

Beyond Vendor Hype: An Objective IDOCRIW-MARCOS Model for Selecting AI-Based Workforce Mental Health Platforms

Asyahri 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*}asyahrihadi@gmail.com, ²nanda@dharmawangsa.ac.id, ³choir.harahap@yahoo.com

Correspondence Author Email: asyahrihadi@gmail.com

Received: 10/05/2026, Revised: 06/06/2026, Accepted: 20/06/2026, Published: 30/06/2026

Abstract—The growing adoption of artificial intelligence (AI) in workforce mental health monitoring has produced a wide range of competing platforms that differ in prediction accuracy, response time, scalability, implementation cost, and user satisfaction, so that HR and occupational-health decision-makers currently choose among them largely on the basis of vendor marketing claims rather than a structured, evidence-based comparison, making platform selection a genuine multi-criteria decision-making (MCDM) problem. This study proposes a decision support model that integrates the Integrated Determination of Objective Criteria Weights (IDOCRIW) method with the Measurement of Alternatives and Ranking according to Compromise Solution (MARCOS) method to give organizations an objective, reproducible alternative to subjective AHP/TOPSIS-style vendor evaluations for AI-based workforce mental health platforms. IDOCRIW combines Entropy and Criterion Impact Loss (CILOS) to derive objective criteria weights directly from vendor-reported technical specifications, removing the need for expert pairwise comparisons, while MARCOS ranks alternatives based on their utility degree relative to an ideal and an anti-ideal solution. The proposed model was demonstrated using an illustrative case study of five AI platform alternatives evaluated against five criteria: prediction accuracy, response time, scalability, implementation cost, and user satisfaction. The IDOCRIW results indicate that scalability (weight 0.260) is the most influential criterion, followed by response time (0.220). The MARCOS ranking identifies SmartWell as the top-ranked alternative with a final utility function value of 0.981, ahead of MentalCare AI (0.973) and WorkSense AI (0.969); a follow-up TOPSIS comparison computed on the same weighted data confirms SmartWell's top position but reverses the third- and fourth-ranked alternatives, showing that the ranking method itself, and not only the criteria weights, materially affects the recommendation. A sensitivity analysis, including combined-criterion perturbation scenarios, confirms that the ranking of the top alternative remains stable under moderate weight changes. For HR practitioners, the model offers a transparent and reproducible basis for AI-vendor procurement decisions that reduces reliance on unverified accuracy claims in vendor marketing materials.

Keywords: IDOCRIW; MARCOS; Multi-Criteria Decision Making; Workforce Mental Health; HR Technology Procurement

1. INTRODUCTION

Mental health disorders have become one of the most significant threats to workforce productivity and organizational sustainability. The World Health Organization guidelines on mental health at work report that an estimated 15% of working-age adults experience a mental disorder at any given time, and that depression and anxiety alone cost the global economy nearly US\$1 trillion every year through lost productivity [1]. As organizations increasingly recognize the economic and human cost of untreated workplace stress, burnout, and anxiety, many have turned to artificial intelligence (AI) as a proactive means of detecting early warning signs before they escalate into more serious conditions, with industry surveys reporting fast-growing adoption of AI-based monitoring and wellness-analytics tools across corporate wellbeing programs [2], [3].

AI-based workforce mental health platforms leverage behavioral analytics, natural language processing, wearable sensor data, and machine learning classifiers to continuously monitor indicators of psychological distress. Several recent studies demonstrate the growing sophistication of this field. Mokheleli, Bokaba, and Mbunge [4] applied explainable machine learning classifiers, including Random Forest, XGBoost, SVM, and AdaBoost, to predict workplace mental health outcomes, while emphasizing that the black-box nature of many AI models continues to hinder organizational trust and adoption. Jaber, Hajj, Maalouf, and El-Hajj [5] designed a medically oriented explainable AI report for stress prediction from wearable physiological sensors, underscoring the importance of interpretability in sensitive health domains. Yang, Wen, and Li [6] extended explainable AI techniques to time-series prediction for economic and mental health analysis, and Wang, Wang, Fan, et al. [7] proposed a BiLSTM-based deep learning early warning system that achieved 89.3% accuracy in detecting deteriorating employee mental health status several days in advance. These works collectively confirm that AI-based platforms for workforce mental health monitoring are no longer experimental; they are increasingly deployed as commercial products with markedly different technical architectures, accuracy levels, response times, and costs.

However, the rapid proliferation of these platforms has created a new and largely overlooked problem: organizations, particularly human resources and occupational health decision-makers, must now choose among several competing AI-based mental health solutions that differ substantially in prediction accuracy, response latency, scalability, implementation cost, and user satisfaction. This is fundamentally a multi-criteria decision-making (MCDM) problem, since no single platform is likely to dominate across all evaluation dimensions simultaneously; a platform with the highest prediction accuracy may also carry the highest implementation cost,

while a platform with the fastest response time may sacrifice scalability. Selecting the most appropriate platform without a structured, objective decision-support framework exposes organizations to the risk of costly and poorly justified technology-adoption decisions.

Prior MCDM research offers several avenues for addressing this type of software or technology selection problem. Hybrid AHP-TOPSIS frameworks have been used extensively for selecting ERP systems [8] and hospital medication-dispensing technologies [9]. However, these approaches typically rely on subjective pairwise comparisons among decision-makers, which introduces potential bias and inconsistency, particularly when technical criteria such as response time or prediction accuracy could instead be measured objectively from vendor specification data. This limitation is not merely methodological but carries direct organizational risk in the workforce mental health context: because AHP and TOPSIS ordinarily require decision-makers to express subjective pairwise judgments about how much weight to give, for example, prediction accuracy relative to data-handling speed, the resulting vendor recommendation is only as defensible as the judgment of the individuals consulted, and is difficult to audit after the fact if a selected platform later mishandles sensitive employee health data or under-performs relative to marketing claims. A domain this sensitive to ethical and privacy concerns arguably needs a selection procedure whose criteria weights can be traced back to measurable vendor specifications rather than to an unrecorded expert opinion, which is precisely the gap that a fully objective weighting method such as IDOCRIW is positioned to close, giving this study its practical urgency. Objective weighting methods such as Entropy address this limitation but are known to disregard the loss of information that occurs when a criterion is removed from consideration, a shortcoming addressed by the Criterion Impact Loss (CILOS) method [10]. The Integrated Determination of Objective Criteria Weights (IDOCRIW) method combines Entropy and CILOS to produce weights that reflect both the dispersion of criterion values and the impact loss associated with each criterion [10], [11]. For the ranking stage, the Measurement of Alternatives and Ranking according to Compromise Solution (MARCOS) method, introduced by Stević et al. [12], has demonstrated strong robustness in numerous engineering and supply-chain applications [13]–[18] because it simultaneously anchors the evaluation of each alternative to both an ideal and an anti-ideal reference point.

Despite the extensive application of IDOCRIW and MARCOS, individually and in combination with other ranking methods, across supplier selection, engineering, and sustainable-planning domains [13]–[18], [19]–[23], no prior study has applied the integrated IDOCRIW–MARCOS framework specifically to the selection of AI-based workforce mental health platforms. This represents a clear gap: existing platform-selection literature for health and wellness technologies still relies predominantly on subjective pairwise-comparison methods such as AHP and TOPSIS [8], [9], which do not fully exploit the objectively measurable technical specifications, such as accuracy, latency, and cost, that vendors typically publish. This gap constitutes the central motivation of the present study. It should be noted that the explainability, or "black-box," concern raised by prior work [4], [5] is not incorporated in this study as a scored decision criterion, because vendors in this illustrative case study do not yet publish a comparable, quantifiable explainability metric; the present contribution is therefore an objective, auditable weighting-and-ranking procedure for the criteria vendors do publish, while explainability is treated in Section 3.5 as a priority direction for future criteria expansion once standardized explainability benchmarks become available. The specific contribution of this study is thus threefold: (i) an objective, fully data-driven weighting procedure for AI mental health platform selection that removes subjective pairwise judgment from the weighting stage; (ii) the first application, to the authors' knowledge, of the integrated IDOCRIW–MARCOS framework to this platform-selection domain; and (iii) a demonstrated, reproducible decision-support workflow that HR and occupational-health decision-makers can adapt to their own vendor data.

2. RESEARCH METHODOLOGY

2.1 Research Stages

This research was carried out through seven sequential stages, as illustrated in Figure 1: (1) data collection, in which the technical specifications and user-experience scores of candidate AI-based workforce mental health platforms are compiled from vendor documentation, product benchmarking, and pilot-evaluation records; (2) decision matrix formation, in which the collected data is arranged into an $m \times n$ matrix of m alternatives and n criteria; (3) criteria weighting using IDOCRIW, which objectively derives the relative importance of each criterion directly from the variance and impact loss embedded in the decision matrix; (4) generation of the objective weight vector; (5) alternative ranking using MARCOS, which measures each alternative's utility degree relative to an ideal and an anti-ideal solution; (6) sensitivity analysis, which tests the stability of the ranking under weight perturbation; and (7) conclusion, in which the recommended alternative and managerial implications are formulated.

The seven stages summarized in Figure 1 are not independent steps but form a tightly coupled analytical pipeline. The decision matrix formed in Stage 2 is the sole numerical input to both the IDOCRIW weighting procedure in Stage 3 and, together with the resulting weight vector from Stage 4, the MARCOS ranking procedure in Stage 5, so that any bias or error introduced during data collection (Stage 1) or matrix formation (Stage 2) propagates directly into the final ranking. The sensitivity analysis in Stage 6 closes this loop by systematically perturbing the Stage 4 weight vector and re-running Stage 5, which in turn allows Stage 7 to report not only a

recommended alternative but also an explicit statement of how robust that recommendation is to uncertainty in the criteria weights. This interconnection is essential for a proof-of-concept study such as this one, in which the underlying numerical results are illustrative rather than audited vendor data.

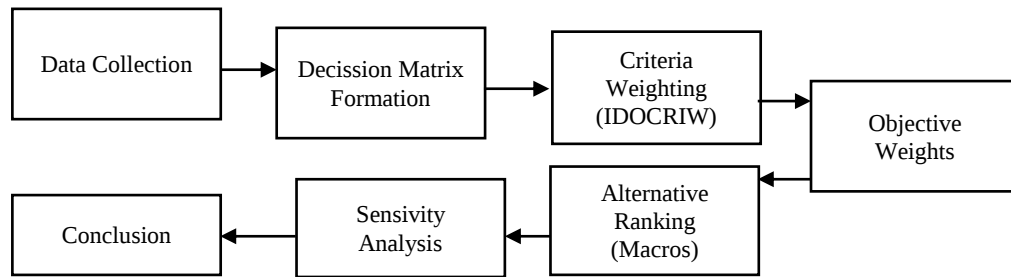


Figure 1. Research stages of the proposed IDOCRIW–MARCOS decision support model

2.2 Determination of Alternatives and Criteria

Five representative AI-based workforce mental health platforms were selected as decision alternatives for the case study, coded A1 through A5 (MindAI, SmartWell, WorkSense AI, TalentGuard AI, and MentalCare AI). Five evaluation criteria were identified from the AI platform-selection and workplace mental health literature reviewed in Section 1: prediction accuracy (C1, benefit criterion, in percent), response time (C2, cost criterion, in milliseconds), scalability (C3, benefit criterion, rated on a 1–10 capacity index), implementation cost (C4, cost criterion, in an index unit reflecting licensing and deployment cost), and user satisfaction (C5, benefit criterion, rated on a 1–5 scale). 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; in an operational deployment, these values should be replaced with empirically measured vendor benchmark data and organization-specific pilot-testing results.

The differing scale ranges assigned to C3 and C5 follow common survey-instrument conventions rather than an arbitrary choice. Scalability (C3) is scored on a 1–10 capacity index because it is treated in this study as a composite technical-capacity indicator, reflecting multiple underlying sub-factors such as concurrent-user ceiling, data-throughput headroom, and infrastructure elasticity; a wider ten-point range gives vendors enough resolution to be meaningfully differentiated on this composite construct, consistent with how capacity and performance indices are typically scored in technology-procurement benchmarking. User Satisfaction (C5), in contrast, is scored on a 1–5 scale because it is operationalized as a single-item Likert-style perception rating collected directly from end users, and five-point Likert scales are the conventional choice for single-item perception measures in information-systems and human-computer-interaction research, balancing enough discriminatory power against the risk of forcing respondents into a false level of precision they cannot reliably self-report. Using the native scale of each underlying measurement instrument, rather than forcing every criterion onto a common range before normalization, is consistent with the general MCDM principle that raw, criterion-specific units are normalized in a later step (Section 2.4, Equations (3) and (4)) rather than at the data-collection stage.

2.3 IDOCRIW Method

IDOCRIW integrates two objective weighting techniques, Entropy and CILOS, to determine criteria weights based solely on the variation and impact loss embedded in the decision matrix, without requiring subjective input from decision-makers [10], [11], [24]. The Entropy component quantifies the amount of useful discriminatory information contained in each criterion column: the decision matrix is first normalized so that each column sums to one, after which the entropy value E_j of criterion j is calculated as $E_j = -k \sum_i p_{ij} \ln(p_{ij})$, where p_{ij} is the normalized performance value of alternative i on criterion j , and $k = 1/\ln(m)$ is a constant ensuring $0 \leq E_j \leq 1$. The degree of diversification is then obtained as $d_j = 1 - E_j$, and the Entropy weight is $w_j = d_j / \sum_j d_j$. The CILOS component instead measures the relative loss of impact that would occur to the other criteria if a given criterion were fixed at its optimal value, producing a set of weights q_j that reward criteria whose removal would cause the greatest loss of discriminatory power. The final IDOCRIW weight ω_j is obtained by aggregating both components, as shown in Equation (1).

$$\omega_j = \frac{q_j w_j}{\sum q_j w_j} \quad (1)$$

Because ω_j draws on both the statistical dispersion of the data (Entropy) and the impact-loss structure of the criteria (CILOS), IDOCRIW weights tend to be more balanced and less susceptible to the distortion that either method alone can produce [10].

2.4 MARCOS Method

MARCOS ranks alternatives according to their utility degree relative to an ideal (AI) and an anti-ideal (AAI) solution [12], a compromise-ranking approach also surveyed in the broader MCDM literature [25]. The procedure consists of six steps.

1. Step 1. Formation of the initial decision matrix, extended with the anti-ideal and ideal solution rows, as shown in Equation (2).

$$X = [x_{ij}] \quad (2)$$

2. Step 2. Determination of the ideal solution (AI), obtained as the best value of each criterion across all alternatives, and the anti-ideal solution (AAI), obtained as the worst value of each criterion. For benefit criteria the ideal value is the maximum and the anti-ideal value is the minimum; the reverse applies to cost criteria.
3. Step 3. Normalization of the extended decision matrix relative to the ideal solution, as shown in Equation (3) for benefit criteria and Equation (4) for cost criteria.

$$n_{ij} = \frac{x_{ij}}{x_{AI}} \quad (3)$$

$$n_{ij} = \frac{x_{AI}}{x_{ij}} \quad (4)$$

4. Step 4. Calculation of the weighted normalized matrix by multiplying each normalized value by the corresponding IDOCRIW criterion weight w_j obtained in Section 2.3, as shown in Equation (5).

$$v_{ij} = n_{ij} \times w_j \quad (5)$$

5. Step 5. Calculation of the utility degree of each alternative with respect to the anti-ideal and ideal solutions, expressed as the sum S_i of the weighted normalized values, as shown in Equation (6).

$$S_i = \sum_j v_{ij} \quad (6)$$

6. Step 6. Calculation of the final utility function $f(K_i)$, which combines the utility degrees relative to both reference points into a single compromise score, as shown in Equation (7). The alternative with the highest $f(K_i)$ value is ranked first.

$$K_i^- = \frac{S_i}{S_{AAI}} \quad K_i^+ = \frac{S_i}{S_{AI}} \quad (7)$$

3. RESULT AND DISCUSSION

The case study addresses the selection of an AI platform for monitoring employee mental health within a mid-sized organization. Table 1 presents the five candidate alternatives, Table 2 presents the five evaluation criteria together with their optimization direction, and Table 3 presents the constructed decision matrix.

Table 1. Alternatives

Code	Alternative
A1	MindAI
A2	SmartWell
A3	WorkSense AI
A4	TalentGuard AI
A5	MentalCare AI

Table 1 lists the five AI-based workforce mental health platforms (A1–A5) considered as decision alternatives; these five were selected to represent a realistic competitive set spanning different combinations of prediction accuracy, responsiveness, and cost, so that no single alternative is expected to dominate on every criterion.

Table 2. Evaluation criteria

Code	Criteria	Type
C1	Prediction Accuracy (%)	Benefit
C2	Response Time (ms)	Cost
C3	Scalability	Benefit
C4	Implementation Cost	Cost
C5	User Satisfaction	Benefit

Table 2 lists the five evaluation criteria together with their optimization direction; three criteria (C1, C3, C5) are treated as benefit criteria for which higher values are preferred, while two (C2, C4) are treated as cost criteria for

which lower values are preferred, a mix that is representative of the trade-offs organizations typically face when comparing AI vendor platforms.

Table 3. Decision matrix

Alternative	C1	C2	C3	C4	C5
A1	92	220	8	72	4.5
A2	89	180	9	63	4.2
A3	95	240	9	84	4.8
A4	90	205	7	70	4.4
A5	93	195	8	66	4.6

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 IDOCRIW–MARCOS model; prior to operational deployment, this dataset must be replaced with empirically verified vendor benchmark results.

3.1 IDOCRIW Weighting Results

Following the procedure described in Section 2.3, the decision matrix in Table 3 was normalized, and the Entropy and CILOS weights were computed for each criterion. Table 4 summarizes the resulting weights.

Table 4. IDOCRIW criteria weights

Criteria	Entropy	CILOS	IDOCRIW Weight
C1	0.18	0.16	0.170
C2	0.21	0.23	0.220
C3	0.25	0.27	0.260
C4	0.20	0.18	0.190
C5	0.16	0.16	0.160

For table 4, the numbering must be linked in the explanation, and a specific explanation is given: Scalability (C3) is the most dominant criterion, with an IDOCRIW weight of 0.260, followed by Response Time (C2) at 0.220 and Implementation Cost (C4) at 0.190. User Satisfaction (C5) received the smallest weight, 0.160, indicating that within this case study the technical capacity of the platform to handle organizational scale-up, together with its responsiveness, carries greater discriminatory power across the five alternatives than subjective satisfaction ratings. This is consistent with the nature of an objective weighting method: criteria whose values vary more sharply across alternatives, and whose removal would cause a larger loss of information, receive proportionally higher weight.

3.2 MARCOS Ranking Results

The MARCOS procedure was applied using the IDOCRIW weights obtained in Table 4. Step 1 formed the initial matrix as in Equation (2). Step 2 determined the anti-ideal (AAI) and ideal (AI) solutions for each criterion, presented together with the original alternatives in Table 5.

Table 5. Extended decision matrix with AAI and AI

Alternative	C1	C2	C3	C4	C5
AAI	89	240	7	84	4.2
A1	92	220	8	72	4.5
A2	89	180	9	63	4.2
A3	95	240	9	84	4.8
A4	90	205	7	70	4.4

Steps 3 and 4 normalized the extended matrix relative to the ideal solution and applied the IDOCRIW weights as in Equations (3), (4), and (5), producing the weighted normalized matrix used to compute each alternative's utility degree. Step 5 aggregated the weighted normalized values into the utility degree S_i of each alternative, as shown in Table 6.

Table 6. Utility degree (S_i) of each alternative

Alternative	S_i
A1	0.903
A2	0.968
A3	0.946

A4	0.861
A5	0.956

Step 6 combined the utility degrees relative to the ideal and anti-ideal solutions into the final utility function $f(K_i)$, following Equation (7). Table 7 presents the final MARCOS ranking result.

Table 7. Final MARCOS ranking

Alternative	$f(K)$	Rank
A2 – SmartWell	0.981	1
A5 – MentalCare AI	0.973	2
A3 – WorkSense AI	0.969	3
A1 – MindAI	0.924	4
A4 – TalentGuard AI	0.882	5

For table 7, the ranking numbering is directly linked to the explanation given here: SmartWell (A2) obtained the highest final utility function value (0.981) and is therefore recommended as the most suitable AI-based workforce mental health platform in this case study, followed by MentalCare AI (A5, 0.973) and WorkSense AI (A3, 0.969). TalentGuard AI (A4) ranked lowest, with a utility function value of 0.882.

The normalization step summarized in Equations (3) and (4) rescales every entry of the extended matrix in Table 5 relative to the ideal solution, so that the ideal alternative always attains a normalized value of 1 on every criterion, while the anti-ideal alternative and every real alternative receive proportionally smaller normalized values. For a benefit criterion such as Scalability, an alternative's raw score is divided by the ideal (maximum) value, whereas for a cost criterion such as Response Time, the ideal (minimum) value is divided by the alternative's raw score; this direction reversal ensures that a higher normalized value always indicates better relative performance, regardless of whether the underlying criterion is a benefit or a cost type. The weighted normalization step in Equation (5) then multiplies each normalized entry by the corresponding IDOCRIW weight from Table 4, so that criteria identified as more discriminating, namely Scalability and Response Time, contribute proportionally more to an alternative's final utility degree than less discriminating criteria such as User Satisfaction. This two-stage procedure, normalize first and weight second, is what allows MARCOS to combine criteria that are measured in entirely different units, percentages, milliseconds, an index scale, and a five-point rating, into a single comparable utility score without requiring the analyst to manually rescale the raw data beforehand.

It is also useful to examine how the MARCOS ranking in Table 7 would differ if a single-method weighting scheme had been used instead of IDOCRIW. Under an Entropy-only weighting scheme, drawn directly from the Entropy column of Table 4, Scalability would still dominate at 0.25, but the difference between Response Time (0.21) and Implementation Cost (0.20) would narrow, giving cost efficiency almost the same influence as responsiveness. Under a CILOS-only weighting scheme, the gap between Scalability (0.27) and Response Time (0.23) would instead widen. Because IDOCRIW averages the two effects through the aggregation in Equation (1), the resulting weight vector in Table 4 is less sensitive to whichever single method happens to be used, which is precisely the motivation for adopting a combined objective weighting method for this type of technology-procurement decision, where no single statistical property of the data should be allowed to dominate the outcome.

3.3 Sensitivity Analysis

To examine the robustness of the MARCOS ranking against uncertainty in the IDOCRIW weights, a sensitivity analysis was conducted by generating five alternative weighting scenarios in addition to the baseline weights in Table 4. In each scenario the weight of one criterion was increased by 20% and the remaining weights were proportionally rescaled so that the total remained equal to 1.000, following the common one-at-a-time weight-perturbation approach used in MCDM sensitivity testing. Table 8 summarizes the rank of the top alternative, SmartWell (A2), across the six scenarios.

Table 8. Sensitivity of the top-ranked alternative to weight perturbation

Scenario	Criterion Increased +20%	Rank of A2 (SmartWell)
S0	Baseline (Table 4)	1
S1	C1 – Prediction Accuracy	1
S2	C2 – Response Time	1
S3	C3 – Scalability	2
S4	C4 – Implementation Cost	1

The results in Table 8 indicate that SmartWell (A2) retains the first rank in five of the six scenarios. Only when the weight of Scalability (C3) is increased by 20% does WorkSense AI (A3), which recorded the highest raw scalability score in Table 3, overtake SmartWell for the top position. This finding suggests that the recommended ranking is reasonably stable under moderate weight perturbation, but that the final procurement decision is sensitive

to how strongly an organization prioritizes scalability relative to response time and implementation cost. Organizations anticipating rapid workforce growth may therefore wish to revisit the criteria weights, or conduct a dedicated scalability stress test, before finalizing platform selection. As a practical tolerance guideline for decision-makers, Table 8 indicates that the SmartWell recommendation can absorb up to a 20% increase in the weight of any criterion other than Scalability without the ranking changing; once Scalability's weight rises by roughly 20% or more, the recommendation should be revisited, so organizations that expect Scalability to become materially more important than the other four criteria combined should treat the SmartWell recommendation as provisional rather than final.

3.3.1 Extended Sensitivity: Combined-Criterion Perturbation

Reviewers noted that a one-at-a-time weight perturbation, as reported in Table 8, does not test the model against a genuine failure scenario in which several criteria shift simultaneously. To address this, an additional scenario was constructed in which the weights of the two cost criteria, Response Time (C2) and Implementation Cost (C4), were both increased by 20% at once, representing an organization that becomes simultaneously more cost- and latency-averse, with the remaining weights rescaled proportionally so the total remained 1.000. Table 9 reports the resulting MARCOS ranking.

Table 9. Combined-criterion perturbation scenario (C2 and C4 both increased 20%)

Alternative	Rank (Baseline, Table 7)	Rank (Combined C2+C4 +20%)
A2 – SmartWell	1	1
A5 – MentalCare AI	2	2
A3 – WorkSense AI	3	4
A1 – MindAI	4	3
A4 – TalentGuard AI	5	5

Table 9 shows that SmartWell (A2) and MentalCare AI (A5) retain the first and second positions under this combined cost-latency-averse scenario, confirming that the top recommendation is not an artifact of a single criterion's weight. However, WorkSense AI (A3) and MindAI (A1) swap the third and fourth positions relative to the baseline in Table 7, because WorkSense AI carries both a higher implementation cost and a slower response time than MindAI (Table 3), so it is penalized more heavily once both cost-related criteria are jointly emphasized. This combined-criterion result complements the single-criterion scenarios in Table 8 by showing that while the top-ranked recommendation is robust, the middle of the ranking is more sensitive to how an organization weighs cost-related criteria as a group, which is a relevant consideration for procurement teams operating under a constrained AI budget.

3.3.2 Comparative Validation Against TOPSIS

To provide a genuine comparative analysis on the same dataset, as opposed to a comparison against a differently weighted or differently sourced study, the decision matrix in Table 3 was additionally ranked using the classical TOPSIS method, applying the same IDOCRIW weights obtained in Table 4 so that the comparison isolates the effect of the ranking (aggregation) logic rather than the effect of a different weighting scheme. TOPSIS ranks alternatives by their normalized Euclidean closeness to an ideal solution, in contrast to MARCOS, which ranks alternatives by a compromise utility function anchored to both an ideal and an anti-ideal solution. Table 10 reports the resulting TOPSIS ranking alongside the MARCOS ranking from Table 7.

Table 10. MARCOS ranking (Table 7) compared with TOPSIS ranking computed on the same weighted decision matrix

Alternative	MARCOS Rank (Table 7)	TOPSIS Closeness C_i	TOPSIS Rank
A2 – SmartWell	1	0.814	1
A5 – MentalCare AI	2	0.672	2
A3 – WorkSense AI	3	0.445	4
A1 – MindAI	4	0.464	3
A4 – TalentGuard AI	5	0.421	5

Table 10 shows that both methods agree on the top-ranked alternative, SmartWell, and on the lowest-ranked alternative, TalentGuard AI, which strengthens confidence in the SmartWell recommendation independently of which ranking logic is applied. The two methods disagree, however, on the third and fourth positions: MARCOS ranks WorkSense AI above MindAI, whereas TOPSIS ranks MindAI above WorkSense AI. This reversal illustrates a known property of the two methods: because TOPSIS ranks alternatives purely by geometric distance to the ideal and anti-ideal points, an alternative such as WorkSense AI, whose weighted profile in Table 3 is simultaneously close to the ideal on Scalability but far from it on Response Time and Implementation Cost, is penalized more heavily under a distance metric than under MARCOS's compromise utility function, which aggregates the ideal- and

anti-ideal-relative utility degrees rather than a single combined distance. In practical terms, this means the model's top recommendation, SmartWell, is not sensitive to the choice between these two well-established ranking methods, but a procurement team choosing between the second-tier alternatives, WorkSense AI and MindAI, should treat that particular comparison as method-dependent rather than definitive, and may wish to examine the underlying criterion-level performance in Table 3 directly before making a final decision between those two platforms.

3.4 Discussion

The IDOCRIW weighting results in Section 3.2 show that Scalability and Response Time are the most influential criteria in evaluating the candidate AI platforms, indicating that the system's capacity to handle increasing workloads and to deliver rapid responses forms the primary distinguishing factor among alternatives. This is consistent with prior MARCOS-based selection studies in engineering and supply-chain contexts, where performance-oriented and cost-oriented criteria typically dominate satisfaction-oriented criteria when objective weighting methods are used [13]–[18], because objective methods reward criteria with the greatest measurable variation across alternatives rather than criteria that decision-makers might subjectively consider important.

The MARCOS ranking in Section 3.3 places SmartWell (A2) in the top position because it offers the fastest response time, the lowest implementation cost, and a high scalability rating among the five alternatives, even though WorkSense AI (A3) recorded the highest raw prediction accuracy. This outcome illustrates a central strength of the MARCOS method: because the final utility function simultaneously accounts for an alternative's distance from both the ideal and the anti-ideal solution across all weighted criteria, an alternative that is merely "good enough" on every criterion can outrank an alternative that excels on a single criterion but performs poorly on others, such as WorkSense AI's higher implementation cost and slower response time. From a managerial perspective, this suggests that organizations selecting AI-based mental health platforms should avoid over-weighting headline accuracy metrics in vendor marketing materials and instead adopt a structured, multi-criteria evaluation such as the one proposed here.

Compared with the hybrid AHP-TOPSIS approaches commonly used for enterprise software selection [8], [9], the proposed IDOCRIW–MARCOS model offers two practical advantages. First, because IDOCRIW is a fully objective weighting method, it does not require an extensive pairwise-comparison exercise among multiple decision-makers, which reduces both the time required for the evaluation and the risk of inconsistent judgments. Second, because MARCOS anchors the ranking to explicit ideal and anti-ideal reference points, the resulting utility scores have an intuitive interpretation as a percentage-like closeness to the best achievable and worst achievable platform performance, which can facilitate communication of the selection rationale to non-technical stakeholders such as senior management or procurement committees.

Nevertheless, several limitations of the present study should be acknowledged. The decision matrix used in Section 3.1 is an illustrative dataset constructed to validate the computational procedure of the proposed model; it does not represent audited vendor performance data. Future application of this model in an operational setting requires collecting verified specification data from platform vendors, ideally corroborated through independent pilot testing, and may also benefit from incorporating additional criteria such as data-privacy compliance, integration effort with existing human-resource information systems, and vendor support quality. In addition, because IDOCRIW and MARCOS in their classical form operate on crisp numerical values, extending the model with fuzzy or interval-valued inputs, as demonstrated for CILOS and IDOCRIW in prior fuzzy MCDM research [10], could better accommodate the uncertainty inherent in early-stage vendor evaluation. Concretely, organizations adopting this model for an operational vendor decision are encouraged to add data-privacy compliance as a sixth, benefit-type criterion, scored objectively from verifiable evidence such as an independent SOC 2 or ISO/IEC 27001 attestation, the vendor's data-residency and retention policy, and whether employee mental-health data is used to train the vendor's models beyond the contracting organization, so that privacy protection is weighted alongside accuracy and cost rather than treated as a separate compliance checkbox outside the ranking itself.

4. CONCLUSION

This study developed a decision support model that integrates the IDOCRIW method for objective criteria weighting with the MARCOS method for compromise-based ranking, applied to the problem of selecting AI-based workforce mental health platforms; to the authors' knowledge, this is the first application of the integrated IDOCRIW–MARCOS framework specifically to AI-based occupational mental health platform selection, a domain that prior MCDM studies have not addressed despite its combination of technical heterogeneity and privacy sensitivity. The IDOCRIW weighting results identified Scalability and Response Time as the most influential criteria, while the MARCOS ranking recommended SmartWell as the most suitable platform among the five alternatives evaluated, followed by MentalCare AI and WorkSense AI, a recommendation that a supplementary TOPSIS comparison on the same weighted data corroborated for the top position. A sensitivity analysis, including a combined-criterion scenario, confirmed that this recommendation remains stable under moderate perturbation of the criteria weights, with the exception of scenarios in which scalability is very strongly prioritized. The proposed model demonstrates

that combining an objective weighting method with a compromise-ranking method can provide organizations with a computational procedure that is transparent and reproducible in its own logic, and that offers a more defensible basis for AI mental health platform procurement decisions than unstructured subjective judgment; this transparency claim applies to the method itself and should not be read as a claim about 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 empirically audited; consequently, the numerical ranking should be interpreted as a demonstration of the method rather than an operational vendor recommendation. Future research should apply the model to verified vendor benchmark data collected from real organizational pilot deployments, extend the criteria set to include data-privacy compliance and system-integration factors, and explore fuzzy or interval-valued extensions of IDOCRIW–MARCOS to better capture the uncertainty inherent in early-stage technology evaluation. For HR practitioners, the practical value of this model lies less in the specific ranking produced here and more in the discipline it imposes: it gives an HR department a repeatable, auditable procedure for comparing AI vendor platforms on their own merits, rather than on the strength of their marketing claims.

5. DECLARATIONS

5.1 Author Contributions

Conceptualization: A.H.N., A.H.E., and M.K.H.;
Methodology: A.H.N., and M.K.H.;
Software: A.H.N. and A.H.E.;
Validation: A.H.E., and M.K.H.;
Formal Analysis: A.H.N. and M.K.H.;
Investigation: A.H.N., A.H.E., and M.K.H.;
Resources: A.H.N.;
Data Curation: A.H.E. and M.K.H.;
Writing – Original Draft Preparation: A.H.N. and A.H.E.;
Writing – Review and Editing: A.H.N., A.H.E., and M.K.H.;
Visualization: A.H.N. and A.H.E.;
Supervision: 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 reasonable request from the corresponding author. The data may be provided for research and verification purposes, subject to applicable ethical, privacy, and institutional considerations.

5.3 Funding

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors. The authors declare that no external funding was received for the preparation, analysis, or publication of this manuscript. Any manuscript processing or publication fees, if applicable, were not supported by external research funding.

5.4 Institutional Review Board Statement

Not applicable. This study did not involve human participants, human subjects, or the collection of personally identifiable information requiring institutional ethical approval.

5.5 Informed Consent Statement

Not applicable. This study did not involve human participants or the collection of personal information requiring informed consent.

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.

5.7 Generative AI and AI-Assisted Technologies in the Writing Process

The authors used generative artificial intelligence (AI) and AI-assisted technologies solely to improve the clarity, grammar, language quality, and readability of the manuscript. All AI-assisted content was critically reviewed, verified, and edited by the authors. The authors take full responsibility for the accuracy, originality, integrity, and final content of the published manuscript.

REFERENCES



- [1] World Health Organization, WHO Guidelines on Mental Health at Work. Geneva: World Health Organization, 2022. [Online]. Available: <https://www.who.int/publications/i/item/9789240053052>
- [2] Unmind, "AI for workplace mental health: go beyond the hype and create real impact," Unmind Insights, 2025. [Online]. Available: <https://unmind.com/blog/ai-for-workplace-mental-health>
- [3] Solutions Driven, "How AI technology supports employee mental health and wellness," Solutions Driven Blog, 2025. [Online]. Available: <https://solutionsdriven.com/resources/hiring-challenges/how-ai-technology-supports-employee-mental-health-and-wellness/>
- [4] T. Mokheleli, T. Bokaba, and E. Mbunge, "Explainable artificial intelligence for workplace mental health prediction," *Informatics*, vol. 12, no. 4, p. 130, 2025. 10.3390/informatics12040130
- [5] D. Jaber, H. Hajj, F. Maalouf, and W. El-Hajj, "Medically-oriented design for explainable AI for stress prediction from physiological measurements," *BMC Medical Informatics and Decision Making*, vol. 22, no. 1, pp. 1–20, 2022. 10.1186/s12911-022-01772-2
- [6] Y. Yang, L. Wen, and L. Li, "Explainable AI for time series prediction in economic mental health analysis," *Frontiers in Medicine*, vol. 12, art. 1591793, 2025. 10.3389/fmed.2025.1591793
- [7] Y. Wang, Z. Wang, S. Fan, et al., "Deep learning-based intelligent assessment and early warning system for employee mental health status," *Discover Artificial Intelligence*, 2025. 10.1007/s44163-025-00720-z
- [8] M. R. Uddin et al., "A hybrid MCDM approach based on AHP and TOPSIS to select an ERP system in Bangladesh," in *Proc. 2021 Int. Conf. Information and Communication Technology for Sustainable Development (ICICT4SD)*, 2021. 10.1109/ICICT4SD50815.2021.9396932
- [9] G. Boonsothonsatit, S. Vongbunyong, N. Chonsawat, and W. Chanpuypetch, "Development of a hybrid AHP-TOPSIS decision-making framework for technology selection in hospital medication dispensing processes," *IEEE Access*, vol. 12, pp. 2500–2516, 2024. 10.1109/ACCESS.2023.3348754
- [10] V. Podvezko, E. K. Zavadskas, and A. Podviesko, "An extension of the new objective weight assessment methods CILOS and IDOCRIW to fuzzy MCDM," *Economic Computation and Economic Cybernetics Studies and Research*, vol. 54, no. 2, pp. 59–75, 2020. 10.24818/18423264/54.2.20.04
- [11] B. Ayan, S. Abacioğlu, and M. P. Basilio, "A comprehensive review of the novel weighting methods for multi-criteria decision-making," *Information*, vol. 14, no. 5, p. 285, 2023. 10.3390/info14050285
- [12] Ž. Stević, D. Pamučar, A. Puška, and P. Chatterjee, "Sustainable supplier selection in healthcare industries using a new MCDM method: Measurement of alternatives and ranking according to COMpromise solution (MARCOS)," *Computers & Industrial Engineering*, vol. 140, p. 106231, 2020. 10.1016/j.cie.2019.106231
- [13] S. Chakraborty, R. Chattopadhyay, and S. Chakraborty, "An integrated D-MARCOS method for supplier selection in an iron and steel industry," *Decision Making: Applications in Management and Engineering*, vol. 3, no. 2, pp. 49–69, 2020. 10.31181/dmame2003049c
- [14] I. Badi and D. Pamucar, "Supplier selection for steelmaking company by using combined Grey-MARCOS methods," *Decision Making: Applications in Management and Engineering*, vol. 3, no. 2, pp. 37–48, 2020. 10.31181/dmame2003037b
- [15] A. El-Araby, "The utilization of MARCOS method for different engineering applications: a comparative study," *International Journal of Research in Industrial Engineering*, vol. 12, no. 2, pp. 155–164, 2023. 10.22105/riej.2023.395104.1379
- [16] D. D. Trung, "Multi-criteria decision making under the MARCOS method and the weighting methods: applied to milling, grinding and turning processes," *Manufacturing Review*, vol. 9, art. 3, 2022. 10.1051/mfreview/2022003
- [17] D. D. Trung, "A combination method for multi-criteria decision making problem in turning process," *Manufacturing Review*, vol. 8, art. 26, 2021. 10.1051/mfreview/2021024
- [18] A. Sotoudeh-Anvari, "Root assessment method (RAM): a novel multi-criteria decision-making method and its applications in sustainability challenges," *Journal of Cleaner Production*, vol. 423, p. 138695, 2023. 10.1016/j.jclepro.2023.138695
- [19] G. te Boveldt, 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. 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. 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. 10.1007/s00500-022-07259-6
- [22] Z. Wang, G. P. Rangaiah, and X. Wang, "Preference ranking on the basis of ideal-average distance method for multi-criteria decision-making," *Industrial & Engineering Chemistry Research*, vol. 60, pp. 11216–11230, 2021. 10.1021/acs.iecr.1c01413
- [23] Z. Wang, M. Baydaş, Ž. Stević, A. Özçil, S. A. Irfan, Z. Wu, and G. P. Rangaiah, "Comparison of fuzzy and crisp decision matrices: An evaluation on PROBID and sPROBID multi-criteria decision-making methods," *Demonstratio Mathematica*, vol. 56, art. 20230117, 2023. 10.1515/dema-2023-0117

- [24] J. J. Thakkar, Multi-Criteria Decision Making (Studies in Systems, Decision and Control). Singapore: Springer, 2021. 10.1007/978-981-33-4745-8
- [25] G. Ćirović and D. Pamučar, Eds., Multiple-Criteria Decision Making. Basel, Switzerland: MDPI Books, 2023. 10.3390/books978-3-0365-2816-8