

Analysis of Food Security Index Predictions in Indonesia Using Machine Learning Approach

Frederic Morado Saragih[♥], Wahyu Catur Wibowo

Faculty of Computer Science, Universitas Indonesia, Depok, Indonesia

[♥]Corresponding author email: frederic.morado@ui.ac.id

Article history: submitted: March 14, 2025; accepted: June 1, 2025; available online: July 6, 2025

Abstract. Food is one of the basic human needs that should always be available. To fulfill the role of in a region, the concept of food security is established to measure sufficiency, availability and quality of food. Food security for a country is expressed using Food Security Index (FSI). FSI score for a country reflects its ability for survival. It is therefore very important to measure the score and be able to predict future scores to enable control and improvement. To realize the improvement of Indonesia's food security, a model is needed to predict the Food Security Index in Indonesia. This paper explores the model using data from the Indonesian Food Security and Vulnerability Atlas (FSVA) at the Regency and City levels in 2018-2024 period with a total of 3,598 records. We evaluated Multiple Linear Regression, Least Absolute Shrinkage and Selection Operator, Random Forest, eXtreme Gradient Boosting, Support Vector Regression, and Ensemble Machine Learning models for predicting the FSI score. The models are evaluated using r-squared (R^2), Root Mean Square Error (RMSE), and Mean Absolute Error (MAE). The results shows that the XGBoost method is the best method for predicting the Food Security Index in Indonesia with an R^2 value of R^2 of 0.978, RMSE of 0.024, and MAE of 0.016. In addition, the XGBoost method predicts the average national Food Security Index score in 2025 and 2026 to be 75.14 respectively.

Keywords: data mining; model evaluation; food security and vulnerability atlas

INTRODUCTION

Food is one of the primary needs that must always be met by humans. In Indonesia, the right of citizens to obtain food is regulated in Article 27 paragraph 2 of the 1945 Constitution which states that “Tiap- tiap warga negara berhak atas pekerjaan dan penghidupan yang layak bagi kemanusiaan”. As a basic need and a form of human rights, food plays a very important role in life in an area. To fulfill the role of food for an area, the concept of food security was formed to measure sufficiency, availability and quality of food. In Article 1 of Law Number 18 of 2012 concerning food, food security is a condition of fulfilling food for the state up to individuals, which is reflected in the availability of sufficient food, both in quantity and quality, safe, diverse, nutritious, evenly distributed, and affordable and does not conflict with religion, beliefs, and culture of the community, to be able to live healthily, actively, and productively in a sustainable manner (Perum Bulog, 2014). On the global scale, food security is integral to the United Nations Sustainable Development Goals

(SDGs), particularly Goal 2, Zero Hunger, which targets the eradication of hunger, achievement of food security and improved nutrition, and the promotion of sustainable agriculture (Pristiandaru, 2023; United Nations, 2015).

In Indonesia, the level of food security is periodically assessed through the Food Security Index (FSI) as documented in the Food Security and Vulnerability Atlas (FSVA) (Sabarella et al., 2022). In 2023, the national average FSI score reached 74.43 (out of 100), based on evaluations across 416 regencies and 98 cities. Despite this moderate score, food security in Indonesia remains a complex issue influenced by a wide array of determinants. Direct factors include food consumption patterns and access to healthcare, while indirect factors encompass food availability, political stability, distribution infrastructure, and socioeconomic conditions (B. Saragih, 2022). In addition, conditions where not everyone has the ease of obtaining the food they need, which causes large-scale hunger and malnutrition in the world, will be an obstacle to achieving food security (Soekarwo, 2021). The



act of food waste, where food that is not eaten by the community often ends up in the trash, is also an obstacle to achieving food security ([Shopnil et al., 2023](#)). Fulfilling food needs is a fundamental aspect for humans to survive, so maintaining food security is very important for a country ([Andaiyani et al., 2024](#)). In addition, the use of technology accelerates the growth of local food agro-industry and various innovations significantly in accordance with market needs. The Food Security Index (FSI) Prediction in Indonesia can be a means for the government to allocate resources appropriately and efficiently in increasing food security. The FSI prediction in Indonesia is expected to be a means for the Government to make policies related to increasing food security, such as increasing domestic production, accelerating food development and infrastructure, developing food estate areas (food production centers), and strengthening national food reserves ([Frisnoiry et al., 2024](#)).

This study is related to previous studies, where research conducted by [Huang et al \(2020\)](#) provided results that using the ensemble machine learning method can provide a lower error rate on test data (test dataset) when compared to using a single method, where using ensemble machine learning obtained an R^2 value of 0.50, Root Mean Square Error (RMSE) of 16.01, and Mean Absolute Error (MAE) of 11.97. In addition, research conducted by [Phyo et al \(2022\)](#) provided results that using voting regression (VR) provided a lower average Mean Absolute Percentage Error (MAPE) value compared to the single prediction method, which was 4.28%. Based on the results of both studies, the ensemble machine learning method can provide better model performance than a single model. This study seeks to answer the following research question: Which machine learning model provides the most accurate prediction of the Food Security Index (FSI) in Indonesia. Based on the description above, this

study aimed to develop an machine learning-based model to predict the Food Security Index in Indonesia.

METHODS

Research Data

This study uses Food Security Index data sourced from the Indonesian Food Security and Vulnerability Atlas (FSVA) website at the regency and city levels in Indonesia. The Food Security Index data collected was from 514 regions (consisting of 416 regencies and 98 cities in Indonesia) each year in the 2018-2024 period with a total of 3,598 data. The attributes used in collecting Food Security Index data can be seen in [Table 1](#).

[Table 1](#) explains that there are 12 attributes and nine (9) of them are indicators that influence the Food Security Index score, which is divided into three aspects, namely aspects of food availability, aspects of food access, and aspects of food utilization. Food availability aspects include the availability of food from domestic production, reserves, imports, and aid. Food access refers to the ability to obtain enough food through one or a combination of food sources. Because sufficient food availability in a place may not be available to households with limited economic, physical, or social access, access to food is important in the food security framework. Food utilization refers to the amount of food consumed by a household and the capacity of each person to consume food and metabolize nutrients properly. The concept of food utilization encompasses various aspects, including preservation, preparation, and compliance with safety standards related to food and beverage products. Food utilization is also related to hygiene standards, eating habits (especially for individuals with specific nutritional needs), allocation of food resources within the household, and the health conditions of family members ([Bapanas, 2023](#)).

Table 1. Attributes used in the research data

No	Attributes	Information
1	Region	Name of the regencies and cities where the Food Security Index data were collected
2	Year	Time period in year of Food Security Index data collection
3	Food Security Index (FSI) score	Food Security Index score in Regency and City areas in Indonesia with a range between 0 and 100
4	Ratio of normative consumption per capita to food availability (NCPR)	Ratio of normative per capita consumption to net production of rice, corn, sweet potatoes, cassava, and sago, as well as government's rice stock area
5	Percentage of the population living below the poverty line	Data on the population living below the poverty line in percentage units (%)
6	Percentage of households with a proportion of expenditure on food of more than 65 percent of total expenditure	Data on the percentage of households where the proportion of expenditure on food is more than 65 percent of total household expenditure, both for food and non-food
7	Percentage of households without access to electricity	Percentage of households that do not have access to electricity, either Perusahaan Listrik Negara (PLN) or non-PLN, such as generators
8	The average length of schooling for women is over 15 years	Data on average length of schooling (total years of schooling up to highest level of education completed and highest grade ever occupied) by women aged 15 years and over
9	Percentage of households without access to clean water	Data on the percentage of households that use unprotected water sources, unprotected springs, surface water and rainwater as their main water source for drinking
10	The ratio of the number of residents per health worker to the population density level	Data on the number of health workers (including general practitioners, specialist doctors, dentists, midwives, public health workers, nutritionists, physical therapists, and medical personnel) per population density level
11	Percentage of toddlers with below standard height (stunting)	Data on the percentage of children under five years whose height is less than -2 Standard Deviations with a height-for-age index based on specific references for height for age and gender
12	Life expectancy at birth	Estimated average life expectancy of newborns assuming no change in mortality patterns throughout their lifetime in years

Food Security Index

The Food Security Index (FSI) is an index that uses a set of indicators used in calculating the composite score of a region's food security conditions. The FSI value can indicate the level of food and nutritional security of each region (regency, city, or province), as well as the relative ranking between regions. One of the food security index calculations was carried out using composite analysis, where the FSVA technical working group agreed to use a weighting method. The purpose of this composite analysis is to determine the relative importance of indicators to each component of food security. The regency/city composite score is calculated by adding up the results of multiplying each standardized indicator value by the indicator weight. using the z-score and scale distance (0 to 100) with the [Equation 1](#).

$$Y(j) = \sum_{i=1}^9 a_i x_{ij} \quad \dots 1)$$

Note: $Y(j)$: j-th regency/city composite score, a_i : weight of each i-th indicator, x_{ij} : standardization value of each i-th indicator in the j-th regency/city, i: 1st, 2nd, ..., 9th indicator, j: regencies 1, 2, ..., 416/cities 1, 2, ..., 98.

Data Mining

Data Mining is a data processing process that aims to obtain new information and can be used as a guide in decision making ([Suntoro, 2019](#)). One of the processes in data mining is The Cross Industry Standard Process for Data Mining (CRISP-DM). CRISP-DM is a data mining standard that was first introduced in 1999 ([Martínez-Plumed et al., 2019](#)). The stages of data mining using CRISP-DM include business understanding, data understanding, data preparation, modeling, evaluation, and deployment.

Machine Learning

Machine learning is a field of study that develops methods to automatically learn and

complete certain tasks that are usually performed by humans. Learning in this case is related to how to complete various tasks or make accurate predictions based on previously learned patterns ([Shalev-Shwartz & Ben-David, 2014](#)). Machine learning has a main feature in the form of a self-learning concept where this refers to the application of statistical modeling to detect patterns and improve performance based on empirical data and information that is carried out without direct programming command.

Machine learning is divided into four main categories, namely supervised learning, unsupervised learning, semi-supervised learning, and reinforcement learning. However, this study focuses more on supervised learning, which is learning that concentrates on learning patterns by connecting relationships between variables with known outcomes and working with labeled data sets. Supervised learning works by feeding a variety of sample data features represented as "X" and the correct output data values represented as "y". The fact that the output values and features are known makes the data set qualify as "labeled". The algorithm then analyzes the patterns contained in the data and builds a model that can produce similar rules when applied to new data. The machine parses the patterns and rules of the data and then creates a model that is an equivalent of the algorithm. This model is used to generate output with new data based on the rules taken from the training data. Once the model is created, it can be used on new data and tested to evaluate its accuracy ([Theobald, 2017](#)).

Multiple Linear Regression (MLR)

One of the most widely used statistical techniques across a variety of disciplines, such as psychology, sociology, marketing, and health research, is multiple linear regression. Controlling confounding factors and measuring their impact on the dependent variable are made possible by multiple linear regressions. More thorough and complete evaluations are made possible by its ability to allow researchers to look at the effects of

several independent factors on a single dependent variable all at once (Y. Huang et al., 2024).

The values of the independent variables are entered into the equation produced by linear regression which can be used to predict future values of the dependent variable (Maharadja et al., 2021). The equation of the line in multiple linear regression can be seen with the [Equation 2](#).

$$y = \beta_0 + \beta_1 x_1 + \dots + \beta_n x_n + \epsilon \quad \dots 2)$$

Note: y : predicted value of the dependent variable, β_0 : constant, $\beta_1 x_1$: regression coefficient on the first independent variable (x_1), $\beta_n x_n$: regression coefficient on the n th independent variable (x_n), ϵ : model error (namely how much variation there is in the prediction of the y value).

Least Absolute Shrinkage and Selection Operator (LASSO)

LASSO is an analytical method that integrates variable selection and normalization to enhance the predictive accuracy and interpretability of the resulting statistical model. The LASSO model is especially beneficial for handling high-dimensional data, where the number of variables significantly exceeds the number of observations. Conventional linear regression models often struggle with overfitting when they become overly complex and too tailored to the training data, leading to poor performance on unseen data. By selecting the most relevant variables and simplifying the model, LASSO effectively mitigates this issue (Christakis et al., 2024).

The parameter estimates for the LASSO method can be obtained using [Equation 3](#).

Random Forest (RF)

Random Forests (RF) were originally developed to process large datasets while preserving strong statistical performance. This approach is highly effective for predictive tasks due to its exceptional accuracy (Doz et al., 2023). RF is a special type of ensemble learning algorithm that is an extension of the

bagging method because it uses bagging and feature randomness to create uncorrelated decision trees. The randomness in the model structure helps reduce variance, with predictions generated by averaging the outputs of the employed trees. The random selection of training instances and input features creates diversity among the decision trees, enabling them to capture different aspects of the relationship between input features and the target variable. This reduces correlation between trees, making the random forest more resilient to noise and outliers. In regression tasks, the predictions of individual decision trees are averaged to produce the final output. Meanwhile, for classification, the predicted results are taken from several decision trees based on majority voting (Özen, 2024).

$$\hat{\beta}^{lasso} = \sum_{i=1}^n (y_i - \beta_0 - \sum_{k=1}^p \beta_k x_{ik})^2 + \lambda \sum_{k=1}^p |\beta_k| \quad \dots 3)$$

Note: y_i : dependent variable on the i -th observation, β_0 : constant, β_k : coefficient of the i -th independent variable, x_{ik} : independent variable, n : number of observations, p : number of independent variables used in a model, λ : regularization parameter.

eXtreme Gradient Boosting (XGBoost)

XGBoost is an algorithm that can find the best solution for various problems, specifically in prediction, regression, and classification tasks. The basic principle of this algorithm is to gradually refine the learning parameters with the aim of minimizing the loss function. Each tree learns from the remains of the previous tree. By using a more regular model, XGBoost builds a better regression tree structure, which improves performance and allows to reduce model complexity to avoid overfitting (Yulianti et al., 2022). XGBoost combining

many decision trees in an iterative manner to minimize errors ([Bentéjac et al., 2021](#); [Wijaya et al., 2024](#)) and has ability to modeling complex nonlinear relationships ([Li, 2022](#)).

The final prediction result of XGBoost is calculated by combining the prediction results from all regression trees, which can be expressed using [Equation 4](#).

$$\hat{y}_i = \sum_{k=1}^n f_k(x_i), f_k \in F \quad \dots 4)$$

Note: F: regression tree space, f_k : corresponds to a tree, $f_k(x_i)$: result of tree k, $\hat{y}_i$: the i-th predicted value of instance x_i .

Support Vector Regression (SVR)

Support Vector Regression (SVR) is one of the applications of Support Vector Machine (SVM) that seeks to minimize error by identifying the optimal hyperplane and reducing the gap between predicted and actual values ([Ahadian & Parand, 2022](#)). The basic principle of SVR is to map the feature vectors on low-dimensional sample data to high dimensions and perform regression analysis on the sample data in high dimensions using kernel functions ([Zhang et al., 2022](#)). Kernel functions are used by SVR to transform non-linear inputs into a larger feature space, which is then solved linearly. There are three kernels that are most often used in the SVR method, namely polynomial kernels, linear kernels, and Radial Basis Function (RBF) kernels ([Saputra et al., 2019](#)). The SVR function can be expressed mathematically with the [Equation 5](#).

$$f(x) = \omega \times x + b \quad \dots 5)$$

Note: ω : function coefficient, x: input feature vector, b: bias constant.

Ensemble Machine Learning

Ensemble Machine Learning is a method used to combine two or more machine learning algorithms to obtain superior performance when compared to using a single machine

learning method ([Mienye & Sun, 2022](#)). Ensemble machine learning aims to integrate multiple machine learning algorithms within a unified framework. Thus, the complementary information from each algorithm is effectively utilized to enable better overall model performance ([Dong et al., 2020](#)). The fundamental idea behind ensemble machine learning is the recognition that machine learning models have limitations and can make mistakes, and that there are limitations to a single machine learning algorithm, such as high variance, high bias, and low accuracy ([Rincy & Gupta, 2020](#)). There are three methods in ensemble machine learning, i.e. boosting, bagging, and stacking. The ensemble method in this study uses voting regression mechanism, where the final prediction is the arithmetic mean of predictions from five base models (MLR, LASSO, RF, XGBoost, and SVR). Equal weights were used for all base models to maintain model interpretability and reduce overfitting risk.

Data Preprocessing

The Food Security Index data that has been collected is subjected to a data preprocessing stage, as shown in [Figure 1](#).

[Figure 1](#) shows a flowchart of data preprocessing. The preprocessing stage begins by reading the raw data and checking if the collected data contains empty or missing data values, then the empty or missing data values are replaced (replace missing value) using the average of each attribute in the data. Furthermore, a check is carried out whether there is duplicate data or not. If duplicate data is found, the data is deleted because it will affect the performance of the model formed. Data that is no longer duplicated is then checked for data outliers. If an outlier is found in the data, the outlier data will be removed. The Region data attribute in [Table 1](#) which has a categorical data type is changed to a numeric data type through the label encoding process because most machine learning algorithms can only process numeric data, and the region data is arranged alphabetically. The label encoding process will provide a unique numeric value

for each category in the variable. After the label encoding process is complete, the next step is the data normalization stage for each

data attribute in [Table 1](#) using the min-max method using the formula that can be seen in the [Equation 6](#).

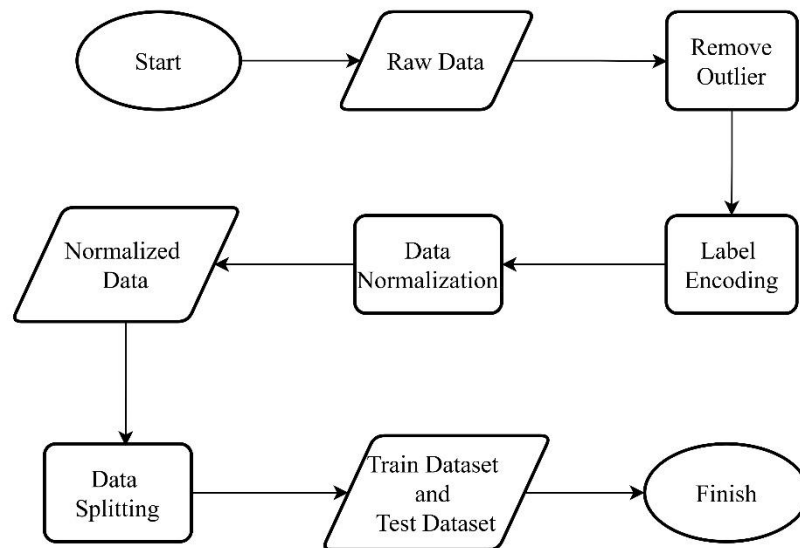


Figure 1. The flowchart of data preprocessing

$$x_{new} = \frac{x_{old} - x_{min}}{x_{max} - x_{min}} \quad \dots 6)$$

Note: x_{new} : data value after normalization, x_{old} : data value, x_{max} : maximum value of the data attribute, x_{min} : minimum value of the data attribute.

The normalized data is then divided into training data and test data randomly with a percentage ratio of training data and test data of 70:30. The percentage ratio of 70 percent for training data and 30 percent for testing data was chosen because it is applied in many studies in the field of machine learning and data mining. In addition, the percentage ratio of 70 percent for training data and 30 percent for testing data provides a good balance between the amount of data used in model training and the amount of data used for model performance evaluation. Train dataset is used to build a model to predict the Food Security Index. Meanwhile, test datasets are used to test the model formed in the overall data training process.

Data Training and Data Testing

A visualization of the process design is presented in [Figure 2](#). The training data and test data that have been formed are trained using training data to form a model using a single machine learning-based prediction algorithm, namely multiple linear regression, Random Forest, Extreme Gradient Boosting (XGBoost), Support Vector Regression (SVR), and Least Absolute Shrinkage and Selection Operator (LASSO).

The training stage begins by applying a single machine learning-based prediction algorithm in the stage of forming a model using previously entered training data to then predict the Food Security Index. pMLR, pLASSO, pXGB, pSVR, and pRF are predictions formed by each model. The Food Security Index (FSI) prediction model uses a single machine learning-based prediction algorithm that has been formed previously in the first scenario, a model ensemble is carried out with the voting regression method which produces the final prediction results (indicated by pVR).

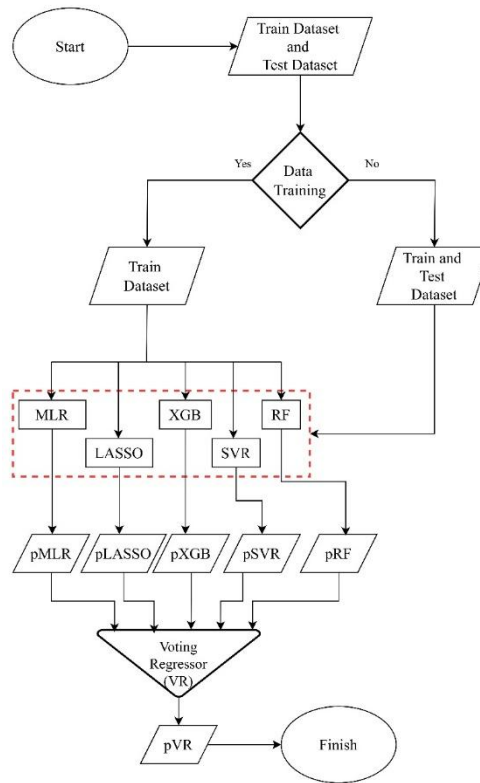


Figure 2. Flowchart of data training and data testing

Model Evaluation

The results of model predictions using a single machine learning-based prediction algorithm and ensemble machine learning are tested for model feasibility by conducting an evaluation in the form of measuring the proportion of variance in the dependent variable (predicted variable) that can be explained by the independent variable (predictor variable) in the model using R-squared (R^2) that can be seen mathematically with [Equation 7](#).

$$R^2 = 1 - \frac{\sum (y_i - \hat{y}_i)^2}{\sum (y_i - \bar{y})^2} \quad \dots 7)$$

Note: R^2 : coefficient of determination, y : actual value; $\hat{y}$: predicted value, $\bar{y}$: average value, i : data sequence.

In addition, the calculation of the error rate between the actual data and the predicted data is carried out by calculating the Root Mean Squared Error (RMSE) and Mean Absolute Error (MAE) values on the training data and

test data. RMSE calculates the root value of the average of the squared differences. Squaring emphasizes if there is a large distance for one data point because it deviates from the average value which helps in assessing the algorithm better. Meanwhile, MAE refers to the average total absolute error calculated, where the absolute error is the total amount of error in the measurement. The RMSE and MAE values on the training data indicate the suitability of the developed model. Meanwhile, the RMSE and MAE values on the test data indicate the performance of the developed model. [Equation 8](#) and [9](#) are equations to mathematically calculate the RMSE and MAE values respectively.

$$RMSE = \sqrt{\frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{n}} \quad \dots 8)$$

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad \dots 9)$$

Note: RMSE: Root Mean Square Error, MAE: Mean Absolute Error, y : actual value, $\hat{y}$: predicted value, $\bar{y}$: average value, i : data sequence, n : amount of data.

In addition to holdout validation, we applied 5-fold cross-validation to each model to enhance the robustness of performance evaluation. The cross-validation was stratified by year to preserve temporal consistency. The average R^2 , RMSE, and MAE across folds were recorded and discussed alongside the holdout validation results.

Experimental Scenario

This study uses two experimental scenarios. The first scenario is prediction using a single machine learning algorithms to obtain prediction results for each algorithm used and determine the best algorithm. The second scenario is prediction using an ensemble machine learning algorithm with a voting regression method, which works by combining predictions from a single machine learning model in the first scenario by averaging the results of the previously formed model evaluation in the form of R -squared (R^2), RMSE, and MAE values using a simple averaging technique that distributes weights uniformly on each single machine learning model in the first scenario to make final predictions and determine the best algorithm.

RESULTS AND DISCUSSION

Relevance of Indicators in the Food Security Index in Indonesia

Indonesian Food Security Index (FSI) data that used in this study based of [Table 1](#) consists of 12 attributes, where attribute number 7, namely "Percentage of households without access to electricity" and attribute number 8, namely "The average length of schooling for women is over 15 years" have data that tends to be stable so that both indicators are no longer relevant in measuring the FSI. Both indicators also have quite small feature importance so that

these indicators are less relevant in predicting the FSI in Indonesia ([Manikas et al., 2023](#)). Attribute number 11, namely "Percentage of toddlers with below standard height (stunting)" has a fairly small feature importance and this indicator is an outcome of various factors so that it is not relevant to be used as an indicator of the FSI in Indonesia ([Bühler et al., 2018](#); [De Sanctis et al., 2021](#)).

Indonesia's Food Security Index (FSI) Average Score

This study uses Indonesia FSI score data in 2018-2024 by taking the average FSI score from regencies and cities in Indonesia. The average scores of the FSI at the regency and city levels in Indonesia can be seen in [Figure 3](#).

[Figure 3](#) shows the average score of the food security index at the district and city levels from 2018 to 2024. In general, there was an increase in the average score from 2018 to 2020 and an insignificant increase in the average score in 2021. However, there was an insignificant decrease in the average score in 2022 due to the prolonged COVID-19 pandemic in Indonesia. In 2023, the average score of the food security index in Indonesia experienced a significant increase. This could happen because conditions in Indonesia became more stable after the COVID-19 pandemic and entered the COVID-19 endemic period ([F. M. Saragih & Saragih, 2020](#); [Utoro et al., 2025](#)). This also happened in 2024 where there was an insignificant increase.

Data Preprocessing

Data preprocessing was done by reading the raw data consisting of 3,598 data and checking whether there was empty or missing data. This check did not find any empty or duplicated data. Then, an outlier check was carried out using the interquartile range technique. The results of the outlier check on the Food Security Index data are presented in [Table 2](#).

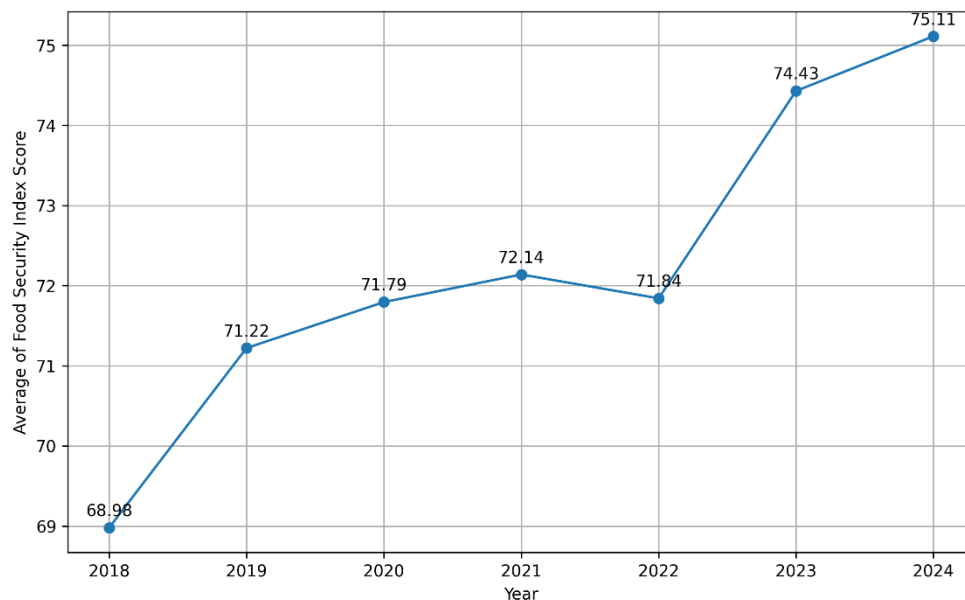


Figure 3. Average Food Security Index (FSI) Score in Indonesia (2018-2024)

Table 2. Outlier Checking Results

No	Attributes	Number of Outliers
1	Food Security Index score	128
2	Percentage of population living below the poverty line	242
3	Percentage of households with the proportion of expenditure on food is more than 65 % of total expenditure	11
4	Percentage of households without electricity access	440
5	The average length of schooling for women is over 15 years	17
6	Percentage of households without access to clean water	36
7	The ratio of population per health worker to population density	236
8	Percentage of toddlers with height below the standard (stunting)	1
9	Life expectancy at birth	12

[Table 2](#) calculates the number of data containing outliers in each attribute of the food security index data in Indonesia. Data attributes containing outliers with a total of 1,123 data are then removed. After that, data balancing is carried out because there is an imbalance in the amount of data available each year and a balanced food security index is produced, where each year there are 312 data so that the total is 2,184 data. Balanced data is processed by label encoding by changing the name of the categorical area to numeric. Then, a data normalization process is carried out which aims to standardize the range of data on each attribute and uses the holdout validation

technique which is divided into training data and test data randomly with a percentage ratio of train dataset and test dataset of 70:30, where there were 1,528 data in train dataset and 656 data in test dataset. Also, we applied 5-fold cross-validation to each model to enhance the robustness of performance evaluation.

Hyperparameter Tuning

The data is divided into training data and testing data and then the training and prediction process is carried out in two experimental scenarios. The training process is carried out without and using hyperparameter tuning that can be seen in [Table 3](#).

Table 3. Hyperparameter Tuning

Algorithm	Hyperparameter	Value
Multiple Linear Regression (MLR)	fit_intercept	True
Least Absolute Shrinkage and Selection Operator (LASSO)	alpha	0.1
	max_iter	1,000
Random Forest (RF)	n_estimators	300
	max_depth	None
	min_samples_split	2
Extreme Gradient Boosting (XGBoost)	n_estimators	300
	max_depth	3
	learning_rate	0.1
Support Vector Regression (SVR)	C	10
	kernel	radial basis function
	gamma	scale

Note: Hyperparameter tuning was performed based on the method used in this study, but not the ensemble machine learning method and represents the best hyperparameters for each method used in this study.

[Table 3](#) lists the hyperparameters used in the data training process. From the several hyperparameters used, grid search techniques are used to find the best combination of hyperparameters for each machine learning algorithm. The data training process was carried out using the hyperparameters listed in [Table 3](#) and then used to predict the Indonesian FSI.

Model Evaluation

The Food Security Index (FSI) model formed at the training and testing stages was tested for feasibility by calculating the r-squared (R^2), Root Mean Squared Error (RMSE), and Mean Absolute Error (MAE) values. The results of the model's evaluation of the Indonesian Food Security Index using holdout validation are presented in [Table 4](#).

[Table 4](#) presents the results of the model evaluation for several machine learning methods, divided into four scenarios: one with hyperparameter tuning and one without, also one with outlier removal and one without. Holdout validation (using a training and test dataset) is used on this study. Based on the R^2 value, the XGBoost method demonstrates the highest capability in explaining data variability across all scenarios, both with and without

hyperparameter tuning with R^2 values ranging from 0.928 to 0.999. Moreover, the Random Forest (RF) and Support Vector Regression (SVR) methods also yield high R^2 values, although they do not surpass XGBoost in all experimental and evaluation scenarios. Conversely, the Multiple Linear Regression (MLR) and Least Absolute Shrinkage and Selection Operator (LASSO) methods exhibit lower R^2 values, indicating a weaker ability to explain data variability. The ensemble method, which integrates multiple individual machine learning models in the second scenario, produces competitive R^2 values, although still lower than those of RF, SVR, and XGBoost.

Regarding Root Mean Squared Error (RMSE) and Mean Absolute Error (MAE) values in [Table 4](#), XGBoost remains the best-performing method across all experimental and evaluation scenarios, with RMSE values ranging from 0.002 to 0.049. It is followed by RF, SVR, the ensemble method, MLR, and LASSO. The ensemble approach provides relatively good RMSE and MAE values, albeit not as strong as RF, SVR, and XGBoost.

If look at the training time in [Table 4](#), the MLR and LASSO models are the two models that provide the fastest training time both

without and using hyperparameter tuning. However, both models have poor performance. RF and ensemble experienced a fairly drastic increase in training time after hyperparameter tuning with an insignificant increase in performance in the RF method. The ensemble model has the highest training time both

without and using hyperparameter tuning which provides performance that is not as good as XGBoost. XGBoost achieves the best balance between high performance and relatively fast training time (0.304 seconds without hyperparameter tuning and 0.326 seconds with hyperparameter tuning).

Table 4. Model evaluation results (holdout validation)

Algorithm	Train Dataset			Test Dataset			Training Time (second)
	R ²	RMSE	MAE	R ²	RMSE	MAE	
Without Outlier Removal and Without Hyperparameter Tuning							
MLR	0.675	0.094	0.07	0.693	0.089	0.067	0.008
LASSO	0.0	0.165	0.119	0.0	0.16	0.115	0.004
RF	0.996	0.01	0.006	0.974	0.026	0.017	2.664
XGBoost	0.999	0.003	0.002	0.978	0.024	0.016	1.488
SVR	0.851	0.064	0.056	0.841	0.064	0.055	0.082
Ensemble	0.878	0.058	0.046	0.873	0.057	0.045	3.11
Without Outlier Removal and With Hyperparameter Tuning							
MLR	0.675	0.094	0.07	0.693	0.089	0.067	0.007
LASSO	0.0	0.165	0.119	0.0	0.16	0.115	0.003
RF	0.996	0.01	0.006	0.974	0.026	0.017	8.539
XGBoost	0.994	0.012	0.009	0.978	0.024	0.016	0.232
SVR	0.898	0.053	0.044	0.841	0.064	0.055	0.116
Ensemble	0.89	0.055	0.043	0.873	0.057	0.045	8.902
With Outlier Removal and Without Hyperparameter Tuning							
MLR	0.311	0.154	0.113	0.302	0.151	0.11	0.004
LASSO	0.0	0.185	0.145	0.0	0.181	0.144	0.005
RF	0.985	0.022	0.015	0.912	0.053	0.037	1.583
XGBoost	0.999	0.002	0.002	0.934	0.046	0.033	0.304
SVR	0.84	0.074	0.06	0.785	0.084	0.064	0.076
Ensemble	0.82	0.079	0.06	0.77	0.086	0.066	1.951
With Outlier Removal and With Hyperparameter Tuning							
MLR	0.311	0.154	0.113	0.302	0.151	0.11	0.003
LASSO	0.0	0.185	0.145	0.0	0.181	0.144	0.009
RF	0.986	0.022	0.015	0.912	0.053	0.037	4.489
XGBoost	0.988	0.02	0.015	0.943	0.043	0.029	0.326
SVR	0.878	0.064	0.054	0.81	0.079	0.062	0.166
Ensemble	0.83	0.076	0.058	0.79	0.083	0.063	5.423

Note: R² value close to 1 is better. Smaller RMSE, MAE, and Training Time values are better, respectively.

The presence of outliers in each attribute of the Indonesian FSI dataset influences model performance. Typically, the larger the number of outliers, the more difficult it is for the model to produce accurate predictions. To address this issue, outlier removal was implemented to enhance model performance. However, outlier removal also has weaknesses, because it can eliminate important information regarding certain regions and cities in Indonesia with extreme food security indicators as indicated by a decrease in the R^2 value along with an

increase in the RMSE and MAE values in models formed using models other than RF and XGBoost. These extreme values may provide essential insights for measuring the Food Security Index in these areas.

In addition to holdout validation, we applied 5-fold cross-validation to each model to enhance the robustness of performance evaluation. The results of the model's evaluation of the Indonesian FSI using 5-fold Cross Validation are presented in [Table 5](#).

Table 5. Model Evaluation Results (5-Fold Cross Validation)

Algorithm	5-Fold Cross Validation			Execution Time (second)
	R ²	RMSE	MAE	
Without Outlier Removal				
MLR	0.662	0.094	0.071	0.035
LASSO	0.0	0.164	0.119	0.028
RF	0.953	0.034	0.022	47.189
XGBoost	0.959	0.031	0.02	1.255
SVR	0.815	0.068	0.055	0.578
Ensemble	0.856	0.061	0.048	50.624
With Outlier Removal				
MLR	0.259	0.155	0.114	0.034
LASSO	0.0	0.185	0.147	0.048
RF	0.86	0.068	0.046	27.758
XGBoost	0.835	0.068	0.043	1.204
SVR	0.758	0.089	0.068	0.614
Ensemble	0.75	0.09	0.067	29.308

Note: R^2 value close to 1 is better. Smaller RMSE, MAE, and Execution Time values are better, respectively.

[Table 5](#) is the result of model evaluation in predicting food security index using 5-fold cross validation with and without outlier removal. This table adds a cross-validation dimension that provides a more robust estimate of the model's performance in general. The R^2 , RMSE, and MAE values from 5-fold Cross Validation are generally slightly more conservative because they are tested on 5 different subsets. This makes 5-fold Cross Validation more accurate in generalizing new data, compared to using holdout validation which is prone to bias.

The application of outlier removal resulted in a decrease in performance in the overall model when compared to without applying outlier removal based on the R^2 , RMSE, and MAE values of each model. The decrease in performance when outliers are removed indicates that extreme data is a key indicator in forming food security patterns and is more pronounced in cross-validation because of a more comprehensive evaluation of model stability. When viewed from the complex relationship between the features that determine the food security index, non-linear

models (such as XGBoost and RF) are more suitable compared to linear models. The XGBoost method shows high performance and relatively fast execution time, so it is considered ideal for real-world applications.

Although the proposed XGBoost model shows strong predictive performance, there are several limitations. First, the model only uses FSI data from 2018 to 2024 and thus may not fully capture unexpected structural changes outside this range. Second, Indonesia's vast and diverse geographical conditions may limit the applicability of the model to other contexts. Third, some indicators used in the model, such as stunting, are outcome variables that can be influenced by multiple factors, raising concerns about reverse causality. Compared with related literature (e.g., [Bentéjac et al., 2021](#); [Li, 2022](#)), our study confirms the superior accuracy of XGBoost in complex

socio-economic prediction tasks. Future research is expected to benefit from incorporating time-series deep learning models and integrating climate variability and food system resilience factors.

Future Food Security Index Prediction

This section contains predictions of the average national Food Security Index score at the Regency and City levels in Indonesia in 2025 and 2026 using the food security index data available in 2024 using without outlier removal and hyperparameter tuning scenarios because these scenarios generally obtained better model evaluations based on [Table 4](#) and [Table 5](#). The average score of the food security index prediction for each model can be seen in [Figure 4](#).

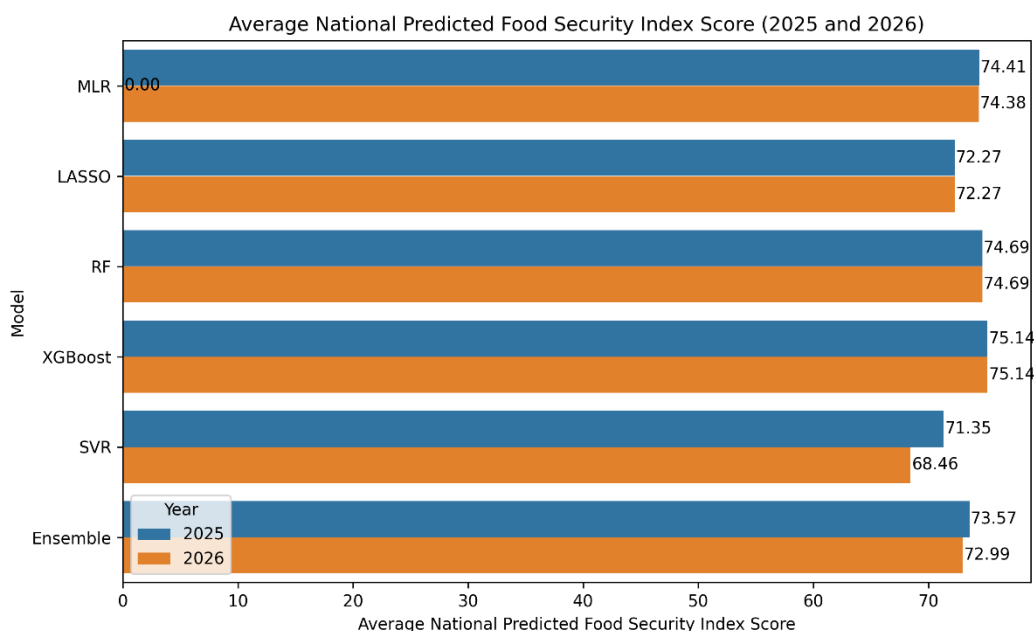


Figure 4. Average National Predicted Food Security Index Score (2025 and 2026)

[Figure 4](#) displays a comparison of the average prediction of the national Food Security Index in Indonesia for 2025 and 2026 based on several Machine Learning models using the FSI data in 2024. For information, the average score of the Food Security Index in Indonesia in 2024 is 75.11. The model formed using the LASSO, RF, and XGBoost methods

provides the same average national predicted Food Security Index score from 2025 to 2026. Meanwhile, the model formed using the MLR, SVM, and Ensemble methods provides a downward trend from 2025 to 2026. The model formed using XGBoost predicts an average score of the national FSI that is closest to the average score of the national FSI in 2024.

CONCLUSION

The analysis reveals that XGBoost is the most effective model for predicting Indonesia's Food Security Index (FSI), achieving an R^2 of 0.978, RMSE of 0.024, and MAE of 0.016. Its predictions for 2025 and 2026 closely align with the 2024 FSI average. Although the ensemble method proposed in this study performs well, it is slightly less accurate than XGBoost. The FSI indicators require revision to reflect current conditions. Specifically, attributes such as "Percentage of households without access to electricity" and "Average length of schooling for women over 15 years" show stable trends and may introduce redundancy. Additionally, the indicator "Percentage of stunted toddlers" should be reconsidered, as it results from multiple contributing factors.

REFERENCES

- Ahadian, P., & Parand, K. (2022). Support vector regression for the temperature-stimulated drug release. *Chaos, Solitons & Fractals*, 165, 112871. <https://doi.org/10.1016/j.chaos.2022.112871>
- Andaiyani, S., Marwa, T., & Nurhaliza, S. (2024). Ekonomi Biru dan Ketahanan Pangan: Studi Empiris Provinsi Kepulauan di Indonesia. *Jurnal Ekonomi Indonesia*, 13(1), 69–85. <https://doi.org/10.52813/jei.v13i1.324>
- Bapanas. (2023). *Peta Ketahanan dan Kerentanan Pangan Tahun 2023*. <https://fsva.badanpangan.go.id/?910c4c71bfed3c0a49daa1b18116a3c9>
- Bentéjac, C., Csörgő, A., & Martínez-Muñoz, G. (2021). A comparative analysis of gradient boosting algorithms. *Artificial Intelligence Review*, 54, 1937–1967. <https://doi.org/10.1007/s10462-020-09896-5>
- Bühler, D., Hartje, R., & Grote, U. (2018). Matching food security and malnutrition indicators: evidence from Southeast Asia. *Agricultural Economics*, 49(4), 481–495. <https://doi.org/10.1111/agec.12430>
- Christakis, I., Sarri, E., Tsakiridis, O., & Stavrakas, I. (2024). Investigation of LASSO Regression Method as a Correction Measurements' Factor for Low-Cost Air Quality Sensors. *Signals*, 5(1), 60–86. <https://doi.org/10.3390/signals5010004>
- De Sanctis, V., Soliman, A., Alaaraj, N., Ahmed, S., Alyafei, F., & Hamed, N. (2021). Early and Long-term Consequences of Nutritional Stunting: From Childhood to Adulthood. *Acta Bio-Medica : Atenei Parmensis*, 92(1), e2021168. <https://doi.org/10.23750/abm.v92i1.11346>
- Dong, X., Yu, Z., Cao, W., Shi, Y., & Ma, Q. (2020). A survey on ensemble learning. *Frontiers of Computer Science*, 14, 241–258. <https://doi.org/10.1007/s11704-019-8208-z>
- Doz, D., Cotič, M., & Felda, D. (2023). Random Forest Regression in Predicting Students' Achievements and Fuzzy Grades. *Mathematics*, 11(19). <https://doi.org/10.3390/math11194129>
- Frisnoiry, S., Pratiwi, I. A., Tarigan, N. C. W., & Sulaiman, R. P. (2024). Konsep, Indeks, Pendekatan, dan Strategi Ketahanan Pangan. *Esensi Pendidikan Inspiratif*, 6(2).
- Huang, J.-C., Tsai, Y.-C., Wu, P.-Y., Lien, Y.-H., Chien, C.-Y., Kuo, C.-F., Hung, J.-F., Chen, S.-C., & Kuo, C.-H. (2020). Predictive Modeling of Blood Pressure during Hemodialysis: A Comparison of Linear Model, Random Forest, Support Vector Regression, XGBoost, LASSO Regression, and Ensemble Method. *Computer Methods and Programs in Biomedicine*, 195, 105536. <https://doi.org/10.1016/j.cmpb.2020.105536>
- Huang, Y., Xu, W., Sukjairungwattana, P., & Yu, Z. (2024). Learners' continuance intention in multimodal language learning education: An innovative multiple linear regression model. *Heliyon*, 10(6). <https://doi.org/10.1016/j.heliyon.2024.e2>

- 8104
- Li, Z. (2022). Extracting spatial effects from machine learning model using local interpretation method: An example of SHAP and XGBoost. *Computers, Environment and Urban Systems*, 96, 101845. <https://doi.org/10.1016/j.compenvurbsys.2022.101845>
- Maharadja, A. N., Maulana, I., & Dermawan, B. A. (2021). Penerapan Metode Regresi Linear Berganda untuk Prediksi Kerugian Negara Berdasarkan Kasus Tindak Pidana Korupsi. *Journal of Applied Informatics and Computing*, 5(1), 95–102. <https://doi.org/10.30871/jaic.v5i1.3184>
- Manikas, I., Ali, B. M., & Sundarakani, B. (2023). A systematic literature review of indicators measuring food security. *Agriculture & Food Security*, 12(1), 10. <https://doi.org/10.1186/s40066-023-00415-7>
- Martínez-Plumed, F., Contreras-Ochando, L., Ferri, C., Hernández-Orallo, J., Kull, M., Lachiche, N., Ramírez-Quintana, M. J., & Flach, P. (2019). CRISP-DM Twenty Years Later: From Data Mining Processes to Data Science Trajectories. *IEEE Transactions on Knowledge and Data Engineering*, 33(8), 3048–3061. <https://doi.org/10.1109/TKDE.2019.2962680>
- Mienye, I. D., & Sun, Y. (2022). A survey of ensemble learning: Concepts, algorithms, applications, and prospects. *IEEE Access*, 10, 99129–99149. <https://doi.org/10.1109/ACCESS.2022.3207287>
- Özen, F. (2024). Random forest regression for prediction of Covid-19 daily cases and deaths in Turkey. *Heliyon*, 10(4). <https://doi.org/10.1016/j.heliyon.2024.e25746>
- Perum Bulog. (2014). *Ketahanan pangan*. Bulog. <https://www.bulog.co.id/beraspangan/ketahanan-pangan/>
- Phyo, P.-P., Byun, Y.-C., & Park, N. (2022). Short-Term Energy Forecasting Using Machine-Learning-Based Ensemble Voting Regression. In *Symmetry* (Vol. 14, Issue 1). <https://doi.org/10.3390/sym14010160>
- Pristiandaru, D. L. (2023, May 7). *Mengenal Tujuan 2 SDGs: Tanpa Kelaparan*. Kompas. Retrieved from <https://lestari.kompas.com/read/2023/05/07/080000986/mengenal-tujuan-2-sdgs-tanpa-kelaparan>
- Rincy, T. N., & Gupta, R. (2020). Ensemble learning techniques and its efficiency in machine learning: A survey. *2nd International Conference on Data, Engineering and Applications (IDEA)*, 1–6. <https://doi.org/10.1109/IDEA49133.2020.9170675>
- Sabarella, Komalasari, W. B., Wahyuningsih, S., Sehusman, Manurung, M., Supriyati, Y., Rinawati, Saida, M. D. N., Seran, K., Arief, M., Firmansyah, R., & Amara, V. D. (2022). *Statistik Ketahanan Pangan Tahun 2022*. Pusat Data dan Sistem Informasi Pertanian Sekretariat Jenderal Kementerian Pertanian.
- Saputra, G. H., Wigena, A. H., & Sartono, B. (2019). Penggunaan Support Vector Regression dalam Pemodelan Indeks Saham Syariah Indonesia dengan Algoritma Grid Search. *Indonesian Journal of Statistics and Its Applications*, 3(2), 148–160. <https://doi.org/10.29244/ijsa.v3i2.172>
- Saragih, B. (2022). *Dapur Tahan Pangan, Sehat dan Bergizi (Pengentasan Kerawanan Pangan dan Gizi di Indonesia)*. Yogyakarta: Penerbit Deepublish.
- Saragih, F. M., & Saragih, B. (2020). Gambaran Kebiasaan Makan Masyarakat Pada Masa Pandemi Covid-19. *ResearchGate*.
- Shalev-Shwartz, S., & Ben-David, S. (2014). *Understanding Machine Learning: From Theory to Algorithms*. United States: Cambridge University Press.
- Shopnil, M. M. M., Husna, A., Sultana, S., & Islam, M. N. (2023). An Ensemble ML

- Model to Predict the Wastage of Food: Towards Achieving the Food Sustainability. 2023 *International Conference on Next-Generation Computing, IoT and Machine Learning (NCIM)*, 1–6. <https://doi.org/10.1109/NCIM59001.2023.10212669>
- Soekarwo. (2021, May 10). *Sistem dan Upaya Memperkuat Ketahanan Pangan*. Dewan Pertimbangan Presiden. Retrieved from <https://wantimpres.go.id/id/2021/05/sistem-dan-upaya-memperkuat-ketahanan-pangan/>
- Suntoro, J. (2019). *Data Mining: Algoritma dan Implementasi dengan Pemrograman PHP*. Elex Media Komputindo.
- Theobald, O. (2017). *Machine Learning For Absolute Beginners* (2nd ed.). London: Scatterplot Press.
- United Nations. (2015). *Goal 2 | Department of Economic and Social Affairs*. United Nations. <https://sdgs.un.org/goals/goal2>
- Utoro, P. A. R., Pujokaroni, A. S., Aini, Q., & Saragih, B. (2025). Factors and comparative analysis of COVID-19's impact on household food security in rural and urban regions. *Nutrición Clínica y Dietética Hospitalaria*, 45(1). <https://doi.org/10.12873/451utoro>
- Wijaya, H., Hostiadi, D. P., & Triandini, E. (2024). Optimization of XGBoost Algorithm Using Parameter Tuning in Retail Sales Prediction. *Jurnal Nasional Pendidikan Teknik Informatika: JANAPATI*, 13(3), 769–786. <https://doi.org/10.23887/janapati.v13i3.82214>
- Yulianti, S. E. H., Soesanto, O., & Sukmawaty, Y. (2022). Penerapan Metode Extreme Gradient Boosting (XGBOOST) pada Klasifikasi Nasabah Kartu Kredit. *Journal of Mathematics: Theory and Applications*, 4(1), 21–26. <https://doi.org/10.31605/jomta.v4i1.1792>
- Zhang, Y., Wang, Q., Chen, X., Yan, Y., Yang, R., Liu, Z., & Jiahong, F. (2022). The Prediction of Spark-Ignition Engine Performance and Emissions Based on the SVR Algorithm. *Processes*, 10, 312. <https://doi.org/10.3390/pr10020312>