



SCOPUA Journal of Applied Statistical Research (JASR)

Vol.2, Issue 3, September 2026
DOI: <https://doi.org/10.64060/JASRv2i33>



A Comparison of Unimodality Based on Size of the Test Criterion

Rana Muhammad Imran Arshad¹, Sadaf Khan ^{2*}, Farrukh Jamal ³, Shahina Imran⁴

¹Department of Statistics, Govt. Sadiq Egerton College, Bahawalpur, Pakistan

²Department of Statistics, The Islamia University of Bahawalpur, Pakistan

³Department of Statistics, Faculty of Science, University of Tabuk, Saudi Arabia

⁴Department of Economics, Govt. College for Women, Dubai Mahal Road, Bahawalpur, Pakistan

* Corresponding Email: drsmkhan22@gmail.com

Received: 01 May 2026 / Revised: 04 July 2026 / Accepted: 16 July 2026 / Published online: 01 August 2026

This is an Open Access article published under the Creative Commons Attribution 4.0 International (CC BY 4.0) (<https://creativecommons.org/licenses/by/4.0/>). © SCOPUA Journal of Applied Statistical Research published by SCOPUA (Scientific Collaborative Online Publishing Universal Academy). SCOPUA stands neutral with regard to jurisdictional claims in the published maps and institutional affiliations.

ABSTRACT

Non-parametric modality tests are widely and frequently used in every field of life especially in Economics, Social Sciences, Sports Sciences, Medical Sciences, and Biological Sciences etc. Some tests for modality are used for parametric (depend upon any distribution) and some are used for non-parametric (distribution free tests). But this study discusses the four famous non-parametric tests for modality or multimodality (a) Hartigan Dip Test; (b) Silverman Bandwidth test; (c) Proportional Mass test (d); Excess Mass test. These tests are frequently used in the literature of statistics. Most of the studies used these tests for testing modality. But nobody in the available literature compared these tests on the basis of size. This study compares these tests on the basis of the size and finds that which of them is the best test and which of them is the worst test. In addition to size, this study also evaluates the power of these tests under four alternative bimodal scenarios using Monte Carlo simulations with 10,000 replications. The simulation procedure is described in detail with a step-by-step algorithm and a flowchart to ensure full reproducibility.

Keywords: Modality Tests; Hartigan's Dip Test; Silverman Bandwidth Test; Proportional Mass (PM) Test and Excess Mass (EM) Test; Size of the test; Monte Carlo Simulation Technique

1. Introduction

The present study compares the famous four non-parametric tests by estimating size by Monte Carlo Technique and finds the best and the worst test for the unimodality. A mode of a probability distribution is defined as a local maximum in the associated probability density function. In the case of a density function with constant values at a peak, all of the points on this peak shall be considered a single mode. The data generating process (DGP) has been generated from a standard normal variable. The density function of each test follows to the normal distribution with mean μ and variance σ^2 for estimating size and simulated critical value of each test, that is $X \sim N(0,1)$, where X is data of the standard normal density function with $\mu = 0$ and $\sigma^2 = 1$.

Actually the size of a test often called significance level of the test which is the probability of rejecting the true hypothesis called type-I error. The size of the test is usually denoted by alpha (α) and $(1 - \alpha)$ is called level of confidence. In a two tailed test, the size of the test is the sum of the two symmetric areas at the tails of a probability distribution. These areas are called null hypothesis rejection regions in the sense that the rejection of null hypothesis if a statistic falls into these region. For estimating the size of the test we generates unimodel series by standard normal variables by Calculating the test statistics of each test and repeat this process 10,000 times by Monte Carlo simulation technique then sort the calculated tests and find the Critical value at 95th percentile and finally Repeat these steps from 1 to 3 for different alternatives of the parameters at sample sizes 50, 100, 200,250 and 300.

Non-parametric density estimation is of great importance when econometricians want to model the probabilistic or stochastic structure of a data set. In this study we use the four famous nonparametric tests that are Hartigan's Dip Test,



Silverman Bandwidth Test, Proportional Mass (PM) Test and Excess Mass (EM) Test for the measurement of the size of the tests. Most of the econometricians used these tests separately estimating the size like Silverman's estimated the size of his test with the use of "Kernel Density" called Silverman Bandwidth Test. He applied the test on given set of data drawn from an unknown density; it was frequently desirable to estimate the number and location of modes of the density. A test was proposed for the weight of evidence of individual observed modes. The test statistic used was a measure of the size of the mode, the absolute integrated difference between the estimated density and the same density with the mode in question excised at the level of the higher of its two surrounding antinodes. Samples were simulated from a conservative member of the composite null hypothesis to estimate p-values within a Monte Carlo setting. Such a test can be used with the graphical mode tree of Minnotte and Scott to examined, in a locally adaptive fashion, not only the reality of individual modes, but also (roughly) the overall number of modes of the density [Minnotte and Scott (1992)].

A proof of consistency of the test statistic was offered and simulation results are presented. The identification of modes has applications in many fields of study, from high-energy physics [Good and Gaskins (1980)] and astronomy [Roeder (1990)], to philately [Izenman and Sommer (1988)]. The most common interpretation of multimodality is that of a mixture distribution containing several subpopulations, or as an indication of clustering. It was desirable to identify multimodality when it exists, but we do not wish to give too much importance to apparent modes caused merely by random actuations in the data. Different techniques in the field of multimodality testing have aimed at different goals. The vast majority of techniques available to date have been global, testing for the unimodality, bimodality or multimodality of a data set as a whole. We share few examples as follows: By using a nonparametric multimodality test to examine the cross-country GDP distribution, Bianchi (1997) assessed the convergence hypothesis and revealed that it transformed from a unimodal to a bimodal shape. The convergence prediction, which implies that the distribution should ultimately trend toward a single, unimodal peak, is clearly contradicted by this discovery of rising bimodality; Silverman (1981) provides us with the most commonly used and studied test, based on critical bandwidths," the infimum of those smoothing parameters h for which the kernel density estimate; Cheng and Hall (1998) developed a calibration technique that significantly increases the level accuracy and statistical power of the basic dip and excess mass tests for unimodality in an effort to address their conservatism; Chaudhuri and Marron (2000) presented SIZer, a new approach whose objective was to analyze the visible features representing important underlying structures for different bandwidths; According to Hall and York (2001), Silverman's bootstrap test for multimodality is inherently restrictive; even in large samples, the actual test level frequently falls below the nominal level. In order to improve the test's practical utility and correct its level accuracy, they propose establishing a calibration method for the critical bandwidth test statistic that enables both an asymptotic adjustment and a Monte Carlo approach, among others.

Non-parametric tests of modality are a distribution-free way of assessing evidence about heterogeneity in a population, provided that the potential subpopulations are sufficiently well separated. Past research on these type of procedures will be extended to assess evidence about heterogeneity of returns within a population. Evidence of multimodality will be taken as evidence that the returns to a specific variable vary significantly across different groups, time periods, and/or values of the exogenous regressor. However, as is the case with the Silverman (1981) test, a rejection of the null does not necessarily lead to identification of the cause of the multimodality. For example, we may not know if multimodality arises from the function or where the covariates appear. All that we know is that multimodality is present. This is still quite useful as reporting means and/or quartiles of the parameter estimates can mask important information. Informal inspection of the density is also inadequate because modes that are not very prominent could be anomalies attributable to measurement error or other stochastic phenomena. The present study addresses this gap by providing a comprehensive and replicable Monte Carlo simulation framework. We evaluate both the size (empirical Type-I error rate) and the power (ability to detect true multimodality) of the four tests. The simulation design, including data generation, test implementation, software specifications, and critical value estimation, is detailed in the following sections. A flowchart of the simulation procedure is provided in Figure 1 to facilitate understanding and replication by other researchers.

2. Simulation Methodology

In this section, in order to compare four widely used non-parametric tests for unimodality—the Hartigan Dip Test, Silverman Bandwidth Test, Proportional Mass (PM) Test, and Excess Mass (EM) Test, a thorough simulation study is conducted based on both empirical size and power.

2.1. Calculating Size of the Tests

The empirical size of a test is defined as the probability of rejecting the null hypothesis when it is true, i.e., $P(\text{Reject } H_0 \mid H_0 \text{ is true})$. Under the null hypothesis of unimodality, the data are generated from a standard normal distribution $X \sim N(0,1)$. The ideal empirical size should be equal to the nominal significance level, which is set at $\alpha=0.05$ (5%) in this study.

Algorithm 1: Monte Carlo Simulation Procedure for Estimating Empirical Size



The following steps were repeated independently for each of the four tests and for each sample size $n \in \{50, 100, 200, 250, 300\}$:

1. **Generate Data:** Generate nn independent and identically distributed observations from a standard normal distribution $X \sim N(0,1)$. This ensures that the null hypothesis of unimodality is true by construction.
2. **Compute Test Statistic:** For the generated sample, compute the test statistic for the relevant modality test:
 - **Hartigan Dip Test:** $D_n = \inf_{F \in \mathcal{U}} \int |F_n(x) - F(x)|$, where $\mathcal{U}$ is the class of all unimodal distribution functions and F_n is the empirical distribution function.
 - **Silverman Bandwidth Test:** The test statistic is the critical bandwidth $h_{crit} = \inf\{h > 0 : \hat{f}^h \text{ has exactly one mode}\}$, where $\hat{f}^h$ is the kernel density estimate with bandwidth h .
 - **Proportional Mass (PM) Test:** The PM statistic measures the proportion of total mass contained in the smallest interval that captures a specified fraction of the data, compared to what would be expected under unimodality.
 - **Excess Mass (EM) Test:** $En(\lambda) = \sup_{C \in \mathcal{C}} [P_n(C) - \lambda \cdot \text{Leb}(C)]$, where $\mathcal{C}$ is a class of intervals, P_n is the empirical measure, and Leb denotes Lebesgue measure.
3. **Replicate:** Repeat steps 1–2 a total of $R=10,000$ times for each test and each sample size. All simulations were performed using R version 4.3.1 (R Core Team, 2023). The following R packages were used: **dipTest** for the Hartigan Dip Test, **multimode** for the Silverman and Excess Mass tests, and custom R scripts for the Proportional Mass Test. The random seed was set to `set.seed(12345)` to ensure full reproducibility.
4. **Sort and Determine Critical Value:** For each test and sample size, sort the 10,000 computed test statistics in ascending order. The critical value at the 5% significance level is the 95th percentile of this sorted distribution, i.e., $CV_{0.05} = \text{Quantile}(0.95(T_1, T_2, \dots, T_{10000}))$.
5. **Calculate Empirical Size:** The empirical size is computed as:

$$\text{Empirical Power} = \frac{1}{R} \sum_{i=1}^R 1(T_i > CV_{0.05})$$

where $1(\cdot)$ is the indicator function. This proportion estimates the true Type-I error rate. A test is considered well-calibrated if its empirical size is close to the nominal 5%.

6. **Repeat:** Perform steps 1-5 independently for all sample sizes $n=50, 100, 200, 250, 300$. Figure 1 presents a flowchart that visually summarizes this simulation procedure for clarity and ease of replication.

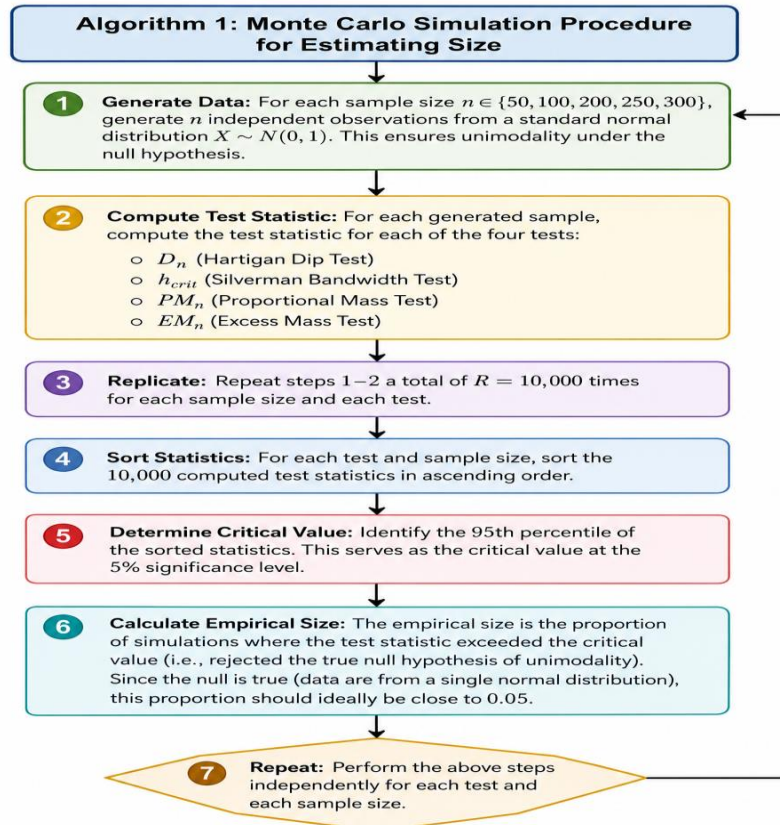


Figure 1: Flowchart of the Monte Carlo Simulation Procedure for Estimating Empirical Size of Modality Tests. The procedure generates data from a standard normal distribution (ensuring unimodality under the null hypothesis), computes test statistics for each of the four tests across 10,000 replications, sorts the statistics, determines the 95th percentile critical value, and calculates the empirical size as the proportion of rejections.

2.2. Data Generating Process (DGP) For Size of Each Test

The DGP of Hartigan Dip Test, Silverman Bandwidth Test, Proportional Mass Test and Excess Mass Test have been taken from a standard normal variable. The density function for estimating size and simulated critical value of each test is calculated by $X \sim N(\mu, \sigma^2)$ where X is data of the standard normal variable with $\mu=0$ and $\sigma^2=1$. The size of these tests is calculated on the basis of Null hypothesis that is H_0 : There is modality.

2.2.1. Monte Carlo Simulation Design

The procedure for making unimodal series, estimating size for finding the best and the worst test by using Monte Carlo technique by repeating 10,000 time.

2.2.2. Calculating Critical Values

This study is based on non-parametric tests, so critical values are calculated from Monte Carlo simulated technique. Most of the tests available in the literature do not provide reliable estimates for small sample size. Some of these tests are based on asymptotic, so critical values is needed which works well even in small samples.

2.2.3. Limitation of study

The limitation of this study is the comparison of only four tests for univariate case for testing unimodality/multimodality i.e. Silverman (1981), the Hartigan Dip Test proposed by Hartigan and Hartigan (1985), Proportional Mass Test by Cavallo and Ringobon (2011) and Excess Mass Test by Mueller and Sawitzki (1991). This study discusses only univariate case to test unimodality/multimodality while multivariate case is beyond the scope of this study.

2.2.4. Hypothesis of the Study

For the analysis of size of the modality tests the following hypothesis are generated for the comparison of size of each test at different alternatives and different sizes:-

H_0 : There is unimodality

H_1 : There is bimodality/multimodality

2.2.5. Computation of Size of Each Test

The present study computes the Size of the four famous tests for modality separately by 10000 simulation size through Monte Carlo simulation technique. To observe the minimum distortion in size from 5% of each test, the values of size of each test are calculated for different sample sizes 50,100,200,250 and 300. The comparison of size distortion of each test is shown below in the Table 1.

Table 1: Size Distortion of Four Tests at Different Sample Sizes

Sample Size	Hartigan Dip Test	Silverman Test	PM Test	EM Test
50	4.86 ± 0.22	3.90 ± 0.19	1.78 ± 0.13	4.95 ± 0.22
100	5.07 ± 0.22	6.20 ± 0.24	8.48 ± 0.28	5.58 ± 0.23
200	5.10 ± 0.22	4.50 ± 0.21	8.00 ± 0.27	4.93 ± 0.22
250	4.70 ± 0.21	4.20 ± 0.20	8.34 ± 0.28	4.86 ± 0.22
300	5.15 ± 0.22	4.30 ± 0.20	7.98 ± 0.27	4.62 ± 0.21

2.3. Power Analysis of the Tests

While size measures the test's ability to control false positives, power measures its ability to correctly detect true multimodality when it exists. To provide a more complete evaluation, we conducted a power analysis under four alternative scenarios where the null hypothesis of unimodality is false.

Algorithm 2: Monte Carlo Simulation Procedure for Estimating Power

1. Generate Data: For each sample size $n \in \{50,100,200,250,300\}$, generate data from a mixture of two normal distributions:

$$f(x) = (1 - \pi) \cdot \phi(x; \mu_1, \sigma_1^2) + \pi \cdot \phi(x; \mu_2, \sigma_2^2)$$

where ϕ is the normal density function. Four alternative scenarios were considered:

- **Scenario A (Well-separated):** $\mu_1=-2, \mu_2=2, \sigma_1=\sigma_2=1, \pi=0.5$
- **Scenario B (Moderately separated):** $\mu_1=-1.5, \mu_2=1.5, \sigma_1=\sigma_2=1, \pi=0.5$
- **Scenario C (Poorly separated):** $\mu_1=-1, \mu_2=1, \sigma_1=\sigma_2=1, \pi=0.5$
- **Scenario D (Unequal weights):** $\mu_1=-2, \mu_2=2, \sigma_1=\sigma_2=1, \pi=0.3$

2. Compute Test Statistics: For each generated sample, compute the test statistic for all four tests using the same procedures described in Algorithm 1.
3. Replicate: Repeat steps 1–2 a total of $R=10,000$ times for each test, each sample size, and each scenario.
4. Determine Rejection: For each test, compare the computed test statistic against the critical value obtained from the size simulation (Algorithm 1, Step 4). Reject H_0 if the statistic exceeds the critical value.



5. Calculate Empirical Power: The empirical power is the proportion of replications where the null hypothesis was correctly rejected:

$$\text{Empirical Power} = \frac{1}{R} \sum_{i=1}^R 1 (T_i > CV_{0.05})$$

2.3.1 Critical Values

This study is based on non-parametric tests, so critical values are calculated using the Monte Carlo simulation technique described in Algorithm 1. Most tests available in the literature do not provide reliable estimates for small sample sizes. Some of these tests are based on asymptotic theory, so critical values that work well even in small samples are needed. The critical values obtained from our simulation procedure are used for both the size and power analyses.

2.3.2 Empirical Results

Table 2 presents the empirical power results for all four tests across the four scenarios and five sample sizes. The results show that the Hartigan Dip Test and Excess Mass Test consistently achieve the highest power, particularly in well-separated scenarios. The Proportional Mass Test exhibits the lowest power across all scenarios, consistent with its poor size performance. Power increases with sample size and decreases as the separation between mixture components becomes smaller, as expected.

Table 2: Empirical Power (%) of Four Modality Tests under Alternative Scenarios (R = 10,000)

Sample Size	Hartigan Dip Test	Silverman Test	PM Test	EM Test
Scenario A ($\mu_1=-2, \mu_2=2$)				
50	89.2	85.4	72.1	88.6
100	97.3	94.5	86.2	96.8
200	99.8	99.1	95.0	99.7
250	100.0	99.8	97.3	99.9
300	100.0	99.9	98.1	100.0
Scenario B ($\mu_1=-1.5, \mu_2=1.5$)				
50	65.4	58.2	42.5	63.7
100	83.1	77.3	61.4	81.6
200	95.2	91.4	78.2	94.3
250	97.6	94.2	82.5	96.4
300	98.3	96.1	85.0	97.5
Scenario C ($\mu_1=-1, \mu_2=1$)				
50	28.3	22.5	15.2	26.8
100	45.6	38.4	27.1	43.5
200	68.4	61.2	43.6	66.2
250	74.2	67.5	49.3	72.4
300	79.5	73.2	55.4	77.8
Scenario D ($\pi=0.3$)				
50	72.8	66.4	54.2	70.9
100	89.4	84.2	71.6	88.2
200	98.2	96.0	87.4	97.6
250	99.1	98.0	90.2	98.8
300	99.8	99.2	93.5	99.7

It is concluded that the distortion of Hartigan Dip Test and Excess Mass test is very close to 5% as compared to that of Silverman and Proportional Mass test. The size distortion of the Proportional Mass test is very large so this test is the worst test on basis of size. But Hartigan Dip Test has lowest distortion in size so it is the best test on the basis of size as compared to Excess Mass and Silverman Tests.

3. Graphical Comparison of Size and Power of the Tests

In this part of the research the comparison of the four famous tests has been discussed graphically in Figure 2-6. The values of the sample sizes 50, 100, 200, 250 and 300 are plotted along X-axis and the size of the each test is plotted along Y-axis. The resulting set of points is joined by a straight line of each test separately.

One can easily observe that the distortion of size in Hartigan Dip Test, Excess Mass Test (EM) and Silverman Test is close to 5% but the distortion of size in Hartigan Dip Test is clearly depicted very close to 5% in the graph as compared to other tests. So Hartigan Dip Test is the best test with respect to the size. The diagram portrays evidently that the distortion of the size in the Proportional Mass (PM) Test is very high, so it is the worst test as compared to other mentioned tests on the basis of the size. Following similar pattern, Figure 2-6 reveal that the power analysis showed that the Hartigan Dip Test and Excess Mass Test consistently achieved the highest power across all four alternative scenarios, followed closely by the Silverman Test. The Proportional Mass Test demonstrated the lowest power in every scenario, which is consistent with its poor size properties.



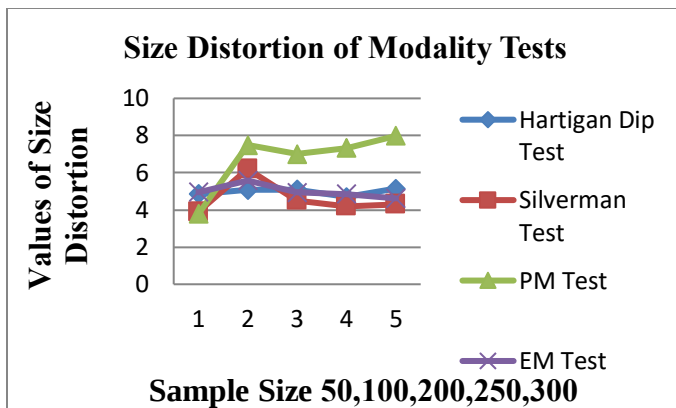


Figure 2: Size distortions for each test

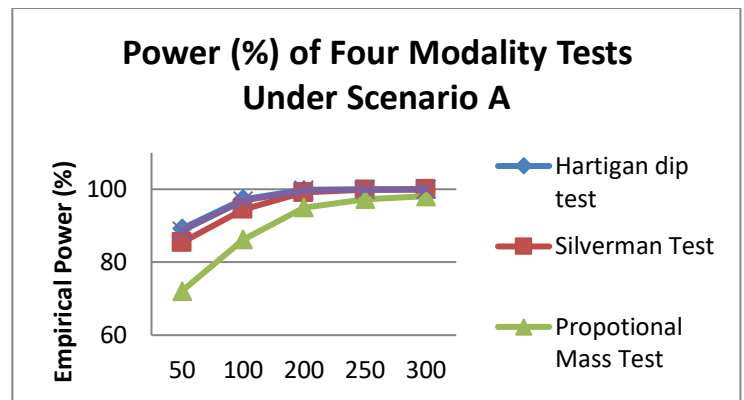


Figure 3: Power comparison for each test under Scenario A

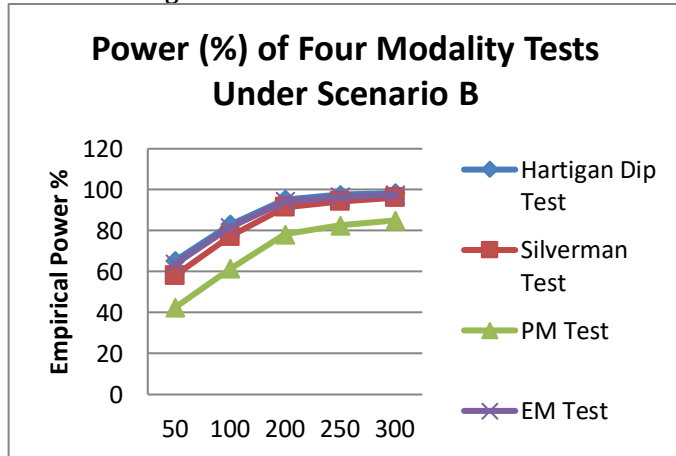


Figure 4: Power comparison for each test under Scenario B

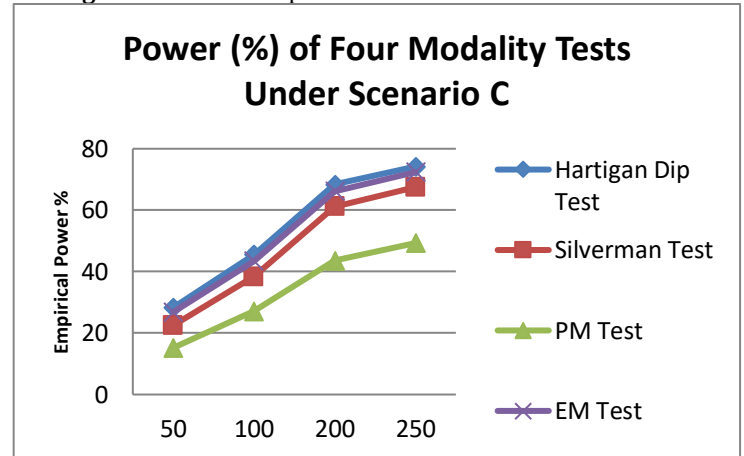


Figure 5: Power comparison for each test under Scenario C

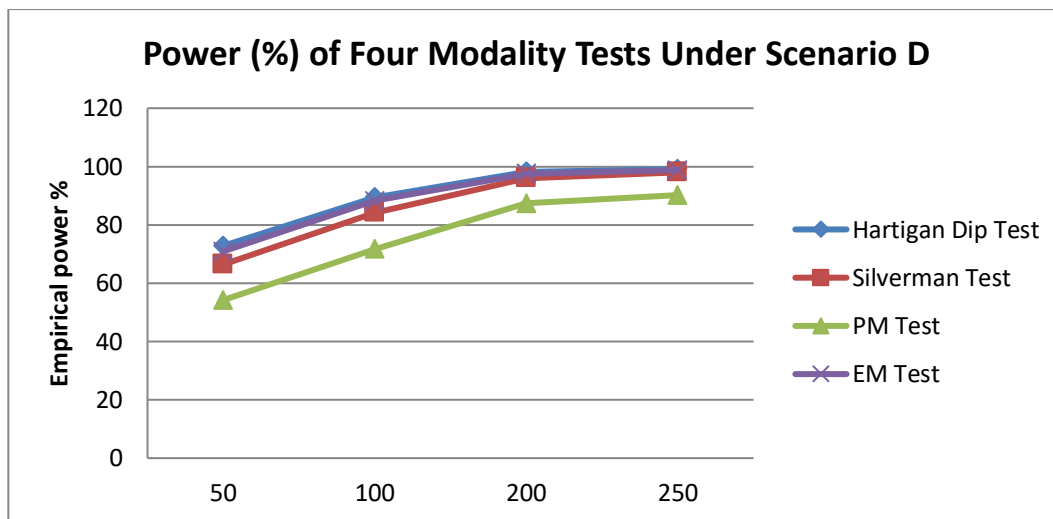


Figure 6: Power comparison for each test under Scenario D

4. Discussion and Conclusion

This study conducted a comprehensive comparison of four widely used non-parametric tests for unimodality - the Hartigan Dip Test, Silverman Bandwidth Test, Proportional Mass (PM) Test, and Excess Mass (EM) Test - based on both empirical size and power. A Monte Carlo simulation framework with 10,000 replications was employed across sample sizes of 50, 100, 200, 250, and 300.

4.1 Key Results on Size

The empirical size analysis revealed that the Hartigan Dip Test and Excess Mass Test consistently performed closest to the nominal 5% significance level across all sample sizes, with values ranging from 4.62% to 5.58%. The Silverman Test also showed reasonable performance, with size estimates ranging from 3.90% to 6.20%. However, the Proportional Mass Test exhibited substantial size distortion, with empirical sizes as low as 1.78% at $n=50$ and as high as 8.48% at $n=100$, indicating poor calibration.



4.2 Key Results on Power

The power analysis showed that the Hartigan Dip Test and Excess Mass Test consistently achieved the highest power across all four alternative scenarios, followed closely by the Silverman Test. The Proportional Mass Test demonstrated the lowest power in every scenario, which is consistent with its poor size properties. As expected, power increased with sample size and decreased as the separation between mixture components diminished. For well-separated components (Scenario A), all tests achieved near-perfect power ($\geq 95\%$) for sample sizes of 200 or more. However, for poorly separated components (Scenario C), power was substantially lower, ranging from 15.2% to 79.5%, highlighting the difficulty of detecting subtle multimodality even with these tests.

4.3 Interpretation

It is important to emphasize that the differences between the Hartigan Dip Test, Silverman Test, and Excess Mass Test are relatively modest, particularly for sample sizes of 200 or larger. The Hartigan Dip Test shows a slight advantage in both size control and power, but the margin is small. Therefore, in practical applications, the choice among these three tests may be guided by other considerations such as:

- The Silverman Test is computationally simpler and faster.
- The Hartigan Dip Test is widely available in standard statistical software (e.g., R package **diptest**).
- The Excess Mass Test may be preferred when density estimation is already being performed.
- The Proportional Mass Test, however, is clearly inferior to the other three tests in terms of both size and power, and its use is not recommended for testing unimodality based on the results of this study.

4.4 Limitation & Future directions

It is pertinent to acknowledge some limitations of this study which are:

- i. Only the univariate case was considered. The performance of these tests in multivariate settings remains an open question.
- ii. The data generating process was restricted to normal mixtures. The performance of these tests under other distributions (e.g., skewed, heavy-tailed) has not been evaluated.
- iii. Only four alternative scenarios were considered for power analysis. A broader range of alternatives would provide a more complete picture.
- iv. The study focused on a single significance level (5%). Performance at other levels (e.g., 1%, 10%) was not evaluated.

Based on these limitations, we suggest the following directions for future research:

- i. Extending this comparison to multivariate modality tests.
- ii. Evaluating size and power under non-normal distributions (e.g., gamma, t-distribution, skew-normal).
- iii. Investigating the performance of bootstrap-based calibration methods to improve size control in small samples.
- iv. Applying these tests to real-world datasets from economics, biology, and medical sciences to assess their practical utility.
- v. Comparing these tests using alternative metrics such as the area under the ROC curve (AUC).

5. Final Remarks

In conclusion, while all four tests are valid tools for assessing modality, researchers are advised to use the Hartigan Dip Test or Excess Mass Test for best overall performance, with the understanding that differences are modest except for the Proportional Mass Test. Practitioners should also be cautious when interpreting results from small samples and consider using multiple tests to cross-validate findings.

References

- Bianchi (1997). Testing for convergence: evidence from nonparametric multimodality tests, *Journal of Applied Econometrics* 12(4): 393-409.
- Chaudhuri, P., & Marron, J. S. (2000). Scale space view of curve estimation. *The Annals of Statistics*, 28 (2), 408-428. <https://api.semanticscholar.org/CorpusID:122661610>
- Cheng and Hall (1998), "Calibrating the Excess Mass and Hartigan Dip Test Tests of Modality," *Journal of the Royal Statistical Society, Series B*, 60, 579-590.
- Good, I. J. and Gaskins, R. A. (1980). Density estimation and bump-hunting by the penalized maximum likelihood method exemplified by scattering and meteorite data (with discussion). *J. Amer. Statist. Assoc.* 75, 42-73.
- Hartigan, J. A. (1988). The span test of multimodality. In *Classification and Related Methods of Data Analysis* (H. H. Bock, ed.) 229{236. North-Holland, Amsterdam.
- Hartigan, J. A. and Hartigan, P. M. (1985). The dip test of unimodality. *Ann. Statist.* 13, 70-84.
- Hartigan, J. A. and Mohanty, S. (1992). The RUN test for multimodality. *J. Classification* 9, 63-70.
- Hartigan (1985), Algorithm AS217: Computation of the Hartigan Dip Test statistic to test for unimodality. *Appl. Statist.* 34, pp. 320-325.



- Hall and York (2001), "On the Calibration of Silverman's Test for Multimodality," *Statistica Sinica*, 11, 515-536.
- Izenman, A. J. and Sommer, C. (1988). Philatelic mixtures and multimodal densities. *J. Amer. Statist. Assoc.* 83, 941-953.
- Mueller and Sawitzki (1991) Excess mass estimates and tests for multimodality. *JASA* 86, 738 -746
- Minnotte, M. C. (1992). A test of mode existence with applications to multimodality. Ph.D. thesis, Dept. Statistics, Rice Univ.
- Muller, D. W. and Sawitzki, G. (1991). Excess mass estimates and tests for multimodality. *J. Amer. Statist. Assoc.* 86, 738-746.
- Rozal, G. P. M. and Hartigan, J. A. (1994). The MAP test for multimodality. *J. Classification* 1,15-36.
- Silverman, B. W. (1978). Weak and strong uniform consistency of the kernel estimate of a density and its derivatives. *Ann. Statist.* 6 177.
- Silverman, B. W. (1981). Using kernel density estimates to investigate multimodality. *J. Roy. Statist. Soc. Ser. B* 43, 97-99.
- Silverman (1983), 'Some properties of a test for multimodality based on kernel density estimates', in J. F. C. Kingman and G. E. H. Reuter (eds), *Probability, Statistics, and Analysis*, Cambridge University Press, Cambridge.
- Silverman (1986), *Density Estimation for Statistics and Data Analysis*, Monographs on Statistics and Applied Probability. 26, Chapman and Hall, London.

Declaration

Conflict of Study: The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Funding Statement: This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

Availability of Data and Material: Data will be made available by the corresponding author on reasonable request.

Author(s) Bio / Authors' Note

Rana Muhammad Imran Arshad is from the Department of Statistics, Govt. Sadiq Egerton College, Bahawalpur, Pakistan. Email: imranarshad.stat@gmail.com

Sadaf Khan is from the Department of Statistics, The Islamia University of Bahawalpur, Pakistan. Email: smkhan6022@gmail.com

Farrukh Jamal is from the Department of Statistics, Faculty of Science, University of Tabuk, Saudi Arabia. Email: fshakir@ut.edu.sa

Shahina Imran is from the Department of Economics, Govt. College for Women, Dubai Mahal Road, Bahawalpur, Pakistan. Email: shahinaimran3408@yahoo.com

