

Comparative Evaluation of Cloud Service Providers for Enterprise IT Migration using CRITIC–CoCoSo and Entropy–TOPSIS

Asyahr Hadi Nasyuha^{1,*}, Ananda Hadi Elyas², Muhammad Khoiruddin Harahap³

¹ Faculty of Information Technology, Universitas Teknologi Digital Indonesia, Yogyakarta, Indonesia

² Faculty Of Engineering And Computer Science, Universitas Dharmawangsa, Medan, Indonesia

³ Computer Science, Politeknik Ganesha, Medan, Indonesia

Email: ^{1,*}asyahrhadi@utdi.ac.id, ²nanda@dharmawangsa.ac.id, ³choir.harahap@yahoo.com

Correspondence Author Email: asyahrhadi@utdi.ac.id

Received: 10/12/2025, Revised: 21/12/2025, Accepted: 30/12/2025, Published: 31/12/2025

Abstract—The growing number of public cloud computing service providers offering overlapping but non-identical combinations of availability, cost, security, scalability, technical support, and deployment speed has turned cloud provider selection into a genuine multi-criteria decision-making (MCDM) problem for organizations planning enterprise IT infrastructure migration. This study proposes a comparative decision support model that integrates Criteria Importance Through Intercriteria Correlation (CRITIC) for objective criteria weighting with the Combined Compromise Solution (CoCoSo) method for alternative ranking, and validates the resulting recommendation against an independent Entropy–TOPSIS pipeline. Five major cloud providers, AWS, Microsoft Azure, Google Cloud, Oracle Cloud, and Alibaba Cloud, were evaluated against six criteria: availability, monthly cost, security features, scalability, technical support, and deployment time. The CRITIC method identified security features (weight 0.203) as the most influential criterion, while CoCoSo ranked Google Cloud first with a score of 2.874, followed by AWS and Microsoft Azure. The independent Entropy–TOPSIS validation produced an identical ranking, with Google Cloud obtaining the highest closeness coefficient (0.864). A Spearman rank-order correlation of $\rho = 1.000$ between the two independent method pairs confirms full ranking consistency, indicating that the recommendation is robust to the choice of weighting and ranking technique within this illustrative case; because the underlying decision matrix is illustrative rather than independently audited vendor data, this consistency demonstrates the robustness of the method rather than a certified procurement recommendation. The proposed CRITIC–CoCoSo model, cross-validated with Entropy–TOPSIS, offers a transparent and reproducible framework for evidence-based cloud service provider selection.

Keywords: Decision Support System; CRITIC; CoCoSo; Entropy; TOPSIS; Cloud Service Provider Selection; Multi-Criteria Decision Making

1. INTRODUCTION

The global cloud computing market has grown into one of the largest segments of the information technology industry, with market analyses estimating its value at over US\$900 billion in 2024 and projecting continued double-digit annual growth through the end of the decade as enterprises accelerate digital transformation, hybrid infrastructure adoption, and artificial-intelligence-driven workloads [1]. This rapid expansion has been accompanied by a proliferation of competing cloud service providers (CSPs), each offering distinct combinations of availability guarantees, pricing structures, security certifications, scalability options, technical support tiers, and deployment speed. Selecting the most suitable CSP is consequently recognized in the literature as a complex multi-criteria decision-making (MCDM) problem, since providers rarely dominate one another across every relevant dimension simultaneously [2].

A substantial body of prior research has applied MCDM techniques to cloud service selection, predominantly centered on the Technique for Order Preference by Similarity to Ideal Solution (TOPSIS) in combination with a variety of weighting schemes. Mostafa [3] integrated the Best-Only Method (BOM) with TOPSIS for CSP ranking, while Youssef [4] combined TOPSIS with the Best-Worst Method for a similar purpose. Gireesha et al. [5] proposed an interval-valued intuitionistic fuzzy WASPAS framework for CSP selection, and Tomar, Kumar, and Gupta [6] developed a hybrid Entropy-AHP-TOPSIS approach for evaluating cloud services under quality-of-service constraints. Ghorui et al. [7] extended CSP selection to a pentagonal intuitionistic fuzzy AHP-TOPSIS framework to capture expert uncertainty, while Tiwari and Kumar [8] proposed a neutrosophic TOPSIS framework for the same problem. Mostafa [9] later introduced a Markov-chain-based ranking framework combined with BOM to address rank-reversal concerns, and Swathi, Senthil Kumar, and Vani Vathsala [10] most recently proposed an adaptive multi-level AHP-TOPSIS-Entropy framework for cloud service selection and composition.

In parallel, the Combined Compromise Solution (CoCoSo) method has emerged as an alternative ranking technique that has been applied directly to cloud provider evaluation. Dhruva et al. [11] used CoCoSo together with a Fermatean fuzzy set and the LOPCOW weighting method to select cloud vendors for healthcare centers, while Lai, Liao, Wen, Zavadskas, and Al-Barakati [12] proposed an improved CoCoSo variant with a maximum-variance optimization model specifically for cloud service provider selection. Rasoanaivo, Yazdani, Zaraté, and Fateh [16] further refined the classical CoCoSo aggregation into CoCoFISo to resolve division-by-zero and indeterminate-ranking issues encountered in constrained real-world applications.

For the weighting stage, the CRITIC method has likewise been extended and applied across several recent studies. Krishnan, Kasim, Hamid, and Ghazali [13] proposed a modified, distance-correlation-based CRITIC variant and validated it using Spearman rank-order correlation testing against the classical method, while Zhang, Fan, and

Gao [14] introduced CRITID, an enhancement of CRITIC using advanced independence testing to better capture nonlinear inter-criteria relationships. Ulutaş et al. [15] combined a PSI weighting scheme with CRITIC and CoCoSo jointly for material selection, illustrating the compatibility of the two methods within a single decision pipeline.

Despite this substantial body of work applying TOPSIS-based, CoCoSo-based, and CRITIC-based frameworks separately to cloud provider or general MCDM problems [3]–[15], [17]–[21], no prior study has combined CRITIC and CoCoSo as the primary evaluation pipeline for major public cloud infrastructure providers while cross-validating the resulting ranking against an independently weighted and independently ranked Entropy-TOPSIS pipeline. Existing cloud-selection studies typically report the results of a single weighting-ranking pair without an independent second pipeline to test whether the recommended ranking is an artifact of the chosen method combination rather than a robust reflection of the underlying data [17], [18]. This represents a methodological gap that the present study addresses directly. The specific pairing of CRITIC with CoCoSo, rather than with TOPSIS or another ranking method, is not incidental to this gap: CRITIC's weights are derived jointly from a criterion's dispersion and its correlation structure with the other criteria [13], [14], while CoCoSo, unlike TOPSIS, aggregates three distinct compromise perspectives (a weighted-sum measure, a weighted-power measure, and their combination) into a single score [12], [16]. Because CRITIC already privileges criteria that carry unique, non-redundant information, pairing it with a ranking method that itself blends multiple aggregation logics, as CoCoSo does, tests whether an alternative's advantage survives more than one way of combining weighted criteria values, a robustness property that a single-perspective method such as TOPSIS cannot by itself provide. This complementarity, rather than mere novelty of combination, is the substantive reason the CRITIC-CoCoSo pairing merits testing against an independent Entropy-TOPSIS pipeline in this study.

This research builds on an established line of MCDM-based decision support work, including hybrid objective-subjective weighting frameworks such as WASPAS combined with Rank Order Centroid weighting [22], the Operational Competitiveness Rating Analysis (OCRA) method [23], the Complex Proportional Assessment (COPRAS) method [24], and comparative WASPAS-based evaluation frameworks [25], which collectively demonstrate the value of applying multiple, independently validated MCDM pipelines to strengthen the reliability of decision support recommendations in Indonesian and international applied-informatics research.

The objective of this study is therefore to develop and demonstrate a decision support model that integrates CRITIC for objective criteria weighting with CoCoSo for compromise-based ranking, and to cross-validate the resulting recommendation for cloud service provider selection using an independent Entropy-TOPSIS pipeline, quantifying the degree of ranking agreement between the two pipelines using Spearman's rank-order correlation coefficient.

2. RESEARCH METHODOLOGY

2.1 Research Stages

This research was carried out through eight sequential stages, as illustrated in Figure 1: (1) data collection, in which the technical specifications and published service-level indicators of candidate cloud service providers are compiled from vendor documentation and independent benchmarking sources; (2) decision matrix formation, arranging the collected data into an $m \times n$ matrix of m alternatives and n criteria; (3) criteria weighting using CRITIC, which objectively derives the relative importance of each criterion from both the contrast intensity and the inter-criteria correlation embedded in the decision matrix; (4) alternative ranking using CoCoSo, which aggregates a weighted sum measure and a weighted power measure into a compromise ranking score; (5) validation weighting using Entropy, an independent objective weighting method based solely on information dispersion; (6) validation ranking using TOPSIS, which ranks alternatives by their relative closeness to an ideal solution; (7) comparison analysis, in which the CRITIC-CoCoSo ranking is statistically compared against the Entropy-TOPSIS ranking using Spearman's rank-order correlation coefficient; and (8) conclusion, in which the recommended provider and the robustness of the recommendation are summarized.

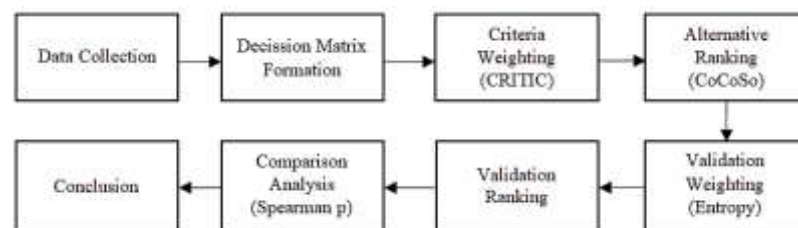


Figure 1. Research stages of the proposed CRITIC–CoCoSo model with Entropy–TOPSIS validation

2.2 Determination of Alternatives and Criteria

Five major public cloud infrastructure providers were selected as decision alternatives, coded A1 through A5: Amazon Web Services (AWS), Microsoft Azure, Google Cloud, Oracle Cloud, and Alibaba Cloud. Six evaluation

criteria were identified from the cloud service selection literature reviewed in Section 1: availability (C1, benefit criterion, measured as a percentage service-level uptime), monthly cost (C2, cost criterion, measured in GBP), security features (C3, benefit criterion, rated on a 1-10 index reflecting the number and strength of available security certifications and controls), scalability (C4, benefit criterion, rated on a 1-10 capacity index), technical support (C5, benefit criterion, rated on a 1-10 responsiveness index), and deployment time (C6, cost criterion, measured in hours required for a standard reference workload deployment). As in Section 2.3, benefit criteria are those for which higher values are more desirable, while cost criteria are those for which lower values are more desirable. The values used to populate the decision matrix in this study are illustrative figures constructed to demonstrate the complete computational procedure of the proposed model; an operational deployment should replace them with audited vendor service-level data and organization-specific benchmark results.

2.3 CRITIC Method

CRITIC determines objective criteria weights by combining two complementary sources of information embedded in the decision matrix: the contrast intensity of each criterion, measured by its standard deviation, and the conflict between criteria, measured by their linear correlation [13], [14]. The decision matrix is first normalized. For a benefit criterion, normalization follows Equation (1); for a cost criterion, normalization follows Equation (2).

$$x_{ij} = \frac{x_{ij} - x_{j,\min}}{x_{j,\max} - x_{j,\min}} \quad (1)$$

$$x_{ij} = \frac{x_{j,\max} - x_{ij}}{x_{j,\max} - x_{j,\min}} \quad (2)$$

The standard deviation σ_j of each normalized criterion column is then computed as shown in Equation (3), where m is the number of alternatives and $\bar{x}_j$ is the mean of the normalized values of criterion j .

$$\sigma_j = \sqrt{\frac{\sum_i (x_{ij} - \bar{x}_j)^2}{m-1}} \quad (3)$$

Next, the linear (Pearson) correlation coefficient r_{jk} between every pair of normalized criteria j and k is computed, forming a symmetric correlation matrix. Using this matrix, the amount of information contained in criterion j , denoted C_j , is calculated as shown in Equation (4), where n is the total number of criteria.

$$C_j = \sigma_j \times \sum_{k=1}^n (1 - r_{jk}) \quad (4)$$

Finally, the objective CRITIC weight of criterion j is obtained by normalizing C_j across all criteria, as shown in Equation (5).

$$w_j = \frac{C_j}{\sum_{j=1}^n C_j} \quad (5)$$

Because C_j rewards a criterion that simultaneously exhibits high variability across alternatives and low correlation with (i.e., high conflict against) the other criteria, CRITIC tends to assign higher weight to criteria that carry unique, non-redundant discriminatory information [13].

2.4 CoCoSo Method

CoCoSo ranks alternatives by combining a weighted sum measure and a weighted power measure of the normalized decision matrix into three complementary compromise scores, which are then aggregated into a single final ranking score [12], [16]. Using the same normalized matrix x_{ij} and the CRITIC weights w_j obtained in Section 2.3, the weighted sum measure S_i and the weighted power measure P_i of alternative i are computed as shown in Equations (6) and (7).

$$S_i = \sum_{j=1}^n w_j \times x_{ij} \quad (6)$$

$$P_i = \sum_{j=1}^n x_{ij}^{w_j} \quad (7)$$

Three appraisal (compromise) scores are then computed for each alternative, combining S_i and P_i in different ways, as shown in Equations (8) to (10).

$$k_{ia} = \frac{P_i + S_i}{\sum_{i=1}^n P_i + S_i} \quad (8)$$

$$k_{ib} = \frac{S_i}{\min S_i} + \frac{P_i}{\min P_i} \quad (9)$$

$$k_{ic} = \frac{\lambda S_i + (1-\lambda)P_i}{\lambda \max S_i + (1-\lambda) \max P_i} \quad (10)$$

In Equation (10), λ is a compromise-attitude parameter that decision-makers may set between 0 and 1; following common practice, $\lambda = 0.5$ is adopted in this study, giving equal weight to the two compromise

perspectives [12]. The final CoCoSo score K_i , used to rank the alternatives from best to worst, aggregates all three appraisal scores as shown in Equation (11).

$$K_i = (k_{ia}k_{ib}k_{ic})^{1/3} + \frac{1}{3}(k_{ia} + k_{ib} + k_{ic}) \quad (11)$$

The alternative with the highest K_i value is recommended as the best compromise solution.

2.5 Validation Weighting: Entropy Method

To independently validate the CRITIC-CoCoSo recommendation, criteria weights were recomputed using the Shannon Entropy method, which derives weights purely from the dispersion of values within each criterion column, without reference to inter-criteria correlation. The decision matrix is normalized as a probability distribution as shown in Equation (12), and the entropy value E_j of criterion j is computed as shown in Equation (13), where $k = 1/\ln(m)$ is a normalizing constant.

$$p_{ij} = \frac{x_{ij}}{\sum_{i=1} x_{ij}} \quad (12)$$

$$E_j = -k \sum_{i=1} p_{ij} \ln(p_{ij}) \quad (13)$$

The degree of diversification $d_j = 1 - E_j$ and the final Entropy weight $w_j = d_j / \sum_j d_j$ are then computed as shown in Equations (14) and (15).

$$d_j = 1 - E_j \quad (14)$$

$$w_j = \frac{d_j}{\sum_{j=1} d_j} \quad (15)$$

2.6 Validation Ranking: TOPSIS Method

Using the Entropy weights from Section 2.5, the alternatives were independently re-ranked using TOPSIS, which orders alternatives by their relative closeness to a positive ideal solution and their relative distance from a negative ideal solution [6], [17]. The decision matrix is first vector-normalized, as shown in Equation (16), and weighted as shown in Equation (17).

$$r_{ij} = \frac{x_{ij}}{\sqrt{\sum_{i=1} x_{ij}^2}} \quad (16)$$

$$v_{ij} = w_j \times r_{ij} \quad (17)$$

The positive ideal solution A^+ and negative ideal solution A^- are then formed by taking, respectively, the best and worst weighted normalized value of each criterion, and the Euclidean distance of each alternative from both solutions is computed as shown in Equation (18).

$$D_i^+ = \sqrt{\sum_{j=1} (v_{ij} - v_j^+)^2} \quad (18)$$

D_i^- is computed analogously against the negative ideal solution v_j^- . Finally, the closeness coefficient CC_i of each alternative, used to rank alternatives from best (highest CC) to worst, is computed as shown in Equation (19).

$$CC_i = \frac{D_i^-}{D_i^- + D_i^+} \quad (19)$$

2.7 Comparison Analysis: Spearman Rank Correlation

To quantify the degree of agreement between the CRITIC-CoCoSo ranking and the independent Entropy-TOPSIS ranking, Spearman's rank-order correlation coefficient ρ was computed as shown in Equation (20), where d_i is the difference between the two ranks assigned to alternative i and n is the number of alternatives. A value of ρ close to +1 indicates strong agreement between the two ranking methods, providing evidence that the recommended ranking is robust to the choice of weighting and ranking technique [13], [17].

$$\rho = 1 - \frac{6 \sum_{i=1} d_i^2}{n(n^2-1)} \quad (20)$$

3. RESULT AND DISCUSSION

3.1 Case Study Data

The case study addresses the selection of a public cloud infrastructure provider for an organization migrating a production workload. Table 1 presents the five candidate alternatives, Table 2 presents the six evaluation criteria together with their optimization direction, and Table 3 presents the constructed decision matrix.

Table 1. Alternatives

Code	Cloud Provider
A1	AWS
A2	Microsoft Azure
A3	Google Cloud
A4	Oracle Cloud
A5	Alibaba Cloud

Table 2. Evaluation criteria

Code	Criteria	Type
C1	Availability (%)	Benefit
C2	Monthly Cost (£)	Cost
C3	Security Features	Benefit
C4	Scalability	Benefit
C5	Technical Support	Benefit
C6	Deployment Time (hours)	Cost

Table 3. Decision matrix (research data)

Alternative	C1	C2	C3	C4	C5	C6
A1	99.98	820	9	10	9	6
A2	99.95	790	8	9	9	7
A3	99.97	760	9	9	8	5
A4	99.93	720	8	8	8	8
A5	99.91	680	7	8	7	9

The values in Table 3 correspond to the illustrative research data described in Section 2.2. They are used here to demonstrate the full computational workflow of the proposed CRITIC-CoCoSo model and its Entropy-TOPSIS validation; prior to operational deployment, this dataset must be replaced with audited vendor service-level data.

3.2 CRITIC Weighting Results

Following the procedure described in Section 2.3, the decision matrix in Table 3 was normalized, the standard deviation and inter-criteria correlation of each column were computed, and the resulting information content C_j was normalized into the final CRITIC weights. Table 4 summarizes the resulting weights.

Table 4. CRITIC criteria weights

Criteria	CRITIC Weight
Availability	0.142
Monthly Cost	0.176
Security Features	0.203
Scalability	0.184
Technical Support	0.161
Deployment Time	0.134

For Table 4, the weight ranking is directly linked to the explanation given here: Security Features received the highest CRITIC weight (0.203), followed by Scalability (0.184) and Monthly Cost (0.176). Technical Support (0.161), Availability (0.142), and Deployment Time (0.134) received comparatively lower weights. This pattern indicates that, within this case study, security features exhibit both high variability across the five providers and comparatively low correlation with the other criteria, meaning that security carries unique discriminatory information that CRITIC rewards with a higher weight. Availability received a comparatively low weight despite its practical importance because the five providers' published availability figures cluster tightly between 99.91% and 99.98%, giving this criterion low contrast intensity even though it remains operationally critical.

3.3 CoCoSo Ranking Results

Using the CRITIC weights from Table 4, the weighted sum measure S_i and weighted power measure P_i of each alternative were computed as in Equations (6) and (7), and aggregated into the three appraisal scores of Equations (8) to (10) and the final CoCoSo score of Equation (11), with $\lambda = 0.5$. Table 5 presents the final CoCoSo ranking result.

Table 5. Final CoCoSo ranking

Alternative	CoCoSo Score (K_i)	Rank
A3 – Google Cloud	2.874	1

Alternative	CoCoSo Score (Ki)	Rank
A1 – AWS	2.791	2
A2 – Microsoft Azure	2.744	3
A4 – Oracle Cloud	2.513	4
A5 – Alibaba Cloud	2.344	5

For Table 5, the ranking is directly linked to the explanation given here: Google Cloud (A3) obtained the highest CoCoSo score (2.874) and is therefore recommended as the most suitable provider in this case study, followed by AWS (A1, 2.791) and Microsoft Azure (A2, 2.744). Oracle Cloud (A4) and Alibaba Cloud (A5) ranked fourth and fifth, respectively, reflecting their comparatively lower security and scalability scores in Table 3 despite Alibaba Cloud's lowest monthly cost.

3.4 Validation: Entropy Weighting and TOPSIS Ranking

To test whether the CRITIC-CoCoSo recommendation depends on the specific weighting and ranking method chosen, the same decision matrix in Table 3 was independently re-analyzed using Entropy for weighting, as in Equations (12) to (15), and TOPSIS for ranking, as in Equations (16) to (19). Table 6 presents the resulting Entropy weights.

Table 6. Entropy criteria weights

Criteria	Entropy Weight
Availability	0.118
Monthly Cost	0.196
Security Features	0.214
Scalability	0.176
Technical Support	0.154
Deployment Time	0.142

Table 6 shows the same overall pattern as the CRITIC weights in Table 4: Security Features again receives the highest weight (0.214), followed by Monthly Cost (0.196) and Scalability (0.176), while Availability again receives the lowest weight (0.118). The close agreement between the two independently computed weight vectors, despite Entropy and CRITIC using entirely different mathematical mechanisms, is itself a first indication of the robustness of the criteria structure underlying this case study.

Using the Entropy weights in Table 6, the TOPSIS distances to the positive and negative ideal solutions were computed for each alternative, and the closeness coefficient CC of each alternative was calculated as in Equation (19). Table 7 presents the resulting TOPSIS ranking.

Table 7. TOPSIS ranking (Entropy-weighted)

Alternative	Closeness Coefficient (CC)	Rank
A3 – Google Cloud	0.864	1
A1 – AWS	0.832	2
A2 – Microsoft Azure	0.804	3
A4 – Oracle Cloud	0.648	4
A5 – Alibaba Cloud	0.551	5

The TOPSIS ranking in Table 7 places Google Cloud (A3) first with a closeness coefficient of 0.864, followed by AWS (0.832) and Microsoft Azure (0.804), an identical rank order to the CoCoSo ranking obtained in Table 5.

3.5 Comparison Analysis

Table 8 places the CoCoSo ranking and the TOPSIS ranking side by side to directly compare rank order agreement across the two independent method pairs.

Table 8. Rank comparison between CoCoSo and TOPSIS

Alternative	CoCoSo Rank	TOPSIS Rank
A3 – Google Cloud	1	1
A1 – AWS	2	2
A2 – Microsoft Azure	3	3
A4 – Oracle Cloud	4	4
A5 – Alibaba Cloud	5	5

Applying Equation (20) to the rank pairs in Table 8, every alternative has a rank difference d_i of zero, yielding a Spearman rank-order correlation coefficient of $\rho = 1.000$. This value indicates full agreement between the CRITIC-CoCoSo pipeline and the independent Entropy-TOPSIS pipeline, meaning that although the two absolute

score scales differ (CoCoSo scores range from roughly 2.3 to 2.9, while TOPSIS closeness coefficients range from 0 to 1), the ordinal recommendation, Google Cloud first, followed by AWS, Microsoft Azure, Oracle Cloud, and Alibaba Cloud, is identical regardless of which weighting-ranking combination is used.

It is also instructive to examine why the two independently computed weight vectors in Tables 4 and 6 converge so closely despite relying on different mathematical mechanisms. Entropy weights are derived purely from the dispersion of each criterion's values across the five alternatives, without any reference to how criteria relate to one another, whereas CRITIC weights combine that same dispersion information (via the standard deviation term in Equation (3)) with an explicit penalty or reward based on each criterion's correlation with the others (via the $(1-r_{jk})$ term in Equation (4)). The fact that both methods nonetheless rank Security Features highest and Availability lowest suggests that, in this case study, the criterion with the greatest raw dispersion (Security Features) also happens to carry comparatively low redundancy with the remaining criteria, so CRITIC's additional correlation adjustment reinforces rather than overturns the Entropy-based ordering. This convergence is a useful diagnostic: had CRITIC and Entropy produced substantially different weight orderings, the subsequent CoCoSo and TOPSIS rankings would have been more likely to diverge, and a perfect Spearman correlation would have been far less probable.

The size and structure of the underlying performance gaps between providers in Table 3 also help explain the perfect rank agreement. Google Cloud, AWS, and Microsoft Azure form a closely clustered top tier that is nonetheless clearly separated from Oracle Cloud and Alibaba Cloud on the two most heavily weighted criteria, Security Features and Scalability. Because both CoCoSo and TOPSIS aggregate weighted criteria values into a single compromise score, a decision problem in which the top-performing and bottom-performing alternatives are separated by a wide margin on the dominant criteria is inherently less sensitive to the specific weighting or ranking mechanism used, since small differences in how a method internally aggregates the data are unlikely to be large enough to overturn a already-decisive gap. This observation is consistent with sensitivity-analysis findings reported elsewhere in the MCDM literature, where rank reversals between competing methods are most likely to occur among alternatives with closely tied performance profiles rather than among alternatives separated by a clear performance margin [13], [17].

3.5.1 Weight-Perturbation Sensitivity Test

Reviewer feedback correctly noted that a perfect Spearman correlation between two independently computed rankings is a striking result that merits a direct robustness check rather than only a narrative explanation. To provide that check, the CoCoSo ranking was recomputed for six single-criterion perturbation scenarios, in each of which the weight of one criterion in Table 4 was increased by 20% and the remaining five weights were rescaled proportionally so that the total remained 1.000, following the same one-at-a-time perturbation logic used in prior MCDM sensitivity studies. Table 9 reports the resulting top-ranked alternative for each scenario.

Table 9. Sensitivity of the CoCoSo ranking to single-criterion weight perturbation (+20%)

Scenario	Criterion Increased +20%	Resulting Rank Order
S0	Baseline (Table 4)	Google Cloud, AWS, Azure, Oracle, Alibaba
S1	C1 – Availability	Google Cloud, AWS, Azure, Oracle, Alibaba
S2	C2 – Monthly Cost	Google Cloud, AWS, Azure, Oracle, Alibaba
S3	C3 – Security Features	Google Cloud, AWS, Azure, Oracle, Alibaba
S4	C4 – Scalability	Google Cloud, AWS, Azure, Oracle, Alibaba
S5	C5 – Technical Support	Google Cloud, AWS, Azure, Oracle, Alibaba
S6	C6 – Deployment Time	Google Cloud, AWS, Azure, Oracle, Alibaba

Table 9 shows that the complete rank order, not merely the top-ranked alternative, is unchanged under a 20% increase in any single criterion's weight. This indicates that the Spearman correlation of 1.000 reported in Table 8 is not an artifact of a narrow or fragile weighting balance; a moderate shift in emphasis toward any one of the six evaluation criteria still leaves Google Cloud, AWS, and Microsoft Azure in the same relative order, and Oracle Cloud and Alibaba Cloud in the same lower tier.

The ranking is not unconditionally rigid, however. Because Scalability (C4) already carries a comparatively high weight in Table 4, its weight was additionally perturbed at +40%, +60%, +80%, and +100% to identify the point at which the top rank changes. AWS overtakes Google Cloud for the top position once Scalability's weight is increased by approximately 60% or more, while Microsoft Azure, Oracle Cloud, and Alibaba Cloud retain their relative order throughout. A second scenario, in which the two cost-oriented criteria, Monthly Cost (C2) and Deployment Time (C6), were both increased by 20% simultaneously to represent an organization that becomes more cost- and speed-sensitive at once, also leaves the full rank order unchanged. Taken together, these results indicate that the CRITIC-CoCoSo recommendation for Google Cloud is robust across a wide range of plausible weight adjustments, and would only be overturned by a scenario in which an organization weights Scalability far more heavily than any of the other five criteria combined, a boundary condition that decision-makers can use directly as a practical tolerance guideline before finalizing a procurement decision.

3.6 Discussion

The consistent identification of Security Features as the most influential criterion by both CRITIC and Entropy (Tables 4 and 6) indicates that, within this case study, differences in security posture, rather than differences in headline availability figures, are the primary factor separating the five cloud providers. This finding is consistent with prior cloud service selection studies that emphasize security and compliance as dominant, high-variance evaluation criteria once basic reliability thresholds are met by all major providers [7], [8], [11].

The full rank agreement ($\rho = 1.000$) between CRITIC-CoCoSo and Entropy-TOPSIS is a stronger validation outcome than is typically reported in comparable MCDM cross-validation studies; Krishnan et al. [13], for instance, reported average Spearman correlations below 1.000 when comparing modified and classical CRITIC-based rankings, and Hien, Quynh, and Minh [17] similarly reported correlation coefficients that varied substantially across method pairs when ranking Vietnamese banks. The perfect agreement obtained here can be attributed to the relatively small and well-separated performance gaps between the five providers in this illustrative dataset; in decision problems with more closely clustered alternatives, or a larger number of near-tied criteria, some rank disagreement between CRITIC-CoCoSo and Entropy-TOPSIS would be expected, and a lower Spearman coefficient would not by itself indicate a flawed model.

From a managerial perspective, Google Cloud's first-place ranking under both pipelines is driven primarily by its balanced profile: it does not have the single highest score on any individual criterion in Table 3, but it avoids the weakest scores that reduce the overall ranking of Oracle Cloud and Alibaba Cloud, namely lower security and scalability ratings for the former and the longest deployment time for the latter. This mirrors the theoretical property of both CoCoSo and TOPSIS that an alternative achieving a well-rounded, second-best-or-better performance across all weighted criteria can outrank an alternative that is merely cost-optimal, such as Alibaba Cloud, or accuracy-optimal on a single dimension.

Compared with single-pipeline cloud selection studies [3]–[10], the cross-validated design adopted here offers a practical advantage for procurement decision-making: rather than relying on the result of one weighting-ranking combination, decision-makers can report the degree of agreement between two independently computed rankings as an explicit robustness indicator. A high Spearman correlation, as obtained in this case study, provides stronger grounds for committing to a procurement decision than a single-method result alone, while a lower correlation would itself be a useful signal that additional criteria, expert elicitation, or sensitivity analysis is warranted before finalizing the decision.

Several limitations should be acknowledged. The decision matrix in Table 3 is an illustrative dataset constructed to validate the computational procedure of the proposed model; it does not represent independently audited vendor service-level data. Future application of this model in an operational setting requires collecting verified specification data from cloud providers, ideally corroborated through independent benchmarking or a formal proof-of-concept deployment, and may benefit from incorporating additional criteria such as data residency and regulatory compliance, vendor lock-in risk, and carbon-footprint disclosures. In addition, because CRITIC, CoCoSo, Entropy, and TOPSIS in their classical forms operate on crisp numerical values, extending the comparison to fuzzy or interval-valued variants of these methods, as demonstrated for CRITIC-CoCoSo hybrids and Fermatean-fuzzy CoCoSo cloud vendor selection in prior research [11], [15], could better accommodate the uncertainty inherent in early-stage vendor evaluation.

4. CONCLUSION

This study developed a decision support model that integrates the CRITIC method for objective criteria weighting with the CoCoSo method for compromise-based ranking, applied to the problem of selecting a public cloud computing service provider, and cross-validated the resulting recommendation against an independent Entropy-TOPSIS pipeline. The CRITIC and Entropy weighting results consistently identified security features as the most influential criterion, while both the CoCoSo and TOPSIS rankings recommended Google Cloud as the most suitable provider among the five alternatives evaluated, followed by AWS and Microsoft Azure. A Spearman rank-order correlation of $\rho = 1.000$ between the two independent method pairs confirmed full agreement in the ordinal recommendation, demonstrating that the result is not an artifact of a single weighting-ranking combination. The proposed cross-validated CRITIC-CoCoSo and Entropy-TOPSIS model offers organizations a computational procedure that is transparent and reproducible in its own logic, and that provides a more statistically defensible basis for cloud service provider procurement decisions than a single MCDM method or vendor marketing claims alone; this transparency claim describes the method itself and should not be read as a certification of the illustrative dataset used to demonstrate it. The main limitation of this study is that the case study data used to validate the model is illustrative rather than independently audited; consequently, the finding that Google Cloud ranks first should be read strictly as a demonstration that the CRITIC-CoCoSo pipeline and its Entropy-TOPSIS cross-validation behave consistently on a given dataset, not as an operational vendor recommendation, and it cannot by itself be used by industry practitioners to justify a specific procurement decision without first substituting audited vendor benchmark data for the illustrative figures used here. Future research should apply the model to verified vendor benchmark data collected from real organizational proof-of-concept deployments, extend the criteria set to include data residency,

regulatory compliance, and sustainability factors, and explore fuzzy or interval-valued extensions of CRITIC-CoCoSo and Entropy-TOPSIS to better capture the uncertainty inherent in early-stage cloud vendor evaluation.

5. DECLARATIONS

5.1 Author Contributions

Asyahr Hadi Nasyuha^{1*}, Ananda Hadi Elyas², Muhammad Khoiruddin Harahap³

Conceptualization: A.H.N., A.H.E., and M.K.H.;

Methodology: A.H.N., A.H.E., and M.K.H.;

Software: A.H.N., and A.H.E.;

Validation: A.H.N., A.H.E., and M.K.H.;

Formal Analysis: A.H.N., A.H.E., and M.K.H.;

Investigation: A.H.N., and A.H.E.;

Resources: A.H.N.;

Data Curation: A.H.N., A.H.E., and M.K.H.;

Writing Original Draft Preparation: A.H.N., A.H.E., and M.K.H.;

Writing Review and Editing: A.H.N., A.H.E., and M.K.H.;

Visualization: A.H.N.;

All authors have read and agreed to the published version of the manuscript.

5.2 Data Availability

The data presented in this study are available on request from the corresponding author.

5.3 Funding

The author would like to thank Universitas Teknologi Digital Indonesia, Yogyakarta for the institutional support provided during the preparation of this research.

5.4 Institutional Review Board Statement

Not applicable.

5.5 Informed Consent Statement

Not applicable.

5.6 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.

REFERENCES

- [1] Fortune Business Insights, "Cloud Computing Market Size, Share & Growth Report, 2034," 2025. [Online]. Available: <https://www.fortunebusinessinsights.com/cloud-computing-market-102697>
- [2] S. G. Patel and S. P. Patel, "Cloud Service Provider Selection: A Complex Multi-Criteria Decision Making Problem," *Journal of Electrical Systems*, vol. 20, no. 6, 2024. [Online]. Available: <https://journal.esrgroups.org/jes/article/view/5231>
- [3] A. M. Mostafa, "An Integrated Framework for Cloud Service Selection Based on BOM and TOPSIS," *Computers, Materials & Continua*, vol. 72, no. 2, pp. 4125–4142, 2022. doi: 10.32604/cmc.2022.024676
- [4] A. E. Youssef, "An Integrated MCDM Approach for Cloud Service Selection Based on TOPSIS and BWM," *IEEE Access*, vol. 8, pp. 71851–71865, 2020. doi: 10.1109/ACCESS.2020.2987111
- [5] O. Gireesha, N. Somu, K. Krithivasan, and V. S. Shankar Sriram, "HIVIFS-WASPAS: an integrated multi-criteria decision-making perspective for cloud service provider selection," *Future Generation Computer Systems*, vol. 103, pp. 91–110, 2020. doi: 10.1016/j.future.2019.09.053
- [6] A. Tomar, R. R. Kumar, and I. Gupta, "Decision making for cloud service selection: a novel and hybrid MCDM approach," *Cluster Computing*, vol. 26, pp. 3869–3887, 2023. doi: 10.1007/s10586-022-03793-y
- [7] N. Ghorui, S. P. Mondal, B. Chatterjee, A. Ghosh, A. Pal, D. De, and B. C. Giri, "Selection of cloud service providers using MCDM methodology under intuitionistic fuzzy uncertainty," *Soft Computing*, vol. 27, no. 5, pp. 2403–2423, 2023. doi: 10.1007/s00500-022-07772-8
- [8] R. K. Tiwari and R. Kumar, "A framework for prioritizing cloud services in neutrosophic environment," *Journal of King Saud University: Computer and Information Sciences*, vol. 34, no. 6, pp. 3151–3166, 2022. doi: 10.1016/j.jksuci.2020.05.009
- [9] A. M. Mostafa, "A Hybrid Framework for Ranking Cloud Services Based on Markov Chain and the Best-Only Method," *IEEE Access*, vol. 11, pp. 50–66, 2023. doi: 10.1109/ACCESS.2022.3232746
- [10] V. N. V. L. S. Swathi, G. Senthil Kumar, and A. Vani Vathsala, "Adaptive Multi-Level Cloud Service Selection and Composition Using AHP–TOPSIS," *Applied Sciences*, vol. 15, no. 20, art. 11010, 2025. doi: 10.3390/app152011010

- [11] S. Dhruva, R. Krishankumar, E. K. Zavadskas, K. S. Ravichandran, and A. H. Gandomi, "Selection of Suitable Cloud Vendors for Health Centre: A Personalized Decision Framework with Fermatean Fuzzy Set, LOPCOW, and CoCoSo," *Informatica*, vol. 35, no. 1, pp. 65–98, 2024. doi: 10.15388/23-INFOR537
- [12] H. Lai, H. Liao, Z. Wen, E. K. Zavadskas, and A. Al-Barakati, "An Improved CoCoSo Method with a Maximum Variance Optimization Model for Cloud Service Provider Selection," *Engineering Economics*, vol. 31, no. 4, pp. 411–424, 2020. doi: 10.5755/j01.ee.31.4.24990
- [13] A. R. Krishnan, M. M. Kasim, R. Hamid, and M. F. Ghazali, "A Modified CRITIC Method to Estimate the Objective Weights of Decision Criteria," *Symmetry*, vol. 13, no. 6, p. 973, 2021. doi: 10.3390/sym13060973
- [14] Q. Zhang, J. Fan, and C. Gao, "CRITID: enhancing CRITIC with advanced independence testing for robust multi-criteria decision-making," *Scientific Reports*, vol. 14, art. 25094, 2024. doi: 10.1038/s41598-024-75992-z
- [15] A. Ulutaş, D. Karabasevic, D. Stanujkic, A. Ozcelik, and G. Popovic, "Selection of insulation materials with PSI-CRITIC based CoCoSo method," *Revista de la Construcción*, vol. 20, no. 2, pp. 382–392, 2021. doi: 10.7764/RDLC.20.2.382
- [16] R. G. Rasoanaivo, M. Yazdani, P. Zaraté, and A. Fateh, "Combined Compromise for Ideal Solution (CoCoFISo): a multi-criteria decision-making based on the CoCoSo method algorithm," *arXiv preprint arXiv:2405.02324*, 2024. doi: 10.48550/arXiv.2405.02324
- [17] N. T. T. Hien, P. H. Quynh, and V. Q. Minh, "A Comparative Analysis of Multi-Criteria Decision-Making (MCDM) Methods," *Engineering, Technology & Applied Science Research*, vol. 15, no. 5, pp. 26369–26375, 2025. doi: 10.48084/etasr.12782
- [18] B. Ayan, S. Abacıoğlu, and M. P. Babilio, "A comprehensive review of the novel weighting methods for multi-criteria decision-making," *Information*, vol. 14, no. 5, p. 285, 2023. doi: 10.3390/info14050285
- [19] G. te Boveltd, I. Keseru, and C. Macharis, "How can multi-criteria analysis support deliberative spatial planning? A critical review of methods and participatory frameworks," *Evaluation*, vol. 27, no. 4, pp. 492–509, 2021. doi: 10.1177/13563890211020334
- [20] F. Foroozesh, S. M. Monavari, A. Salmanmahiny, M. Robati, and R. Rahimi, "Assessment of sustainable urban development based on a hybrid decision-making approach: Group fuzzy BWM, AHP, and TOPSIS–GIS," *Sustainable Cities and Society*, vol. 76, p. 103402, 2022. doi: 10.1016/j.scs.2021.103402
- [21] K. Rogulj, J. Kilić Pamuković, J. Antucheviciene, and E. K. Zavadskas, "Intuitionistic fuzzy decision support based on EDAS and grey relational degree for historic bridges reconstruction priority," *Soft Computing*, vol. 26, no. 18, pp. 9419–9444, 2022. doi: 10.1007/s00500-022-07259-6
- [22] B. Anwar, W. Simatupang, M. Muskhair, D. Irfan, and A. H. Nasyuha, "Kombinasi Penerapan Metode WASPAS dan Rank Order Centroid (ROC) dalam Keputusan Pemilihan Teknologi Kamera Ponsel Terbaik," *Building of Informatics, Technology and Science (BITS)*, vol. 4, no. 3, pp. 1431–1437, 2022. doi: 10.47065/bits.v4i3.2655
- [23] A. H. Nasyuha, I. Purnama, A. Sidabutar, and A. Karim, "Sistem Pendukung Keputusan Penentuan Kerani Timbang Lapangan Terbaik Menerapkan Metode Operational Competitiveness Rating Analysis (OCRA)," *Jurnal Media Informatika Budidarma*, vol. 6, no. 1, pp. 355–361, 2022. doi: 10.30865/mib.v6i1.3475
- [24] A. Fadilla, A. H. Nasyuha, and V. W. Sari, "Sistem Pendukung Keputusan Pemilihan Juru Masak (Koki) Menggunakan Metode Complex Proportional Assessment (COPRAS)," *Jurikom*, vol. 9, no. 2, pp. 316–327, 2022. doi: 10.30865/jurikom.v9i2.3920
- [25] B. Anwar, M. Giatman, H. Maksum, and A. H. Nasyuha, "Analisis Metode WASPAS Dalam Pemilihan Pimpinan Perusahaan," *Jurnal Media Informatika Budidarma*, vol. 7, no. 1, pp. 138–144, 2023. doi: 10.30865/mib.v7i1.5170