



Benchmarking Aggregation-Based MCDM Methods: A Comparative Analysis of Ranking Consistency, Stability, and Robustness

Sumanto^{1*}✉, Mochamad Wahyudi¹✉, Lise Pujiastuti²✉

¹Faculty of Engineering and Informatics, Universitas Bina Sarana Informatika, DKI Jakarta, Indonesia

²Study Information System, STMIK Antar Bangsa, Tangerang, Indonesia

ABSTRACT

Keywords:

Multi-Criteria
Decision Making;
Decision Support
Systems;
Ranking Stability;
Ranking
Robustness;
Employee
Selection;

Multi-Criteria Decision Making (MCDM) methods are widely applied in Decision Support Systems (DSS) to address selection problems involving multiple criteria. However, differences in aggregation mechanisms may affect ranking consistency, stability, and robustness under changing decision conditions. This study presents a benchmarking analysis of five aggregation-based MCDM methods, namely Simple Additive Weighting (SAW), Weighted Aggregated Sum Product Assessment (WASPAS), Additive Ratio Assessment (ARAS), Complex Proportional Assessment (COPRAS), and Multi-Attributive Ideal-Real Comparative Analysis (MAIRCA), for employee selection. A dataset of ten candidates evaluated using six criteria is employed as the experimental basis. The methods are first applied using the same decision matrix and criterion weights to establish baseline rankings. Subsequently, 36 sensitivity scenarios are generated by increasing and decreasing individual criterion weights by 0.05, 0.10, and 0.15. Ranking consistency is evaluated using Spearman rank correlation, while stability and robustness are assessed using the Ranking Stability Index and Ranking Retention Rate. SAW, WASPAS, ARAS, and COPRAS achieve perfect consistency with the reference ranking, with a Spearman coefficient of 1.0000, while MAIRCA obtains 0.9879. ARAS and COPRAS achieve the highest stability, with an RSI of 97.22%, and robustness, with an RRR of 86.11%. SAW, WASPAS, and MAIRCA achieve an RSI of 95.56% and an RRR of 77.78%. Ranking changes are mainly limited to an exchange between Candidates 1 and 9, while Candidate 3 remains first across all scenarios. These findings suggest that ARAS and COPRAS exhibit comparatively higher ranking stability under the experimental conditions considered in this study.

Publication Dates:

Online First Date:
November 11, 2026

Published Date:
March 20, 2027

This is an open access article under the CC BY license



**Corresponding Author:**

Sumanto

Universitas Bina Sarana Informatika, Indonesia

Kramat Raya No.98, RT.2/RW.9 Kwitang, Central Jakarta, DKI Jakarta, Indonesia

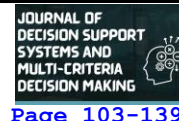
Email: sumanto@bsi.ac.id

I. Introduction

Multi-Criteria Decision Making (MCDM) plays an important role in the development of Decision Support Systems (DSS) by providing a systematic framework for evaluating and comparing alternatives based on multiple, often conflicting, criteria[1]. In practical decision-making environments, alternatives rarely can be assessed using a single performance indicator because decision problems generally involve several dimensions, such as cost, benefit, quality, risk, efficiency, and other relevant attributes. MCDM methods address this complexity by processing the decision matrix, criterion weights, and performance values to obtain an overall preference value for each alternative[2]. Consequently, MCDM enables DSS to transform multidimensional evaluation data into priorities or rankings, thereby supporting decision-makers in identifying the most preferable alternatives and making decisions in a more structured and transparent manner.

Among the various MCDM approaches, aggregation-based methods are widely adopted because their computational structures are relatively simple, intuitive, and easy to implement and integrate into DSS environments[3], [4]. These methods generally aggregate normalized performance values and criterion weights into a comprehensive preference score that can subsequently be used to rank alternatives. However, despite their computational simplicity, aggregation-based MCDM methods may produce different rankings when the underlying decision conditions change[5]. Variations in the decision matrix, criterion weights, or alternative performance values can alter the aggregated preference scores and, consequently, the relative positions of alternatives. This condition highlights the importance of evaluating not only the final ranking but also the ranking consistency, stability, and robustness of MCDM methods under changes in decision inputs.

Despite the widespread application of aggregation-based MCDM methods in Decision Support Systems, their effectiveness cannot be assessed solely based on the final ranking produced under a single decision scenario[6]. Different MCDM methods may generate different rankings when applied to the same decision problem, while even a single method may produce changes in alternative positions when the decision matrix or criterion weights are slightly modified[7], [8]. This raises an important research problem concerning the extent to which an MCDM method can maintain reliable ranking outcomes under different decision conditions. Therefore, a comprehensive evaluation is required to examine three related aspects: ranking consistency, which assesses whether a method produces comparable rankings across different evaluation

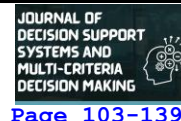


approaches or scenarios; ranking stability, which examines the degree to which rankings change in response to small variations in decision inputs; and ranking robustness, which evaluates the ability of a method to preserve its ranking structure when the decision matrix or criterion weights are subjected to perturbations.

These three aspects provide complementary perspectives for benchmarking MCDM methods. Ranking consistency reflects the agreement of ranking outcomes, ranking stability captures the sensitivity of rankings to relatively small changes, while ranking robustness represents the resistance of the ranking structure to systematic perturbations in the decision environment. A method that produces a high-quality ranking under the initial conditions may not necessarily demonstrate stable or robust behavior when its inputs change. Consequently, comparing MCDM methods based only on their original ranking may provide an incomplete assessment of their reliability for DSS applications. A systematic benchmarking framework that simultaneously evaluates consistency, stability, and robustness is therefore needed to identify the strengths and limitations of aggregation-based MCDM methods and to provide a more comprehensive basis for selecting an appropriate method for decision-support applications.

Previous studies on MCDM methods have predominantly focused on comparing the resulting rankings generated by different methods, often using a single dataset to assess their relative performance. In many cases, rank correlation measures have been employed as the primary evaluation tool to determine the degree of agreement between ranking results. Although such approaches are useful for identifying similarities and differences among methods, they provide a limited assessment of the overall reliability of MCDM methods because they do not necessarily capture how rankings behave under changes in decision conditions. In particular, previous research has rarely evaluated ranking consistency, ranking stability, and ranking robustness simultaneously, making it difficult to obtain a comprehensive understanding of method behavior under different scenarios. Furthermore, there remains a lack of a dedicated benchmarking framework for aggregation-based MCDM methods that systematically integrates these three evaluation dimensions to assess and compare the reliability of alternative ranking approaches in Decision Support Systems.

The objectives of this study are to systematically benchmark aggregation-based MCDM methods under a range of systematic perturbation scenarios to examine their behavior under changing decision conditions[9]-[11]. The study specifically aims to evaluate ranking consistency, ranking stability, and ranking robustness across different scenarios and identify the extent to which each method can maintain reliable ranking outcomes. Furthermore, the study seeks to analyze the sensitivity of each MCDM method to changes in criterion weights and alternative performance values, thereby determining which input factors have the greatest influence on ranking changes. Finally, this study aims to identify the most reliable aggregation-based MCDM method for DSS by considering its overall ability to produce consistent, stable, and robust



rankings, providing a comprehensive basis for method selection in practical multi-criteria decision-making applications.

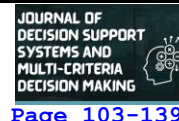
This study contributes a benchmarking framework specifically designed for aggregation-based MCDM methods, providing a structured approach for comparing their performance under diverse decision conditions. The proposed framework introduces a multi-dimensional evaluation that simultaneously considers ranking consistency, stability, and robustness, rather than relying solely on the similarity of ranking results. In addition, the study develops a systematic perturbation and simulation procedure to examine how changes in criterion weights and alternative performance values influence the resulting rankings. Finally, the study provides empirical evidence to support the selection of reliable MCDM methods for Decision Support Systems (DSS), enabling researchers and decision-makers to identify methods that demonstrate stronger and more dependable ranking behavior under varying decision scenarios.

II. Related Work

2.1. MCDM in Decision Support Systems

Multi-criteria decision making is a decision-making approach designed to evaluate and prioritize several alternatives based on multiple criteria that may have different characteristics and levels of importance. Unlike single-criterion decision-making, MCDM considers the simultaneous contribution of several criteria to obtain a comprehensive assessment of each alternative. The general MCDM process consists of several stages, including the identification of alternatives and criteria, construction of the decision matrix, normalization of criterion values, determination or application of criterion weights, aggregation of weighted performance values, and generation of the final ranking. Through this structure, MCDM provides a systematic mechanism for transforming multidimensional evaluation data into preference scores that can support objective and structured decision-making.

The fundamental representation of an MCDM problem is the decision matrix, which contains the performance values of alternatives with respect to the selected criteria. The matrix can be expressed as a set of alternatives evaluated across criteria, where each element represents the performance of an alternative on a particular criterion. Alternatives represent the available options or decision objects being evaluated, while criteria represent the attributes used to measure and differentiate their performance. Criteria can generally be classified into benefit and cost criteria. A benefit criterion is an attribute for which a higher value indicates better performance, whereas a cost criterion represents an attribute for which a lower value is preferable. This distinction is important because benefit and cost criteria require appropriate treatment during the normalization and aggregation processes to ensure that the resulting preference scores accurately reflect the decision-maker's objectives.



The integration of MCDM with DSS enables systematic decision analysis by combining structured data processing with mathematical evaluation techniques. Within a DSS, the decision matrix serves as the primary input, while MCDM methods process the performance values and criterion weights to generate preference scores and alternative rankings. This integration allows DSS to handle complex decision problems involving multiple alternatives, heterogeneous criteria, and conflicting objectives in a more transparent and reproducible manner. Furthermore, the modular structure of MCDM methods makes them relatively easy to incorporate into DSS architectures, enabling decision-makers to compare alternatives, perform scenario analysis, and evaluate the effects of changes in decision inputs. Therefore, MCDM functions as an important analytical component of DSS, particularly for applications requiring systematic prioritization and ranking of alternatives under multiple criteria.

2.2. Aggregation-Based MCDM Methods

Aggregation-based MCDM methods are a class of multi-criteria decision-making approaches that determine the overall preference of alternatives by aggregating their performance across multiple criteria into a single preference or utility value. The general principle is to transform the original performance values into comparable scales, incorporate the relative importance of criteria through weighting, and then combine the resulting weighted values to obtain an overall assessment for each alternative. The alternatives can subsequently be ranked according to their aggregated scores, with higher or lower values interpreted according to the specific formulation of the method. This computational structure makes aggregation-based methods relatively straightforward to understand, implement, and integrate into DSS.

A typical aggregation-based MCDM procedure begins with the construction of a decision matrix, followed by normalization to eliminate differences in measurement units among criteria. The normalized values are then combined with criterion weights to represent the relative importance of each criterion. The weighted values are aggregated using the mathematical formulation of the selected MCDM method, producing a final preference score for every alternative. Although the underlying principles are similar, different aggregation mechanisms can produce different ranking outcomes because each method may respond differently to criterion weights, performance distributions, and the relative differences among alternatives. Therefore, the choice of an aggregation-based MCDM method can influence the resulting decision priorities.

The simplicity of aggregation-based MCDM methods provides an important advantage for DSS applications, particularly when decision problems involve a large number of alternatives and criteria. However, their ranking results may be sensitive to changes in the decision matrix, criterion weights, and alternative performance values. A small perturbation in these inputs can potentially change preference scores and alter the positions of alternatives in the final ranking. Consequently, evaluating aggregation-based MCDM methods should not be



limited to their ability to generate a ranking under a single baseline condition. Their consistency, stability, and robustness under systematic changes in decision inputs are also important indicators for determining whether a method can provide reliable ranking outcomes in practical DSS environments.

The aggregation-based MCDM methods considered in this study include Simple Additive Weighting (SAW), Weighted Aggregated Sum Product Assessment (WASPAS), Additive Ratio Assessment (ARAS), Complex Proportional Assessment (COPRAS), and Multi-Attributive Ideal-Real Comparative Analysis (MAIRCA). These methods were selected because they employ aggregation-based mechanisms to combine the performance of alternatives across multiple criteria and generate an overall preference score for ranking. Although they share the general objective of transforming multidimensional performance information into a single ranking, each method adopts a different aggregation mechanism and mathematical structure, which may lead to variations in ranking outcomes. SAW applies a weighted additive aggregation, WASPAS combines weighted sum and weighted product approaches, ARAS evaluates alternatives relative to an optimal reference alternative, COPRAS considers the proportional significance of alternatives based on beneficial and non-beneficial criteria, while MAIRCA compares observed performance with theoretically expected performance. These methodological differences provide an appropriate basis for benchmarking their ranking behavior under systematic changes in decision inputs.

The comparison of SAW, WASPAS, ARAS, COPRAS, and MAIRCA is particularly relevant for examining whether differences in aggregation mechanisms affect ranking reliability. In the proposed benchmarking framework, each method is applied to the same decision matrix and criterion-weight configurations to establish comparable baseline rankings. The inputs are subsequently subjected to systematic perturbations involving changes in criterion weights and alternative performance values. The resulting rankings are evaluated from three complementary perspectives: ranking consistency, which examines the agreement of ranking outcomes; ranking stability, which measures the degree of ranking changes under small variations in inputs; and ranking robustness, which evaluates the ability of a method to preserve its ranking structure under broader or systematic perturbations. This comparative evaluation provides a more comprehensive assessment of the reliability of the five aggregation-based MCDM methods for integration into DSS.

2.3. Critical Review and Research Gap

To identify the research gap, previous studies on MCDM are critically reviewed based on the methods employed and the extent to which they evaluate ranking consistency, stability, robustness, and benchmarking performance. Existing research has provided valuable contributions through comparisons of multiple MCDM methods, rank-correlation analysis, sensitivity analysis, and evaluations under varying decision conditions. However, these studies differ in terms of



evaluation dimensions, perturbation strategies, and methodological scope, with many focusing on only one or two aspects of ranking reliability. A systematic comparison of the existing literature is therefore necessary to determine the limitations of current approaches and clarify the position of the proposed study. Table 1 summarizes selected relevant studies and compares their methodological characteristics and evaluation dimensions with those addressed in the proposed benchmarking framework.

Table 1. Comparative Analysis of MCDM Methods

Title	Methods	Consistency	Stability	Robustness	Benchmarking
A Comparative Analysis of MCDM Methods for Resource Selection in Mobile Crowd Computing [12].	Several MCDM	✓	✓	—	✓
A Novel Method: Integrative Reference Point Approach (IRPA) [13].	Several MCDM	✓	✓	—	✓
Sustainable Evaluation of Major Third-Party Logistics Providers [14].	SAW, WASPAS, COPRAS, ARAS, MARCOS, CODAS	✓	—	—	✓
Proposed Study	Aggregation-based MCDM	✓	✓	✓	✓

As summarized in Table 1, previous studies have demonstrated substantial efforts to compare MCDM methods and assess the reliability of their ranking outcomes. Several studies have employed rank-correlation measures, such as Spearman's rho and Kendall's tau, to examine the consistency or sensitivity of rankings, while others have incorporated sensitivity analysis and parameter variations to investigate ranking stability and robustness. Comparative studies have also evaluated multiple aggregation-based methods, including SAW, WASPAS, ARAS, COPRAS, and MAIRCA, under common decision-making conditions. These contributions provide important foundations for understanding the behavior of MCDM methods; however, the evaluation dimensions remain fragmented across studies.

Based on this review, a clear methodological gap can be identified in the absence of a comprehensive benchmarking framework specifically designed for aggregation-based MCDM methods. Existing studies tend to emphasize either comparison of ranking outcomes, sensitivity to criterion weights, or robustness under particular perturbation



conditions, rather than simultaneously evaluating how methods maintain ranking reliability across multiple types of changes. The proposed study addresses this limitation by establishing a systematic benchmarking framework for SAW, WASPAS, ARAS, COPRAS, and MAIRCA, incorporating controlled perturbations of criterion weights and alternative performance values. The resulting rankings are evaluated jointly in terms of consistency, stability, and robustness, enabling a more comprehensive assessment of the reliability of aggregation-based MCDM methods for Decision Support Systems. This framework is expected to provide stronger empirical evidence for selecting MCDM methods based not only on their baseline ranking performance but also on their ability to maintain reliable outcomes under changing decision conditions.

III. Methodology

Methodology is the approach or systematic framework used in a study to explain how the research is carried out to achieve its goals and answer the research questions. It includes the research design, data sources and characteristics, data processing stages, the methods or analysis techniques used, testing procedures, and the indicators used to evaluate the research results.

3.1. Research Framework

A research framework is a structured conceptual representation that describes the relationships among the main components, variables, methods, processes, and evaluation stages involved in a research study. It provides a systematic overview of how the research problem is addressed, beginning with the required inputs, followed by the methodological procedures, analytical processes, and evaluation stages, and ultimately leading to the expected research outcomes. A research framework therefore serves as a roadmap that ensures each stage of the study is logically connected to the research objectives and contributes to answering the research questions. The overall research framework is presented in Figure 1.

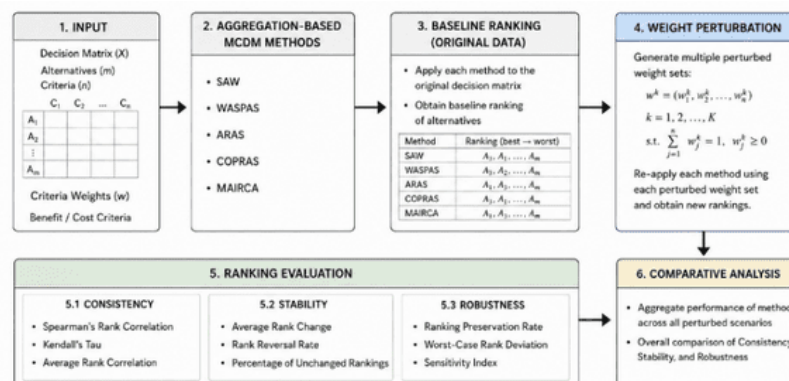


Figure 1. Research Framework



The research framework in Figure 1 is the process of benchmarking aggregation-based MCDM methods for Decision Support Systems, starting with creating a decision matrix with alternatives, criteria, weights, and criteria types, then applying five MCDM methods, namely SAW, WASPAS, ARAS, COPRAS, and MAIRCA, to get a baseline ranking based on the original data. Next, weight perturbation is carried out by creating several weight variations to test how each method responds to changes in criteria weights. The rankings generated from each scenario are then evaluated based on three main dimensions: consistency using rank correlation, stability based on changes and reversals in rankings, and robustness based on the ability to maintain rankings under weight variations. The evaluation results are then compared overall to identify the method with the most consistent, stable, and robust performance, making it a reliable basis for selecting MCDM methods in DSS.

3.2. Aggregation-Based MCDM Methods

Aggregation-based MCDM refers to a group of decision-making methods that evaluate and rank alternatives by aggregating their performance across multiple criteria into a comprehensive preference score. The general procedure involves constructing a decision matrix, classifying criteria as benefit or cost, normalizing the performance values, applying criterion weights, and aggregating the weighted values according to the mathematical structure of the selected method. Although the methods follow a similar overall principle, each method uses a different mechanism for normalization, aggregation, and preference calculation. These differences can influence the resulting ranking, particularly when the decision matrix or criterion weights are modified. In this study, five aggregation-based MCDM methods are employed, namely SAW, WASPAS, ARAS, COPRAS, and MAIRCA, to provide a comparative basis for evaluating ranking consistency, stability, and robustness.

The SAW is one of the most widely used aggregation-based MCDM methods for evaluating and ranking alternatives based on multiple criteria [15], [16]. The fundamental principle of SAW is to aggregate the normalized performance values of each alternative by considering the relative importance of each criterion. The method can accommodate both benefit and cost criteria and transforms heterogeneous performance values into a comparable scale before calculating the overall preference value. Due to its simple computational structure and intuitive interpretation, SAW is frequently applied in Decision Support Systems where multiple alternatives must be prioritized according to several criteria.

The SAW procedure consists of two main computational stages. First, the decision matrix is constructed to represent the performance of each alternative with respect to the considered criteria, as presented in (1). Second, the original decision matrix is transformed into normalized values according to the type of criterion. Benefit criteria are normalized by considering higher performance values as preferable, whereas cost criteria are normalized by considering lower performance values as preferable. The normalization process is performed using (2).



After obtaining the normalized decision matrix, each normalized value is multiplied by its corresponding criterion weight and aggregated across all criteria. The resulting preference value for each alternative is calculated using (3). The alternatives are then ranked according to their preference values, with a higher value indicating a more preferable alternative.

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

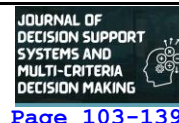
$$r_{ij} = \begin{cases} \frac{x_{ij}}{\max_j x_{ij}}; \text{benefit} \\ \frac{\min_j x_{ij}}{x_{ij}}; \text{cost} \end{cases} \quad (2)$$

$$V_i = \sum_{j=1}^n w_j * r_{ij} \quad (3)$$

The final ranking generated by SAW represents the overall performance of each alternative after considering the contribution and importance of all criteria. The direct additive aggregation mechanism makes SAW computationally simple and relatively easy to implement within a DSS. However, because the final preference value depends directly on normalized performance and criterion weights, changes in either the decision matrix or weighting structure may influence the resulting ranking. Therefore, SAW provides an important reference method in this study for evaluating how a simple additive aggregation mechanism behaves under systematic perturbation scenarios.

The WASPAS is an aggregation-based MCDM method that combines the principles of the Weighted Sum Model (WSM) and the Weighted Product Model (WPM) [17]–[19]. The method was developed to obtain a more comprehensive assessment by incorporating both additive and multiplicative aggregation mechanisms into a single preference calculation. WASPAS can be applied to decision problems involving multiple benefit and cost criteria and is particularly useful when the relative importance of criteria and the performance of alternatives need to be considered simultaneously.

The computational procedure of WASPAS begins with the normalization of the decision matrix according to the characteristics of benefit and cost criteria. The normalization process ensures that performance values measured on different scales become comparable and is performed using (4). The normalized values are subsequently combined with the corresponding criterion weights to construct the weighted assessment of each alternative. The weighted sum and weighted product components are then integrated according to the WASPAS aggregation mechanism to obtain the final preference value using (5). The alternatives are subsequently ordered according to their final preference values, where a higher value represents a more preferable alternative.



$$n_{ij} = \begin{cases} \frac{x_{ij}}{\max_j x_{ij}}; \text{benefit} \\ \frac{\min_j x_{ij}}{x_{ij}}; \text{cost} \end{cases} \quad (4)$$

$$Q_i = 0.5 \sum_{j=1}^n w_j * n_{ij} + 0.5 \prod_{j=1}^n n_{ij}^{w_j} \quad (5)$$

The integration of additive and multiplicative components gives WASPAS a different response mechanism compared with purely additive approaches such as SAW. This characteristic may affect the ranking when the relative performance of alternatives or the importance of criteria changes. Consequently, WASPAS is included in this study to determine whether its hybrid aggregation structure provides consistent and stable rankings under systematic perturbations. Its performance is subsequently evaluated together with the other MCDM methods based on ranking consistency, stability, and robustness.

The ARAS is an aggregation-based MCDM method that evaluates alternatives by comparing their overall performance with an optimal alternative or reference condition[20]–[22]. The central principle of ARAS is to determine the relative utility of each alternative based on its performance across multiple criteria. By incorporating an optimal reference into the assessment process, ARAS provides a relative measure of alternative performance and enables the alternatives to be ranked according to their degree of utility.

The ARAS procedure consists of several sequential stages. First, the decision matrix is normalized according to the type of criterion, with appropriate treatment for benefit and cost criteria. The normalization process is performed using (6). The normalized values are then incorporated into the weighting process to determine the contribution of each criterion. The preference value of each alternative is subsequently calculated using (7). Based on these values, the optimization value is determined using (8), which represents the relative performance of each alternative with respect to the optimal reference. Finally, the final ARAS value is obtained using (9), and the alternatives are ranked according to the resulting values.

$$x_{ij} = \frac{1}{x_{ij}^*}; \bar{x}_{ij} = \frac{x_{ij}}{\sum_{i=1}^m x_{ij}} \quad (6)$$

$$d_{ij} = w_j * \bar{x}_{ij} \quad (7)$$

$$S_i = \sum_{i=1}^m d_{ij} \quad (8)$$

$$K_i = \frac{S_i}{S_0} \quad (9)$$

The final ARAS value indicates the relative utility of an alternative compared with the optimal reference condition. A higher final value indicates a more favorable alternative. The use of an optimal reference distinguishes ARAS from direct aggregation methods because changes in performance values may influence not only the alternative's own score but also its relative position with respect to the reference.



Therefore, ARAS provides a useful methodological representation for the benchmarking analysis, particularly for examining the sensitivity of ranking outcomes to changes in performance values and criterion weights.

The COPRAS is an aggregation-based MCDM method that evaluates alternatives by considering the proportional significance of their performance across multiple criteria[17], [23], [24]. A distinctive characteristic of COPRAS is its explicit consideration of both benefit and cost criteria in the assessment process. Benefit criteria contribute positively to the performance of an alternative, whereas cost criteria represent undesirable characteristics that should be minimized. By incorporating these two components into the assessment, COPRAS determines the relative significance and utility of each alternative.

The COPRAS procedure begins with the construction and normalization of the decision matrix according to the characteristics of each criterion. The normalization process is performed using (10) to transform the original performance values into comparable values. The normalized values are subsequently multiplied by their respective criterion weights using (11). Based on the weighted normalized matrix, the performance indices of the alternatives are calculated using (12) and (13) by considering the contributions of benefit and cost criteria. The utility value of each alternative is then determined using (14), followed by the calculation of the final utility value using (15). The alternatives are subsequently ranked according to their final utility values.

$$r_{ij} = \frac{x_{ij}}{\sum_{i=1}^m x_{ij}} \quad (10)$$

$$d_{ij} = w_j * r_{ij} \quad (11)$$

$$S_{+i} = \sum_{j=1}^m d_{ij}^{+} \quad (12)$$

$$S_{-i} = \sum_{j=1}^m d_{ij}^{-} \quad (13)$$

$$Q_i = S_{+i} + \frac{\sum_{i=1}^m S_{-i}}{S_{-i} \sum_{i=1}^m (\frac{1}{S_{-i}})} \quad (14)$$

$$U_i = \left[\frac{Q_i}{\max Q_i} \right] * 100 \quad (15)$$

The final utility value obtained through COPRAS reflects the relative significance of each alternative after simultaneously considering beneficial and non-beneficial criteria. An alternative with a higher utility value is considered more preferable because it provides a more favorable balance between benefit and cost performance.

The MAIRCA is an aggregation-based MCDM method that evaluates alternatives by measuring the gap between their theoretical and actual performance across multiple criteria[25]-[27]. A distinctive characteristic of MAIRCA is its comparison between the expected contribution of each criterion and the observed contribution obtained from the decision matrix. The theoretical contribution represents the expected performance of an alternative based on the assigned criterion



weight, whereas the actual contribution reflects the performance of the alternative under each criterion. By evaluating the difference between these two components, MAIRCA determines the overall deviation of each alternative from its expected performance.

The MAIRCA procedure begins with the determination of the alternative preferences using (16), which represents the preference structure of the alternatives based on the available decision criteria. Subsequently, the theoretical evaluation matrix is constructed using (17) to represent the expected evaluation of each alternative according to the assigned criterion weights. The realistic evaluation matrix is then obtained using (18) by incorporating the actual performance values of the alternatives. The total gap between the theoretical and realistic evaluations is calculated using (19), representing the overall deviation of each alternative from its expected performance. Finally, the final value of each alternative is determined using (20) by aggregating the corresponding gap values across all criteria. The alternatives are subsequently ranked according to their final values, where a lower final value indicates a more preferable alternative.

$$p_{ij} = \frac{1}{m} \quad (16)$$

$$tp_{ij} = p_{ij} * w_j \quad (17)$$

$$tr_{ij} = \begin{cases} tp_{ij} * \left(\frac{x_{ij} - \min x_{ij}}{\max x_{ij} - \min x_{ij}} \right); \text{benefit} \\ tp_{ij} * \left(\frac{x_{ij} - \max x_{ij}}{\min x_{ij} - \max x_{ij}} \right); \text{cost} \end{cases} \quad (18)$$

$$g_{ij} = tp_{ij} - tr_{ij} \quad (19)$$

$$Q_i = \sum_{j=1}^n g_{ij} \quad (20)$$

The final value obtained through MAIRCA represents the overall deviation between the expected and actual performance of each alternative. An alternative with a lower total gap is considered more preferable because its actual performance is closer to the theoretical performance determined by the criterion weights. Therefore, MAIRCA provides an aggregation-based mechanism for identifying alternatives that most closely satisfy the expected contribution of the considered criteria.

The five methods provide different aggregation mechanisms while maintaining the common objective of transforming multi-criteria performance information into an alternative ranking. SAW represents a direct additive aggregation approach, WASPAS combines additive and multiplicative aggregation, ARAS evaluates alternatives relative to an optimal reference, COPRAS explicitly balances benefit and cost contributions, and MAIRCA assesses the deviation between theoretical and real performance. These methodological differences are important for the proposed benchmarking experiment because they may produce different responses to perturbations in criterion weights and performance values. Consequently, the five methods are applied under identical baseline and



perturbation scenarios, and their resulting rankings are subsequently evaluated in terms of consistency, stability, and robustness.

IV. Experimental Design

Experimental design refers to a systematic framework established to evaluate the performance and reliability of the proposed method through a series of structured experimental scenarios. In MCDM research, the experimental design may involve multiple datasets, criterion-weight configurations, benchmark methods, and decision-making conditions to assess the consistency, stability, and sensitivity of the resulting rankings. Each method is evaluated under the same experimental conditions to ensure a fair and objective comparison of their performance. Therefore, the experimental design provides empirical evidence regarding the effectiveness, robustness, and generalizability of the proposed method.

4.1. Decision-Making Dataset

The Decision-Making Dataset serves as the primary empirical input for evaluating the performance of the selected aggregation-based MCDM methods. It consists of a set of alternatives that are assessed according to multiple decision criteria, with each criterion representing a specific attribute relevant to the decision problem. The dataset includes both benefit and cost criteria, allowing the evaluation to capture different preference directions during the MCDM process. Each alternative is represented by a vector of performance values across the criteria, forming the decision matrix used as the baseline configuration. The original decision matrix and criterion weights are first used to generate the baseline rankings using SAW, WASPAS, ARAS, COPRAS, and MAIRCA. Subsequently, controlled perturbations are introduced to criterion weights and performance values to generate alternative decision scenarios, which serve as the basis for evaluating the consistency, stability, and robustness of each MCDM method.

The dataset used in this study was adopted from [28], the study presents an employee admission selection dataset consisting of candidate assessment data evaluated according to six criteria: Education (K1), Work Experience (K2), Creativeness (K3), English Language Proficiency (K4), Technical Skills (K5), and Leadership (K6). These criteria represent different aspects of candidate qualifications and competencies used to support employee selection decisions. The original assessment data are used to establish the baseline ranking and subsequently serve as the basis for systematic perturbation scenarios involving changes in criterion weights and performance values. The assessment data used in this study are presented in Table 2.

Table 2. Dataset [28]

Alternatives	K1	K2	K3	K4	K5	K6
Cand A1	4	5	85	90	8	4
Cand A2	3	3	75	80	7	3



Cand A3	5	7	90	95	9	5
Cand A4	4	4	70	85	6	3
Cand A5	3	2	60	70	5	2
Cand A6	5	6	80	92	9	4
Cand A7	4	4	78	88	8	4
Cand A8	2	1	50	65	4	2
Cand A9	4	5	80	89	7	5
Cand A10	3	3	65	75	6	3

Based on the decision-making dataset shown in Table 2, the criteria weights are set to represent the relative importance of each evaluation criterion in the employee recruitment selection problem. The criteria weights and the basic results obtained from these MCDM methods are shown in Table 3.

Table 3. Criteria Weight Value

K1	K2	K3	K4	K5	K6
0.1650	0.1808	0.1618	0.1608	0.1644	0.1673

The ranking results are obtained by ordering the alternatives based on their respective preference values generated from the decision-making dataset. This ranking represents the initial or baseline ranking of the alternatives before any perturbation is introduced into the decision matrix or criterion weights. The baseline ranking is important because it serves as the reference for the subsequent benchmarking analysis, particularly in evaluating how the positions of alternatives change under different perturbation scenarios. The ranking results derived from the decision-making dataset using the selected aggregation-based MCDM methods are presented in Table 4.

Table 4. Alternative Ranking Results

Alternatives	Rank
Cand A3	1
Cand A6	2
Cand A9	3
Cand A1	4
Cand A7	5
Cand A4	6
Cand A2	7
Cand A10	8
Cand A5	9
Cand A8	10

Overall, the decision-making dataset and the resulting baseline rankings provide the initial foundation for the benchmarking experiment. By applying the same alternatives, criteria, and criterion weights across all selected aggregation-based MCDM methods, the resulting rankings can be compared under a common decision environment. These baseline results establish a reference point for examining the effects of systematic perturbations introduced in the subsequent analysis. The changes in ranking positions across different scenarios will then be used to assess the consistency, stability, and robustness of SAW, WASPAS, ARAS, COPRAS,



and MAIRCA, thereby providing a comprehensive basis for comparing their reliability in Decision Support Systems.

4.2. Aggregation-based MCDM Method Application

The aggregation-based MCDM methods are applied to the decision-making dataset to obtain the preference scores and baseline rankings of the evaluated alternatives. In this stage, the same decision matrix, criterion classification, and initial criterion weights are used across all methods to ensure a fair and consistent comparison. The five selected methods, namely SAW, WASPAS, ARAS, COPRAS, and MAIRCA, are independently applied according to their respective normalization, weighting, and aggregation mechanisms. Although the methods share the common objective of aggregating multi-criteria performance into an overall preference measure, their computational structures may produce different preference scores and ranking orders. The resulting baseline rankings are subsequently used as the reference configuration for the systematic perturbation experiments and the evaluation of ranking consistency, stability, and robustness.

The implementation of the SAW method begins with the decision-making dataset presented in Table 2 and the criterion weights provided in Table 3. The performance values of each alternative are first normalized according to the type of criterion. For benefit criteria, the normalization process assigns higher normalized values to alternatives with higher performance, whereas for cost criteria, lower performance values receive higher normalized values. The resulting normalized values are then multiplied by their corresponding criterion weights to obtain the weighted normalized values. These values are subsequently summed across all criteria to calculate the overall preference score for each alternative. The alternatives are then ranked in descending order of their preference scores, where the alternative with the highest score represents the most preferred option. The resulting preference scores and ranking constitute the baseline results of the SAW method and are presented as the reference for the subsequent analysis. The preference values and ranking results obtained using the SAW method are presented in Table 5.

Table 5. Alternative Ranking Results using SAW Method

Alternatives	Final Value SAW	Rank
Cand A3	1.0000	1
Cand A6	0.9180	2
Cand A9	0.8510	3
Cand A1	0.8460	4
Cand A7	0.8040	5
Cand A4	0.7150	6
Cand A2	0.6750	7
Cand A10	0.6300	8
Cand A5	0.5350	9
Cand A8	0.4320	10



The implementation of the WASPAS method begins with the decision-making dataset presented in Table 2 and the criterion weights provided in Table 3. The performance values of each alternative are first normalized according to the type of criterion. For benefit criteria, the normalization process assigns higher normalized values to alternatives with higher performance, whereas for cost criteria, lower performance values receive higher normalized values. The resulting normalized values are then multiplied by their corresponding criterion weights to obtain the weighted performance values. These values are subsequently used to calculate the WSM and WPM components for each alternative. The two components are then aggregated to obtain the overall WASPAS preference score. The alternatives are ranked in descending order of their preference scores, where the alternative with the highest score represents the most preferred option. The resulting preference scores and ranking constitute the baseline results of the WASPAS method and are presented as the reference for the subsequent analysis. The preference values and ranking results obtained using the WASPAS method are presented in Table 6.

Table 6. Alternative Ranking Results using WASPAS Method

Alternatives	Final Value WASPAS	Rank
Cand A3	1.0000	1
Cand A6	0.9160	2
Cand A9	0.8480	3
Cand A1	0.8440	4
Cand A7	0.7990	5
Cand A4	0.7100	6
Cand A2	0.6660	7
Cand A10	0.6240	8
Cand A5	0.5220	9
Cand A8	0.4100	10

The implementation of the ARAS method begins with the decision-making dataset presented in Table 2 and the criterion weights provided in Table 3. The performance values of each alternative are first combined with an optimal reference alternative to establish the assessment framework. The decision matrix is then normalized according to the type of criterion, where benefit criteria are normalized so that higher performance receives a higher value, while cost criteria are transformed so that lower performance is considered more favorable. The normalized values are subsequently multiplied by their corresponding criterion weights to obtain the weighted normalized values. These weighted values are then aggregated across all criteria to calculate the overall optimality or utility degree of each alternative relative to the optimal reference alternative. The alternatives are ranked in descending order of their utility degrees, where the alternative with the highest value represents the most preferred option. The resulting preference values and ranking constitute the baseline results of the ARAS method and are presented as the reference for the subsequent analysis. The preference



values and ranking results obtained using the ARAS method are presented in Table 7.

Table 7. Alternative Ranking Results using ARAS Method

Alternatives	Final Value ARAS	Rank
Cand A3	1.0000	1
Cand A6	0.9177	2
Cand A9	0.8507	3
Cand A1	0.8462	4
Cand A7	0.8044	5
Cand A4	0.7149	6
Cand A2	0.6749	7
Cand A10	0.6302	8
Cand A5	0.5352	9
Cand A8	0.4317	10

The implementation of the COPRAS method begins with the decision-making dataset presented in Table 2 and the criterion weights provided in Table 3. The performance values of each alternative are first normalized according to the type of criterion. For benefit criteria, higher performance values are assigned higher normalized values, whereas for cost criteria, lower performance values receive higher normalized values. The normalized values are then multiplied by their corresponding criterion weights to obtain the weighted normalized values. These weighted values are subsequently aggregated separately for benefit and cost criteria to determine the positive and negative sums for each alternative. The relative significance of each alternative is then calculated by considering the positive and negative sums, resulting in the overall COPRAS preference value. The alternatives are ranked in descending order of their preference values, where the alternative with the highest value represents the most preferred option. The resulting preference values and ranking constitute the baseline results of the COPRAS method and are presented as the reference for the subsequent analysis. The preference values and ranking results obtained using the COPRAS method are presented in Table 8.

Table 8. Alternative Ranking Results using COPRAS Method

Alternatives	Final Value COPRAS	Rank
Cand A3	100.00	1
Cand A6	91.21	2
Cand A9	84.27	3
Cand A1	83.47	4
Cand A7	78.78	5
Cand A4	70.14	6
Cand A2	65.44	7
Cand A10	61.40	8
Cand A5	51.40	9
Cand A8	40.83	10

The implementation of the MAIRCA method begins with the decision-making dataset presented in Table 2 and the criterion weights provided in Table 3. The performance values of each alternative are first



normalized according to the type of criterion. For benefit criteria, higher performance values receive higher normalized values, whereas for cost criteria, lower performance values are considered more favorable. The criterion weights are then used to determine the theoretical rating of each alternative, representing the expected contribution of each criterion to the decision-making process. These theoretical ratings are subsequently compared with the corresponding real ratings obtained from the normalized performance values to determine the gap between the expected and actual assessment values. The resulting differences are aggregated across all criteria to obtain the overall MAIRCA assessment value for each alternative. The alternatives are ranked according to their assessment values, where a smaller total gap indicates a more favorable alternative. The resulting assessment values and ranking constitute the baseline results of the MAIRCA method and are presented as the reference for the subsequent analysis. The assessment values and ranking results obtained using the MAIRCA method are presented in Table 9.

Table 9. Alternative Ranking Results using MAIRCA Method

Alternatives	Final Value MAIRCA	Rank
Cand A3	0.0000	1
Cand A6	0.0142	2
Cand A1	0.0251	3
Cand A9	0.0254	4
Cand A7	0.0320	5
Cand A4	0.0490	6
Cand A2	0.0549	7
Cand A10	0.0649	8
Cand A5	0.0815	9
Cand A8	0.1000	10

Overall, the implementation of SAW, WASPAS, ARAS, COPRAS, and MAIRCA provides the baseline preference values and ranking results for the employee selection dataset. Although all methods use the same decision matrix, criterion weights, and criterion classifications, their different aggregation mechanisms may produce variations in preference values and ranking positions. These baseline rankings serve as the reference configuration for the subsequent benchmarking analysis, in which systematic perturbations are introduced to the criterion weights and alternative performance values. The resulting changes in ranking outcomes are then analyzed to evaluate the consistency, stability, and robustness of each MCDM method and to identify the method that provides the most reliable ranking performance for Decision Support Systems.

V. Result and Discussion

This section presents the results of the benchmarking analysis and discusses the ranking behavior of the selected aggregation-based MCDM methods, namely SAW, WASPAS, ARAS, COPRAS, and MAIRCA. The analysis begins by establishing the baseline ranking results obtained under the



initial decision-making configuration, followed by a comparison of the ranking positions of the evaluated alternatives. Subsequently, systematic changes in criterion weights are examined to identify their effects on alternative rankings, while ranking consistency, stability, and robustness are evaluated to assess the reliability of each method under different decision conditions. Finally, a comparative analysis is conducted to synthesize the results across all evaluation dimensions and identify the MCDM method with the most reliable overall ranking performance for DSS.

5.1. Comparison of Alternative Rankings

The comparison of alternative rankings is conducted to examine the extent to which the selected aggregation-based MCDM methods reproduce the original ranking of employee candidates. The original ranking is used as the reference, while the rankings generated by SAW, WASPAS, ARAS, COPRAS, and MAIRCA are compared against it. This comparison provides an initial assessment of ranking agreement before the subsequent analyses of consistency, stability, and robustness. The analysis focuses on the ranking positions of each candidate and identifies changes across the evaluated methods. The graphical comparison is presented in Figure 2.

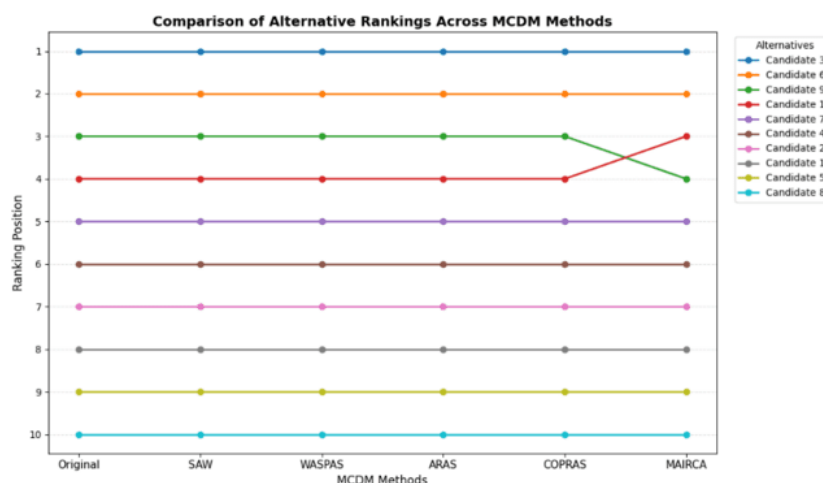


Figure 2. Comparison Ranking Alternative

Figure 2 illustrates that the rankings generated by SAW, WASPAS, ARAS, and COPRAS completely overlap with the original ranking across all ten candidates. In contrast, MAIRCA produces a slight deviation involving Candidate 9 and Candidate 1. Candidate 9 changes from the third position in the original ranking to the fourth position under MAIRCA, while Candidate 1 moves from the fourth position to the third position. The remaining candidates maintain their respective ranking positions. This pattern indicates that the overall ranking structure is highly similar across the evaluated methods, although MAIRCA produces a minor rearrangement among two alternatives.



The SAW method produces a ranking that is identical to the original ranking. Candidate 3 obtains the first position, followed by Candidate 6, Candidate 9, Candidate 1, Candidate 7, Candidate 4, Candidate 2, Candidate 10, Candidate 5, and Candidate 8. No changes in ranking position are observed for any candidate. This result indicates that the additive aggregation mechanism of SAW preserves the relative ordering of the alternatives under the baseline decision-making configuration. The complete agreement between the SAW and original rankings provides an initial indication of strong ranking consistency.

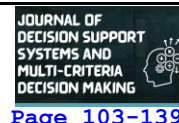
The WASPAS method also generates exactly the same ranking as the original ranking. Candidate 3 remains in first position and Candidate 6 in second position, while Candidate 9 and Candidate 1 occupy the third and fourth positions, respectively. The remaining candidates likewise retain their original positions. Despite WASPAS incorporating both weighted sum and weighted product components, the combination of these aggregation mechanisms does not result in any change in the ranking order for the current dataset. Therefore, WASPAS demonstrates complete agreement with the original ranking under the baseline conditions.

The ARAS method produces the same ranking order as the original ranking for all ten candidates. Candidate 3 achieves the highest ranking, followed by Candidate 6 and Candidate 9, while Candidate 8 remains in the tenth position. No candidate experiences a change in ranking position. This result indicates that the utility-based assessment mechanism of ARAS produces an ordering that is consistent with the reference ranking. The absence of ranking changes suggests that the relative performance of the candidates is sufficiently preserved through the ARAS assessment process.

The COPRAS method similarly reproduces the original ranking without any changes. Candidate 3 remains the highest-ranked alternative, followed by Candidate 6, Candidate 9, and Candidate 1. Candidates 7, 4, 2, 10, 5, and 8 also maintain their original positions from fifth to tenth, respectively. Thus, the proportional assessment mechanism used by COPRAS results in the same alternative ordering as the original ranking. The complete overlap between the COPRAS and original rankings indicates a high level of agreement under the baseline configuration.

The MAIRCA method produces the only observable difference in the ranking comparison. Candidate 1 moves upward from fourth to third position, while Candidate 9 moves downward from third to fourth position. This exchange does not affect the ranking positions of the other eight candidates, which remain identical to the original ranking. The result suggests that the comparison between theoretical and real assessment values in MAIRCA produces a slightly different relative evaluation for Candidates 1 and 9. Nevertheless, the magnitude of the difference is limited, as only two alternatives exchange positions.

Overall, the comparison demonstrates that four methods—SAW, WASPAS, ARAS, and COPRAS—fully reproduce the original ranking, whereas MAIRCA introduces only one pairwise rank exchange between Candidate 1 and Candidate 9. Therefore, the baseline ranking results show a very high degree of agreement among the evaluated methods. However, this



baseline similarity does not necessarily imply that the methods will respond identically to changes in the decision environment. Consequently, the subsequent analyses examine ranking consistency, stability, and robustness under systematic changes in criterion weights and alternative performance values.

5.2. Alternative Rankings Based on Changes in Criteria Weights

The analysis was conducted by systematically varying the weights of each criterion to evaluate the sensitivity and robustness of the resulting alternative rankings under different decision-making preferences[29]–[31]. For each criterion, its initial weight was independently increased and decreased by 0.05, 0.10, and 0.15, generating a series of weight-variation scenarios. Thus, the analysis considers both positive and negative adjustments, namely +0.05, -0.05, +0.10, -0.10, +0.15, and -0.15, applied separately to each criterion while the weights of the remaining criteria were adjusted accordingly to preserve the required total weight. This procedure was repeated for all criteria in the decision matrix, allowing the effect of changes in the importance of each criterion to be examined individually. The resulting rankings were compared with the original ranking to identify changes in alternative positions and assess the ranking stability of each MCDM method. This sensitivity analysis evaluates the robustness of the methods under moderate and substantial changes in criterion importance.

To evaluate the stability of the ranking results, a sensitivity analysis was conducted by applying several scenarios involving changes in the criterion weights of the SAW method. Each scenario represents a variation in the relative importance of the criteria while maintaining the overall weighting structure, allowing the effect of weight changes on the ranking positions of the alternatives to be examined[32]–[34]. The resulting changes in the alternative rankings across the different weighting scenarios are presented in Figure 3. This analysis provides an overview of the robustness of the SAW ranking and identifies alternatives whose positions are more sensitive to changes in criterion importance.

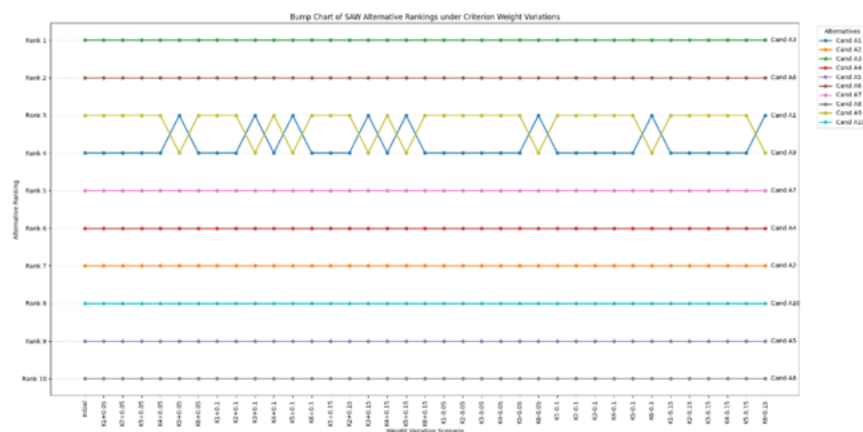


Figure 3. SAW Method Ranking Sensitivity Analysis



The sensitivity analysis results in Figure 3 show that changes in criteria weights within the range of +0.05 to +0.15 and -0.05 to -0.15 generally don't cause significant changes to the ranking structure of the alternatives. Initially, A3 was ranked first, followed by A6, A9, A1, A7, A4, A10, A2, A8, and A5. Most scenarios kept this order, which indicates a high level of ranking stability against weight changes. Changes were only seen in a few scenarios involving K5 and K6, as well as an increase in K3's weight, causing A1 to drop from 4th to 3rd place and A9 to go up from 3rd to 4th, while in the K6-0.05, K6-0.10, and K6-0.15 scenarios, A1 also dropped to 3rd place.

To evaluate the stability of the WASPAS ranking results, a sensitivity analysis was conducted using several scenarios involving changes in the criterion weights. These scenarios were designed to examine the effect of variations in the relative importance of the criteria on the ranking positions of the alternatives. The resulting ranking variations provide an indication of the sensitivity and robustness of the WASPAS method under different weighting conditions. The changes in alternative rankings obtained across the respective weight variation scenarios are presented in Figure 4.



Figure 4. WASPAS Method Ranking Sensitivity Analysis

The sensitivity analysis results for the WASPAS method in Figure 4 show that changing the criteria weights by +0.05 to +0.15 and -0.05 to -0.15 produces very stable rankings for most alternatives. In the initial condition, A3 remains in first place, followed by A6, A9, A1, A7, A4, A10, A2, A5, and A8. Ranking changes only occur for A1 and A9, especially when the weight of K5 or K3 is increased, causing A1 to move up from 4th to 3rd place and A9 to drop from 3rd to 4th place. A similar situation happens when the weight of K6 is decreased, making A1 move to 3rd place and A9 to 4th. Meanwhile, A2, A4, A5, A6, A7, A8, and A10 keep their positions the same across all scenarios.

To evaluate the stability of the ARAS ranking results, a sensitivity analysis was conducted using several scenarios involving changes in the criterion weights. These scenarios were designed to examine the impact of variations in the relative importance of the



criteria on the ranking positions of the alternatives. The ranking changes obtained from the respective weight variation scenarios are presented in Figure 5.



Figure 5. ARAS Method Ranking Sensitivity Analysis

The sensitivity analysis results of the ARAS method in Figure 5 show that changes in the criteria weights by +0.05 to +0.15 and -0.05 to -0.15 produce a very stable ranking for most alternatives. Initially, A3 ranked first, followed by A6, A9, A1, A7, A4, A10, A2, A5, and A8. Almost all scenarios maintained this order, with limited changes mainly for A1 and A9. When the weight of K5 was increased by 0.10 and 0.15, A1 rose from 4th to 3rd place, while A9 dropped to 4th. The same pattern occurred when the weight of K6 was decreased by 0.05, 0.10, and 0.15, causing A1 to be in 3rd place and A9 in 4th. Meanwhile, A2, A4, A5, A6, A7, A8, and A10 didn't change positions in any scenario.

To evaluate the stability of the COPRAS ranking results, a sensitivity analysis was performed using several scenarios involving changes in the criterion weights. The resulting ranking variations indicate the sensitivity of COPRAS to changes in the weighting structure and provide insight into the robustness of the obtained preferences. The changes in alternative rankings across the different weight variation scenarios are presented in Figure 6.



Figure 6. COPRAS Method Ranking Sensitivity Analysis



The sensitivity analysis results using the COPRAS method in figure 6 show that changing the criteria weights by +0.05 to +0.15 and -0.05 to -0.15 results in rankings that are very stable across almost all scenarios. Initially, A3 was ranked first, followed by A6, A9, A1, A7, A4, A10, A2, A5, and A8. Most weight changes didn't affect the positions of the alternatives, with changes only happening for A1 and A9. When the weight of K5 was increased by 0.10 and 0.15, A1 moved up from 4th to 3rd place, while A9 dropped to 4th. The same pattern happened when the weight of K6 was decreased by 0.05, 0.10, and 0.15, which caused A1 to be in 3rd place and A9 in 4th. Meanwhile, A2, A4, A5, A6, A7, A8, and A10 kept their positions in all scenarios.

To evaluate the stability of the MAIRCA ranking results, a sensitivity analysis was conducted through several scenarios involving changes in the criterion weights. These scenarios were designed to examine how variations in the relative importance of the criteria affect the ranking positions of the alternatives. The ranking changes obtained from each weight variation provide an indication of the robustness and sensitivity of the MAIRCA method to changes in the weighting structure. The resulting changes in the alternative rankings across the evaluated scenarios are presented in Figure 7, allowing the consistency of the MAIRCA ranking results to be assessed under different weighting conditions.

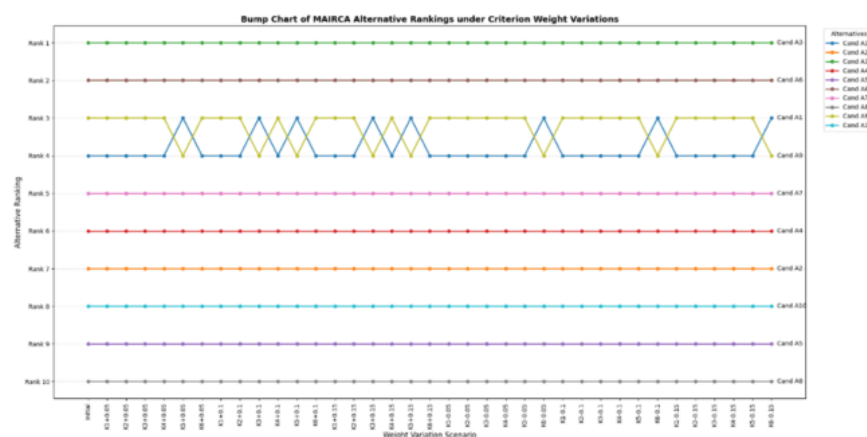


Figure 7. MAIRCA Method Ranking Sensitivity Analysis

The sensitivity analysis results on the MAIRCA method in Figure 7 show that changing the criteria weights by +0.05 to +0.15 and -0.05 to -0.15 produces a relatively stable ranking pattern. Initially, A3 was in first place, followed by A6, A9, A1, A7, A4, A10, A2, A5, and A8. Most scenarios didn't change the positions of the alternatives, with changes mainly happening to A1 and A9. Increasing the weight of K5 by 0.05, 0.10, and 0.15 caused A1 to move up from 4th to 3rd place, while A9 dropped to 4th. A similar situation occurred when the weight of K3 was increased by 0.10 and 0.15, and when the weight of K6 was decreased by 0.05, 0.10, and 0.15, which also led to a swap between A1 and A9. Meanwhile, A2, A4,



A5, A6, A7, A8, and A10 consistently maintained their positions across all scenarios.

5.3. Ranking Consistency Analysis

The comparison of alternative rankings is performed to examine the degree of agreement among the ranking results generated by the selected aggregation-based MCDM methods. Since each method applies a different computational and aggregation mechanism, differences in preference values may occur even when the same decision matrix and criterion weights are used. Therefore, comparing the resulting ranking orders provides an initial indication of whether the methods produce similar or different decision outcomes. In this study, the original ranking is used as the reference ranking, while the rankings generated by SAW, WASPAS, ARAS, COPRAS, and MAIRCA are evaluated against this reference. The degree of agreement is measured using Spearman's rank correlation coefficient, where a coefficient closer to 1 indicates stronger agreement between two ranking orders, while a lower coefficient indicates greater differences in the relative positions of alternatives[35]–[37]. This analysis is important because ranking similarity provides an initial assessment of the consistency of the methods before their behavior is further examined under changes in criterion weights and alternative performance values. The correlation results between the original ranking and each MCDM method are presented in Figure 3.

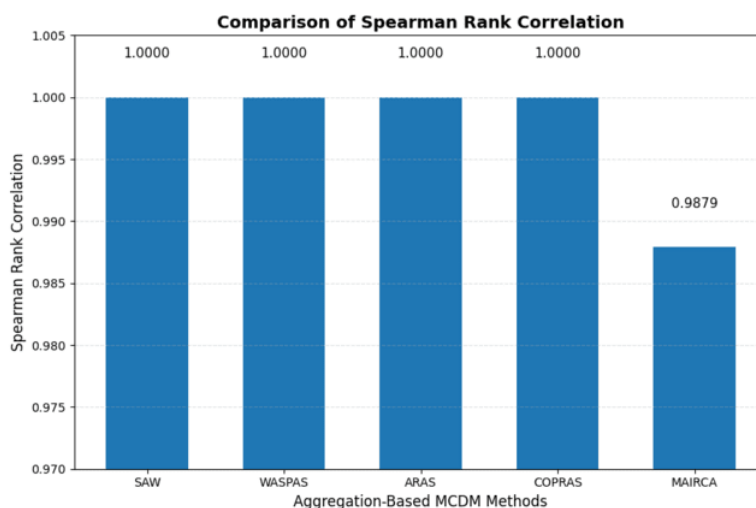


Figure 8. Spearman Rank Correlation of MCDM Methods

The Spearman Rank Correlation results presented in Figure 3 show that SAW, WASPAS, ARAS, and COPRAS achieve a perfect correlation coefficient of 1.0000 with the original ranking, indicating that these methods produce exactly the same alternative ordering as the reference ranking. Meanwhile, MAIRCA obtains a correlation coefficient of 0.9879, which still indicates a very strong ranking agreement, despite minor differences in the ordering of alternatives. Overall, the results in



Figure 3 demonstrate that all evaluated MCDM methods exhibit a high level of ranking consistency with the original ranking, with SAW, WASPAS, ARAS, and COPRAS achieving complete agreement.

5.4. Ranking Stability Analysis

Ranking Stability analysis is carried out as part of ranking sensitivity analysis to evaluate how changes in criteria weights affect the consistency of the order of alternatives produced by each MCDM method. In sensitivity analysis, weight changes are made through several scenarios of increasing and decreasing weights to observe possible shifts in the ranking of alternatives. Furthermore, the stability of rankings across all these scenarios is evaluated using the Ranking Stability Index (RSI), which represents the proportion of ranking positions that remain the same compared to all observed positions [38], [39]. A higher RSI value indicates that the method is increasingly able to maintain the ranking structure even when weights change, thus reflecting a better level of robustness and sensitivity to variations in criteria weights. This analysis is used to compare the stability of the SAW, WASPAS, ARAS, COPRAS, and MAIRCA methods based on changes in rankings that occur across all sensitivity scenarios.

Based on the sensitivity analysis results, the SAW method shows a very high level of ranking stability with a RSI of 95.56%. Out of a total of 360 ranking positions analyzed across 36 weight-change scenarios, 344 positions kept their initial ranking, while 16 positions changed. These changes only involved A1 and A9, which swapped positions from 4th and 3rd place to 3rd and 4th place. This shift occurred in several scenarios, namely when the weights of K3 and K5 were increased and when the weight of K6 was decreased at certain change levels. Meanwhile, the other alternatives maintained their positions across all scenarios. These results show that SAW has a relatively low sensitivity to weight changes, so the ranking structure it produces can be considered stable and robust against variations in weight preferences.

In the WASPAS method, the analysis results show a stability pattern that is almost the same as SAW, with an RSI value of 95.56%. Out of 360 observed ranking positions, 344 positions remained stable and only 16 positions changed. The ranking shifts were limited to A1 and A9, which swapped positions in eight different weight change scenarios. This happened when the weights of K3 or K5 increased to a certain level, and when the weight of K6 decreased. No changes occurred for the other alternatives, so most of the ranking structure stayed consistent across all tested weight variations. The high RSI value indicates that WASPAS is quite robust, since changes in criterion weights only have a limited effect on the order of alternatives. Thus, the WASPAS ranking results can be considered fairly reliable, even if there are changes in the importance levels of some criteria.

The ARAS method produced an RSI value of 97.22%, which is one of the highest stability levels in this analysis. Out of 360 ranking positions, 350 positions remained unchanged, while only 10 positions shifted. These changes only involved A1 and A9, swapping places between



3rd and 4th rank. Ranking changes appeared in scenarios where the weight of K5 was increased by 0.10 and 0.15, and when the weight of K6 was decreased by 0.05, 0.10, and 0.15. Meanwhile, changes in K3 did not affect the ARAS ranking. Alternative A3 also stayed in first place across all scenarios, showing that its position was relatively unaffected by the changes in weight. The high RSI shows that ARAS has very strong ranking stability and is relatively less sensitive to changes in criteria weights.

The sensitivity analysis results using the COPRAS method showed an RSI value of 97.22%, the same as ARAS and the highest value compared to other methods in this analysis. Out of the 360 ranking positions evaluated, 350 positions remained stable, and only 10 positions experienced changes. These changes only occurred with A1 and A9, which swapped rankings between the 3rd and 4th positions. This shift happened when the weight of K5 was increased by 0.10 and 0.15, and when the weight of K6 was decreased by 0.05, 0.10, and 0.15. On the other hand, changes to the weight of K3 or other change scenarios didn't affect the order of alternatives. A3 also stayed in the first position in all scenarios, showing the consistency of the top alternative despite weight changes.

The MAIRCA method shows a high level of stability with an RSI value of 95.56%. Out of a total of 360 ranking positions analyzed, 344 positions remained stable, while 16 positions experienced changes. Just like SAW and WASPAS, ranking changes in MAIRCA only happened with A1 and A9, which swapped between 3rd and 4th place. These shifts were found in scenarios where the weights of K3 and K5 were increased to a certain level, as well as in all tested scenarios of reducing K6's weight. Other alternatives didn't see any rank changes, including A3, which stayed in first place. This shows that MAIRCA is relatively insensitive to weight changes and can maintain its ranking structure in most scenarios. The RSI value of 95.56% thus indicates that MAIRCA has very good stability, even though its sensitivity is slightly higher compared to ARAS and COPRAS.

5.5. Ranking Robustness Analysis

Ranking Robustness Analysis is conducted to evaluate the ability of MCDM methods to maintain the consistency of ranking results when criterion weights are modified under different sensitivity scenarios. This analysis provides an indication of how effectively a method can withstand variations in criterion weights without producing substantial changes in the ranking structure. The robustness of a method can be assessed based on the consistency of alternative positions, the number of ranking changes, and the ability to preserve highly ranked alternatives under different weighting conditions. A smaller degree of ranking variation indicates a higher level of robustness. In this study, Ranking Robustness Analysis is employed to compare the resilience of the SAW, WASPAS, ARAS, COPRAS, and MAIRCA methods across 36 weight-variation scenarios and to identify the method that provides the most consistent decision outcomes.



Ranking robustness analysis is further evaluated using the Ranking Retention Rate, which measures the percentage of sensitivity scenarios in which the complete initial ranking order is preserved without any changes in alternative positions[40], [41]. A higher Ranking Robustness Analysis indicates greater robustness to variations in criterion weights because the resulting ranking remains consistent despite changes in the relative importance of the criteria. Therefore, RRR is used as an indicator for comparing the robustness of the SAW, WASPAS, ARAS, COPRAS, and MAIRCA methods across all weight-variation scenarios considered in this study.

For the SAW method, the Ranking Retention Rate is 77.78%, indicating that the complete initial ranking is maintained in 28 out of 36 sensitivity scenarios. The remaining eight scenarios result in changes in the ranking order, particularly under specific variations in the weights of K3, K5, and K6. These changes are limited to an exchange of positions between A1 and A9, while the other alternatives retain their original positions. This result indicates that SAW has a relatively high level of ranking stability under variations in criterion weights, although several specific weighting conditions can still lead to changes in the relative positions of alternatives.

The WASPAS method also achieves a Ranking Retention Rate of 77.78%, with the initial ranking being fully retained in 28 of the 36 scenarios. Eight scenarios produce changes in the ranking structure, particularly when the weights of K3 and K5 are increased or the weight of K6 is decreased. The ranking changes are limited to A1 and A9, whereas the remaining eight alternatives maintain their original positions throughout the sensitivity analysis. This finding indicates that WASPAS exhibits a sensitivity pattern similar to that of SAW and is able to preserve the initial ranking structure across most weight-variation scenarios. Therefore, WASPAS demonstrates a good level of ranking retention and robustness against changes in criterion weights.

For the ARAS method, the Ranking Retention Rate reaches 86.11%, representing the highest value among the evaluated methods, together with COPRAS. The initial ranking is completely maintained in 31 of the 36 sensitivity scenarios, while only five scenarios produce changes in the ranking order. These changes occur under specific increases in the weight of K5 and decreases in the weight of K6, with the ranking shift limited to an exchange between A1 and A9. The remaining alternatives preserve their original positions, while A3 consistently remains in first place across all scenarios. The high RRR value demonstrates that ARAS has a very strong ranking retention capability and is relatively insensitive to the weight variations applied in the sensitivity analysis.

The COPRAS method obtains a Ranking Retention Rate of 86.11%, indicating that the complete initial ranking is preserved in 31 of the 36 sensitivity scenarios. Only five scenarios result in changes in the ranking order, whereas the other 31 scenarios produce a ranking structure identical to the initial condition. The observed changes are limited to A1 and A9, particularly when the weight of K5 is increased or the weight of K6 is decreased under specific variation levels. A3 consistently



occupies the first position across all scenarios, while the other alternatives also maintain their respective positions. These results demonstrate that COPRAS has a very high level of ranking stability and retention, indicating strong robustness against variations in criterion weights.

The MAIRCA method achieves a Ranking Retention Rate of 77.78%, with the complete initial ranking maintained in 28 of the 36 sensitivity scenarios. Eight scenarios produce changes in the ranking order, particularly under increases in the weights of K3 and K5 and decreases in the weight of K6. Similar to SAW and WASPAS, the ranking changes are limited to A1 and A9, while the remaining alternatives retain their original positions. This finding indicates that although several sensitivity scenarios can affect the ranking order, the overall ranking structure generated by MAIRCA remains consistent across most weight-variation conditions. With an RRR of 77.78%, MAIRCA demonstrates a good level of ranking retention and robustness, although its stability is lower than that of ARAS and COPRAS.

5.6. Comparative Analysis

A comparative analysis is conducted to evaluate the performance and consistency of the investigated multi-criteria decision-making methods under the same decision-making conditions[42]. The comparison is intended to identify similarities and differences among the obtained results, particularly in terms of alternative scores, ranking positions, ranking consistency, and the ability of each method to produce reliable decision outcomes. In this experiment, the SAW, WASPAS, ARAS, COPRAS, and MAIRCA methods are applied to the same decision matrix and evaluated using identical alternatives, criteria, and initial criterion weights to ensure a fair comparison. The resulting rankings are subsequently compared to determine the extent to which different MCDM approaches agree or differ in selecting and prioritizing alternatives. Furthermore, the analysis considers the effect of criterion weight variation on the resulting rankings to examine the robustness of each method under different decision-making scenarios. Rank-based statistical measures are employed to quantify the degree of agreement between methods, while ranking changes across scenarios are analyzed to identify methods that are more consistent and less sensitive to variations in criterion importance. Through this comparative analysis, the study provides a comprehensive assessment of the relative behavior of the evaluated MCDM methods and establishes an empirical basis for determining which methods demonstrate stronger consistency, ranking robustness, and reliability for multi-criteria decision-making.

The comparative analysis evaluates the performance of the investigated MCDM methods from three complementary perspectives, namely ranking consistency, ranking stability, and ranking robustness. Ranking Consistency Analysis is performed using the Spearman rank correlation coefficient to measure the degree of agreement between the rankings produced by different MCDM methods, where a coefficient closer to one indicates stronger consistency among the resulting rankings. Ranking



Stability Analysis is conducted using the Ranking Stability Index to assess the extent to which the ranking positions are preserved when the criterion weights are varied across different experimental scenarios, with a higher index value indicating greater stability. Meanwhile, Ranking Robustness Analysis is employed to evaluate the ability of each method to maintain reliable and competitive alternative rankings under changes in decision-making conditions, particularly with respect to variations in criterion importance. The combination of these three analyses provides a comprehensive evaluation because consistency reflects the agreement among methods, stability reflects the preservation of rankings under-weight variations, and robustness reflects the resilience of the obtained ranking outcomes against changes in the decision environment. The results of the comparative analysis, including the correlation coefficients, stability indices, and robustness measures obtained for each method, are presented in Table 10.

Table 10. Comparative Analysis of Ranking Consistency, Stability, and Robustness

Method	Spearman Correlation	Ranking Stability Analysis	Ranking Retention Rate
SAW	1.0000	95.56%	77.78%
WASPAS	1.0000	95.56%	77.78%
ARAS	1.0000	97.22%	86.11%
COPRAS	1.0000	97.22%	86.11%
MAIRCA	0.9879	95.56%	77.78%

The overall results of the methods show a very high level of ranking consistency, with a Spearman Correlation value of 1.0000 for SAW, WASPAS, ARAS, and COPRAS, while MAIRCA got a value of 0.9879, indicating that all methods produce alternative rankings that are very consistent with the reference ranking. In terms of stability, ARAS and COPRAS showed the best performance with a Ranking Stability Analysis of 97.22%, while SAW, WASPAS, and MAIRCA each reached 95.56%. The Ranking Retention Rate results also showed that ARAS and COPRAS have the highest robustness at 86.11%, while SAW, WASPAS, and MAIRCA got 77.78%. These findings suggest that ARAS and COPRAS not only produce consistent rankings but are also better at maintaining alternative positions when there are variations in criterion weights. Meanwhile, although MAIRCA has a slightly lower correlation compared to the other methods, its stability level remains high, so overall all the methods show relatively robust ranking performance and can be relied on in various weight change scenarios.

5.7. Discussion

The experimental results demonstrate a strong overall agreement among the evaluated aggregation-based MCDM methods in the employee selection problem. Under the baseline decision-making configuration, SAW, WASPAS, ARAS, and COPRAS reproduce exactly the same ranking as the reference ranking, with Candidate 3 consistently identified as the most preferred alternative and Candidate 8 as the least preferred alternative.



This consistency indicates that the relative performance structure of the alternatives is sufficiently distinct to produce a similar ordering despite differences in the mathematical mechanisms used by the methods. MAIRCA also produces a highly comparable result, with only a pairwise exchange between Candidate 1 and Candidate 9. The difference is therefore not associated with a substantial change in the overall decision structure, but rather with the relative evaluation of two closely positioned alternatives. This finding is important because it demonstrates that the selected methods reach a highly similar decision outcome when applied to the same dataset and weighting configuration. However, the similarity observed under the baseline condition should not be interpreted as evidence that all methods have identical behavior. Their differences become more evident when the decision environment is systematically perturbed through changes in criterion weights.

The sensitivity analysis provides further evidence that the ranking structure is generally resistant to moderate and relatively substantial changes in criterion importance. Across the 36 weight-variation scenarios, the majority of alternatives preserve their original positions, while the observed changes are concentrated almost exclusively on Candidates 1 and 9. This pattern suggests that these two candidates have relatively close overall preference levels and are therefore more susceptible to changes in the importance assigned to particular criteria. In contrast, Candidate 3 consistently occupies the first position across the tested scenarios, indicating that its superior performance is sufficiently strong across the evaluation criteria to withstand variations in criterion weights. Similarly, the positions of Candidates 2, 4, 5, 6, 7, 8, and 10 remain highly consistent. From a decision-making perspective, this behavior indicates that the principal decision outcome is not easily overturned by changes in preference assumptions. The sensitivity analysis also reveals that increases in K3 and K5 and decreases in K6 are the most influential conditions for the observed exchange between Candidates 1 and 9. Therefore, these criteria can be considered important drivers of the relative positioning of candidates, particularly when alternatives have similar overall performance levels.

The ranking stability results reinforce the findings of the sensitivity analysis and reveal differences in the resilience of the methods. ARAS and COPRAS achieve the highest Ranking Stability Index of 97.22%, with 350 of the 360 evaluated ranking positions remaining unchanged. This indicates that only a very small proportion of ranking positions are affected by the tested weight variations. In comparison, SAW, WASPAS, and MAIRCA achieve an RSI of 95.56%, corresponding to 344 unchanged positions out of 360. Although the difference between 97.22% and 95.56% is relatively small, it provides evidence that ARAS and COPRAS are slightly less sensitive to changes in criterion weights in this particular decision environment. The difference can be associated with the internal aggregation structures of the methods. SAW relies on additive aggregation, while WASPAS combines weighted sum and weighted product mechanisms, and MAIRCA evaluates the deviation between



theoretical and real ratings. ARAS incorporates an optimal reference alternative through its utility degree, whereas COPRAS separately considers beneficial and non-beneficial criteria in determining relative significance. These mechanisms may influence how changes in one criterion propagate through the final preference value, explaining why ARAS and COPRAS preserve the ranking structure under a slightly broader range of weight configurations.

The robustness analysis provides a complementary perspective by examining whether the complete initial ranking is retained across each sensitivity scenario. ARAS and COPRAS again demonstrate the strongest performance, achieving a Ranking Retention Rate of 86.11%, meaning that the complete ranking order is preserved in 31 of the 36 scenarios. SAW, WASPAS, and MAIRCA retain the complete ranking in 28 scenarios, resulting in an RRR of 77.78%. The distinction between ranking stability and ranking retention is particularly relevant. A method may maintain most individual ranking positions while still experiencing a change in the overall ranking order because of a single pairwise exchange. This is precisely what occurs in the present experiment, where most ranking changes are limited to Candidates 1 and 9. Thus, ARAS and COPRAS demonstrate a stronger ability not only to preserve individual ranking positions but also to maintain the entire decision ordering under changes in criterion weights. Nevertheless, the robustness results should be interpreted within the context of the experimental design. The superiority of ARAS and COPRAS is demonstrated for the present dataset and the specified perturbation range, rather than establishing that these methods are universally superior in every MCDM application. Different datasets, criterion distributions, weighting schemes, or perturbation magnitudes may produce different relative performance patterns.

Taken together, the consistency, stability, and robustness analyses indicate that ARAS and COPRAS provide the strongest overall performance among the evaluated methods for this employee selection problem. Both methods achieve a perfect Spearman correlation of 1.0000 with the reference ranking, an RSI of 97.22%, and an RRR of 86.11%. SAW and WASPAS also demonstrate strong performance, with perfect ranking consistency and high stability, but their lower retention rates indicate slightly greater sensitivity to specific weight configurations. MAIRCA maintains a high level of stability but exhibits a minor baseline ranking deviation and a lower Spearman correlation of 0.9879. From the perspective of Decision Support Systems, these findings suggest that selecting an MCDM method should not be based solely on the ranking obtained under one initial configuration. A method that produces a reasonable baseline ranking but changes substantially when decision preferences are modified may provide less reliable support in practical environments. In this experiment, ARAS and COPRAS offer a more resilient ranking structure because they maintain both the relative ordering of individual alternatives and the complete ranking sequence across most sensitivity scenarios. The results therefore support their use when decision makers require ranking outcomes that remain reliable despite reasonable changes in the importance assigned to employee selection criteria.



VI. Conclusion

This study evaluated the consistency, stability, and robustness of five aggregation-based Multi-Criteria Decision Making (MCDM) methods, namely SAW, WASPAS, ARAS, COPRAS, and MAIRCA, in an employee selection decision-making problem. The experimental design employed the same decision-making dataset, criterion weights, and evaluation conditions for all methods, followed by systematic perturbations of criterion weights to examine the sensitivity of the resulting rankings. The baseline analysis showed that SAW, WASPAS, ARAS, and COPRAS generated rankings identical to the reference ranking, while MAIRCA produced only a minor exchange between Candidates 1 and 9. These findings demonstrate a high level of agreement among the investigated methods under the initial decision-making configuration.

The sensitivity analysis further showed that the ranking structures generated by all methods were generally resistant to changes in criterion weights. Most alternatives maintained their original positions across the 36 weight-variation scenarios, while ranking changes were primarily limited to Candidates 1 and 9. The Ranking Stability Index (RSI) confirmed this high level of stability, with ARAS and COPRAS achieving the highest value of 97.22%, compared with 95.56% for SAW, WASPAS, and MAIRCA. The Ranking Retention Rate (RRR) produced a similar pattern, where ARAS and COPRAS retained the complete initial ranking in 31 of 36 scenarios, corresponding to 86.11%, while SAW, WASPAS, and MAIRCA achieved 77.78%.

Overall, ARAS and COPRAS demonstrated the strongest performance across the three evaluation dimensions. Both methods achieved a Spearman correlation of 1.0000 with the reference ranking, an RSI of 97.22%, and an RRR of 86.11%. These results indicate that ARAS and COPRAS were able to preserve both individual ranking positions and the complete ranking structure under a broader range of criterion-weight variations than the other evaluated methods. Nevertheless, the results should be interpreted within the scope of the dataset and experimental scenarios considered in this study. The findings indicate that ARAS and COPRAS provide particularly reliable alternatives for employee selection when decision makers require ranking outcomes that remain consistent under changes in criterion importance.

The findings also highlight the importance of evaluating MCDM methods beyond their baseline ranking results. Although several methods can produce identical rankings under the initial conditions, their responses to changes in criterion weights may differ considerably. Therefore, ranking consistency alone is insufficient to determine the reliability of an MCDM method. Combining correlation-based consistency analysis with stability and robustness measures provides a more comprehensive evaluation of method performance. For Decision Support Systems, this approach can help decision makers identify methods that are not only capable of producing plausible rankings but also maintain



reliable decision outcomes when preferences or criterion priorities change.

CRediT Authorship Contribution Statement

Sumanto: Conceptualization, Methodology, Validation, Writing - Original Draft Preparation, Writing - Review & Editing. **Mochamad Wahyudi:** Conceptualization, Data curation, Formal analysis, Writing - Original Draft Preparation, Writing - Review & Editing. **Lise Pujiastuti:** Investigation, Software, Writing - Original Draft Preparation.

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.

Funding Statements

The author(s) received no financial support for the research, authorship, and/or publication of this article.

Data Availability

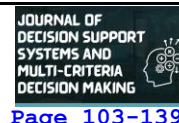
The data used to support the findings of this study are available from the corresponding author upon request.

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

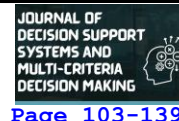
The authors used Generative AI to improve the writing clarity of this paper. They reviewed and edited the AI-assisted content and take full responsibility for the final publication.

References

- [1] T. Fujita, "The Hyperfuzzy VIKOR and Hyperfuzzy DEMATEL Methods for Multi-Criteria Decision-Making," *Spectr. Decis. Mak. Appl.*, vol. 3, no. 1 SE-Articles, pp. 292-315, Jan. 2026, doi: [10.31181/sdmap31202654](https://doi.org/10.31181/sdmap31202654).
- [2] I. Badi, M. B. Bouraima, and C. Kiprotich Kiptum, "Evaluating the Barriers to Logistics Outsourcing through a Fuzzy Multi-Criteria Decision-Making Model," *Spectr. Mech. Eng. Oper. Res.*, vol. 3, no. 1, pp. 65-78, Apr. 2026, doi: [10.31181/smeor31202653](https://doi.org/10.31181/smeor31202653).
- [3] J. M. L. G. Leite, L. G. Marujo, and E. F. Arruda, "Novel Time Aggregation-Based Algorithms for Markov Decision Processes," *IEEE Trans. Automat. Contr.*, vol. 70, no. 12, pp. 8353-8360, 2025, doi: [10.1109/TAC.2025.3583262](https://doi.org/10.1109/TAC.2025.3583262).
- [4] M. Behera and A. C. Panda, "A q-rung orthopair fuzzy Gaussian aggregation-based decision-making framework for sustainable solid waste management," *Appl. Soft Comput.*, vol. 199, p. 115381, 2026, doi: [10.1016/j.asoc.2026.115381](https://doi.org/10.1016/j.asoc.2026.115381).
- [5] R. S. Kanwal, S. M. Qurashi, N. Kausar, and F. T. Zahra, "A Hamacher Aggregation-Based Pythagorean Fuzzy Z-Number MCDM Framework for Sustainable Urban Green Space Planning," *J. Contemp. Decis. Sci.*, vol. 3,



- no. 1, pp. 1-15, 2026, doi: [10.67334/cds31202637](https://doi.org/10.67334/cds31202637).
- [6] T. Widodo, D. Darwis, and S. Nassor, "Hybrid Modification Preference Selection Index with Optimized Pairwise Ratio Analysis Method for Multi-Criteria Decision Making: An Integrated Weighting and Ranking Approach," *J. Decis. Support Syst. Multi-Criteria Decis. Mak.*, vol. 1, no. 1, pp. 84-102, 2026, doi: [10.67449/jodesma.v1i1.5](https://doi.org/10.67449/jodesma.v1i1.5).
- [7] G. Patel, S. Das, and R. Das, "Evaluation of optimal normalization techniques in multi-criteria decision-making to rank CMIP6 climate models," *Theor. Appl. Climatol.*, vol. 156, no. 7, p. 385, 2025, doi: [10.1007/s00704-025-05617-6](https://doi.org/10.1007/s00704-025-05617-6).
- [8] D. Štreimikienė, A. Bathaei, T. Baležentis, and J. Štreimikis, "Multi-criteria decision analysis of Circular Economy performance in the Baltic States: a comparative evaluation," *J. Bus. Econ. Manag.*, vol. 26, no. 5, pp. 1050-1070, 2025, doi: [10.3846/jbem.2025.24717](https://doi.org/10.3846/jbem.2025.24717).
- [9] S. Sintaro, P. Palupiningsih, and J. Wang, "Comparative Analysis of Objective Weighting Approaches in Multi-Criteria Decision Making Using WASPAS," *J. Decis. Support Syst. Multi-Criteria Decis. Mak.*, vol. 1, no. 1, pp. 16-42, 2026, doi: [10.67449/jodesma.v1i1.2](https://doi.org/10.67449/jodesma.v1i1.2).
- [10] D. A. Megawaty et al., "DAM Weighting: A New Approach to Determining Criteria Weighting in Multi-Criteria Decision Making," *Evergreen*, vol. 13, no. 2, pp. 788-800, 2026, doi: [10.5109/7432655](https://doi.org/10.5109/7432655).
- [11] M. Mesran et al., "Modification of Weighted Aggregated Sum Product Assessment Method to Improve Objective Weighting Accuracy in Multi-Criteria Decision Making," *Evergreen*, vol. 12, no. 2, pp. 1213-1225, Jun. 2025, doi: [10.5109/7363505](https://doi.org/10.5109/7363505).
- [12] P. K. Pramanik, S. Biswas, S. Pal, D. Marinković, and P. Choudhury, "A Comparative Analysis of Multi-Criteria Decision-Making Methods for Resource Selection in Mobile Crowd Computing," *Symmetry*, vol. 13, no. 9, p. 1713, 2021. doi: [10.3390/sym13091713](https://doi.org/10.3390/sym13091713).
- [13] E. Aytaç Adali, "A Novel MCDM Method: The Integrative Reference Point Approach," *Informatica*, vol. 36, no. 3, pp. 625-655, 2025, doi: [10.15388/25-INFOR594](https://doi.org/10.15388/25-INFOR594).
- [14] C.-N. Wang, N.-A.-T. Nguyen, and T.-T. Dang, "Sustainable Evaluation of Major Third-Party Logistics Providers: A Framework of an MCDM-Based Entropy Objective Weighting Method," *Mathematics*, vol. 11, no. 19, p. 4203, 2023. doi: [10.3390/math11194203](https://doi.org/10.3390/math11194203).
- [15] 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).
- [16] H. Gaspars-Wieloch and D. Gawroński, "How can one improve SAW and max-min multi-criteria rankings based on uncertain decision rules?," *Oper. Res. Decis.*, vol. 34, no. 1, pp. 131-148, 2024, doi: [10.37190/ord240107](https://doi.org/10.37190/ord240107).
- [17] 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).
- [18] A. Biswas, K. H. Gazi, P. Bhaduri, and S. P. Mondal, "Site Selection for Girls Hostel in a University Campus by MCDM based Strategy," *Spectr. Decis. Mak. Appl.*, vol. 2, no. 1, pp. 68-93, Jan. 2025, doi: [10.31181/sdmap21202511](https://doi.org/10.31181/sdmap21202511).
- [19] M. Hashemi-Tabatabaei, M. Amiri, and M. Keshavarz-Ghorabae, "An Expected Value-Based Symmetric-Asymmetric Polygonal Fuzzy Z-MCDM Framework for



- Sustainable-Smart Supplier Evaluation," *Information*, vol. 16, no. 3. 2025. doi: [10.3390/info16030187](https://doi.org/10.3390/info16030187).
- [20] S. L. Weng, H. L. Weng, F. L. Pei, S. H. Jaaman, and L. B. Swee, "Multi-Criteria Decision Making for the Selection of E-Commerce Platforms using AHP-TOPSIS Model," *J. Adv. Res. Appl. Sci. Eng. Technol.*, vol. 1, no. 2, pp. 120-136, 2024, doi: [10.37934/araset.60.1.120136](https://doi.org/10.37934/araset.60.1.120136).
- [21] S. Dündar, "Prioritizing the Regional Investors 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).
- [22] R. Raj, A. Singh, V. Kumar, T. De, and S. Singh, "Assessing the e-commerce last-mile logistics' hidden risk hurdles," *Clean. Logist. Supply Chain*, vol. 10, no. 1, p. 100131, 2024, doi: [10.1016/j.clscn.2023.100131](https://doi.org/10.1016/j.clscn.2023.100131).
- [23] H. Gholami et al., "Mapping flood risk using a workflow including deep learning and MCDM- Application to southern Iran," *Urban Clim.*, vol. 59, p. 102272, 2025, doi: [10.1016/j.uclim.2024.102272](https://doi.org/10.1016/j.uclim.2024.102272).
- [24] I. M. Hezam, A. R. Mishra, P. Rani, A. Saha, F. Smarandache, and D. Pamucar, "An integrated decision support framework using single-valued neutrosophic-MASWIP-COPRAS for sustainability assessment of bioenergy production technologies," *Expert Syst. Appl.*, vol. 211, p. 118674, Jan. 2023, doi: [10.1016/j.eswa.2022.118674](https://doi.org/10.1016/j.eswa.2022.118674).
- [25] I. Z. Mukhametzhanov and D. Pamucar, "Equivalence of MCDM Methods and Synthesis of Solution Based on Ratings Obtained in Different Models," *Decis. Mak. Appl. Manag. Eng.*, vol. 8, no. 2, pp. 1-20, Aug. 2025, doi: [10.31181/dmame8220251473](https://doi.org/10.31181/dmame8220251473).
- [26] 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).
- [27] S. Hadian, E. Shahiri Tabarestani, and Q. B. Pham, "Multi attributive ideal-real comparative analysis (MAIRCA) method for evaluating flood susceptibility in a temperate Mediterranean climate," *Hydrol. Sci. J.*, vol. 67, no. 3, pp. 401-418, Feb. 2022, doi: [10.1080/02626667.2022.2027949](https://doi.org/10.1080/02626667.2022.2027949).
- [28] Y. Rahmanto, J. Wang, S. Setiawansyah, A. Yudhistira, D. Darwis, and R. R. Suryono, "Optimizing Employee Admission Selection Using G2M Weighting and MOORA Method," *Paradig. - J. Komput. dan Inform.*, vol. 27, no. 1, pp. 1-10, Mar. 2025, doi: [10.31294/p.v27i1.8224](https://doi.org/10.31294/p.v27i1.8224).
- [29] 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).
- [30] M. Mohamed, S. Ayman, and J. Ye, "Assessment of Cybersecurity in Industry 4.0 using Delphi-Based Factor Relationships and Comprehensive Distance-Based Ranking Methods under Uncertainty," *Artif. Intell. Cybersecurity*, vol. 1, pp. 21-36, Jun. 2024, doi: [10.61356/j.aics.2024.1296](https://doi.org/10.61356/j.aics.2024.1296).
- [31] D. A. Megaraty, H. Sulistiani, Setiawansyah, A. Qurania, Y. Yadin, and R. Oktaviani, "Integration Method Based on the Removal Effects of Criteria Weighting and MOORA Method: Wi-Fi Router Selection Case Study," in *2024 International Conference on Informatics, Multimedia, Cyber and Information System (ICIMCIS)*, 2024, pp. 241-246. doi: [10.1109/ICIMCIS63449.2024.10956545](https://doi.org/10.1109/ICIMCIS63449.2024.10956545).
- [32] I. Boukrouh, F. Tayalati, and A. Azmani, "A Comprehensive Framework for Supplier Selection: Using Subjective, Objective, and Hybrid Multi-Criteria



- Decision-Making Techniques With Sensitivity Analysis," *IEEE Access*, vol. 12, pp. 145550-145569, 2024, doi: [10.1109/ACCESS.2024.3462348](https://doi.org/10.1109/ACCESS.2024.3462348).
- [33] M. R. Seikh and P. Chatterjee, "Determination of best renewable energy sources in India using SWARA-ARAS in confidence level based interval-valued Fermatean fuzzy environment," *Appl. Soft Comput.*, vol. 155, p. 111495, 2024, doi: [10.1016/j.asoc.2024.111495](https://doi.org/10.1016/j.asoc.2024.111495).
- [34] S. Dündar, "Performance evaluation of IPARD-II rural development programs with integrated DIBR-RAWEC methods," *Pamukkale Üniversitesi Mühendislik Bilim. Derg.*, vol. 31, no. 3, pp. 339-350, 2025.
- [35] T. Van Dua, D. Van Duc, N. C. Bao, and D. D. Trung, "Integration of objective weighting methods for criteria and MCDM methods: application in material selection," *EUREKA Phys. Eng.*, no. 2, pp. 131-148, Mar. 2024, doi: [10.21303/2461-4262.2024.003171](https://doi.org/10.21303/2461-4262.2024.003171).
- [36] Gokhan Basar and Oguzhan Der, "Multi-objective optimization of process parameters for laser cutting polyethylene using fuzzy AHP-based MCDM methods," *Proc. Inst. Mech. Eng. Part E J. Process Mech. Eng.*, vol. 239, no. 4, pp. 2295-2309, Feb. 2025, doi: [10.1177/09544089251319202](https://doi.org/10.1177/09544089251319202).
- [37] T. Van Dua and D. D. Trung, "MEPSI (Muttriss Enhanced Preference Selection Index): a novel method for ranking alternatives," *EUREKA Phys. Eng.*, vol. 2024, no. 6, pp. 169-178, Nov. 2024, doi: [10.21303/2461-4262.2024.003408](https://doi.org/10.21303/2461-4262.2024.003408).
- [38] J. Wang, S. Setiawansyah, and V. Saputra, "Hybrid Entropy and CRADIS Method Approach in Decision Support System for Selecting the Best Employees," *Build. Informatics, Technol. Sci.*, vol. 7, no. 4, pp. 2657-2668, Mar. 2026, doi: [10.47065/bits.v7i4.8985](https://doi.org/10.47065/bits.v7i4.8985).
- [39] T. Van Dua, "A Novel Approach for Criteria Weight Determination: A Case Study in Machine Ranking," *Eng. Technol. Appl. Sci. Res.*, vol. 16, no. 1, pp. 31333-31337, Feb. 2026, doi: [10.48084/etasr.15778](https://doi.org/10.48084/etasr.15778).
- [40] Hemn Othman Hama Ali and Karzan Mahdi Ghafour, "Enhancing Workforce Stability through Data-Driven Selection: An AHP-TOPSIS Approach for Healthcare HRM: Enhancing Workforce Stability through Data-Driven Selection," *Acad. J. Int. Univ. Erbil*, vol. 3, no. 2, pp. 973-994, Apr. 2026, doi: [10.63841/iue32673](https://doi.org/10.63841/iue32673).
- [41] M. J. Naeini, M. Shakerian, S. Yazdanirad, and S. M. Mousavi, "Application of the SWARA-TOPSIS method for prioritizing turnover intention factors among nurses: a case study in an Isfahan hospital, Iran," *BMC Nurs.*, vol. 24, no. 1, p. 1415, 2025, doi: [10.1186/s12912-025-04066-w](https://doi.org/10.1186/s12912-025-04066-w).
- [42] M. I. Takaendengan, "Performance Comparison of Multi-Criteria Decision-Making Methods in Decision Support Systems," *J. Comput. Data Sci.*, vol. 1, no. 1, pp. 88-108, 2026, doi: [10.67449/jcoda.v1i1.5](https://doi.org/10.67449/jcoda.v1i1.5).