



Performance Comparison of Multi-Criteria Decision-Making Methods in Decision Support Systems

Mahardika Inra Takaendengan*✉

Faculty of Mathematics and Natural Sciences, Universitas Sam Ratulangi, Manado, Indonesia

ABSTRACT

Keywords:

CRISUS
Multi-Criteria
Decision-Making
Comparison
WASPAS
Spearman Rank
Correlation

This study proposes an integrated Multi-Criteria Decision-Making (MCDM) framework based on the CRISUS weighting method and five ranking approaches, namely Simple Additive Weighting (SAW), Multi-Objective Optimization on the basis of Ratio Analysis (MOORA), Weighted Aggregated Sum Product Assessment (WASPAS), Grey Relational Analysis (GRA), and Multi-Attribute Utility Theory (MAUT), for evaluating and ranking alternatives across multiple criteria. The study aims to assess the consistency and robustness of alternative rankings when different MCDM methods are applied using the same decision matrix and criterion weights. CRISUS is first used to determine the relative importance of the evaluation criteria, and the resulting weights are subsequently incorporated into each ranking method to calculate final preference scores and rankings. The results indicate that Alternative-BM and Alternative-AV consistently achieve the highest rankings across the evaluated methods, while Alternative-JI and Alternative-AR remain among the lowest-ranked alternatives. Although several differences occur in the middle-ranking positions, the overall ranking patterns remain highly consistent. Spearman Rank Correlation analysis confirms this consistency, with SAW, MOORA, WASPAS, and MAUT obtaining correlation coefficients of 0.972, while GRA obtains 0.965. These results indicate a very strong positive relationship between the reference ranking and the rankings produced by the evaluated methods. Therefore, the proposed CRISUS-based MCDM framework demonstrates strong ranking stability and robustness and can provide a reliable basis for multi-criteria evaluation and decision-making.

I. Introduction

Decision Support Systems (DSS) have become an important technological approach for assisting decision makers in addressing complex problems involving multiple alternatives, criteria, and competing preferences[1], [2]. In many organizational and operational environments, decision makers are required to evaluate alternatives

*Corresponding Author: mahardika@unsrat.ac.id (Mahardika Inra Takaendengan)

DOI: <https://doi.org/10.67449/jcoda.v1i1.5>

Copyright © 2026 by the authors.

This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by-sa/4.0/>)



based on several criteria that differ in importance, scale, and characteristics. Relying solely on individual judgment can make the decision-making process difficult to reproduce and potentially inconsistent, particularly when the number of alternatives and criteria increases. MCDM methods provide a systematic framework for addressing these problems by processing quantitative and qualitative information, incorporating the relative importance of criteria, and producing a ranking of alternatives[3], [4]. Consequently, the integration of MCDM methods into DSS can provide a more structured basis for evaluating alternatives and supporting rational decision-making.

The effectiveness of an MCDM-based DSS depends substantially on the procedures used to determine criterion weights and rank alternatives. Criteria weighting represents the relative contribution or importance of each criterion in the decision-making process, whereas the ranking method determines how the weighted information is aggregated to obtain the final preference of each alternative[5], [6]. Errors or inconsistencies in either component can influence the final decision. In particular, different MCDM ranking methods employ different normalization techniques, mathematical transformations, and aggregation mechanisms. Although these methods may use the same decision matrix and criteria weights, they can generate different preference scores and alternative rankings. Therefore, the selection of an MCDM method should not be based solely on its ability to produce a ranking, but should also consider the consistency, stability, sensitivity, and robustness of the resulting decisions.

The use of objective criteria weighting can provide a more consistent foundation for comparing MCDM ranking methods[7], [8]. Unlike subjective weighting approaches that depend heavily on expert judgments, objective weighting derives criterion importance from the information contained in the decision matrix. One method that can be used for this purpose is CRISUS (CRiterion Importance based on SUM of Squares). CRISUS determines the relative importance of criteria based on the statistical characteristics of the observed data, allowing criterion weights to be generated without directly relying on subjective assessments[9]–[11]. In a comparative MCDM study, applying the same CRISUS weights to all ranking methods is particularly important because it establishes a common weighting condition. Thus, differences in the resulting rankings can be examined primarily from the perspective of the ranking mechanisms rather than being influenced by different weighting methods.

Among the numerous MCDM techniques available, Simple Additive Weighting (SAW), Multi-Objective Optimization on the Basis of Ratio Analysis (MOORA), Weighted Aggregated Sum Product Assessment (WASPAS), Grey Relational Analysis (GRA), and Multi-Attribute Utility Theory (MAUT) represent distinct approaches to alternative evaluation. SAW calculates an overall preference score through weighted summation after normalization. MOORA applies a ratio-based approach to evaluate the performance of alternatives according to beneficial and non-beneficial criteria. WASPAS integrates the Weighted Sum Model and Weighted Product



Model to combine additive and multiplicative aggregation mechanisms. GRA evaluates the degree of similarity between alternatives and a reference sequence through grey relational coefficients and grades. Meanwhile, MAUT transforms criterion performance into utility values and aggregates these utilities according to their respective weights. These differences provide an appropriate basis for examining whether the mathematical structure of an MCDM method influences ranking agreement and decision stability.

A comparison of MCDM methods is particularly relevant because the highest-ranked alternative obtained by one method may not necessarily remain the highest-ranked alternative when another method is applied. Such differences can become more pronounced when alternatives have similar performance across several criteria or when criterion weights are changed. A method that produces a ranking that changes substantially under relatively small variations in criterion weights may be less stable for decision support applications. Conversely, a method that maintains a similar ranking under different scenarios may provide greater robustness. Therefore, evaluating MCDM performance requires more than a direct comparison of the final ranking. Ranking correlation can be used to measure the degree of agreement among methods, while sensitivity analysis can examine the effect of changes in criterion weights. Ranking stability and robustness analysis can further determine whether the preferred alternatives remain consistent under different decision scenarios.

Previous MCDM studies have frequently concentrated on the application of a particular method to a specific case study or on comparing final rankings between several techniques. Although such comparisons provide useful information, they may not fully explain the behavioral characteristics and reliability of each method when the same objective weights are applied. In addition, differences in weighting methods, datasets, normalization procedures, and evaluation indicators can make it difficult to determine whether observed ranking differences originate from the weighting process or from the ranking algorithms themselves. A controlled comparison using one common objective weighting method can therefore provide a clearer experimental setting for evaluating the relative performance of alternative ranking techniques.

Based on this consideration, this study conducts a systematic performance comparison of SAW, MOORA, WASPAS, GRA, and MAUT using CRISUS as a common objective criteria weighting method within a decision support framework. The study evaluates the ranking results generated by the five methods and examines their degree of agreement using ranking correlation analysis. Furthermore, sensitivity analysis is performed by varying criterion weights to determine the extent to which the rankings respond to changes in criterion importance. Ranking stability and robustness are also assessed to identify the methods that consistently maintain their preferred alternatives across different scenarios. Through this comprehensive evaluation, the study aims to provide empirical evidence regarding the relative performance and reliability of the five MCDM



methods and to support the selection of an appropriate ranking method for MCDM-based Decision Support Systems.

II. Methodology

2.1. Research Framework

The research framework is designed to systematically compare the performance of five Multi-Criteria Decision-Making (MCDM) methods, namely SAW, MOORA, WASPAS, GRA, and MAUT, using CRISUS as a common objective weighting method. The framework begins with the identification of the decision problem, followed by data collection and the construction of the decision matrix. The criteria are then classified into benefit and cost attributes to ensure that the subsequent calculations are performed according to their respective characteristics.

The research procedure consists of several sequential stages, beginning with the preparation of the decision matrix and identification of decision criteria. The criteria are classified into benefit and cost attributes according to their influence on the decision. The CRISUS method is then applied to calculate the objective weight of each criterion. The resulting weight vector is subsequently used as the input for SAW, MOORA, WASPAS, GRA, and MAUT. Each method produces a preference score and ranking of alternatives. The resulting rankings are then compared using correlation, sensitivity, stability, and robustness analyses. The overall research framework is presented in Figure 1.

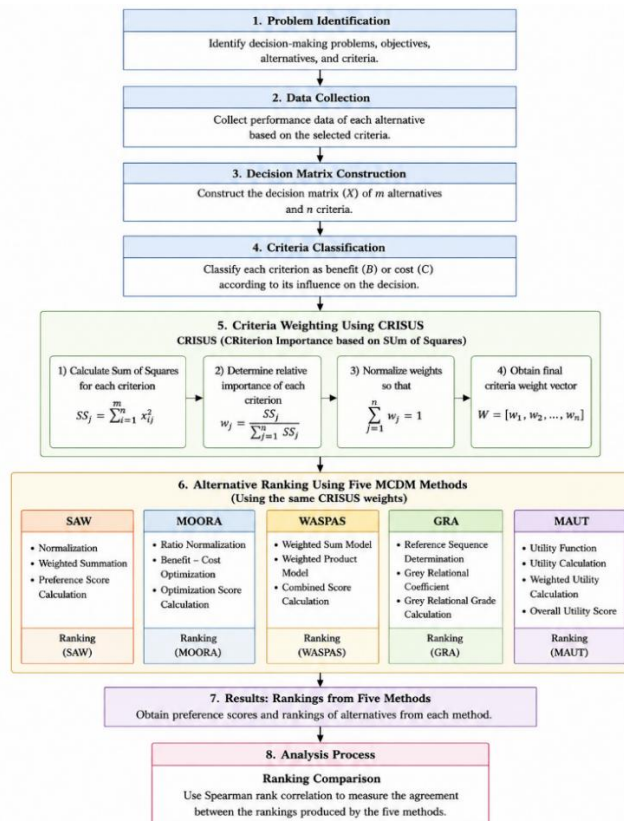


Figure 1. Research Framework

The research framework presents a systematic process for comparing the performance of five MCDM methods. The process begins with problem identification, followed by data collection and decision matrix construction, where the alternatives are evaluated based on predefined criteria. The criteria are then classified as benefit or cost attributes before their objective weights are calculated using the CRISUS method. The resulting CRISUS weights are consistently applied to SAW, MOORA, WASPAS, GRA, and MAUT to generate preference scores and alternative rankings. The rankings produced by the five methods are subsequently compared using Spearman rank correlation to measure the level of agreement and identify similarities or differences among the methods, providing a basis for determining the most consistent MCDM approach for decision support.

2.2. CRISUS Method

CRISUS method is an objective-criteria weighting method used to determine criterion importance based on the characteristics and distribution of values in the decision matrix[12]. The process begins with the construction of the decision matrix as shown in (1), which serves as the input for the weighting procedure. The matrix values are then normalized according to the criterion type, namely benefit or cost



criteria, using (2). Next, the normalized values are converted into proportional values using (3). The criterion importance measure is subsequently obtained by calculating the sum of squared proportional values using (4), while the dispersion of each criterion is determined through the standard deviation in (5). Finally, the criterion importance and dispersion measures are combined and normalized to generate the final objective criterion weights using (6).

$$X = [x_{ij}]_{m \times n} \quad (1)$$

$$r_{ij} = \begin{cases} \frac{x_{ij}}{\sqrt{\sum_{i=1}^m x_{ij}^2}}; \text{benefit criteria} \\ 1 - \frac{x_{ij}}{\sqrt{\sum_{i=1}^m x_{ij}^2}}; \text{cost criteria} \end{cases} \quad (2)$$

$$s_{ij} = \frac{r_{ij}}{\sum_{i=1}^m r_{ij}} \quad (3)$$

$$\rho_j = \sum_{i=1}^m s_{ij}^2 \quad (4)$$

$$\sigma_j = \sqrt{\frac{\sum_{i=1}^m (s_{ij} - \bar{s}_j)^2}{m}} \quad (5)$$

$$w_j = \frac{\rho_j * \sigma_j}{\sum_{j=1}^n \rho_j * \sigma_j} \quad (6)$$

The CRISUS calculation produces an objective weight for each criterion based on its relative importance within the decision matrix. The resulting weights indicate the contribution of each criterion to the decision-making process, with higher weights representing criteria with greater importance. These weights are subsequently used consistently in the SAW, MOORA, WASPAS, GRA, and MAUT methods for the comparative ranking analysis.

2.3. SAW Method

The SAW method is an MCDM approach that determines the preference of each alternative by calculating the weighted sum of normalized performance values across all criteria [13], [14]. The process begins with the decision matrix established in (1), followed by normalization to convert the values into a comparable scale, the normalization is calculated using (7). The normalized values are then multiplied by the corresponding CRISUS criterion weights and summed to obtain the preference score of each alternative using (9).

$$r_{ij} = \begin{cases} \frac{x_{ij}}{\max_j x_{ij}}; \text{benefit criteria} \\ \frac{\min_j x_{ij}}{x_{ij}}; \text{cost criteria} \end{cases} \quad (7)$$

$$V_i = \sum_{j=1}^m w_j * r_{ij} \quad (8)$$

The resulting SAW preference scores provide the basis for determining the performance ranking of all alternatives. The alternative with the highest score is considered the most preferred, while the remaining alternatives are ranked in descending order according to their preference scores.



2.4. MOORA Method

The MOORA method is an MCDM approach that evaluates alternatives by considering the ratio between each alternative's performance and the overall performance of the alternatives for each criterion[15], [16]. The process begins with the decision matrix presented in (1). The decision matrix is then normalized using the vector normalization procedure in (9) to obtain comparable values across criteria. The normalized values are multiplied by the corresponding CRISUS weights, and the weighted values of the benefit criteria are maximized while those of the cost criteria are minimized. The final MOORA score for each alternative is calculated using (10) by subtracting the weighted values of the cost criteria from the weighted values of the benefit criteria. A higher MOORA score indicates better alternative performance.

$$x_{ij}^* = \frac{x_{ij}}{\sqrt{\sum_{i=1}^m x_{ij}^2}} \quad (9)$$

$$Y_i^* = \sum_{j=1}^g x_{ij}^* * w_j - \sum_{j=g+1}^n x_{ij}^* * w_j \quad (10)$$

The resulting MOORA scores are used to rank the alternatives in descending order. The highest-scoring alternative is considered the most preferred alternative under the MOORA approach

2.5. WASPAS Method

The WASPAS method combines the Weighted Sum Model (WSM) and Weighted Product Model (WPM) to evaluate the performance of alternatives[17]. The process begins with the decision matrix presented in (1), followed by normalization of the decision matrix according to the type of criterion using (11). The two components are then combined to obtain the final WASPAS preference score using (12), where the contribution of the WSM and WPM components can be controlled by the aggregation coefficient. A higher preference score indicates a better alternative.

$$N_{ij} = \begin{cases} \frac{x_{ij}}{x_{jmax}}; \text{benefit criteria} \\ \frac{x_{jmin}}{x_{ij}}; \text{cost criteria} \end{cases} \quad (11)$$

$$Q_i = 0,5 \sum n_{ij} w_j + 0,5 \prod_{j=1}^n n_{ij}^{w_j} \quad (12)$$

The resulting WASPAS scores are used to rank the alternatives in descending order, with the alternative achieving the highest score considered the most preferred.

2.6. GRA Method

The GRA method is an MCDM approach that evaluates alternatives based on their degree of relationship with an ideal reference sequence[18], [19]. The process begins with the decision matrix presented in (1), followed by normalization according to the type of criterion to obtain comparable values. The normalized values are then used to establish the reference sequence and calculate the absolute differences between each alternative and the reference sequence using (13). The grey relational coefficient is subsequently calculated using (14), with the



distinguishing coefficient controlling the contrast among alternatives. The weighted grey relational grade is then calculated using the CRISUS criterion weights in (15).

$$X_{norm} = \frac{X_{ij} - X_{min}}{X_{max} - X_{min}} \quad (13)$$

$$V_{ij} = x_{ij} \cdot w_j \quad (14)$$

$$GRG_i = \frac{1}{n} \sum_{j=1}^n V_{ij} \quad (15)$$

The resulting grey relational grades are used to rank the alternatives in descending order, with the alternative achieving the highest grade considered the most preferred.

2.7. MAUT Method

The MAUT method is an MCDM approach that evaluates alternatives by transforming the performance of each criterion into a utility value on a comparable scale [20], [21]. The process begins with the decision matrix presented in (1), followed by normalization according to the type of criterion using (16). For benefit criteria, higher performance values result in higher utility values, whereas for cost criteria, lower values are more desirable. The utility value for each criterion is calculated using (17). The resulting utility values are then multiplied by the corresponding CRISUS criterion weights and aggregated to obtain the overall utility score of each alternative using (18).

$$x_{ij} = \begin{cases} \frac{x_{ij} - \min x_{ij}}{\max x_{ij} - \min x_{ij}}; \text{benefit criteria} \\ 1 + \frac{\min x_{ij} - x_{ij}}{\max x_{ij} - \min x_{ij}}; \text{cost criteria} \end{cases} \quad (16)$$

$$u_{ij} = \frac{e((r_{ij}^*)^2) - 1}{1.71} \quad (17)$$

$$u_{(x)} = \sum_{j=1}^n u_{ij} * w_j \quad (18)$$

The resulting MAUT utility scores are used to rank the alternatives in descending order, with the alternative achieving the highest score considered the most preferred.

2.8. Dataset

The dataset used in this study is adopted from the study by [22]. In the present study, the dataset is reused as a benchmark dataset to provide a consistent basis for comparing the performance of SAW, MOORA, WASPAS, GRA, and MAUT using CRISUS as the common objective weighting method. The dataset contains the decision alternatives and criteria required to construct the decision matrix, as presented in Table 1, which is subsequently processed through the proposed comparative MCDM framework.

Table 1. Dataset

Alternatives	C01	C02	C03	C04	C05
Alternative-DL	4	2800000	5	4	4

*Corresponding Author: mahardika@unsrat.ac.id (Mahardika Inra Takaendengan)

DOI: <https://doi.org/10.67449/jcoda.v1i1.5>

Copyright © 2026 by the authors.

This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by-sa/4.0/>)



Alternative-SL	5	3500000	4	5	4
Alternative-JI	3	2200000	3	3	3
Alternative-AL	5	3800000	5	4	5
Alternative-RV	4	3200000	4	4	4
Alternative-AR	2	2600000	3	2	2
Alternative-BM	5	3600000	5	5	5
Alternative-RN	3	2900000	3	4	3
Alternative-TO	4	3400000	4	5	4
Alternative-AV	5	4000000	5	5	5
Alternative-SI	4	2700000	3	3	3
Alternative-TN	5	3300000	4	4	4

Based on the dataset presented in Table 1, twelve alternatives are evaluated using five criteria with different performance values. The dataset provides the decision matrix required for the subsequent MCDM analysis, where each criterion is considered according to its respective characteristics. The data are first processed using the CRISUS method to determine the objective criterion weights, and the resulting weights are then applied consistently to SAW, MOORA, WASPAS, GRA, and MAUT to obtain and compare the alternative rankings.

III. Result and Discussion

The performance comparison of MCDM methods is an important aspect in developing reliable DSS, particularly when different methods may produce varying alternative rankings from the same decision data. This study focuses on comparing five MCDM methods, namely SAW, MOORA, WASPAS, GRA, and MAUT, using CRISUS as a common objective criteria weighting method. By applying the same dataset and criterion weights to all methods, the study provides a consistent basis for examining differences in ranking results and identifying the method that produces the most reliable and consistent decision outcomes. The comparison is conducted through the resulting alternative rankings and their agreement with the reference ranking, providing insights into the suitability of each method for supporting multi-criteria decision-making within DSS.

3.1. Calculating Criteria Weights Using the CRISUS Method

The calculation of criteria weights using the CRISUS method is performed to determine the relative importance of each criterion objectively based on the information contained in the decision matrix. This method utilizes the characteristics and distribution of the data for each criterion, allowing the resulting weights to be determined without relying on subjective judgments from decision makers. In this study, CRISUS is employed as the common weighting method for all selected MCDM methods, ensuring that each method is evaluated under the same weighting conditions.

The calculation begins with the decision matrix consisting of 12 alternatives and 5 criteria as the primary input. Each value in the



decision matrix is processed according to the characteristics of its corresponding criterion to produce comparable values. The CRISUS calculation then proceeds through several stages, including decision matrix processing, normalization, proportion calculation, determination of criterion importance, and measurement of data variation. All stages of the CRISUS calculation are performed using (1)-(6) according to the formulation described in the methodology section.

After the values for each calculation stage have been obtained, the criterion importance and variation measures are used to determine the relative weight of each criterion. The resulting criterion weights are then normalized so that they can be consistently applied in the subsequent decision-making process, with the total weight equal to one. This normalization ensures that the weights are represented on a common scale and can be appropriately combined with the performance values of the alternatives. The final results of the CRISUS weighting process are presented in Table 2.

Table 2. Criteria Weight Results Using the CRISUS Method

C01	C02	C03	C04	C05
0.2437	0.0640	0.2104	0.2376	0.2443

The results in Table 2 show that C05 has the highest weight at 0.2443, followed closely by C01 with a weight of 0.2437 and C04 with 0.2376. Meanwhile, C03 obtains a weight of 0.2104, while C02 has the lowest weight at 0.0640. These results indicate that the five criteria have different levels of relative importance within the decision-making process, with C05, C01, and C04 contributing more substantially than C03 and C02 based on the CRISUS weighting results. The resulting criterion weights are subsequently used as weighting parameters in the ranking process using SAW, MOORA, WASPAS, GRA, and MAUT. All five methods use the same CRISUS weights to prevent differences in ranking results from being influenced by different weighting approaches. Under these consistent weighting conditions, the rankings generated by each method can be compared objectively based on the calculation characteristics of each MCDM method. The resulting rankings are then used to evaluate the level of agreement between the MCDM methods and the reference ranking through ranking comparison.

3.2. Final Score Calculation Using MCDM Method

The calculation of final scores is conducted to evaluate and rank each alternative using the five selected MCDM methods, namely SAW, MOORA, WASPAS, GRA, and MAUT. At this stage, the normalized decision matrix is combined with the CRISUS criterion weights obtained from the previous calculation to provide a consistent basis for evaluating all alternatives. Each MCDM method applies its specific calculation procedure to aggregate the performance values across the criteria and produce a final preference score for each alternative. Although the same decision matrix and CRISUS weights are used, differences in the



calculation mechanisms of the five methods may produce different preference scores and rankings. The resulting final scores are therefore used to establish the ranking of alternatives and subsequently compare the ranking performance of SAW, MOORA, WASPAS, GRA, and MAUT against the reference ranking.

The SAW method is applied to determine the final preference score of each alternative by aggregating the normalized performance values according to the CRISUS criterion weights. The calculation begins by normalizing the decision matrix based on the characteristics of each criterion. Benefit criteria are normalized by dividing the value of an alternative by the maximum value of the corresponding criterion, while cost criteria are normalized by dividing the minimum value by the value of the alternative. The normalization process is calculated using (7). The normalized values are then multiplied by their respective CRISUS weights and summed to obtain the final preference score using (8).

The calculation is applied to all alternatives using the same decision matrix and CRISUS weights to ensure consistency with the comparative MCDM framework. After obtaining the final preference scores, the alternatives are arranged in descending order, where the alternative with the highest score receives the first rank. The complete results of the SAW calculation, including the final preference scores and ranking of each alternative, are presented in Table 3.

Table 3. Final Score of the SAW Method

Alternatives	Final Score
Alternative-DL	0.8412
Alternative-SL	0.8853
Alternative-JI	0.6256
Alternative-AL	0.9255
Alternative-RV	0.7928
Alternative-AR	0.4706
Alternative-BM	0.9751
Alternative-RN	0.6577
Alternative-TO	0.8377
Alternative-AV	0.9712
Alternative-SI	0.6625
Alternative-TN	0.8402

The MOORA method is applied to evaluate the performance of each alternative by considering the relative contribution of each criterion. The calculation begins with the decision matrix, where the performance values of all alternatives are normalized using the vector normalization procedure in (9). This process converts the values of different criteria into comparable scales while maintaining the relative performance of each alternative. The normalized values are then multiplied by the corresponding CRISUS criterion weights to incorporate the relative importance of each criterion in the decision-making process.

The weighted normalized values are subsequently aggregated according to the characteristics of the criteria. The weighted values



of the benefit criteria are summed, while the weighted values of the cost criteria are subtracted to obtain the overall MOORA optimization score. This calculation is performed using (10). The final scores are then arranged in descending order to determine the ranking of the alternatives, and the complete MOORA results are presented in Table 4.

Table 4. Final Score of the MOORA Method

Alternatives	Final Score
Alternative-DL	0.2639
Alternative-SL	0.2785
Alternative-JI	0.1862
Alternative-AL	0.2928
Alternative-RV	0.2467
Alternative-AR	0.1325
Alternative-BM	0.3107
Alternative-RN	0.1989
Alternative-TO	0.2623
Alternative-AV	0.3084
Alternative-SI	0.2001
Alternative-TN	0.2629

The WASPAS method combines the Weighted Sum Model (WSM) and Weighted Product Model (WPM) to obtain the final preference value of each alternative. The calculation begins by normalizing the decision matrix according to the characteristics of each criterion. The normalized values are then combined with the CRISUS criterion weights to calculate the weighted sum component and the weighted product component (11). These calculations provide two different perspectives for evaluating the performance of each alternative using (12).

The results of the WSM and WPM components are subsequently combined to obtain the final WASPAS preference score. A higher preference score indicates a better-performing alternative. The alternatives are then ranked in descending order based on their final scores, with the highest-scoring alternative receiving the first rank. The complete results of the WASPAS calculation, including the final preference scores and ranking of each alternative, are presented in Table 5.

Table 5. Final Score of the WASPAS Method

Alternatives	Final Score
Alternative-DL	0.8393
Alternative-SL	0.8812
Alternative-JI	0.6228
Alternative-AL	0.9206
Alternative-RV	0.7925
Alternative-AR	0.4638
Alternative-BM	0.9720
Alternative-RN	0.6549
Alternative-TO	0.8349
Alternative-AV	0.9668
Alternative-SI	0.6594



Alternative-TN 0.8376

The GRA method is applied to determine the degree of relationship between each alternative and the ideal reference sequence. The calculation begins with the normalized decision matrix obtained from the performance values of all alternatives. Based on the normalized values, the ideal reference sequence is determined, representing the best performance for each criterion. The absolute difference between each alternative and the reference sequence is then calculated using (13). This step provides the basis for measuring the degree of similarity between each alternative and the ideal condition.

The next stage calculates the grey relational coefficient for each alternative and criterion using (14). The coefficient reflects the closeness of each alternative's performance to the ideal reference sequence, where a higher coefficient indicates a stronger relationship with the ideal performance. The distinguishing coefficient is incorporated into this calculation to control the influence of differences between alternatives. The resulting grey relational coefficients are then combined with the CRISUS criterion weights to obtain the weighted grey relational grade using (15).

The final grey relational grades are used to evaluate the overall performance of the alternatives. A higher grey relational grade indicates that an alternative has a stronger overall relationship with the ideal reference and is therefore considered more preferable. The alternatives are subsequently ranked in descending order based on their grey relational grades. The complete results of the GRA calculation, including the final scores and ranking of the alternatives, are presented in Table 6.

Table 6. Final Score of the GRA Method

Alternatives	Final Score
Alternative-DL	0.1431
Alternative-SL	0.1591
Alternative-JI	0.0484
Alternative-AL	0.1827
Alternative-RV	0.1249
Alternative-AR	0.0028
Alternative-BM	0.1972
Alternative-RN	0.0692
Alternative-TO	0.1422
Alternative-AV	0.2000
Alternative-SI	0.0682
Alternative-TN	0.1419

The MAUT method is applied to evaluate the utility of each alternative based on its performance across the defined criteria. The calculation begins by transforming the original performance values into utility values so that the different criteria can be evaluated on a common scale. The utility value for each criterion is determined



according to whether the criterion represents a benefit or cost attribute. This transformation is calculated using (16), where higher utility values represent better performance after considering the characteristics of each criterion.

The utility values obtained from the previous stage are then combined with the corresponding CRISUS criterion weights to determine the overall utility value of each alternative. The weighted utility values across all criteria are aggregated using (17) to obtain the final MAUT score. A higher final utility score indicates a more preferred alternative because it represents stronger overall performance across the evaluated criteria.

The final MAUT score is subsequently used to rank all alternatives in descending order. The alternative with the highest utility score is assigned the first rank, while the alternative with the lowest score receives the last rank. The complete MAUT calculation results, including the final utility scores and ranking of the alternatives, are presented in Table 7.

Table 7. Final Score of the MAUT Method

Alternatives	Final Score
Alternative-DL	0.4698
Alternative-SL	0.6015
Alternative-JI	0.1142
Alternative-AL	0.7800
Alternative-RV	0.2806
Alternative-AR	0.0311
Alternative-BM	0.9424
Alternative-RN	0.1282
Alternative-TO	0.4378
Alternative-AV	0.9405
Alternative-SI	0.1385
Alternative-TN	0.4437

In addition, the discussion may include an evaluation of the performance of the proposed method, model, or system based on the testing and validation results. Authors may elaborate on the strengths, limitations, and implications of the findings, enabling readers to gain a more comprehensive understanding of the study's contributions and its potential directions for future research.

3.3. Ranking Comparison Results

The ranking comparison was done to find out the level of consistency and match of the results obtained from the method used with several other comparative MCDM methods. The comparison was carried out by looking at the position of each alternative in each method, so we can see whether alternatives that get high rankings in the proposed method also tend to have similar rankings in other methods. The evaluation results show that there are some alternatives with relatively consistent ranking positions, especially for alternatives with dominant performance values

*Corresponding Author: mahardika@unsrat.ac.id (Mahardika Inra Takaendengan)

DOI: <https://doi.org/10.67449/jcoda.v1i1.5>

Copyright © 2026 by the authors.

This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by-sa/4.0/>)



in most criteria. Even so, ranking differences are still found in some alternatives due to the characteristics of each method in performing normalization, weighting, and aggregation of performance values. The SAW method ranks alternatives based on the accumulation of performance values that have been normalized and multiplied by the criteria weights, while methods like MOORA, WASPAS, GRA, MAUT, and others use different calculation mechanisms, which can lead to changes in the position of certain alternatives. These differences are important to analyze because they show how sensitive each method is to the criteria values and the weights used. By comparing all the ranking results, you can find out which method produces a ranking pattern closest to the proposed method, as well as which alternatives have the most stable rankings across different MCDM approaches.

This analysis also provides a basis for assessing the reliability of the ranking results, since the consistency of the alternatives' positions across several methods can be an indication that the decision results are not just influenced by the characteristics of a single method. Overall, the ranking comparison shows a fairly good level of agreement between methods, even though there are some variations in certain ranking positions that reflect the differences in calculation mechanisms of each method. A complete comparison of the ranking results from the methods used and the comparison MCDM method is presented in Table 8.

Table 8. Ranking Comparison Results

Alternative	Rank ORI	Rank SAW	Rank MOORA	Rank WASPAS	Rank GRA	Rank MAUT
Alternative-AV	1	2	2	2	1	2
Alternative-BM	2	1	1	1	2	1
Alternative-AL	3	3	3	3	3	3
Alternative-SL	4	4	4	4	4	4
Alternative-TN	5	6	6	6	7	6
Alternative-TO	6	7	7	7	6	7
Alternative-DL	7	5	5	5	5	5
Alternative-RV	8	8	8	8	8	8
Alternative-SI	9	9	9	9	10	9
Alternative-RN	10	10	10	10	9	10
Alternative-JI	11	11	11	11	11	11
Alternative-AR	12	12	12	12	12	12

Based on the ranking comparisons obtained from each MCDM method, the next step is to perform a correlation analysis using Spearman Rank Correlation to measure how well the rankings from the proposed method match those from the comparison method. This analysis is used because the data being compared are ordinal rankings, so the Spearman coefficient can show how consistent the order of alternatives is across methods. The Spearman correlation coefficient ranges from -1 to 1, where values closer to 1 indicate a stronger positive relationship and that the two methods produce increasingly similar alternative rankings, while values near 0



indicate a weak relationship. Thus, correlation comparison is not only used to see alternatives that experience a change in position, but also to evaluate the overall consistency level of ranking results. The Spearman correlation calculation is done by comparing the ranks of the proposed method against each of the other MCDM methods, resulting in a correlation value for each pair of methods. These values are then used as a basis to assess which method has a ranking pattern closest to the proposed method. The higher the correlation value obtained, the higher the similarity in ranking patterns between the two methods, indicating that the ranking results are relatively stable when using different MCDM approaches. The Spearman correlation results from all those method comparisons were then visualized to give a clearer picture of how well each method matches up, as shown in Figure 2.

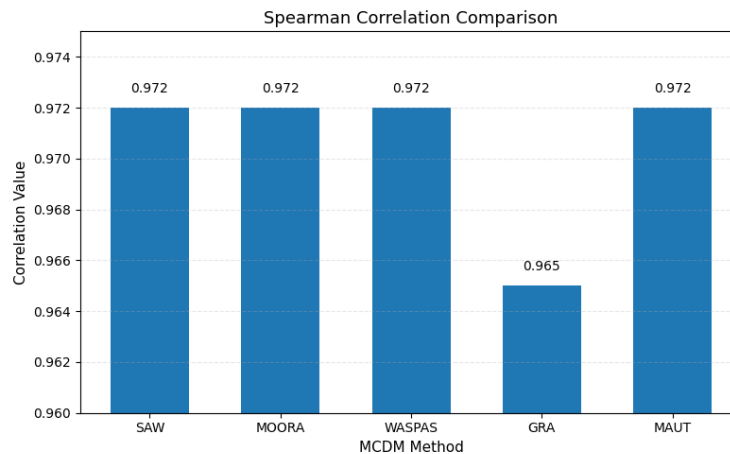


Figure 2. Comparison Results of Spearman Correlation Values

The comparison of Spearman correlation values shows that all MCDM methods have a very high correlation with the ranking results used as a reference. The SAW, MOORA, WASPAS, and MAUT methods each produced a correlation value of 0.972, while the GRA method got 0.965. Correlation values close to 1 indicate that the ranking patterns produced by each method are very strongly matched, so changes in alternative positions between methods are relatively small. The highest value of 0.972 achieved by SAW, MOORA, WASPAS, and MAUT shows a higher level of ranking consistency compared to GRA, even though the difference is relatively small. Meanwhile, the GRA correlation value of 0.965 still remains at a very strong level, indicating that this method also produces a ranking pattern that is almost the same as the reference method. Overall, these results show that the alternative rankings are highly consistent and stable across various MCDM methods, so the decision results can be considered relatively robust against using different ranking methods.



3.4. Discussion

The results demonstrate that the integration of CRISUS criterion weights with five different MCDM methods produces highly consistent evaluations of the alternatives, despite differences in the mathematical mechanisms used by each method. The final scores indicate that Alternative-BM and Alternative-AV consistently occupy the two highest positions across almost all approaches, while Alternative-AL and Alternative-SL remain in the third and fourth positions, respectively. This consistency is particularly evident in the ranking results, where Alternative-BM obtains the first rank in SAW, MOORA, and WASPAS and the second rank in GRA, whereas Alternative-AV obtains the first rank in GRA and the second rank in SAW, MOORA, WASPAS, and MAUT. The small interchange between these two alternatives can be explained by the different aggregation mechanisms applied by each method. SAW emphasizes the weighted sum of normalized performances, MOORA distinguishes between benefit and cost contributions, WASPAS combines additive and multiplicative aggregation, GRA evaluates the closeness of alternatives to an ideal reference sequence, and MAUT transforms criterion performance into utility values before aggregation. Nevertheless, these methodological differences do not substantially alter the overall ordering of the alternatives. Alternative-AL also demonstrates strong stability by maintaining the third position across all five methods, while Alternative-SL consistently occupies the fourth position. At the lower end of the ranking, Alternative-JI and Alternative-AR consistently receive the eleventh and twelfth positions, respectively, indicating that their relatively weak performance across the evaluated criteria remains evident regardless of the MCDM technique applied. These findings suggest that the underlying performance structure of the alternatives is sufficiently strong to produce similar decision outcomes even when different ranking mechanisms are employed.

The comparison also reveals that the differences among the methods are concentrated mainly in the middle-ranking alternatives rather than at the extreme positions. For example, Alternative-TN changes from rank 5 in the reference ranking to rank 6 in SAW, MOORA, WASPAS, and MAUT and rank 7 in GRA, while Alternative-TO shifts between ranks 6 and 7. Similarly, Alternative-DL is ranked seventh in the reference ranking but improves to fifth in all five MCDM methods, indicating that the reference approach and the tested methods place somewhat different emphasis on the performance characteristics of this alternative. Alternative-SI and Alternative-RN also show only minor changes, with their positions varying between ninth and tenth. These relatively small variations are important because they demonstrate that the choice of MCDM technique mainly affects alternatives with comparable performance rather than fundamentally changing the identification of the strongest and weakest alternatives. The consistency observed in the final scores further supports this interpretation. For instance, Alternative-BM records the highest SAW score of 0.9751, the highest MOORA score of 0.3107, the highest WASPAS



score of 0.9720, and a very high MAUT score of 0.9424, while Alternative-AV achieves the highest GRA score of 0.2000 and a MAUT score of 0.9405 that is very close to Alternative-BM. This pattern indicates that both alternatives possess strong overall performance across the criteria and remain competitive regardless of whether the evaluation emphasizes additive performance, multiplicative relationships, ideal-reference similarity, or utility-based aggregation. Conversely, Alternative-AR consistently obtains the lowest score and last rank, with scores of 0.4706 in SAW, 0.1325 in MOORA, 0.4638 in WASPAS, 0.0028 in GRA, and 0.0311 in MAUT. The consistency of this result indicates that its lower performance is not an artifact of a particular MCDM procedure. Therefore, the observed ranking stability provides practical evidence that the decision model is not excessively dependent on a single ranking technique and that the CRISUS weighting scheme provides a stable basis for evaluating the alternatives.

The Spearman Rank Correlation analysis provides further evidence regarding the robustness of the ranking results. The obtained coefficients are 0.972 for SAW, MOORA, WASPAS, and MAUT, while GRA produces a coefficient of 0.965. All coefficients are very close to 1, indicating a very strong positive relationship between the reference ranking and the rankings generated by the five MCDM methods. The identical coefficient of 0.972 for four methods suggests that SAW, MOORA, WASPAS, and MAUT generate highly similar ordering patterns, even though their internal computational procedures are fundamentally different. The slightly lower value obtained by GRA does not indicate weak agreement; rather, it reflects a small number of positional differences, particularly among the middle-ranked alternatives. This result is consistent with the ranking table, where GRA differs from the other methods mainly in the positions of Alternative-AV, Alternative-BM, Alternative-TN, Alternative-TO, and Alternative-SI. From a decision-making perspective, the high correlation values strengthen the reliability of the proposed evaluation because the identification of the leading and lowest-performing alternatives remains largely unchanged across different MCDM approaches. The findings therefore indicate that the ranking results are robust and relatively insensitive to the choice of aggregation method. In practical terms, Alternative-BM and Alternative-AV can be considered the strongest candidates, while Alternative-AR and Alternative-JI consistently occupy the lowest positions. The high level of agreement among the methods also suggests that the CRISUS-based weighting structure successfully captures the relative importance of the evaluation criteria and contributes to stable alternative prioritization. However, the small ranking variations observed among the middle positions should not be disregarded, particularly when decision makers need to distinguish between alternatives with closely related performance. In such cases, additional sensitivity or robustness analysis can provide further evidence regarding whether these positional differences are meaningful or merely a consequence of the mathematical characteristics of each MCDM method.

*Corresponding Author: mahardika@unsrat.ac.id (Mahardika Inra Takaendengan)

DOI: <https://doi.org/10.67449/jcoda.v1i1.5>

Copyright © 2026 by the authors.

This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by-sa/4.0/>)



IV. Conclusion

This study applied the CRISUS weighting method in combination with five MCDM approaches, namely SAW, MOORA, WASPAS, GRA, and MAUT, to evaluate and rank alternatives based on multiple criteria. The results demonstrate that the use of consistent criterion weights and a common decision matrix produces highly similar ranking patterns across the evaluated methods. Alternative-BM and Alternative-AV consistently emerged as the strongest alternatives, while Alternative-AL and Alternative-SL maintained stable positions in the upper part of the ranking. Conversely, Alternative-JI and Alternative-AR consistently occupied the lowest positions, indicating that their relatively lower performance was not strongly influenced by the selection of the ranking method. Although several differences were observed among the middle-ranked alternatives, these changes were relatively limited and did not substantially affect the overall decision pattern.

The Spearman Rank Correlation analysis further confirmed the consistency of the obtained rankings. SAW, MOORA, WASPAS, and MAUT achieved a correlation coefficient of 0.972, while GRA obtained 0.965. These values indicate a very strong positive relationship between the reference ranking and the rankings generated by all five MCDM methods. The results demonstrate that the proposed evaluation framework provides stable and robust alternative prioritization despite differences in the aggregation mechanisms of the individual MCDM methods. The findings also indicate that the integration of CRISUS criterion weights with multiple MCDM techniques can provide a reliable basis for comparative decision-making. Therefore, the proposed framework can support decision makers in identifying priority alternatives while reducing dependence on a single MCDM ranking procedure. Future research may extend this study by incorporating additional MCDM methods, alternative objective weighting techniques, larger datasets, and sensitivity or robustness analysis to further examine the stability of the rankings under changes in criterion weights and decision-making conditions.

CRedit Authorship Contribution Statement

M. I. Takaendengan: Writing - original draft, Writing - review & editing, Validation, Software, Methodology, Data curation, Conceptualization, Investigation.

Declaration of Competing Interest

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.



Data Availability

The data used to support the findings of this study are available from the corresponding author upon request.

Declaration of Generative AI and AI-assisted Technologies in The Writing Process

The authors used Generative AI to improve the writing clarity of this paper. They reviewed and edited the AI-assisted content and take full responsibility for the final publication.

References

- [1] S. Dünder, "Prioritizing the Regional Preferences of Turkish Investors Regarding Foreign Direct Investment by ARLON Method," *Konya J. Eng. Sci.*, vol. 13, no. 3, pp. 927-946, 2025, doi: [10.36306/konjes.1648279](https://doi.org/10.36306/konjes.1648279).
- [2] A. Ajithkumar and P. GaneshKumar, "A greener approach for selecting organic PCM in solar thermal collectors using the MCDM framework," *Sol. Energy*, vol. 305, p. 114283, 2026, doi: [10.1016/j.solener.2025.114283](https://doi.org/10.1016/j.solener.2025.114283).
- [3] K. Aliyeva, A. Aliyeva, R. Aliyev, and M. Özdeşer, "Application of Fuzzy Simple Additive Weighting Method in Group Decision-Making for Capital Investment," *Axioms*, vol. 12, no. 8, 2023, doi: [10.3390/axioms12080797](https://doi.org/10.3390/axioms12080797).
- [4] A. M. Yiğit, "Supplier Selection with AHP, TOPSIS and WASPAS as Multi-criteria Decision-making Techniques: An Application in Office Chairs Industry," *Uluslararası İşletme Bilim. ve Uygulamaları Derg.*, vol. 5, no. 2, pp. 116-130, 2025.
- [5] A. Suryadi, A. Adam Muiz, and A. Alpan Hidayat, "Implementation of a Decision Support System for Selecting the Best Supplier Using the SAW Method," *bit-Tech*, vol. 7, no. 3, pp. 928-935, Apr. 2025, doi: [10.32877/bt.v7i3.2255](https://doi.org/10.32877/bt.v7i3.2255).
- [6] D. Abdul, J. Wenqi, A. Tanveer, and I. M. Sharaf, "Analyzing the sustainability factors to prioritize sustainable energy source: A novel integrated MCDM model," *Energy*, vol. 337, p. 138637, 2025, doi: [10.1016/j.energy.2025.138637](https://doi.org/10.1016/j.energy.2025.138637).
- [7] Y. Liu, "A novel decision making approach based on Picture Hesitant Fuzzy Environment for optimizing art education strategies," *Sci. Rep.*, 2026, doi: [10.1038/s41598-026-48827-2](https://doi.org/10.1038/s41598-026-48827-2).
- [8] P. Rani, A. R. Mishra, A. M. Alshamrani, A. F. Alrasheedi, and D. Pamucar, "Adoption of Smart Manufacturing Technologies in Small and Medium Enterprises Using Picture Fuzzy Combinative Distance-Based Assessment Model," *J. Knowl. Econ.*, 2025, doi: [10.1007/s13132-025-02693-x](https://doi.org/10.1007/s13132-025-02693-x).
- [9] I. Adalar and Ö. IŞIK, "CRiterion Importance Based on SUM of Squares (CRISUS): A Novel Objective Weighting Method and Its Implementation in Multidimensional Sustainability Performance Measurement.," *Econ. Comput. Econ. Cybern. Stud. Res.*, vol. 59, no. 2, 2025, doi: [10.24818/18423264/59.2.25.13](https://doi.org/10.24818/18423264/59.2.25.13).
- [10] S. Bektaş, "Financial Performance Ranking With A New Hybrid Decision Making Model: Application Of The CRISUS-RAWEC Hybrid Model For The Reinsurance Industry," *Gümüşhane Univ. J. Soc. Sci.*, vol. 17, no. 1, pp. 233-249, 2026.
- [11] S. GÜL and S. BEKTAŞ, "ESG Performance Ranking with a Novel Hybrid



- Decision-Making Model: Application of the CRISUS-PIV Hybrid Model for the Banking Sector," *Üçüncü Sektör Sos. Ekon. Derg.*, vol. 60, no. 4, pp. 3557-3579, 2025, doi: [10.63556/tisej.2025.1582](https://doi.org/10.63556/tisej.2025.1582).
- [12] S. Setiawansyah, "Decision Support System for Selecting the Best Outsourcing Employee Using CRISUS and WASPAS," *JEECS (Journal Electr. Eng. Comput. Sci.)*, vol. 11, no. 1, pp. 36-49, 2026, doi: [10.54732/jeeecs.v11i1.4](https://doi.org/10.54732/jeeecs.v11i1.4).
- [13] 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, p. 100553, 2025, doi: [10.1016/j.dajour.2025.100553](https://doi.org/10.1016/j.dajour.2025.100553).
- [14] D. Sedghiyan, A. Ashouri, N. Maftouni, Q. Xiong, E. Rezaee, and S. Sadeghi, "RETRACTED: Prioritization of renewable energy resources in five climate zones in Iran using AHP, hybrid AHP-TOPSIS and AHP-SAW methods," *Sustain. Energy Technol. Assessments*, vol. 44, p. 101045, Apr. 2021, doi: [10.1016/j.seta.2021.101045](https://doi.org/10.1016/j.seta.2021.101045).
- [15] M. Bharadwaj, L. R. Sankargupta, R. Madurai Elavarasan, E. Devaraj, and A. Nandagopal, "Multi-Criteria Decision Analysis framework: Transitioning from coal to Large-Scale Photovoltaics in Australia for achieving SDG 7," *Appl. Energy*, vol. 405, p. 127208, 2026, doi: [10.1016/j.apenergy.2025.127208](https://doi.org/10.1016/j.apenergy.2025.127208).
- [16] A. R. Isnain and Y. Rahmanto, "Employee Performance Evaluation Using the Standard Method of Deviation Multi-Objective Optimization by Ratio Analysis," *J. Inf. Technol. Softw. Eng. Comput. Sci.*, vol. 2, no. 4, pp. 202-210, Oct. 2024, doi: [10.58602/itsecs.v2i4.164](https://doi.org/10.58602/itsecs.v2i4.164).
- [17] I. Tronnebati, F. Jawab, Y. Frichi, and J. Arif, "Green Supplier Selection Using Fuzzy AHP, Fuzzy TOSIS, and Fuzzy WASPAS: A Case Study of the Moroccan Automotive Industry," *Sustainability*, vol. 16, no. 11, pp. 1-22, 2024. doi: [10.3390/su16114580](https://doi.org/10.3390/su16114580).
- [18] N. Hendrastuty, S. Setiawansyah, M. G. An'ars, F. A. Rahmadianti, V. H. Saputra, and M. Rahman, "G2M weighting: a new approach based on multi-objective assessment data (case study of MOORA method in determining supplier performance evaluation)," *Indones. J. Electr. Eng. Comput. Sci.*, vol. 38, no. 1, pp. 403-416, 2025, doi: [10.11591/ijeecs.v38.i1.pp403-416](https://doi.org/10.11591/ijeecs.v38.i1.pp403-416).
- [19] H. Lu, Y. Zhao, X. Zhou, and Z. Wei, "Selection of Agricultural Machinery Based on Improved CRITIC-Entropy Weight and GRA-TOPSIS Method," *Processes*, vol. 10, no. 2, p. 266, Jan. 2022, doi: [10.3390/pr10020266](https://doi.org/10.3390/pr10020266).
- [20] S. Setiawansyah and A. Sulistiyawati, "Penerapan Metode Logarithmic Percentage Change-Driven Objective Weighting dan Multi-Attribute Utility Theory dalam Penerimaan Guru Bahasa Inggris," *J. Artif. Intell. Technol. Inf.*, vol. 2, no. 2, pp. 62-75, 2024, doi: [10.58602/jaiti.v2i2.119](https://doi.org/10.58602/jaiti.v2i2.119).
- [21] A. Yudhistira, S. Setiawansyah, T. Ardiansah, S. Maryana, Y. Yadin, and R. Oktaviani, "Development of Multi-Attribute Utility Theory Methods in Dynamic Decision Models Using Change-Data Driven," *Evergreen*, vol. 11, no. 4, pp. 3279-3289, Dec. 2024, doi: [10.5109/7326962](https://doi.org/10.5109/7326962).
- [22] D. A. Megawaty, H. Sulistiani, S. Setiawansyah, A. Qur'Ania, Y. Rahmanto, and P. Palupiningsih, "Modification of the Complex Proportional Assessment Method: A New Methodology for Decision Support," *Evergreen*, vol. 13, no. 1, pp. 229-239, 2026, doi: [10.5109/7411063](https://doi.org/10.5109/7411063).