

Toward A Robust and Learning-Based Decision Support Model: Hybrid Preference Preservation Inverse Optimization-AHP Framework

Afaf Dadda ^{1,*}, Nada Baddou ², Rokaya Douiri ¹, Latifa Ezzine ¹

¹Laboratory of Modeling and Optimization of Industrial Systems, University Moulay Ismail, Morocco

²Laboratory of Industrial Techniques, Sidi Mohamed Ben Abdellah University, FST Fez, Morocco

Abstract Industry 5.0 enhances Industry 4.0 paradigm by directing the evolution to more resilient, human-centric and sustainable ecosystem. In this context, Decision Makers(DM) require the use of a more reliable Decision Support Model(DSM). Multi-Criteria Decision Making(MCDM) framework offers great and consistent support. More recently, owed to the advance in Artificial Intelligence (AI) techniques, several researchers try improving the classical MCDM frameworks by proposing hybrid AI-MCDM models. These models can identify “the best decision” based on prediction, using historical decisions Data. Nevertheless, two major limitations are listed: (i) in the case of small data set the AI prediction models are not recommended. (ii) where transparency and trust are crucial, these models can be perceived as “black boxes” with low interpretability or explicability. To overcome these limitations, this study proposes a new collaborative framework, that incorporates a hybrid MCDM model based on “Inverse optimization (IO) and Human-In-The-Loop (HITL) validation”, adequate for various dataset sizes, gives a better explicability and improves the Decision Trust. In this paper, the improved IO-MCDM framework is applied to one of the famous MCDM method “Analytic Hierarchic Process (AHP)”. Entitled “Preference Preservation Inverse Optimization-AHP (PPIO-AHP)”. The proposed framework, in addition to enhancing the learning based on historical previous best decisions, improves AHP consistency and Preference Preservation Mechanism. Finally, a “sustainable project evaluation” is established. The result emphasizes the transformative role of “human-Inverse optimization-MCDM” collaboration, in progressing Industry 5.0’s vision of smart, sustainable, and human-centered decision making.

Keywords Inverse Optimization (IO), Multi-Criteria Decision-Making (MCDM), Analytic Hierarchy Process (AHP), Human-In-The-Loop (HITL), Decision Support System (DSS), Decision Trust Index (DTI)

DOI: 10.19139/soic-2310-5070-4111

1. Introduction

The assessment of industrial projects has multifaceted aspects, frequently including several constraints (technical, economic, environmental, etc.). Throughout the project lifecycle, various decisions are made and stakeholders should successfully choose the appropriate parameters [1] and [2]. Selecting the appropriate decision (Alternatives, choices, etc. . .) is a complex multi-criteria decision making process. In fact, since its emergence in the 1970s until now, various scientific studies [2], [3], [4] and [5], have addressed their decision problems by the use of MCDM methods, due to their simplicity, adaptability and their potential for hybridization with other methods. Despite its wide application, the challenge of handling dynamic decision-making environment and the exponential growth in pairwise comparisons, introducing a degree of inconsistency and complexity, remains underexplored and very limited [6], [7] and [36]. To overcome these limits, researchers combine MCDM and Artificial Intelligence (AI) techniques (such as machine learning (ML) and/or deep Learning (DL)) to improve their ability to analyze and

*Correspondence to: Afaf Dadda (Email: a.dadda@ensam.umi.ac.ma). Laboratory of Modeling and Optimization of Industrial Systems, University Moulay Ismail, Morocco.

predict the best decision using historical decisions [8], [9] and [10]. In the major cases this Hybrid AI-MCDM models can be perceived as “black boxes” lacking interpretability or explainability, and which may thus limit their adoption for strategic decision making problems [11] and [37]. Thus, hence the need for more explainable decision models adapted to strategic background, including the multi-criteria measurement and able to strengthen trust, auditability and acceptance of decision-making actors. Recent progress in Inverse Optimization (IO) [12], [13], [14] and Robust Optimization [15] are promising and may help overcome the classical limits of MCDM and hybrid AI-MCDM. Historically, IO is used in optimization problems to infer the objective function from observed optimal solutions [14]. This study presents a new collaborative framework in line with the principles of the industry 5.0-namely human centricity [25], [26] and [27], sustainability, resilience, and collaboration between humans and Decision Support Systems (DSS) in adaptive environments. From a practical point of view, this study proposes a learning and robust decision making framework, the following research questions are addressed:

RQ1: How can IO be used to overcome the limitations of existing MCDM approaches, learn from previous best decisions, and thereby define a Decision Trust Index (DTI) to strengthen the decision-maker’s confidence?

RQ2: How can a “human in the loop architecture” be structured within an inverse optimization-based MCDM framework to support expert validation, interpretation, and iterative learning?

This paper is structured as follows; Section 2 presents recent MCDM and inverse-optimization studies, then reviews the Human-in-the Loop concept for an interactive decision process. Section 3 details the proposed methodological framework, combining IO-MCDM, HITL and Decision Trust Index. Section 4 illustrates the application of the model through a sustainable case study, presents the results, and discusses the implication for industry 5.0. Section 5 concludes by highlighting the contributions, summarizing the key findings and future work.

2. Related works

2.1. Emerging Multi-Criteria Decision Making (MCDM) Frameworks.

To quantify the exponential growth in MCDM applications and identify the most widely used MCDM methods in the current research, this section is based on significant MCDM’s review papers (most cited in Google scholar dataset), published in the last three years in high-impact indexed journals.

- Authors in [20] provide a rich, 45-year bibliometric overview connecting MCDM methods (AHP, TOPSIS, VIKOR, PROMETHEE, and ANP) and metaheuristics, identify emerging hybrid methods integrating AI, fuzzy logic, and swarm intelligence with MCDM, and highlight a methodological shift toward data-driven, sustainability-oriented and context-specific decision models.
- The paper [3] highlights a global review of more than 23,494 documents (from 1977-2022), identifies main methods (AHP, TOPSIS, ELECTRE, VIKOR). The authors concluded that the scientific production growth rate of 14.18 % per year, demonstrates the academic community’s interest in researching and publishing publications on multi-criteria decision-making approaches.
- The review [4] consolidates insights from 3,655 peer-reviewed articles to highlight exponential growth in MCDM applications, particularly in sustainable energy, urban planning, and healthcare optimization. And concludes that MCDM’s potential lies in embracing inclusivity, advancing into emerging technologies like Blockchain, the metaverse, and fostering collaboration across underrepresented regions and domains.
- The review [5] surveys more than 190 papers from 2010 to 2022, constructs a general structure of MCDM techniques, provides a comprehensive review and analysis of it, and clarifies the current practices used in sustainable transport, finally shows that VIKOR, AHP, ANP, PROMETHEE, and hybrid methods were found to be the most widely used MCDM methods for low-carbon transport and green logistics.
- Paper [2] integrates MCDM into healthcare and medical decision support systems, identifies more than 140 MCDM papers, published during 2013 and 2022 (available in the Scopus database) and conclude that the most popular MCDM tool is Analytic Hierarchy Process (AHP).

Finally, the above review studies show the increasing use of MCDM methods in last years, across various field and demonstrate their ability to be hybridized with other methods.

2.2. Inverse Optimization and MCDM.

The main objective of Inverse Optimization (IO) is to infer the objective function of an optimization model from observed optimal decisions [12] and [14]. So, rather than solving a traditional optimization problem, inverse optimization reverses the process and defines a Forward Optimization model FOP as in [12]:

$$\text{FOP}(y, \gamma) = \min_x \{f(x, y, \gamma) \mid x \in X(y, \gamma)\} \quad (1)$$

Where $y \in Y$ and $\gamma \in \Theta$, γ controls the objective function $f(x, y, \gamma)$ and the feasible set $X(y, \gamma)$. Consider y as an input representing expert preferences, and γ as the parameter to estimate in (IO) problem. Typical objective function modeling choices include linear, quadratic, and convex separable forms. So for the convex form, the feasible set $X(y, \gamma)$ is defined by m convex constraints, and the optimal solution is learned from an observed data set, $\{(\hat{x}_i, \hat{y}_i)\}_{i=1}^N$, $N \geq 1$ decisions and solution will be :

$$X^{\text{opt}}(y, \gamma) = \arg \min_{x \in X(y, \gamma)} f(x, y, \gamma) \quad (2)$$

The authors in [13] provide a comprehensive review of inverse optimization (IO) applications, and find that the use of the IO concept has grown reflecting rising interest. Papers [15], [35] extend classical inverse optimization by introducing a robust framework that accounts for uncertainty and imperfect observations in the decision-maker's behavior, this, Robust inverse optimization (RIO) recognizes that such decisions may be noisy or suboptimal.

Combining Inverse Optimization (IO) and Multi-Criteria Decision Making methods (MCDM) is an emerging field of research. Recently, only a few researches hybridized MCDM method with Inverse Optimization (IO). For example, Hung, H et al [16] and Flachs et al [18] have both proposed a new framework based on inverse optimization and PROMETHEE. The author in [18] aim to provide a rigorous and efficient algorithmic framework able to answer a crucial practical question: "How to improve one's ranking position, referring to previous preferences?". In [16] Hung, H et al, integrate the Multi-Actor Multi-Criteria Analysis(MAMCA) with inverse mixed-integer linear optimization, to enhances the PROMETHEE method with Inverse Optimization to explore stakeholder compromise. These results are empirically validated and pave the way for many future applications and research in the field of decision support.

2.3. Human-In-the-Loop and Interactive Decision Making

Recently, several research papers discuss how Human-In-The-Loop(HITL) can be an increasingly important concept and how it can enhance the performance of various models [17], [21] and [22]. Indeed, in multi-criteria decision making area, as in [16],], integrating the human as an active agent who continuously refines and enriches the obtained results(decisions, recommendation, etc..) would improve decision-maker's confidence and mitigate systemic risks during the deployment of new support systems.

On the other hand, relevant works show that integrating the AI-interactive decision making model will influence the decision making process, as in [23], by reformulating the optimization problem as a constrained one. And so, human preferences are used to restrict search space (Pareto front). In addition, authors in [24] present an interactive multi-objective optimization approach in which the decision maker intervenes at each iteration, gives feedback (preference, choices, directions) and the algorithm updates the search. The model allows human interaction and the decision-maker can formulate the preferences influencing the selection of optimal solutions at each control iteration. Other works such as [22], [23], show how human can influence optimization and constitute a solid basis for building such dynamic framework, in which decision-maker dynamically adjust the constraints of a multi-criteria optimization model in iterative way. Finally, in [28], the authors demonstrate that systems integrating a human feedback, allows a better uncertainties management, improve the acceptance, reliability and operational robustness of deployed solutions.

Table 1. The main steps of human intervention in the multi-criteria decision-making model

Decision Stage	Human intervention
Stage 1: Previous Decision Making	Human assumes the preparation of the database including alternatives, criteria, and preferences.
Stage 2: Inference stage	The human interacts directly during model execution (e.g., the expert co-leads and the decision-maker validates questionable decisions).
Stage 3: Post-decision	Human feedback for efficiency and reliability validation.

2.4. Decision Trust Index(DTI) in MCDM literature

The notion of “Decision Trust Index(DTI)” can be defined as post-decision validation metric designed to quantify how much confidence a decision maker can place in a recommended ranking, especially when multiple MCDM approaches produce different results for same set of alternatives. Hence, the DTI is proposed to address this issue by measuring the reliability of an aggregated ranking. According to a brief review based on “google scholar data base”, the DTI is surprisingly underdeveloped in the MCDM literature and most MCDM papers focus on producing a ranking, while relatively few assess how much confidence decision-makers should place in the recommended alternative [41]. The literature review in last 3 decades can be classified into three generations:

- First Generation “Robustness and Confidence Measures(1990s-2010)”: Researchers did not use the term Decision Trust Index, but they introduced metrics that essentially quantify the reliability of recommendation as SMAA Confidence Factor [38].
- Second Generation “Robust Decision Analysis (2010-2020)”: Researchers realized that MCDM ranking can change dramatically due to small perturbations in (criteria, normalization methods, aggregation operators or expert judgments). For Sensitivity-Based Trust Measures, the existing methods are (Robust Ordinal Regression, Stochastic Ordinal Regression, Robust PROMETHEE, Robust ELECTRE) [39].
- Third Generation” Reliability-Aware MCDM(2022-Present)”: The most recent trend focuses on reliability itself. Recent MCGDM studies explicitly estimate “Expert Reliability”, “Source Credibility”, “Evidence Reliability” and “Monte Carlo Stability Indices”, therefore the recommendation trust becomes “Trust=f(ranking quality, expert reliability)” [40].

This literature has many trust-related indicators (confidence factor, acceptability index, robustness index, reliability score), but there is no standardized Decision Trust Index that integrates avenue for a methodological contribution, especially if you are developing a new MCDM framework.

This section highlights the important opportunity to improve the multi-criteria decision making model through their integration with the paradigm of Inverse Optimization(IO) and the Human-In-The-Loop (HITL) to ensure acceptability, auditability and safe decisions by defining a Decision Trust Index(DTI) related to MCDM new framework.

3. The Proposed "Preference-Preserving Inverse Optimization–AHP with HITL" (PPIO-AHP) Framework

Under Industry 5.0 dynamic environments, it is important that a DSM be adaptive to change. This expectation reinforces the importance of rethinking the framework in terms of adaptively and transparency. The proposed collaborative human-inverse optimization-MCDM framework can be applied to various MCDM methods, and constitutes a general HITL framework. The optimization layer is independent of the specific MCDM method used. The general IO-MCDM model assumes a set of alternatives and criteria, together with a human-adjustable decision

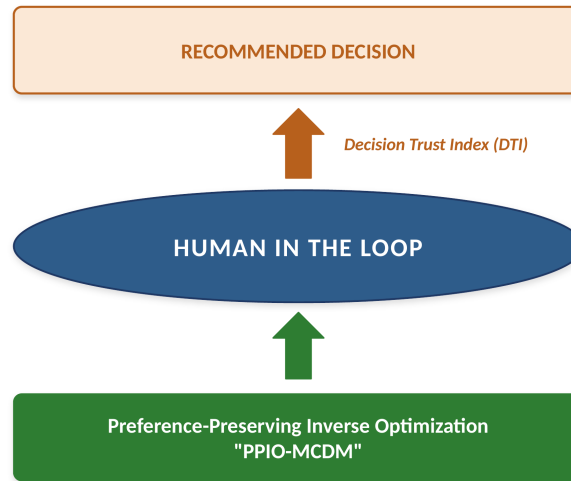


Figure 1. Proposed General Inverse MCDM HITL Framework

vector/matrix.

IO operates on machine-generated criterion evaluation, human-feedback on “weights or preferences” and an MCDM aggregation operator $F(\cdot)$, the decision score is therefore $S_i = F(a_i, w, X)$ where F can be instantiated by different MCDM methods(TOPSIS, VIKOR, AHP, ELECTRE..).

To enhance decision-maker confidence in the recommendations generated by classical MCDM models, a Decision Trust Index (DTI) is introduced as an additional evaluation layer. Unlike traditional approaches that rely solely on current judgements, the proposed DTI exploits historical decision cases that have been previously validated by domain experts.

In this paper, the proposed framework is applied to Analytic Hierarchy Process(AHP), who since its introduction by Thomas L. Saaty [31] has established itself as one of the most influential MCDM methods (as shown in Table.1). Over the decades, AHP has experienced numerous extensions intend to overcome its limitations in terms of subjectivity and uncertainty [32] and [33]. Fuzzy-AHP was proposed to model imprecise preferences (probabilistic). Moreover, the integration of AHP with the AI techniques and Machine learning(ML) allowed to develop hybrid-AHP, able to adjust dynamically the weights making this method more data- driven and adaptive [29].

The proposed solution introduces an intrinsic modification of hybrid AHP and implements a full aggregation scheme capable of producing a complete ranking of alternatives. The proposal enhances (i) the ability to automatically correct the inconsistency of judgments, and (ii) the possibility of producing more reliable recommendations, which are validated by the decision-maker in two ways (simple and preference-preserving inverse AHP).

The framework is divided into two key interconnected components (As shown in Figure 1) includes:

- **The first component** combines the Inverse Optimization (IO) with an MCDM method. The suggested framework formulates the problem as an inverse optimization model in which the pairwise comparison matrix is processed as an unknown latent parameter inferred from historical decision data. The framework follows the data-driven inverse optimization paradigm, where observed decisions are used to estimate hidden parameters of a forward optimization model under imperfect model-data fit and the pairwise comparison matrix becomes the unknown parameter vector of the forward optimization model FOP(A).
- **The second component:** The Human-In-the-Loop (HITL) interactive collaboration mainly involves interactions between the expert and the decision-maker(s). Throughout all the decision process, both human(s) collaborate in order to ensure their role in the loop; from the first step in which the expert assume

the matrix and data preparation, to the final step when the decision-maker should validate the final decision (As shown in Figure 2).

Notation and hypotheses

- Let $A = \{a_1, \dots, a_n\}$ be a set of n alternatives.
- Let $C = \{c_1, \dots, c_m\}$ a set of m criteria.
For each alternative (choice, solution, etc.) a_i , exists an evaluation vector $\mathbf{A} = (a_{i1}, a_{i2}, \dots, a_{im})$, and local priority vector (performance) p_{ij} under criterion c_j , where p_{ij} denotes the local priority of alternative a_i under criterion.
- Suppose (AHP) is appropriate, when criteria are reasonably independent and trade-offs between them are acceptable.
- The weight that best reports on the observed choices :

$$\mathbf{w} = (w_1, w_2, \dots, w_m), \quad w_j > 0, \quad \sum_{j=1}^m w_j = 1. \quad (3)$$

AHP Classical Decision Model In the classical AHP framework, the decision maker provide pairwise comparisons represented by the reciprocal matrix: $A = a_{ij}$ such that $a > 0$, $a_{ij} = 1/a_{ji}$ and $a_{ii} = 1$. The criteria priority vector: $\mathbf{w} = (w_1, w_2, \dots, w_m)^T$ is obtained from the principal eigenvalue problem:

$$\mathbf{A}\mathbf{w} = W\lambda_{max} \quad (4)$$

The resulting priority vector is then used to evaluate alternatives through a multi-criteria utility function or AHP global Score (additive weighting):

$$S_i(a, w) = \sum_j^m w_j P_{ij} \quad (5)$$

Alternatives are ranked by S_i and the higher score indicates higher overall preference. So the best alternative match:

$$\max_{a_i \in A} \sum_j^m w_j P_{ij} \quad (6)$$

And accordant to inverse optimization it match

$$\min_{a_i \in A} - \sum_j^m w_j P_{ij} \quad (7)$$

p_{ij} are linear (i.e. convex), the weighted sum (score S_i) is convex.

3.1. The Proposed PPIO-AHP Model

To directly address the AHP limitations related to subjective and potentially inconsistent judgments, the proposed PPIO-AHP model explores using Inverse Optimization IO to infer the pairwise comparison matrices themselves from historical decision data. The proposed inverse optimization model includes: (i) Decision reconstruction, which allowed the optimal solution learning from an observed data set. (ii) Consistency penalty, which tolerates the value of CR to be less than 0.1. (iii) Preference preservation Mechanism based on minimal cognitive adjustment. The distance term $d(A, A^{(0)})$ penalizes excessive deviation from the original matrix. So the solution A^* directly follows the data-driven inverse optimization structure and the Preference Preservation Analysis:

$$A^* = \arg \min_A \frac{1}{n} \sum_{k=1}^N \underbrace{l(a^{(k)}, a^*(A, u_k))}_{\text{(i) Decision Reconstruction}} + \underbrace{\lambda_1 \max(0, CR(A) - 0.1)^2}_{\text{(ii) Consistency Penalty}} + \underbrace{\lambda_2 d(A, A^{(0)})}_{\text{(iii) Preference}} \quad (8)$$

Where $l(\cdot)$ the loss function type Hinge Ranking Loss which better measures the discrepancy between observed and predicted decisions. For each alternative $a^* \neq a^{(k)}$

$$l = \sum_{a^* \neq a^{(k)}} \max(0, 1 - S(a^{(k)})^{(k)} + S(a^*)^{(K)}) \quad (9)$$

- CR(A) denotes Saaty's Consistency Ratio and subject to $CR(A) \leq 0.1$
- And $A^{((0))}$ is the original judgment matrix provided by decision maker.
- $d(A, A^{((0))})$ measures the deviation from the initial cognitive preferences, $\lambda_1 \in [0.1; 0.3]$ and $\lambda_2 \in [0.3; 0.6]$ are regularization parameters.

The proposed framework is designed as a generic inverse optimization model that can be applied to various decision-making contexts. As a result, the regularization coefficients (λ_1 and λ_2) are treated as configurable hyper parameters rather than fixed model constants. Their values should be determined by the domain expert in collaboration with the decision maker to reflect the characteristics and priorities of the specific application (step 3 Figure2). These values are therefore designed as practical starting points.

Consistency Preservation Strategy: The proposed framework introduces a Preference Preservation Mechanism which preserve the decision maker preference through a proposed cognitive adjustment. The distance term: $d(A, A^{((0))})$ penalizes excessive deviation from the original matrix. In this study, the Frobenius norm is adopted:

$$d(A, A^{(0)}) = \|A - A^{(0)}\|_{F^2} = \sum_{i,j} (X_{i,j} - X_{i,j}^{(0)})^2 \quad (10)$$

This strategy guarantees that consistency improvement is achieved with the smallest possible alteration of the initial judgments. The Preference Preservation Index (PPI) is defined as:

$$PPI = 1 - d(A^*, A^{(0)})/d_{max} \quad (11)$$

A high PPI value indicates strong preservation of the original decision-maker cognition.

Consistency Ratio computation: The consistency ratio is computed according to Saaty's formulation. For each pair comparison matrix ($m \times m$), the Consistency Index (CI) is obtained as:

$$CI = (\lambda_{max} - m)/(m - 1) \quad (12)$$

$$CR = CI/RI \quad (13)$$

Where λ_{max} is the maximum eigenvalue of matrix A and RI is Saaty's Random Index (Notice that frequently, RI values for $m > 10$ are obtained from extended Monte Carlo estimates available in the literature). The constraint: $CR \leq 0.1$ ensures acceptable logical consistency.

Sensitivity test PPIO-AHP: Since the proposed PPIO-AHP model learns A^* from historical decisions D, the inferred matrix may depend on the composition of the observed dataset. Variations in the observed decision set or in the associated ranking may affect the reconstructed preference model. The purpose of sensitive analysis is to evaluate how the inferred pairwise comparison matrix responds under similar variations. To assess this effect, a leave-one-out (LOO) strategy was employed. For each $(k \in 1, \dots, N)$, a reduced dataset was constructed by removing the corresponding decision-context pair: $D_{-k} = D / (a^{((k))}, u_k)$. The inverse optimization problem was then solved again using the reduced dataset. The sensitivity was quantified using the Frobenius distance between the full-data solution and the leave one-out solution $S(K) = \|A^{(-k)} - A^*\|_F$. The Frobenius norm was selected because it provides a natural measure of the overall discrepancy between two pairwise comparison matrices while preserving computational simplicity. A small value of $S(k)$ indicates that the reconstructed model remains stable after removing observed (kth) observation, whereas larger values of $S(k)$ reveal higher sensitivity to the corresponding observed decision.

In this framework, the LOO procedure is not intended to prove insensitivity to historical data; rather, it is used to

Table 2. Decision Trust

DTI	Trust level
<0.5	Low
0.5 – 0.7	Moderate
0.7-0.85	High
>0.85	very High

identify whether the reconstructed preference model is particularly sensitive to a specific set of past decisions or whether the learned preference structure remains globally consistent. Therefore, robustness is interpreted as the ability of the reconstructed preference model to maintain stable preference structures under small perturbations induced by the removal of individual historical decision while remaining sufficiently adaptive to incorporate informative historical decisions.

Decision Trust Index: The original assumption is that recommendations revealing a high degree of similarity with past accepted decisions are more likely to be perceived as reliable and truthful by decision makers. The DTI is therefore defined, in this study, as a quantitative measure of the consistency between the current MCDM-derived ranking and the ranking associated with historical validated cases. Finally, by using simple trust metric based solely on ranking stability ($\text{Stability Ratio} = \text{Number of unchanged alternatives ranking} / \text{Total alternatives}$, Spearman's " ρ ", Kendall's Tau " τ ")

A simple Decision Trust Index (DTI) can be defined as: $DTI = (\text{Stability Ratio} + \rho + \tau) / 3$. A higher DTI value indicates a stronger agreement between the current recommendation and the organization's historical decision patterns, thereby providing an additional level of assurance without altering the decision-maker's original preference. The Trust Level (TL) can be for example defined between 0 and 1 as explained in Table 2.

Consequently, the proposed framework not only preserves the interpretability and consistency properties of the classical AHP methodology but also incorporates organizational experience into the decision-making process, resulting in improved trust, transparency and acceptance of the final recommendation.

3.2. A Step-by-Step PPIO-AHP Methodology

To implement the proposed hybrid Inverse AHP, the proposed methodology is detailed in six steps:

Step1 : The Performance Matrix

The objective is to define the MCDM performance Matrix. The Expert collaborates with the Decision Maker (DM) to:

- **Identify the key evaluation dimensions:** Financial assessment, environmental impact, technical risk
- **Structure the appropriate criteria framework:** incorporates a maximum of criteria, categorized to address explicit facets of "problem" performance and effect. For example: the technical criteria include efficiency and reliability measurements, the Environmental criteria assess the ecological sustainability and footprint, the social criteria consider the risk and opportunities of local population and the financial criteria measure the financial viability.
- **Collect and integrate real data sources to define the multi-criteria decision making performance Matrix:** Based on the decision-makers preferences an initial pairwise comparison matrix $A^{(0)}$ is built and may contain inconsistencies.

Step2: Forward AHP Evaluation

The priority vector of $A^{(0)}$ is computed using the AHP eigenvalue method and the alternatives utilities are then evaluated. The initial ranking is obtained.

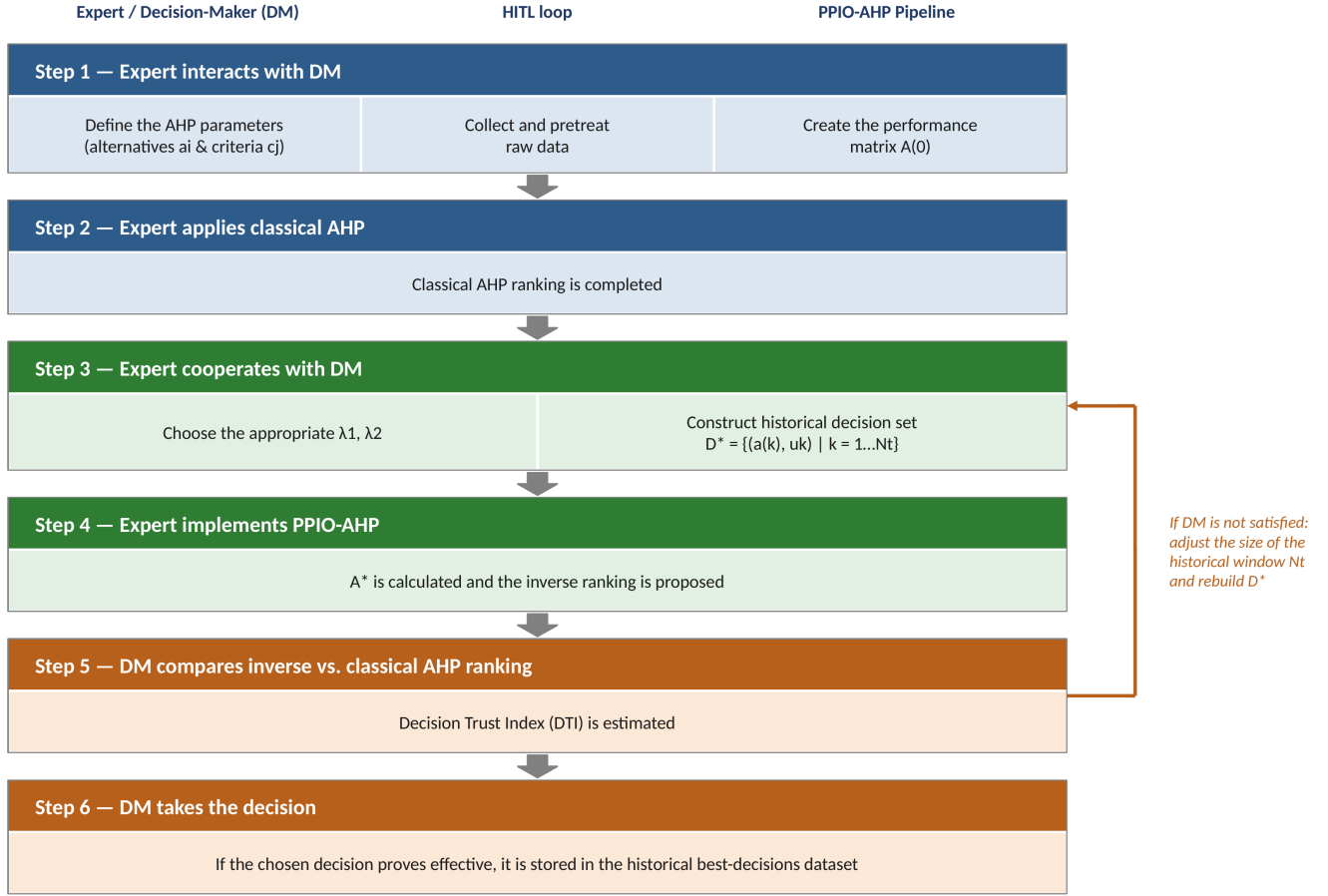


Figure 2. PPIO-AHP HITL Framework

Step3: Historical Validate Decision Acquisition

Historical ranking is collected to construct dataset D . The Observed Decision Dataset assume that a historical dataset of decision is available:

$$D = (a^{(k)}, u_k) \mid k = 1 \dots N_t \quad (14)$$

Where $a^{(k)}$ denotes the validate selected alternative, and u_k represents the nearest(similar) contextual information associated with decision instance k . Therefore, inverse optimization seeks parameters that render observed decisions approximately optimal under the forward model.

Step4: PPIO-AHP is executed

The inverse AHP problem is solved to infer the latent matrix A^* minimizing by ranking reconstruction error, inconsistency and cognitive deviation.

Step5: Decision Trust Index and human validation

Once the inverse ranking is achieved, the Decision Trust Index(DTI) is estimated and the Expert discuss with the Decision Maker(DM) the proposed PPIO-AHP result:

- If he validates the above inverse AHP ranking, the decision is taken.
- If If he does not validate this ranking, so a feedback is operationalized by adjusting the size of the historical window. The updated is used to rebuild D^* and the PPIO-AHP problem is resolved (Steps 3 to 5) until the DM validates the ranking.

Step6: Final Ranking DM take decision and in the future, if the selected PPIO-AHP decision is demonstrate its effectiveness as one of the best decisions, then it will be automatically storage in the D data-set of the historical best decisions.

3.3. *Proposed Interactive Human in The Loop and PPIO-AHP framework*

Industry 5.0 redefines industrial automation by placing greater emphasis on cooperation between humans and systems. Rather than replacing humans, it places them at the center of intelligent decision-making loops. Human-in-the-Loop (HITL) approaches align with this shift, enabling continuous interaction between systems and human. At the core of this study, a Human-In-the-Loop and PPIO-AHP pipeline is defined, as illustrated in Figure 2. Human (Expert & DM) rather than acting solely as data validators, intervene after the reconstruction of the preference model to assess the consistency and practical relevance of the inferred pairwise comparison matrix. Based on this assessment, Human (Expert & DM) can confirm the reconstructed preference structure or request refinement by re-examining historical decision data used as input to the inverse optimization model. It is important to note that, in the current version of the proposed framework, human feedback is integrated at the decision support level, while expert feedback determines whether the reconstructed preference model is accepted or if the reconstruction process should be repeated using revised historical information. in the future it's suggested that an optimization-based feedback mechanism. he human intervention is mainly during all the decision making process (as shown in Table 1) and the decision support model learns from human feedback.

- The first human collaboration: The Expert interacts with DM to define the AHP parameters (Alternatives/Criteria) then the Expert collects and pretreats data.
- The second human intervention: The Expert collaborates with DM to build "Dataset D" of the K previous best decisions.
- The final human collaboration: Once the inverse ranking is accomplished and the DTI is estimated then Decision maker chooses the most trustful ranking.

Consequently, the proposed decision support framework establishes a robust human-centric teaming component to support decision-maker acceptance. The bi-level optimization formalizes preference elicitation as a reverse optimization problem, adjusting variables to reproduce observed decisions. Then, the stability (robust versus unstable preferences) and sensitivity (impact of perturbations). Due to the proposed framework, the hybrid decision support model become more reliable.

4. Case study: Experimental results

The Human in the loop collaboration is essential for the industry 5.0 as human co-evolution-technology, based on symbiosis and responsibility, Decision-maker confidence then becomes an indicator of decision-making performance.

In this paper, the case study focuses on the prioritization of strategic renewable energy projects, integrating energy water-environmental (NEXUS) considerations [30]. In the renewable energy project selection process, it is important to categorize a comprehensive set of criteria to ensure a harmonious and well-organized decision.

This study considers 9 strategic projects, inspired from real case (Moroccan Solar Plant). A thorough evaluation is crucial to optimize the sustainable project choice, as illustrated in (Table 3). The ten selected criteria are aligned with (IEA, World Bank, UN SDGs), and are classified into three key categories: Technical, environmental/social and economic. Each group detailed a precise aspect to assess the performance of the considered project (Table 4).

Table 3. List of Alternatives

Code	Project
A_1	Large-scale photovoltaic solar park (Ain BNI MATHAR ONEE)
A_2	AIT BAHA park.
A_3	ECARE
A_4	Noor I project (MASEN)
A_5	Noor II project (MASEN 2024/ACWA 2025)
A_6	Noor III project
A_7	Hybrid CSP/PV Noor Midelt I project
A_8	Noor Midelt II project
A_9	Sud-Project

This type of issue is widely documented in MCDA works applied to sustainable energy policies IEA (Energy policy indicators), World Bank (Sustainable infrastructure).

Table 4. List of Criteria

Code	Criteria
C_1	DNI(Kwh/m2/years)
C_2	Cost « investment operating » (MDH)
C_3	CO_2 reduction
C_4	Capacity of Production (MW)
C_5	Storage (average hour)
C_6	Superficies (Ha)
C_7	Project Life time (years)
C_8	Distance (Km)
C_9	Land Slope (%)
C_{10}	Job opportunities

The decision maker is usually a national public authority (Ex: MASEN, IRESEN, ONEE.), responsible of energy planning and possessing a history of decisions made on similar projects.

The classical AHP process includes three “priorities levels” and each level is validating by the “Consistence Index (CI)”, which is a key concept used to measure how consistent the decision maker’s pairwise comparisons in AHP. In the above case study, the criterion C_1, C_2, C_3 are considered more important than C_4 to C_{10} .

Table 5. Initial matrix

Project	C_1	C_2	C_3	C_4	C_5	C_6	C_7	C_8	C_9	C_{10}
A_1	2290	4600	2	470	1	160	20	6	2.1	60
A_2	2200	30	3	3	5	24	25	20	5	7
A_3	2050	19	4	1	2	2	20	25	4	2
A_4	2635	9000	3	160	3	480	20	10	1.1	80
A_5	2635	1050	3	200	7	680	25	10	1.1	80
A_6	2635	8300	3	150	7.5	750	25	10	1.1	80
A_7	2630	20000	2	800	5	3000	25	20	1.5	150
A_8	2630	10000	3	400	2	1500	25	20	1.5	75
A_9	2630	24000	3	1000	5	4500	25	20	1.5	175

4.1. PPIO AHP application

By formulating the behavior of the decision-maker as an inverse optimization problem, the proposed method allows us to retrieve implicit preference structures consistent with historical decisions.

To perform a controlled evaluation based on scenarios of the proposed method, this case study suppose a DM's history choice over the last 10 years and so construct a determinist synthetic dataset ($R^{(1)}$ to $R^{(10)}$).

Subsequently, ten historical scenarios were created by changing the position of a part of the alternatives in each scenario. Each ranking thus represents a distinct hypothetical preference history of the decision-maker.

So for each ranked alternatives $R^{(t)}$ it exist a chosen alternative $y^{(t)}$: for $t = 1, \dots, K$ it exist a set of ranked alternatives $R^{(t)}$ and the chosen alternative $y^{(t)}$:

$R^{(1)} = A_3 > A_1 > A_2 > A_7 > A_8 > A_4 > A_5 > A_6 > A_9$ and $y^{(1)}$ is A_3 ;
 $R^{(2)} = A_3 > A_1 > A_7 > A_2 > A_8 > A_4 > A_5 > A_6 > A_9$ and $y^{(2)}$ is A_3 ;
 $R^{(3)} = A_1 > A_2 > A_3 > A_7 > A_8 > A_4 > A_5 > A_6 > A_9$ and $y^{(3)}$ is A_1 ;
 $R^{(4)} = A_5 > A_2 > A_1 > A_8 > A_7 > A_3 > A_4 > A_6 > A_9$ and $y^{(4)}$ is A_5 ;
 $R^{(5)} = A_2 > A_1 > A_3 > A_8 > A_7 > A_5 > A_4 > A_6 > A_9$ and $y^{(5)}$ is A_2 ;
 $R^{(6)} = A_6 > A_3 > A_7 > A_2 > A_8 > A_4 > A_5 > A_1 > A_9$ and $y^{(6)}$ is A_6 ;
 $R^{(7)} = A_3 > A_7 > A_4 > A_2 > A_8 > A_9 > A_5 > A_6 > A_1$ and $y^{(7)}$ is A_3 ;
 $R^{(8)} = A_1 > A_5 > A_9 > A_3 > A_2 > A_6 > A_7 > A_8 > A_4$ and $y^{(8)}$ is A_1 ;
 $R^{(9)} = A_4 > A_5 > A_6 > A_1 > A_3 > A_7 > A_2 > A_8 > A_9$ and $y^{(9)}$ is A_4 ;
 $R^{(10)} = A_9 > A_3 > A_2 > A_1 > A_7 > A_8 > A_4 > A_5 > A_6$ and $y^{(10)}$ is A_9 .

These controlled ranking variations allow the proposed learning mechanism to infer the evolution of the decision maker's preferences and update the ranking accordingly. Then the objective is to find out an A^* while the total score be consistent with the observed preferences, preserve the decision-maker first preference $A^{(0)}$ and consistency ($0\% < IC < 10\%$).

Table 6. The inverse optimized Ranking PPIO-AHP

Alternatives Code	Classic AHP Ranking	Inverse Ranking	Observed change of Score
A_1	2	2	Similar
A_2	3	3	Similar
A_3	1	1	Similar
A_4	6	6	Similar
A_5	7	7	Similar
A_6	8	8	Similar
A_7	4	5	Change
A_8	5	4	Change
A_9	9	9	Similar

In this case study the optimized inverse ranking proposes a set of best recommended alternatives $s_{A3} > s_{A1} > s_{A2}$, the priority calculated converges towards the empirical preference of the decision maker over 10 years (as showed in Table 6). In this case its observed that PPIO-AHP keeps the same 3 first best alternatives and 70% of similar ranking for the 6 last alternatives. Generally, the comparison between the rankings produced by the proposed framework and those obtained using classical AHP. When the optimized ranking remains concordant with the classical AHP ranking, the agreement can be interpreted as an additional source of confidence for the decision maker, since the recommendation is simultaneously supported by the original preference model and by the historical learning process. Conversely, when the proposed framework yields a ranking that differs from the classical AHP result, and the interaction human- PPIO-AHP has done regulate the hyper parameter/ Historical decision set D^* without significant modification of this results; so the obtained divergence should not be always interpreted as an

inconsistency. Rather, it may reveal that the historical best decisions convey additional knowledge not captured by the original pairwise comparisons alone. In this case, the divergence becomes informative, as it highlights situations where the decision maker's initial preferences and the historical evidence do not fully align.

4.2. Sensitivity test of PPIO-AHP Ranking

To explain the proposed sensitivity analysis, the experiment was directed using a dataset composed of ten observed PPIO-AHP decision scenarios. First, the inverse optimization model was solved using the complete dataset to reconstruct the consensus pairwise comparison matrix A_1^* and the corresponding baseline ranking $R^{10} = (A_3 > A_1 > A_2 > A_8 > A_7 > A_4 > A_5 > A_6 > A_9)$. Subsequently, a leave-one-out (LOO) procedure was achieved by removing each decision scenario individually and reconstructing a new consensus matrix. For illustration, consider the case in which the removed observation corresponds to a decision favoring alternative A_9 . After solving the inverse optimization problem with the reduced dataset, the resulting ranking becomes $R^9 = (A_3 > A_1 > A_2 > A_8 > A_5 > A_4 > A_7 > A_6 > A_9)$, where the only change is the interchange between alternatives A_5 and A_7 . The confidence intervals associated with the priority weights of the leading alternatives A_3 overlap across the LOO realizations, indicating that the first ranked alternative is conserved despite the removal of an individual historical decision. According to DTI classification presented in Table 2, the reconstructed ranking is therefore categorized as a "Very High Trust/Very High Stability", with a DTI value close to 0.90.

Generally, the LOO results indicate that the proposed framework exhibits a controlled sensitivity to historical knowledge. Removing a single historical decision produce only limited modifications in the reconstructed preference structure, suggesting that the reconstructed preference structure is not arbitrarily unstable when one historical decision is removed, yet it remains sufficiently responsive to the informational content of past best decisions. This observation is particularly important for the interpretation of robustness in the proposed framework. Unlike conventional robustness analyses, where low sensitivity to all data variations is generally expected, the present model is explicitly designed to learn from historical decisions.

Consequently, some degree of variation under the LOO procedure should not be interpreted as a weakness; rather, it confirms that the historical learning component actively contributes to the reconstruction of the optimized AHP matrix. In this sense, the LOO analysis should be understood as a measure of historical influence rather than as a test of strict invariance.

5. Conclusions and perspectives

This paper proposes a trust-enhanced Analytic Hierarchy Process(AHP) framework that integrates historical decision knowledge through Inverse Optimization into classical multi criteria decision-making approaches, to enhance the reliability, interpretability, and verifiability of decision making process. The results of the study address two critical research questions regarding the enhancement of classical MCDM models by Inverse Optimization (IO) and the integration of the Human-In-the-Loop (HITL).

Regarding the RQ1, this research has revealed that IO can effectively overcome the inherent limitations of classical AHP. By inferring latent preference structures directly from historical decision-making data, the IO-based approach confirmed that inverse optimization could reconstruct the priorities of policymakers that better align with observed actions, thus improving the reliability of the model and the traceability of decisions. The concept of a Decision Trust Index (DTI) was introduced to quantify the level of agreement between the current recommendation and historical validated decisions.

Regarding RQ2, the study proposes a unified human in the loop framework where expert feedback and behavioral data inform the optimization process. This design enables the integration of human expertise into the decision making process. The proposed solution offers decision-makers not only a recommended ranking but also an interpretable indication of its trustworthiness, thereby improving the acceptance and perceived reliability of the decision support process.

Overall, the proposed framework provides a methodologically rigorous and ethically transparent paradigm for modeling complex decision-making behaviors.

Despite its contributions, this research has some limitations: (i) The absence of uncertainty handling mechanism, such as fuzzy or interval-based pairwise comparisons (ii) Moreover, total aggregation approaches generally assume criteria independence and full compensability. (iii) an additional limitation is due to the unused AI techniques, for example automated data collection and generative AI techniques.

As future perspectives:

- The proposed approach will be extended to additional aggregation methods (as the partial and local aggregation) and future works would propose a new hybrid IO-MCDM based on further methods such as “TOPSIS, ELECTRE or VIKOR”. Thus, will generalize the integration of inverse reasoning across various decision-making models. Such extensions would supplementary contribute to the development of auditable and ethically transparent decision support systems, thereby bridging the gap between prescriptive and descriptive analyses in multi-criteria decision analysis.
- The integration of real-time dynamic adjustment of the HITL interaction with new inverse-multi-criteria optimization models. The use of Large Language Model and Generative AI, would be an additional enhancement for a future data-driven Hybrid IO-decision support systems.

REFERENCES

1. Ika, L. A., & Munro, L. T. *Tackling grand challenges with projects: Five insights and a research agenda for project management theory and practice*, International Journal of Project Management, 40(6), 601-607.2022
2. Chakraborty, S., Raut, R. D., Rofin, T. M., & Chakraborty, S. *A comprehensive and systematic review of multi-criteria decision-making methods and applications in healthcare*, Healthcare Analytics, 4, 100232.2023
3. Basílio, M. P., Pereira, V., Costa, H. G., Santos, M., & Ghosh, A. *A systematic review of the applications of multi-criteria decision aid methods (1977–2022)*, A systematic review of the applications of multi-criteria decision aid methods (1977–2022). Electronics, 11(11), 1720.2022
4. Kumar, R., & Pamucar, D. *A comprehensive and systematic review of multi-criteria decisionmaking (MCDM) methods to solve decision-making problems: two decades from 2004 to 2024*, Spectrum of Decision Making and Applications, 2(1), 177-196.2025
5. Tian, G., Lu, W., Zhang, X., Zhan, M., Dulebenets, M. A., Aleksandrov, A., ... & Ivanov, M. *A survey of multi-criteria decision-making techniques for green logistics and low-carbon transportation systems.*, Environmental Science and Pollution Research, 30(20), 57279-57301.2023.
6. Demir, G., Chatterjee, P., & Pamucar, D. *Sensitivity analysis in multi-criteria decision making: A state-of-the-art research perspective using bibliometric analysis*, Expert Systems with Applications, 237, 121660.2024.
7. Więckowski, J., & Sałabun, W. *Sensitivity analysis approaches in multi-criteria decision analysis: A systematic review*,
8. Derakhshanfar, N., Heydari, H., Mirzaei, A., & Pourhajikazem, A. *An intelligent framework for energy optimization in IoT networks using LSTM and multi-criteria decision making*, Sustainable Computing: Informatics and Systems, 101246.2025.
9. Luo, Y., Zhang, Y., Du, C., Zhang, H., & Liu, Y. *Handover algorithm based on Bayesian-optimized LSTM and multi-attribute decision making for heterogeneous networks*, Ad Hoc Networks, 157, 103454.2024.
10. Kostopoulos, G., Davrazos, G., & Kotsiantis, S. *Explainable artificial intelligence-based decision support systems: A recent review*, Electronics, 13(14), 2842.2024
11. ŞAHİN, E., Arslan, N. N., & Özdemir, D. *Unlocking the black box: an in-depth review on interpretability, explainability, and reliability in deep learning*, Neural Computing and Applications, 37(2), 859-965 .2025
12. Chan, T. C., Mahmood, R., & Zhu, I. Y. *Inverse optimization: Theory and applications.*, Operations Research, 73(2), 1046-1074.2025
13. Ghaffar, A. R. A., Melethil, A., & Adhami, A. Y. *A bibliometric analysis of inverse optimization.*, Journal of King Saud University-Science, 35(7), 102825. 2023
14. Ahuja, R. K., & Orlin, J. B. (2001). Inverse optimization. Operations research, 49(5), 771-783. *Inverse optimization*, Operations research, 49(5), 771-783.20201
15. Ghobadi, K., Lee, T., Mahmoudzadeh, H., & Terekhov, D. (2018). Robust inverse optimization. Operations Research Letters, 46(3), 339-344. *Robust inverse optimization.*, Operations Research Letters, 46(3), 339-344. 2018
16. Huang, H., De Smet, Y., Macharis, C., & Doan, N. A. V. *Collaborative decision-making in sustainable mobility: identifying possible consensus in the multi-actor multi-criteria analysis based on inverse mixed-integer linear optimization*, International Journal of Sustainable Development & World Ecology, 28(1), 64-74. 2021
17. Kumar, S., Datta, S., Singh, V., Datta, D., Singh, S. K., & Sharma, R. (2024). Applications, challenges, and future directions of human-in-the-loop learning. IEEE Access, 12, 75735-75760. *Applications, challenges, and future directions of human-in-the-loop learning.*, IEEE Access, 12, 75735-75760 .2024
18. Flachs, A., & De Smet, Y. (2025). Inverse optimization on the evaluations of alternatives in the Promethee II ranking method. Omega, 136, 103325. *Inverse optimization on the evaluations of alternatives in the Promethee II ranking method*, Omega, 136, 103325 .2025
19. Mohajerin Esfahani, P., Shafieezadeh-Abadeh, S., Hanasusanto, G. A., & Kuhn, D. *Data-driven inverse optimization with imperfect information*, Mathematical Programming, 167(1), 191-234. 2018
20. Rocha, C. M. M., Benitez, A. S., & Buelvas, D. A. *Review and bibliographic analysis of metaheuristic methods in multicriteria decision-making: A 45-year perspective across international, Latin American, and Colombian contexts*, Journal of Applied Mathematics, 2024(1), 5577682.2024.

21. Wu, X., Xiao, L., Sun, Y., Zhang, J., Ma, T., & He, L. (2022). A survey of human-in-the-loop for machine learning. *Future Generation Computer Systems*, 135, 364-381. *A survey of human-in-the-loop for machine learning*, *Future Generation Computer Systems*, 135, 364-381 .2022
22. Shabur, M. A., Shahriar, A., & Ara, M. A. *From automation to collaboration: exploring the impact of industry 5.0 on sustainable manufacturing*, *Discover sustainability*, 6(1), 341.2025
23. Li, S., Puig, X., Paxton, C., Du, Y., Wang, C., Fan, L., ... & Zhu, Y. *Pre-trained language models for interactive decision-making*, *Advances in Neural Information Processing Systems*, 35, 31199-31212.2022
24. Wang, X., Zheng, J., Hou, Z., Liu, Y., Zou, J., Xia, Y., & Yang, S. *Unlocking the black box: an in-depth review on interpretability, explainability, and reliability in deep learning*, *Swarm and Evolutionary Computation*, 89, 101638 .2024
25. Lu, Y., Zheng, H., Chand, S., Xia, W., Liu, Z., Xu, X., ... & Bao, J. (2022). Outlook on human-centric manufacturing towards Industry 5.0. *Journal of manufacturing systems*, 62, 612-627 *Outlook on human-centric manufacturing towards Industry 5.0*, *Journal of manufacturing systems*, 62, 612-627 . 2022
26. Wang, B., Zheng, P., Yin, Y., Shih, A., & Wang, L. *Toward human-centric smart manufacturing: A human-cyber-physical systems (HCPS) perspective*, *Journal of Manufacturing Systems*, 63, 471-490 .2022
27. Longo, F., Padovano, A., & Umbrello, S. *Value-oriented and ethical technology engineering in industry 5.0: A human-centric perspective for the design of the factory of the future*, *Applied sciences*, 10(12), 4182 . 2020
28. Nicodeme, C. *Build confidence and acceptance of AI-based decision support systems-Explainable and liable AI*, In 2020 13th international conference on human system interaction (HSI) (pp. 20-23). IEEE . 2020
29. Kumar, R., Dwivedi, S. B., & Gaur, S. *A comparative study of machine learning and Fuzzy-AHP technique to groundwater potential mapping in the data-scarce region*, *Computers & Geosciences*, 155, 104855 .2021
30. Neto, R. D. C. S., Bohner, J. P., Birch, R. S., Junges, I., Mussi, C. C., Soares, S. V., ... & Salgueirinho Osório de Andrade Guerra, J. B. *Water, Energy and Food Nexus: A Project Evaluation Model*, *Water*, 16(16), 2235 .2024
31. Tavana, M., Soltanifar, M., & Santos-Arteaga, F. J. *Analytical hierarchy process: Revolution and evolution*, *Annals of operations research*, 326(2), 879-907.2023
32. Ibrishesh, I., Ho, S. B., & Chai, I. *Overcoming scalability issues in analytic hierarchy process with ReDCCahp: An empirical investigation*, *Arabian Journal for Science and Engineering*, 43(12), 7995-8011 . 2018
33. Alharbi, A., Alosaimi, W., Alyami, H., Alouffi, B., Almulihi, A., Nadeem, M., ... & Khan, R. A. *Selection of data analytic techniques by using fuzzy AHP TOPSIS from a healthcare perspective*, *BMC Medical Informatics and Decision Making*, 24(1), 240 . 2024
34. Wang, Q., Ma, Y., Zhao, K., & Tian, Y. (2022). A comprehensive survey of loss functions in machine learning. *Annals of Data Science*, 9(2), 187-212. *A comprehensive survey of loss functions in machine learning*, *Annals of Data Science*, 9(2), 187-212 .2022
35. Guo, S., & Xu, H. (2021). Statistical robustness in utility preference robust optimization models. *Mathematical Programming*, 190(1), 679-720. *Statistical robustness in utility preference robust optimization models*, *Mathematical Programming*, 190(1), 679-720.2021
36. Bonissone, P. P., Subbu, R., & Lizzi, J. *Multicriteria decision making (MCDM): a framework for research and applications*, *IEEE Computational Intelligence Magazine*, 4(3), 48-61 .2009
37. Wang, X., Zheng, J., Hou, Z., Liu, Y., Zou, J., Xia, Y., & Yang, S. *Unlocking the black box: an in-depth review on interpretability, explainability, and reliability in deep learning*, *Swarm and Evolutionary Computation*, 89, 101638 .2024
38. Song, S., & Zhang, Y. *An SMAA-based approach for multi-attribute fixed resource allocation problem with random attribute values and uncertain weights*, *Journal of Modelling in Management*, 20(4), 1246-1264.2025
39. Li, W., Meng, W., Kwok, L. F., & Ip, H. H. *Enhancing collaborative intrusion detection networks against insider attacks using supervised intrusion sensitivity-based trust management model*, *Journal of Network and Computer Applications*, 77, 135-145.2017
40. Kaya, R., Salhi, S., & Spiegler, V. *A novel integration of MCDM methods and Bayesian networks: the case of incomplete expert knowledge*, *Annals of Operations Research*, 320(1), 205-234.2023
41. Majid Mohammadi, Jafar Rezaei. *Ensemble ranking: Aggregation of rankings produced by different multi-criteria decision-making methods*, *Omega*, Volume 96,102254.2020