

ESG Factors and Interest Rate Stability: A Comparative Analysis of Statistical and Machine Learning Approaches

Khusnia Nurul Khikmah^{1*} and A'yunin Sofro²

¹Department of Mathematics, Faculty of Mathematics and Natural Sciences, Universitas Palangka Raya, Palangka Raya, Central Kalimantan, 74874 Indonesia

²Department of Actuarial Science, Faculty of Mathematics and Natural Sciences, Universitas Negeri Surabaya, Surabaya, East Java, 60231 Indonesia

Received: 11 March 2026

Revised: 8 May 2026

Accepted: 12 June 2026

*Corresponding author: khusniank@gmail.com

ABSTRACT – This study aims to compare the performance of basic statistical approaches and machine learning in identifying environmental, social, and governance (ESG) issues that influence the stability of the rupiah through the Indonesian interest Rate. This study proposes two approaches: machine learning analysis (long short-term memory (LSTM) and multiple long short-term memory (M-LSTM)) and fundamental statistical models (autoregressive integrated moving average (ARIMA), transfer function, and Koyck regression) to forecast the Indonesian interest rate time series data based on ESG factors. The analysis uses a completely randomized design to capture complex patterns, utilizing interest rate data, PM2.5 air quality data, and rainfall index data. This study provides novelty through data splitting scenarios in its empirical analysis and a comparative study of approaches to the model. The analysis results indicate that ESG factors, specifically air quality (PM2.5) and the rainfall index, have a significant influence on Indonesian interest rate values, as determined by the transfer function results, Koyck regression, and M-LSTM model analyses. All proposed model approaches demonstrate good forecasting accuracy, empirically proven by MAPE values with MAPE<10%. The MAPE values of the basic statistical models (Koyck regression, transfer function, and ARIMA) are 0.2842%, 1.0350%, and 2.1245%, outperforming machine learning models (LSTM and M-LSTM) with MAPE values of 4.660% and 7.7353%.

Keywords– ESG, climate risk, forecasting, Indonesia interest rate, machine learning.

I. INTRODUCTION

Climate change (environmental), environmental issues (social), and governance (ESG) are the main forces driving global finance. These global challenges require the involvement of various parties, such as investors, companies, regulators, and governments, to implement ESG principles in maintaining the stability of the Rupiah, controlling inflation, and promoting economic growth through the Indonesian Interest Rate. ESG factors are not only a social responsibility that will have a significant impact on finance [1] but also on the environment [2]. However, environmental problems resulting from climate change risks are becoming increasingly worrying and require attention because they have a significant impact on the survival of humans and future generations [3]. As a result of these conditions, structured steps are necessary to prevent issues for all levels of society. Another impact that can be caused by increasingly frequent disasters due to climate change and the transition to a low-carbon economy is its fundamental influence on the national economy, such as inflation and the overall stability of the Rupiah [4]. Fluctuations in the value of the Rupiah due to climate risk pose a complex challenge. The dynamic national economy, influenced by various factors such as macroeconomic conditions, geopolitical events, investor sentiment, and incomplete, unstructured ESG data, makes it challenging to integrate data into models for analysis. Therefore, a model capable of capturing data volatility and predicting Indonesia's Interest Rate values due to ESG factors is crucial.

Developments in statistics, particularly through the latest approaches with high levels of data complexity, offer new opportunities that need to be studied to address these issues. A renewed approach with machine learning statistics in modelling time series data, such as neural networks and their developments, as has been done in [5], [6], [7], which have not only one input but multiple inputs, has proven to be effective in modelling non-linear, complex, and noisy time series data. The approach proposed in this study, which considers the volatility history of previous time series data, aims to identify the risks associated with climate change and obtain benefits. In addition, volatility is the most important consideration in this study because it can estimate the potential future losses from the Indonesian interest Rate. Therefore, this study proposes two strategic approaches, long short-term memory (LSTM) and multiple long short-term memory (M-LSTM), for the machine learning statistical approach and the basic model approach, autoregressive integrated moving average (ARIMA), transfer function, and Koyck regression. The proposed model will be analysed and tested for significance using a completely randomised design to capture complex patterns in historical Indonesian interest Rate data and ESG factors, thereby enhancing the accuracy of Indonesian interest Rate forecasting.

This study divides the research objectives into two specific and general objectives, taking into account the contributions generated. The specific objectives are to determine the significance of the relationship between ESG factors and Indonesia's Interest Rate fluctuations, and to compare the performance of machine learning statistics and basic model

statistics in predicting Indonesia's interest rate. Moreover, to determine the accuracy of ESG factor forecasts in relation to Indonesia's interest rate fluctuations, based on the mean absolute percentage error (MAPE) accuracy measure. The general objective is to develop an ESG-based Indonesia interest rate forecasting model with several proposed approaches, with the best accuracy based on the analysis results. In general, this study presents a novel approach to empirical studies, integrating multiple inputs to compare forecasts between basic model statistics and machine learning statistics in the context of ESG factors and Indonesia's interest rate performance. Leveraging the power of the proposed approach to forecast Indonesia's interest rate performance in the context of climate risk, this study integrates ESG data into the approach, utilising multiple inputs. Thus, the significance of the relationship between the two can be clearly determined through the model. Therefore, this research is fundamental because business actors require accurate and relevant information to make informed decisions about the appropriate national economic conditions, particularly regarding how ESG factors impact Indonesia's Interest Rate performance, thereby minimising climate risk. In addition, regulators also need effective and accurate tools to monitor and manage systemic risks due to climate change that affect the national economic sector.

II. LITERATURE REVIEW

A. Koyck Regression

Koyck regression, also known as distributed-lag regression, is a basic statistical forecasting method that helps illustrate the relationship between dependent and independent variables distributed over a specific time period. This method considers the role of time, which separates the prediction of the dependent variable from changes in the value of the independent variable. In general, the Koyck regression model is defined as follows [8].

$$Y_t = \alpha + \sum_{i=0}^p \beta_i X_{t-i} + e_t$$

$$\beta_i = a_0 + a_1 i + a_2 i^2 + \dots + a_m i^m$$

$$Y_t = \alpha + \sum_{i=0}^p (a_0 + a_1 i + a_2 i^2 + \dots + a_m i^m) \beta_i X_{t-i} + e_t$$

Where Y_t is the value of the dependent variable at time t , α is the intercept, β_i is the coefficient of the i^{th} independent variable, X_{t-i} is the independent variable at time $t - i$, a_m is the coefficient of the independent variable with the m^{th} degree of polynomial, and e_t is the error at time t .

B. Autoregressive Integrated Moving Average (ARIMA)

Autoregressive integrated moving average, or ARIMA, is a basic statistical forecasting method for time series data, also known as the Box-Jenkins method. This method uses past and present values to generate predictions. In general, the ARIMA model can be denoted as $ARIMA(p, d, q)$, where p represents the order of the autoregressive process, q represents the order of the moving average process, and d represents the order of the differencing process. Mathematically, the $ARIMA(p, d, q)$ model is defined as follows [9], [10].

$$\phi_p(1 - B)^d Y_t = \mu + \theta_q(B) e_t$$

ϕ_p is the coefficient of the autoregressive component with order p , θ_q is the coefficient of the moving average component with order q , d is the order of the differencing, B is the backshift operator, Y_t is the time series data at time t , and e_t is the residual at time t .

C. Transfer Function

The transfer function, also known as autoregressive integrated moving average with exogenous (ARIMAX), is an extended forecasting method of ARIMA where the predicted value of the output time series data Y_t is based on the past values of the input time series data X_t . If $\omega_s(B) = \omega_0 - \omega_1 B - \omega_2 B^2 - \dots - \omega_s B^s$ and $\delta_r(B) = 1 - \delta_1 B - \delta_2 B^2 - \dots - \delta_r B^r$, then mathematically the transfer function equation can be written as follows [11], [12].

$$Y_t = \frac{\omega_s(B)}{\delta_r(B)} x_{t-b} + \frac{\theta_q(B)}{\phi_p(B)} a_t$$

b is the lag recorded in x_{t-b} , s is the length of time the input series influences the output series, and r is the length of time the output series at time t is influenced by the output series at time $(k - t)$. A comprehensive explanation of the transfer function analysis stages can be found in [10].

D. Long Short-Term Memory (LSTM)

Long short-term memory (LSTM) is a statistical machine learning method for forecasting that extends recurrent neural networks (RNNs). This method is capable of producing good forecasts with an architecture consisting of several layers. The first layer is the input layer, followed by the hidden layer, and finally the output layer [13]. First, the input gate is a gate that takes the previous output and the new input. This gate also returns a value of 0 or 1. Suppose the input value is symbolised by in_t , with the sigmoid activation function σ , for the input gate weight matrix W_i , and the bias of the input gate b_i . Then, if Y_t is the input at time t and h_{t-1} is the hidden state at time $t - 1$, the value of in_t can be defined as follows.

$$in_t = \sigma(W_i Y_t + W_i h_{t-1} + b_i)$$

For a memory cell that has a candidate value $\bar{C}_t$ with input gate weights from memory cell W_c , where the bias of the memory cell is b_c , and h_{t-1} is the state at time $t - 1$. The candidate value can then be defined as follows.

$$\bar{C}_t = \tanh(W_c Y_t + W_c h_{t-1} + b_c)$$

Second, the forget gate takes input at time t and output at time $(t - 1)$, which then uses the sigmoid activation function σ . Suppose the forget gate is symbolised by f_t , with the weight matrix for the forget gate W_f , and the bias for the forget gate b_f . Then, if there is input at time t symbolised by Y_t and h_{t-1} is the hidden state at time $t - 1$, the value of the forget gate can be defined as follows.

$$f_t = \sigma(W_f Y_t + W_f h_{t-1} + b_f)$$

With the state value of the memory cell updated from the previous process, denoted as C_t , defined as follows.

$$C_t = in_t * \bar{C}_t + f_t * C_{t-1}$$

Where in_t is the input gate at time t , $\bar{C}_t$ is the candidate value of the memory cell, f_t is the forget gate at time t , and C_{t-1} is the cell state at time t . Finally, the output gate. This gate controls the number of states that pass through. If we start from the output gate symbolised by o_t , where σ is the sigmoid activation function with the weight matrix for the output gate W_o and b_o is the bias for the output gate. Then, if there is an input at time t symbolised by Y_t , with h_{t-1} being the hidden state at time $t - 1$ for the state of the updated memory cell C_t and V_o being the weight matrix for the output gate, then the value of the output gate can be defined as follows.

$$o_t = \sigma(W_o Y_t + W_o h_{t-1} + V_o C_t + b_o)$$

Where the final output of the cell is symbolised by h_t , then the value of h_t can be calculated based on the following formula.

$$h_t = o_t * \tanh(C_t)$$

o_t is the value of the output gate and C_t is the state of the updated memory cell [13], [14]. Furthermore, the activation function itself is used in LSTM to generate the output of each layer. Mathematically, in the explanation above, the activation functions, specifically the sigmoid and tanh activation functions, are defined as follows.

$$\sigma = \frac{1}{1 + e^{-x}}; \tanh(x) = 2\sigma(2x) - 1$$

x is the input data, σ is the value of the sigmoid activation function, $\tanh$ is the value of the tanh activation function, and e is the exponential value [15]. In general, the architecture of LSTM is shown [16].

E. Multiple Long Short-Term Memory (M-LSTM)

Multiple long short-term memory (M-LSTM) is a forecasting method in statistical machine learning that extends the LSTM method, differing primarily in the input section. In M-LSTM, multiple inputs are used. However, in general, the architecture of both is the same. Clearly, the distinguishing input between the two can be seen based on previous research (Xu et al. 2020; Miran et al. 2021). The LSTM method with a single input, where the input used is the data at time $t - 1$, symbolised by Y_{t-1} , then in the M-LSTM method or with multiple inputs with data at time $t - 1$, there are three inputs, which can be symbolised as Y_{1t-1} , Y_{1t-2} , Y_{1t-3} . In general, the architecture of M-LSTM is shown in [5].

F. Forecasting Accuracy Measures

Forecasting accuracy measures are calculations used to determine the accuracy level of the forecasting method. This study employs the mean absolute percentage error (MAPE) as a measure of accuracy. Mathematically, the MAPE value can be calculated using the formula proposed by [19].

$$MAPE = \frac{1}{n} \sum_{i=1}^n \left| \frac{e_i}{y_i} \right| \times 100\%$$

The i -th observation data is symbolised by y_i and the estimated value at time i is symbolised by $\hat{y}_i$, where e_i is the difference between y_i and $\hat{y}_i$ or $e_i = y_i - \hat{y}_i$.

G. T-Test

The hypothesis test used in this study is the t-test. This test is one of the tests used to test the difference in the means of two samples. In general, the paired t-test has hypotheses $H_0: \mu_a = A_0$ or $H_1: \mu_a \neq A_0$, with t-test statistics where n is the sample size and n_1 and n_2 are the sample sizes of group 1 and group 2. Then, the t-test statistics can be defined as follows [5].

$$t = \frac{\bar{x}_1 - \bar{x}_2}{s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$$

The variance value $s_p = \sqrt{\frac{(n_1-1)s_1^2 + (n_2-1)s_2^2}{n_1 + n_2 - 2}}$ for s_1^2 and s_2^2 are the sample variances of groups 1 and 2. For $\bar{x}_1$ and $\bar{x}_2$ are the means of groups 1 and 2 with degrees of freedom $d_f = n_1 + n_2 - 2$.

III. METHODOLOGY

The empirical data used in this study are divided into three categories. Indonesia Interest Rate (Y) data is historical monthly Indonesia Interest Rate data from the Bank of Indonesia, ESG data (X_1) is monthly air quality (PM25) data from the Air Quality Index (AQI), and climate risk data (X_2) is rainfall index data sourced from the Indonesian Central Statistics Agency database with analysis steps as below.

1. Data preprocessing by collecting the proposed data for analysis, performing data cleaning, handling missing data: linear interpolation and forward-filling, normalising data according to the model input, and testing data stationarity and data transformation.
2. Data exploration
3. Splitting Data

Table 1 Splitting Data

Data	Training Data	Testing Data
1	80%	20%
2	75%	25%
3	70%	30%

4. Modelling with machine learning statistics by preparing the data input from step (1). Then, determine the model hyperparameters and building a model using training data with LSTM and M-LSTM.
5. Modelling with basic model statistics by preparing input data from step (1) and building a model using training data with ARIMA, transfer function, and Koyck regression.
6. Comparison of analysis results, by comparing the accuracy of model predictions based on analysis results using training data and identifying significant factors.
7. Comparative analysis of model performance: to compare the forecasting accuracy across models, a descriptive comparison of MAPE values was conducted. Additionally, an exploratory analysis of variance (ANOVA) was performed to examine the MAPE values, which are presented as exploratory and descriptive guidance only, not as confirmatory statistical inference.
8. Model validation with test data
9. Interpretation and recommendations

IV. RESULTS AND DISCUSSIONS

This research began with the collection of data related to Indonesia's monthly interest rates from Bank Indonesia, ESG data related to monthly air quality (PM25) data from the Air Quality Index (AQI), and climate risk data related to rainfall index data sourced from the Indonesian Central Statistics Agency database, taken from January 2015 to February 2025. Next, data preprocessing was carried out in accordance with missing data checks, utilising linear interpolation and forward filling methods. The handling of missing data using both proposed methods yielded the same value. Therefore, the method that will be used next is linear interpolation, which will then be used for analysis.

This research considers two study scenarios. First, a comparative study of basic statistical methods and machine learning statistics. Second, a study of data splitting. The data, which had undergone missing data handling, was split into three training and testing data comparison scenarios, as shown in Table 1. Furthermore, both basic statistical methods and machine learning statistics will use the proposed data splitting scenario. The three basic statistical model methods proposed in this study were tested for data stationarity using the Augmented Dickey-Fuller (ADF) test, where the null hypothesis is H_0 : non-stationary data and the alternative hypothesis is H_1 : stationary data. With a significance level of 5%, if the p-value is greater than 0.05, it means that H_0 cannot be rejected, or the data is non-stationary. Meanwhile, if the p-value ≤ 0.05 , it means that H_0 is rejected, or the data is stationary. The results of the ADF test on the Indonesian Interest Rate (Y) data are presented in Table 2.

Table 2 ADF Test Results for Indonesia Interest Rate (Y) Data

Data	Differencing	Value	Decision
1	2	0.01	Reject H_0
2	2	0.01	Reject H_0
3	2	0.01	Reject H_0

The first analysis using this basic statistical model approach, where the data that has undergone the ADF test in Table 2 will be modelled using ARIMA, where the model candidates obtained are based on the values of the autocorrelation function (ACF), partial autocorrelation function (PACF), and extended autocorrelation function (EACF) plots. The ARIMA model candidates obtained are presented in Table 3.

Table 3 ARIMA Model Candidates

Data	Differencing	Value
1.	ARIMA(0,2,2)	-102.18
	ARIMA(2,2,0)	-99.41
	ARIMA(1,2,1)	-103.85
	ARIMA(2,2,2)	-100.02
2.	ARIMA(0,2,1)	-87.65
	ARIMA(2,2,0)	-87.48
	ARIMA(1,2,1)	-90.7
	ARIMA(2,2,1)	-89.21
3.	ARIMA(0,2,1)	-79.29
	ARIMA(2,2,0)	-79.53
	ARIMA(1,2,1)	-81.48
	ARIMA(2,2,1)	-79.57

The best model, selected from the model candidates, is the one with the smallest AIC value. Therefore, in each proposed data division scenario, the best model obtained is the same, namely *ARIMA*(1,2,1). Next, each model undergoes parameter estimation and model diagnostics. The formality tests performed in this model diagnostic process are normality tests, residual mean tests, and autocorrelation tests. For normality tests with a significance level of $\alpha = 5\%$, if the p -value is greater than 0.05, it means that H_0 cannot be rejected, indicating that the residuals are normally distributed. If the p -value < 0.05 , it means that H_0 is rejected or the residuals are not normally distributed. The residual mean test was conducted with a significance level of $\alpha = 5\%$. If the p -value is greater than 0.05, it means that H_0 cannot be rejected, indicating that the residual mean is equal to zero. If the p -value < 0.05 , it means that H_0 is rejected or the residual mean is not equal to zero. The last, the autocorrelation test was conducted with a significance level of $\alpha = 5\%$, where a p -value greater than 0.05 indicates that H_0 cannot be rejected, suggesting no signs of autocorrelation. If p -value < 0.05 , it means that H_0 is rejected or there are signs of autocorrelation. The complete diagnostic results of this model are presented in Table 4.

Table 4 Estimation of Parameters

Data	Parameter Estimation	
	Estimated Value	p-value
1	$\phi_1 = 0.5799501$	$6.96 \times 10^{-12*}$
	$\theta_1 = 0.9999981$	$2.22 \times 10^{-16*}$
2	$\phi_1 = 0.5839122$	$7.05 \times 10^{-11*}$
	$\theta_1 = 0.9999987$	$2.22 \times 10^{-16*}$
3	$\phi_1 = 0.407994$	$2.70 \times 10^{-2*}$
	$\theta_1 = 0.846132$	$2.11 \times 10^{-12*}$

Table 5 Diagnostics Model of Table 4

Data	Parameter Estimation					
	Normality Test**		Mean Residual Test***		Autocorrelation Test****	
	p-val	Decision	p-val	Decision	p-val	Decision
1	1.44×10^{-13}	Reject H_0	9.39×10^{-1}	Do Not Reject H_0	9.47×10^{-1}	Do Not Reject H_0
2	1.35×10^{-12}	Reject H_0	9.31×10^{-1}	Do Not Reject H_0	9.79×10^{-1}	Do Not Reject H_0
3	5.99×10^{-12}	Reject H_0	6.56×10^{-1}	Do Not Reject H_0	9.83×10^{-1}	Do Not Reject H_0

Note: * Significant at $\alpha = 5\%$ in the dataset. **Kolmogorov-Smirnov test. *** t-test. ****Ljung-Box test.

Based on Table 4-5, the estimated values for each data division scenario are statistically significant at $\alpha = 0.05$. Additionally, the model's diagnostic results confirm that the mean residual value is equal to 0 and there are no signs of autocorrelation. Therefore, this model is suitable for forecasting. Mathematically, the best model obtained using the backshift operator decomposition for the *ARIMA* model is following this equation.

$$\text{Data 1: } Y_t = \mu + 2.5799501Y_{t-1} - 2.1599002Y_{t-2} + 0.5799501Y_{t-3} + e_t - 0.9999981e_{t-1}$$

$$\text{Data 2: } Y_t = \mu + 2.5839122Y_{t-1} - 2.1678244Y_{t-2} + 0.5839122Y_{t-3} + e_t - 0.9999987e_{t-1}$$

$$\text{Data 3: } Y_t = \mu + 2.407994Y_{t-1} - 1.815988Y_{t-2} + 0.407994Y_{t-3} + e_t - 0.846132e_{t-1}$$

The equation above can then be used to forecast the training data, which is subsequently applied to the test data. The goodness of this research model is evaluated using the MAPE measure. The accuracy of the *ARIMA* model's forecast results, validated on the test data in this study, is shown in Table 6.

Table 6 MAPE of ARIMA Model Forecasts

Data	MAPE (%)
1	2.4439
2	1.9741
3	1.9556

The findings presented in Table 6 indicate that forecast accuracy increases as the amount of training data used decreases. Based on previous research [19], MAPE accuracy values can be categorised into four categories, with MAPE < 10% indicating an excellent model forecast. Therefore, using three forecast data scenarios, the *ARIMA*(1,2,1) model categorises the Indonesian interest Rate data as excellent.

In the second analysis, using the basic statistical model approach, the data tested using the ADF (Augmented Dickey-Fuller) in Table 2 will be modelled using a transfer function, or *ARIMAX*. In accordance with the stages in modelling a transfer function, the first input variable, X_{1t} , in this study relates to air quality data (pm25), and the output variable, Y_t , is the Indonesian interest Rate data. In the second analysis, the second input variable, X_{2t} , is rainfall index data, and the output variable, Y_t , is the Indonesian Interest Rate data. The input data from X_{1t} and X_{2t} will first be modelled using *ARIMA*. The missing data handling results are divided into the same proportions as in the standard *ARIMA* model, with training data and test data ratios of 80:20, 75:25, and 70:30. Next, the training data is tested for stationarity using the Augmented Dickey-Fuller (ADF) test, where the hypothesis is that H_0 is non-stationary and H_1 is stationary. With a significance level of 5%, a p-value greater than 0.05 fails to reject the null hypothesis (H_0), indicating non-stationary data. A $p\text{-value} \leq 0.05$ rejects H_0 , indicating stationary data. The results of the ADF test on the input data X_{1t} and X_{2t} are Rejecting H_0 with all of the p-value of first differencing is 0.01.

The input data X_t , which has undergone the ADF test in Table 2, will be modelled using *ARIMA*. The candidate models obtained are based on the values of the autocorrelation function (ACF), partial autocorrelation function (PACF), and extended autocorrelation function (EACF) plots. The candidate *ARIMA* models obtained are presented in Table 7.

Table 7 AIC of *ARIMA* Candidates Model

Data	Model	AIC
X_{1t}	1 <i>ARIMA</i> (0,1,1)	894.044
	<i>ARIMA</i> (1,1,1)	895.367
	<i>ARIMA</i>(1, 1, 0)	893.703
	<i>ARIMA</i> (2,1,0)	894.963
	2 <i>ARIMA</i> (0,1,1)	835.691
	<i>ARIMA</i> (1,1,1)	837.026
	<i>ARIMA</i>(1, 1, 0)	835.345
	<i>ARIMA</i> (2,1,0)	836.658
	3 <i>ARIMA</i> (0,1,1)	783.886
	<i>ARIMA</i> (1,1,1)	785.455
	<i>ARIMA</i>(1, 1, 0)	783.629
	<i>ARIMA</i> (2,1,0)	785.144
X_{2t}	1 <i>ARIMA</i>(0. 1. 1)	1600.712
	<i>ARIMA</i> (1,1,1)	1602.543
	<i>ARIMA</i> (1,1,0)	1639.593
	<i>ARIMA</i> (2,1,0)	1630.408
	2 <i>ARIMA</i>(0. 1. 1)	1492.310
	<i>ARIMA</i> (1,1,1)	1494.144
	<i>ARIMA</i> (1,1,0)	1528.341
	<i>ARIMA</i> (2,1,0)	1519.970
	3 <i>ARIMA</i>(0. 1. 1)	1399.172
	<i>ARIMA</i> (1,1,1)	1401.014
	<i>ARIMA</i> (1,1,0)	1432.610
	<i>ARIMA</i> (2,1,0)	1424.948

The best model for the input X_t , selected from the candidate models, is the one with the smallest AIC value. Therefore, in each data distribution scenario for the two proposed inputs, the best models obtained were the same: *ARIMA*(1,1,0) for X_{1t} and *ARIMA*(0,1,1) for X_{2t} . Next, parameter estimation and model diagnostics were performed for each model. The formal tests performed in this model diagnostic process were the normality test, the residual mean test, and the autocorrelation test. The hypotheses used in the previous *ARIMA* model explanation were used.

Table 8 Estimation of Parameters

Data	Parameter Estimation	
	Estimated Value	p-value
X_{1t}	1 $\phi_1 = -0.14566$	1.48×10^{-1}
	2 $\phi_1 = -0.151681$	1.47×10^{-1}
	3 $\phi_1 = -0.14775$	1.72×10^{-1}
X_{2t}	1 $\theta_1 = 1.00$	$2 \times 10^{-16*}$
	2 $\theta_1 = 0.99$	$2 \times 10^{-16*}$
	3 $\theta_1 = 0.99$	$2 \times 10^{-16*}$

Table 9 Diagnostics Model of Table 8

Data		Parameter Estimation					
		Normality Test**		Mean Residual Test***		Autocorrelation Test****	
		p-val	Decision	p-val	Decision	p-val	Decision
X_{1t}	1	2.99×10^{-14}	Reject H_0	9.85×10^{-1}	Do Not Reject H_0	0.51×10^{-2}	Do Not Reject H_0
	2	3.36×10^{-15}	Reject H_0	9.85×10^{-1}	Do Not Reject H_0	1.13×10^{-1}	Do Not Reject H_0
	3	2.22×10^{-16}	Reject H_0	9.85×10^{-1}	Do Not Reject H_0	1.01×10^{-1}	Do Not Reject H_0
X_{2t}	1	5.55×10^{-16}	Reject H_0	9.96×10^{-1}	Do Not Reject H_0	1.00	Do Not Reject H_0
	2	6.66×10^{-16}	Reject H_0	9.51×10^{-1}	Do Not Reject H_0	1.00	Do Not Reject H_0
	3	2.20×10^{-16}	Reject H_0	9.30×10^{-1}	Do Not Reject H_0	1.00	Do Not Reject H_0

Note: * Significant at $\alpha = 5\%$ in the dataset. **Kolmogorov-Smirnov test. *** t-test. ****Ljung-Box test.

The estimated values presented in Table 8 for each input and data sharing scenario, for input X_{2t} , are significant at $\alpha = 5\%$. Furthermore, the diagnostic results (see Table 9) for all model inputs meet the assumption of a residual mean of 0 and no autocorrelation. Therefore, this model is suitable for use. Next, the pre-whitening of input X_t is converted to α_t . This pre-whitening value of input α_t is applied to output Y_t , denoted by β_t . Mathematically, the best model obtained using the backshift operator is $ARIMA(1,1,0)$ for X_{1t} and $ARIMA(0,1,1)$ for X_{2t} with α_t is input and β_t is output are presented following this equation.

For X_{1t} Data 1.

$$\begin{aligned}\alpha_t &= \mu - X_t + 0.85434X_{t-1} + 0.14566X_{t-2} + e_t \\ \beta_t &= \mu - Y_t + 0.85434Y_{t-1} + 0.14566Y_{t-2} + e_t\end{aligned}$$

For X_{1t} Data 2.

$$\begin{aligned}\alpha_t &= \mu - X_t + 0.848319X_{t-1} + 0.151681X_{t-2} + e_t \\ \beta_t &= \mu - Y_t + 0.848319Y_{t-1} + 0.151681Y_{t-2} + e_t\end{aligned}$$

For X_{1t} Data 3.

$$\begin{aligned}\alpha_t &= \mu - X_t + 0.85225X_{t-1} + 0.14775X_{t-2} + e_t \\ \beta_t &= \mu - Y_t + 0.85225Y_{t-1} + 0.14775Y_{t-2} + e_t\end{aligned}$$

For X_{2t} Data 1.

$$\begin{aligned}\alpha_t &= \mu - X_t + X_{t-1} + e_t - e_{t-1} \\ \beta_t &= \mu - Y_t + Y_{t-1} + e_t - e_{t-1}\end{aligned}$$

For X_{2t} Data 2.

$$\begin{aligned}\alpha_t &= \mu - X_t + X_{t-1} + e_t - 0.99e_{t-1} \\ \beta_t &= \mu - Y_t + Y_{t-1} + e_t - 0.99e_{t-1}\end{aligned}$$

For X_{2t} Data 3.

$$\begin{aligned}\alpha_t &= \mu - X_t + X_{t-1} + e_t - 0.99e_{t-1} \\ \beta_t &= \mu - Y_t + Y_{t-1} + e_t - 0.99e_{t-1}\end{aligned}$$

The pre-whitening results of the input α_t and output β_t are calculated based on the values from the cross-correlation plot. Based on the cross-correlation plot, both inputs are significant at the first lag, resulting in a value of $b = 1$, indicating that a one-month delay in X_t affects the value of Y_t . The selection of order $r = 2$ is based on the fit of the cross-correlation plot to the impulse function with $s = 0$. The residual values of b , s , and r are then checked based on the ACF, PACF, and EACF values to obtain the order (p_n, q_n) . The obtained transfer function model has the order $(b, r, s) = (1, 2, 0)$ and the noise series model $(p_n, q_n) = (1, 0)$. The obtained transfer function model is followed this equation.

For X_{1t} Data 1.

$$Y_t = \frac{1}{1 + 0.4262B + 0.5490B^2}x_{t-1} + \frac{1}{1 - 0.9994B}a_t$$

For X_{1t} Data 2.

$$Y_t = \frac{1}{1 + 0.9605B - 0.1053B^2}x_{t-1} + \frac{1}{1 - 0.9994B}a_t$$

For X_{1t} Data 3.

$$Y_t = \frac{1}{1 + 0.35017B - 0.9993B^2}x_{t-1} + \frac{1}{1 - 0.9993B}a_t$$

For X_{2t} Data 1.

$$Y_t = \frac{1}{1 + 0.8106B + 0.2322B^2}x_{t-1} + \frac{1}{1 - B}a_t$$

For X_{2t} Data 2.

$$Y_t = \frac{1}{1 + 0.7084B + 0.3293B^2}x_{t-1} + \frac{1}{1 - 0.9991B}a_t$$

For X_{2t} Data 3.

$$Y_t = \frac{1}{1 + 1.3116B + -0.2856B^2}x_{t-1} + \frac{1}{1 - 0.09993B}a_t$$

The equations above can then be used to forecast training data, which is then applied to the test data. The MAPE measure is used to evaluate the goodness of fit of the research model. The accuracy of the transfer function model forecasts validated on the test data in this study is shown in Table 10.

Table 10 MAPE of Transfer Function Model Forecasts

Data	MAPE (%)
1	0.12319
X_{1t} 2	0.13588
3	0.14244
1	1.54733
X_{2t} 2	1.76314
3	2.49827

The findings presented in Table 8 indicate that forecast accuracy decreases as the amount of training data used decreases. Based on previous research [19], MAPE accuracy can be categorised into four categories, with MAPE <10% indicating excellent model forecasting. Therefore, using three forecast data scenarios, the model's input X_{1t} is air quality data (PM2.5) and its input X_{2t} is rainfall index data, with the output variable Y_t being the Indonesian interest rate data, which is categorised as excellent.

The third analysis, employing a basic statistical model approach in this study, is Koyck regression. This method is proposed to examine the influence of predictor variables, namely air quality data (PM2.5) and rainfall index data, on the dependent variable, the Indonesian Interest Rate in Indonesia. Mathematically, the Koyck regression model, obtained using the three data distribution scenarios, is employed for parameter estimation and model diagnostics. With a significance level of $\alpha = 5\%$, this model diagnostic will reject H_0 if the $p - value < \alpha$, meaning there is sufficient statistical evidence to conclude that there is a significant relationship between the independent variable (X_1) or (X_2) and the dependent variable (Y). If the $p - value > \alpha$, it will accept H_0 , meaning there is insufficient statistical evidence to conclude that there is a significant relationship between the independent variable (X_1) or (X_2) and the dependent variable (Y).

Table 11 Parameter Estimates for the Koyck Regression Model

Data	Estimated Parameter	
	Estimated Value	p-value
1	$\alpha(1 - \lambda) = 0.0667947$	4.40×10^{-1}
	$\beta_0 = 0.9957643$	$2.00 \times 10^{-16*}$
	$\lambda = -0.0006773$	3.58×10^{-1}
X_1 2	$\alpha(1 - \lambda) = 0.0601214$	5.00×10^{-1}
	$\beta_0 = 0.9965110$	$2.00 \times 10^{-16*}$
	$\lambda = -0.0006656$	4.03×10^{-1}
3	$\alpha(1 - \lambda) = 0.0535969$	5.63×10^{-1}
	$\beta_0 = 0.9970804$	$2.00 \times 10^{-16*}$
	$\lambda = -0.0006412$	4.76×10^{-1}
1	$\alpha(1 - \lambda) = 0.2061418$	5.58×10^{-1}
	$\beta_0 = 0.9318711$	$1.86 \times 10^{-11*}$
	$\lambda = 0.0003710$	6.68×10^{-1}
X_2 2	$\alpha(1 - \lambda) = 0.1824076$	5.77×10^{-1}
	$\beta_0 = 0.9373913$	$2.16 \times 10^{-12*}$
	$\lambda = 0.0003425$	6.76×10^{-1}
3	$\alpha(1 - \lambda) = 0.1580410$	6.17×10^{-1}
	$\beta_0 = 0.9416830$	$1.64 \times 10^{-12*}$
	$\lambda = 0.0003363$	6.85×10^{-1}

Table 12 Diagnostics Model of Transfer Function

Data	Coefficient of Determination		Hypothesis Testing**	
	R^2	Adjusted R^2	p-value	Decision
1	0.9723	0.9718	2.22×10^{-16}	Reject H_0
X_{1t} 2	0.9730	0.9724	2.22×10^{-16}	Reject H_0
3	0.9730	0.9723	2.22×10^{-16}	Reject H_0
1	0.8643	0.8614	2.22×10^{-16}	Reject H_0
X_{2t} 2	0.8817	0.8790	2.22×10^{-16}	Reject H_0
3	0.8833	0.8804	2.22×10^{-16}	Reject H_0

Note: * significant at $\alpha = 5\%$. ** Wald Test

Based on Table 11-12 above, the estimated values obtained indicate an excellent coefficient of determination, supported by the test that the proposed independent variables significantly influence the dependent variable, the Indonesian Interest Rate. Furthermore, the goodness-of-fit model, based on the R^2 and adj R^2 values, shows values greater than > 0.85 . This means that more than 85% of the variation in the Indonesian Interest Rate can be explained by air quality (PM2.5) and the rainfall index included in the model. Therefore, mathematically, the obtained Koyck regression model has the equation presented below.

For X_1 Data 1.

$$Y_t = 0.0667947 + 0.9957643X_t - 0.0006773Y_{t-1} + e_t$$

For X_1 Data 2.

$$Y_t = 0.0601214 + 0.9965110X_t - 0.0006656Y_{t-1} + e_t$$

For X_1 Data 3.

$$Y_t = 0.0535969 + 0.9970804X_t - 0.0006412Y_{t-1} + e_t$$

For X_2 Data 1.

$$Y_t = 0.2061418 + 0.9318711X_t + 0.0003710Y_{t-1} + e_t$$

For X_2 Data 2.

$$Y_t = 0.1824076 + 0.9373913X_t + 0.0003425Y_{t-1} + e_t$$

For X_2 Data 3.

$$Y_t = 0.1580410 + 0.9416830X_t + 0.0003363Y_{t-1} + e_t$$

Based on the equation above, the model can then be used to forecast training data, which is then applied to test data. The model's goodness-of-fit (MAPE) is used to evaluate the model's goodness-of-fit in this study. The forecast accuracy of the Koyck regression model for training and test data in this study is shown in Table 13.

Table 13 MAPE of Koyck Regression Model Forecasts

Data	MAPE (%)
X_1	1 0.244
	2 0.226
	3 0.219
X_2	1 0.357
	2 0.328
	3 0.331

Table 13 shows that forecast accuracy increases with decreasing training data, a trend similar to that observed in ARIMA forecasting. Previous research [19] found that MAPE accuracy can be categorised into four categories, with MAPE <10% indicating superior forecasting. Therefore, using three forecast data scenarios for the Koyck regression model, the Indonesian Interest Rate data, based on both air quality data (PM25PM2.5) and rainfall index data, is categorised as excellent.

The second approach is through machine learning statistics, where the first analysis in this section uses LSTM. Modelling with LSTM in practice does not have standard rules, such as basic model statistics, for determining its hyperparameter values. Generally, hyperparameters in LSTM include the number of epochs, batch size, and number of neurons or unit. If Khikmah et al. (2025) proposed to conduct a one-by-one analysis and then test the hypothesis for the significance of the difference in the prediction results with ANOVA, this study proposes to conduct further research by tuning hyperparameters using grid search. This method the proposed parameter combination by evaluating all combinations of hyperparameter values using cross-validation techniques on the training data to identify the hyperparameter combination that yields the best accuracy. Furthermore, this value is referred to as hyperparameter initialisation. The hyperparameter initialisation to be tuned on the training data is presented in full in Table 14.

Table 14 Hyperparameter Initiaization of the LSTM Architecture

Hyperparameter	Specification Value
Neuron/ Unit	4; 20; 50
Batch Size	10; 20; 32; 40; 100
Epoch	10; 50; 100; 200
LSTM layer	1 LSTM Layer
Cross validation	3
Number of dense layer	1 Dense Layer
Activation Function	Relu
Optimizer	Adam

Research by [20], [21], [22] shows that the number of neurons or unit, batch size, and epochs affect the forecasting power of LSTMs. Therefore, the hyperparameter tuning applied in Table 12 in this study was applied to the Indonesian Interest Rate training data in each normalised data sharing scenario. The best model hyperparameter combination obtained based on grid search for each data sharing scenario is presented in Table 15.

Table 15 Hyperparameter Tuning Results with Grid Search

Data	Hyperparameter		
	Neuron/ Unit	Batch Size	Epoch
1	20	20	200
2	20	20	200
3	50	20	200

The values obtained in Table 15 were then used for modelling and forecasting on the Indonesian Interest Rate training data, where the proposed empirical data forecasting results were compared based on MAPE accuracy. Based on the forecasting results validated using the test data, the error for the 80:20 data was 4.264%, followed by the 75:25 data with a forecast error of 4.551%, and the 70:30 data with a forecast error of 5.165%. See Table 16. These results indicate that the forecasting error increases with decreasing training data, consistent with the forecasting results using the transfer function.

Table 16 MAPE of LSTM Model Forecasts

Data	MAPE (%)
1	4.264
2	4.551
3	5.165

The second analysis, based on the statistical machine learning approach, uses M-LSTM. This method, an extension of LSTM, similarly lacks standard rules for hyperparameter values in practice. Therefore, similar to LSTM, M-LSTM employs the same hyperparameter initialisation criteria as those in Tables 14-15. The hyperparameter specifications used by the M-LSTM method with LSTM are then applied to modelling and forecasting the Indonesian interest Rate training data, combined with air quality data (PM2.5) and the PM25 rainfall index data. The forecast comparison criterion used is MAPE. Validated forecast results on the test data show an error of 6.902% for the 80:20 data, 7.703% for the 75:25 data, and 8.601% for the 70:30 data (see Table 17). Similar to LSTM, M-LSTM has an increasing forecast error as the training data decreases.

Table 17 MAPE of M-LSTM Model Forecasts

Data	MAPE (%)
1	6.902
2	7.703
3	8.601

Next, this assumption is tested by hypothesis testing the forecast MAPE values using analysis of variance. This analysis aims to detect MAPE values based on the levels of each factor and their interactions. This hypothesis testing was conducted at a significance level of $\alpha = 0.05$, as fully presented in Table 16.

Table 18 Results of The Analysis of Variance

Source of Variance	DF	SS	MS	F-count	p-val
Approach	1	117.90	117.90	2860.863	3.49×10^{-4}
Data Splitting	2	0.70	0.35	8.553	1.05×10^{-1}
Model	4	19.32	4.83	117.197	8.48×10^{-3}
Approach*Data Splitting	2	1.05	0.53	12.768	7.26×10^{-2}
Data Splitting*Model	8	0.68	0.09	2.063	3.67×10^{-1}
Approach *Model	1	6.60	6.60	160.129	6.19×10^{-3}
Residual	2	0.08	0.04		

Note: *) significant at $\alpha=5\%$; *) interaction

To explore the sources of variation in MAPE values, a three-factor analysis of variance (ANOVA) was conducted with factors, approach, data splitting, and model. The results, presented in Table 18, should be interpreted cautiously due to the limited degrees of freedom (only two residual degrees of freedom). As such, these findings are exploratory rather than confirmatory. Furthermore, this empirical study also proposed a t-test hypothesis test. With a significance level of $\alpha = 5\%$, the p-value of the t-test obtained was 3.468×10^{-8} (p-value < α), meaning H_0 is rejected. This test is shown as exploratory descriptive statistics only. They should not be overinterpreted as confirmatory hypothesis tests and the p-values are shown for transparency but lack sufficient statistical power for definitive conclusions. Instead, the primary comparison of model performance is based on descriptive MAPE values (see Table 19).

Table 19 MAPE of All Model Forecasts

Approach	Model	Best MAPE Value
Statistical Fundamental	ARIMA	2.1245
	Transfer Function	1.0350
	Koyck Regression	0.2842
Machine Learning	LSTM	4.6600
	M-LSTM	7.7353

The MAPE values summarized in Table 19 above indicate that the fundamental model statistics have a smaller MAPE value than the machine learning statistics, indicating that the baseline model statistics have better accuracy than the machine learning statistics. According to [19], they fall into the category of good forecast accuracy (MAPE < 10%). This aligns with the findings of research [23] that the baseline model statistics are suitable for use when the data used is time series data with a specific time period or not big data. However, the forecast accuracy results of the machine learning statistics are still in the good category (MAPE < 10%). In theory, machine learning statistics should have high accuracy when used for big data forecasting due to the learning outcomes achieved with large datasets. The results of the forecast accuracy with limited time period data from the machine learning statistics, which fall into the good category, also align with the results of a thesis study by [24], where LSTM and M-LSTM have good forecast accuracy on data with a limited number of observations ($n < 30$). The analytical results obtained in this study have significant research strengths, providing important insights into national economic interventions aimed at increasing awareness of the Indonesian Interest Rate. This study also highlights the significant influence of independent variables on the dependent variable (the Indonesian interest rate).

Furthermore, this study sheds light on Indonesia's interest rate fluctuations from a statistical perspective, using both basic models and machine learning. These perspectives provide insights not only into their application in empirical

studies. The results of these studies provide experience and insight into the comparison of forecasting methods, indicating that basic statistical models are also worthy of consideration for forecasting, as demonstrated in this study by the smaller MAPE values. However, this study also provides important insights, namely that statistical machine learning actually has good forecasting accuracy when used not only on big data, as theoretically, but also on data with limited timeframes. This suggests that statistical machine learning possesses significant strengths that can be further developed in future research for forecasting purposes.

Furthermore, this research is not without limitations. The study is based on values initialized by the author, and the number of possible hyperparameter combinations is substantial. However, in general, this research provides insight into the evaluation of basic statistical models and machine learning methods for Indonesia's Interest Rate fluctuations, both with and without ESG and environmental factors, and with the Indonesia interest rate alone.

The novelty of this comprehensive evaluation, which includes comparative studies, data sharing scenarios, multi-input analysis, and hyperparameter tuning, provides empirical evidence for the superior performance of the proposed method. The consistently high accuracy of the forecast MAPE values achieved by the proposed model underscores its potential as an automated tool for Indonesia's interest rate forecasting, facilitating policymakers and other researchers in maintaining national economic stability. This research also makes a significant contribution to the literature, particularly in statistics. Further studies can be developed through the integration of machine learning into economic issues to maintain national economic stability, particularly in the context of Indonesia's interest rate fluctuations. It demonstrates the importance of selecting appropriate models for specific domains.

Statistics plays a crucial role in understanding, measuring, and predicting the impacts of various human activities on the economy and the environment, and serves as a tool for identifying trends, patterns, and correlations in simple to complex data. The comparative study in this research examines the results of various statistical approaches to enrich understanding by highlighting the differences and similarities in the effects of environmental, social, and governance (ESG) factors on national economic conditions, as reflected in Indonesian interest rates. It is important to emphasize that this research does not claim that air quality (PM2.5) and rainfall indices are the only factors influencing Indonesian interest rates. We fully acknowledge that other macroeconomic variables, such as inflation, gross domestic product (GDP), and foreign exchange reserves, also play important roles in determining monetary policy and interest rate fluctuations. The primary contribution of this research is to quantify the statistical relationship between two measurable environmental factors (PM2.5 and rainfall) within a comparative forecasting framework involving basic statistical models (ARIMA, transfer function, Koyck regression) and machine learning models (LSTM and M-LSTM), not to prove full causality between environmental factors and interest rates. Nevertheless, statistical and empirical evidence obtained from the analysis indicates that environmental and ESG factors have a significant influence on Indonesian interest rates, as evidenced by the results of the transfer function model, Koyck regression, and M-LSTM. These findings allow us to measure the effectiveness of various environmental policies, such as air pollution control and water resource management.

The study of ESG and environmental effects, analyzed using various statistical approaches proposed in this research, provides new and interesting insights for discussion. The Indonesian interest rate is Bank Indonesia's benchmark interest rate, serving as the primary monetary policy instrument for controlling inflation and maintaining national economic stability [26]. The statistical forecasts conducted in this study the environmental impact on these economic issues. The results, which are significantly influenced by air quality data and the rainfall index, obtained in this study, evaluate how long-term environmental quality stability is needed to achieve Indonesia's interest rate stability. Furthermore, this relationship can be described as complex, involving a feedback mechanism where all three factors mutually influence one another. Changes in monetary policy that affect the Indonesian interest rate can influence innovation in green technology, putting pressure on natural resources, which in turn affects environmental quality and air quality. Therefore, mitigating the impact by integrating ESG environmental data into economic models is crucial.

The integration of environmental and ESG data into economic models conducted in this study, using five statistical models, demonstrates the urgency of this research. In the context of comparative studies, this provides findings from two perspectives, both in terms of scientific literature and meta-analysis, which offer a more comprehensive and measurable understanding of environmental and ESG factors that moderate the relationship between monetary policy, as represented by the Indonesian interest Rate, and its sustainability and environmental effects. The development of a forecasting model, with comparisons and a preliminary holistic study conducted in this research, enables the capture of the complexity of non-linear relationships, interactions between environmental issues, and dynamic economic feedback with the Indonesian interest Rate, utilizing fundamental statistical insights and machine learning statistics.

V. CONCLUSIONS AND SUGGESTIONS

The Indonesian interest rate forecasted through a comparative study of the two approaches showed promising results. In accordance with the research objectives, ESG factors namely, air quality (PM2.5) and rainfall index significantly affect the Indonesian interest rate value, as indicated by the analysis results of the transfer function model, Koyck regression, and M-LSTM. Generally, all proposed approach models have good forecast accuracy, as evidenced empirically by MAPE values with $MAPE < 10\%$. The MAPE values of basic statistical models (Koyck regression, transfer function, and ARIMA)

of 0.2842%, 1.0350%, and 2.1245% outperformed machine learning (LSTM and M-LSTM) by 4.660% and 7.7353%. Thus, the initial hypothesis of the study is that the development of an ESG-based Indonesia interest rate forecasting model is worthy of use as a decision-making tool, as the results of the forecast accuracy analysis have also shown significant differences through analysis of variance. The results of this study, encompassing the empirical study, can provide new insights for various parties, particularly researchers in the field of statistics who aim to understand the performance of different approaches. Furthermore, the findings of this study can serve as reference material and tools for policymakers in making decisions to maintain national economic stability, which is affected by ESG conditions and requires the development of monetary policies that internalize climate risks, particularly those related to PM2.5 and rainfall variables. This study has limitations in data coverage from January 2015 to February 2025 and location (Indonesia). Therefore, further research can be expanded to incorporate global data or explore hybrid model interpretability techniques to enhance the transparency of predictions and integrate macroeconomic and ESG variables within a single model. This would provide a clearer understanding of the interactions between environmental conditions, monetary policy, and economic stability. This study contributes methodologically by comparing forecasting approaches. It also paves the way for more holistic economic models that include environmental sustainability factors.

REFERENCES

- [1] G. Zhou, L. Liu, and S. Luo, "Sustainable development, ESG performance and company market value: Mediating effect of financial performance," *Bus. Strategy Environ.*, vol. 31, no. 7, pp. 3371–3387, 2022.
- [2] F. Wang and Z. Sun, "Does the environmental regulation intensity and ESG performance have a substitution effect on the impact of enterprise green innovation: evidence from China," *Int. J. Environ. Res. Public Health*, vol. 19, no. 14, p. 8558, 2022.
- [3] K. Abbass, M. Z. Qasim, H. Song, M. Mursheed, H. Mahmood, and I. Younis, "A review of the global climate change impacts, adaptation, and sustainable mitigation measures," *Environmental science and pollution research*, vol. 29, no. 28, pp. 42539–42559, 2022.
- [4] L. Sun, S. Fang, S. Iqbal, and A. R. Bilal, "Financial stability role on climate risks, and climate change mitigation: implications for green economic recovery," *Environmental Science and Pollution Research*, vol. 29, no. 22, pp. 33063–33074, 2022.
- [5] K. N. Khikmah, "Kajian Kinerja Multiple Long Short-Term Memory untuk Analisis Data Deret Waktu Takstasioner," Master Thesis, IPB University, Bogor, 2024.
- [6] O. Assaf, G. Di Fatta, and G. Nicosia, "Multivariate LSTM for stock market volatility prediction," in *International Conference on Machine Learning, Optimization, and Data Science*, Springer, 2021, pp. 531–544.
- [7] N.-M. Dang-Quang and M. Yoo, "An efficient multivariate autoscaling framework using bi-lstm for cloud computing," *Applied Sciences*, vol. 12, no. 7, p. 3523, 2022.
- [8] U. Turğut, D. Güler, and S. Engindeniz, "The Analysis of the Relation between Production and Price of Sunflower by Koyck Model," *Tarım Ekonomisi Dergisi*, vol. 29, no. 1, pp. 57–64, 2023.
- [9] S. Siami-Namini, N. Tavakoli, and A. S. Namin, "A comparison of ARIMA and LSTM in forecasting time series," in *2018 17th IEEE international conference on machine learning and applications (ICMLA)*, IEEE, 2018, pp. 1394–1401.
- [10] K. N. Khikmah, K. Sadik, and I. Indahwati, "TRANSFER FUNCTION AND ARIMA MODEL FOR FORECASTING BI RATE IN INDONESIA," *BAREKENG: Jurnal Ilmu Matematika dan Terapan*, vol. 17, no. 3, pp. 1359–1366, 2023.
- [11] E. Aivazidou and I. Politis, "Transfer function models for forecasting maritime passenger traffic in Greece under an economic crisis environment," *Transportation Letters*, vol. 13, no. 8, pp. 591–607, 2021.
- [12] A. Faricha *et al.*, "Comparison study of transfer function and artificial neural network for cash flow analysis at Bank Rakyat Indonesia," *International Journal of Electrical & Computer Engineering (2088-8708)*, vol. 12, no. 6, 2022.
- [13] V. K. R. Chimmula and L. Zhang, "Time series forecasting of COVID-19 transmission in Canada using LSTM networks," *Chaos Solitons Fractals*, vol. 135, p. 109864, 2020.
- [14] H. Chung and K. Shin, "Genetic algorithm-optimized long short-term memory network for stock market prediction," *Sustainability*, vol. 10, no. 10, p. 3765, 2018.
- [15] B. Karlik and A. V. Olgac, "Performance analysis of various activation functions in generalized MLP architectures of neural networks," *International Journal of Artificial Intelligence and Expert Systems*, vol. 1, no. 4, pp. 111–122, 2011.
- [16] C. Olah, "Understanding LSTM Networks," <https://colah.github.io/posts/2015-08-Understanding-LSTMs/>.
- [17] S. M. Miran, S. J. Nelson, D. Redd, and Q. Zeng-Treitler, "Using multivariate long short-term memory neural network to detect aberrant signals in health data for quality assurance," *Int. J. Med. Inform.*, vol. 147, p. 104368, 2021.
- [18] X. Xu, X. Rui, Y. Fan, T. Yu, and Y. Ju, "A multivariate long short-term memory neural network for coalbed methane production forecasting," *Symmetry (Basel)*, vol. 12, no. 12, p. 2045, 2020.
- [19] B. C. Blasco, J. J. M. Moreno, A. P. Pol, and A. S. Abad, "Using the R-MAPE index as a resistant measure of forecast accuracy," *Psicothema*, vol. 25, no. 4, pp. 500–506, 2013.
- [20] H. Alizadegan, B. Rashidi Malki, A. Radmehr, H. Karimi, and M. A. Ilani, "Comparative study of long short-term memory (LSTM), bidirectional LSTM, and traditional machine learning approaches for energy consumption prediction," *Energy Exploration & Exploitation*, vol. 43, no. 1, pp. 281–301, 2025.
- [21] X. Chen, B. Li, J. Wang, Y. Zhao, and Y. Xiong, "Integrating EMD with multivariate LSTM for time series QoS prediction," in *2020 IEEE International Conference on Web Services (ICWS)*, IEEE, 2020, pp. 58–65.
- [22] K. Smagulova and A. P. James, "A survey on LSTM memristive neural network architectures and applications," *Eur. Phys. J. Spec. Top.*, vol. 228, no. 10, pp. 2313–2324, 2019.
- [23] M. Naeem *et al.*, "Trends and future perspective challenges in big data," in *Advances in intelligent data analysis and applications: Proceeding of the sixth euro-China conference on intelligent data analysis and applications, 15–18 October 2019, Arad, Romania*, Springer, 2021, pp. 309–325.

- [24] K. N. Khikmah, K. Sadik, and K. A. Notodiputro, "Simulation and Empirical Studies of Long Short-Term Memory Performance to Deal with Limited Data," *Jurnal Online Informatika*, vol. 10, no. 1, pp. 216–226, 2025.
- [25] A. Hudaya and F. Firmansyah, "Financial stability in the Indonesian monetary policy analysis," *Cogent Economics & Finance*, vol. 11, no. 1, p. 2174637, 2023.



© 2026 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/>).