

LEMBAR PENGESAHAN TULISAN ILMIAH

Pada Hari, Tanggal : Senin, 15 juni 2026

Telah Disahkan Tulisan Ilmiah

JUDUL: *Smoothed Empirical Likelihood Approach to Quantile Differences in Paired Data: A Simulation and Application Study on Indonesian Economic Data.*

DISUSUN OLEH :

Eka Putri Nur Utami, Yuriska Destania, Kusman Sadik , Anang Kurnia

Disahkan Oleh:
Dekan Sekolah Sains Data, Matematika, dan Informatika

Prof. Dr. Ir. Agus Buono, M.Si., M.Kom
NIP 196607021993021001

Smoothed Empirical Likelihood Approach to Quantile Differences in Paired Data: A Simulation and Application Study on Indonesian Economic Data

Eka Putri Nur Utami¹, Yuriska Destania^{1,2}, Kusman Sadik¹, Anang Kurnia¹

¹School of Data Science, Mathematics and Informatics, IPB University, Bogor, Indonesia

²Mathematics Education Study Program, Muhammadiyah University of Bengkulu, Indonesia

ABSTRACT

This study aims to conduct inference on quantile differences in paired data using the smoothed empirical likelihood (SEL) approach. Conventional mean-based approaches are limited in their ability to handle outliers and are less effective in comprehensively representing changes in distribution. The SEL method addresses the issue of the discontinuity of indicator functions in standard empirical likelihood by integrating kernel-based smoothing techniques, thereby enhancing computational stability and coverage accuracy. Through simulation studies involving various distribution scenarios (normal, exponential, and lognormal), sample sizes, and correlation levels, the performance of SEL is evaluated based on coverage probability and average interval length, compared to bootstrap, t-test, and Wilcoxon test methods. The simulation results demonstrate that SEL consistently yields stable coverage probabilities close to the nominal target and achieves high estimation efficiency with narrow confidence intervals, particularly for skewed data distributions. In an empirical application using Indonesia's Gross Regional Domestic Product (PDRB) data for 2016 and 2025, the SEL method proves to be more robust in representing regional economic dynamics than the mean-based method, which tends to be upwardly biased due to the influence of extreme values in metropolitan provinces. Furthermore, the confidence intervals generated by SEL are almost twice as precise as those from the t-test, making it a more accurate analytical tool for development planning aimed at regional equity.

Keywords:

Paired data, Quantile differences, Kernel refinement

1. INTRODUCTION

Macroeconomic indicators such as Gross Regional Domestic Product (PDRB), income inequality indices, and regional productivity are frequently characterized by non-normal traits, including extreme right skewness, heavy tails, and highly localized outliers [1], [2]. In many empirical distributions within the field of economics, these features undermine the validity of classical Gaussian assumptions [1], [3]. Conventional statistical methods that rely on the mean, such as the standard *t-test*, often provide biased representations of typical welfare because they are disproportionately sensitive to extreme values. Consequently, mean based metrics may lead to misleading conclusions regarding structural growth shifts and regional progress.

To address these distributional complexities, quantile-based inference has emerged as a robust and vital alternative, as it allows for a comprehensive characterization of the conditional distribution [2], [4]. By capturing the behavior of tails and central tendencies separately, quantile regressions originally introduced by Koenker and Bassett provide a more nuanced understanding of economic relationships [4]–[6]. This approach is particularly valuable in program evaluation and poverty analysis, where the impact of a variable may vary significantly for regions at different levels of the economic hierarchy [7], [8].

The Empirical Likelihood (EL) method, pioneered by Owen, represents a significant advancement in nonparametric statistics by combining the flexibility of distribution-free approaches with the efficiency of likelihood based frameworks [9], [10]. EL is highly advantageous because it allows for the construction of range-preserving confidence regions that are determined automatically by the data's configuration [11]. Unlike traditional methods, it avoids the need for explicit

studentization or the estimation of complex variance components, making it robust for skewed data [11]–[14].

However, applying standard EL to quantiles faces methodological challenges due to the non differentiable indicator function in the estimating equations, which often leads to computational instabilities and poor coverage accuracy in finite samples [14], [15]. To resolve these limitations, the Smoothed Empirical Likelihood (SEL) approach replaces the indicator function with a smooth kernel-based function [9], [11]. This smoothing makes the constraint function differentiable, which not only improves numerical optimization but also ensures the test statistic is Bartlett correctable, thereby enhancing accuracy in small samples.

In regional economic studies and longitudinal surveys, data frequently arise in paired formats, such as GRDP figures recorded for the same administrative units across two different years [1], [10]–[21]. Assuming independence in these settings is a fundamental violation of core statistical assumptions, which can result in invalid inferences and severe under coverage of confidence intervals. Effective inference for paired samples requires accounting for the inherent temporal or spatial dependencies between the observation units [1].

Theoretical expansions of SEL for paired quantile differences allow researchers to model bivariate joint distributions while effectively profiling out marginal nuisance parameters [22]. This specialized framework is essential for isolating shifts at specific distribution point for instance, analyzing the lower quartile to evaluate progress in low-income provinces versus using the median to measure robust central growth trends. Furthermore, SEL for paired data provides a statistically accurate alternative to traditional approximations when dealing with heavy tailed macroeconomic distributions.

Recent methodological innovations have introduced even more refined tools for distribution-free inference, such as the Smoothed Jackknife Empirical Likelihood (JEL), which avoids estimating link variables to increase computational efficiency [23], [24]. Additionally, advanced two-stage EL procedures have been developed to handle unknown intercepts in predictive models without sacrificing sample size [2]. New developments in Self Normalized Quantile Empirical Saddlepoint Approximation (SNQESA) also provide analytical solutions that achieve second-order coverage accuracy without the need for bandwidth selection [18], [25].

In the specific context of the Indonesian economy, significant regional disparities persist, with economic output and development heavily concentrated on Java Island compared to Sumatra or Kalimantan [3], [5], [8]. High economic growth in metropolitan centers often masks the stagnation in other regions, creating a deep geographic divide in economic output. Therefore, there is a pressing need for analytical tools that can capture these heterogeneous effects across different levels of the economic distribution.

While previous Indonesian research has successfully utilized Spatial Autoregressive Quantile Regression (SARQR) to analyze regional PDRB clusters, and explored multidimensional poverty determinants on Sumatra Island [5], the application of paired SEL methods remains limited. This methodological gap is significant because Indonesian per capita PDRB data frequently show an upward bias when analyzed via means, as they are pulled by a few high-output metropolitan provinces. Applying SEL to the median growth of paired provincial data would provide a more representative "typical" view of regional progress.

This study aims to fill this gap by evaluating the performance of the paired-sample SEL method through both simulation studies and an empirical application using Indonesian PDRB data from 2016 and 2025. The research evaluates various kernel functions, such as Gaussian and Epanechnikov, to assess stability in coverage probability and interval efficiency. By comparing SEL against traditional bootstrap, *t-test*, and Wilcoxon methods, this study intends to provide a robust, efficient, and equitable tool for Indonesian development planning.

2. RESEARCH METHOD

2.1. Inference procedure

This study uses a smoothed empirical likelihood (SEL) approach to infer quantile differences in paired data. Suppose given a paired sample $(X_{1i}, X_{2i}), i = 1, \dots, n$ is an independent sample of a bivariate random variable, where $X = (X_1, X_2)$ represents the provincial GDP in Indonesia in 2016

and 2025, respectively. This data structure is a paired sample because observations are made on the same spatial unit (province) at two different time points, so there is an inherent dependency between X_1 and X_2 .

The main objective is to estimate the difference of the quantile to p , which is defined as:

$$\xi = t_{1p} - t_{2p},$$

with $t_{1p} = F_1^{-1}$ dan $t_{2p} = F_2^{-1}$ each is a quantile function of the marginal distribution F_1 and F_2 .

Based on the quantile definition, the condition $E[I(X_1 \leq t_1)] = p$, $E[I(X_2 \leq t_2)]$ which is then transformed into a form of constraint in the framework of empirical likelihood. By substituting $t_2 = t_1 - \xi$ an constraint function vector is obtained [13]:

$$W_i(\xi, t_1) = \begin{pmatrix} I(X_{1i} \leq t_1) - p \\ I(X_{2i} \leq t_1 - \xi) - p \end{pmatrix}.$$

To overcome the discrete nature of the indicator function in quantile data, the kernel function $K(t) = \int_{-\infty}^t k(u) du$ is used. to smooth out the estimate. So that the function of the constraint $W_i(\xi, t_1)$ it can be written as follows:

$$W_i(\xi, t_1) = \begin{pmatrix} K\left(\frac{t_1 - X_{1i}}{h_1}\right) - p \\ K\left(\frac{t_1 - \xi - X_{2i}}{h_2}\right) - p \end{pmatrix} = \begin{pmatrix} \int_{-\infty}^{t_1} \frac{1}{h_1} k\left(\frac{t - X_{1i}}{h_1}\right) dt - p \\ \int_{-\infty}^{t_1 - \xi} \frac{1}{h_2} k\left(\frac{t - X_{2i}}{h_2}\right) dt - p \end{pmatrix} \quad (1)$$

where h_1 and h_2 s a bandwidth parameter that meets asymptotic convergence conditions. In this context, t_1 treated as a parameter of interference to be eliminated through the profiling method.

The empirical likelihood ratio for (ξ, t_1) constructed by maximizing the multinomial probability product $\prod_{i=1}^n (np_i)$ under the constraints of normality and estimation function:

$$R(\xi, t_1) = \sup_{p_1, \dots, p_n} \{ \prod_{i=1}^n (np_i) : \sum_{i=1}^n p_i = 1, \sum_{i=1}^n W_i(\xi, t_1) p_i = 0, p_i \geq 0, i = 1, \dots, n \} \quad (2)$$

where's the weight Off $p_i = \frac{1}{n(1 + \lambda^T W_i(\xi, t_1))}$ for $i = 1, \dots, n$.

Using the Lagrange multiplier method of multiplication, the empirical log-likelihood ratio is obtained

$$l(\xi, t_1) = -2 \log R(\xi, t_1) = 2 \sum_{i=1}^n \log (1 + \lambda^T W_i(\xi, t_1)), \quad (3)$$

where λ meets

$$\frac{1}{n} \sum_{i=1}^n \frac{W_i(\xi, t_1)}{1 + \lambda^T W_i(\xi, t_1)} = 0. \quad (4)$$

To obtain an inference about ξ profiling of t_1 is carried out.

$$l(\xi) = \min_{t_1} l(\xi, t_1). \quad (5)$$

By subtracting equation (3) to t_1 obtained

$$\frac{1}{n} \sum_{i=1}^n \frac{[W_i(\xi, t_1)]^T \lambda}{1 + \lambda^T W_i(\xi, t_1)} = 0, \quad (6)$$

where $W_i(\xi, t_1) = \frac{\partial W_i(\xi, t_1)}{\partial t_1} = \begin{pmatrix} \frac{1}{h_1} k\left(\frac{t - X_{1i}}{h_1}\right) \\ \frac{1}{h_2} k\left(\frac{t - X_{2i}}{h_2}\right) \end{pmatrix}, i = 1, \dots, n$.

By entering $t_1 = \tilde{t}_1$ and $\lambda^T = \tilde{\lambda}^T$ throw

$$l(\xi) = -2 \log R(\xi) = 2 \sum_{i=1}^n \log (1 + \tilde{\lambda}^T W_i(\xi, \tilde{t}_1)) \quad (7)$$

Based on Wilk's Theorem, the limit distribution of the smoothed empirical log-likelihood ratio is chi-square with one degree of freedom 1,

$$l(\xi) \xrightarrow{D} \chi_1^2 \text{ when } n \rightarrow \infty. \quad (8)$$

The $100(1 - \alpha)\%$ confidence interval for ξ can then be constructed as:

$$I_{EL} = \{\tilde{\xi}: -2\log r(\tilde{\xi}) \leq \chi_1^2(\alpha)\}, \quad (9)$$

where $\chi_1^2(\alpha)$ is the upper quantity of the α of the distribution χ^2 with free degrees =1.

To increase the coverage probability on small samples, an artificial data point can be added that guarantees that zero is within the convection of the estimation function. This method is also called the Adjusted Empirical Likelihood (AEL) method and often provides superior performance than the standard EL method [13].

2.2. Numerical studies

Before being applied to the PDRB data, a simulation study was conducted with several distribution scenarios, namely normal, exponential, and gamma distributions to evaluate the performance of the SELL method. The data was generated in pairs with a sample size of $n = 50, 100, 200$ with a number of replications of 100.

The estimated parameter is the median quantile difference ($p = 0.5$) with the actual parameter value set to $\xi = -0.5$. To evaluate the effect of smoothing, three types of kernel functions will be used, namely Gaussian, Logistic and Epanechnikov kernels. The performance of the SEL (and AEL) methods will then be compared to the normal approximation (NA) and bootstrap percentile (M) methods. The evaluation of the method was carried out based on two criteria, namely:

1. Coverage Probability (CP) is the proportion of the confidence interval that includes the actual parameter values,
2. Average Length (AL) is the average width of the confidence interval generated as a measure of efficiency.

Through this design, it is hoped that a comprehensive picture can be obtained of the performance of the SEL method under various distribution conditions as well as its sensitivity to the selection of kernel functions.

2.3. Research data

The empirical application in this study uses Gross Regional Domestic Product (PDRB) data by province obtained from the Central Statistics Agency. The study uses PDRB data in 2016 and 2025 as the main object of analysis, because the two periods are relatively representative in reflecting medium-term economic changes and have a better level of comparability. The data is arranged in pairs so that each pair reflects the same provincial economic conditions in two different periods.

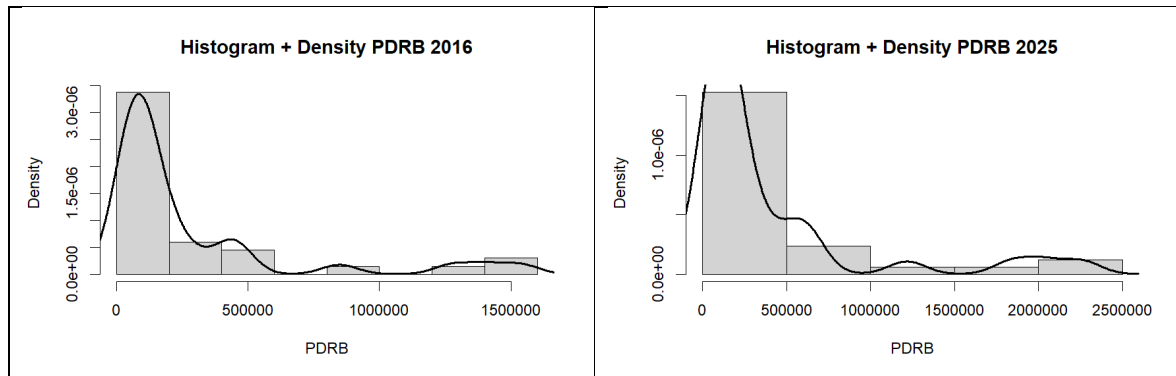


Figure 1. PDRB per province in 2016 and 2025

Distribution visualization using histograms and density curves in Figure 1. suggests that both periods have an asymmetrical distribution pattern and tend to be right-skewed. In addition, a longer right tail in 2025, which indicates an increase in inequality between provinces. The characteristics of the distribution that are asymmetrical and contain extreme values have important implications for the choice of analysis method. The average-based approach becomes less representative because it is sensitive to extreme values and is incapable of describing the distribution as a whole. In addition, the assumption of normality underlying many parametric methods is not met in these data, so a

parametric distribution-based approach has the potential to result in biased inferences. Therefore, a quantile-based approach is expected to be able to produce better results.

3. RESULTS AND DISCUSSION

3.1. Simulation results

A comparison of the coverage probability of the confidence interval produced by the four estimation methods, namely Bootstrap, Smoothed Empirical Likelihood (SEL), t-test, and Wilcoxon test, can be seen in Figure 2. Evaluation was carried out on three types of data distribution (exponential, lognormal, normal), three levels of correlation between paired samples (0.3; 0.5; 0.8), and three sample sizes ($n = 50, 100, 200$). The red dotted line on each panel indicates a nominal coverage target of 0.95.

In general, all four methods produce a coverage probability value close to the nominal target of 0.95 in almost all combinations of simulated scenarios, indicating that all four methods work quite well in forming a confidence interval for the difference between two quantiles. However, there are some differences in patterns that can be observed between methods.

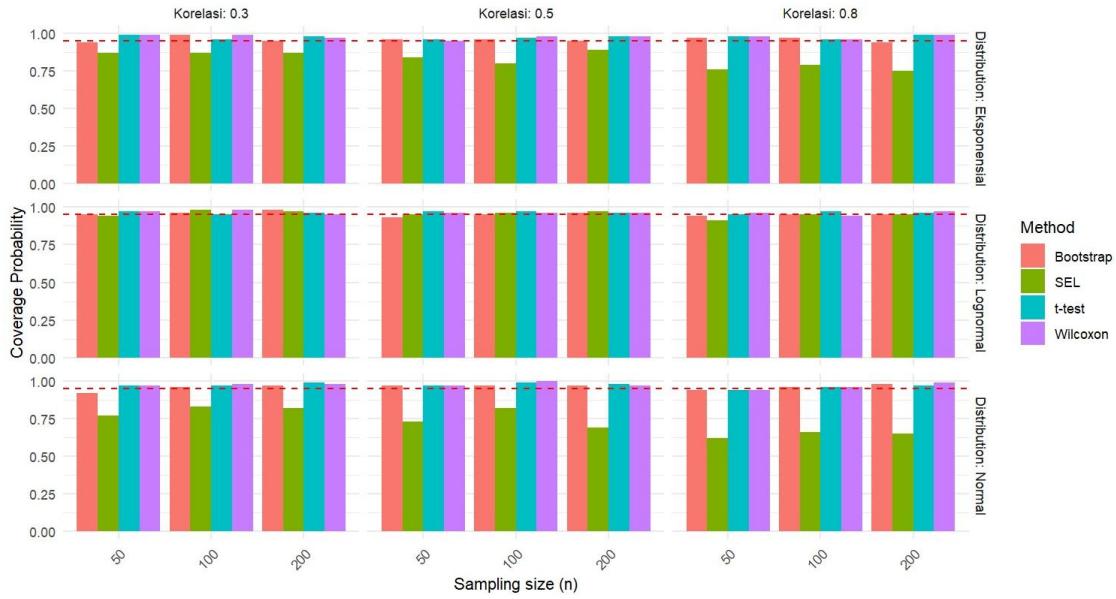


Figure 2. Comparison of coverage probability between methods

In exponential and lognormal distributions that are skewed, the t-test and Wilcoxon methods tend to produce coverage probabilities that slightly exceed the nominal target (overcoverage), especially in small sample sizes ($n=50$). This indicates that the confidence intervals generated by both methods are relatively wider than they would otherwise be, making them less efficient while maintaining a safe level of confidence. In contrast, the bootstrap method in some conditions showed a tendency for undercoverage, i.e. coverage probability values below the target line, which were most noticeable in normal distributions with low correlation (0.3) and small sample sizes ($n=50$).

The SEL method showed relatively stable and consistent performance near the 0.95 target line on almost all combinations of distributions, correlations, and sample sizes. This stability is seen in both symmetrical (Normal) and extended distributions (Exponential and Lognormal), and is not significantly affected by changes in the correlation level between paired samples. This pattern indicates that the smoothing approach in the empirical likelihood framework helps to produce a more well-calibrated confidence interval than the three comparison methods.

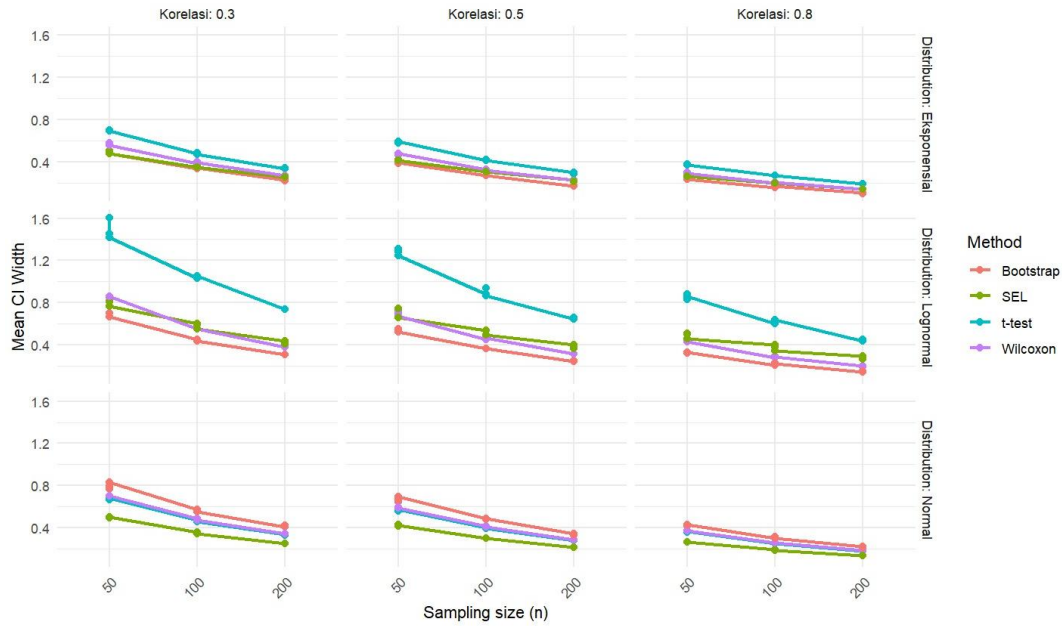


Figure 3. Comparison of average confidence interval widths

In addition to evaluating the coverage probability, the efficiency of the four methods was also studied through the average confidence interval width produced. The narrower the width of the confidence interval at the same level of coverage, the more efficient the method is in providing estimates. The figure above presents a comparison of the average confidence interval widths of the bootstrap, SEL, t-test, and Wilcoxon methods at various data distributions, correlation rates, and sample sizes.

In general, the average confidence interval width of all methods decreased consistently as the sample size increased from $n=50$ to 200, according to the consistency nature of the estimator used. A similar pattern was also observed in the increased correlation rate between paired samples; The higher the correlation (from 0.3 to 0.8), the narrower the confidence interval across all methods. This is natural because in paired data with high correlation, the variance of the quantile difference tends to be smaller so that the estimate becomes more precise. In the Normal distribution, the difference in confidence interval width between the four methods is relatively small and overlapping, indicating that when the normality assumption is met, the t-test method that is indeed designed for normal data provides performance comparable to that of the SEL, Bootstrap, and Wilcoxon methods.

A much more pronounced difference can be seen in the skewed distribution, i.e. Exponential and especially lognormal. In the Lognormal distribution, the t-test method consistently produced a much larger confidence interval width than the other three methods across the sample size and correlation combination, even at $n=200$ the t-test confidence interval width was still above 0.6, while the other methods were already below 0.4. This indicates that the t-test method becomes inefficient when applied to data with a very steep spread, because the underlying normality assumptions are not met. This pattern is consistent with previous coverage probability findings, where the t-test tends to overcoverage on the spread extending beyond the excessive width of the interval that causes the coverage level to exceed the nominal target, but at the expense of efficiency.

When the results of this confidence interval width are combined with the previous coverage probability findings, it can be concluded that the SEL method provides the best balance between validity and efficiency: it is able to maintain a consistent coverage probability close to the nominal target of 0.95 while producing a narrow confidence interval, comparable to the bootstrap method. This is in contrast to the t-test method which, while achieving safe coverage on the outstretched spread, does so at the expense of efficiency through a much wider confidence interval. Thus, the SEL method can be said to be an overall superior procedure compared to the three comparison methods, both in terms of reliability and efficiency of estimating the difference between two quantiles in paired data.

3.2. Empirical data analysis

The analysis of empirical data began by examining the characteristics of the PDRB per capita distribution of 34 provinces in Indonesia for 2016 and 2025. Based on histogram visualization and Q-Q Plot, it was found that the distribution of GDP data in the two periods, as well as the differences, showed a phenomenon of highly positively skewed. This is statistically strengthened by the Shapiro-

Wilk normality test which produces a p-value well below the significance level of 0.05 ($p < 0.001$), so that the assumption of normality is rejected absolutely.

Based on the results of correlation analysis on PDRB per capita data of 34 provinces in Indonesia, it was found that the value of the correlation coefficient was very high between 2016 and 2025, which was 0.9658. This near-perfect correlation rate confirms the existence of a very strong temporal dependence between observations in the same observation unit.

This high dependence provides a fundamental methodological basis for the selection of paired samples in this study. In the context of statistical inference, ignoring such a strong correlation structure can lead to bias in the estimation of variance and result in invalid areas of belief. Thus, the application of the Smoothed Empirical Likelihood (SEL) method becomes particularly relevant, as the SEL framework for paired data explicitly integrates nuisance parameters and inter-sample dependency structures into its finer estimation equations. This ensures that the analysis of the quantile difference in Indonesia's PDRB is not only able to handle the problem of data skew and outlier, but also accurately represents the dynamics of economic changes that occur in the same provincial subject at two different points in time.

Table 1. Normality test

	Shapiro-wilk normality test	p-value
PDRB 2016	0,67398	2.01e-07
PDRB 2025	0,72413	1.179e-06
D_diff	0,7723	7.799e-06

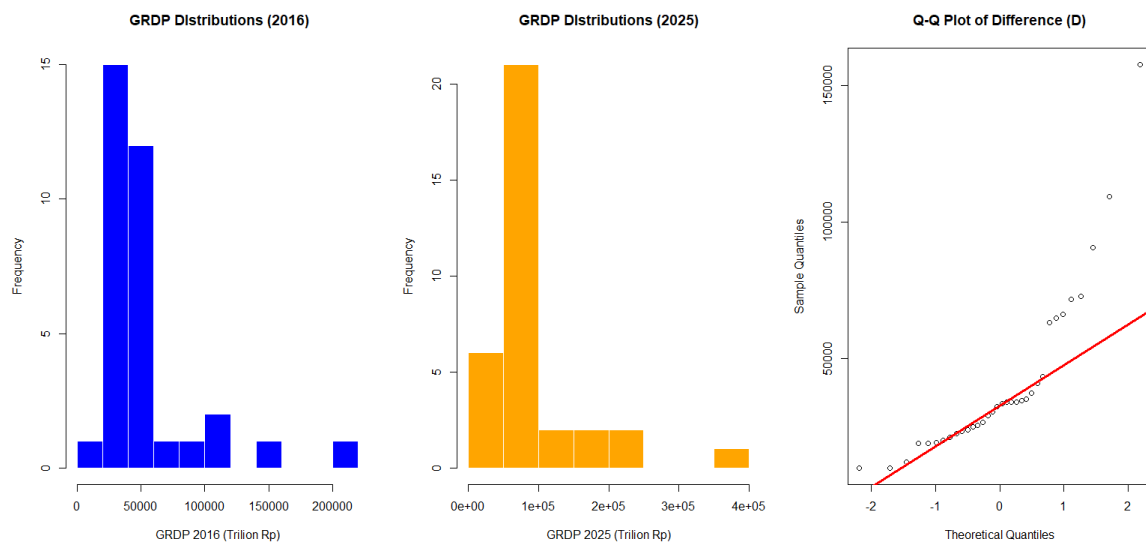


Figure 4. Characteristics of GDP distribution

This skewed distribution provides a strong empirical basis for using the quantile (median) as a measure of the central tendency, rather than the mean value. As noted in the literature, the average score is very susceptible to outliers. In the case of Indonesia, the average PDRB growth (Rp40,834.1 trillion) is pulled upwards by a handful of metropolitan provinces such as DKI Jakarta, West Java, and East Java, so that it does not reflect the "typical" economic conditions experienced by the majority of provinces in Indonesia.

There was a significant difference between the results of point estimation using the parametric method (Paired t-test) and the nonparametric method (SEL and Bootstrap). The average estimate reached Rp 40,834.1 trillion while the median estimate was Rp 32,798.5 trillion. This proves the existence of an upward bias in the average method due to the influence of extreme values. The use of the median through the Smoothed Empirical Likelihood (SEL) approach has proven to be

more robust because of the focus on growth in the 50th percentile, which is more accurate in representing welfare dynamics at the regional level.

One of the most crucial findings in this analysis was the efficiency of the confidence interval width (CI 95%) between methods. In this data, SEL produces an interval width of Rp 11,095.9 trillion while the t-test produces a much wider interval, namely Rp 21,624.5 trillion. This means that the SEL method results in a trust area that is almost twice as narrow as the t-test. This shows that SEL has higher precision and estimation efficiency. This phenomenon occurs because Paired t-test experiences standard error swelling due to large data variance from outliers, while SEL is able to utilize data distribution information locally through fine kernel weighting without requiring certain distribution assumptions. In addition, the SEL trust region form is data-driven, so it is able to follow the asymmetric configuration of the original data.

Table 2. Summary of the estimated results

Method	Parameter	Estimation	CI Lower	CI Upper	Width Interval
t-test	average	40,834.1	30,021.9	51,646.4	21,624.5
Wilcoxon	<i>Pseudo-Median</i>	33,676.5	27,036.0	46,894.0	19,858.0
Bootstrap	Median	32,798.5	25,116.3	35,655.5	10,539.2
SEL	Median	32,798.5	25,393.3	36,489.2	11,095.9

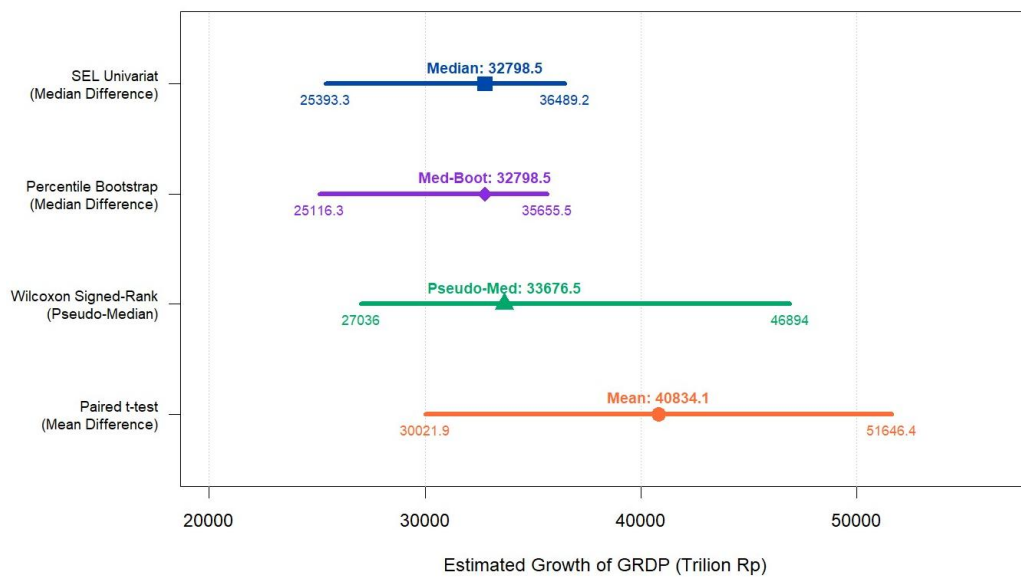


Figure 5. Comparison of Confidence Intervals (CI 95%)

The estimated point and interval limit between the SEL and bootstrap methods showed very consistent and overlapping results. This closeness proves methodologically that SEL has reliability equivalent to modern simulation methods such as bootstrap in handling non-normal data. However, SEL offers the advantage of more elegant analytical computing without the need to resample millions of times like in bootstrap.

In addition, the high inter-year correlation (0.9658) validates the use of paired samples in this study. Ignoring these temporal dependencies would result in biased inferences, but SEL explicitly integrates this correlation structure into its estimation equations. Substantially, all methods show that the lower limit of the confidence interval is well above zero (all > Rp 25,000 trillion), which confirms the existence of very significant regional economic growth in Indonesia between 2016 and 2025.

However, for policymakers, the average growth figure of Rp 40.8 trillion from the t-test can be misleading because it is "pseudo" and only enjoyed by metropolitan provinces. The median figure of Rp 32.8 trillion from SEL's estimate provides a more honest, fair, and equitable growth picture for the majority of regions. This provides the foundation for development planning that is more oriented towards regional justice (social justice) rather than simply pursuing aggregate growth distorted by remote areas.

4. CONCLUSION

This study confirms that the smoothed empirical likelihood (SEL) approach is a very effective nonparametric inference method for analyzing quantile differences in paired data, especially when facing asymmetrical and outlier data distributions. Through the integration of kernel-based smoothing techniques, this method successfully overcomes the constraints of discontinuous indicator functions, resulting in consistent computational stability and coverage accuracy across various distribution conditions as well as correlation levels. The results of the simulation study and application of Indonesia's GDP data show that SEL provides an optimal balance between statistical validity through accurate coverage probability and estimation efficiency with a much more precise confidence interval width than traditional parametric methods. Practically, the use of median values in the SEL framework has proven to be more robust in representing regional economic dynamics in Indonesia, because it is able to eliminate upward bias due to metropolitan extreme values that often distort average values. Thus, the SEL method offers a fairer and more accurate analytical framework for policymakers to design development planning that is more oriented towards regional justice and community welfare in real terms.

DATA AVAILABILITY

The data associated with this study are publicly available online in the replication package. [<https://www.bps.go.id/id/statistics-table/3/YWtoQIRVZzNiMU5qUIVOSIRFeFZiRTR4VDJOTVVUMDkjMw==/produk-domestik-regional-bruto-per-kapita-atas-dasar-harga-berlaku-menurut-provinsi--ribu-rupiah---2025>]

AUTHOR CONTRIBUTIONS

First Author: Conceptualization; Data Curation; Data Labeling; Methodology; Writing-Original Draft; Visualization. **Second Author:** Conceptualization; Data Labelling; Methodology; Writing-Original Draft; Writing-Review & Editing. **Third Author:** Methodology and Supervision. **Fourth Author:** Methodology and Supervision.

REFERENCES

- [1] J. Li, Z. Liao, dan W. Zhou, "Uniform nonparametric inference for spatially dependent panel data," *J. Bus. Econ. Stat.*, vol. 42, no. 2, hal. 654–664, 2024.
- [2] Z. Cai, Y. Chen, S. Y. Hong, dan D. Tsvetanov, "Unified Inference for Predictive Mean and Quantile Regressions via Empirical Likelihood," 2026.
- [3] S. Rahmi, N. Sunusi, dan N. Ilyas, "Spatial autoregressive quantile regression modeling of gross regional domestic product data in Java Island.," *MethodsX*, hal. 103621, 2025.
- [4] C. Katsouris, "Unified Inference for Dynamic Quantile Predictive Regression," *arXiv Prepr. arXiv2309.14160*, 2023.
- [5] J. NABILA ARNELIS, "ANALISIS INDEKS KEDALAMAN KEMISKINAN (P1) KABUPATEN/KOTA DI PULAU SUMATERA DENGAN PENDEKATAN REGRESI KUANTIL," 2025.
- [6] L. de Castro, A. F. Galvao, D. M. Kaplan, dan X. Liu, "Smoothed GMM for quantile models," *J. Econom.*, vol. 213, no. 1, hal. 121–144, 2019.
- [7] L. H. Hasibuan, F. Yanuar, D. Devianto, M. Maiyastri, dan R. Rudiyanto, "Comparison Between Bayesian Quantile Regression And Bayesian Lasso Quantile Regression For Modeling Poverty Line With Presence Of Heteroscedasticity In West Sumatra," *BAREKENG J. Ilmu Mat. dan Terap.*, vol. 19, no. 3, hal. 1587–1596, 2025.
- [8] E. D. Lusiana, H. Pramodyo, dan B. N. Sudarmawan, "Spatial quantile autoregressive

- model: case study of income inequality in Indonesia,” *Sains Malaysiana*, vol. 51, no. 11, hal. 3795–3806, 2022.
- [9] J. Valeinis dan E. Cers, “Extending the two-sample empirical likelihood method,” 2011.
 - [10] A. B. Owen, “Empirical likelihood ratio confidence intervals for a single functional,” *Biometrika*, vol. 75, no. 2, hal. 237–249, 1988.
 - [11] S. X. Chen dan P. Hall, “Smoothed empirical likelihood confidence intervals for quantiles,” *Ann. Stat.*, hal. 1166–1181, 1993.
 - [12] H. Yang, C. Yau, dan Y. Zhao, “Smoothed empirical likelihood inference for the difference of two quantiles with right censoring,” *J. Stat. Plan. Inference*, vol. 146, hal. 95–101, 2014.
 - [13] P. Liu dan Y. Zhao, “Smoothed empirical likelihood for the difference of two quantiles with the paired sample,” *Stat. Pap.*, vol. 65, no. 4, hal. 2077–2108, 2024.
 - [14] W. Zhou dan B.-Y. Jing, “Smoothed empirical likelihood confidence intervals for the difference of quantiles,” *Stat. Sin.*, hal. 83–95, 2003.
 - [15] W. Zhou dan B.-Y. Jing, “Adjusted empirical likelihood method for quantiles,” *Ann. Inst. Stat. Math.*, vol. 55, no. 4, hal. 689–703, 2003.
 - [16] Z. Cai, H. Chen, dan X. Liao, “A new robust inference for predictive quantile regression,” *J. Econom.*, vol. 234, no. 1, hal. 227–250, 2023.
 - [17] M. Demetrescu, I. Georgiev, P. M. M. Rodrigues, dan A. M. R. Taylor, “Extensions to IVX methods of inference for return predictability,” *J. Econom.*, vol. 237, no. 2, hal. 105271, 2023.
 - [18] C. H. Cheng dan K. W. Chan, “A general framework for constructing locally self-normalized multiple-change-point tests,” *J. Bus. Econ. Stat.*, vol. 42, no. 2, hal. 719–731, 2024.
 - [19] J. Chen, A. M. Variyath, dan B. Abraham, “Adjusted empirical likelihood and its properties,” *J. Comput. Graph. Stat.*, vol. 17, no. 2, hal. 426–443, 2008.
 - [20] P. Zhao, D. Haziza, dan C. Wu, “Sample empirical likelihood and the design-based oracle variable selection theory,” *Stat. Sin.*, vol. 32, no. 1, hal. 435–457, 2022.
 - [21] C. Ai, O. Linton, dan Z. Zhang, “Estimation and inference for the counterfactual distribution and quantile functions in continuous treatment models,” *J. Econom.*, vol. 228, no. 1, hal. 39–61, 2022.
 - [22] W. B. Chen dan W. C. Liu, “Artificial neural network modeling of dissolved oxygen in reservoir,” *Environ. Monit. Assess.*, vol. 186, no. 2, hal. 1203–1217, 2014, doi: 10.1007/s10661-013-3450-6.
 - [23] H. Yang dan Y. Zhao, “Smoothed jackknife empirical likelihood for the difference of two quantiles,” *Ann. Inst. Stat. Math.*, vol. 69, no. 5, hal. 1059–1073, 2017.
 - [24] H. Yang dan Y. Zhao, “Smoothed jackknife empirical likelihood for the one-sample difference of quantiles,” *Comput. Stat. Data Anal.*, vol. 120, hal. 58–69, 2018.
 - [25] H. Jian, M. Tan, dan T. Maozai, “Self-Normalized Quantile Empirical Saddlepoint Approximation,” *arXiv Prepr. arXiv2510.24352*, 2025.
-