probl. Eng.*, vol. 2021, 2021, doi: 10.1155/2021/4832864.
- [15] B. Vrigazova, "The Proportion for Splitting Data into Training and Test Set for the Bootstrap in Classification Problems," *Bus. Syst. Res.*, vol. 12, no. 1, pp. 228–242, 2021, doi: 10.2478/bsrj-2021-0015.
- [16] E. Bartz, T. Bartz-Beielstein, M. Zaefferer, and O. Mersmann, *Hyperparameter Tuning for Machine and Deep Learning with R: A Practical Guide*. 2023. doi: 10.1007/978-981-19-5170-1.
- [17] V. Verma, "Exploring Key XGBoost Hyperparameters: A Study on Optimal Search Spaces and Practical Recommendations for Regression and Classification," *Int. J. All Res. Educ. Sci. Methods*, vol. 12, no. 10, pp. 3259–3266, 2024, doi: 10.56025/ijaresm.2024.1210243259.
- [18] D. Chicco, M. J. Warrens, and G. Jurman, "The coefficient of determination R-squared is more informative than SMAPE, MAE, MAPE, MSE and RMSE in regression analysis evaluation," *PeerJ Comput. Sci.*, vol. 7, pp. 1–24, 2021, doi: 10.7717/PEERJ-CS.623.
- [19] J. Gao, "R-Squared (R²) – How much variation is explained?," *Res. Methods Med. Heal. Sci.*, vol. 5, no. 4, pp. 104–109, 2024, doi: 10.1177/26320843231186398.
- [20] M. Anjum, K. Khan, W. Ahmad, A. Ahmad, M. N. Amin, and A. Nafees, "New SHapley Additive ExPlanations (SHAP) Approach to Evaluate the Raw Materials Interactions of Steel-Fiber-Reinforced Concrete," *Materials (Basel)*, vol. 15, no. 18, 2022, doi: 10.3390/ma15186261.
- [21] Y. Gebreyesus, D. Dalton, D. De Chiara, M. Chinnici, and A. Chinnici, "AI for Automating Data Center Operations: Model Explainability in the Data Centre Context Using Shapley Additive Explanations (SHAP)," *Electron.*, vol. 13, no. 9, 2024, doi: 10.3390/electronics13091628.