

Comparison of Ordinal Logistic Regression and Artificial Neural Network in Stunting Prevalence Classification

May Risnawati¹, M. Fathurahman^{2*}, Surya Prangga²

¹Laboratory of Computational Statistics, Mulawarman University, Samarinda, Indonesia

²Department of Statistics, Mulawarman University, Samarinda, Indonesia

*Corresponding author: fathur@fmipa.unmul.ac.id

Received: 3 January 2025

Revised: 16 October 2025

Accepted: 30 October 2025

ABSTRACT

The prevalence of under-five stunting is one of the crucial health problems in Indonesia. Stunting is a growth and development disorder in children due to chronic malnutrition and repeat infections that can hurt children's physical and cognitive development. This study aims to analyse the accuracy of the classification of the prevalence of stunting on regencies/cities in Indonesia, in 2022 using two methods, namely Ordinal Logistic Regression (OLR) and Artificial Neural Network (ANN). OLR is development of logistic regression applied to response variables with more the two categories that have levels or ranks, while ANN is a method that mimics the function of the biological nervous system and is designed for complex information processing. This study used two proportions of data splitting namely 80:20 and 90:10. Each method produce two models, OLR 1 and OLR 2 for the OLR method, and ANN 1 and ANN 2 for the ANN method. The results show that the ANN 1 model with 80:20 data proportion performs better than other models with an accuracy of 63.37%.

Keywords—ANN, Classification, OLR, Stunting.

I. INTRODUCTION

Stunting is a condition where children under five years old show a lower stature or height than their peers [1]. Stunting can inhibit the growth of a child's height, resulting in a shorter stature compared to children of the same age and impairing motor skills. In addition, stunting can reduce Intelligence Quotient (IQ), which affects children's learning ability [2]. The state of stunting can be measured using the z-score (standard deviation/SD-score), where height-for-age < -2 SD score from the median of WHO child growth standards [3]. Based on the prevalence data of stunted children under five collected by the World Health Organization (WHO), Indonesia is the third country with the highest prevalence in the South-East Asia Regional (SEAR) area after India (38.4%) and Timor Leste (50.5%). The average prevalence of stunting in Indonesia from 2005-2017 was 36.4% [4]. The prevalence of stunting shows a decreasing trend in Indonesia, from 37.22% in 2013 to 30.8% in 2018. In 2019 it was 27.7%, 24.4% in 2021, and 21.6% in 2022. Despite the decline, stunting is still considered a severe problem in Indonesia. The prevalence rate is still higher than the target of the National Medium-Term Development Plan for 2020-2024, which is 14%, and the WHO target is below 20% [5].

One of the government's efforts to reduce the prevalence of *stunting* among children under five is the issuance of Presidential Regulation Number 72 of 2021, which focuses on the Acceleration of *Stunting* Reduction [6]. This regulation is the legal basis for the National Strategy to Accelerate Stunting Reduction which has been implemented since 2018. The National Strategy for Accelerating Stunting Reduction includes specific interventions and sensitivity, implemented in a convergent, holistic, integrative, and quality manner through multisectoral cooperation at the center, regions, and villages [7]. This classification of stunting severity is important to support this strategy, as it helps health workers, and the government identify and prioritize interventions in areas with a prevalence of stunted children under five [8].

Machine learning plays a role in supporting this strategy, mainly through data classification. Machine learning is a subset of Artificial Intelligence (AI) that emphasizes how computers can learn from data through complex patterning and simulation. Classification is a data analysis method that aims to develop models that describe essential classes. These models, known as classifiers, are used to predict the labels of categorical classes [9], [10]. In this research, the classification methods used are Ordinal Logistic Regression (OLR) and Artificial Neural Network (ANN).

Logistic regression is a method in statistical analysis that has been popular over the past three decades in healthcare analysis and medical information retrieval [11]. This method is designed to model and predict the probability of an event between 0 and 1 [12]. Logistic regression is a method used to model the relationship between binary or polytomous response variables with one or more categorical or continuous predictor variables [13]. OLR is a generalization of logistic regression applied to response variables with more than two categories considered to have a particular order [14].

Artificial Neural Network (ANN) is a method inspired by the function of the biological nervous system, especially the information processing capabilities of human brain cells [15]. The framework of an ANN is described by the number of neurons, the number of layers, and the layer size, which consists of the number of neurons in each layer. ANN is a practical and flexible modelling method that can apply pattern information to new data, handle input that may contain noise, and produce reliable and logical estimates [16].

This research was conducted to compare ANN and OLR methods for classifying stunting prevalence in regencies/cities in Indonesia in 2022.

II. LITERATURE REVIEW

A. Ordinal Logistic Regression

Logistic regression is a technique used to analyze the relationship between binary (having only two categories) or polytomous (having more than two categories) response variables with one or more categorical or continuous predictor variables [13]. Ordinal Logistic Regression (OLR) is a logistic regression with an ordinal scale categorical dependent variable. OLR does not allow a linear relationship or correlation between independent variables. Multicollinearity can be detected by looking at the Variance Inflation Factor (VIF) value [17], [18]. If there is a response variable Y_i and $\mathbf{x}_i^T = [x_{i1} \ x_{i2} \ \dots \ x_{ip}]$ states a vector of p predictor variables on i -th subject, the logit model formed is:

$$\text{logit} = [P(Y_i \leq j | \mathbf{x}_i)] = \ln \left[\frac{P(Y_i \leq j | \mathbf{x}_i)}{1 - P(Y_i \leq j | \mathbf{x}_i)} \right] = \alpha_j + \mathbf{x}_i^T \boldsymbol{\beta} \quad (1)$$

where $\boldsymbol{\beta} = [\beta_1 \ \beta_2 \ \dots \ \beta_p]^T$ is a vector of parameters describing the effect of the predictor variable. $P(Y_i \leq j | \mathbf{x}_i)$ is the probability that the value of the response variable y_i is less than or equal to j , which depends on the value of the predictor variable ($\mathbf{x}_i$). The OLR model used is expressed as below [19]:

$$P(Y_i \leq j | \mathbf{x}_i) = \frac{\exp(\alpha_j + \mathbf{x}_i^T \boldsymbol{\beta})}{1 + \exp(\alpha_j + \mathbf{x}_i^T \boldsymbol{\beta})}, \quad j = 1, 2, \dots, J - 1 \quad (2)$$

Parameter estimation of the OLR model is done using the Maximum Likelihood Estimation (MLE) method. Suppose there is a response variable with 3 categories ($J = 3$), then the OLR model formed is as follows:

$$\begin{aligned} \text{logit}[P(Y_i \leq 1 | \mathbf{x}_i)] &= \ln \left[\frac{P(Y_i \leq 1 | \mathbf{x}_i)}{1 - P(Y_i \leq 1 | \mathbf{x}_i)} \right] = \alpha_1 + \mathbf{x}_i^T \boldsymbol{\beta} \\ \text{logit}[P(Y_i \leq 2 | \mathbf{x}_i)] &= \ln \left[\frac{P(Y_i \leq 2 | \mathbf{x}_i)}{1 - P(Y_i \leq 2 | \mathbf{x}_i)} \right] = \alpha_2 + \mathbf{x}_i^T \boldsymbol{\beta} \end{aligned} \quad (3)$$

Then the probability value of each category on the j -th response variable is:

$$\begin{aligned} \pi_1(\mathbf{x}_i) &= \frac{\exp(\alpha_1 + \mathbf{x}_i^T \boldsymbol{\beta})}{1 + \exp(\alpha_1 + \mathbf{x}_i^T \boldsymbol{\beta})} \\ \pi_2(\mathbf{x}_i) &= \frac{\exp(\alpha_2 + \mathbf{x}_i^T \boldsymbol{\beta})}{1 + \exp(\alpha_2 + \mathbf{x}_i^T \boldsymbol{\beta})} - \frac{\exp(\alpha_1 + \mathbf{x}_i^T \boldsymbol{\beta})}{1 + \exp(\alpha_1 + \mathbf{x}_i^T \boldsymbol{\beta})} \\ \pi_3(\mathbf{x}_i) &= 1 - \frac{\exp(\alpha_2 + \mathbf{x}_i^T \boldsymbol{\beta})}{1 + \exp(\alpha_2 + \mathbf{x}_i^T \boldsymbol{\beta})} \end{aligned} \quad (4)$$

Parameter estimates obtained from the equation must go through a significance test which consists of two, namely simultaneous tests and partial tests. The simultaneous tests are used to determine the effect of predictor variables together or simultaneously on the response variable [17], [20]. The hypothesis used in the simultaneous test is as follows:

$$\begin{aligned} H_0 : \beta_1 = \beta_2 = \dots = \beta_p &= 0 \\ H_1 : \text{there is at least one } \beta_k &\neq 0, \quad k = 1, 2, \dots, p \end{aligned} \quad (5)$$

The simultaneous tests use the likelihood ratio test with the test statistic G which is formulated as [13]:

$$G = -2 \ln \left[\frac{\left(\frac{n_1}{n} \right)^{n_1} \left(\frac{n_2}{n} \right)^{n_2} \left(\frac{n_3}{n} \right)^{n_3}}{\prod_{i=1}^n [\pi_1(x_i)^{y_{1i}} \pi_2(x_i)^{y_{2i}} \pi_3(x_i)^{y_{3i}}]} \right] \quad (6)$$

with $n_1 = \sum_{i=1}^n y_{1i}$, $n_2 = \sum_{i=1}^n y_{2i}$, $n_3 = \sum_{i=1}^n y_{3i}$, $n = n_1 + n_2 + n_3$. The critical region at the α significant level for the hypothesis in Equation (5) is H_0 rejected if the value of $G > \chi_{\alpha, db}^2$ with $db = p$ is the degree of freedom.

Partial test is used to determine the effect of each predictor variable on the response variable [20]. The hypothesis of partial test is as follows:

$$\begin{aligned} H_0 : \beta_k &= 0 \\ H_1 : \beta_k &\neq 0, \quad k = 1, 2, \dots, p \end{aligned} \quad (7)$$

The partial test uses the Wald test with the W test statistic which is formulated as follows [13]:

$$W = \frac{\hat{\beta}_k}{SE(\hat{\beta}_k)} \quad (8)$$

where $\hat{\beta}_k$ is the parameter estimate of β_k and $SE(\hat{\beta}_k)$ the standard error of $\hat{\beta}_k$. The critical region at the α significant level for the hypothesis in Equation (7) is H_0 rejected if the value of $|W| > Z_{\frac{\alpha}{2}}$.

B. Artificial Neural Network

Artificial Neural Network (ANN) is an analysis technique inspired by how the human brain calculates complex processes different from conventional digital computers [21]. ANN is composed of several layers, namely [22]:

1. *Input layer*: Neurons that receive external input from outside the network (data source). These neurons do not perform any processing on the data, they only forward the data to the next layer.

2. *Hidden layer*: neurons that have no direct interaction with the “outside world” but only with other neurons in the network.
3. *Output layer*: neurons that output the processed information from the network to the external environment.

The main elements used to build an ANN model consists of $x_1, x_2, \dots, x_p$ Representing the information (input) received by neurons, $\mathbf{W} = (w_1, w_2, \dots, w_p)$ is a vector of synapse weights that modify the information received and b is the bias (intercept or threshold) of a neuron. Mathematically, this can be expressed as [23]:

$$z = \mathbf{W}^T \mathbf{x} = w_1 x_1 + w_2 x_2 + \dots + w_p x_p. \quad (9)$$

An activation function is a function used in ANN to determine the output of each neuron. Its adequacy lies in its ability to make the system learn and solve complex tasks, making the neural network more powerful [24].

The optimizer is key component of the model training process. Adam is the most popular and widely used optimizer in deep learning. Adam is an efficient method for stochastic optimization because it only requires first-order gradients with small memory requirements. Adams algorithm combines Adagrad and RMSProp methods that adjust the learning rate of each parameter according to the estimated 1st and 2nd moment of each gradient [25].

Learning rate is the parameter that controls how much the weights and biases are corrected in the ANN during the training process. It plays a role in determining how certain inputs to the network will produce the desired output. The learning rate value usually uses a very small number, such as 0.001 or less [23].

C. Confusion Matrix

The level of classification accuracy can also be seen using the Apparent Correct Classification Rate (ACCR) [17]. ACCR measures the proportion of training samples that are correctly classified by the classification function [26].

An example of a confusion matrix for three-category classification is shown in Table 1.

Table 1 Confusion matrix for three categories

Actual Group	Number of observations	Prediction Group		
		$\hat{n}_1$	$\hat{n}_2$	$\hat{n}_3$
1	n_1	n_{11}	n_{12}	n_{13}
2	n_2	n_{21}	n_{22}	n_{23}
3	n_3	n_{31}	n_{32}	n_{33}

where n_r is the number of observations in group r , with $r = 1, 2, 3$. Classification accuracy using ACCR is formulated as follows:

$$\text{ACCR} = \frac{n_{11} + n_{22} + n_{33}}{n_1 + n_2 + n_3} \times 100\%. \quad (10)$$

D. Stunting Prevalence

Stunting prevalence is the percentage of children under the age of five with a height less than or below the average expected height for their age. The formula for calculating stunting prevalence is as follows:

$$\text{Stunting prevalence} = \frac{\text{Number of stunted children}}{\text{Number of children measured}} \times 100\%. \quad (11)$$

The following is a description of the stunting prevalence categories based on the prevalence threshold set by the WHO-UNICEF Technical Advisory Group on Nutrition Monitoring (TEAM) and released in 2018 [8]:

Table 2 Stunting prevalence categories and thresholds

Category	Prevalence Threshold (%)
Very low	$y < 2,5$
Low	$2,5 \leq y < 10$
Medium	$10 \leq y < 20$
High	$20 \leq y < 30$
Very high	$y \geq 30$

III. METHODOLOGY

The data used in this study are secondary. Data on the prevalence of stunting toddlers and data on stunting handling indicators were obtained from the Ministry of Health website <https://kesmas.kemkes.go.id/>. The population used in this study is 514 districts/cities in Indonesia, while the sample used is 503 districts/cities in Indonesia in 2022. The variables in this study consist of response variables (y) and predictor variables (x), presented in Table 3.

Table 3 Variables research

Notation	Variable Name	Definition
y	Prevalence of stunting in toddler	The percentage of children under five who are underweight or below the average expected height for their age. Stunting prevalence categories: 0 = Medium (2,5% $\leq y < 20\%$) 1 = High (20% $\leq y < 30\%$) 2 = Very high ($y \geq 30\%$)
x_1	Immunization	Percentage of children aged 12-23 months who received complete basic immunization
x_2	Birth assistance by healthcare workers in healthcare facilities	Percentage of ever-married women (PPK) aged 15-49 years whose last childbirth was assisted by a skilled health worker at a health facility
x_3	Modern family planning (FP)	Percentage of women of reproductive age (15-49 years) or their partners who are sexually active and want to delay having children or do not want to have more children and use modern methods of contraception
x_4	Exclusive breastfeeding	Percentage of infants less than six months of age who are exclusively breastfed
x_5	Complimentary feeding (CF)	Percentage of children 6-23 months of age receiving complementary feeding
x_6	Access to safe drinking water	Percentage of households that have access to safe drinking water source services
x_7	Access to proper sanitation	Percentage of households that have access to proper and sustainable sanitation services
x_8	Early Childhood Education (ECE)	Percentage of gross enrollment rate (APK) of early childhood education 3-6 years old
x_9	JKN/jamkesda ownership	Percentage of the population who have JKN/jamkesda
x_{10}	Recipient of KPS/KKS or food assistance	Percentage of households receiving KPS/KKS or food assistance (bottom 40 percent of the population)

The data analysis technique used for classifying stunting in toddlers is the OLR and ANN method. The steps in the data analysis process by researchers are as follows:

1. Performing data pre-processing
2. Perform descriptive statistical analysis
3. Detecting multicollinearity in predictor variables
4. Data labeling is performed, namely changing the value of the variable of stunting toddlers into three categories: medium, high, and very high.
5. The dataset is divided into training data and testing data, with the proportion of division divided into two scenarios, namely 80:20 and 90:10.
6. Perform classification using the OLR method.
7. Performing classification using the ANN method
8. Comparing the accuracy of OLR and ANN models.

IV. RESULTS AND DISCUSSIONS

A. Data Preprocessing

Data preprocessing involves checking missing data and duplicating data. Variables with a proportion of missing data, less than 5%, will be processed by removing all rows that have missing values from the dataset. Before the removal of missing data, the number of research units was 514, after the removal, the number of research units became 503. In addition, no duplicate data was found in the dataset.

B. Descriptive Statistical Analysis

Figure 1 shows a map of the prevalence of stunting in 514 districts/cities in Indonesia, which is categorized based on the prevalence threshold of stunting in Table 2.

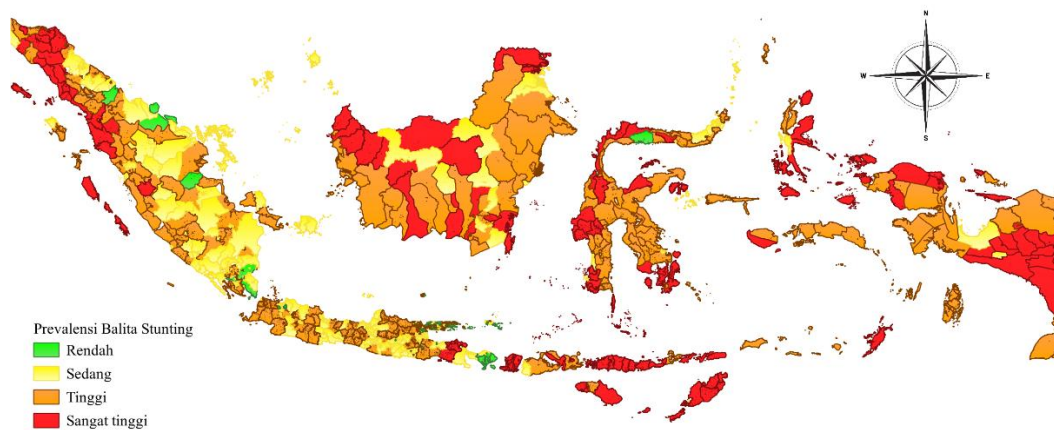


Figure 1 Distribution of the prevalence of stunting by district/city in Indonesia in 2022

Based on Figure 1, most districts/cities in Indonesia are categorized as having a high prevalence of stunting. Several districts/cities with a very high prevalence of stunting are in Papua, East Nusa Tenggara, and Aceh, as well as several districts/cities on Sulawesi Island and Kalimantan Island. In addition, there are a small number of districts/cities with a very high category in Java Island. Meanwhile, several districts/cities are in the medium category, especially in Sumatra Island and Java Island. On the other hand, several regions have reduced the prevalence of stunting toddlers to the low category, including big cities such as Surabaya and Denpasar and several other regions in various provinces.

C. Detecting Multicollinearity

Multicollinearity is detected by looking at the Variance Inflation Factor (VIF) value. A VIF value of ≥ 10 can indicate multicollinearity between predictor variables. Based on Table 4, the results show that the VIF value for each predictor variable is less than 10, so there is no indication of multicollinearity among the predictor variables

Table 4 VIF value of predictor variables

Variable	VIF Value
Immunization (x_1)	1,22
Birth assistance by healthcare workers in healthcare facilities (x_2)	1,91
Modern Family Planning (FP) (x_3)	1,46
Exclusive breastfeeding (x_4)	1,11
Complimentary Feeding (CF) (x_5)	1,12
Access to safe drinking water (x_6)	1,77
Access to proper sanitation (x_7)	2,01
Early Childhood Education (ECE) (x_8)	1,37
JKN/jamkesda ownership (x_9)	1,26
Recipient of KPS/KKS or food assistance (x_{10})	1,27

D. Data Labeling

Data labeling was performed on the prevalence of stunting among children under five (y), classified into three categories: 0 for medium, 1 for high, and 2 for very high. The category is medium if the prevalence of stunting is between 2.5% and less than 20%, high if the prevalence of stunting is between 20% and less than 30%, and very high if the prevalence of stunting reaches 30% or more. Based on the available data, districts/cities with a prevalence of stunting among children under five that were initially in the low category will be combined with the medium category because the data in the low category is relatively small. The results of the classification of the prevalence of stunting among children under five (y) are as follows: 170 districts/cities in category 0 (medium), 217 districts/cities in category 1 (high), and 116 districts/cities in category 3 (very high).

E. Splitting data

The results of the division of training data and testing data with a proportion of 80:20, consisting of 402 training data and 101 testing data. The training data consisted of 136 districts/cities in category 0 (medium), 173 districts/cities in category 1 (high), and 93 districts/cities in category 2 (very high). The testing data included 34 districts/cities in category 0 (medium), 44 districts/cities in category 1 (high), and 23 districts/cities in category 2 (very high). The result of the

division of training data and testing data is a proportion of 90:10, consisting of 452 training data and 51 testing data. The training data consisted of 153 districts/cities in category 0 (medium), 195 districts/cities in category 1 (high), and 104 districts/cities in category 2 (very high). The testing data included 17 districts/cities in category 0 (medium), 22 districts/cities in category 1 (high), and 12 districts/cities in category 2 (very high).

F. Prevalence Stunting Prevalence Classification Using OLR

1. Parameter estimation

Classification of the prevalence of stunting toddlers in Indonesia using OLR begins with the formation of an OLR model written as follows:

$$\text{logit}[P(Y_i \leq 1 | \mathbf{x}_i)] = \alpha_1 + \beta_1 x_{i1} + \beta_2 x_{i2} + \beta_3 x_{i3} + \beta_4 x_{i4} + \beta_5 x_{i5} + \beta_6 x_{i6} + \beta_7 x_{i7} + \beta_8 x_{i8} + \beta_9 x_{i9} + \beta_{10} x_{i10} \quad (12)$$

$$\text{logit}[P(Y_i \leq 2 | \mathbf{x}_i)] = \alpha_2 + \beta_1 x_{i1} + \beta_2 x_{i2} + \beta_3 x_{i3} + \beta_4 x_{i4} + \beta_5 x_{i5} + \beta_6 x_{i6} + \beta_7 x_{i7} + \beta_8 x_{i8} + \beta_9 x_{i9} + \beta_{10} x_{i10} \quad (13)$$

The parameter estimation of the OLR model in Equations (12) and (13) can be obtained using the MLE method and Fisher Scoring iteration. There are two OLR models used to classify the prevalence of stunting toddlers in Indonesia, namely the OLR 1 model, which is built based on the proportion of training and testing data, namely 80:20, and the OLR 2 model, which uses the proportion of data, namely 90:10. Based on the results of the OLR model parameter estimation in, the logit model for the OLR 1 model is displayed in Equation (14).

$$\begin{aligned} \text{logit}[P(Y_i \leq 1 | \mathbf{x}_i)] &= -5,1825 + 0,0035x_{i1} + 0,0247x_{i2} + 0,0195x_{i3} - 0,0107x_{i4} - 0,0128x_{i5} + 0,0310x_{i6} + \\ &\quad 0,0196x_{i7} - 0,0004x_{i8} - 0,0112x_{i9} - 0,0188x_{i10} \\ \text{logit}[P(Y_i \leq 2 | \mathbf{x}_i)] &= -2,8952 + 0,0035x_{i1} + 0,0247x_{i2} + 0,0195x_{i3} - 0,0107x_{i4} - 0,0128x_{i5} + 0,0310x_{i6} + \\ &\quad 0,0196x_{i7} - 0,0004x_{i8} - 0,0112x_{i9} - 0,0188x_{i10} \end{aligned} \quad (14)$$

The OLR 2 model can be seen in Equation (15).

$$\begin{aligned} \text{logit}[P(Y_i \leq 1 | \mathbf{x}_i)] &= -5,0335 + 0,0028x_{i1} + 0,0260x_{i2} + 0,0178x_{i3} - 0,0078x_{i4} - 0,0138x_{i5} + 0,0209x_{i6} + \\ &\quad 0,0270x_{i7} - 0,0028x_{i8} - 0,0125x_{i9} - 0,0128x_{i10} \\ \text{logit}[P(Y_i \leq 2 | \mathbf{x}_i)] &= -2,7733 + 0,0028x_{i1} + 0,0260x_{i2} + 0,0178x_{i3} - 0,0078x_{i4} - 0,0138x_{i5} + 0,0209x_{i6} + \\ &\quad 0,0270x_{i7} - 0,0028x_{i8} - 0,0125x_{i9} - 0,0128x_{i10} \end{aligned} \quad (15)$$

2. Simultaneous parameter testing

Simultaneous testing is used to determine the effect of predictor variables on response variables together or simultaneously. The hypothesis used for the simultaneous test is as follows:

$$H_0 : \beta_1 = \beta_2 = \dots = \beta_{10} = 0$$

$$H_1 : \text{there is at least one } \beta_k \neq 0, k = 1, 2, \dots, 10$$

The test statistic used is the G test statistic in Equation (6), with the critical region H_0 rejected if the $G > \chi^2_{0,05,10}$. The calculation results can be seen in Table 5.

Table 5 Simultaneous test result

Model	G	db	$\chi^2_{0,05,10}$	Decision
Model OLR 1	103,8236	10	18,31	H_0 is rejected
Model OLR 2	107,5529	10		H_0 is rejected

Based on Table 5, it can be seen that the G value of the OLR 1 and OLR 2 models is greater than the value of $\chi^2_{0,05,10}$, so it is decided that H_0 is rejected, so it can be concluded that simultaneously the predictor variables have a significant effect on the prevalence of stunting in toddlers in Indonesia.

Table 6 Results of partial testing of OLR Model 1

Parameter	W	$Z_{0,025}$	Decision
$\hat{\beta}_1$	0,7260	1,96	H_0 accepted
$\hat{\beta}_2$	2,9400		H_0 rejected
$\hat{\beta}_3$	2,8480		H_0 rejected
$\hat{\beta}_4$	-1,9680		H_0 rejected
$\hat{\beta}_5$	-1,5370		H_0 accepted
$\hat{\beta}_6$	2,8960		H_0 rejected
$\hat{\beta}_7$	2,2360		H_0 rejected
$\hat{\beta}_8$	0,0450		H_0 accepted
$\hat{\beta}_9$	1,7460		H_0 accepted
$\hat{\beta}_{10}$	-3,1070		H_0 rejected

3. Partial parameter testing

Partial testing is used to determine the effect of predictor variables on response variables individually (partially). The hypotheses used for partial testing are as follows:

$$H_0 : \beta_k = 0$$

$$H_1 : \beta_k \neq 0, k = 1, 2, \dots, 10$$

The test statistic used is the W test statistic in Equation (7), with the critical area H_0 rejected if the value of $|W| > Z_{\frac{\alpha}{2}}$. The results are shown in Tables 6 and 7.

In Table 6, six predictor variables significantly influence the prevalence of stunted toddlers in Indonesia: birth assistance by healthcare workers in healthcare facilities, modern Family Planning (FP), exclusive breastfeeding, access to safe drinking water, access to proper sanitation, and recipients of KPS/KKS or food assistance.

Table 7 Partial test of OLR Model 2

Parameter	W	$Z_{0,025}$	Decision
$\hat{\beta}_1$	0,6150	1,96	H_0 accepted
$\hat{\beta}_2$	3,2560		H_0 rejected
$\hat{\beta}_3$	2,7400		H_0 rejected
$\hat{\beta}_4$	-1,5380		H_0 accepted
$\hat{\beta}_5$	-1,7840		H_0 accepted
$\hat{\beta}_6$	2,1570		H_0 rejected
$\hat{\beta}_7$	3,2040		H_0 rejected
$\hat{\beta}_8$	-0,3860		H_0 accepted
$\hat{\beta}_9$	-2,0720		H_0 rejected
$\hat{\beta}_{10}$	-2,2580		H_0 rejected

In Table 7, six predictor variables significantly affect the prevalence of stunted toddlers in Indonesia: birth assistance by healthcare workers in healthcare facilities, modern Family Planning (FP), access to safe drinking water, access to proper sanitation, ownership of JKN/Jamkesda, and recipients of KPS/KKS or food assistance.

4. Classification of Prevalence of Stunting Toddlers Using the OLR 1 Model

Prediction of classification of the prevalence of stunting toddlers in districts/cities in Indonesia using the OLR 1 model, which is built from 80% of the total data, namely 402 training data. Classification of the prevalence of stunting toddlers using the model in Equation (14). The results of the classification prediction on the training data are displayed in the form of a confusion matrix in Table 8.

Table 8 Confusion matrix of the OLR 1 model on the training data

Actual Group	Number of observation	Prediction Group		
		Medium	High	Very high
Medium	136	65	69	2
High	173	44	116	13
Very high	93	5	60	28

Based on Table 8, the ACCR value is obtained using Equation (10) as follows:

$$\begin{aligned} \text{ACCR} &= \frac{65 + 116 + 28}{136 + 173 + 93} \times 100\% \\ &= \frac{209}{402} \times 100\% \\ &= 51,99\%. \end{aligned}$$

Model performance evaluation is carried out using testing data, with the results of the confusion matrix prediction shown in Table 9.

Table 9 Confusion matrix of the OLR 1 model on the testing data

Actual Group	Number of observations	Prediction Group		
		Medium	High	Very high
Medium	34	16	18	0
High	44	9	31	4
Very high	23	2	17	4

Based on Table 9, the ACCR value is obtained as follows:

$$\begin{aligned} \text{ACCR} &= \frac{16 + 31 + 4}{34 + 44 + 23} \times 100\% \\ &= \frac{51}{101} \times 100\% \\ &= 50,50\%. \end{aligned}$$

The ACCR value obtained shows that the model can predict new data by 50.50%.

5. Classification of Prevalence of Stunting Toddlers Using the OLR 2 Model

Prediction of classification of the prevalence of stunting toddlers in districts/cities in Indonesia using the OLR 2 model, which is built from 90% of the total data, namely 452 training data. Classification of the prevalence of stunting toddlers

using the model in Equation (15). The results of the classification prediction on the training data are displayed in the form of a confusion matrix in Table 10.

Table 10 Confusion Matrix of the OLR 2 Model on the training data

Actual Group	Number of observations	Prediction Group		
		Medium	High	Very high
Medium	153	72	79	2
High	195	48	131	16
Very high	104	5	70	29

In Table 10, the ACCR value can be seen as follows:

$$\begin{aligned} \text{ACCR} &= \frac{72 + 131 + 29}{153 + 195 + 104} \times 100\% \\ &= \frac{235}{452} \times 100\% \\ &= 51,33\%. \end{aligned}$$

Model performance evaluation is carried out using testing data, with the results of the Confusion matrix prediction shown in Table 11.

Table 11 Confusion matrix of the OLR 2 model on the testing data

Actual Group	Number of observations	Prediction Group		
		Medium	High	Very high
Medium	17	7	10	0
High	22	4	16	2
Very high	12	0	9	3

Based on Table 11, the ACCR value is obtained as follows:

$$\begin{aligned} \text{ACCR} &= \frac{7 + 16 + 3}{17 + 22 + 12} \times 100\% \\ &= \frac{26}{51} \times 100\% \\ &= 50,98\%. \end{aligned}$$

The ACCR value obtained shows that the model can predict new data by 50.98%.

G. Stunting Prevalence Classification Using ANN

To understand the structure of the model used, the architectural illustration of the ANN applied in this study can be seen in Figure 2.

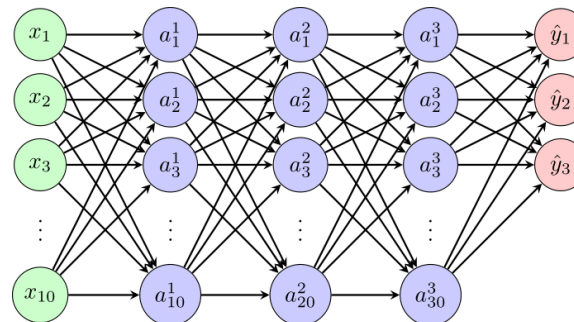


Figure 2 ANN model design

In Figure 2, the ANN network architecture has three main layers: the input, hidden, and output. The input layer has 10 neurons representing 10 input variables or predictor variables. Each neuron in this layer receives one feature from the data used in the model. This network has three hidden layers, each having 10, 20, and 30 neurons. Each neuron in the first hidden layer is connected to each neuron in the input layer, and each neuron in the second hidden layer is connected to each neuron in the first hidden layer. The hidden layer functions to perform computations to recognize complex patterns in the data. Furthermore, the output layer has three neurons representing three output categories. Each neuron in this layer outputs a probability value for each category predicted by the model. The activation function used in the hidden layer is ReLU, while the output layer uses SoftMax to produce probability values. The models built will be divided into ANN model 1 for the 80:20 data sharing proportion and ANN model 2 for the 90:10 data sharing proportion.

After designing the ANN network architecture, the next step is to compile the model. This compilation process aims to determine several important parameters used during model training. In this model, using the Adam optimizer, the loss function is categorical cross-entropy, and the evaluation metric is accuracy in monitoring model performance during training.

The compiled model is then used to train the training data. This training process will be carried out for 200 epochs, which means the model will see the entire dataset 200 times to minimize the loss function and improve its accuracy. The

training process in ANN involves two phases, namely forward propagation and backpropagation. These two phases work together to improve the weights and biases in the network so that the model can make more accurate predictions.

1. Classification of Prevalence of Stunting Toddlers Using ANN Model 1

The prediction results of classifying the prevalence of stunting toddlers in districts/cities in Indonesia using 402 training data are presented in the confusion matrix in Table 12.

Table 12 Confusion Matrix of ANN model 1 on training data

Actual Group	Number of observations	Prediction Group		
		Medium	High	Very high
Medium	136	82	46	8
High	173	22	130	21
Very high	93	12	29	52

Based on Table 12, the ACCR value can be seen as follows:

$$\begin{aligned} \text{ACCR} &= \frac{82 + 130 + 52}{136 + 173 + 93} \times 100\% \\ &= \frac{264}{402} \times 100\% \\ &= 65.67\%. \end{aligned}$$

The model evaluation results using testing data are summarized in the Confusion matrix presented in Table 13.

Table 13 Confusion Matrix of ANN model 1 on testing data

Actual Group	Number of observations	Prediction Group		
		Medium	High	Very high
Medium	34	21	13	0
High	44	8	32	4
Very high	23	5	7	11

Based on Table 13, the ACCR value is obtained as follows:

$$\begin{aligned} \text{ACCR} &= \frac{21 + 32 + 11}{34 + 44 + 23} \times 100\% \\ &= \frac{64}{101} \times 100\% \\ &= 63.37\%. \end{aligned}$$

The ACCR value obtained shows that the model can predict new data by 63.37%.

2. Classification of Prevalence of Stunting Toddlers Using ANN 2 Model

The results of predicting the prevalence classification of stunting toddlers in districts/cities in Indonesia using 452 training data are summarized in the confusion matrix presented in Table 14.

Table 14 Confusion Matrix of the ANN 2 model on training data

Actual Group	Number of observations	Prediction Group		
		Medium	High	Very high
Medium	153	92	57	4
High	195	61	109	25
Very high	104	17	38	49

Based on Table 14, the ACCR values can be seen as follows:

$$\begin{aligned} \text{ACCR} &= \frac{92 + 109 + 49}{153 + 195 + 104} \times 100\% \\ &= \frac{250}{452} \times 100\% \\ &= 55.31\%. \end{aligned}$$

The model evaluation results using testing data are summarized in the Confusion matrix presented in Table 15.

Table 15 Confusion Matrix of the ANN 2 model on testing data

Actual Group	Number of observations	Prediction Group		
		Medium	High	Very high
Medium	17	10	7	0
High	22	6	13	3
Very high	12	1	5	6

In Table 15, the ACCR values are obtained as follows:

$$\begin{aligned}\text{ACCR} &= \frac{10 + 13 + 6}{17 + 22 + 12} \times 100\% \\ &= \frac{29}{51} \times 100\% \\ &= 56.86\%.\end{aligned}$$

The ACCR value obtained shows that the model can predict new data by 56.86%.

H. Comparison of Classification Methods in Classifying the Prevalence of Toddler Stunting in Districts/Cities in Indonesia in 2022

Table 16 compares the accuracy of the classification results of the prevalence of toddler stunting in districts/cities in Indonesia in 2022 using the OLR and ANN methods.

Table 16 Comparison of classification method accuracy

Method	Proportion	Accuracy	
		Training	Testing
OLR	80:20	51,99%	50,50%
	90:10	51,33%	50,98%
ANN	80:20	65,67%	63,37%
	90:10	55,31%	56,86%

Based on Table 16, the ANN method with a training and testing data proportion of 80:20 has the highest accuracy value, 63.37%. This shows that the ANN method with a training and testing data proportion of 80:20 works better than the OLR method in classifying the prevalence of stunted toddlers in districts/cities in Indonesia in 2022.

V. CONCLUSIONS AND SUGGESTIONS

Based on the results of the classification of the prevalence of stunting toddlers in districts/cities in Indonesia in 2022 using the ANN and OLR methods with data proportions of 80:20 and 90:10, the ANN method with data proportions of 80:20 shows better performance than the OLR method. The ANN method, with a data proportion of 80:20, obtained an accuracy of 63.37%, while the proportion of 90:10 only reached 56.86%. OLR, with a data proportion of 80:20, obtained an accuracy of 50.50%, while at a data proportion of 90:10, it obtained an accuracy of 50.98%. This indicates that ANN, especially at a proportion of 80:20, provides better classification performance than OLR. For future research, the selection of architecture design parameters in ANN is generally done through trial-and-error procedures, so it is recommended to explore the ANN architecture, such as variations in the number of neurons and the number of hidden layers, in order to produce optimal model performance. In addition to OLR and ANN methods, several other methods were considered to improve the model's accuracy. Ensemble methods such as Random Forest, Gradient Boosting, or XGBoost can handle unbalanced data to provide more optimal results.

REFERENCES

- [1] Kementerian Kesehatan Republik Indonesia, *Profil Kesehatan Indonesia Tahun 2019*. Jakarta: Kementerian Kesehatan Republik Indonesia, 2019.
- [2] L. C. Sasmita, "PREVENTION OF CHILDHOOD STUNTING PROBLEMS WITH THE MAYANG-WATI PROGRAM," *Jurnal Layanan Masyarakat (Journal of Public Services)*, vol. 5, no. 1, pp. 140–150, May 2021, doi: 10.20473/jlm.v5i1.2021.140-150.
- [3] WHO, *Reducing stunting in children: equity considerations for achieving the Global Nutrition Targets 2025*. WHO, 2018.
- [4] Pusat Data dan Informasi, "Buletin Jendela Data dan Informasi Kesehatan : Semester I, 2018," Jakarta, 2018.
- [5] Kementerian Kesehatan Republik Indonesia, "Status Gizi SSGI 2022," Jakarta, 2022.
- [6] Y. G. Putra, Y. Yuhana, C. Fafilaya, E. Kurniatin, U. Sultan, and A. Tirtayasa, "Facilitation of Non-Formal Education for Families at Risk of Stunting Through Mobilization of Family Planning Instructors," *Jurnal Pemikiran Administrasi Negara*, vol. 15, no. 2, pp. 534–542, 2023.
- [7] Badan Kependudukan dan Keluarga Berencana Nasional, "Laporan Percepatan Penurunan Stunting Tahun 2022 dan Rencana Aksi Tahun 2023," Jakarta, 2023.
- [8] M. De Onis *et al.*, "Prevalence thresholds for wasting, overweight and stunting in children under 5 years," *Public Health Nutr*, vol. 22, no. 1, pp. 175–179, Jan. 2019, doi: 10.1017/S1368980018002434.
- [9] R. Qamar and B. A. Zardari, "Artificial Neural Networks: An Overview," *Mesopotamian journal of Computer Science*, vol. 2023, pp. 130–139, 2023.
- [10] J. Han, M. Kamber, J. Pei, and M. Kaufmann, "[DATA MINING: CONCEPTS AND TECHNIQUES 3RD EDITION] 2 Data Mining: Concepts and Techniques Third Edition."
- [11] N. R. Panda, J. K. Pati, J. N. Mohanty, and R. Bhuyan, "A Review on Logistic Regression in Medical Research," Apr. 01, 2022, *MedSci Publications*. doi: 10.55489/njcm.134202222.
- [12] D. G. Kleinbaum and M. Klein, *Logistic Regression: A Self-Learning Text, Second Edition*. New York: Springer, 2002.
- [13] D. W. Hosmer, Stanley. Lemeshow, and R. X. Sturdivant, *Applied Logistic Regression Third edition*. New Jersey: John Wiley & Sons, Inc., 2013.
- [14] A. Agresti, *An Introduction to Categorical Data Analysis, 3rd Edition*. New Jersey: John Wiley & Sons, Inc., 2019. [Online]. Available: <http://www.wiley.com/go/wsp>
- [15] A. Géron, *Hands-on machine learning with Scikit-Learn, Keras, and TensorFlow. Concepts, tools, and techniques to build intelligent systems*. California: O'Reilly Media, Inc., 2019. [Online]. Available: <http://oreilly.com>

- [16] N. Shahid, T. Rappon, and W. Berta, "Applications of artificial neural networks in health care organizational decision-making: A scoping review," *PLoS One*, vol. 14, no. 2, Feb. 2019, doi: 10.1371/journal.pone.0212356.
- [17] F. F. Addini, D. Haryanto, and R. A. Maolani, "Klasifikasi Tingkat Risiko Kerugian Kecelakaan berdasarkan Karakteristik Pengemudi dengan Analisis Regresi Logistik Ordinal," *Jurnal Matematika Integratif*, vol. 18, no. 2, p. 167, Dec. 2022, doi: 10.24198/jmi.v18.n2.41317.167-177.
- [18] R. Kusumawardhani, Z. D. Rizqiena, and S. P. Astuti, *Ekonometrika*. Yogyakarta: CV Gerbang Media Aksara, 2021.
- [19] A. Agresti, *Categorical data analysis: 3rd Edition*. New Jersey: John Wiley & Sons, Inc, 2013.
- [20] M. Alwi Aliu and A. Fitrianto, "Pemodelan Regresi Logistik Ordinal Backward dengan Imputasi K-Nearest Neighbour pada Indeks Pembangunan Manusia di Indonesia Tahun 2021," *Jurnal Statistika dan Aplikasinya*, vol. 7, no. 1, 2023.
- [21] S. S. Haykin, *Neural networks and learning machines*. Prentice Hall/Pearson, 2009.
- [22] D. Galih Pradana, M. L. Alghifari, M. Farhan Juna, and S. Dwisiwi Palaguna, "Klasifikasi Penyakit Jantung Menggunakan Metode Artificial Neural Network," *Indonesian Journal of Data and Science (IJODAS)*, vol. 3, no. 2, pp. 55–60, 2022.
- [23] O. A. Montesinos López, A. Montesinos López, and J. Crossa, *Multivariate Statistical Machine Learning Methods for Genomic Prediction*. Cham: Springer International Publishing, 2022. doi: 10.1007/978-3-030-89010-0.
- [24] R. M. Tharsanee, R. S. Soundariya, A. S. Kumar, M. Karthiga, and S. Sountharajan, "Deep convolutional neural network-based image classification for COVID-19 diagnosis," in *Data Science for COVID-19 Volume 1: Computational Perspectives*, Elsevier, 2021, pp. 117–145. doi: 10.1016/B978-0-12-824536-1.00012-5.
- [25] D. P. Kingma and J. Lei Ba, "ADAM: A METHOD FOR STOCHASTIC OPTIMIZATION."
- [26] A. C. Rencher, *Methods of Multivariate Analysis Second Edition*. John Willey & Sons Inc., 2002.



© 2025 by the authors. This work is licensed under a Creative Commons Attribution-ShareAlike 4.0 International License (<http://creativecommons.org/licenses/by-sa/4.0/>).