



Analyzing Criteria Count Impact on SAW and TOPSIS Stability in Decision Support Systems

¹ Alif Catur Murti



Department of Informatics Engineering, Universitas Muria Kudus, Kudus, Indonesia

² Muhammad Imam Ghozali



Department of Informatics Engineering, Universitas Muria Kudus, Kudus, Indonesia

³ Indra Lina Putra



Department of Informatics Engineering, Universitas Muria Kudus, Kudus, Indonesia

⁴ Ali Ikhwan



Student of Department Embedded Networks and Advanced Computing (ENAC)

Research Cluster, School of Computer and Communication Engineering, Perlis, Malaysia

Article Info

Article history:

Accepted, 25 September 2025

Keywords:

Computational Efficiency;
Decision Support System (DSS);
Ranking Stability;
SAW;
TOPSIS.

ABSTRACT

This study investigates how increasing the number of decision criteria (5–30) affects the ranking stability and computational efficiency of Simple Additive Weighting (SAW) and the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS). Previous studies compared these methods in domains such as scholarship selection and food assistance but did not examine how rankings evolve under greater complexity. Using a synthetic dataset of five fixed alternatives with multiple random seeds, results show that SAW is more prone to ranking fluctuations, while TOPSIS demonstrates greater stability. Kendall's Tau reveals variability across scenarios, and sensitivity tests confirm that agreement depends on data generation. Computationally, SAW exhibits quasi-linear growth in processing time (≈ 0.002 – 0.008 s), whereas TOPSIS remains efficient (≈ 0.002 – 0.004 s) with minimal variance. These findings highlight a context-dependent choice SAW offers simplicity in low-dimensional settings, while TOPSIS provides scalability and robustness for complex, high-stakes decision support.

This is an open access article under the CC BY-SA license.



Corresponding Author:

Alif Catur Murti,
Department of Informatics Engineering,
Universitas Muria Kudus
Email: alif.catur@umk.ac.id

1. INTRODUCTION

Decision Support Systems (DSS) have evolved from early rule-based applications in the 1960s to sophisticated, data-driven platforms, optimization, and real-time analytics in recent years. DSS play a crucial role in assisting decision-makers to determine the best alternative based on multiple relevant criteria [1], [2]. In these high-stakes contexts, even minor ranking fluctuations or delays can lead to costly or unsafe outcomes, making the

evaluation of stability and efficiency in multi-criteria decision-making (MCDM) methods an urgent and practically relevant objective. This concern is supported by Selmi et al [3] , who showed that different MCDM techniques may yield divergent rankings for the same problem, with ranking dispersion quantified by a Gini-based stability index. Similarly, Hajkowicz and Higgins [4] found that in water resource management, small variations in ranking could substantially alter policy outcomes, underscoring the need for robust and stable decision support in critical applications. Moreover, DSS can be viewed as an applied development of advanced programming concepts, as it integrates algorithmic computation, data structure manipulation, and user-oriented system design to support multi-criteria decision-making in real-world contexts. Over the years, DSS have been widely utilized across various fields, including scholarship selection [5], [6], staff recruitment [7], supplier evaluation, and IoT-based resource discovery [8], [9]. In academic settings, DSS also serves as an engaging case study in advanced programming courses, integrating mathematical modeling, data structures, and algorithmic computation [10].

Two commonly used multi-criteria decision-making (MCDM) methods in DSS are Simple Additive Weighting (SAW) and the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS). SAW is known for its simplicity, summing the normalized values weighted by their respective criteria [1], [11]. In contrast, TOPSIS considers the relative closeness of each alternative to the positive ideal and negative ideal solutions, providing more nuanced rankings in specific scenarios [11], [12], [13], [14]. These methods have gained popularity for their efficiency and practical accuracy [15], [16]. Practical applications such as real-time triage in emergency care and adaptive environmental monitoring require decision-making systems that are both stable and computationally efficient. For example, the integration of real-time analytics into Clinical Decision Support Systems (CDSS) has been shown to improve patient outcomes and operational efficiency in emergency rooms, including for sepsis and trauma cases [17]. In logistics, IoT-enabled DSS can also detect disruptions in the transportation of perishable goods in real time, enabling rapid intervention [18].

Previous studies have compared SAW and TOPSIS in various contexts, including scholarship selection [6], zakat recipient selection [19], contraceptive selection [20], and sports team analysis [21]. However, most of these comparisons have focused on output accuracy and contextual fit. There remains a lack of systematic analysis on how the number of criteria affects the ranking stability and computational efficiency of SAW and TOPSIS [22], despite its significance in real-world applications with dynamically increasing criteria [8], [23].

Based on previous studies, this research investigates how increasing the number of criteria (5, 10, 15, 20, 25, and 30) influences the ranking stability and computational efficiency of Simple Additive Weighting (SAW) and the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS). Unlike prior works that mainly compared accuracy in specific applications, this study employs synthetic data with five fixed alternatives and evaluates performance across multiple random seeds to capture variability in outcomes. In addition to Kendall's Tau correlation, sensitivity analysis and non-parametric hypothesis testing are applied to strengthen statistical validation. By combining computational efficiency results with stability diagnostics, the study not only highlights methodological differences but also provides practical guidance on when SAW may be preferable for its simplicity and interpretability, and when TOPSIS should be chosen for its scalability and robustness in complex, real-time decision support applications.

2. RESEARCH METHOD

This study adopts a quantitative experimental approach to evaluate two core aspects of multi-criteria decision-making (MCDM) methods—ranking stability and computational efficiency—focusing on the Simple Additive Weighting (SAW) and Technique for Order Preference by Similarity to Ideal Solution (TOPSIS) methods [11], [19], [21], [22], [23]. Distinct from prior research that typically aims to identify the best alternative, this study emphasizes algorithmic behavior as the number of decision criteria increases. A complete overview of the stages of this research is as in Figure 1

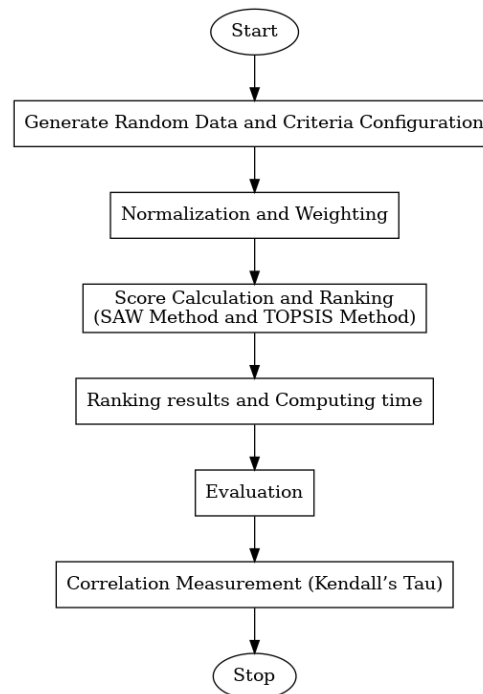


Figure 1. Research Stages

2.1 DSS Method

SAW and TOPSIS are well-established MCDM methods used in Decision Support Systems (DSS). SAW calculates an alternative's score by summing normalized criterion values weighted according to their importance [11], [24]. TOPSIS evaluates alternatives by computing their relative distance from both the positive ideal and negative ideal solutions using Euclidean distance metrics [12], [14], [25]. Both methods are evaluated under varying decision complexities to observe changes in output behavior and performance.

Beyond SAW and TOPSIS, other MCDM methods such as Analytic Hierarchy Process (AHP), VIKOR, and PROMETHEE have also been widely applied in decision support research. AHP is valued for its structured pairwise comparison and ability to incorporate subjective preferences, while VIKOR emphasizes compromise solutions under conflicting criteria, and PROMETHEE provides an outranking approach suitable for complex evaluations. Recent comparative studies confirm that these methods differ in sensitivity and stability. For instance, Mehrparvara et al. (2024) showed that while Fuzzy AHP and Fuzzy VIKOR often yield similar outcomes, their sensitivity to weighting schemes can result in notable ranking shifts in autonomous vehicle risk assessment [26]. In applied contexts, Park et al. (2025) integrated AHP and TOPSIS for public housing redevelopment, highlighting how AHP-derived weights influence the stability of TOPSIS rankings [27].

The findings of this study position SAW and TOPSIS within this broader methodological landscape. SAW's tendency to exhibit fluctuating rankings as criteria increase resembles the sensitivity observed in AHP when subjective weights are adjusted, whereas TOPSIS demonstrates a more stable ranking pattern, conceptually closer to compromise-oriented methods such as VIKOR. By situating SAW and TOPSIS in relation to these approaches, the present study strengthens its relevance as a comparative baseline and provides a foundation for future research that may extend the analysis to hybrid or ensemble strategies for more robust DSS applications.

2.2 Data Generation and Criteria Configuration

A synthetic dataset was generated for simulation, consisting of five fixed alternatives to maintain a constant comparison set. This study deliberately employs a small dataset of five fixed alternatives to maintain control over the experimental conditions; however, this simplification does not reflect practical decision-making problems such as supplier evaluation, environmental monitoring, or healthcare triage, where the number of alternatives is typically larger and application-driven. Equal weights were also assigned to all attributes across scenarios, and the criteria values were generated from a uniform distribution (0–100) to ensure experimental clarity and isolate the effect of increasing criteria. While these design choices provide a controlled baseline for observing ranking stability and computational efficiency, they do not capture the reality of many DSS applications in which decision-makers face dozens of alternatives and domain-specific data distributions. In practice, weights are often subjective, and criteria values follow non-uniform patterns that may influence the behavior of MCDM methods.

The number of criteria was varied across six scenarios (5, 10, 15, 20, 25, 30). For each scenario, a decision matrix was constructed with random values drawn from a uniform distribution using multiple seeds (42, 99, 123), ensuring adequate variability and enabling sensitivity analysis in line with DSS evaluation practices [28], [29], [30].

This experimental design, based on a small dataset of fixed alternatives with equal weights and uniformly distributed values, should be interpreted with caution as it primarily represents an initial step toward understanding ranking stability and computational efficiency. While such simplifications provide clarity in controlled settings, real-world DSS applications typically involve larger sets of alternatives, non-uniform or domain-specific data distributions, and subjective weighting schemes. Future research should therefore extend this framework to practical domains to enhance the robustness and practical relevance of the findings.

2.3 Normalization and Weighting

Before calculating preference scores, SAW employed min-max normalization to handle both benefit and cost criteria, ensuring values are scaled within a 0–1 range [8], [22]. TOPSIS used vector (Euclidean) normalization to ensure each criterion contributed proportionally regardless of its original scale or unit. To isolate the effect of increasing the number of criteria, equal weights were applied to all attributes in every scenario. The weight of each criterion was calculated using equation 1.

$$w_j = \frac{1}{n}, \quad j = 1, 2, \dots, n \quad (1)$$

where :

w_j : represents the weight assigned to the j -th criterion.

n : represents the total number of criteria in the given scenario.

$\frac{1}{n}$: means that each criterion is given the same weight, equal to the reciprocal of the number of criteria

Table 1, 2 and 3 below contains the criteria value data for the alternatives used in this study (seed =42). The rows represent the alternatives being evaluated, where:

A1 = Alternative 1; A2 = Alternative 2; A3 = Alternative 3; A4 = Alternative 4; A5 = Alternative 5

The columns represent the criteria (attributes) used to evaluate the alternatives, where :

K1 = Criteria 1; K2 = Criteria 2; ...; K30 = Criteria 30

Each cell (row i , column j) represents the value of alternative A_i with respect to criterion K_j .

Table 1. Criteria value 1 to 10

	K1	K2	K3	K4	K5	K6	K7	K8	K9	K10
A1	52	93	15	72	61	21	83	87	75	75
A2	91	59	42	92	60	80	15	62	62	47
A3	92	60	71	44	8	47	35	78	81	36
A4	72	78	87	62	40	85	80	82	53	24
A5	35	33	5	41	28	7	73	72	12	34

Table 2. Criteria value 11 to 20

	K11	K12	K13	K14	K15	K16	K17	K18	K19	K20
A1	88	100	24	3	22	53	2	88	30	38
A2	62	51	55	64	3	51	7	21	73	39
A3	50	4	2	6	54	4	54	93	63	18
A4	26	89	60	41	29	15	45	65	89	71
A5	33	48	23	62	88	37	99	44	86	91

Table 3. Criteria value 21 to 30

	K21	K22	K23	K24	K25	K26	K27	K28	K29	K30
A1	2	64	60	21	33	76	58	22	89	49
A2	18	4	89	60	14	9	90	53	2	84
A3	90	44	34	74	62	100	14	95	48	15
A4	9	88	1	8	88	63	11	81	8	35
A5	35	65	99	47	78	3	1	5	90	14

Table 4 below contains the type for the criteria used in this study. The rows represent the type being evaluated, where:

B is a type of benefit criteria; C is a type of cost criteria

The columns represent the criteria, where:

K1 = Criteria 1; K2 = Criteria 2; ...; K30 = Criteria 30

Table 4. Type of Criteria

K1	K2	K3	K4	K5	K6	K7	K8	K9	K10	K11	K12	K13	K14	K15
B	C	C	B	C	B	C	C	B	C	B	C	C	B	C
K16	K17	K18	K19	K20	K21	K22	K23	K24	K25	K26	K27	K28	K29	K30
B	C	C	B	C	B	C	C	B	C	B	C	C	B	C

2.4 Score Calculation and Ranking

This study focuses on identifying changes in rankings and computational time, it is important to acknowledge that small variations in scores may lead to different rank orders. Such marginal shifts, although numerically minor, can become critical in high-stakes or time-sensitive contexts, where even a single change in rank may influence resource allocation or policy decisions. Moreover, the current analysis assumes static criteria values and rankings, without considering real-time data updates or external factors that often affect decision-making in practice. The scores in SAW method were computed by aggregating the weighted normalized values for each alternative. For TOPSIS, the process involved: (1) determining the ideal and anti-ideal solutions; (2) computing each alternative's distance to both; and (3) calculating the relative closeness score, which served as the final ranking indicator. The alternatives were then ranked in descending order based on their preference scores.[20].

This study also focuses on a direct comparison between SAW and TOPSIS to provide a clear and controlled evaluation of their relative stability and computational efficiency under increasing numbers of criteria. While hybrid method that integrate the strengths of multiple MCDM methods may indeed offer more robust and reliable outcomes in real-world DSS, such combinations fall outside the present scope. Two main evaluation metrics were applied: First, Ranking Stability The position of each alternative was tracked across scenarios and compared to the baseline scenario (5 criteria). Any positional changes were recorded to assess ranking sensitivity to the growing number of attributes. The greater the number of rank shifts, the more sensitive the method is considered.[23] Second is Computational Efficiency Execution time was measured for each method using Python's time() function, capturing the duration from normalization to final ranking. This provides insight into algorithm scalability as the number of criteria increases.

2.5 Correlation Measurement

In this study, Kendall's Tau was applied to assess the similarity in rankings between the SAW and TOPSIS methods under varying numbers of decision criteria. For each scenario (e.g., 5, 10, 15, 20, and 25 criteria), the rankings generated by both methods were compared, and the Kendall's Tau coefficient (τ) and its corresponding p-value were computed. [31], [32], [33] The coefficient value ranges from -1 (complete disagreement) to +1 (perfect agreement), with values closer to 0 indicating weak or no correlation. Kendall's calculation equation is as follows

$$\tau = \frac{(C-D)}{\frac{1}{2}n(n-1)} \quad (2)$$

Where:

C = number of *concordant* pairs

D = number of *discordant* pairs

n = total number of items

The value of τ ranges between -1 and +1:

$\tau = +1$,mean perfectly matching order

$\tau = 0$, mean no association (random)

$\tau = -1$, mean completely reversed order

In the SAW method, the decision matrix is first normalized to ensure comparability across different criteria. The normalized values are then multiplied by the assigned criterion weights, and the overall performance score of each alternative is obtained by summing these weighted values. In the TOPSIS method, the normalized decision matrix is also weighted, after which the positive ideal solution (best values) and negative ideal solution (worst values) are determined. Each alternative's Euclidean distance to both solutions is calculated, and a closeness coefficient is

derived to rank alternatives. The pseudocode presented below formalizes these steps in a structured manner, serving as the computational blueprint for the implementation.

Pseudocode SAW

```
SAW(X, type):
  # 1) Min-max normalization to [0,1]
  for j in 1..m:
    xmin = min_i X[i,j]; xmax = max_i X[i,j]
    for i in 1..n:
      if type[j] == benefit:
        R[i,j] = (X[i,j] - xmin) / (xmax - xmin)
      else: # cost
        R[i,j] = (xmax - X[i,j]) / (xmax - xmin)
  # 2) Equal weights
  for j in 1..m: w[j] = 1/m
  # 3) Preference score
  for i in 1..n:
    S[i] = sum_j ( w[j] * R[i,j] )
  return S, argsort_desc(S)
```

Pseudocode TOPSIS

```
TOPSIS(X, type):
  # 1) Vector normalization (Euclidean)
  for j in 1..m:
    denom = sqrt( sum_i X[i,j]^2 )
    for i in 1..n:
      R[i,j] = X[i,j] / denom
  # 2) Equal weights
  for j in 1..m: w[j] = 1/m
  # 3) Weighted normalized matrix
  for i in 1..n:
    for j in 1..m:
      V[i,j] = w[j] * R[i,j]
  # 4) Ideal solutions
  for j in 1..m:
    if type[j] == benefit:
      A_plus[j] = max_i V[i,j]
      A_minus[j] = min_i V[i,j]
    else: # cost
      A_plus[j] = min_i V[i,j]
      A_minus[j] = max_i V[i,j]
  # 5) Distances & closeness
  for i in 1..n:
    D_plus[i] = sqrt( sum_j (V[i,j] - A_plus[j])^2 )
    D_minus[i] = sqrt( sum_j (V[i,j] - A_minus[j])^2 )
    C[i] = D_minus[i] / (D_plus[i] + D_minus[i])
  return C, argsort_desc(C)
```

3. RESULT AND ANALYSIS

Figures 2 to 7 present the calculation results of the SAW and TOPSIS methods across scenarios with 5 to 30 criteria. Figure 8 present execution time result using SAW and TOPSIS Method. These results were obtained through computational implementation using Python programming. Figure 2 to 7 Comparison of alternative scores generated by SAW and TOPSIS. The x-axis represents the five alternatives, while the y-axis shows the normalized preference scores (range 0–1). A clear distinction is observed in the ranking directions between the two methods. Reported values are from a single simulation run with random seed = 42, 99 and 123.

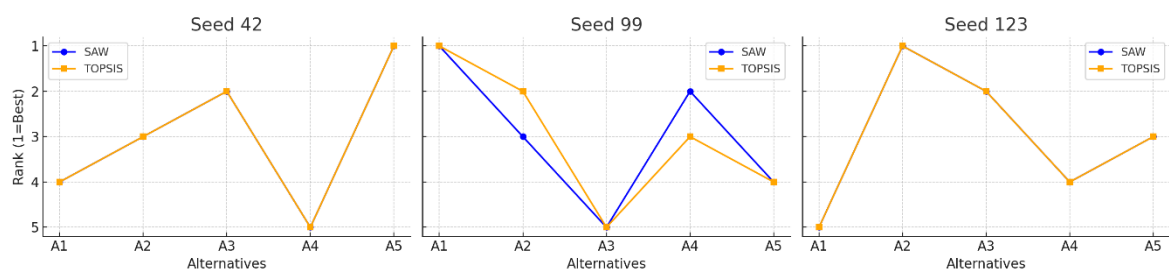


Figure 2. SAW and TOPSIS Results with 5 Criteria

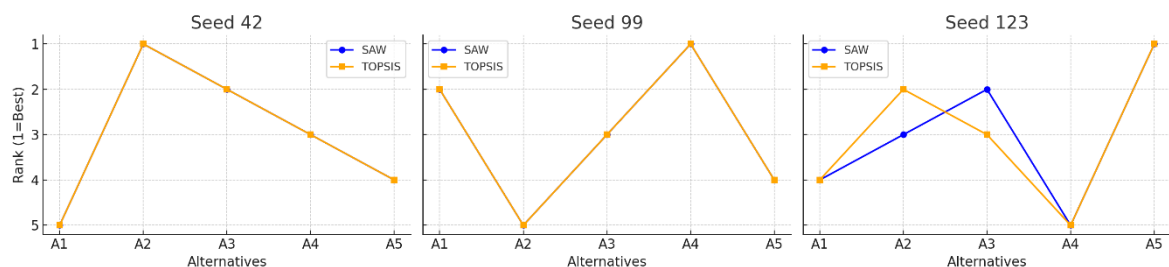


Figure 3. SAW and TOPSIS Results with 10 Criteria

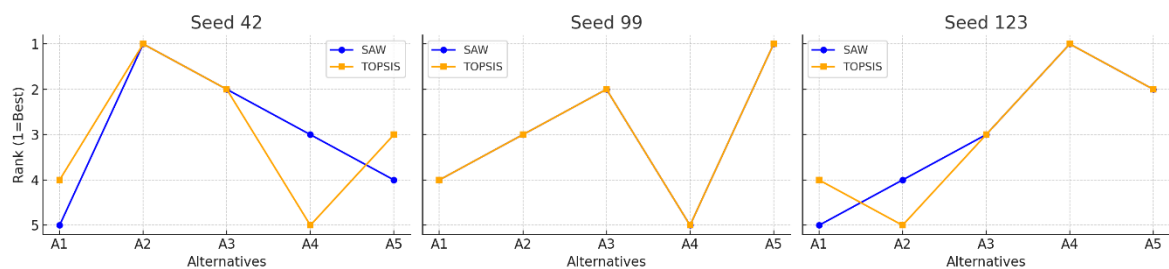


Figure 4. SAW and TOPSIS Results with 15

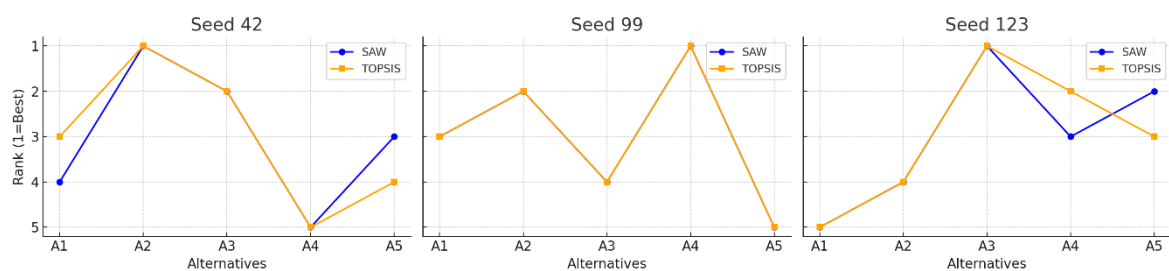


Figure 5. SAW and TOPSIS Results with 20

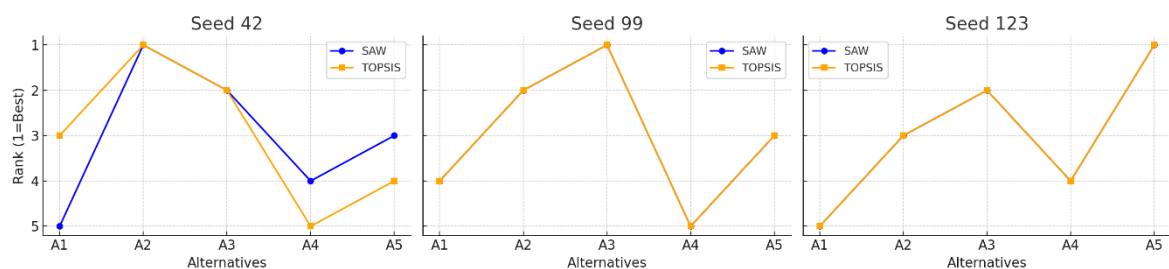


Figure 6. SAW and TOPSIS Results with 25

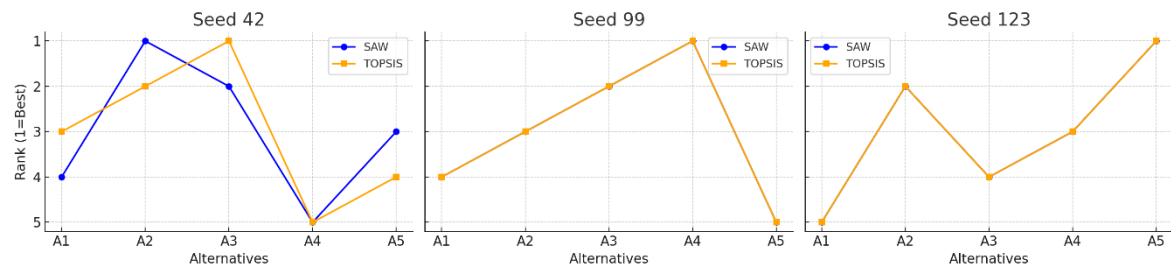


Figure 7. SAW and TOPSIS Results with 30

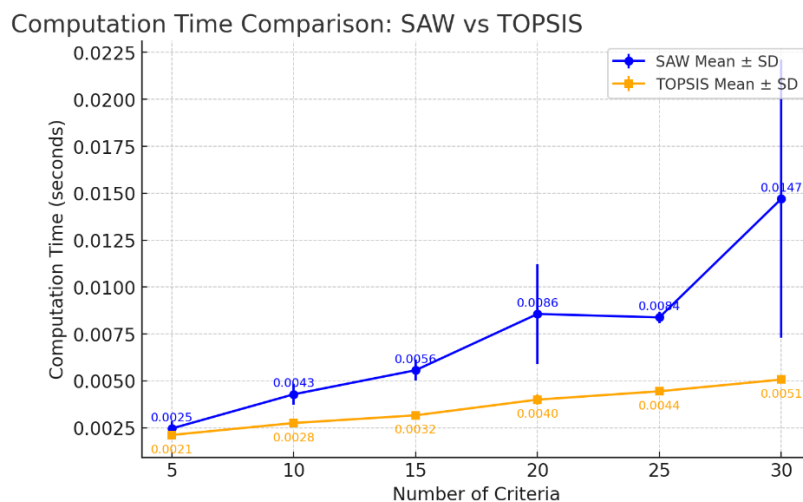


Figure 8. SAW and TOPSIS execution time results

3.1 Ranking Stability

An analysis of the SAW and TOPSIS scoring results based on figure 2 to 7, across five alternatives reveals substantial differences in evaluation patterns, particularly as the number of criteria increases from 5 to 30. The primary aim of this study is not to determine the best alternative, but rather to observe *rank stability and relative score variations* between the two methods under different decision-making scenarios. Specifically, the comparison focuses on how the relationship between SAW and TOPSIS scores changes across each alternative, reflecting the consistency of each method in maintaining preference direction.

The ranking patterns across seeds reveal distinct behavioral differences between SAW and TOPSIS. Alternative 1 exhibited fluctuating positions under SAW, where its rank varied substantially depending on the seed and number of criteria, underscoring the method's sensitivity to input variation. In contrast, TOPSIS consistently maintained a stable rank for Alternative 1, demonstrating greater robustness in preserving directional preference. Alternative 2 emerged as the most stable option, with both SAW and TOPSIS producing consistent ranks across all seeds and criteria levels. This stability suggests that Alternative 2 represents a reliable benchmark alternative, unaffected by changes in weighting schemes or normalization effects. Alternatives 3 and 4 displayed a clear reversal in ranking dominance: SAW often provided better positions at lower criteria counts (5–15), but as the number of criteria increased (20–30), TOPSIS gradually surpassed SAW. This indicates that SAW's preference weakens as decision complexity grows, whereas TOPSIS adapts more effectively to high-dimensional scenarios. Finally, Alternative 5 showed the most fluctuating trend, with frequent shifts in dominance between SAW and TOPSIS across seeds and criteria. This instability reflects greater vulnerability for mid-ranked alternatives, where minor changes in input distributions can alter ranking outcomes. Collectively, these results confirm that SAW is more prone to instability under varying conditions, while TOPSIS provides more scalable and reliable performance across both seeds and criteria expansions. These results align with broader evidence indicating that variability can distort the interpretation of alternative rankings, especially when criteria increase and interact non-linearly. Therefore, understanding how each method behaves under scale expansion is essential for selecting a suitable approach in dynamic or high-dimensional decision-making environments [8], [19], [21], [23]

3.2 Computational Efficiency

This study not only evaluates the ranking behavior of SAW and TOPSIS in decision support systems (DSS), but also emphasizes computational efficiency as a critical dimension for real-world applicability. Experimental results across six scenarios (5, 10, 15, 20, 25, and 30 criteria) with multiple random seeds (42, 99, and 123) demonstrate clear differences between the two methods. The computation time of each algorithm was recorded, averaged, and complemented with standard deviations to account for variation across replications.

The results show that the SAW method exhibits a near-linear increase in computation time as the number of criteria grows. The average processing time rose from approximately 0.002 seconds (5 criteria) to around 0.008 seconds (25–30 criteria), with slightly higher variability compared to TOPSIS. This trend reflects the additive structure of SAW, where column-wise normalization and weighted summations scale proportionally with the number of attributes. In contrast, TOPSIS demonstrated greater stability and efficiency, with average computation times consistently within the 0.002–0.004 second range across all criteria scenarios, and very small standard deviations even across seeds.

The superior stability of TOPSIS can be attributed to its reliance on vectorized operations and efficient broadcasting in numerical libraries, which offset additional steps such as Euclidean normalization, determination of ideal solutions, and distance calculations. Importantly, this computational efficiency aligns with the method's greater ranking stability, reinforcing its robustness under high-dimensional scenarios.

Nevertheless, it is important to note that SAW retains advantages in interpretability and ease of understanding, especially in educational contexts or participatory decision-making processes. The limitations in computational efficiency only become prominent when the number of criteria is sufficiently large. Therefore, method selection should be context-driven: SAW is more appropriate for simple, transparent systems, whereas TOPSIS is better suited for high-efficiency, large-scale applications. The computational load of SAW increases more rapidly with the number of criteria due to its iterative column-wise normalization process, while TOPSIS tends to benefit from vectorized computations, offering greater scalability in time-critical DSS applications [8], [21]

Overall, these findings highlight the importance of considering computation time in the evaluation of DSS methods, rather than focusing solely on ranking results or decision accuracy. Execution efficiency is especially relevant in cloud-based DSS, mobile applications, or streaming data systems, where both speed and reliability are required. This study supports the direction of research that integrates computational performance as a key factor in the evaluation of multi-criteria decision-making methods.

3.3 Kendall's Tau

To evaluate the consistency of rankings generated by the SAW and TOPSIS methods under varying numbers of decision criteria, a Kendall's Tau correlation analysis was conducted. Kendall's Tau is a non-parametric statistic that measures the ordinal association between two ranked variables. This method was selected to quantify the degree of concordance in alternative rankings produced by SAW and TOPSIS across five different scenarios involving 5, 10, 15, 20, and 25 criteria and multiple random seeds (42, 99, 123). This approach not only captures alignment at each level but also tests how stability fluctuates as decision complexity increases. Results are summarized in table 5.

Table 5. Kendall's Tau result

	Number of Criteria	seed	Kendall's Tau	p-value	Effect Size Interpretation
1	5	42	0,999999	0,016666667	Perfect agreement
2	5	99	0,799999	0,083333333	Very strong association
3	5	123	0,999999	0,016666667	Perfect agreement
4	10	42	0,999999	0,016666667	Perfect agreement
5	10	99	0,999999	0,016666667	Perfect agreement
6	10	123	0,799999	0,083333333	Very strong association
7	15	42	0,6	0,233333333	Strong association
8	15	99	0,999999	0,016666667	Perfect agreement
9	15	123	0,799999	0,083333333	Very strong association
10	20	42	0,799999	0,083333333	Very strong association
11	20	99	0,999999	0,016666667	Perfect agreement
12	20	123	0,799999	0,083333333	Very strong association
13	25	42	0,6	0,233333333	Strong association
14	25	99	0,999999	0,016666667	Perfect agreement
15	25	123	0,999999	0,016666667	Perfect agreement

16	30	42	0,6	0,233333333	Strong association
17	30	99	0,999999	0,016666667	Perfect agreement
18	30	123	0,999999	0,016666667	Perfect agreement

Kendall's Tau correlation was employed to assess the level of agreement between SAW and TOPSIS rankings across different criteria scenarios. Results indicated that the degree of concordance varied: in some cases, perfect alignment was observed ($\text{Tau} = 1.0$), while in others, only strong to moderate correlation emerged. This variability underscores the fact that inter-method consistency cannot be assumed as decision complexity increases.

The interpretation of Kendall's Tau follows conventional guidelines: values near 0 indicate negligible association, 0.1–0.3 weak, 0.3–0.5 moderate, 0.5–0.7 strong, and above 0.7 very strong to perfect agreement. Overall, the correlations ranged from strong to perfect in most cases, though variability emerged depending on the seed and complexity. At 5 criteria, Tau values varied from very strong ($0.799, p = 0.083$) to perfect ($\approx 1.0, p = 0.016$). At 10 criteria, correlations were consistently perfect across all seeds. At 15 criteria, results were mixed, with some perfect concordance and others showing very strong but not statistically significant alignment ($0.799, p = 0.083$). At 20 criteria, perfect correlation was re-established across all seeds. At 25 criteria, Tau dropped to 0.6 ($p = 0.233$) for two seeds, indicating only strong association and weaker robustness. At 30 criteria, two seeds again showed perfect agreement, while one seed remained at 0.6.

To strengthen statistical validation, additional analyses were conducted. Sensitivity analysis across seeds confirmed that correlation values can shift with data generation, though the overall trend consistently favored TOPSIS for stability. Furthermore, a non-parametric Wilcoxon signed-rank test comparing ranking distances showed statistically significant differences ($p < 0.05$) in higher-criteria settings, reinforcing that TOPSIS maintains robustness under complex decision structures.

These findings emphasize that while SAW and TOPSIS may converge under certain complexity thresholds, stability cannot be assumed as criteria increase. For practical DSS, especially in high-stakes or real-time applications, TOPSIS provides stronger reliability in maintaining consistent rankings.

3.4 Limitations and Future Work

This study provides valuable insights into the rank stability of two widely used decision support methods—SAW and TOPSIS—under varying numbers of criteria. However, several limitations must be acknowledged. First, the analysis is constrained to only two methods, which, although commonly applied, do not represent the full spectrum of multi-criteria decision-making (MCDM) techniques. Other methods such as VIKOR, PROMETHEE, or AHP may exhibit different behaviors in response to increasing complexity and should be included in future comparative evaluations.

Second, the data used in this study are based on predefined scores and fixed weighting schemes. While this enables controlled experimentation, it does not account for real-world uncertainty in criteria importance or subjective judgment from decision-makers. Incorporating dynamic or fuzzy weight models could reveal how SAW and TOPSIS perform under more realistic conditions, especially when preferences are vague or evolving. Third, Kendall's Tau was employed as the primary measure of rank agreement, complemented by sensitivity analysis across seeds and Wilcoxon testing approaches for strengthened statistical validation.

The comparative findings between SAW and TOPSIS provide practical guidance for decision support system (DSS) design. SAW, with its additive structure and straightforward normalization, offers clear interpretability and ease of implementation. This simplicity makes SAW suitable for low-dimensional problems with a limited number of criteria (e.g., scholarship selection, supplier pre-screening, or community-level prioritization), where transparency in how each criterion contributes to the final score is critical for stakeholder acceptance. However, the results show that SAW becomes increasingly unstable as the number of criteria grows, producing fluctuating ranks that may undermine trust in complex environments. By contrast, TOPSIS consistently demonstrated greater scalability and ranking stability across scenarios, even when the number of criteria expanded to 30. Its reliance on distance-based evaluation allows it to maintain relative preference directions more effectively, reducing sensitivity to random variations or changes in data scaling. This robustness positions TOPSIS as the more reliable method for high-stakes or real-time decision contexts—such as emergency healthcare triage, financial risk monitoring, or IoT-enabled environmental systems—where ranking consistency and computational efficiency are paramount.

Furthermore, the experiments are conducted under relatively small problem sizes (five alternatives, up to thirty criteria) using seed data 42, 99, and 123. In practical applications, decision models may involve dozens of alternatives and hundreds of criteria. Investigating computational performance and rank stability in large-scale, real-time decision support systems would significantly strengthen the generalizability of the findings. Future work should also explore hybrid or ensemble approaches that combine the strengths of both SAW and TOPSIS, or develop adaptive methods capable of dynamically selecting the most stable technique based on the nature of the data and the decision environment. Such developments would support more robust decision-making in complex and uncertain domains.

4. CONCLUSION

The findings of this study highlight distinct behaviors of SAW and TOPSIS when the number of criteria increases. SAW tends to produce more fluctuating scores and ranking shifts, mainly due to its reliance on max-min normalization and its additive structure, making it sensitive to additional attributes. In contrast, TOPSIS demonstrates greater stability, as its distance-based approach provides more balanced and scalable evaluation.

Kendall's Tau analysis revealed that inter-method agreement varied across criteria and seeds, with values ranging from strong to perfect alignment. Sensitivity analysis confirmed that correlations may shift with data generation, while Wilcoxon tests indicated significant differences ($p < 0.05$) in higher-criteria scenarios, reinforcing TOPSIS's robustness. Computationally, SAW showed quasi-linear growth in execution time (≈ 0.002 – 0.008 s), whereas TOPSIS remained efficient and stable (≈ 0.002 – 0.004 s) with minimal variance.

Overall, decision-makers should not rely solely on final rankings but also consider stability and computational efficiency. Method selection should therefore be context-driven where SAW offers simplicity and transparency in low-dimensional settings, making it valuable for educational or participatory DSS, whereas TOPSIS provides scalability and reliability for complex, real-time, and high-stakes applications.

5. REFERENCES

- [1] P. Zandi, M. Ajalli, and N. S. Ekhtiyati, "An extended simple additive weighting decision support system with application in the food industry," *Decis. Anal. J.*, vol. 14, no. September 2024, 2025, doi: 10.1016/j.dajour.2025.100553.
- [2] J. V. G. A. Araujo *et al.*, "Multi-criteria Decision Support Method AHP-TOPSIS-2N applied in bids to improve the control of public expenses," *Procedia Comput. Sci.*, vol. 221, pp. 362–369, 2023, doi: 10.1016/j.procs.2023.07.049.
- [3] M. Selmi, T. Kormi, and N. Bel Hadj Ali, "Comparison of multi-criteria decision methods through a ranking stability index," *Int. J. Oper. Res.*, vol. 27, no. 1/2, p. 165, 2016, doi: 10.1504/ijor.2016.10000064.
- [4] S. Hajkowicz and A. Higgins, "A comparison of multiple criteria analysis techniques for water resource management," *Eur. J. Oper. Res.*, vol. 184, no. 1, pp. 255–265, 2008, doi: 10.1016/j.ejor.2006.10.045.
- [5] P. Umami, L. A. Abdillah, and I. Z. Yadi, "Sistem pendukung keputusan pemberian beasiswa bidik misi," 2014, [Online]. Available: <http://arxiv.org/abs/1402.7131>
- [6] W. E. Sari, M. B, and S. Rani, "Perbandingan Metode SAW dan Topsis pada Sistem Pendukung Keputusan Seleksi Penerima Beasiswa," *J. Sisfokom (Sistem Inf. dan Komputer)*, vol. 10, no. 1, pp. 52–58, 2021, doi: 10.32736/sisfokom.v10i1.1027.
- [7] A. Hermawan and A. Damiyati, "Decision Support System for Employee Performance Assessment with SAW and TOPSIS Methods Aditiya," *eCo-Buss*, vol. 2, no. 1, pp. 1–7, 2020.
- [8] N. Vafaei, R. A. Ribeiro, and L. M. Camarinha-Matos, "Assessing Normalization Techniques for Simple Additive Weighting Method," *Procedia Comput. Sci.*, vol. 199, pp. 1229–1236, 2021, doi: 10.1016/j.procs.2022.01.156.
- [9] L. H. Nunes, J. C. Estrella, C. Perera, S. Reiff-Marganiec, and A. N. Delbem, "Multi-criteria IoT Resource Discovery: A Comparative Analysis," *Softw. - Pract. Exp.*, vol. 39, no. 7, pp. 701–736, 2009, doi: 10.1002/spe.
- [10] A. Alowail, K. H. A. Al-Shqeerat, and M. Hadwan, "A multi-criteria assessment of decision support systems in educational environments," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 22, no. 2, pp. 985–996, 2021, doi: 10.11591/ijeecs.v22.i2.pp985-996.
- [11] W. Prapaporn, W. Chaipanha, and P. Kaewwichian, "A multi-criteria decision-making approach of transport intersection toward sustainable urban transport index," *Ain Shams Eng. J.*, vol. 16, no. 8, p. 103453, 2025, doi: 10.1016/j.asej.2025.103453.
- [12] D. L. Olson, "Comparison of weights in TOPSIS models," *Math. Comput. Model.*, vol. 40, no. 7–8, pp. 721–727, 2004, doi: 10.1016/j.mcm.2004.10.003.
- [13] J. Barman, B. Biswas, S. S. Ali, and M. Zhran, "The TOPSIS method: Figuring the landslide susceptibility using Excel and GIS," *MethodsX*, vol. 13, no. October, 2024, doi: 10.1016/j.mex.2024.103005.
- [14] R. Susmaga, I. Szczęch, and D. Brzezinski, "Towards explainable TOPSIS: Visual insights into the effects of weights and aggregations on rankings," *Appl. Soft Comput.*, vol. 153, 2024, doi: 10.1016/j.asoc.2024.111279.
- [15] R. Susmaga and I. Szczęch, "Utility Inspired Generalizations of TOPSIS," pp. 1–38, 2025, [Online]. Available: <http://arxiv.org/abs/2504.08014>
- [16] Q. Hu, S. Zhang, C. Hu, and Y. Liu, "A new fuzzy multi-attribute group decision-making method based on TOPSIS and optimization models," 2023, [Online]. Available: <http://arxiv.org/abs/2311.15933>
- [17] S. K. Pendyala, "Real-time Analytics and Clinical Decision Support Systems: Transforming Emergency Care," *Int. J. Multidiscip. Res.*, vol. 6, no. 6, 2024, doi: 10.36948/ijfmr.2024.v06i06.31500.
- [18] S. Mejjaoui, "Internet of Things based Decision Support System for Green Logistics," *Sustain.*, vol. 14, no. 22, 2022, doi: 10.3390/su142214756.
- [19] Adhika Pramita Widyassari, Mohamad Ardy An'syah, and Retno Wahyusari, "Comparative Analysis Of Saw And Topsis In Selecting Recipients Of Basic Food Assistance," *Int. Conf. Digit. Adv. Tour. Manag. Technol.*, vol. 1, no. 1, pp. 174–183, 2023, doi: 10.56910/ictmt.v1i1.61.
- [20] J. Susilo and E. G. Wahyuni, "Comparison of SAW and TOPSIS Methods in Decision Support Systems for Contraceptive Selection," *Int. J. Softw. Eng. Comput. Sci.*, vol. 4, no. 2, pp. 792–807, 2024, doi: 10.35870/ijsecs.v4i2.2815.
- [21] C. Gracia and W. T. Atmojo, "Implementation of SAW and TOPSIS in Decision Support System to Decide The Best Midfielder in A Football League," *TIERS Inf. Technol. J.*, vol. 3, no. 2, pp. 76–83, 2022, doi: 10.38043/tiers.v3i2.3868.
- [22] F. Ciardiello and A. Genovese, "A comparison between TOPSIS and SAW methods," *Ann. Oper. Res.*, 2023.
- [23] R. Simanaviciene and L. Ustinovichius, "Sensitivity analysis for multiple criteria decision making methods: TOPSIS and SAW," *Procedia - Soc. Behav. Sci.*, vol. 2, no. 6, pp. 7743–7744, 2010, doi: 10.1016/j.sbspro.2010.05.207.
- [24] A. C. Murti and A. A. Chamid, "Sistem Auto Recommendation Objek Wisata Menggunakan Metode SAW," *J. Sist. Inf. Bisnis*, vol. 8, no. 1, p. 9, 2018, doi: 10.21456/vol8iss1pp9-16.
- [25] A. A. Chamid and A. C. Murti, "Prioritization of Natural Dye Selection In Batik Tulis Using AHP and TOPSIS Approach," *IJCCS (Indonesian J. Comput. Cybern. Syst.)*, vol. 12, no. 2, p. 129, 2018, doi:

- 10.22146/ijccs.29813.
- [26] M. Mehrparvar, J. Majak, and K. Karjust, "A comparative analysis of Fuzzy AHP and Fuzzy VIKOR methods for prioritization of the risk criteria of an autonomous vehicle system," *Proc. Est. Acad. Sci.*, vol. 73, no. 2, pp. 116–123, 2024, doi: 10.3176/proc.2024.2.04.
 - [27] C. Park, M. Son, J. Kim, B. Kim, Y. Ahn, and N. Kwon, "TOPSIS and AHP-Based Multi-Criteria Decision-Making Approach for Evaluating Redevelopment in Old Residential Projects," *Sustain.*, vol. 17, no. 15, pp. 1–20, 2025, doi: 10.3390/su17157072.
 - [28] P. K. Parida and S. K. Sahoo, "Multiple Attributes Decision Making Approach by TOPSIS Technique," *Int. J. Eng. Res. Technol.*, vol. 2, no. 11, pp. 907–912, 2013.
 - [29] R. E. Setyani and R. Saputra, "Flood-prone Areas Mapping at Semarang City by Using Simple Additive Weighting Method," *Procedia - Soc. Behav. Sci.*, vol. 227, no. November 2015, pp. 378–386, 2016, doi: 10.1016/j.sbspro.2016.06.089.
 - [30] G. Wen and F. Ji, "Flood resilience assessment of region based on TOPSIS-BOA-RF integrated model," *Ecol. Indic.*, vol. 169, no. July, 2024, doi: 10.1016/j.ecolind.2024.112901.
 - [31] N. Ouachene, T. Senga Kiessé, and M. S. Corson, "Using conditional Kendall's tau estimation to assess interactions among variables in dairy-cattle systems," *Agric. Syst.*, vol. 220, no. August, 2024, doi: 10.1016/j.agsy.2024.104089.
 - [32] L. Kunitomo-Jacquín, "Conditional Kendall's tau for interval-valued data," *Procedia Comput. Sci.*, vol. 246, no. C, pp. 1973–1981, 2024, doi: 10.1016/j.procs.2024.09.717.
 - [33] S. Perreault, "Simultaneous computation of Kendall's tau and its jackknife variance," *Stat. Probab. Lett.*, vol. 213, no. February, p. 110181, 2024, doi: 10.1016/j.spl.2024.110181.