

Bias Reduced Linearization for Cluster Robust Standard Errors: An Application to Panel Data Analysis of Indonesia's Human Development Index

A. Muthiah Nur Angriany ^{1*}, A. Miftah Nabila Muthiah ²,
Andi Harismahyanti A ³, Raupong ⁴

^{1,2,4}Departement of Statistics

Faculty of Mathematics and Natural Science, Hasanuddin University, Perintis Kemerdekaan
Street Km. 10, Makassar, 90245, South Sulawesi, Indonesia

³Departement of Physics and Mathematics

Faculty of Mathematics and Natural Science, Tadulako University,
Soekarno Hatta Street km. 9, Palu City, 94148, Central Sulawesi, Indonesia

E-mail Correspondence Author: *muthiahnur@unhas.ac.id✉

Abstract

The Human Development Index (HDI) is a vital metric for evaluating the quality of life across the domains of health, education, and standard of living. Accurate estimation in this context necessitates a robust statistical approach, particularly when dealing with panel data characterized by limited clusters. This study employs the Cluster Robust Standard Error (CRSE) method, grounded in Bias-Reduced Linearization (BRL), to address bias in standard error estimations and facilitate more valid statistical inferences for modeling the HDI in Indonesia from 2018–2022. Utilizing secondary data from the Central Bureau of Statistics (BPS), the analysis compares common, fixed, and random effect models, with the final model selection informed by the Chow and Hausman tests. The findings reveal that the average years of schooling, Gross Regional Domestic Product (GRDP) at constant prices, and access to safe drinking water significantly enhance the HDI, while poverty and open unemployment rates exert a significantly negative effect. The application of BRL-adjusted cluster-robust standard errors improves the reliability of statistical inference by adjusting the variance-covariance matrix and standard errors to account for intra-cluster correlation. Importantly, the BRL adjustment does not alter the FEM coefficient estimates; rather, it provides bias-reduced standard errors and more cautious *t*- and *F*-test results. This enhances the reliability of the *t*- and *F*-test results, thereby facilitating more accurate policy-making conclusions regarding human development.

Keywords: BRL, cluster robust standard error, fixed effect, human development index, panel data.



<https://doi.org/10.30598/parameter.v5i2pp301-318>



This article is an open access article distributed under the terms and conditions of the [Creative Commons Attribution-ShareAlike 4.0 International License](#).

1. INTRODUCTION

The Human Development Index (HDI) serves as an integrated measure for assessing human well-being, encompassing three key aspects: health, education, and living standards. It provides a holistic measure of human welfare and is widely employed by researchers and policymakers to assess and compare development levels across regions and countries [1]. Understanding the socioeconomic determinants of HDI is essential for formulating policies that promote inclusive and sustainable development [2]. Regression analysis is an essential statistical method for exploring how a dependent variable relates to one or more independent variables [3]. Typically, parameter estimation in regression models is performed using the Ordinary Least Squares (OLS) method, which relies on the fulfillment of several classical assumptions including normality, homoscedasticity, independence of residuals, and absence of multicollinearity [4]. When these assumptions are violated, the model may still yield unbiased coefficients; however, the estimated standard errors can become inefficient and unreliable, leading to invalid statistical inferences [5].

Panel data regression provides an improved analytical approach for studying data observed across both individuals and periods, as it captures cross-sectional and temporal variations simultaneously [6]. The Common Effect Model (CEM), Fixed Effect Model (FEM), and Random Effect Model (REM) represent the three standard approaches in panel regression. Among these, FEM is often preferred for socioeconomic research because it controls unobserved heterogeneity by allowing individual-specific intercepts [7]. However, real-world panel data often exhibit heteroscedasticity and autocorrelation within groups or clusters, which can result in biased standard error estimates and inaccurate significance testing [8]. The Cluster Robust Standard Error (CRSE) method provides a solution by modifying the covariance matrix of regression estimators to account for intracluster correlation [9]. The CRSE offers valid inferences even when observations within the same cluster are correlated. In situations with a limited number of clusters, standard CRSE estimators may yield smaller standard errors than appropriate, creating an overly positive bias in hypothesis evaluation [10]. Bell and McCaffrey [11] introduced the Bias Reduced Linearization (BRL) estimator, which incorporates leverage adjustments derived from the projection matrix into the OLS regression. The BRL approach produces more accurate and less biased estimates of standard errors, particularly in samples with a limited number of clusters. Subsequent studies have confirmed that BRL significantly improves inference reliability in small-sample or unbalanced cluster designs [12][13]. It should be emphasized that BRL is an adjustment to the variance-covariance estimator rather than an alternative estimator for regression coefficients. Therefore, its main contribution lies in improving statistical inference through bias-reduced cluster-robust standard errors, while the underlying FEM coefficient estimates remain unchanged. The application of the CRSE with BRL adjustment is particularly relevant for analyzing the Human Development Index in Indonesia during the 2018–2022 period. Indonesia's HDI data consists of repeated observations across 34 provinces, where each province forms a distinct cluster observed over multiple years. Such clustered panel data structures are prone to intra-cluster correlations. Therefore, the CRSE-BRL approach offers an appropriate and statistically robust framework to model the relationship between HDI and its explanatory variables, including the poverty rate, Gross Regional Domestic Product (GRDP), open unemployment rate, average years of schooling, literacy rate, access to sanitation, and safe drinking water [14]. This study aims to develop an optimal panel data regression model for analysing Indonesia's HDI using cluster-robust standard errors with BRL

correction. It then evaluates the effects of key socioeconomic factors on HDI and examines how BRL-adjusted inference can improve the reliability of statistical conclusions when conventional OLS-based assumptions are not satisfied. A review of widely cited empirical studies on HDI in Indonesia indicates that most research relies on standard panel estimators and conventional cluster-robust inference, while the explicit use of BRL-adjusted cluster-robust standard errors for provincial HDI panel data is rarely discussed. Accordingly, this study's main contribution is to introduce BRL-adjusted CRSE to strengthen the inference in HDI panel modelling and to provide more credible evidence for regional development planning.

2. METHOD

2.1. Research Design

The study utilizes a quantitative approach with a causal explanatory framework to examine the impact of socioeconomic variables on HDI across Indonesia's provinces from 2018 to 2022. Panel data regression models are implemented to capture both spatial and temporal differences [3][6]. The analysis utilizes panel regression approaches comprising the Common Effect Model, Fixed Effect Model, and Random Effect Model. These models were selected sequentially using the Chow Test and Hausman Test to determine the best-fitting specification for the data [15]. All statistical analyses and model estimations were conducted in RStudio using econometric packages for panel data and robust covariance estimation, while Microsoft Excel was used for initial data auditing and management. The framework follows established econometric theory to ensure statistically valid inference that is robust to finite-sample bias.

2.2. Data Sources

The analysis was conducted using a balanced panel dataset derived from the the BPS-Statistic Indonesia website, it is available at www.bps.go.id. The observation period spans five years, from 2018 to 2022, across 34 provinces, generating a total sample size of 170 observations ($N = 34, T = 5$). The selection of variables is theoretically motivated by endogenous growth models and capability approaches found in the development economics literature [16]. Detailed operational definitions, measurement units, and relevant literature sources for each variable are summarized in **Table 1**.

Table 1. Data Sources and Variable Definitions

Variable Name	Notation	Unit	Definition and Literature Source
Human Development Index	Y	Index	Composite index of health, education, and living standards.
Poverty Rate	X_1	%	Percentage of residents living under the nationally defined poverty level.
GRDP (Constant Prices)	X_2	Trillion IDR	Gross Regional Domestic Product valuation adjusted for inflation.
Open Unemployment Rate	X_3	%	Percentage of the labor force that is unemployed.
Avg. Years of Schooling	X_4	Years	Average duration of formal education attended by the population.
Literacy Rate	X_5	%	Percentage of the population, aged 15 years or older, with literacy ability.
Proper Sanitation Access	X_6	%	Percentage of families with access to proper sanitation infrastructure.

Variable Name	Notation	Unit	Definition and Literature Source
Safe Water Access	X_7	%	Percentage of households with access to protected drinking water sources.

2.3. Analytical Procedures

The analysis follows eleven sequential steps to guarantee that the final model meets the Best Linear Unbiased Estimator (BLUE) requirements or is properly adjusted for any assumption breaches.

1. Exploratory Data Analysis (EDA)

The initial stage involved a comprehensive descriptive analysis to examine the central tendency and dispersion of data. A correlation matrix was computed to detect early signs of multicollinearity between the model's independent variables, ensuring that the predictors provided distinct information regarding the variance in HDI and allowed for an initial assessment of the relationships between variables [17].

2. Panel Data Regression Model

This study estimates the relationship between HDI and its determinants using the general panel regression framework, as shown in Equation (1) [18]:

$$Y_{it} = \beta_0 + \sum_{k=1}^K \beta_k X_{kit} + \varepsilon_{it}, i = 1, \dots, N; t = 1, \dots, T \quad (1)$$

where Y_{it} represents the HDI for province i at time t , β_0 is the intercept, β_k represents the slope coefficient for the K independent variables, and ε_{it} is the error term. Panel data regression models were classified into three types. Panel regression approaches such as Common Effect Model (CEM), Fixed Effect Model (FEM), and Random Effect Model (REM) assume differences among individuals and across time periods [19]. Each model offers a distinct approach to addressing variations among cross-sectional units and across time periods. Equations for these three models are presented as Equation (2), (3), and (4).

$$Y_{it} = \alpha + \sum_{k=1}^p \beta_k X_{kit} + \varepsilon_{it} \quad (2)$$

$$Y_{it} = \beta_{0i} + \sum_{k=1}^K \beta_k X_{kit} + \varepsilon_{it} \quad (3)$$

$$Y_{it} = \alpha + \sum_{k=1}^p \beta_k X_{kit} + \mu_i + \varepsilon_{it} \quad (4)$$

where α is a common intercept, β_{0i} is the province-specific intercept (fixed effect), μ_i is the unobserved random province effect, and the remaining terms follow the definitions in Equation (1). The CEM in Equation (2) is used as a baseline but may yield biased estimates when unobserved province characteristics correlate with the explanatory variables. The FEM in Equation (3) controls time-invariant unobserved heterogeneity by allowing β_{0i} to vary across provinces [21] [22]. The REM in Equation (4) treats the unobserved province effect as random (μ_i) and assumes it is independent of the regressors, while ε_{it} captures idiosyncratic shocks [23].

3. Model Selection Tests

The process of selecting the most suitable model specification begins with statistical tests. The Chow test is first applied to compare CEM and FEM by examining the null hypothesis that intercepts remain constant across provinces. The test statistic was calculated using Equation (5) [18], [24]:

$$F = \frac{(R_{CEM}^2 - R_{FEM}^2)/(N-1)}{\frac{R_{FEM}^2}{n-N-K}} \sim F_{(N-1, n-N-K)} \quad (5)$$

Subsequently, the Hausman test was applied to determine the relationship between the FEM and REM. This test evaluates whether the unique errors are correlated with the regressors by using the statistic in Equation (6):

$$W = (\mathbf{b} - \boldsymbol{\beta})^T [V(\mathbf{b}) - V(\boldsymbol{\beta})]^{-1} (\mathbf{b} - \boldsymbol{\beta}) \sim \chi_K^2 \quad (6)$$

where $\mathbf{b}$ and $\boldsymbol{\beta}$ are the coefficient vectors from the fixed- and random-effects models, respectively.

4. Fixed Effect Model Estimation

The Fixed Effect Model (FEM) is estimated using the within-transformation approach, which eliminates the unit-specific intercept (β_{0i}) by centering each variable around its time mean for every individual. The time-averaged values of each variable for individual i are calculated using Equation (7).

$$\hat{y}_i = \frac{1}{T} \sum_{t=1}^T y_{it}, \hat{X}_{ki} = \frac{1}{T} \sum_{t=1}^T X_{kit}, \hat{\varepsilon}_i = \frac{1}{T} \sum_{t=1}^T \varepsilon_{it} \quad (7)$$

For each individual i , the transformed variables are defined as deviations from their respective means in Equation (8):

$$y_{it}^* = y_{it} - \hat{y}_i, X_{kit}^* = X_{kit} - \hat{X}_{ki}, \varepsilon_{it}^* = \varepsilon_{it} - \hat{\varepsilon}_i \quad (8)$$

Since β_{0i} is constant over time, it is eliminated through the within-transformation. Thus, the transformed model is given by Equation (9):

$$y_{it}^* = \sum_{k=1}^K \beta_k X_{kit}^* + \varepsilon_{it}^* \quad (9)$$

In matrix notation, this can be written as Equation (10):

$$\mathbf{y}^* = \mathbf{X}^* \boldsymbol{\beta} + \boldsymbol{\varepsilon}^* \quad (10)$$

Equation (10) has the same structure as the classical linear regression model, allowing the estimation of $\boldsymbol{\beta}$ using the Ordinary Least Squares (OLS) method. The OLS estimator is formulated as Equation (11).

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}^{*'} \mathbf{X}^*)^{-1} \mathbf{X}^{*'} \mathbf{y}^* \quad (11)$$

To demonstrate unbiasedness formulated as Equation (12):

$$E(\hat{\boldsymbol{\beta}}) = \boldsymbol{\beta} \quad (12)$$

The variance-covariance matrix of $\hat{\beta}$ is formulated as Equation (13):

$$\text{Cov}(\hat{\beta}) = \sigma^2(\mathbf{X}^*'\mathbf{X}^*)^{-1} \quad (13)$$

Thus, under standard OLS assumptions, FEM provides an unbiased and efficient estimator of the slope parameter β [25].

5. Classical Assumption Diagnostics

The validity of the chosen model is verified through a series of diagnostic tests [26].

1. Normality: The distribution of the residuals was assessed using the Kolmogorov–Smirnov test to verify that they are normally distributed.
2. Homoskedasticity: The Breusch–Pagan test was applied to check for constant variance in the error terms.
3. Multicollinearity: Each predictor's Variance Inflation Factor (VIF) is computed, and values above 10 signal severe multicollinearity that requires correction.
4. Autocorrelation: The Durbin–Watson test was conducted to check for serial correlation in the residuals [27].

6. Initial Parameter Evaluation

Before applying the robust corrections, a standard regression model was evaluated. This involves calculating the determination coefficient (R^2) applied to quantify goodness of fit, performing an F-test to test the combined effect of independent variables, followed by t-tests to check the significance of each variable individually under standard conditions [28].

7. Cluster Identification

Prior to implementing robust estimation techniques, observations were systematically organized into clusters that reflected the inherent structure of the dataset. In this study, the clustering unit is defined as the province ($G = 34$), recognizing that observations within the same province over time are likely to be correlated [29].

8. Matrix Aggregation by Cluster

The regression residuals and design matrix of the independent variables were aggregated and stacked according to the identified clusters. This structural reorganization is necessary for computing the block diagonal covariance matrix required for cluster-robust estimators [7].

9. Bias Reduced Linearization (BRL) Estimation

The core methodological intervention involves calculating the bias-reduced linearization estimator, also known as CR2. This method corrects the downward bias often found in standard robust estimators, under conditions of limited cluster availability [30]. CR2 was selected because the dataset consists of 34 provincial clusters observed over five years, where observations within the same province may be correlated over time. Compared with CR0 and CR1, BRL/CR2 provides a leverage-adjusted variance-covariance matrix, making it more appropriate for finite-cluster inference in the fixed effects model. Although the wild cluster bootstrap is also a valid alternative, BRL/CR2 was chosen because this study focuses on adjusted standard errors

and coefficient-level inference. The CR2 covariance matrix is computed using [Equation \(15\)](#):

$$\hat{V}_{CR2}(\hat{\beta}) = (X'X)^{-1} \left(\sum_{g=1}^G X'_g A_g \hat{e}_g \hat{e}'_g A'_g X_g \right) (X'X)^{-1} \quad (15)$$

In this [Equation \(15\)](#), X_g represents the matrix of regressors for cluster g , and $\hat{e}_g$ denotes the vector of residuals for the cluster. The term A_g is the leverage adjustment matrix, calculated as $A_g = (I_g - H_{gg})^{-\frac{1}{2}}$, where H_{gg} is the diagonal block of the hat matrix corresponding to cluster g [\[29\]](#).

10. Robust Model Evaluation

The model is re-evaluated using the variance-covariance matrix derived from the Bias-Reduced Linearization (BRL/CR2) estimator. This adjustment produces robust standard errors that remain consistent even when model errors exhibit heteroscedasticity or intracluster correlation. The BRL correction particularly improves the inference accuracy in situations with a minimal or constrained number of clusters or when the cluster sizes are unbalanced [\[11\]](#). This re-evaluation recalculates two key statistics: the robust F statistic for joint (simultaneous) significance testing and the robust t statistic for individual (partial) significance testing. These adjusted statistics provide valid statistical inferences, even in the presence of intracluster dependence or heteroskedasticity [\[10\]](#).

a. Robust F-Test (Simultaneous Significance)

The robust F-test determines whether all explanatory variables collectively have a significant impact on the response variable when using cluster-robust variance estimation. The hypotheses are:

$H_0: \beta_1 = \beta_2 = \dots = \beta_K = 0$ (no joint significance)

$H_1: \text{At least one } \beta_k \neq 0$ (joint significance exists)

The robust F-statistic is computed as [Equation \(16\)](#):

$$F_{robust} = \frac{\left(\frac{R^2}{K} \right)}{\left(\frac{1 - R^2}{n - K - 1} \right)} \sim F_{K, n-K-1} \quad (16)$$

where R^2 is the coefficient of determination; K is the number of explanatory variables; and n is the total number of observations. The error variance $\hat{\sigma}_{BRL}^2$ is derived from the BRL covariance matrix $\hat{V}_{CR2}$. The null hypothesis H_0 is rejected if $F_{robust} > F_{critical}$ or if the p-value $< \alpha$ (typically 0.05). This indicates that the independent variables collectively exert a statistically significant influence on the dependent variable.

b. Robust t-Test (Partial Significance)

The robust t-test examines the impact of each individual predictor while incorporating the BRL-adjusted standard errors into the analysis. For the k -th coefficient, the hypotheses are as follows:

$H_0: \beta_k = 0$ vs. $H_1: \beta_k \neq 0$

The test statistic is formulated as [Equation \(17\)](#):

$$t_{k,robust} = \frac{\hat{\beta}_k}{\sqrt{\widehat{SE}_{cr_2}(\hat{\beta}_k)}} \sim t_{n-K-1} \quad (17)$$

where $\widehat{Var}(\hat{\beta}_k)$ is the k -th diagonal element of the BRL-adjusted covariance matrix $\hat{V}_{CR_2}$. The null hypothesis is rejected if $|t_{k,robust}| > t_{critical}$ or if the p-value $< \alpha$, confirming that the k -th variable significantly affects the dependent variable.

11. Comparative Analysis

Finally, a comparative analysis was conducted between the initial model estimates (Step 5) and CRSE-corrected estimates (Step 9). This comparison focuses on the differences in standard error magnitudes, the shifts in the statistical significance of coefficients (p-values), and the reliability of the simultaneous and partial test results, thereby highlighting the importance of robust estimation in regional panel data analysis [31].

3. RESULTS AND DISCUSSION

3.1. Data Exploration

Descriptive statistical analysis of the Human Development Index (HDI) in Indonesia for the period 2018–2022 shows that, in general, most provinces fall into medium to high human development categories according to UNDP standards. Based on the data in Table 2, the HDI variable, as the dependent variable, has a distribution that tends to be homogeneous, as reflected by the relatively small standard deviation. In contrast, significant fluctuations were found in the independent variables, particularly in variable X_6 , which exhibited a very wide value range and a high standard deviation. This phenomenon indicates the presence of sharp socioeconomic heterogeneity among provinces in Indonesia, where the factors supporting human development are not yet evenly distributed across all regions.

Table 2. Descriptive Statistics of Each Variable

Variable	Min	Max	Mean	Standard Deviation
Y	60.060	81.650	71.170	3.917
X_1	3.420	27.430	10.475	5.408
X_2	25.030	1953.490	326.580	462.131
X_3	1.400	10.950	5.201	1.797
X_4	6.520	11.310	8.624	0.935
X_5	76.790	99.870	96.350	4.177
X_6	33.750	97.120	77.490	11.596
X_7	49.370	99.860	83.380	10.458

Further data exploration through the correlation plot in Figure 1 provides an overview of the direction and strength of the relationships between the variables. Based on Figure 1, variables X_4 (0.79) and X_6 (0.72) have the strongest positive correlation with HDI, while X_1 shows a negative correlation of -0.66 . The other independent variables (X_2, X_3, X_5, X_7) are positively correlated with a moderate strength. The presence of correlations among independent variables, such as X_6 with X_7 (0.67) and X_4 with X_5 (0.63), indicates potential multicollinearity that needs to be mitigated in developing the panel regression model to maintain the reliability of the result interpretation.



Figure 1. Correlation Plot between Variables

3.2. Selection of Panel Data Regression Model

An initial analysis using panel data modeling was conducted to determine the most accurate specification for describing the characteristics of the open unemployment rate in Indonesia. Three estimation approaches were used: the Common Effect Model (CEM), Fixed Effect Model (FEM), and Random Effect Model (REM). A summary of the parameter estimation results for the three models based on [Equations \(2\), \(3\), and \(4\)](#) is presented in [Table 3](#).

Table 3. Summary of CEM, FEM, and REM Estimation Results

Variable	CEM Coefficient	REM Coefficient	FEM Coefficient
Intercept	52.570	50.869	
X_{1it}	-0.163	-0.167	-0.138
X_{2it}	0.002	0.002	0.002
X_{3it}	-0.201	-0.064	-0.070
X_{4it}	2.372	2.749	2.778
X_{5it}	-0.073	-0.026	-0.024
X_{6it}	0.127	0.001	-0.000
X_{7it}	-0.031	0.006	0.007

The optimal model was selected using a two-step evaluation process. Initially, the Chow Test was employed to assess the CEM and FEM models. The results derived from [Equation \(5\)](#), yielded an F_{value} of 458.450 and a p – value of 0.000. Since the p – value was less than the significance level $\alpha = 0.05$, the null hypothesis (H_0) was rejected. This suggests that the FEM model is more suitable than the CEM model for this study. Subsequently, a Hausman Test was performed to compare the FEM and REM models. According to [Equation \(6\)](#) the test produced a chi-square (χ^2) statistic of 18.029, with a p – value of 0.012. As the p – value was below $\alpha = 0.05$, H_0 was rejected. From this selection process, it was concluded that the fixed effect model (FEM) is the most appropriate for analyzing the Human Development Index in Indonesia. This conclusion implies that the regional differences in Indonesia are fixed and correlated with the model's explanatory variables. Consequently, the FEM model was chosen as an efficient and consistent estimator for further analysis.

3.3. Fixed Effect Model Panel Data Regression Estimation

The fixed effects model was employed to estimate panel data regression for factors affecting the Human Development Index (HDI) in Indonesia from 2018 to 2022, utilizing the within-group method as outlined in Equation (11). Based on an analysis of 170 observations (34 provinces over five periods), the estimated model is shown in Equation (18).

$$\hat{y}_{it} = \beta_{0i} - 0.138X_{1,it} + 0.002X_{2,it} - 0.070X_{3,it} + 2.778X_{4,it} - 0.024X_{5,it} - 0.000X_{6,it} + 0.007X_{7,it} \quad (18)$$

The main characteristic of the fixed effects model is the presence of a different intercept β_0 for each cross-sectional unit, which in this study represents each province in Indonesia. Table 4 presents the estimated values of province-specific intercepts or constants.

Table 4. Province intercept in Indonesia

Province	β_{0i}	Province	β_{0i}
Aceh	50.348	Kepulauan Riau	50.498
Bali	52.834	Lampung	50.904
Banten	50.387	Maluku	46.619
Bengkulu	51.145	North Maluku	46.581
DI Yogyakarta	57.004	West Nusa Tenggara	51.549
DKI Jakarta	49.836	East Nusa Tenggara	48.834
Gorontalo	51.219	Papua	47.153
Jambi	50.750	West Papua	49.057
West Java	49.596	Riau	49.949
Central Java	52.609	West Sulawesi	47.773
East Java	50.998	South Sulawesi	51.357
West Kalimantan	50.116	Central Sulawesi	48.799
South Kalimantan	50.662	Southeast Sulawesi	49.768
Central Kalimantan	50.011	North Sulawesi	49.832
East Kalimantan	51.852	West Sumatra	50.369
North Kalimantan	48.859	South Sumatra	50.663
Kepulauan Bangka Belitung	51.935	North Sumatra	47.945

Table 4 shows that each province has a varying constant value, indicating the presence of region-specific effects on HDI achievement that are not captured by the independent variables in the model. Overall, this model exhibited a strong performance, with a coefficient of determination (R^2) of 0.926. This indicates that 92.60% of the variation in the dependent variable (HDI) can be explained by the independent variables in the model, while the remaining 7.40% is explained by factors outside the scope of this study.

3.4. Classical Assumption Testing

The panel data regression model's validity is assessed through a series of classical assumption tests to confirm that the estimated parameters are the Best Linear Unbiased Estimators (BLUE). The initial phase of testing involves a normality test using the Kolmogorov-Smirnov method, which aims to verify if the model's residuals follow a normal distribution. The null hypothesis (H_0) posits that the errors are normally distributed. From the best model's results, a test statistic value of $D = 0.0903$ was obtained, with a p – value of 0.125. Since this significance value exceeds $\alpha = 0.05$, the null hypothesis H_0 is accepted. This indicates that the errors in this regression model meet the normal distribution assumption. The subsequent step involves checking for

multicollinearity among the independent variables by examining the Variance Inflation Factor (VIF) values. The results, summarized in [Table 5](#), reveal that all independent variables have VIF values below 10. This consistency suggests that there is no strong linear correlation among the explanatory variables, ensuring that the relationships between variables remain within acceptable limits and do not compromise the stability of the regression parameter estimates.

Table 5. VIF Values of the Model

Variable	VIF Value
X ₁	1.859
X ₂	1.426
X ₃	1.712
X ₄	2.268
X ₅	2.068
X ₆	2.710
X ₇	2.341

Next, a Breusch–Pagan test was conducted to detect heteroscedasticity. The structure of the error variance was tested to determine whether it is homoscedastic or has a non-constant variance. Based on the test results, the obtained Breusch-Pagan statistic was 27.116 with a p – value of 0.000. Because the p – value is smaller than the significance level of 0.05, the null hypothesis is rejected, indicating the presence of heteroscedasticity in this panel data model. The final stage of classical assumption testing is the autocorrelation test using the Durbin (DW) method. This test is intended to detect correlations between errors during consecutive observation periods. The test results showed a DW value of 1.729 with a p-value of 0.033. The significance value, which is less than 0.05, leads to rejection of the null hypothesis, indicating the presence of positive autocorrelation in the model’s residuals. Although indications of heteroscedasticity and autocorrelation were found, this is common in panel data with both spatial and temporal dimensions; thus, handling through standard error adjustment becomes essential to maintain the consistency of the analysis results in the subsequent stages. Based on these diagnostic results, the use of BRL/CR2 is justified by the presence of heteroscedasticity, positive autocorrelation, and the clustered panel structure of the data. The dataset contains 34 provincial clusters observed repeatedly over five years, so observations within the same province are likely to be dependent over time. Conventional CR0 and CR1 standard errors may be less reliable in this finite-cluster setting because they do not fully account for cluster leverage. Therefore, BRL/CR2 was applied to adjust the variance-covariance matrix and obtain bias-reduced cluster-robust standard errors. This adjustment improves the reliability of the t- and F-test results without changing the FEM coefficient estimates.

3.5. Cluster Robust Standard Error with Bias-Reduced Linearization

Parameter analysis of the regression model in the initial stage was conducted using Ordinary Least Squares (OLS) estimation to identify the significance of the parameters both simultaneously and partially. A simultaneous test (F-test) was performed to determine the joint effect of all independent variables on the dependent variable. The results show that $F_{\text{value}} = 237.569 > F_{0.05,7,129} = 2.081$ and $p\text{ – value} = 0.000 < \alpha = 0.05$, therefore H_0 is rejected. It can be concluded that simultaneously, all variables

used in the study affect the Human Development Index (Y). To determine the individual contributions of each variable, a partial t-test was conducted, as presented in [Table 6](#).

Table 6. Partial Test Result

Variable	Standard Error	t-value	p – value	Decision
X_1	0.038	-3.637	0.000*	Reject H_0
X_2	0.001	2.448	0.016*	Reject H_0
X_3	0.121	-3.266	0.001*	Reject H_0
X_4	0.151	18.367	0.000*	Reject H_0
X_5	0.037	-0.652	0.516	Accept H_0
X_6	0.004	-0.058	0.954	Accept H_0
X_7	0.004	1.623	0.107	Accept H_0

Based on [Table 6](#), variables X_1, X_2, X_3 , and X_4 show a significant effect on the Human Development Index (HDI). However, considering the presence of bias in the standard error due to heteroskedasticity and autocorrelation, a variance-covariance adjustment was made using the Cluster Robust Standard Error (CRSE) method of the Bias-Reduced Linearization (BRL) or CR2 type. This adjustment was made because the standard Ordinary Least Squares (OLS) estimation tends to yield less accurate results for panel data with a limited number of clusters.

3.5.1. Implementation of the Bias-Reduced Linearization Method

Bias-Reduced Linearization (BRL) or CR2 adjustment is applied to obtain bias-reduced cluster-robust variance estimates and to reduce finite-sample distortion caused by the limited number of provincial clusters. The BRL estimate was calculated using [Equation \(15\)](#).

$$\hat{V}_{CR2}(\hat{\beta}) = \begin{bmatrix} 0,0039 & 0,0000 & -0,0007 & -0,0091 & 0,0014 & 0,0001 & 0,0000 \\ 0,0000 & 0,0000 & 0,0000 & -0,0001 & 0,0000 & 0,0000 & 0,0000 \\ -0,0007 & 0,0000 & 0,0007 & 0,0002 & -0,0002 & -0,0000 & 0,0000 \\ -0,0091 & -0,0001 & 0,0002 & 0,0430 & -0,0061 & -0,0004 & -0,0003 \\ 0,0014 & 0,0000 & -0,0002 & -0,0061 & 0,0020 & 0,0000 & 0,0001 \\ 0,0001 & 0,0000 & -0,0000 & -0,0004 & 0,0000 & 0,0000 & -0,0000 \\ 0,0000 & 0,0000 & 0,0000 & -0,0003 & 0,0001 & -0,0000 & 0,0000 \end{bmatrix}$$

Therefore, the Standard Error value shows in [Equation \(19\)](#).

$$SE_{CR2}(\hat{\beta}_j) = \sqrt{Var(\hat{\beta}_j)} = \begin{bmatrix} 0,0623 \\ 0,0007 \\ 0,0258 \\ 0,2075 \\ 0,0449 \\ 0,0035 \\ 0,0031 \end{bmatrix} \quad (19)$$

3.5.2. Parameter Estimation Results with the Bias-Reduced Linearization Method

The parameter estimation results of the fixed effects model (FEM), following the adjustment of the variance-covariance matrix using the Cluster Robust Standard Error (CRSE) method type CR2, also known as Bias-Reduced Linearization (BRL), are presented in [Table 7](#).

Table 7. FEM Estimation Result with CR2

Variable	Coefficient	Standard Error	t-value	p – value
X_1	-0.1383	0.0623	-2.2198	0.0282*
X_2	0.0016	0.0007	2.3378	0.0209*
X_3	-0.0696	0.0258	-2.6919	0.0080*
X_4	2.7785	0.2075	13.3910	0.0000*
X_5	-0.0239	0.0449	-0.5321	0.5956
X_6	-0.0002	0.0035	-0.0664	0.9472
X_7	0.0068	0.0031	2.2019	0.0294*

Table 7 shows that the estimated regression coefficients remain the same after adjustment using the CRSE method with BRL. This confirms that BRL does not re-estimate or modify the FEM coefficients. Instead, it adjusts the estimated variance-covariance matrix, which subsequently affects the standard errors, t-values, and p-values used for statistical inference. The t-values in **Table 7**, obtained based on **Equation (17)**, indicate that variables X_1, X_2, X_3, X_4 , and X_7 have a significant effect on variable Y at a significance level of $\alpha = 5\%$. The adjustment with BRL aims to address issues of heteroscedasticity and autocorrelation within clusters, which can result in biased standard error values. With this correction, the t-test results become more reliable as they take into account dependencies within clusters. In addition, the simultaneous significance test based on **Equation (16)** yields an $F_{\text{value}} = 153$, showing that the model simultaneously has a significant effect on the HDI. The differences between the General OLS and BRL results are presented in **Table 8**.

Table 8. Comparison of Standard Error Value

Parameter	General OLS			Bias-Reduced Linearization		
	Standar Error	t-value	p-value	Standar Error	t-value	p-value
β_1	0.0380	-3.6369	0.0004*	0.0623	-2.2198	0.0282*
β_2	0.0007	2.4479	0.0157*	0.0007	2.3378	0.0209*
β_3	0.1212	-3.2664	0.0013*	0.0258	-2.6919	0.0080*
β_4	0.1513	18.3669	0.0000*	0.2075	13.391	0.0000*
β_5	0.0367	-0.6515	0.5159	0.0449	-0.5321	0.5956
β_6	0.0040	-0.0584	0.9535	0.0035	-0.0664	0.9472
β_7	0.0042	1.6226	0.1071	0.0031	2.2019	0.0294*

*) Significant at $\alpha=5\%$

Table 8 presents a comparison between the conventional OLS-based standard errors and the standard errors obtained after applying the Bias-Reduced Linearization (BRL/CR2) adjustment. The comparison shows that the BRL adjustment affects the standard errors, t-values, and p-values, while the fixed effect model coefficient estimates remain unchanged. This occurs because BRL does not re-estimate the regression coefficients; rather, it modifies the estimated variance-covariance matrix to account for heteroscedasticity and intra-cluster correlation. After applying the BRL adjustment, several standard errors changed, which subsequently affected the statistical significance of some variables. For example, the safe water access variable, which was not statistically significant under the conventional OLS-based inference, became significant after the BRL adjustment. Therefore, the BRL results should be interpreted as providing more reliable cluster-robust inference for panel data with intra-cluster dependence, not as producing coefficient estimates that are superior to conventional OLS or FEM estimates.

3.6. Best Model Parameter Estimation

The final model was obtained using a backward elimination procedure guided by BRL-adjusted robust p-values and theoretical considerations. In the full FEM specification with BRL-adjusted cluster-robust standard errors, Literacy Rate (X_5) and Access to Proper Sanitation (X_6) had BRL-adjusted p-values above the 5% significance level. Therefore, these variables were considered for removal to obtain a more parsimonious model. Their exclusion was also evaluated in relation to the theoretical framework of human development, so the final specification was not determined solely by statistical insignificance. The final model consistently uses the fixed effects model (FEM) with BRL-adjusted cluster-robust standard errors. The parameter estimation results are presented in [Table 9](#).

Table 9. Parameter Estimation Results of the Best Model

Parameter	coefficient	Cluster Robust Standard Error	t-value	P-Value
β_1	-0.1358	0.0561	-2.4215	0.0168*
β_2	0.0017	0.0007	2.4386	0.0168*
β_3	-0.0706	0.0246	-2.8637	0.0049*
β_4	2.7284	0.1415	19.2763	0.0000*
β_7	0.0072	0.0027	2.7048	0.0077*

*) Significant at $\alpha = 5\%$

The application of the fixed effects model results in varying intercept values between regions (β_0) which represent region-specific effects or unobserved variables that are permanent in each province. The values of these constants are listed in [Table 10](#).

Table 10. Provincial Constant Values of the Best Model

Province	β_{0i}	Province	β_{0i}
Aceh	48.3830	Kepulauan Riau	48.5733
Bali	50.9556	Lampung	48.9067
Banten	48.4245	Maluku	44.6548
Bengkulu	49.1716	North Maluku	44.6110
DI Yogyakarta	55.1248	West Nusa Tenggara	49.7312
DKI Jakarta	47.9041	East Nusa Tenggara	46.8866
Gorontalo	49.1638	Papua	45.5142
Jambi	48.7658	West Papua	47.0070
West Java	47.5693	Riau	47.9585
Central Java	50.6502	West Sulawesi	45.8655
East Java	49.0594	South Sulawesi	49.4861
West Kalimantan	48.1934	Central Sulawesi	46.8141
South Kalimantan	48.6700	Southeast Sulawesi	47.8696
Central Kalimantan	48.0157	North Sulawesi	47.8560
East Kalimantan	49.9016	West Sumatra	48.3898
North Kalimantan	46.9400	South Sumatra	48.6345
Kepulauan Bangka	49.9365	North Sumatra	45.9732
Belitung			

Based on [Table 9](#) and [Table 10](#), the FEM model equation for Indonesia ($i = 28$) is formulated as shown in [Equation \(20\)](#).

$$\hat{y}_{it} = \beta_{0i} - 0.1358X_{1,it} + 0.0017X_{2,it} - 0.0706X_{3,it} + 2.7284X_{4,it} + 0.0072X_{7,it} \quad (20)$$

The coefficient interpretation shows that average years of schooling has the strongest positive association with the Human Development Index (HDI) in Indonesia. A one-year higher average years of schooling is associated with a 2.7284-point higher HDI, holding other variables constant. GRDP at constant prices is also positively associated with HDI, with a coefficient of 0.0017. In addition, access to safe drinking water is positively associated with HDI, with a coefficient of 0.0072. These results indicate that provinces with higher educational attainment, stronger regional economic performance, and better access to safe drinking water tend to have higher HDI values. The findings of this study are consistent with the broader literature on human development. The United Nations Development Programme (UNDP, 2022) identifies education and living standards as fundamental components of the Human Development Index (HDI) [1], which supports the positive association between average years of schooling and HDI found in this study. Pham (2021) provides panel evidence showing that educational advancement is strongly correlated with higher HDI in ASEAN countries [16]. Similarly, Rao (2021), using cross-country panel data, indicates that the association between income growth and human development may be limited when it is not accompanied by effective social investment and poverty reduction strategies [26]. Conversely, poverty rate and open unemployment rate are negatively associated with HDI. A one-percentage-point higher poverty rate is associated with a 0.1358-point lower HDI, while a one-percentage-point higher open unemployment rate is associated with a 0.0706-point lower HDI, holding other variables constant. These findings suggest that provinces with higher poverty and unemployment levels tend to have lower HDI values. Therefore, the results should be interpreted as statistical associations rather than direct causal effects.

4. CONCLUSION

Based on the results and discussion, the best panel data regression model for the Human Development Index (HDI) in Indonesia during the 2018–2022 period is the fixed effects model (FEM) with BRL-adjusted cluster-robust standard errors. The estimated FEM coefficients indicate that one-unit increases in $X_{1,it}$ and $X_{3,it}$ are associated with decreases in HDI of 0.1358 and 0.0706 units, whereas one-unit increases in $X_{2,it}$, $X_{4,it}$, and $X_{7,it}$ are associated with increases in HDI of 0.0017, 2.7284, and 0.0072 units, with province-specific intercepts β_{0i} capturing unobserved heterogeneity across the 34 provinces. The Bias-Reduced Linearization (BRL/CR2) adjustment is applied to the variance-covariance matrix in order to obtain bias-reduced cluster-robust standard errors and more reliable statistical inference in the presence of intra-cluster correlation. It should be emphasized that BRL does not change the FEM coefficient estimates; rather, it affects the standard errors, t-values, p-values, and the corresponding inference. After applying the BRL-adjusted cluster-robust standard errors, poverty rate, open unemployment rate, GRDP at constant prices, average years of schooling, and access to safe drinking water remain statistically significant, while literacy rate and access to proper sanitation are excluded from the final model because they are not statistically significant. Substantively, the findings suggest that policies aimed at improving HDI in Indonesia should prioritize reducing poverty and unemployment, improving educational attainment, expanding access to safe drinking water, and strengthening regional economic performance. These policy directions are supported by the final FEM results after correcting the standard errors for intra-cluster dependence at the provincial level.

Funding Information

This study did not receive any specific grants from funding agencies in the public, commercial, or not-for-profit sectors. The authors declare no conflicts of interest.

Conflict of Interest Statement

The authors declare that they have no conflict of interest.

Informed Consent

Not applicable. This study used secondary data from publicly available sources and did not involve individual human participants, patients, or identifiable personal data.

Ethical Approval

Not applicable. This research was based solely on secondary data obtained from publicly available sources and did not involve experiments on humans or animals; therefore, formal ethical approval was not required.

REFERENCES

- [1] U. N. D. Programme, *Human Development Report 2022: Uncertain Times, Unsettled Lives – Shaping Our Future in a Transforming World*. New York, NY, USA: UNDP, 2022. [Online]. Available: <https://hdr.undp.org>
- [2] B. C. Sarma, "Spatial analysis of human development and economic growth in Asia," *Asian Economic Journal*, vol. 33, no. 2, pp. 141–159, 2020.
- [3] D. N. Gujarati and D. C. Porter, *Basic Econometrics*, 5th ed. New York, NY, USA: McGraw-Hill, 2009.
- [4] J. M. Wooldridge, *Introductory Econometrics: A Modern Approach*, 7th ed. Boston, MA, USA: Cengage Learning, 2020.
- [5] H. L. White, "A heteroskedasticity-consistent covariance matrix estimator and a direct test for heteroskedasticity," *Econometrica*, vol. 48, no. 4, pp. 817–838, 1980, doi: 10.2307/1912934.
- [6] B. H. Baltagi, *Econometric Analysis of Panel Data*, 6th ed. Chichester, U.K.: John Wiley & Sons, 2021.
- [7] A. C. Cameron and D. L. Miller, "A Practitioner's Guide to Cluster-Robust Inference," *Journal of Human Resources*, vol. 50, no. 2, pp. 317–372, 2015, [Online]. Available: https://cameron.econ.ucdavis.edu/research/Cameron_Miller_JHR_2015_February.pdf
- [8] J. E. Pustejovsky and E. Tipton, "Small sample methods for cluster-robust variance estimation and hypothesis testing in fixed effect models," *Journal of Business and Economic Statistics*, vol. 40, no. 4, pp. 1431–1446, 2022, [Online]. Available: <https://jepusto.com/files/Pustejovsky-Tipton-201601.pdf>
- [9] W. K. Newey and K. D. West, "A simple, positive semi-definite, heteroskedasticity and autocorrelation consistent covariance matrix," *Econometrica*, vol. 55, no. 3, pp. 703–708, 1987, doi: 10.2307/1913610.
- [10] C. Young and M. Shah, "Cluster-robust standard errors for few clusters," *Econom J*, vol. 27, no. 1, pp. 23–41, 2024, doi: 10.1093/ectj/utad005.
- [11] R. M. Bell and D. F. McCaffrey, "Bias reduction in standard errors for linear regression with multi-stage samples," *Surv Methodol*, vol. 28, no. 2, pp. 169–181, 2002, [Online]. Available: <https://www150.statcan.gc.ca/n1/pub/12-001-x/2002002/article/9058-eng.pdf>
- [12] G. W. Imbens and M. Kolesar, "Robust Standard Errors in Small Samples: Some Practical Advice," 2016. [Online]. Available: https://oar.princeton.edu/bitstream/88435/pr1xv3r/1/REST_a_00552.pdf
- [13] J. G. MacKinnon, M. Ø. Nielsen, and M. D. Webb, "Cluster-robust inference: A guide to empirical practice," 2022. [Online]. Available: <https://arxiv.org/abs/2205.03285>
- [14] P. Kennedy, *A Guide to Econometrics*, 7th ed. Malden, MA, USA: Blackwell Publishing, 2008.

- [15] A. H. A, A. Najiha, and A. I. Yunita, "Comparison of FEM-LSDV Panel Regression with Classical Panel Regression Models in Analyzing Economic Growth in Indonesia," *Journal of Applied Informatics and Computing (JAIC)*, vol. 9, no. 5, pp. 2450–2460, 2025, [Online]. Available: <https://jurnal.polibatam.ac.id/index.php/JAIC/article/view/10318/3082>
- [16] T. T. Pham, "Human development, inequality, and growth: Evidence from ASEAN," *Int J Soc Econ*, vol. 48, no. 3, pp. 432–449, 2021.
- [17] L. D. Cavallo and M. Rossi, "How does human capital affect growth? Evidence from Latin America," *Econ Model*, vol. 79, pp. 101–111, 2019.
- [18] C. Hsiao, *Analysis of Panel Data*, 3rd ed. Cambridge, UK: Cambridge University Press, 2014.
- [19] A. H. A. and N. Nur'eni, "Analyzing the impact of inflation, exports and unemployment on economic growth in indonesia: A fixed effects least squares dummy variable panel regression approach," *Priviet Social Sciences Journal*, vol. 5, no. 8, pp. 338–348, Aug. 2025, doi: 10.55942/PSSJ.V5I8.541.
- [20] N. Indrasari and P. Yulianto, "Determinants of poverty in Papua Province: A panel data regression approach (2021–2023)," *Jurnal Statistika dan Aplikasinya*, vol. 9, no. 1, 2025, doi: 10.21009/jsa.09112.
- [21] R. Aprilianti, M. Sari, and D. Putri, "Fixed effect model analysis in panel data regression," *Jurnal Ilmiah Statistika dan Aplikasinya*, vol. 8, no. 2, pp. 115–122, 2022.
- [22] I. P. Hutagalung and O. Darnius, "Analisis Regresi Data Panel Dengan Pendekatan Common Effect Model (CEM), Fixed Effect Model (FEM) dan Random Effect Model (REM)," vol. 5, no. No. 2, pp. 217–226, 2022.
- [23] N. Prawoto and A. Basuki, "Factors affecting poverty in Indonesia: A panel data approach," *Quality - Access to Success*, vol. 23, no. 186, 2022, doi: 10.47750/qas/23.186.20.
- [24] A. M. H. Supriyanto, "Regional disparity in Indonesia's human development index: A panel data approach," *Procedia Economics and Finance*, vol. 23, pp. 1190–1195, 2015.
- [25] A. Firdaus and Y. Dawood, "Determinants of poverty in Indonesia: An empirical evidence using panel data regression," *International Journal of Global Operations Research*, vol. 2, no. 4, 2021, doi: 10.47194/ijgor.v2i4.90.
- [26] R. B. Rao, "Economic growth and human development nexus: Panel evidence," *World Dev*, vol. 139, p. 105293, 2021.
- [27] J. H. Stock and M. W. Watson, *Introduction to Econometrics*, 4th ed. Harlow, UK: Pearson, 2020.
- [28] F. R. Lestari, "Determinants of Human Development Index in Indonesia," *Journal of Economic Development Policy*, vol. 7, no. 2, pp. 121–135, 2021.
- [29] J. G. MacKinnon and M. D. Webb, "The wild bootstrap for few (treated) clusters," *Econom J*, vol. 21, no. 2, pp. 114–135, 2018.
- [30] M. D. Webb, "Reworking wild bootstrap-based inference for clustered errors," *Econom J*, vol. 17, no. 1, pp. 16–48, 2014.
- [31] M. E. Santos and S. Alkire, "Training material on the Multidimensional Poverty Index," 2015.

