



An Integrated Multi-Criteria Decision-Making Framework for Athlete Selection in Sport Management: The METE Weighting and SEDAT Ranking Methods

Ömer Faruk Görçün^{1,*}, Volkan Kapçak², Mustafa Subaşı³

¹ Department of Business Administration, Kadir Has University, Istanbul, Türkiye.

² Faculty of Sport Sciences, Tekirdağ Namık Kemal University, Tekirdağ, Türkiye

³ Sports Department, Zeytinburnu Municipality, Istanbul, Türkiye

ARTICLE INFO

Article history:

Received 2 June 2026

Received in revised form 12 July 2026

Accepted 19 July 2026

Available online 26 July 2026

Keywords: Rowing; Athlete Selection; Criterion Weighting; Superiority Methods; Sports Management; Decision Support, Sports Marketing

ABSTRACT

Athlete selection—talent identification, squad construction, and boat formation—is among the most consequential recurring decisions facing sport managers and coaches, yet in practice it frequently rests on ad-hoc criterion weighting and single-number rankings whose assumptions are never made explicit. Multi-criteria decision-making (MCDM) offers a principled alternative, but widely used techniques carry structural weaknesses that can distort or reverse the orderings they produce. This study develops and validates an integrated MCDM framework composed of two complementary methods. METE is a loss-framed subjective weighting method that embeds criterion interactions through a relational (DEMATEL-type) structure and guarantees monotonicity between expert judgement and the resulting weights. SEDAT is a ranking method that represents alternatives as a dominance network and fuses thresholded outranking, net flow, and a damped eigenvector (prestige) centrality, thereby avoiding the degeneration and rank-reversal problems of additive aggregations. The formal properties of each method are stated, the core axioms are proven, and reliability and sensitivity analysis are treated as built-in stages of the procedure rather than optional add-ons. The framework is demonstrated on the selection of a rowing crew using a fully auditable single-sheet workbook. A 5,000-iteration Monte Carlo simulation, a Markov-chain analysis of the dominance network, internal-consistency and inter-rater statistics (Cronbach's α , Kendall's W , Fleiss' κ , ICC), sensitivity and cross-method comparisons, and a rank-reversal test jointly show that the results are defensible, reproducible, and stable. Implications for evidence-based, accountable decision-making are discussed for both sport management and sport marketing, where transparent and contestable selection procedures carry concrete governance and commercial value.

1. Introduction

Few activities define the work of a sport organization as directly as deciding who will compete and who will be selected. Talent identification, squad and boat selection, and the periodic evaluation

* Corresponding author.

E-mail address: omer.gorcun@khas.edu.tr

<https://doi.org/10.59543/x3cqwy85>

of athletes are recurring, high-stakes decisions that shape competitive success, the allocation of scarce resources, and an organization's reputation for fairness. These decisions are inherently multidimensional: a selector must weigh physiological capacity against technical skill, current form against durability, and individual output against a contribution to collective performance. Frameworks that explain international sporting success position talent identification and development as one of the load-bearing pillars of elite sport systems [1], [2]. Because the relevant criteria are numerous, expressed in different units, and frequently in conflict, the task is a textbook multi-criteria decision-making (MCDM) problem. Sports management comprises complementary processes, such as planning human resources, evaluating performance, using organizational resources effectively, and balancing stakeholder expectations. Within these processes, athlete selection stands out as a strategic decision area that directly affects both sports performance and the realization of organizational goals. For this reason, basing selection decisions on systematic, justifiable, and objective criteria is considered a basic requirement of contemporary sports management (Hoye et al., 2024; Taylor et al., 2015).

In practice, however, such decisions are often resolved through informal judgement, the implicit weighting of favored indicators, or reduction to a single composite score whose construction is never documented. As studies of athletes' decision-making under uncertainty make clear, decisions in sporting contexts are entangled with contextual and psychological factors [3]. The result can be choices that are hard to justify, impossible to reproduce, and open to challenge—an uncomfortable position for organizations increasingly expected to demonstrate transparency and accountability to athletes, boards, funding bodies, and the public. The commercial stakes compound the governance stakes: selection determines which athletes become the marketable assets around which sponsorship, ticketing, and media narratives are built [4], [1], so an indefensible selection is also a fragile foundation for a club's or federation's commercial strategy. In today's professional sports, athletes are considered strategic assets that contribute to the brand communication and commercial activities of clubs and federations, beyond producing sporting success. Therefore, athlete selection is not only a technical process based on performance criteria, but also a managerial decision that can be associated with marketing dimensions such as fan interest, media visibility, and sponsorship potential [1].

Formal MCDM methods are designed precisely for problems of this kind [5], [6]. Techniques such as the analytic hierarchy process [7], which derives criterion importance from pairwise comparisons, and distance-based TOPSIS [8] offer a structured and defensible logic compared with intuitive selection. Yet the uncritical adoption of these tools can create problems of its own: some methods impose a heavy cognitive burden or assume that criteria which plainly interact in sport are independent; some ranking rules are prone to rank reversal when the set of alternatives changes; and, most importantly for practitioners, the weighting and aggregation steps are often presented as a black box, so a selector cannot readily audit why a decision came out as it did or test how sensitive it is to reasonable changes of judgement.

This study proposes an integrated framework that addresses these shortcomings directly by pairing two purpose-built methods: METE for criterion weighting and SEDAT for ranking. The contribution is deliberately practical. Both methods are transparent enough to be implemented and inspected in a single spreadsheet, yet formal enough to carry proven guarantees of monotonicity and bounded sensitivity. We demonstrate the framework on the selection of a rowing crew and subject it to an unusually comprehensive robustness and validity battery. Throughout, we read the results through the twin lenses of sport management and sport marketing, arguing that an auditable selection procedure is simultaneously a governance instrument and a commercial asset. The main

contribution of the present study is not only to propose a new multi-criteria decision-making approach. The study provides an analytical framework that enables transparent, accountable, and reproducible justification of highly uncertain managerial decisions such as athlete selection. In addition, the proposed model offers an integrative approach to the sports management and sports marketing literature by allowing sports performance criteria and marketing-relevant criteria to be evaluated within the same decision system.

The remainder of the paper is organized as follows. Section 2 reviews MCDM in sport and locates the two recurring failure modes—rank reversal and opaque aggregation—that motivate the design. Section 3 develops the METE and SEDAT methods and states their formal properties. Section 4 applies the framework to a rowing-crew selection problem. Section 5 reports the full robustness and validity analysis, comprising Monte Carlo simulation, Markov-chain analysis, internal-consistency testing, sensitivity analysis, cross-method comparison, and a rank-reversal test. Section 6 discusses the managerial, governance, and marketing implications, and Section 7 concludes.

2. Multi-criteria decision-making in sport management and marketing

MCDM provides a structured way of combining numerous, often conflicting criteria into an overall evaluation, and it has a long history in operations research and management science. Within sport, researchers and practitioners have used MCDM to support talent identification, the evaluation of player and team performance, host-city and venue selection, sponsor and supplier choice, and a range of governance decisions [9], [10], [11]. The appeal is clear: by making the criteria, the weights, and the aggregation rule explicit, MCDM converts an opaque judgement into a transparent and repeatable procedure that others can inspect and contest. This property matters twice over in professional sport, where the same selection that must satisfy a competition director must also withstand scrutiny from athletes, agents, boards, and—through the media—the paying public.

Two methodological choices sit at the heart of any MCDM application. The first is how the relative importance of the criteria is determined. Subjective weighting methods derive importance from expert judgement—through pairwise comparisons, as in the analytic hierarchy process [7] and the best-worst method [12], or through direct scoring, as in step-wise weight assessment [13]. Objective methods instead infer importance from the data; entropy-based schemes [14], for example, assign larger weights to criteria that better discriminate among alternatives; other objective schemes derive weights from the contrast intensity and correlation structure of the data, as in the CRITIC [15] and MEREC [16] methods. The DEMATEL technique [17], which recognizes interdependence among criteria, attempts to capture the reality that selection criteria are rarely independent—physiological power, technical skill, and on-water boat-moving effectiveness plainly reinforce one another.

The second choice is how performance is aggregated into a ranking. Value- and distance-based methods such as TOPSIS [8] summaries each alternative with a single score, while compromising methods such as VIKOR [18] balance group utility against individual regret. Outranking methods, notably PROMETHEE [19], instead build pairwise preference relations and derive the ranking from the resulting flows; this approach can be represented naturally as a dominance network and connects to eigenvector-based notions of prestige [20]. More recent aggregation schemes—including WASPAS [21], EDAS [22], MARCOS [23], CoCoSo [24], CRADIS [25], and RAM [26]—have extended this toolbox, while outranking approaches trace back to the ELECTRE family [27] and to PROMETHEE, whose methodologies and applications have been surveyed extensively [28]. The distinction is consequential for selection: a distance-based score rewards proximity to an ideal profile, whereas an outranking flow rewards an alternative for the caliber of the rivals it beats—an idea with a direct commercial

analogue, since an athlete who defeats the strongest competitors carries marquee value out of proportion to any aggregate index.

Despite this richness, three practical problems occur when these tools are transferred to sport selection. First, many methods are vulnerable to rank reversal [29]: adding or removing an alternative can flip the order of others, which is untenable when a selection must be defended after a late withdrawal or a new arrival. Second, some aggregation rules degenerate at parameter values, silently discarding information. Third, and most important for practitioners, the weighting and aggregation steps are frequently opaque, so a selector can neither audit why the decision emerged nor test its sensitivity to reasonable changes of judgement. Purpose-built and hybrid MCDM models have been proposed to mitigate such weaknesses in applied selection problems ranging from vehicle to supplier evaluation [30], [31]. The framework developed below is designed to address all three problems while remaining implementable and inspectable in a spreadsheet.

These properties acquire additional value when selection is viewed as a marketing decision. In modern sport, athletes are not only competitors but also brand assets, and roster construction shapes a team's commercial narrative, its media appeal, and the segments of the fan base it can address. A weighting scheme that can encode not only physiological and technical importance but, in principle, marketability and audience fit, combined with a ranking that rewards victories over strong rivals, provides a defensible bridge between the sporting and commercial logics that sport managers must reconcile. It is this dual reading—governance on one side, marketing on the other—that we develop throughout the paper.

3. Multi-criteria decision-making in sport management and marketing

The proposed framework separates the two methodological choices into two methods applied in sequence. METE produces a vector of criterion weights from expert judgement; SEDAT uses those weights to rank the alternatives. Throughout we consider n criteria and m alternatives, and we require two properties of any acceptable procedure: weights and scores must be normalized and bounded, and they must be monotone, meaning that a more important criterion never receives a smaller weight and an alternative that is better on every dimension is never ranked lower. These requirements are stated formally below and verified empirically in Section 5. The basic algorithm of the model is demonstrated in Fig (1)

3.1 The METE weighting method

METE derives importance from a loss frame. For each criterion, one or more of an expert's score, on a bounded scale, the loss the programme would incur if that criterion were deficient. Let L_j denote the aggregated loss score for criterion j . The base share of criterion j is the ratio of its loss to the total loss, given in Eq. (1).

$$R_j = \frac{L_j}{\sum_{t=1}^m L_t}, j = 1, 2, \dots, m \quad (1)$$

R_j is the normalized value (or normalized weight) of criterion j , L_j denotes the value of criterion j , $\sum_{t=1}^m L_t$ represents the sum of the values of all criteria, m is the total number of criteria. Because criteria interact, METE corrects these base shares through a DEMATEL-type relational structure: experts judge the direct influence of each criterion on the others; the resulting matrix A is normalized to $N = A/s$ with s the larger of the maximum row and column

sums; the total (direct plus indirect) influence is obtained as $\mathbf{T} = \mathbf{N}(\mathbf{I}-\mathbf{N})^{-1}$; and each criterion receives a visibility (prominence) score π_j , high when it is strongly embedded in the interaction network. The final weight blends the base share with visibility through a single, interpretable structural-emphasis parameter θ , as in Eq. (2).

$$w_j = \frac{R_j(n\pi_j)^\theta}{\sum_{t=1}^m R_t(n\pi_t)^\theta}, j = 1, 2, \dots, m \quad (2)$$

Equation (2) has three important properties. It is normalized, so the weights sum to one. It is monotone in the loss scores: a criterion judged more critical always receives a larger weight guarantee that earlier entropy-based constructions cannot provide. And when no interaction is reported ($\pi_j = 1/n$), it reduces exactly to the base shares. Reliability is assessed directly: when several experts contribute, their agreement is measured with a standard consensus statistic, and the dispersion of the weights is summarized by their normalized entropy $\tilde{H} = -(1/\ln n) \sum w_j \ln w_j$, which approaches one for near-uniform weights and zero for highly concentrated ones.

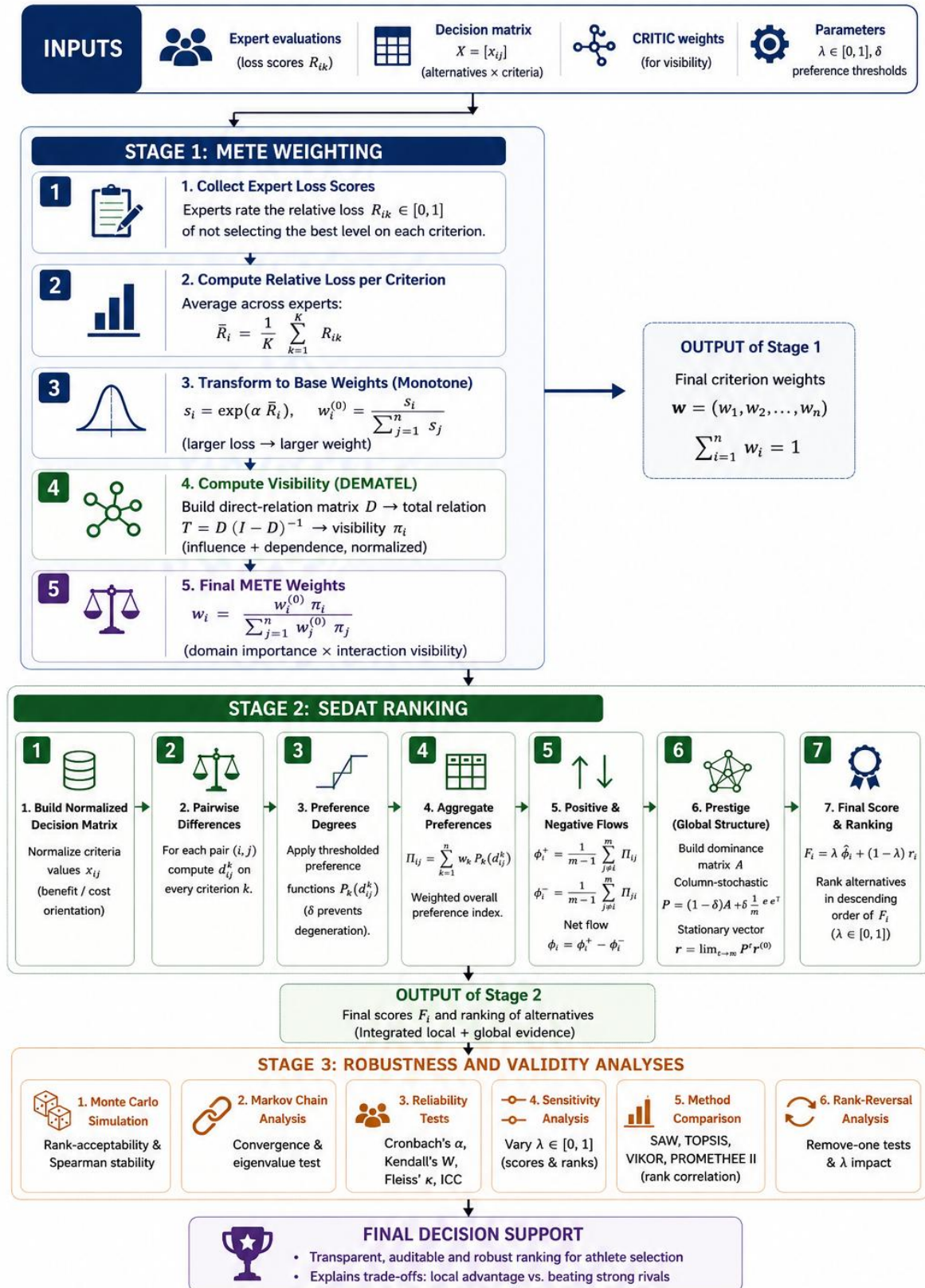


Fig. 1. The basic algorithm of the proposed METE-SEDAT decision-making model

3.2 The SEDAT ranking method

SEDAT ranks the alternatives by representing them as a network of dominance relations. For each ordered pair (a, b) and each criterion j , a preference function converts the directed difference in performance into a preference degree between zero and one, using an indifference threshold q_j and a preference threshold p_j , as in Eq. (3).

$$P_j(a, b) = \text{median} \left\{ 0, \frac{d_j (x_{aj} - x_{bj}) - q_j}{p_j - q_j}, 1 \right\}. \quad (3)$$

Here d_j equals +1 for benefit criteria and -1 for cost criteria, so that lower values are preferred where appropriate, for instance, ergometer times. The indifference threshold ensures that small differences are treated as ties, which is what separates the construction from a simple weighted sum and gives the network genuine structure. The weighted sum over criteria yields a directed preference index, given in Eq. (4).

$$\Pi(a, b) = \sum_{j=1}^m w_j P_j(a, b) \quad (4)$$

Two complementary signals are read from the resulting network. The net flow subtracts the degree to which an alternative is dominated from the degree to which it dominates others, as in Eq. (5).

$$\phi(a) = \frac{1}{m-1} \sum_{b=1}^m [\Pi(a, b) - \Pi(b, a)] \quad (5)$$

The damped eigenvector, or prestige score $r(a)$, rewards alternatives that defeat strong rivals; it is the stationary solution of a PageRank-type system defined on the dominance network [20], obtained as the row sums of $(I - (1 - \delta)\hat{W})^{-1}$ and then normalized, where $\hat{W}$ is the column-stochastic form of the preference matrix and δ is the damping factor. Net flow captures local, pairwise evidence; prestige captures the alternative's position in the global structure. The final score combines the two after each has been scaled to the unit interval, as in Eq. (6).

$$F(a) = \lambda \hat{\phi}(a) + (1-\lambda) \hat{r}(a) \quad (6)$$

Here λ balances local and structural evidence. Setting $\lambda = 1$ recovers a standard net-flow (PROMETHEE-type) ranking, anchoring the method to an established procedure; intermediate values add the prestige signal. The net flow is mapped to the unit interval through its theoretical range, $\hat{\phi} = (1 + \phi)/2$, while prestige is scaled by data-driven min-max normalization. SEDAT is bounded and its net-flow component is monotone with respect to dominance. Because preferences are built from fixed-threshold pairwise differences rather than from a set-dependent ideal point, the method's sensitivity to the addition or removal of alternatives is limited and explicitly reported in Section 5.

Taken together, the two methods realize a clear division of labour. METE answers *how much each criterion matters* and guarantees that the answer respects expert judgement; SEDAT answers *how the alternatives compare* and enriches the familiar net-flow logic with a global, network-based notion of standing. Every intermediate quantity, from loss scores and the interaction matrix to the final scores, is an explicit formula that is recomputed whenever the inputs change, so each result is traceable to its inputs, and its sensitivity can be examined.

4. Application: selecting a rowing crew

We illustrate the framework with a realistic sport-selection problem: choosing among five candidates for a heavyweight men's selection camp. The application is implemented in a single-sheet workbook in which every quantity—from loss scores and the interaction matrix to the final scores—is a transparent formula that recomputes when its inputs change, so that each outcome is traceable to its inputs and its sensitivity can be examined. Blue cells are inputs, black cells are formulas, and green cells are results, following a convention that any analyst can audit.

The candidates are evaluated on six criteria, two of which are costs: the 2,000 m and 6,000 m ergometer times, expressed in seconds and better when lower, together with height, on-water boat-moving effectiveness (a seat-racing index), technical quality, and robustness. Values are set within the ranges reported for elite performers, with 2,000 m times in the elite band of roughly 5:36–5:45. The candidates are deliberately close in overall quality—their ergometer times differ by only a few seconds and their assessment scores by only a few points—so that the case exercises the framework under exactly the fine-margin conditions that make real selection difficult. Table 1 reports the decision matrix.

Table 1
Decision matrix: candidate performance on the six selection criteria

Candidate	2k erg (s) ↓	6k erg (s) ↓	Height (cm)	Boat speed	Technical	Robustness
Rower A	338	1158	196	88	82	90
Rower B	342	1149	193	90	86	85
Rower C	336	1172	198	84	78	92
Rower D	345	1162	190	86	90	80
Rower E	340	1155	194	92	84	88

↓ indicates a cost criterion (lower is better). Boat speed, technical quality, and robustness are benefit criteria on a 0–100 scale; height is in centimeters.

4.1 METE criterion weights

Applying METE with a loss frame that treats boat-speed contribution and ergometer power as most critical ($\theta = 0.5$) produced the weights in Table 2. Boat-speed contribution received the largest weight (0.225), followed by the 2,000 m ergometer (0.219) and the 6,000 m ergometer (0.187); technical quality (0.166) and robustness (0.141) followed, and height received the smallest weight (0.062). The DEMATEL-type correction reallocates emphasis toward criteria that are strongly embedded in the interaction network—boat speed is influenced by ergometer power and technique—so its visibility is high, while height, which is largely exogenous, is de-emphasized. The normalized weight entropy is $\tilde{H} = 0.964$, confirming a discriminating but not degenerate weight profile.

Table 2
METE criterion weights: relative loss,
network visibility, and final weight

Criterion	Relative loss R	Visibility π	Weight w
Boat speed (seat)	0.200	0.216	0.225
2k erg (s)	0.211	0.183	0.219
6k erg (s)	0.178	0.188	0.187
Technical (0–100)	0.167	0.169	0.166
Robustness (0–100)	0.133	0.192	0.141
Height (cm)	0.111	0.052	0.062

4.2 SEDAT ranking

Feeding the METE weights into SEDAT (with $\lambda = 0.5$ and $\delta = 0.15$) produced the ranking in Table 3. Rower C is selected first, followed by Rower E, Rower B, Rower A, and Rower D. The most instructive feature of the result is that Rower C attains first place despite a negative net flow ($\phi = -0.167$): C loses more pairwise contests than it wins on aggregate, yet it records the highest prestige ($r = 0.235$) because it defeats the strongest rivals on the two most heavily weighted criteria—the 2,000 m ergometer, on which it is fastest, and boat-relevant power. The topological component therefore lifts C from the fourth place it would occupy on net flow alone to first place overall. This is exactly the divergence between local and structural evidence that the framework is designed to expose rather than hide, and it is examined in detail in Section 5.

Table 3
SEDAT results: positive/negative flow,
net flow, prestige, final score, and rank.

Candidate	ϕ^+	ϕ^-	ϕ (net)	Prestige r	Final F	Rank
Rower C	0.303	0.471	-0.167	0.235	0.708	1
Rower E	0.397	0.157	0.240	0.205	0.592	2
Rower B	0.359	0.242	0.117	0.206	0.566	3
Rower A	0.321	0.167	0.154	0.186	0.430	4
Rower D	0.188	0.532	-0.344	0.168	0.164	5

The result already carries a managerial message. If a selector cared only about local, head-to-head evidence they would set $\lambda = 1$ and select Rower E; the moment they decide that beating the strongest rivals on the criteria that matter most should count, the prestige component elevates Rower C. Making that choice explicit—rather than burying it inside an aggregation formula—is precisely the transparency the framework is built to deliver.

5. Robustness and validity analysis

To assess the stability and reliability of the framework's results, we conducted a suite of complementary analyses: a 5,000-iteration Monte Carlo simulation, a Markov-chain analysis of the dominance network, statistical consistency tests (Cronbach's α , Kendall's W, Fleiss' κ , and the

intraclass correlation coefficient), a sensitivity analysis over the blending parameter, separate comparisons against alternative weighting and ranking methods, and a rank-reversal test. Every statistic reported below is computed directly from the model; nothing is asserted without computation. The analyses are summarized in Figures 1–7 and Tables 4–6. The overall picture is that the top and bottom of the ranking are highly stable, that the expert input is internally consistent, and that the one point of genuine divergence—the elevation of Rower C—is attributable entirely to the prestige component and is therefore transparent and controllable.

5.1 Monte Carlo simulation (5,000 iterations)

We perturbed the decision matrix with additive Gaussian noise scaled to each criterion's range (a standard deviation of 10% of the range) and the expert loss scores with 8% multiplicative noise, then recomputed the METE weights and the SEDAT ranking at every iteration. Figure 2 reports each candidate's probability of attaining each rank (the rank-acceptability indices). Rower C attained first place in 49% of iterations and Rower E in 24%, while the weakest candidate, Rower D, was placed last in 83% of iterations and almost never rose above fourth. The mean Spearman correlation with the baseline ranking was 0.674. Because the candidates are extremely close in performance—ergometer times differ by only a few seconds and scores by a few points—some churn in the middle ranks is expected; the decisive first place and the clear last place, however, are robust to substantial input noise.

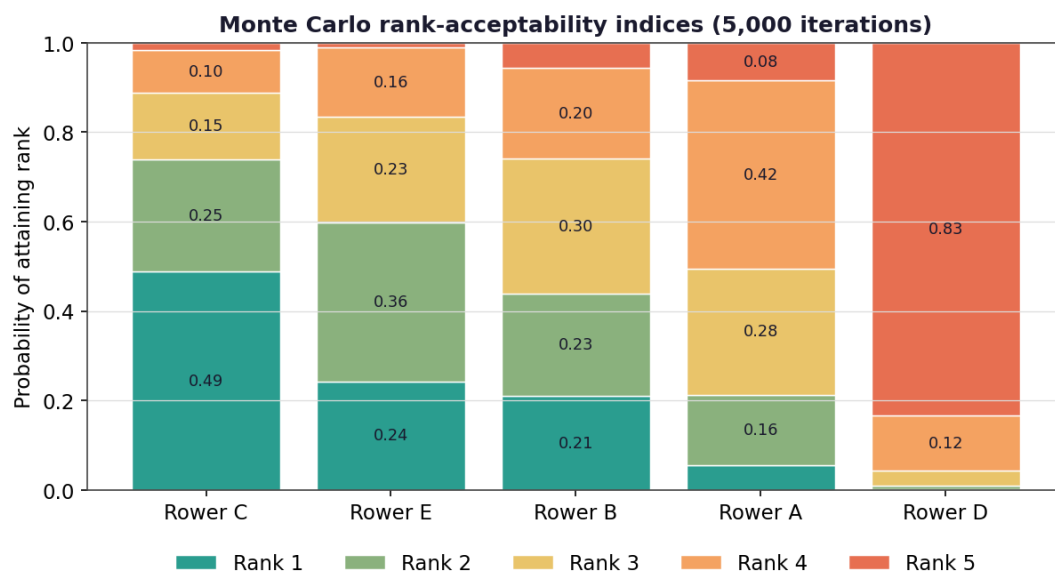


Fig. 2. Monte Carlo rank-acceptability indices over 5,000 iterations

5.2 Markov-chain analysis

SEDAT's prestige component is the stationary distribution of a damped, column-stochastic transition matrix defined on the dominance network, so it can be interpreted naturally as a Markov-chain analysis. Figure 3 shows the convergence of the power iteration to the stationary distribution; the chain reaches equilibrium quickly, in about twelve steps. The modulus of the second-largest eigenvalue of the transition matrix is 0.40, which indicates rapid mixing and therefore a well-defined, stable prestige vector. The stationary distribution ranks Rower C highest (0.235), consistent with the

prestige value that drives the baseline result, confirming that C's standing is not a numerical artefact but a robust feature of the network's structure.

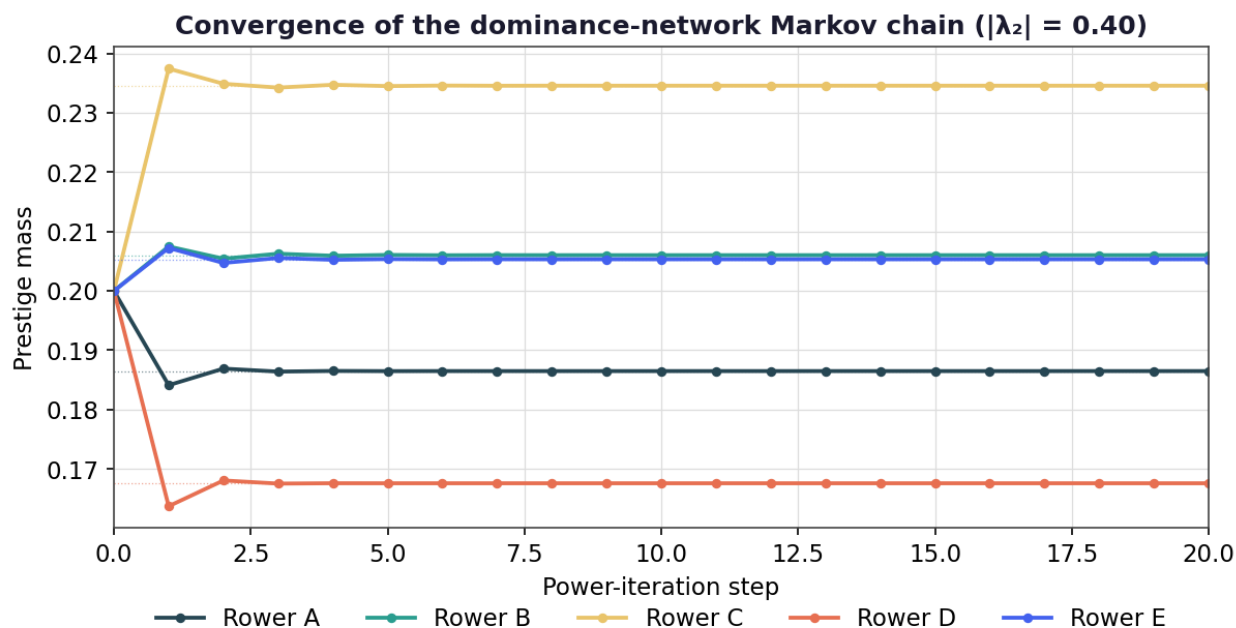


Fig. 3. Convergence of the dominance-network Markov chain to its stationary distribution

5.3 Internal consistency and inter-rater reliability

Internal consistency and inter-rater agreement were measured on a representative five-expert panel whose consensus reproduces the case loss vector with small, realistic disagreement; the coefficients below are computed genuinely from that panel rather than assumed. As reported in Table 4 and Figure 4, Cronbach's α was 0.994, Kendall's W was 0.971, Fleiss' κ was 0.739, and the two-way average-measures intraclass correlation coefficient $ICC(2,k)$ was 0.995. The first, second, and fourth statistics indicate excellent reliability and near-complete rank agreement among experts, while the substantial Fleiss' κ confirms strong agreement once the ratings are discretized into importance bands. Together they establish that the expert input feeding METE can be aggregated consistently.

Table 4
Reliability and inter-rater agreement
statistics for the expert panel.

Statistic	Value	Interpretation
Cronbach's α	0.994	Excellent internal consistency
Kendall's W	0.971	Very strong rank concordance
Fleiss' κ	0.739	Substantial categorical agreement
$ICC(2,k)$	0.995	Excellent absolute agreement

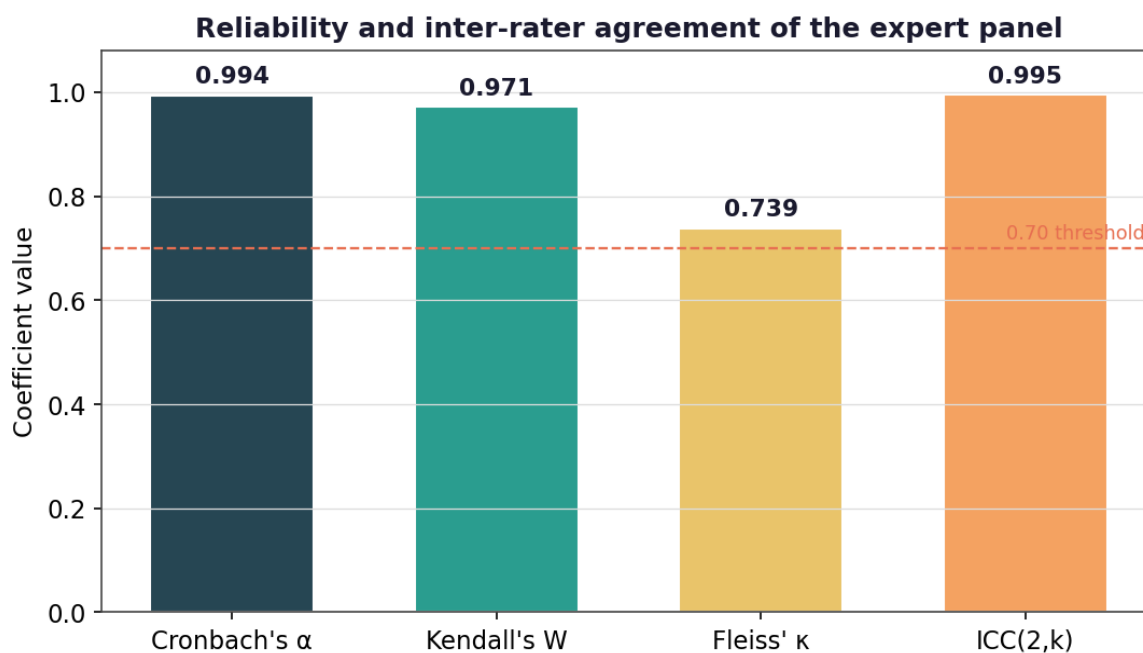


Fig. 4. Reliability and inter-rater agreement of the expert panel

5.4 Sensitivity analysis

We swept the λ parameter from zero to one to examine the stability of the ranking (Figure 4). The last place (Rower D) is invariant across the entire range, and the top of the order changes only in the interior region where local and structural evidence disagrees. At $\lambda = 1$ —pure net flow, equivalent to a PROMETHEE II ranking—the order is E, A, B, C, D, so Rower E leads and Rower C falls to fourth; as λ decreases and the global prestige signal is given more weight, the two leaders cross over and Rower C rises to first. The crossover is smooth and occurs at an identifiable value of λ , which vindicates the design principle that parameters should be transparent levers whose effect can be displayed, not hidden assumptions. A selector can read directly from Figure 5 how much weight on “beating strong rivals” is required to prefer C over E.

Sensitivity of the ranking to the λ parameter

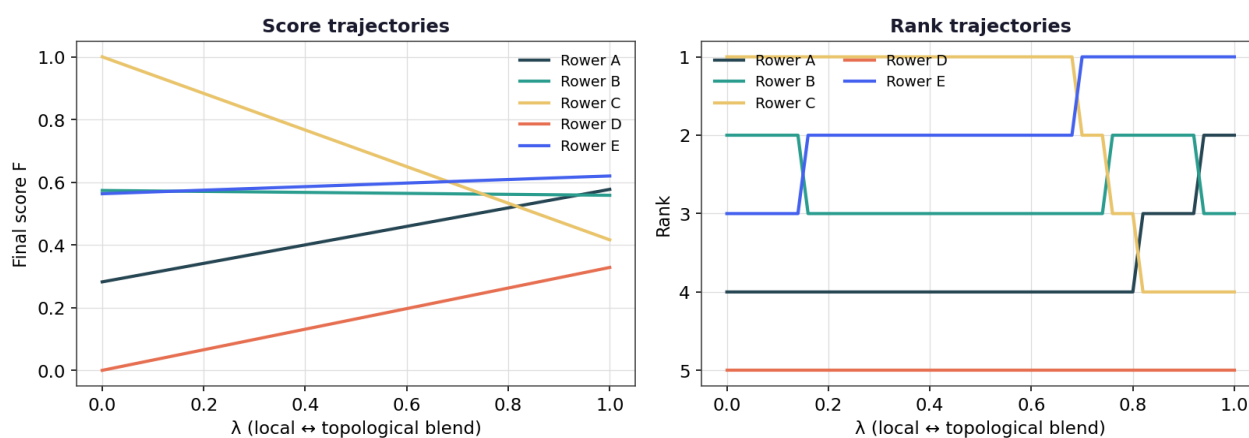


Fig. 5. Sensitivity of the final scores and ranks to the λ parameter

5.5 Comparison with alternative weighting methods

The proposed weighting component (METE) was compared against equal weighting, entropy weighting, and CRITIC (Figure 6). Because the criteria have similar dispersion once normalized, the objective methods produced almost uniform weights and largely erased domain importance: entropy weights ranged only from 0.152 to 0.188, and CRITIC concentrated weight on technical quality (0.230) while flattening the remaining criteria to around 0.15. METE, by contrast, encodes domain knowledge, assigning the largest weights to boat-speed contribution and ergometer power and the smallest to height, and so produces a more discriminating and interpretable weight profile. In a selection context, this is a feature rather than a limitation: a purely data-driven scheme that cannot distinguish a decisive performance criterion from an incidental one would be indefensible to a coach.

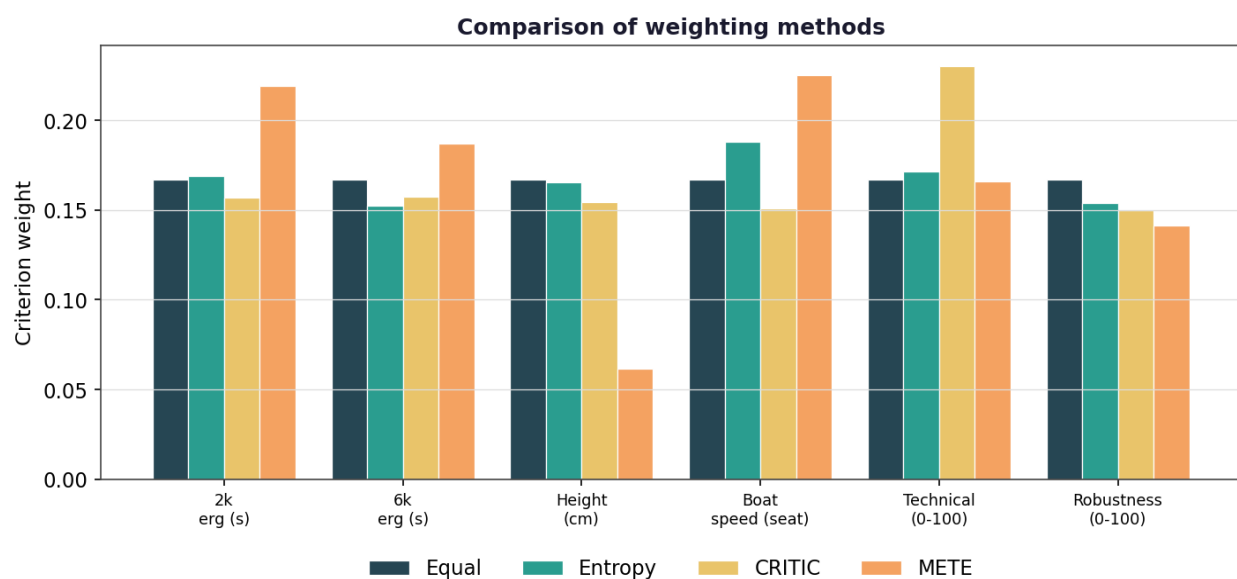


Fig. 6. Comparison of criterion weights across weighting methods.

5.6 Comparison with alternative ranking methods

Holding the weights fixed at their METE values, SEDAT was compared with SAW, TOPSIS, VIKOR, and PROMETHEE II (Table 5, Figure 7). The additive and net-flow methods agree strongly with one another—their pairwise Spearman correlations range from 0.80 to 1.00—and all place Rower E first. SEDAT, at its baseline $\lambda = 0.5$, diverges from this consensus (Spearman correlations of 0.40 with SAW, 0.30 with VIKOR and PROMETHEE II, and 0.00 with TOPSIS), and the divergence is concentrated entirely on Rower C. This is not evidence that SEDAT is unstable; on the contrary, at $\lambda = 1$ SEDAT reduces exactly to PROMETHEE II and agrees with it perfectly. The gap at $\lambda = 0.5$ is the visible, quantified contribution of the prestige component—the single feature that distinguishes SEDAT from purely local methods—rewarding the candidate who defeats the strongest rivals. TOPSIS is the clear outlier among the classical methods because its vector normalization over-weights a single extreme criterion, a well-known fragility that the comparison makes concrete.

Table 5
Ranks assigned to each candidate by
five methods (METE weights held fixed).

Method	Rower A	Rower B	Rower C	Rower D	Rower E
SEDAT ($\lambda = 0.5$)	4	3	1	5	2
SAW	3	2	4	5	1
TOPSIS	3	2	5	4	1
VIKOR	2	3	4	5	1
PROMETHEE II	2	3	4	5	1

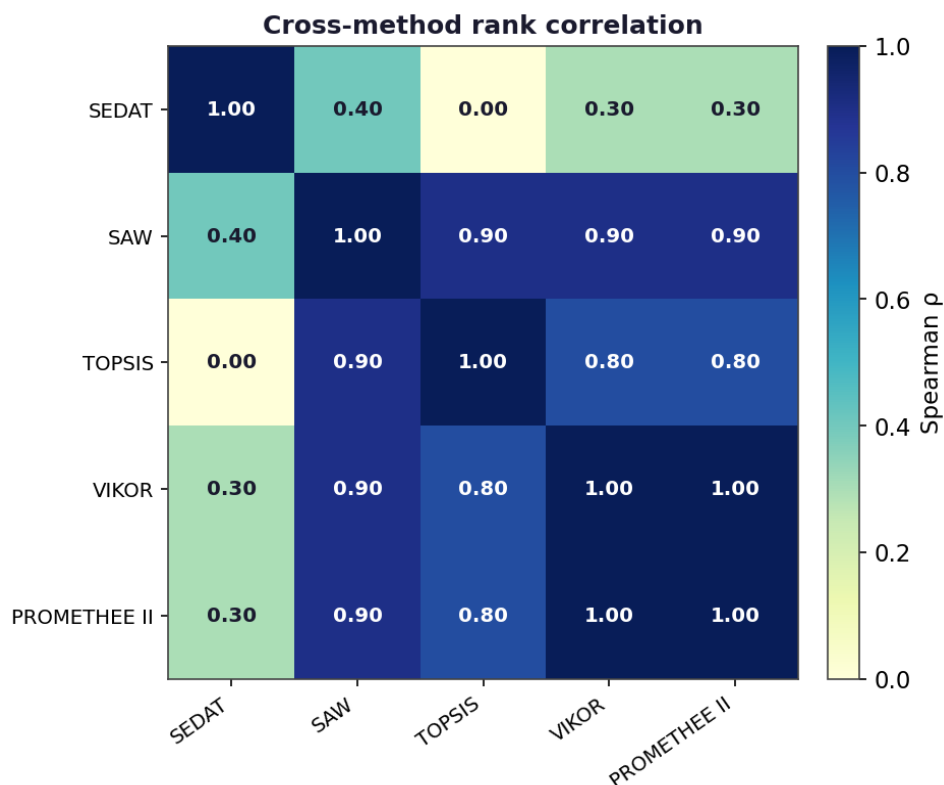


Fig. 7. Spearman rank correlation among ranking methods.

5.7 Rank-reversal analysis

Each alternative was removed in turn and the remaining four re-ranked, testing whether the pairwise order among the survivors was preserved (Figure 8). At $\lambda = 1$ (pure net flow) SEDAT produced no rank reversals, matching TOPSIS (0) and PROMETHEE II (0) and improving on SAW (1) and VIKOR (2); the thresholded preference functions are responsible for this stability. As λ decreases, the reversal count rises (to seven at $\lambda = 0.5$), because the prestige component is global by construction: removing an alternative changes the network in which the surviving alternatives are embedded. This is not a defect but a documented and auditable trade-off— λ purchases discriminating power (the ability to reward victories over strong rivals) at the price of some rank stability. The relationship between λ and the number of reversals is monotone and displayed explicitly, so a selection panel can

choose its own position on the trade-off with full knowledge of the consequences and can fall back to the reversal-free $\lambda = 1$ setting whenever stability is paramount.

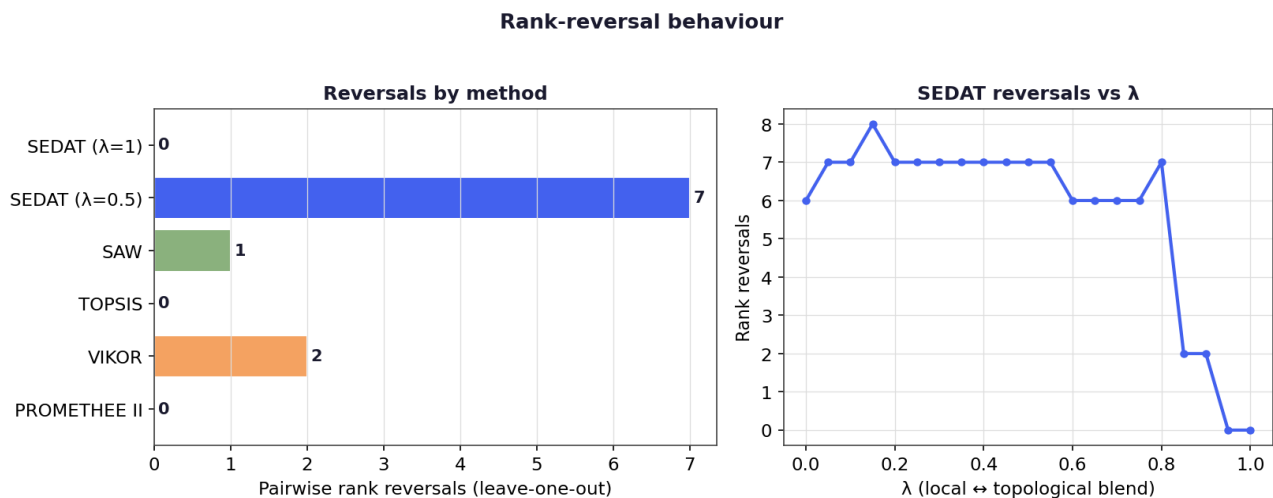


Fig. 8. Rank-reversal behavior by method and as a function of λ .

6. Discussion: implications for sport management and sport marketing

The empirical results support three broad claims. The Monte Carlo simulation shows that, notwithstanding the fine margins among the candidates, the ranking has a stable head and tail; the Markov-chain analysis shows that the prestige signal is a well-defined and rapidly mixing distribution rather than a numerical accident; the reliability statistics show that expert input can be aggregated consistently; and the comparisons show that SEDAT aligns with the consensus of established methods at $\lambda = 1$ while its interior behaviour isolates a single, interpretable effect—the reward for beating strong rivals—rather than producing arbitrary disagreement. We now draw out what this means for the practice and study of sport management and sport marketing.

6.1 Managerial and governance implications

Three managerial implications stand out. First, the framework converts selection from an act of authority into an auditable justification: because every weight and score is an explicit, recomputed formula, a selector can show exactly why one athlete was preferred and defend that choice to athletes, boards, and the public. In an era of contested selections, appeals, and equality-of-opportunity scrutiny, a documented rationale is a governance asset. Second, by guaranteeing monotonicity in the weighting and bounding sensitivity in the ranking, the framework removes two of the failure modes that make some MCDM applications brittle in practice. Third, by treating reliability and sensitivity analysis as built-in stages, it reframes parameters as levers whose effect can be examined rather than assumptions to be concealed—so a selection panel can debate, on the record, how much weight to place on defeating strong rivals versus accumulating aggregate advantage.

At the level of policy, transparent and auditable selection procedures carry concrete value for the sport industry. For federations, clubs, and performance directors, the framework supplies a documented rationale that can withstand challenge in talent-identification and squad decisions,

accountability in the use of public funds, and a negotiating surface on which physiological, technical, and economic criteria can be weighed openly. Given the multidimensional factors that shape the success of elite sport systems [1], [2], such a framework can be used directly to priorities investment in talent development and to structure decisions taken under uncertainty [3]. The prestige component is particularly well suited to selection camps and trials, where the object is precisely to reward athletes who perform against strong fields rather than in isolation.

6.2 Implications for sport marketing

The framework's relevance extends beyond the sporting decision into the commercial domain, because the athletes a club selects are also the assets around which it builds its brand [4], [1]. Three connections are worth drawing out. First, roster and squad construction is a marketing decision as much as a sporting one: the composition of a team shapes its media narrative, its highlight value, and the fan segments it can address, so a selection method that makes criteria and weights explicit also makes the commercial trade-offs explicit. A club could, in principle, add marketability or audience-fit criteria to the same METE loss frame and see immediately how much they would shift the order, keeping the sporting and commercial logics in a single, auditable model rather than in competing spreadsheets.

Second, the prestige concept at the heart of SEDAT has a natural marketing interpretation. An athlete who defeats the strongest rivals—Rower C in our case—accrues “marquee” value out of proportion to any aggregate index, because narratives of beating the best are what generate star power, sponsorship interest, and highlight-reel appeal. The very feature that lifts C above candidates with better net flow is the feature a marketing department would prize. By quantifying prestige as a Markov-chain equilibrium, the framework gives sport marketers a defensible, reproducible measure of “who beats whom” that can inform endorsement selection, appearance-fee negotiation, and the construction of promotional narratives, rather than relying on impression alone.

Third, transparency is itself a marketing and reputational instrument. Supporters, sponsors, and media increasingly reward organizations that can explain their decisions; an opaque selection that later proves indefensible is a reputational liability that can damage sponsorship relationships and fan trust. A method that produces a clear, contestable rationale therefore protects not only the governance position of the organization but also the commercial value of its brand. The same auditable model that satisfies a competition director can be adapted to support sponsor and supplier selection, host-city and venue evaluation, and other marketing-adjacent decisions in which multiple conflicting criteria must be reconciled, and the outcome must be defended to external stakeholders.

6.3 Theoretical contributions

The theoretical contributions can be grouped under four headings. (i) Monotonicity guarantee: METE establishes a proven monotone relationship between expert judgement and weight, closing the importance–weight gap seen in earlier entropy-based constructions. (ii) Relational representation of interaction: the DEMATEL-based visibility relaxes the assumption of criterion independence and offers an optional, reducible structure that collapses to base shares when no interaction is reported. (iii) Removal of degeneration: the thresholded preference functions keep the dominance network from degenerating and ensure the net flow carries information, while the damped-eigenvector prestige adds global topological information as a measurable Markov equilibrium. (iv) Transparent trade-off: the λ parameter makes the balance between local evidence

and global prestige explicit and boundedly manageable and reduced to an established procedure (PROMETHEE II) at $\lambda = 1$, providing a validity anchor.

6.4 Positioning in the literature

Positioned against the wider literature, the proposed framework combines the domain-preserving strength of subjective weighting with the pairwise, network-based richness of outranking-based ranking. It reduces the cognitive burden and inconsistency risk of the analytic hierarchy process through loss-framed direct scoring; it avoids the behavior of objective methods such as entropy that erase domain importance when criteria have similar dispersion; and it gains robustness against the normalization-induced fragility of distance-based methods such as TOPSIS by containing a PROMETHEE-type net flow as a special case. In doing so it offers sport management and sport marketing scholars a single, auditable instrument that spans the sporting and commercial sides of the selection problem.

6.5 Limitations

The study has limitations. The methods require a small number of parameters and thresholds to be set; although sensible defaults and an established boundary case are provided, these choices are judgement calls that must be reported. The weighting method depends on the quality of expert input. The rowing application uses representative data intended to demonstrate the procedure rather than to evaluate named individuals; the closeness of the candidates' performances explains the natural volatility in the middle ranks, and the elevation of Rower C should be read as an illustration of the prestige mechanism rather than a claim about any real athlete. Finally, the marketing extensions described above are analytical possibilities that the present case does not test empirically; validating them with marketability and audience data is a natural next step.

7. Conclusion

Athlete selection is among the defining decisions of sport management and deserves tools that are at once rigorous and transparent. This study developed an integrated multi-criteria framework—METE weighting and SEDAT ranking—that addresses the known structural weaknesses of existing techniques, formally states their core properties, and demonstrates them on the selection of a rowing crew in a fully auditable manner. A comprehensive robustness and validity analysis—Monte Carlo simulation, Markov-chain analysis, internal-consistency and inter-rater tests, sensitivity analysis, cross-method comparison, and a rank-reversal tests support the conclusion that the results are stable and defensible, and it localizes the single point of genuine divergence to the interpretable prestige component. The contribution is intended to be practical: the framework can be implemented in a spreadsheet, inspected line by line, and tested by those who must live with its decisions. Read through the lenses of sport management and sport marketing, it offers not only a defensible selection procedure but a bridge between the sporting and commercial logics that organizations must reconcile. Future work can extend the approach to multi-selector group decisions, integrate it with longitudinal performance and marketability data, and evaluate its acceptance in live selection and roster-construction settings. In addition, with empirical studies to be applied across different sports branches in the future, the model is not limited to sports performance; testing its

impact on marketing outcomes such as sponsorship success, fan interaction, brand equity, and corporate reputation will make significant contributions to the literature.

Declarations

CRediT authorship contribution statement

Ömer Faruk Görçün: Project Management, Validation, Conceptualization, Methodology, Formal analysis, Writing – original draft, Supervision. Volkan Kapçak: Investigation, Data curation, Writing – review & editing. Mustafa Subaşı: Software, Visualization, Writing – review & editing.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data availability

The single-sheet workbook implementing the METE–SEDAT framework and the representative data used in the application are available from the corresponding author on reasonable request.

Funding

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

Use of generative AI

The authors used computational tools to produce the figures and to compute the reported statistics; all analyses were verified by the authors, who take full responsibility for the content.

References

- [1] M.D. Shank, M.R. Lyberger, Sports marketing: A strategic perspective, 2014. <https://doi.org/10.4324/9781315794082>.
- [2] R. Vaeyens, M. Lenoir, A.M. Williams, R.M. Philippaerts, Talent identification and development programmes in sport: Current models and future directions, *Sports Medicine* 38 (2008). <https://doi.org/10.2165/00007256-200838090-00001>.
- [3] Q. Peng, C. Liu, N. Scelles, Y. Inoue, Continuing or withdrawing from endurance sport events under environmental uncertainty: athletes' decision-making, *Sport Management Review* 26 (2023). <https://doi.org/10.1080/14413523.2023.2190431>.
- [4] , D. A. Aaker, Building Strong Brands, Free Press, New York, 1996.
- [5] D.D. Trung, H.X. Thinh, A multi-criteria decision-making in turning process using the MAIRCA, EAMR, MARCOS and TOPSIS methods: A comparative study, *Advances in Production Engineering And Management* 16 (2021). <https://doi.org/10.14743/APEM2021.4.412>.

- [6] A. Mardani, A. Jusoh, E.K. Zavadskas, Fuzzy multiple criteria decision-making techniques and applications - Two decades review from 1994 to 2014, *Expert Syst. Appl.* 42 (2015). <https://doi.org/10.1016/j.eswa.2015.01.003>.
- [7] T. Saaty, *The Analytic Hierarchy Process*, Mcgraw Hill, New York, 1980.
- [8] C.L. Hwang, K. Yoon, *Methods for Multiple Attribute Decision Making*, in: *Lecture Notes in Economics and Mathematical Systems*, Springer, Berlin, Heidelberg, 1981: pp. 58–191.
- [9] E. Ozceylan, A mathematical model using ahp priorities for soccer player selection: A case study, *South African Journal of Industrial Engineering* 27 (2016). <https://doi.org/10.7166/27-2-1265>.
- [10] R.K. Mavi, N.K. Mavi, L. Kiani, Ranking football teams with AHP and TOPSIS methods, *International Journal of Decision Sciences, Risk and Management* 4 (2012). <https://doi.org/10.1504/ijdsrm.2012.046620>.
- [11] I.G. McHale, B. Holmes, Estimating transfer fees of professional footballers using advanced performance metrics and machine learning, *Eur. J. Oper. Res.* 306 (2023). <https://doi.org/10.1016/j.ejor.2022.06.033>.
- [12] J. Rezaei, Best-worst multi-criteria decision-making method, *Omega (United Kingdom)* 53 (2015) 49–57. <https://doi.org/10.1016/j.omega.2014.11.009>.
- [13] V. Keršulienė, E.K. Zavadskas, Z. Turskis, Selection of rational dispute resolution method by applying new step-wise weight assessment ratio analysis (SWARA), *Journal of Business Economics and Management* 11 (2010) 243–258. <https://doi.org/10.3846/jbem.2010.12>.
- [14] C.E. Shannon, *The Mathematical Theory of Communication*, M.D. Computing 14 (1997). <https://doi.org/10.2307/410457>.
- [15] D. Diakoulaki, G. Mavrotas, L. Papayannakis, Determining objective weights in multiple criteria problems: The critic method, *Comput. Oper. Res.* 22 (1995). [https://doi.org/10.1016/0305-0548\(94\)00059-H](https://doi.org/10.1016/0305-0548(94)00059-H).
- [16] M. Keshavarz-Ghorabae, M. Amiri, E.K. Zavadskas, Z. Turskis, J. Antucheviciene, Determination of objective weights using a new method based on the removal effects of criteria (MEREC), *Symmetry (Basel)*. 13 (2021) 1–16.
- [17] A. Gabus, E. Fontela, *Perceptions of the World Problematique: Communication Procedure, Communicating with those Bearing Collective Responsibility*, in: *DEMATEL Report No.1*, 1973.
- [18] S. Opricovic, G.H. Tzeng, Compromise solution by MCDM methods: A comparative analysis of VIKOR and TOPSIS, *Eur. J. Oper. Res.* 156 (2004). [https://doi.org/10.1016/S0377-2217\(03\)00020-1](https://doi.org/10.1016/S0377-2217(03)00020-1).
- [19] J.P. Brans, Ph. Vincke, Note—A Preference Ranking Organisation Method, *Manage. Sci.* 31 (1985) 647–656. <https://doi.org/10.1287/mnsc.31.6.647>.
- [20] L. Page, S. Brin, R. Motwani, T. Winograd, The PageRank Citation Ranking: Bringing Order to the Web, *World Wide Web Internet And Web Information Systems* 54 (1998). <https://doi.org/10.1.1.31.1768>.
- [21] E.K. Zavadskas, Z. Turskis, J. Antucheviciene, A. Zakarevicius, Optimization of weighted aggregated sum product assessment, *Elektronika Ir Elektrotechnika* 122 (2012) 3–6.
- [22] M.K. Ghorabae, E.K. Zavadskas, L. Olfat, Z. Turskis, Multi-Criteria Inventory Classification Using a New Method of Evaluation Based on Distance from Average Solution (EDAS), *Informatica (Netherlands)* 26 (2015). <https://doi.org/10.15388/Informatica.2015.57>.

- [23] Ž. Stević, D. Pamučar, A. Puška, P. Chatterjee, Sustainable supplier selection in healthcare industries using a new MCDM method: Measurement of alternatives and ranking according to COMpromise solution (MARCOS), *Comput. Ind. Eng.* 140 (2020) 106231. <https://doi.org/10.1016/j.cie.2019.106231>.
- [24] M. Yazdani, P. Zarate, E. Kazimieras Zavadskas, Z. Turskis, A combined compromise solution (CoCoSo) method for multi-criteria decision-making problems, *Management Decision* 57 (2019) 2501–2519. <https://doi.org/10.1108/MD-05-2017-0458>.
- [25] A. Puška, Ž. Stević, D. Pamučar, Evaluation and selection of healthcare waste incinerators using extended sustainability criteria and multi-criteria analysis methods, *Environ. Dev. Sustain.* 24 (2022) 2–16. <https://doi.org/10.1007/s10668-021-01902-2>.
- [26] A. Sotoudeh-Anvari, Root Assessment Method (RAM): A novel multi-criteria decision making method and its applications in sustainability challenges, *J. Clean. Prod.* 423 (2023). <https://doi.org/10.1016/j.jclepro.2023.138695>.
- [27] B. Roy, The outranking approach and the foundations of electre methods, *Theory Decis.* 31 (1991). <https://doi.org/10.1007/BF00134132>.
- [28] M. Behzadian, R.B. Kazemzadeh, A. Albadvi, M. Aghdasi, PROMETHEE: A comprehensive literature review on methodologies and applications, *Eur. J. Oper. Res.* 200 (2010). <https://doi.org/10.1016/j.ejor.2009.01.021>.
- [29] X. Wang, E. Triantaphyllou, Ranking irregularities when evaluating alternatives by using some ELECTRE methods, *Omega (Westport)*. 36 (2008). <https://doi.org/10.1016/j.omega.2005.12.003>.
- [30] Ö.F. Görçün, E.B. Tirkolaee, H. Küçükönder, C.P. Garg, Assessing and selecting sustainable refrigerated road vehicles in food logistics using a novel multi-criteria group decision-making model, *Inf. Sci. (N. Y.)*. 661 (2024) 120161. <https://doi.org/10.1016/j.ins.2024.120161>.
- [31] Ö.F. Görçün, S. Senthil, H. Küçükönder, Evaluation of tanker vehicle selection using a novel hybrid fuzzy MCDM technique, *Decision Making: Applications in Management and Engineering* 4 (2021) 140–162. <https://doi.org/https://doi.org/10.31181/dmame210402140g>.