

A Novel Weighted Preference Relation Approach to Detect Outliers in Multi-Criteria Decision Aid Context

Toufik Achir^{*,1}, Baroudi Rouba^{*,2}

Abstract. The problem of outlier detection in the multicriteria decision aid (MCDA) field has not been extensively explored in the current literature. This study presents a novel approach to tackle this challenge, based on two key concepts. Firstly, the degree of importance of a preference relation, which utilizes multicriteria preference indices (derived from the PROMETHEE method) to assess the significance level of a preference relation. Secondly, the similarity of alternatives, which uses the degree of importance to evaluate how similar each alternative is to the others. Based on the distribution of these similarities, outliers are identified using either the Interquartile Range (IQR) method or the Standard Deviations (SD) method. The proposed approach is applied to two real-world scenarios: the world happiness ranking problem and the human development index problem.

Keywords: degree of importance, preference relation, multicriteria decision aid, PROMETHEE, outlier detection, statistical methods.

1. Introduction

Outlier detection, also known as anomaly detection, is a critical data analysis process aimed at identifying observations or instances that deviate significantly from the majority of the data. In various fields, including statistics, machine learning, and data mining, the presence of outliers can impact the accuracy and reliability of analyses, making their detection and handling an essential task. Numerous techniques have been developed to identify outliers. These methods can be categorized into several distinct categories based on their underlying principles and approaches. These categories encompass [24] statistical-based methods [29], distance-based methods [7], density-based methods [6], clustering-based methods [16], graph-based methods [18], ensemble-based methods [19] and learning-based methods [20].

* Lab CSTL, Mostaganem University, 27000 Mostaganem, Algeria.

¹ Faculty of Exact Sciences and Computer Science, Mostaganem University, toufik.achir.etu@univ-mosta.dz

² Faculty of Science and technology, Mostaganem University, baroudi.rouba@univ-mosta.dz

Outlier detection finds applications in diverse fields, including fraud detection [3], network intrusion [13], wireless sensor networks [26] and various other fields. However, outlier detection has not received much attention in the field of Multicriteria Decision Aid (MCDA), as reflected by the limited number of works addressing this aspect.

In the MCDA field, methods are designed to help decision makers (DM) in the comparison of a set of alternatives, denoted as $A = \{a_1, a_2, \dots, a_n\}$, which are evaluated on a set $F = \{f_1, f_2, \dots, f_k\}$ of k conflicting criteria. Three distinct types of problems are considered in the MCDA context [23] : the choosing problem involves comparing a set of alternatives to determine the best one; the ranking problem requires alternatives to be ranked from best to worst; and the sorting problem assigns each alternative to a predefined category. According to Vincke [27], we can differentiate three main families of multicriteria methods: aggregating, outranking, and interactive methods. Interactive methods [15] involve active engagement with the DM. After an initial computation, a preliminary solution is proposed. If the DM finds the current solution unsatisfactory, they provide additional preferences, such as expressing a desire to enhance a specific criterion while accepting degradation in others. This information is then incorporated into the optimization model, leading to the generation of a new solution. The iterative process continues until convergence towards a satisfactory solution is achieved. Aggregation methods [14] involve combining individual criteria into a single comprehensive measure, typically represented as a utility function. Decision-makers assign weights to criteria or use mathematical functions to aggregate criterion values, facilitating the evaluation of alternatives. Outranking methods focus on pairwise comparison of alternatives considering their performance across multiple criteria. These methods aim to identify alternatives that dominate others without the need for aggregating criteria into a single utility function. In this work, our focus will be on a particular family of Outranking methods known as PROMETHEE [5], which are widely utilized in practice. Further details about this method will be provided in the background section.

1.1. Existing challenges

Certainly, most multicriteria methods heavily rely on subjective parameters, such as preference, indifference, weight, and veto thresholds, which serve as representations of decision makers' preferences. These subjective parameters play a crucial role in shaping the outcome of the methods. Changes in these parameters may significantly impact the results generated, potentially leading to the emergence of an alternative or alternatives that are uncommon, diverging from the majority of the other elements. In MCDA context such alternative(s) are termed multicriteria outliers [11]. Therefore, the detection of outliers in MCDA can play a pivotal role in enhancing the results of MCDA methods and guiding the DM to define the optimal parameter values. Identifying and addressing multicriteria outliers contributes to refining the decision-making process, ensuring more robust and accurate outcomes. However, this field, like any other, comes with its set of challenges, including:

- The specific characteristics of the MCDA field, which relies on preference relations.
- An outlier is defined according to the MCDA context.
- Existing outlier detection techniques may need adaptation or customization to suit the MCDA context.

In the domain of MCDA, the identifying of outliers represents a novel research direction that has received limited attention in the literature. To our knowledge, few papers have

addressed this topic [11] [22] [21]. This paper represents a novel contribution to this field. Our contribution is guided by the following research questions:

- (1) In the outranking context, many alternatives may be compared similarly to a particular alternative. For example, if two alternatives a_i and a_j are both preferred to a third alternative a_k , the primary research question is: which alternative, a_i or a_j , holds a greater preference over a_k ? In other words, does the preference of a_i over a_k carry the same level of importance as the preference of a_j over a_k ?
- (2) The second research question is how to compute this degree of importance?
- (3) The third research question inquires: How can we employ this level of importance to identify outliers in the MCDA context, and what is the effectiveness of this integration compared to existing approaches?

By addressing these research questions, the primary goal of this work is to introduce the concept of importance degree of a preference relation. It also provides a simple way to measure this importance degree. Lastly, it discusses the benefits of integrating this new concept in the process of detecting outliers in the MCDA context.

1.2. Our contribution

Our contribution in this paper can be divided into two parts:

Part 1: We propose the concept of importance degree which represents the degree of significance of a preference relation between pairs of alternatives generated by the PROMETHEE method. To calculate the importance degree, we represent each alternative a_i in multidimensional space by a vector of its multicriteria preference indices: $a_i = (\pi(a_i, a_1), \dots, \pi(a_i, a_n), \pi(a_1, a_i), \dots, \pi(a_n, a_i))$, where $\pi(a_i, a_j)$ denotes the preference value of a_i over a_j . The importance degree of the preference relation between two alternatives is quantified based on the Euclidean distance between them.

Part 2: We present a novel idea for detecting outliers in the MCDA field, as well as a new definition of an outlier in the MCDA context. In this approach, we propose to quantify the similarity between alternatives to measure how similar an alternative aligns with the rest of the alternatives. This measure is determined based on the concept of the importance degree of the preference relation. Then, according to the distribution of the similarities of alternatives, we apply an appropriate statistical technique to identify outlier alternatives.

1.3. Paper organization

The remainder of the paper is structured as follows: Section 2 presents essential concepts within the MCDA domain and introduces the PROMETHEE I method. In addition, we provide crucial background details about the methods utilized in this work, including the LOF algorithm, and statistical techniques for detecting outliers. Section 3 gives an overview of existing work in this field. The proposed approach is presented in Section 4 along with an example of application. Section 5 presents an evaluation of the proposed approach using two real-life problems. Section 6 discusses the results. The paper concludes in Section 7.

2. Background

In this section, our aim is to provide essential background information about the concepts and methods employed in this work, aiming to elucidate the underlying problem statement.

2.1. Our contribution

In MCDA, preference relations serve as valuable tools for expressing the preference information of decision makers in the decision-making process. When comparing two alternatives a_i and $a_j \in A$, typically, three preference relations are distinguished [4]:

- a_i is preferred to a_j (denoted as $a_i P a_j$): corresponds to the presence of evident and positive reasons to favor alternative a_i over a_j ;
- a_i is indifferent to a_j (denoted as $a_i I a_j$): corresponds to the presence of evident and positive reasons that justify the equivalence between a_i and a_j ;
- a_i is incomparable to a_j (denoted as $a_i R a_j$): corresponds to a lack of clear and positive reasons that validate either of the two preceding relations.

In this preference structures $\langle P, I, R \rangle$, each preference relation is characterized by several properties that uniquely define it:

$$\begin{cases} P \text{ is asymmetric} \\ I \text{ is reflexive, symmetric} \\ R \text{ is irreflexive, symmetric} \end{cases}$$

In this work, we utilize a particular MCDA method that relies on positive and negative flow scores to generate valued preference relations. This method, known as PROMETHEE I, will be illustrated in the next section.

2.2. PROMETHEE

The PROMETHEE methods have been developed to assist decision-makers in ranking a set of alternatives evaluated on multiple criteria, either partially (PROMETHEE I) or completely (PROMETHEE II). In this section, we will focus on PROMETHEE I. Further information about this method can be found in [5].

In the first step, we calculate the difference between the evaluations of a pair of alternatives (a_i, a_j) in $A \times A$ on each criterion c in the range from 1 to k :

$$d_c(a_i, a_j) = f_c(a_i) - f_c(a_j) \quad \forall c \in 1, \dots, k$$

Where $f_c(a_i)$ is the evaluation of alternative a_i on criterion c . To translate these differences in terms of preference, we apply the preference function P . The value of $P_c[d_c(a_i, a_j)]$ represents the degree of preference of a_i over a_j for the criterion c , where P_c is a positive non-decreasing preference function ranging between 0 and 1. The original PROMETHEE method [5] provides six variants of preference functions, granting decision makers the flexibility to model their preferences.

The second step consists of providing for each pair (a_i, a_j) in $A \times A$ a multicriteria preference index:

$$\pi(a_i, a_j) = \sum_{c=1}^k w_c P_c(a_i, a_j), \quad (\sum_{c=1}^k w_c = 1) \quad (1)$$

where w_c ($c = 1, 2, \dots, k$) are normed weights associated to the criteria. This index is a positive real number between 0 and 1, representing a global measure of the preference of one alternative over another.

In the final step, for each alternative $a_i \in A$ we compute the outranking flow scores, positive and negative. The positive flow score indicates how much alternative a_i is preferred to the other alternatives.

$$\phi^+(a_i) = \sum_{a_j \in A} \pi(a_i, a_j) \quad (2)$$

The negative flow score indicates how much the other alternative are preferred to alternative a_i .

$$\phi^-(a_i) = \sum_{a_j \in A} \pi(a_j, a_i) \quad (3)$$

The partial preorder of PROMETHEE I is then defined as follows:

$a_i P a_j$ (preference)

$$\text{if } \begin{cases} \phi^+(a_i) \geq \phi^+(a_j) \wedge \phi^-(a_i) < \phi^-(a_j) \\ \vee \\ \phi^+(a_i) > \phi^+(a_j) \wedge \phi^-(a_i) \leq \phi^-(a_j) \end{cases}$$

$a_i I a_j$ (indifference)

$$\text{if } \phi^+(a_i) = \phi^+(a_j) \wedge \phi^-(a_i) = \phi^-(a_j)$$

$a_i R a_j$ (incomparability)

otherwise.

2.3. Statistic methods

Our approach is based on statistical methods to detect outliers. We apply the standard deviation method (SD) when the data follows a normal distribution, while for non-normally distributed data, we utilize the interquartile range method (IQR).

2.3.1. Interquartile range method (IQR)

This method [9] identifies outliers based on the spread of the middle 50% of the data. To identify outliers using IQR method, the following procedure is followed:

- Calculate the interquartile range (IQR), which is the difference between the third quartile (Q3) and the first quartile (Q1).

- Define a lower bound ($Q1 - t1 * IQR$) and an upper bound ($Q3 + t1 * IQR$), where ' $t1$ ' is a parameter often set to 1.5 or 3 (depending on the desired sensitivity to outliers). In our approach, we set it to 1.5
- Identify outliers as data points that fall below the lower bound or above the upper bound.

2.3.2. Standard Deviations method (SD)

This method [2] identifies outliers based on their deviation from the mean of the data. SD method is based on the following steps:

- Calculate the mean (μ) and standard deviation (σ) of the data.
- Define a lower bound ($\mu - t2 * \sigma$) and an upper bound ($\mu + t2 * \sigma$), where ' $t2$ ' is a parameter usually set to 2 or 3. We configure this parameter to 1.5 in our approach.
- Identify outliers as data points that fall below the lower bound or above the upper bound.

2.4. LOF (Local Outlier Factor) method

The Local Outlier Factor (LOF) method [6] is an unsupervised outlier detection algorithm used to identify outliers in datasets. Developed by Breunig et al , it focuses on the local density deviation of data points relative to their neighbours. Below, we outline the process by which the LOF method operates:

- 1 *K-distance neighborhood (KDN)*: The concept of K-distance neighborhood plays a vital role in the LOF approach, which relies on local density estimation to uncover clusters and anomalies. This concept is based on two key definitions. First, the K-distance of a point p refers to the distance between p and its K-th closest neighbor in the dataset. Second, the K-distance neighborhood of p consists of all data points whose distance to p is less than or equal to its K-distance.

$$N_K(p) = \{o \mid D(p, o) \leq K - distance(p)\} \quad (4)$$

Where:

- $N_K(p)$ is the K-distance neighborhood of p .
 - $D(p, o)$ is the distance between points p and o .
 - $K-distance(p)$ is the distance to the K-th nearest neighbor of p .
- 2 *Reachability Distance*: The reachability distance of a point p with respect to a neighboring point q is defined as the maximum of the distance between p and q and the K-distance neighborhood of q . This distance measures how reachable p is from q while considering the density of q .

$$reach - distance(p, q) = \max(distance(p, q), KDN(q)) \quad (5)$$

- 3 *Local Reachability Density*: The local reachability density of a point p is calculated as the inverse of the average reachability distance to its neighbors. This value reflects how densely surrounded p is by its neighbors.

$$Lrd(p) = 1 / \left(\frac{\sum_{q \in N_k(p)} reach-distance(p,q)}{|N_k(p)|} \right) \quad (6)$$

where $N_k(p)$ is the set of k -nearest neighbors of p

- 4 *Local Outlier Factor (LOF) Calculation*: The LOF of a point p measures its degree of outlierness compared to its neighbors. It is computed as the ratio of the average local reachability density of p to the local reachability densities of its neighbors. A high LOF indicates that p is less dense compared to its neighbors, suggesting it may be an outlier.

$$LOF(P) = \frac{\sum_{q \in N_k(p)} \frac{Lrd(q)}{Lrd(p)}}{|N_k(p)|} \quad (7)$$

- 5 *Outlier Identification*: Data points with high LOF values are considered outliers. The threshold for determining outliers can be set based on domain knowledge or by analyzing the distribution of LOF scores.

3. Related work

As previously mentioned, outlier detection within the MCDA field has not received significant attention in the existing literature. To the best of our knowledge, only three papers have tackled this particular aspect. The first approach has been proposed by De Smet & al [11]. The authors proposed a model relies on distance metrics introduced by Y. De Smet and L. M. Guzmán in [10]. This distance is determined by defining a profile for each alternative based on their relationships with others, and computing the distance between alternatives based on shared relations. By examining the distribution of distance values across different samplings of alternatives, outliers are identified through bi-modal distributions in these distance values. Nevertheless, determining parameters in this approach, such as the size of subsets and the number of repetitions, could pose a challenge and potentially impact the robustness and reliability of the outlier detection process.

In the second paper, Rouba and Nait Bahloul [22] proposed to define the distribution vector for each alternative, wherein the preference relations of one alternative with the others determine its distribution vector. To detect outliers, the local outlier factor (LOF) algorithm [6] is applied to the distribution vectors of the alternatives. However, the robustness of the outcomes of this approach depends on the value assigned to the parameter k in the LOF method as well as the selected multicriteria method.

The third paper's author utilized statistical techniques to find outliers [21]. First, the PROMETHEE II method is used for generating a net-flow value for each alternative. Then, based on the distribution of the net-flow values, the appropriate statistical technique is applied: the Standard Deviation method is utilized if the distribution follows a normal

distribution; otherwise, the interquartile range method is employed. However, in cases where indifferent alternatives exist (sharing the same net-flow value), such instances can result in a lack of normality.

3.1. Our work positioning

In our research, we aim to tackle the obstacles that the previous studies have encountered. One significant aspect of our approach involves utilizing statistical methods that offer:

- The capability to identify numerous outliers within the dataset.
- The ability to differentiate outliers without erroneously labeling them when they are not present.

This sets our approach apart from the methods discussed in the first [11] and second [22] papers. Furthermore, utilizing net-flow values in the third paper may not be accurate and could result in a lack of normality when encountering scenarios where multiple alternatives possess identical net-flow values. Our approach addresses this by incorporating similarity measures rather than net-flow values, which can distinguish between such alternatives more effectively.

4. The proposed approach

Our contribution is structured into two parts. In the first part, we will introduce the concept of degree of importance of a preference relation. In the second part, we will present an approach for detecting outliers in the MCDA field, utilizing the concept of degree of importance.

4.1. Degree of importance of a preference relation

In this section, we aim to clarify the concept of the importance degree in a preference relation between two alternatives as generated by the PROMETHEE method. Additionally, we describe the methodology used to quantify this importance degree. In the PROMETHEE method, the identification of the type of a preference relation between two alternatives depends on their negative flow score (ϕ^-) and positive flow score (ϕ^+). The process of quantifying the degree of importance a preference relation involves plotting each alternative in an orthogonal plane based on its negative and positive flow scores values, to analyze their preference relations. For clarity, let's denote alternative a_i as a reference point for comparison, as depicted in Figure. 1. To simplify notation, the points (0,0), (1,0), (1,1), and (0,1) have been respectively marked as x , y , z , and w .

The preference relation between a_i and any other potential alternative is determined based on their positions in relative to each other. The graph will be divided into colored regions. Each region indicates a specific preference relation with a_i :

- The dark blue region contains all alternatives that are preferred to a_i .
- The light blue region contains all alternatives for which a_i is preferred.

- The white region contains all alternatives that are incomparable to a_i
- Alternatives that are indifferent to a_i have the same coordinates as a_i .

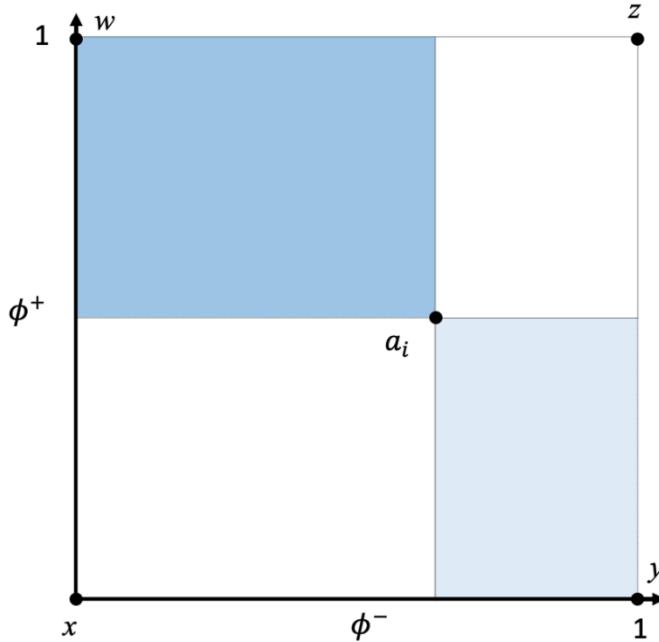


Figure 1. Preference relation regions for each alternative in the graph of ϕ^+ and ϕ^-

In our approach, we define the degree of importance of a preference relation between two alternatives as the normed distance between these two alternatives. As the distance increases, the importance of this relation becomes more significant.

To standardize the distances, we aim to identify the maximum distance for each type of preference relation. In cases of incomparability, it is evident that the maximum distance between two incomparable alternatives is the distance between alternatives x and z . Selecting any alternative within the white regions ensures that the distance from that alternative to a_i does not exceed the distance between x and z . This remains consistent regardless of the chosen point on the plane; the distance between that point and any points within its white regions does not surpass the distance from x to z . Similarly, in cases of preference, the maximum distance between two alternatives, where one is preferred to the other, is the distance between y and w . However, in cases of indifference, any two indifferent alternatives share the same point, indicating that the distance between them is always zero. This suggests that their relation importance is undefined.

The use of positive and negative flow scores to represent alternatives reveals limitations in addressing all preference scenarios. This leads to explore other representations in the following sections.

4.1.1. Analysing the existing representations in literature

In this section, we explore the various representations of alternatives found within the existing literature. For instance, in PROMETHEE GAIA [17], the authors utilize the unicriterion net flows as coordinates to represent alternatives in multidimensional space:

$$\phi^c(a_i) = \frac{1}{n-1} \sum_{j \in A} (\pi^c(a_i, a_j) - \pi^c(a_j, a_i)) \quad (8)$$

where $\phi^c(a_i)$ represent the unicriterion net flows of a_i on criterion c .

Similarly, in [8], the authors utilized noncriterion net flows as coordinates to represent alternatives. Through calculating and analyzing the distances between alternatives, they could evaluate the efficacy of their approach (PROMETHEE γ) in comparison to PROMETHEE I.

Using unicriterion net flows as coordinates may result in different alternatives having identical coordinates. Therefore, this could lead to inaccurate evaluations of the importance of preference relations in certain scenarios. For example, let's consider the alternatives x ($\phi^+ = 0, \phi^- = 0$) and z ($\phi^+ = 1, \phi^- = 1$). In this case, according to formulas (1), (2), (3), and (8), we determine that all their unicriterion net flows are equal to zero ($\forall c \in k, \phi^c(x) = 0$ and $\phi^c(z) = 0$). This implies that in cases of incomparability, the maximum distance equals zero (the distance between x and z). Consequently, we are unable to define the importance of preference relations in these cases.

4.1.2. Our representation of alternatives

In this section, we propose a novel approach to represent alternatives for evaluating the importance of their preference relations. The concept involves presenting each alternative, denoted as a_i , in multidimensional space using its multicriteria preference indices as coordinates:

$$a_i = (\pi(a_i, a_1), \dots, \pi(a_i, a_n), \pi(a_1, a_i), \dots, \pi(a_n, a_i)) \quad (9)$$

where $\pi(a_i, a_j)$ denotes the preference value of a_i over a_j , and the following properties obviously hold for this preference values, $\forall a_i, a_j \in A$:

- (1) $\pi(a_i, a_i) = 0$
- (2) $0 \leq \pi(a_i, a_j) \leq 1$

For example, let's examine alternative y ($\phi^+ = 1, \phi^- = 0$) (Fig.1), where it holds a multicriteria preference index of 1 over every other alternative, while each alternative has a multicriteria preference index of 0 compared to y . Alternative y will, therefore, be represented by the point $(1_{y,1}, 1_{y,2}, \dots, 1_{y,n}, 0_{1,y}, 0_{2,y}, \dots, 0_{n,y})$ Where:

$1_{i,j}$: denotes the multicriteria preference index of alternative i over alternative j being equal to 1

$0_{i,j}$: represents the multicriteria preference index of alternative i over alternative j being equal to 0.

Using multicriteria preference indices, each alternative x , y , z , or w is represented by a unique point in space, which cannot be achieved using unicriterion net flows. Through this new representation, we aim to normalize the importance degree by identifying the maximum possible distance between any two alternatives for each type of preference relation:

- *Incomparability*: As illustrated in Figure 1, the maximum possible distance between any two incomparable alternatives occurs between x and z .
- *Preference*: The maximum distance between an alternative that is preferred over another is found between y and w .
- *Indifference*: The maximum distance remains undefined, as indifferent alternatives are positioned close to each other. As shown in Figure 1, the greatest distances occur in cases of incomparability (x and z) and preference (y and w). To ensure consistency, we define the maximum distance for indifference as equal to that of y and w or x and z (since they are the same), maintaining the importance degree within the range of 0 to 1.

Thus, the maximum distance across all scenarios is established as $[x, z] = [y, w] = \sqrt{2n - 2}$, where n is the number of alternatives and $n > 1$.

Definition 1

The importance of a preference relation of two alternatives a_i and a_j , is denoted as imp_{ij} , where $imp_{ij} \in]0; 1]$, defined by:

In case of indifference

$$imp_{ij} = 1 - \frac{1}{\sqrt{2n-2}} \sqrt{\sum_{k \in A} ((\pi_{ik} - \pi_{jk})^2 + (\pi_{ki} - \pi_{kj})^2)} \quad (10)$$

in other cases:

$$imp_{ij} = \frac{1}{\sqrt{2n-2}} \sqrt{\sum_{k \in A} ((\pi_{ik} - \pi_{jk})^2 + (\pi_{ki} - \pi_{kj})^2)} \quad (11)$$

In cases of indifference, the importance equals 1 when alternatives have similar multicriteria preference indices, with a distance equal to 0 between them. Conversely, in cases of incomparability or preference, the importance equals 1 when alternatives have the maximum distance between them. Therefore, a preference relation between two alternatives a_i and a_j is represented as follows:

$$a_i \text{ } imp_{i,j} \text{ } S \text{ } a_j$$

where S is a preference relation $\in \{P^+, P^-, R, I\}$, and $imp_{i,j}$ is the importance of preference relation between a_i and a_j .

The importance degree satisfies the following properties.

- *Monotonicity*: The degree of importance is monotonic with respect to the distance between alternatives, as represented by their multicriteria preference indices:
 - *In the case of indifference*:
 - As the distance between alternatives decreases, the degree of importance increases. It reaches a value of 1 when the distance is 0.
 - Conversely, as the distance increases, the degree of importance decreases, asymptotically approaching 0.
 - *In other cases*:
 - The degree of importance increases as the distance between alternatives increases, reaching a value of 1 at the maximum distance.
 - Conversely, as the distance decreases, the degree of importance decreases, approaching 0.
- *Reflexivity*: The degree of importance satisfies reflexivity for self-comparisons. When comparing an alternative a_i with itself, the degree of importance is always 1,

$$imp_{i,i} = 1$$

as there is no distance between the preference indices of the alternative.

- *Symmetry*: The importance degree is symmetric by definition:

$$imp_{i,j} = imp_{j,i}$$

This property holds because the importance of a preference relation between a_i and a_j is invariant under their order, as the distance metric used in the calculation is symmetric.

- *Compensation*: According to formulas (10) and (11), the sum involves independent squared differences. Squaring ensures that each term is non-negative, meaning increases in one term cannot be directly canceled by decreases in another, leading to a non-compensatory importance degree.

4.1.3. Illustrative example

In this section, we will present an illustrative example to elucidate the computation of preference relation importance. Let us consider the decision problem represented in Table 1.

Table 1. A decision problem

Max Min	Min	Max	Min	Min	Min	Max
Parameters	$p = 10$	$p = 30$	$q = 0.5$	$p = 5$	$q = 1$	$p = 6$
Type of Criteria	Quasi- criterion	Criterion with linear preference	Criterion with - linear preference and indifference	Level criterion	Usual criterion	Gaussian criterion
a_6	94	96	7	3.6	5	6
a_5	52	72	6	2.0	3	8
a_4	40	80	10	7.5	7	10
a_3	83	60	4	7.2	4	7
a_2	65	58	2	9.7	1	1
a_1	80	90	6	5.4	8	5
Criteria	f_1	f_2	f_3	f_4	f_5	f_6

Using PROMETHEE, the multicriteria preference indices for each alternative are summarized in Table2. The positive and negative flow scores are presented in Table 3.

Table 2. The multicriteria preference indexes

	a_1	a_2	a_3	a_4	a_5	a_6
a_1	-	0.296	0.250	0.269	0.100	0.185
a_2	0.463	-	0.389	0.333	0.296	0.500
a_3	0.235	0.180	-	0.333	0.056	0.429
a_4	0.399	0.506	0.305	-	0.224	0.212
a_5	0.444	0.515	0.487	0.380	-	0.448
a_6	0.287	0.399	0.250	0.431	0.133	-

Table 3. Positive and negative outranking flows

	a_1	a_2	a_3	a_4	a_5	a_6
ϕ^+	0.220	0.396	0.247	0.329	0.455	0.300
ϕ^-	0.366	0.379	0.336	0.349	0.162	0.355

The initial step in computing the degree of importance involves representing each alternative by its vector as defined in formula 1. Table 4 summarizes all alternative vectors.

Table 4. The alternative vectors

<i>Alternative</i>	<i>Vector</i>
a_1	(0.296, 0.250, 0.269, 0.100, 0.185, 0.463, 0.235, 0.399, 0.444, 0.287)
a_2	(0.463, 0.389, 0.333, 0.296, 0.500, 0.296, 0.180, 0.506, 0.515, 0.399)
a_3	(0.235, 0.180, 0.333, 0.056, 0.429, 0.250, 0.389, 0.305, 0.487, 0.250)
a_4	(0.399, 0.506, 0.305, 0.224, 0.212, 0.269, 0.333, 0.333, 0.380, 0.431)
a_5	(0.444, 0.515, 0.487, 0.380, 0.448, 0.100, 0.296, 0.056, 0.224, 0.133)
a_6	(0.287, 0.399, 0.250, 0.431, 0.133, 0.185, 0.500, 0.429, 0.212, 0.448)

By employing formulas 10 and 11, we compute the degree of importance of each preference relation, as shown in table 5.

Table 5. Importance degrees of preference relations

	a_1	a_2	a_3	a_4	a_5	a_6
a_1	I	0.158 R	0.125 P ⁻	0.130 P ⁻	0.247 P ⁻	0.191 P ⁻
a_2	0.158 R	I	0.164 R	0.134 R	0.214 P ⁻	0.201 R
a_3	0.125 P ⁺	0.164 R	I	0.161 R	0.215 P ⁻	0.208 R
a_4	0.130 P ⁺	0.134 R	0.161 R	I	0.183 P ⁻	0.122 P ⁺
a_5	0.247 P ⁺	0.214 P ⁺	0.215 P ⁺	0.183 P ⁺	I	0.220 P ⁺
a_6	0.191 P ⁺	0.201 R	0.208 R	0.122 P ⁻	0.220 P ⁻	I

As an illustration, a_1 is incomparable to a_2 , and to determine the importance degree of this relation, we first compute the Euclidean distance between a_1 and a_2 , represented by their vectors in Table 4.

$$\sqrt{\sum_{k \in A} (\pi_{1k} - \pi_{2k})^2 + (\pi_{k1} - \pi_{k2})^2} = 0.500$$

Given that the number of alternatives in this example is 6, we calculate:

$$\frac{1}{\sqrt{2n-2}} = \frac{1}{\sqrt{2*6-2}} = 0.316$$

Hence, the importance measure is computed as:

$$imp_{1,2} = 0.316 * 0.500 = 0.158$$

According to Definition 1, their relationship is expressed as a_1 0.158 R a_2 .

Table 5 demonstrates that alternative a_5 is preferred to all other alternatives. Even though the relations between a_5 and other alternatives are similar, their levels of importance vary. This disparity enables us to differentiate between similar preference relations.

In the following section, we utilize this concept of importance as the foundation for introducing a new approach of detecting outliers based on preference relations.

4.2. Outlier detection

In this section, we introduce a method for identifying outliers in the MCDA domain. Our approach involves measuring the similarity between alternatives using the concept of the degree of importance of a preference relation. Based on the distribution of these similarities, we employ the appropriate statistical method for outlier detection: either the IQR method or SD method.

4.2.1. Outlier in MCDA field

Generally, an outlier is an observation that deviates markedly from other observations in a dataset [1]. However, the concept of an outlier can vary depending on the field of study or context in which it is used. In the MCDA field, for instance, the evaluation of an alternative depends on its preference relations with the other alternatives, making these preference relations the pivotal element for identifying outliers in this domain. In [22], the authors characterize an outlier as any alternative having a number of specific preference relations with other alternatives (P+, P-, or R) that equal to or greater than a given percentage threshold.

In this study, we propose a new understanding of outliers in the MCDA field, centered around the notion of similarity. This concept quantifies the extent of resemblance of an alternative to the remainder of the dataset.

Definition 2

Let's denote three alternatives as a_i, a_j , and a_k . We consider that alternatives a_i and a_j are similar with respect to a_k . If a_i and a_j have the same preference relation with a_k , we term this similarity ($similarity(a_i, a_j)_{a_k}$), and it is computed as follows:

$$similarity(a_i, a_j)_{a_k} = 1 - \frac{|imp_{i,k} - imp_{j,k}|}{\max(imp_{i,k}, imp_{j,k})} \quad (12)$$

where $imp_{i,k}$ is the importance degree of the preference relation between alternatives a_i and a_k

According to Definition 2, we can deduce the similarity between two alternatives a_i and a_j by averaging their similarities with respect to all other alternatives:

$$similarity(a_i, a_j) = \frac{\sum_{a_k \in A} similarity(a_i, a_j)_{a_k}}{n-2} \quad (13)$$

where n represents the number of alternatives of the set A

Likewise, the similarity of an alternative a_i and all the remaining alternatives within A is given by:

$$similarity(a_i)_A = \frac{\sum_{a_j \in A} similarity(a_i, a_j)}{n-1} \quad (14)$$

As previously defined, the similarity between two alternatives with respect to a third serves as the foundation for both the similarity between two alternatives and the similarity of an alternative. Therefore, we will focus solely on the properties of this similarity in what follows:

- *Monotonicity*: The similarity will increase as the absolute difference $|imp_{i,k} - imp_{j,k}|$ decreases. This is because the term inside the absolute value is subtracted

from 1. Therefore, if the importance degrees $imp_{i,k}$ and $imp_{j,k}$ get closer, the value of the similarity increases, which corresponds to the idea of monotonicity.

- *Reflexivity*: If $a_i = a_j$, then $imp_{i,k} = imp_{j,k}$, and:

$$similarity(a_i, a_i)_{a_k} = 1 - \frac{0}{\max(imp_{i,k}, imp_{i,k})} = 1$$

This satisfies reflexivity, as the similarity of an alternative with itself is always 1.

- *Symmetry*: In the formula (12), the term $|imp_{i,k} - imp_{j,k}|$ is symmetric because the absolute difference remains the same regardless of the order of a_i and a_j . Therefore, the similarity measure satisfies the symmetry property.

$$similarity(a_i, a_j)_{a_k} = similarity(a_j, a_i)_{a_k}$$

- *Compensation*: The similarity between a_i and a_j with respect to a_k depends solely on the difference $|imp_{i,k} - imp_{j,k}|$. Therefore, compensation is not inherent at this level of similarity computation, as no other factors can balance a decrease in similarity due to a larger $|imp_{i,k} - imp_{j,k}|$.
- *Normalization*:
 - If $|imp_{i,k} - imp_{j,k}| = 0$, we have:

$$similarity(a_i, a_j)_{a_k} = 1 - \frac{0}{\max(imp_{i,k}, imp_{j,k})} = 1.$$

- When $|imp_{i,k} - imp_{j,k}|$ approaches its maximum, it gets close to $\max(imp_{i,k}, imp_{j,k})$ but cannot reach it, as the importance degree are always nonzero (Definition 1). Therefore:

$$similarity(a_i, a_j)_{a_k} \rightarrow 0 \text{ as } |imp_{i,k} - imp_{j,k}| \rightarrow \max(imp_{i,k}, imp_{j,k})$$

Thus, $similarity(a_i, a_j)_{a_k} \in]0; 1]$

After studying the properties of the proposed similarity measure, the next step aims to identify outliers in the MCDA, we rely on the similarity of an alternative to all other alternatives, as it represents the extent to which an alternative resembles the others.

Definition 3

A multicriteria outlier is an alternative whose similarity significantly deviates from the general pattern of similarities within the set of alternatives. It is characterized by a near-zero similarity score relative to other alternatives, indicating a substantial lack of resemblance to the rest.

Table 6 presents the similarity scores of alternatives of our example, calculated through the application of formula (14) to the data in Table 5.

Table 6. The similarity of alternatives

Alternative	a_1	a_2	a_3	a_4	a_5	a_6
Similarity(.)	0.432	0.454	0.638	0.631	0.148	0.659

According to definition 3, an outlier is characterized by having a similarity that is significantly closer to zero compared to the rest. To evaluate the degree of closeness of an alternative to being an outlier, we apply statistical methods as follows.

- Test the normality of the alternative similarities
- If similarities follow a normal distribution, we utilize the SD method; otherwise, we employ the IQR method.

To test the normality of the data, we utilize the Shapiro-Wilk test [25]. This test evaluates how well the data fit the characteristics of a normal distribution by examining the null hypothesis that the data are normally distributed. If the p-value associated with the test is less than a chosen significance level (often 0.05), we reject the null hypothesis, indicating that the data do not follow a normal distribution. Conversely, if the p-value exceeds the chosen significance level, there is insufficient evidence to reject the null hypothesis, suggesting that the data may indeed be normally distributed.

For our example, the Shapiro–Wilk test of the similarities in Table 6 leads to a p-value of 0.1556 which means that the alternative similarities follow a normal distribution. In this scenario, the SD method is utilized, yielding the following outcomes: the mean (μ) = 0.493 and standard deviation (σ) = 0.178. Consequently, alternatives with similarities beyond the range of [0.226, 0.760] are identified as outliers. In this instance, alternative a_5 , with a similarity of 0.148, is considered an outlier.

5. Application and results

Our proposed approach will be tested in two phases in this section. In the first phase, our focus is solely on illustrating the effectiveness of the degree of importance concept by integrating it into another approach, while the second phase will assess the entirety of the approach.

5.1. Phase1

In this phase, we integrate the degree of importance concept into the LOF-based outlier detection method as outlined in [22]. By comparing the results before and after integration, we can evaluate the efficacy of this concept using a real-world problem: the world happiness ranking.

5.1.1. Integrating the importance degree

As noted earlier, this technique for outlier detection applies the LOF algorithm to the distribution vectors of alternatives. These vectors are constructed based on the preference relations of each alternative with the set of alternatives A, as follows:

$$V_1(a_i) = |\{a_j \in A, a_i \neq a_j \wedge a_i I a_j\}|$$

$$V_2(a_i) = |\{a_j \in A, a_i \neq a_j \wedge a_i P a_j\}|$$

$$V_3(a_i) = |\{a_j \in A, a_i \neq a_j \wedge a_j P a_i\}|$$

$$V_4(a_i) = |\{a_j \in A, a_i \neq a_j \wedge a_i R a_j\}|$$

By integrating the importance degree, instead of calculating a preference relation as 1 in the distribution vectors, we compute its degree of importance. Consequently, the distribution vector will be as follows:

$$V_1(a_i) = \sum_{a_j \in A, a_i \neq a_j \wedge a_i I a_j} imp_{i,j}$$

$$V_2(a_i) = \sum_{a_j \in A, a_i \neq a_j \wedge a_i P a_j} imp_{i,j}$$

$$V_3(a_i) = \sum_{a_j \in A, a_i \neq a_j \wedge a_j P a_i} imp_{i,j}$$

$$V_4(a_i) = \sum_{a_j \in A, a_i \neq a_j \wedge a_i R a_j} imp_{i,j}$$

5.1.2. The world happiness ranking problem

In this section, we address a real-life problem known as the world happiness ranking. The problem aims to rank countries based on various criteria that contribute to happiness and quality of life. Six key criteria are used to measure happiness, including GDP (Gross Domestic Product per capita), Family (Social support), Life (Healthy Life Expectancy), Freedom (to make life choices), Trust and Generosity. Our experimentations have been conducted on the 2015 world happiness report [28], which includes 158 countries.

Our experimentation proceeded as follows: we conducted experiments in three different scenarios. In each scenario, we applied the LOF-based outlier detection method both with and without the integration of the degree of importance. In the first scenario, we considered only the top 20 ranked alternatives, as well as the last ranked alternative. For the second scenario, we examined the 20 alternatives ranked in the middle, along with both the first and last ranked alternatives. Finally, in the third scenario, we focused on the 20 alternatives ranked lowest along with the first ranked alternative.

With this strategy for data selection, we seek to make the first and last ranked alternatives stand out as outliers depending on the scenario:

- In the first scenario, the last ranked alternative is the outlier.
- In the second scenario, the first and last ranked alternatives are the outliers.
- In the third scenario, the first ranked alternative is the outlier.

Initially, the PROMETHEE method is applied to rank the alternatives, as shown in Appendix A. Thereafter, we use it in each scenario to:

- Establish preference relations (see Appendices B, C, and D).
- Calculate the multicriteria preference indices for each alternative, thus determining the degree of importance of the established preference relations.

The level preference function has been utilized as the preference function, and for simplicity, the first and third quartiles of the evaluation differences have been used as the indifference and preference thresholds, respectively. Table 7 summarizes the parameters utilized for the execution of the PROMETHEE method.

Table 7. Parameters used for the execution of the PROMETHEE method

	GDP	Family	Life	Freedom	Trust	Generosity
Weight	0.1666667	0.1666667	0.1666667	0.1666667	0.1666667	0.1666667
Indifferent threshold (q)	0.105	0.072	0.042	0.044	0.034	0.052
Preference threshold (p)	0.412	0.295	0.216	0.206	0.125	0.205
Type of Criteria	Level criterion	Level criterion	Level criterion	Level criterion	Level criterion	Level criterion
Max/Min	max	max	max	max	max	max

For the sake of sensitivity analysis, the LOF based approach was implemented with various values of the parameter k . In this experiment, Tables 8, 10, and 12 present the distribution matrix of alternatives for the three scenarios, both with and without using the degree of importance concept. The respective results of the LOF algorithm, as shown in Tables 9, 11, and 13, indicate the three highest scores produced

5.1.2.1 Scenario 1 (The top 20 alternatives and the lowest-ranked alternative)

Table 8. The distribution matrix for two cases, with and without the degree of importance, in Scenario 1

Alternatives	Without the degree of importance				With the degree of importance			
	V ₁ (.)	V ₂ (.)	V ₃ (.)	V ₄ (.)	V ₁ (.)	V ₂ (.)	V ₃ (.)	V ₄ (.)
New Zealand	1	14	0	6	0	2.272	0	0.716
Australia	1	12	1	7	0	1.960	0.072	0.755
Canada	1	7	5	8	0	1.488	0.388	0.897
Ireland	1	6	5	9	0	1.333	0.465	0.886
Norway	1	12	1	7	0	2.077	0.082	0.705
Netherlands	1	6	7	7	0	1.296	0.664	0.766
Sweden	1	8	2	10	0	1.600	0.144	1.060
United Kingdom	1	7	5	8	0	1.433	0.538	0.856
Denmark	1	13	1	6	0	2.185	0.063	0.595
Switzerland	1	15	0	5	0	2.422	0	0.505
Luxembourg	1	5	7	8	0	1.111	0.707	0.741
Hong Kong	1	5	3	12	0	1.120	0.313	1.209
Iceland	1	6	4	10	0	1.188	0.435	0.961
Qatar	1	5	0	15	0	1.243	0	1.887
Finland	1	5	10	5	0	1.081	0.998	0.516
Singapore	1	3	0	17	0	0.905	0	2.573
Austria	1	1	15	4	0	0.530	2.239	0.453
Malta	1	2	16	2	0	0.626	2.445	0.183
Germany	1	1	15	4	0	0.523	2.370	0.456
United States	1	1	17	2	0	0.462	3.238	0.208
Burundi	1	0	20	0	0	0	11.694	0

Table 9. The outcomes of the approach proposed in [22] with and without importance degree in Scenario 1

k	Without the degree of importance			With the degree of importance		
	1 st	2 nd	3 rd	1 st	2 nd	3 rd
2	Singapore	Qatar	Finland	Burundi	Singapore	United States
3	Singapore	Qatar	Burundi	Burundi	Singapore	United States
4	Singapore	Burundi	United States	Burundi	United States	Singapore
5	Burundi	United States	Malta	Burundi	United States	Singapore
6	Burundi	United States	Malta	Burundi	United States	Malta
7	Burundi	United States	Malta	Burundi	United States	Malta

5.1.2.2 Scenario 2 (the 20 alternatives ranked in the middle, along with both the first and last ranked alternatives).

Table 10. The distribution matrix for two cases, with and without the importance degree, in Scenario 2

Alternatives	Without the degree of importance				With the degree of importance			
	V ₁ (.)	V ₂ (.)	V ₃ (.)	V ₄ (.)	V ₁ (.)	V ₂ (.)	V ₃ (.)	V ₄ (.)
New Zealand	1	21	0	0	0	9.595	0	0
Slovakia	1	17	1	3	0	3.023	0.384	0.535
Venezuela	1	17	1	3	0	2.553	0.382	0.525
Philippines	1	13	3	5	0	2.004	0.692	0.745
Jamaica	1	11	4	6	0	1.712	0.754	0.854
Vietnam	1	10	4	7	0	1.527	0.792	0.984
Hungary	1	6	7	8	0	1.142	1.141	1.362
Jordan	1	9	4	8	0	1.317	0.768	1.171
Latvia	1	10	3	8	0	1.452	0.665	1.254
Rwanda	1	1	1	19	0	0.300	0.453	3.458
Cambodia	1	1	3	17	0	0.313	0.849	2.776
China	1	5	8	8	0	0.778	1.201	1.046
Myanmar	1	1	1	19	0	0.304	0.446	3.447
Kyrgyzstan	1	2	7	12	0	0.446	1.212	1.460
Bolivia	1	1	12	8	0	0.350	1.938	1.077
Lebanon	1	3	11	7	0	0.545	1.751	0.872
Turkey	1	6	6	9	0	0.860	1.186	1.210
Azerbaijan	1	1	12	8	0	0.348	1.810	1.026
Macedonia	1	2	9	10	0	0.435	1.678	1.308
Russia	1	1	11	9	0	0.333	1.902	1.253
Peru	1	1	10	10	0	0.364	1.743	1.339
Burundi	1	0	21	0	0	0	7.955	0

Table 11. The outcomes of the approach proposed in [22] with and without importance degree in Scenario 2

k	Without the degree of importance			With the degree of importance		
	1 st	2 nd	3 rd	1 st	2 nd	3 rd
2	Burundi	Rwanda	Myanmar	Burundi	New Zealand	Rwanda
3	Burundi	Rwanda	Myanmar	Burundi	New Zealand	Rwanda
4	Burundi	Rwanda	Myanmar	Burundi	New Zealand	Rwanda
5	Burundi	Rwanda	Myanmar	Burundi	New Zealand	Rwanda
6	Burundi	Rwanda	Myanmar	Burundi	New Zealand	Rwanda
7	Burundi	New Zealand	Rwanda	New Zealand	Burundi	Rwanda

5.1.2.3 Scenario 3(the 20 alternatives with the lowest ranks along with the first ranked alternative)

Table 12. The distribution matrix for two cases, with and without the importance degree, in Scenario 3.

Alternatives	Without the degree of importance				With the degree of importance			
	V ₁ (.)	V ₂ (.)	V ₃ (.)	V ₄ (.)	V ₁ (.)	V ₂ (.)	V ₃ (.)	V ₄ (.)
New Zealand	1	20	1	1	0	11.497	0	0
Egypt	1	10	1	9	0	1.857	0.485	1.186
Lesotho	1	13	4	3	0	1.579	0.870	0.323
Sierra Leone	1	7	8	5	0	0.878	1.336	0.500
Ghana	1	17	1	2	0	2.459	0.499	0.208
Pakistan	1	10	2	8	0	1.493	0.601	0.866
Comoros	1	15	2	3	0	2.172	0.605	0.276
Zimbabwe	1	12	1	7	0	1.583	0.553	0.805
Guinea	1	7	8	5	0	0.801	1.367	0.499
Malawi	1	4	14	2	0	0.480	2.228	0.260
Benin	1	9	6	5	0	1.083	1.142	0.596
Afghanistan	1	6	13	1	0	0.741	1.930	0.132
Yemen	1	14	3	3	0	1.875	0.714	0.319
Congo	1	3	15	2	0	0.349	2.456	0.188
(Kinshasa)								
Central African Republic	1	0	19	1	0	0	3.678	0.108
Liberia	1	8	7	5	0	0.909	1.225	0.562
Angola	1	2	7	11	0	0.299	1.413	1.332
Madagascar	1	7	8	5	0	0.940	1.349	0.582
Togo	1	2	16	2	0	0.236	2.626	0.208
Chad	1	2	14	2	0	0.275	2.302	0.396
Burundi	1	0	19	1	0	0	4.127	0.108

Table 13. The outcomes of the approach proposed in [22] with and without importance degree in Scenario 3

k	Without the degree of importance			With the degree of importance		
	1 st	2 nd	3 rd	1 st	2 nd	3 rd
2	Liberia	New Zealand	Egypt	New Zealand	Burundi	Liberia
3	New Zealand	Liberia	Ghana	New Zealand	Burundi	Central African Republic
4	New Zealand	Central African Republic	Burundi	New Zealand	Burundi	Central African Republic
5	New Zealand	Central African Republic	Burundi	New Zealand	Burundi	Central African Republic
6	New Zealand	Central African Republic	Burundi	New Zealand	Burundi	Central African Republic
7	New Zealand	Central African Republic	Burundi	New Zealand	Burundi	Central African Republic

5.2. Phase 2

After testing the effectiveness of the concept of the degree of importance separately. In this section, we illustrate our approach to detect outliers using this concept by considering a real-world problem: the Human Development Index (HDI) problem, as adopted by [12]. The United Nations Development Program (UNDP) has introduced the HDI ranking, which evaluates 179 United Nations countries based on three criteria: life expectancy, education, and income index. In this problem, we aren't focusing on the exact ranking of countries; our goal is to calculate their similarities to detect any outliers.

To implement our method, we initially use the PROMETHEE method to determine the preference relations between alternatives and calculate their degree of importance. Table 14 details the type of generalized criterion and the corresponding parameters for each criterion, as specified by the DM (for additional details about the parameters, please refer to De Smet et al. [12]).

Table 14. The indifference and preference thresholds, as well as the weight and preference function of each criterion

Parameters	Life expectancy	Adult literacy index	GDP
preference threshold	0.704	0.719	0.828
Indifference threshold	0	0	0
Weight of criterion	0.333	0.333	0.333
preference function	V-shape	V-shape	V-shape

We seek to test our technique on two cases, one with and one without outliers, and then evaluate the results. In the first case, we used our method on the list of countries without any outliers (The evaluations of these countries are provided in [12]). In the second case, we added an artificial outlier (*named ARTIF*) by assigning it the minimum value for each criterion, as shown in Table 15. Evaluating our approach on data-sets both with and without outliers helps us assess the sensitivity and performance under different conditions.

Table 15. The evaluations of the artificial outlier (*ARTIF*)

	Life expectancy	Adult literacy index	GDP
ARTIF	0.253	0.274	0.172

The similarities of the alternatives in the two cases are presented in Table 16.

Table 16. The similarities of the alternatives in the two cases.

Country	Case1	Case2	Country	Case 1	Case 2	Country	Case 1	Case2
Iceland	0.296	0.297	Venezuela	0.404	0.405	Guatemala	0.37	0.371
Norway	0.298	0.299	Romania	0.406	0.406	Kyrgyzstan	0.373	0.374
Canada	0.299	0.299	Malaysia	0.406	0.406	Vanuatu	0.366	0.367
Australia	0.3	0.301	Montenegro	0.406	0.406	Tajikistan	0.366	0.367
Ireland	0.309	0.31	Serbia	0.405	0.406	South Africa	0.353	0.354
Netherlands	0.309	0.31	Saint_Lucia	0.406	0.406	Botswana	0.348	0.349
Sweden	0.309	0.309	Belarus	0.406	0.406	Morocco	0.346	0.347
Japan	0.312	0.312	Macedonia TFYR	0.411	0.411	Sao Tome and Principe	0.346	0.347
Luxembourg	0.314	0.314	Albania	0.411	0.411	Namibia	0.341	0.342
Switzerland	0.314	0.315	Brazil	0.411	0.411	Congo	0.335	0.336
France	0.313	0.313	Kazakhstan	0.406	0.406	Bhutan	0.332	0.333
Finland	0.315	0.315	Ecuador	0.411	0.411	India	0.331	0.333
Denmark	0.318	0.318	Russian Federation	0.409	0.409	Lao People	0.331	0.334
Austria	0.317	0.318	Mauritius	0.413	0.413	Solomon Islands	0.325	0.326
United States	0.32	0.321	Bosnia and Herzegovina	0.413	0.413	Myanmar	0.324	0.327
Spain	0.319	0.319	Turkey	0.413	0.413	Cambodia	0.319	0.32
Belgium	0.319	0.319	Dominica	0.412	0.413	Comoros	0.319	0.321
Greece	0.321	0.321	Lebanon	0.412	0.412	Yemen	0.315	0.317
Italy	0.322	0.322	Peru	0.411	0.412	Pakistan	0.314	0.316
New Zealand	0.322	0.323	Colombia	0.413	0.412	Mauritania	0.313	0.315
United Kingdom	0.324	0.324	Thailand	0.411	0.411	Swaziland	0.301	0.305
Hong Kong, China(SAR)	0.327	0.327	Ukraine	0.413	0.414	Ghana	0.304	0.306
Germany	0.325	0.326	Armenia	0.414	0.414	Madagascar	0.307	0.309
Israel	0.331	0.331	Iran	0.414	0.414	Kenya	0.304	0.306
Korea(Republic of)	0.333	0.333	Tonga	0.414	0.413	Nepal	0.306	0.308
Slovenia	0.336	0.336	Grenada	0.414	0.414	Sudan	0.3	0.302

Brunei	0.343	0.344	Jamaica	0.414	0.414	Bangladesh	0.302	0.305
Darussalam						Haiti	0.3	0.302
Singapore	0.342	0.343	Belize	0.409	0.409	Papua New Guinea	0.297	0.299
Kuwait	0.347	0.347	Suriname	0.412	0.414	Cameroon	0.295	0.297
Cyprus	0.344	0.345	Jordan	0.414	0.414	Djibouti	0.295	0.296
United Arab Emirates	0.35	0.35	Dominican Republic	0.412	0.413			
			Saint Vincent and the Grenadines	0.413	0.413	Tanzania	0.291	0.294
Bahrain	0.355	0.355	Georgia	0.412	0.412	Senegal	0.292	0.294
Portugal	0.353	0.353	China	0.411	0.411	Nigeria	0.289	0.291
Qatar	0.357	0.358	Tunisia	0.41	0.412	Lesotho	0.287	0.289
Czech Republic	0.355	0.356	Samoa	0.411	0.412	Uganda	0.289	0.292
Malta	0.357	0.357	Azerbaijan	0.409	0.41	Angola	0.274	0.276
Barbados	0.359	0.359	Paraguay	0.408	0.408	Timor-Leste	0.284	0.287
Hungary	0.372	0.371	Maldives	0.406	0.405	Togo	0.281	0.284
Poland	0.371	0.371	Algeria	0.405	0.405	Gambia	0.276	0.279
Chile	0.371	0.371	El Salvador	0.404	0.404	Benin	0.267	0.27
Slovakia	0.374	0.374	Philippines	0.406	0.406	Malawi	0.27	0.272
Estonia	0.377	0.377	Fiji	0.401	0.401	Zambia	0.266	0.268
Lithuania	0.378	0.378	Sri Lanka	0.402	0.402	Eritrea	0.262	0.265
Latvia	0.381	0.381	Syria	0.399	0.4	Rwanda	0.257	0.259
Croatia	0.38	0.38	Occupied Palestinian Territories	0.398	0.398	Cote d'Ivoire	0.251	0.254
Argentina	0.382	0.382	Gabon	0.376	0.377	Guinea	0.249	0.252
Uruguay	0.383	0.382	Turkmenistan	0.394	0.395	Mali	0.234	0.237
Cuba	0.384	0.384	Indonesia	0.393	0.394	Ethiopia	0.234	0.237
Bahamas	0.389	0.389	Guyana	0.39	0.391	Chad	0.232	0.235
Costa Rica	0.39	0.39	Bolivia	0.392	0.393	Guinea Bissau	0.231	0.235
Mexico	0.394	0.394	Mongolia	0.393	0.392	Burundi	0.233	0.236
Libyan Arab Jamahiriya	0.396	0.396	Moldova	0.393	0.394	Burkina Faso	0.225	0.229
Oman	0.396	0.396	Viet Nam	0.393	0.393	Niger	0.226	0.23
Seychelles	0.401	0.402	Equatorial Guinea	0.339	0.341	Mozambique	0.221	0.224
Saudi Arabia	0.399	0.399	Egypt	0.387	0.388	Liberia	0.224	0.227
Bulgaria	0.402	0.402	Honduras	0.385	0.385	Congo Democratic	0.223	0.226
Trinidad and Tobago	0.401	0.401	Cape Verde	0.379	0.38	Central African	0.213	0.217
Panama	0.402	0.402	Uzbekistan	0.379	0.379	Sierra Leone	0.202	0.205
Antigua and Barbuda	0.404	0.404	Nicaragua	0.376	0.376	ARTIF	-	0.17
Saint Kitts and Nevis	0.404	0.404						

For the first experiment, the similarities of countries did not follow a normal distribution, with a p-value of 5.589e-9. We calculated the upper bound ($Q3 + [1.5 * IQR]$) as 0.5415 and the lower bound ($Q1 - [1.5 * IQR]$) as 0.1695. No country similarities lay outside this interval, indicating an absence of outliers.

In the second experiment, we added the artificial outlier ‘*ARTIF*’ to the sample and retested for normality. The new sample did not follow a normal distribution, with a p-value of 6.882e-9. After calculating the upper and lower bounds, we found that *ARTIF*, with a similarity of 0.170, fell outside the interval [0.171, 0.540], indicating it as an outlier.

6. Discussion

In the initial phase of the experiment, we evaluated the results by defining accuracy as the proportion of times the LOF algorithm correctly identifies the first outlier (or the first and second outliers, depending on the scenario) relative to the total number of variations in parameter k, which remains 7 in all scenarios. The results of this experiment demonstrated that incorporating the degree of importance in preference relations led to optimal performance, with a substantially higher accuracy rate in all three scenarios, as illustrated in Figure. 2. This clearly indicates that when the importance degree is considered, the accuracy significantly improves across all tested scenarios, reaching 100% accuracy, in contrast to the lower rates observed when the importance degree is neglected.

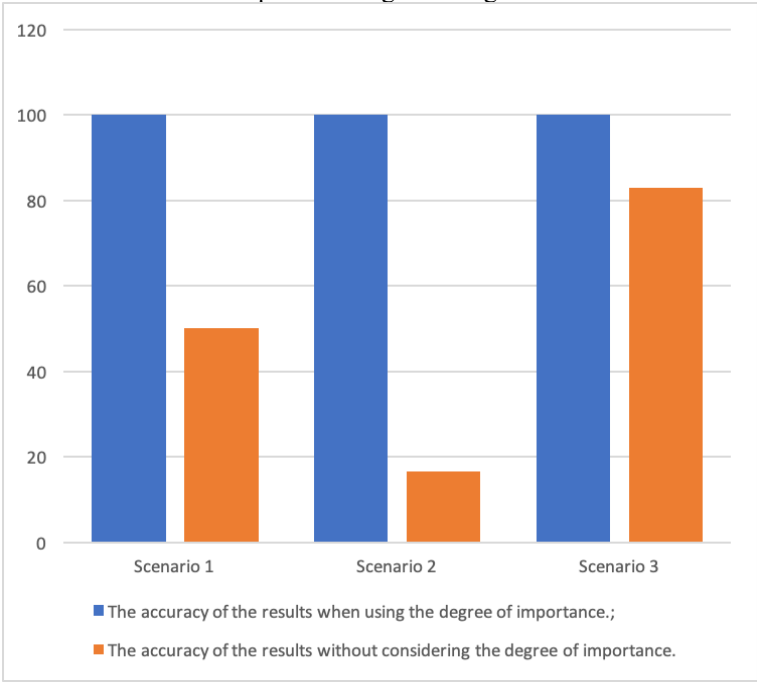


Figure 2. Comparing the accuracy of results in all scenarios

Notably, the scenarios showed a significant drop in accuracy when the degree of importance was not considered, with Scenario 1 falling to 42%, Scenario 2 to 14%, and Scenario 3 to 85%

Additionally, when alternatives are represented in three-dimensional space using their distribution matrix (with the first coordinate omitted due to all alternatives sharing the same value), the impact of incorporating the degree of importance becomes more apparent. As illustrated in Figure. 3, the spatial positions of the alternatives show that including the importance degree causes outliers to be distinctly separated from the main clusters, making them easier to detect. In contrast, when the importance degree is omitted, the alternatives are positioned more closely together, making outlier identification more difficult. This visual evidence supports the notion that accounting for the degree of importance significantly enhances the clarity and effectiveness of outlier detection in multicriteria decision-making contexts.

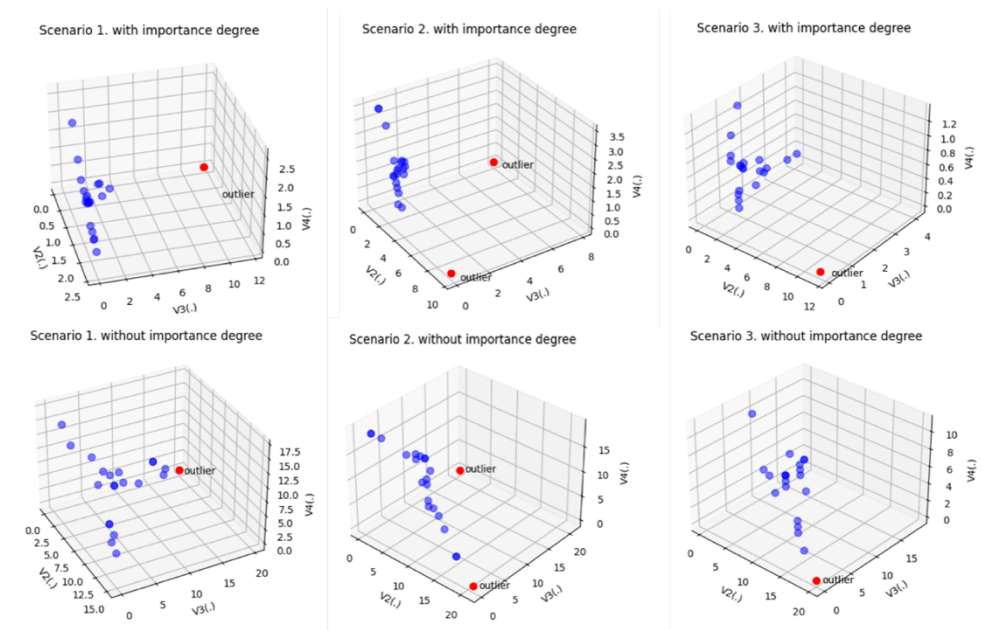


Figure 3. Representing alternatives using their distribution matrix in all scenarios.

In the second phase of the experiment, the results from both scenarios further validated the effectiveness of the proposed approach for detecting outliers within the MCDA framework. In the first scenario, the method was tested on a dataset without any outliers to ensure that no false positives were detected. The dataset, which was consistent and homogeneous in terms of country similarity measures, was analyzed. The findings confirmed that the data did not follow a normal distribution, and no outliers were identified within the defined bounds.

The second scenario introduced an artificial outlier, 'ARTIF', to directly assess the robustness of the method. The successful identification of ARTIF as an outlier, once it fell

outside the calculated bounds, highlights the method's strong ability to detect anomalies. As illustrated in Fig. 4, the outliers are clearly distinguished, reinforcing the method's efficiency in separating anomalies from normal data points.

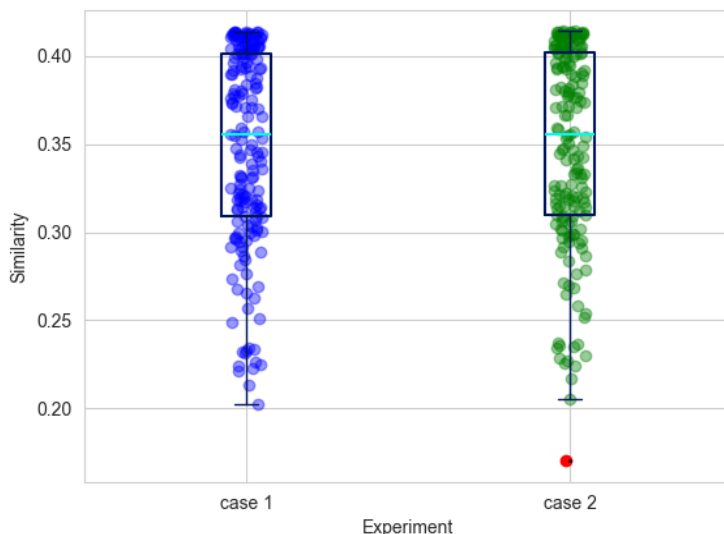


Figure 4. Data representation of the first and second experimentations, with outliers marked in red

7. Conclusion

Identifying outliers within the MCDA domain, which possesses specific characteristics, represents a new avenue of research that has not received adequate attention in the existing literature. In this paper, we address this issue by presenting a new definition of a multicriteria outlier and proposing an approach to detect outliers based on this definition. Our approach relies on two core concepts:

- (1) The importance degree: This concept indicates the significance level of a preference relation between two alternatives. Each alternative is characterized by its multicriteria preference indices as coordinates, and the importance of the relation is determined based on the Euclidean distance between them.
- (2) Using the importance degree, we introduced the similarity concept, which measures how similar each alternative is to the others.

The approach identifies outliers by applying either the IQR or SD method, depending on the distribution of the alternatives' similarities.

The effectiveness of the approach was tested in two phases and experimented using two real-life problems. In the first phase, we focused solely on evaluating the effectiveness of the degree of importance in the world happiness ranking problem. The results of our approach were compared to the approach proposed in [22], both with and without integrating the degree of importance. The results showed 100% accuracy when the degree of importance was used in all scenarios. In the second phase, we tested the overall effectiveness of the approach using

the human development index problem. We evaluated its performance in both outlier-present and outlier-absent scenarios, and it performed well in both situations.

A key advantage of our approach lies in its consideration of the multicriteria nature, ensuring accurate identification of outliers in this context. Furthermore, we have introduced a novel concept in the MCDA domain: the degree of importance. This concept provides an opportunity for various approaches to enhance their effectiveness through integration. However, our approach is based solely on preference relations generated by the PROMETHEE method.

For future work, we aim to extend the application of our approach beyond PROMETHEE to encompass other outranking methods. This includes integrating the concept of the degree of importance into methods such as ELECTRE. Additionally, it is crucial to apply our proposed approach to big data.

References

- [1] Barnett V, Lewis T. Outliers in statistical data John Wiley and Sons. New York 1994.
- [2] Bland J. DG.(1996)‘Statistics notes: measurement error.’ *BMJ* n.d.;312:1654.
- [3] Bolton RJ, Hand DJ. Unsupervised profiling methods for fraud detection. *Credit Scoring Credit Control VII* 2001:235–55.
- [4] Brans J-P, Mareschal B. The PROMCALC & GAIA decision support system for multicriteria decision aid. *Decis Support Syst* 1994;12:297–310. [https://doi.org/https://doi.org/10.1016/0167-9236\(94\)90048-5](https://doi.org/https://doi.org/10.1016/0167-9236(94)90048-5).
- [5] Brans J-P, Vincke P, Mareschal B. How to select and how to rank projects: The PROMETHEE method. *Eur J Oper Res* 1986;24:228–38.
- [6] Breunig MM, Kriegel H-P, Ng RT, Sander J. LOF: identifying density-based local outliers. *Proc. 2000 ACM SIGMOD Int. Conf. Manag. data, 2000*, p. 93–104.
- [7] Dang TT, Ngan HYT, Liu W. Distance-based k-nearest neighbors outlier detection method in large-scale traffic data. *2015 IEEE Int. Conf. Digit. Signal Process., IEEE;* 2015, p. 507–10.
- [8] Dejaegere G, De Smet Y. Promethee γ : A new Promethee based method for partial ranking based on valued coalitions of monocriterion net flow scores. *J Multi-Criteria Decis Anal* 2023;30:147–60.
- [9] Dekking FM, Kraaikamp C, Lopuhaä HP, Meester LE. A Modern Introduction to Probability and Statistics: Understanding why and how. vol. 488. Springer; 2005.
- [10] De Smet Y, Guzmán LM. Towards multicriteria clustering: An extension of the k-means algorithm. *Eur J Oper Res* 2004;158:390–8. <https://doi.org/10.1016/j.ejor.2003.06.012>.
- [11] De Smet Y, Hubinont J-PP, Rosenfeld J. A note on the detection of outliers in a binary outranking relation. *Lect Notes Comput Sci (Including Subser Lect Notes Artif Intell Lect Notes Bioinformatics)* 2017;10173 LNCS:151–9.

- https://doi.org/10.1007/978-3-319-54157-0_11.
- [12] De Smet Y, Nemery P, Selvaraj R. An exact algorithm for the multicriteria ordered clustering problem. *Omega* 2012;40:861–9.
 - [13] Ding Q, Katenka N, Barford P, Kolaczyk E, Crovella M. Intrusion as (anti) social communication: characterization and detection. *Proc. 18th ACM SIGKDD Int. Conf. Knowl. Discov. data Min.*, 2012, p. 886–94.
 - [14] Ishizaka A, Nemery P. Multi-attribute utility theory. *Multi-Criteria Decis Anal Methods Softw* 2013:81–113.
 - [15] Korhonen P. Interactive methods. *Mult Criteria Decis Anal State Art Surv* 2005:641–61.
 - [16] MacQueen J. Some methods for classification and analysis of multivariate observations. *Proc. fifth Berkeley Symp. Math. Stat. Probab.*, vol. 1, Oakland, CA, USA; 1967, p. 281–97.
 - [17] Mareschal B, Brans J-P. Geometrical representations for MCDA. *Eur J Oper Res* 1988;34:69–77. [https://doi.org/https://doi.org/10.1016/0377-2217\(88\)90456-0](https://doi.org/https://doi.org/10.1016/0377-2217(88)90456-0).
 - [18] Moonesinghe HDK, Tan P-N. Outrank: a graph-based outlier detection framework using random walk. *Int J Artif Intell Tools* 2008;17:19–36.
 - [19] Nguyen HV, Ang HH, Gopalkrishnan V. Mining outliers with ensemble of heterogeneous detectors on random subspaces. *Database Syst. Adv. Appl. 15th Int. Conf. DASFAA 2010*, Tsukuba, Japan, April 1-4, 2010, *Proceedings, Part I* 15, Springer; 2010, p. 368–83.
 - [20] Pimentel T, Monteiro M, Veloso A, Ziviani N. Deep active learning for anomaly detection. *2020 Int. Jt. Conf. Neural Networks, IEEE*; 2020, p. 1–8.
 - [21] Rouba B. A net-flow based approach to detect outliers in multicriteria decision problems. *Intell Decis Technol* 2021;15:239–50. <https://doi.org/10.3233/IDT-200046>.
 - [22] Rouba B, Nait-Bahloul S. Towards identifying multicriteria outliers: An outranking relation-based approach. *Int J Decis Support Syst Technol* 2018;10:27–38. <https://doi.org/10.4018/IJDSST.2018070102>.
 - [23] Roy B. *Multicriteria methodology for decision aiding*. vol. 12. Springer Science & Business Media; 2013.
 - [24] Roy B. *Multicriteria methodology for decision aiding*. vol. 12. Springer Science & Business Media; 1996.
 - [25] Shapiro SS, Wilk MB. An analysis of variance test for normality (complete samples). *Biometrika* 1965;52:591–611.
 - [26] Tsou Y-L, Chu H-M, Li C, Yang S-W. Robust distributed anomaly detection using optimal weighted one-class random forests. *2018 IEEE Int. Conf. Data Min., IEEE*; 2018, p. 1272–7.

- [27] Vincke P. Multicriteria decision-aid. John Wiley & Sons; 1992.
- [28] World Happiness Report 2015.
<https://www.kaggle.com/datasets/mathurinache/world-happiness-report?resource=download&select=2015.csv>.
- [29] Yang X, Latecki LJ, Pokrajac D. Outlier detection with globally optimal exemplar-based GMM. Proc. 2009 SIAM Int. Conf. data Min., SIAM; 2009, p. 145–54.

Received 20.10.2024, Accepted 10.03.2025

Appendix A. Ranking of 158 alternatives

Rank	Alternative	Φ	Φ^+	Φ^-	Rank	Alternative	Φ	Φ^+	Φ^-
1	New Zealand	0.7116	0.7208	0.0092	80	China	-0.072	0.2329	0.3049
2	Australia	0.7107	0.7219	0.0112	81	Myanmar	-0.0735	0.3367	0.4102
3	Canada	0.7003	0.7139	0.0135	82	Kyrgyzstan	-0.0743	0.2264	0.3007
4	Ireland	0.693	0.7119	0.0189	83	Bolivia	-0.0899	0.2175	0.3075
5	Norway	0.682	0.6998	0.0178	84	Lebanon	-0.0954	0.2114	0.3069
6	Netherlands	0.6816	0.6973	0.0157	85	Turkey	-0.097	0.2264	0.3235
7	Sweden	0.679	0.6922	0.0132	86	Azerbaijan	-0.0992	0.2264	0.3256
8	United Kingdom	0.6784	0.695	0.0165	87	Macedonia	-0.0994	0.1919	0.2913
9	Denmark	0.6719	0.6879	0.016	88	Russia	-0.1	0.2418	0.3418
10	Switzerland	0.6682	0.6886	0.0205	89	Peru	-0.1056	0.1867	0.2923
11	Luxembourg	0.6158	0.6429	0.0271	90	Algeria	-0.1083	0.2336	0.3419
12	Hong Kong	0.6148	0.6587	0.0439	91	Syria	-0.1124	0.3139	0.4263
13	Iceland	0.6029	0.6406	0.0377	92	Bulgaria	-0.1218	0.2039	0.3257
14	Qatar	0.5816	0.6361	0.0545	93	El Salvador	-0.1231	0.1913	0.3144
15	Finland	0.5711	0.6075	0.0363	94	Iran	-0.1285	0.2389	0.3674
16	Singapore	0.5615	0.6258	0.0644	95	Botswana	-0.1295	0.2087	0.3382
17	Austria	0.5448	0.5873	0.0425	96	Georgia	-0.1322	0.2777	0.41
18	Malta	0.544	0.5878	0.0438	97	Lithuania	-0.143	0.2405	0.3835
19	Germany	0.531	0.5767	0.0457	98	Montenegro	-0.1453	0.2	0.3452
20	United States	0.5108	0.5587	0.0479	99	Honduras	-0.1612	0.1679	0.3291
21	United Arab Emirates	0.5097	0.5657	0.056	100	Albania	-0.1631	0.1775	0.3405
22	Belgium	0.4775	0.5367	0.0592	101	Tajikistan	-0.1661	0.1951	0.3612
23	Oman	0.4103	0.5022	0.0919	102	Romania	-0.167	0.1834	0.3504
24	France	0.3961	0.498	0.1019	103	Ukraine	-0.1706	0.1855	0.3562

25	Uruguay	0.3915	0.476	0.0845	104	Serbia	-0.1887	0.1729	0.3615
26	Kuwait	0.3704	0.484	0.1136	105	Kenya	-0.1951	0.2007	0.3958
27	Japan	0.3422	0.4676	0.1254	106	Uganda	-0.1962	0.1791	0.3753
28	Bahrain	0.3262	0.4457	0.1195	107	South Africa	-0.2055	0.1666	0.3721
29	Slovenia	0.3076	0.4221	0.1145	108	Djibouti	-0.2323	0.2213	0.4537
30	Costa Rica	0.3016	0.4083	0.1067	109	Bosnia and Herzegovina	-0.2401	0.1801	0.4202
31	Thailand	0.3012	0.4533	0.1522	110	Greece	-0.2432	0.2176	0.4608
32	Israel	0.2891	0.4096	0.1205	111	Nepal	-0.2457	0.1661	0.4118
33	North Cyprus	0.2861	0.3971	0.1109	112	Zambia	-0.2462	0.1514	0.3976
34	Uzbekistan	0.2679	0.4645	0.1966	113	Kosovo	-0.2502	0.1616	0.4118
35	Chile	0.2678	0.3867	0.1189	114	Haiti	-0.2525	0.243	0.4955
36	Spain	0.2478	0.3872	0.1394	115	Iraq	-0.2548	0.1737	0.4285
37	Panama	0.222	0.343	0.121	116	Croatia	-0.2582	0.172	0.4302
38	Malaysia	0.2128	0.3514	0.1386	117	Moldova	-0.2633	0.1422	0.4055
39	Bhutan	0.1972	0.4024	0.2053	118	Gabon	-0.2665	0.1552	0.4217
40	Nicaragua	0.1823	0.3653	0.183	119	Palestinian Territories	-0.2786	0.1455	0.4241
41	Saudi Arabia	0.1724	0.3865	0.2142	120	Ivory Coast	-0.2788	0.1702	0.449
42	Paraguay	0.1705	0.3581	0.1876	121	Tanzania	-0.2866	0.1578	0.4443
43	Turkmenistan	0.1679	0.376	0.2081	122	Senegal	-0.3	0.1216	0.4216
44	Trinidad and Tobago	0.1668	0.3613	0.1944	123	Bangladesh	-0.3155	0.1355	0.451
45	Taiwan	0.1651	0.3274	0.1622	124	India	-0.3171	0.1412	0.4584
46	Mauritius	0.1632	0.3442	0.181	125	Mali	-0.3182	0.1183	0.4365
47	Poland	0.1627	0.3276	0.1649	126	Morocco	-0.3204	0.1348	0.4552
48	Dominican Republic	0.1597	0.3276	0.168	127	Niger	-0.3211	0.1471	0.4683
49	Brazil	0.1586	0.3356	0.177	128	Ethiopia	-0.3228	0.152	0.4749
50	Estonia	0.1534	0.3368	0.1834	129	Cameroon	-0.3279	0.115	0.4429

51	Sri Lanka	0.1498	0.3429	0.1931	130	Mozambique	-0.333	0.1365	0.4695
52	Mexico	0.1226	0.3325	0.2099	131	Mauritania	-0.3337	0.1468	0.4805
53	Cyprus	0.1211	0.3282	0.2071	132	Tunisia	-0.3429	0.1309	0.4738
54	Portugal	0.1192	0.3191	0.1999	133	Burkina Faso	-0.349	0.113	0.462
55	Belarus	0.1145	0.3273	0.2128	134	Sudan	-0.3494	0.1352	0.4846
56	Italy	0.1055	0.333	0.2275	135	Nigeria	-0.3501	0.1187	0.4688
57	Czech Republic	0.1	0.3097	0.2097	136	Swaziland	-0.3613	0.1078	0.4691
58	Laos	0.0975	0.4101	0.3126	137	Congo (Brazzaville)	-0.3663	0.1157	0.4821
59	Argentina	0.0941	0.2889	0.1949	138	Armenia	-0.3699	0.1213	0.4912
60	Colombia	0.0767	0.2808	0.2041	139	Egypt	-0.3702	0.1158	0.486
61	Libya	0.064	0.2592	0.1952	140	Lesotho	-0.3703	0.1032	0.4735
62	Indonesia	0.0629	0.3385	0.2756	141	Sierra Leone	-0.3725	0.0959	0.4684
63	Ecuador	0.0621	0.2952	0.2331	142	Ghana	-0.3737	0.1102	0.484
64	South Korea	0.0583	0.2915	0.2333	143	Pakistan	-0.3797	0.1577	0.5374
65	Somaliland region	0.0563	0.3936	0.3374	144	Comoros	-0.3966	0.1378	0.5343
66	Suriname	0.0543	0.2843	0.23	145	Zimbabwe	-0.4024	0.0838	0.4862
67	Mongolia	0.0479	0.3086	0.2607	146	Guinea	-0.42	0.112	0.532
68	Guatemala	0.0472	0.2756	0.2284	147	Malawi	-0.4238	0.1245	0.5482
69	Kazakhstan	0.0423	0.2617	0.2194	148	Benin	-0.4347	0.0871	0.5217
70	Slovakia	0.0412	0.2777	0.2365	149	Afghanistan	-0.444	0.1289	0.5729
71	Venezuela	0.0403	0.2794	0.2391	150	Yemen	-0.4635	0.0786	0.5421
72	Philippines	0.0385	0.2884	0.2499	151	Congo (Kinshasa)	-0.4674	0.0783	0.5457
73	Jamaica	-0.0088	0.2427	0.2515	152	Central African Republic	-0.4742	0.0891	0.5633
74	Vietnam	-0.0306	0.2402	0.2707	153	Liberia	-0.4806	0.0699	0.5504
75	Hungary	-0.0311	0.2416	0.2727	154	Angola	-0.4956	0.0702	0.5658
76	Jordan	-0.0377	0.2326	0.2704	155	Madagascar	-0.5083	0.0671	0.5754

77	Latvia	-0.0503	0.2138	0.2642	156	Togo	-0.5191	0.0581	0.577 2
78	Rwanda	-0.0579	0.3203	0.3783	157	Chad	-0.5581	0.0383	0.596 4
79	Cambodia	-0.0604	0.305	0.3654	158	Burundi	-0.6278	0.0366	0.664 4

To simplify, we present the alternatives by their rank in the following appendices.

Appendix B. The weighted preference matrix of the top 20 alternatives and the lowest-ranked alternative.

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	158
1	1	0.072 ⁺	0.060 ⁺	0.088 ⁺	0.089 ⁺	0.094 ⁺	0.079 ⁺	0.094 ⁺	0.083 ⁺	0.084 ⁺	0.116 ⁺	0.123 ⁺	0.112 ⁺	0.154 ⁺	0.116 ⁺	0.177 ⁺	0.180 ⁺	0.184 ⁺	0.186 ⁺	0.237 ⁺	0.654 ⁺
2	0.072 ⁻	1	0.061 ⁺	0.076 ⁺	0.077 ⁺	0.080 ⁺	0.065 ⁺	0.096 ⁺	0.085 ⁺	0.073 ⁺	0.102 ⁺	0.110 ⁺	0.101 ⁺	0.143 ⁺	0.099 ⁺	0.166 ⁺	0.163 ⁺	0.177 ⁺	0.177 ⁺	0.222 ⁺	0.639 ⁺
3	0.060 ⁻	0.061 ⁻	1	0.062 ⁺	0.105 ⁺	0.070 ⁺	0.080 ⁻	0.085 ⁺	0.096 ⁻	0.091 ⁻	0.100 ⁺	0.116 ⁺	0.098 ⁺	0.155 ⁺	0.094 ⁺	0.176 ⁺	0.154 ⁺	0.167 ⁺	0.165 ⁺	0.211 ⁺	0.627 ⁺
4	0.088 ⁻	0.076 ⁻	0.062 ⁺	1	0.101 ⁻	0.070 ⁺	0.092 ⁺	0.087 ⁺	0.101 ⁻	0.099 ⁻	0.093 ⁺	0.096 ⁺	0.088 ⁺	0.141 ⁺	0.086 ⁺	0.157 ⁺	0.146 ⁺	0.148 ⁺	0.159 ⁺	0.189 ⁺	0.605 ⁺
5	0.089 ⁺	0.077 ⁺	0.105 ⁺	0.101 ⁻	1	0.099 ⁺	0.075 ⁺	0.108 ⁺	0.091 ⁺	0.082 ⁻	0.103 ⁺	0.100 ⁺	0.124 ⁺	0.107 ⁺	0.144 ⁺	0.168 ⁺	0.168 ⁺	0.164 ⁺	0.176 ⁺	0.216 ⁺	0.625 ⁺
6	0.094 ⁻	0.080 ⁻	0.070 ⁻	0.070 ⁺	0.109 ⁻	1	0.094 ⁻	0.072 ⁺	0.105 ⁻	0.112 ⁻	0.099 ⁺	0.113 ⁺	0.082 ⁺	0.154 ⁺	0.084 ⁺	0.176 ⁺	0.128 ⁺	0.148 ⁺	0.138 ⁺	0.189 ⁺	0.609 ⁺
7	0.079 ⁻	0.065 ⁻	0.080 ⁺	0.092 ⁺	0.075 ⁺	0.094 ⁺	1	0.113 ⁺	0.077 ⁺	0.080 ⁺	0.098 ⁺	0.115 ⁺	0.100 ⁺	0.144 ⁺	0.099 ⁺	0.166 ⁺	0.152 ⁺	0.165 ⁺	0.164 ⁺	0.214 ⁺	0.632 ⁺
8	0.094 ⁻	0.096 ⁻	0.085 ⁺	0.087 ⁺	0.108 ⁻	0.072 ⁺	0.113 ⁺	1	0.119 ⁺	0.121 ⁻	0.097 ⁺	0.106 ⁺	0.084 ⁺	0.146 ⁺	0.100 ⁺	0.163 ⁺	0.149 ⁺	0.148 ⁺	0.146 ⁺	0.193 ⁺	0.600 ⁺
9	0.083 ⁺	0.085 ⁺	0.096 ⁺	0.099 ⁺	0.091 ⁺	0.105 ⁺	0.077 ⁺	0.119 ⁺	1	0.063 ⁻	0.106 ⁺	0.105 ⁺	0.107 ⁺	0.114 ⁺	0.102 ⁺	0.145 ⁺	0.164 ⁺	0.167 ⁺	0.177 ⁺	0.215 ⁺	0.631 ⁺
10	0.084 ⁺	0.073 ⁺	0.091 ⁺	0.093 ⁺	0.082 ⁺	0.112 ⁺	0.080 ⁺	0.121 ⁺	0.063 ⁻	1	0.106 ⁺	0.108 ⁺	0.116 ⁺	0.121 ⁺	0.111 ⁺	0.147 ⁺	0.180 ⁺	0.181 ⁺	0.191 ⁺	0.225 ⁺	0.638 ⁺
11	0.116 ⁻	0.105 ⁻	0.100 ⁺	0.093 ⁺	0.103 ⁻	0.099 ⁺	0.098 ⁺	0.097 ⁻	0.096 ⁻	0.106 ⁻	1	0.054 ⁺	0.084 ⁻	0.100 ⁺	0.076 ⁺	0.121 ⁺	0.125 ⁺	0.120 ⁺	0.127 ⁺	0.165 ⁺	0.574 ⁺
12	0.129 ⁺	0.110 ⁺	0.116 ⁺	0.096 ⁺	0.100 ⁻	0.113 ⁺	0.115 ⁺	0.106 ⁺	0.105 ⁻	0.108 ⁻	0.054 ⁺	1	0.089 ⁺	0.084 ⁺	0.091 ⁺	0.106 ⁺	0.136 ⁺	0.119 ⁺	0.138 ⁺	0.160 ⁺	0.567 ⁺
13	0.112 ⁻	0.101 ⁺	0.098 ⁺	0.088 ⁺	0.100 ⁻	0.082 ⁺	0.100 ⁺	0.084 ⁺	0.107 ⁻	0.116 ⁻	0.084 ⁺	0.089 ⁺	1	0.114 ⁺	0.081 ⁺	0.124 ⁺	0.121 ⁺	0.119 ⁺	0.128 ⁺	0.161 ⁺	0.575 ⁺
14	0.154 ⁺	0.143 ⁺	0.155 ⁺	0.141 ⁺	0.124 ⁺	0.154 ⁺	0.144 ⁺	0.146 ⁺	0.114 ⁺	0.121 ⁺	0.100 ⁺	0.084 ⁺	0.114 ⁺	1	0.121 ⁺	0.072 ⁺	0.168 ⁺	0.152 ⁺	0.177 ⁺	0.182 ⁺	0.564 ⁺
15	0.116 ⁻	0.099 ⁻	0.094 ⁻	0.086 ⁻	0.107 ⁻	0.084 ⁻	0.099 ⁻	0.100 ⁻	0.102 ⁻	0.111 ⁻	0.076 ⁻	0.091 ⁺	0.081 ⁺	0.121 ⁺	1	0.147 ⁺	0.105 ⁺	0.121 ⁺	0.121 ⁺	0.160 ⁺	0.574 ⁺
16	0.177 ⁺	0.166 ⁺	0.176 ⁺	0.157 ⁺	0.144 ⁺	0.176 ⁺	0.166 ⁺	0.163 ⁺	0.145 ⁺	0.147 ⁺	0.121 ⁺	0.106 ⁺	0.124 ⁺	0.072 ⁺	0.147 ⁺	1	0.189 ⁺	0.165 ⁺	0.197 ⁺	0.190 ⁺	0.550 ⁺
17	0.180 ⁻	0.163 ⁻	0.154 ⁻	0.146 ⁻	0.168 ⁻	0.128 ⁻	0.152 ⁻	0.149 ⁻	0.164 ⁻	0.180 ⁻	0.125 ⁻	0.136 ⁻	0.121 ⁻	0.168 ⁻	0.105 ⁻	0.189 ⁺	1	0.096 ⁺	0.066 ⁺	0.102 ⁺	0.530 ⁺
18	0.184 ⁻	0.177 ⁻	0.167 ⁻	0.148 ⁻	0.164 ⁻	0.148 ⁻	0.165 ⁻	0.148 ⁻	0.167 ⁻	0.181 ⁻	0.120 ⁻	0.119 ⁻	0.119 ⁻	0.152 ⁻	0.121 ⁻	0.165 ⁻	0.096 ⁺	1	0.087 ⁺	0.010 ⁺	0.517 ⁺
19	0.188 ⁻	0.177 ⁻	0.165 ⁻	0.159 ⁻	0.176 ⁻	0.138 ⁻	0.164 ⁻	0.148 ⁻	0.177 ⁻	0.191 ⁻	0.127 ⁻	0.136 ⁻	0.128 ⁻	0.177 ⁻	0.121 ⁻	0.197 ⁺	0.066 ⁺	0.087 ⁺	1	0.106 ⁺	0.523 ⁺
20	0.237 ⁻	0.222 ⁻	0.211 ⁻	0.189 ⁻	0.216 ⁻	0.189 ⁻	0.214 ⁻	0.193 ⁻	0.215 ⁻	0.225 ⁻	0.169 ⁻	0.160 ⁻	0.161 ⁻	0.182 ⁻	0.160 ⁻	0.190 ⁻	0.160 ⁻	0.199 ⁻	0.106 ⁺	1	0.462 ⁺
158	0.654 ⁻	0.639 ⁻	0.627 ⁻	0.605 ⁻	0.623 ⁻	0.609 ⁻	0.632 ⁻	0.600 ⁻	0.631 ⁻	0.638 ⁻	0.574 ⁻	0.567 ⁻	0.575 ⁻	0.564 ⁻	0.574 ⁻	0.550 ⁻	0.530 ⁻	0.517 ⁻	0.523 ⁻	0.462 ⁻	1

Appendix C. The weighted preference matrix of the 20 alternatives ranked in the middle, along with both the first and last ranked alternatives.

	1	70	71	72	73	74	75	76	77	78	79	80	81	82	83	84	85	86	87	88	89	158
1	1	0.384 ⁺	0.382 ⁺	0.402 ⁺	0.429 ⁺	0.433 ⁺	0.448 ⁺	0.442 ⁺	0.439 ⁺	0.453 ⁺	0.448 ⁺	0.449 ⁺	0.446 ⁺	0.452 ⁺	0.465 ⁺	0.458 ⁺	0.445 ⁺	0.467 ⁺	0.479 ⁺	0.470 ⁺	0.482 ⁺	0.712 ⁺
70	0.384 ⁻	1	0.097 ⁻	0.168 ⁻	0.119 ⁺	0.160 ⁺	0.134 ⁺	0.143 ⁺	0.126 ⁺	0.220 ⁻	0.204 ⁺	0.140 ⁺	0.218 ⁻	0.165 ⁺	0.167 ⁺	0.174 ⁺	0.168 ⁺	0.172 ⁺	0.199 ⁺	0.198 ⁺	0.179 ⁺	0.421 ⁺
71	0.382 ⁻	0.097 ⁻	1	0.124 ⁺	0.099 ⁺	0.119 ⁺	0.120 ⁺	0.099 ⁺	0.100 ⁺	0.216 ⁻	0.179 ⁺	0.108 ⁺	0.212 ⁻	0.138 ⁺	0.148 ⁺	0.142 ⁺	0.131 ⁺	0.142 ⁺	0.151 ⁺	0.163 ⁺	0.146 ⁺	0.426 ⁺
72	0.402 ⁻	0.168 ⁻	0.124 ⁺	1	0.107 ⁺	0.078 ⁺	0.164 ⁺	0.124 ⁻	0.131 ⁻	0.172 ⁻	0.145 ⁻	0.118 ⁺	0.173 ⁻	0.114 ⁺	0.128 ⁺	0.133 ⁺	0.153 ⁺	0.139 ⁺	0.159 ⁺	0.159 ⁺	0.161 ⁺	0.397 ⁺
73	0.429 ⁻	0.119 ⁻	0.099 ⁻	0.107 ⁺	1	0.104 ⁻	0.103 ⁺	0.100 ⁻	0.091 ⁻	0.197 ⁻	0.169 ⁻	0.101 ⁺	0.193 ⁻	0.114 ⁺	0.130 ⁺	0.140 ⁺	0.145 ⁺	0.135 ⁺	0.155 ⁺	0.158 ⁺	0.148 ⁺	0.383 ⁺
74	0.435 ⁻	0.168 ⁻	0.119 ⁻	0.078 ⁺	0.104 ⁻	1	0.159 ⁻	0.112 ⁻	0.129 ⁻	0.165 ⁻	0.145 ⁻	0.114 ⁺	0.170 ⁻	0.108 ⁺	0.100 ⁺	0.130 ⁺	0.135 ⁺	0.122 ⁺	0.138 ⁺	0.158 ⁺	0.146 ⁺	0.378 ⁺
75	0.448 ⁻	0.134 ⁻	0.120 ⁻	0.164 ⁻	0.103 ⁺	0.159 ⁻	1	0.099 ⁺	0.073 ⁻	0.234 ⁻	0.210 ⁻	0.103 ⁻	0.219 ⁻	0.145 ⁻	0.166 ⁺	0.147 ⁺	0.145 ⁻	0.146 ⁺	0.153 ⁺	0.147 ⁻	0.153 ⁺	0.377 ⁺
76	0.442 ⁻	0.143 ⁻	0.099 ⁻	0.124 ⁻	0.100 ⁻	0.112 ⁻	0.099 ⁺	1	0.084 ⁻	0.211 ⁻	0.190 ⁻	0.090 ⁺	0.203 ⁻	0.113 ⁻	0.123 ⁺	0.120 ⁺	0.118 ⁻	0.112 ⁺	0.122 ⁺	0.138 ⁺	0.123 ⁺	0.390 ⁺
77	0.439 ⁻	0.126 ⁻	0.100 ⁻	0.131 ⁻	0.091 ⁻	0.129 ⁻	0.073 ⁺	0.084 ⁺	1	0.223 ⁻	0.201 ⁻	0.079 ⁺	0.215 ⁻	0.132 ⁻	0.151 ⁺	0.131 ⁺	0.132 ⁻	0.130 ⁺	0.138 ⁺	0.144 ⁺	0.136 ⁺	0.389 ⁺
78	0.453 ⁻	0.220 ⁻	0.216 ⁻	0.172 ⁻	0.197 ⁻	0.165 ⁻	0.234 ⁻	0.211 ⁻	0.223 ⁻	1	0.084 ⁻	0.190 ⁻	0.081 ⁻	0.153 ⁻	0.155 ⁻	0.178 ⁻	0.190 ⁻	0.175 ⁻	0.211 ⁻	0.189 ⁻	0.214 ⁻	0.300 ⁺
79	0.448 ⁻	0.204 ⁻	0.197 ⁻	0.145 ⁻	0.168 ⁻	0.145 ⁻	0.210 ⁻	0.190 ⁻	0.201 ⁻	0.084 ⁻	1	0.168 ⁻	0.092 ⁻	0.134 ⁻	0.137 ⁻	0.172 ⁻	0.181 ⁻	0.167 ⁻	0.200 ⁻	0.182 ⁻	0.199 ⁻	0.313 ⁺
80	0.449 ⁻	0.140 ⁻	0.108 ⁻	0.118 ⁻	0.101 ⁻	0.116 ⁻	0.103 ⁻	0.090 ⁻	0.079 ⁻	0.190 ⁻	0.168 ⁻	1	0.186 ⁻	0.097 ⁻	0.121 ⁺	0.094 ⁺	0.092 ⁻	0.095 ⁺	0.105 ⁻	0.111 ⁺	0.105 ⁻	0.357 ⁺
81	0.446 ⁻	0.218 ⁻	0.212 ⁻	0.173 ⁻	0.193 ⁻	0.170 ⁻	0.219 ⁻	0.203 ⁻	0.215 ⁻	0.081 ⁻	0.092 ⁻	0.186 ⁻	1	0.156 ⁻	0.163 ⁻	0.183 ⁻	0.189 ⁻	0.173 ⁻	0.218 ⁻	0.187 ⁻	0.214 ⁻	0.304 ⁺
82	0.452 ⁻	0.165 ⁻	0.138 ⁻	0.114 ⁻	0.114 ⁻	0.108 ⁻	0.145 ⁻	0.113 ⁻	0.132 ⁻	0.153 ⁻	0.134 ⁻	0.097 ⁻	0.156 ⁻	1	0.099 ⁻	0.095 ⁻	0.121 ⁻	0.111 ⁻	0.109 ⁻	0.108 ⁺	0.116 ⁻	0.338 ⁺
83	0.465 ⁻	0.167 ⁻	0.140 ⁻	0.128 ⁻	0.130 ⁻	0.100 ⁻	0.168 ⁻	0.123 ⁻	0.151 ⁻	0.155 ⁻	0.137 ⁻	0.121 ⁻	0.163 ⁻	0.099 ⁻	1	0.116 ⁻	0.123 ⁻	0.113 ⁻	0.128 ⁻	0.157 ⁻	0.123 ⁻	0.350 ⁺
84	0.458 ⁻	0.174 ⁻	0.142 ⁻	0.135 ⁻	0.140 ⁻	0.130 ⁻	0.147 ⁻	0.120 ⁻	0.131 ⁻	0.178 ⁻	0.172 ⁻	0.094 ⁻	0.185 ⁻	0.095 ⁻	0.116 ⁺	1	0.080 ⁻	0.080 ⁻	0.069 ⁻	0.086 ⁻	0.087 ⁻	0.349 ⁺
85	0.454 ⁻	0.168 ⁻	0.131 ⁻	0.153 ⁻	0.145 ⁻	0.135 ⁻	0.188 ⁻	0.132 ⁻	0.190 ⁻	0.181 ⁻	0.092 ⁻	0.189 ⁻	0.121 ⁺	0.123 ⁺	0.080 ⁺	1	0.070 ⁻	0.079 ⁻	0.104 ⁺	0.084 ⁻	0.084 ⁻	0.360 ⁺
86	0.467 ⁻	0.172 ⁻	0.142 ⁻	0.139 ⁻	0.135 ⁻	0.122 ⁻	0.148 ⁻	0.112 ⁻	0.130 ⁻	0.175 ⁻	0.167 ⁻	0.095 ⁻	0.173 ⁻	0.111 ⁻	0.113 ⁻	0.080 ⁻	0.079 ⁻	1	0.093 ⁻	0.104 ⁻	0.090 ⁻	0.388 ⁺
87	0.479 ⁻	0.189 ⁻	0.151 ⁻	0.155 ⁻	0.155 ⁻	0.138 ⁻	0.153 ⁻	0.122 ⁻	0.138 ⁻	0.211 ⁻	0.200 ⁻	0.105 ⁻	0.218 ⁻	0.109 ⁻	0.128 ⁻	0.069 ⁻	0.079 ⁻	0.093 ⁻	1	0.096 ⁻	0.064 ⁺	0.368 ⁺
88	0.470 ⁻	0.198 ⁻	0.163 ⁻	0.155 ⁻	0.156 ⁻	0.147 ⁻	0.138 ⁻	0.123 ⁻	0.141 ⁻	0.189 ⁻	0.182 ⁻	0.111 ⁻	0.187 ⁻	0.108 ⁻	0.157 ⁻	0.086 ⁻	0.108 ⁻	0.104 ⁻	0.096 ⁻	1	0.058 ⁻	0.333 ⁺
89	0.482 ⁻	0.179 ⁻	0.146 ⁻	0.161 ⁻	0.148 ⁻	0.146 ⁻	0.153 ⁻	0.133 ⁻	0.136 ⁻	0.214 ⁻	0.199 ⁻	0.105 ⁻	0.214 ⁻	0.116 ⁻	0.125 ⁻	0.087 ⁻	0.084 ⁻	0.090 ⁻	0.069 ⁻	0.108 ⁻	1	0.364 ⁺
158	0.712 ⁻	0.421 ⁻	0.426 ⁻	0.397 ⁻	0.383 ⁻	0.378 ⁻	0.377 ⁻	0.390 ⁻	0.389 ⁻	0.300 ⁻	0.313 ⁻	0.357 ⁻	0.304 ⁻	0.338 ⁻	0.350 ⁻	0.349 ⁻	0.360 ⁻	0.348 ⁻	0.368 ⁻	0.333 ⁻	0.364 ⁻	1

Appendix D. The weighted preference matrix of the 20 alternatives with the lowest ranks along with the first ranked alternative.

	1	139	140	141	142	143	144	145	146	147	148	149	150	151	152	153	154	155	156	157	158
1	/	0.485 ⁺	0.542 ⁺	0.569 ⁺	0.499 ⁺	0.510 ⁺	0.521 ⁺	0.553 ⁺	0.567 ⁺	0.594 ⁺	0.572 ⁺	0.587 ⁺	0.533 ⁺	0.618 ⁺	0.651 ⁺	0.578 ⁺	0.583 ⁺	0.583 ⁺	0.625 ⁺	0.622 ⁺	0.701 ⁺
139	0.485 ⁻	/	0.130 ^R	0.139 ⁺	0.095 ^R	0.113 ^R	0.100 ^R	0.150 ^R	0.157 ⁺	0.187 ⁺	0.169 ^R	0.178 ⁺	0.121 ^R	0.173 ⁺	0.229 ⁺	0.148 ^R	0.146 ⁺	0.151 ^R	0.199 ⁺	0.181 ⁺	0.266 ⁺
140	0.542 ⁻	0.130 ^R	/	0.071 ⁺	0.118 ⁺	0.104 ^R	0.103 ⁺	0.089 ^R	0.093 ⁺	0.122 ⁺	0.112 ⁺	0.107 ⁺	0.107 ⁺	0.123 ⁺	0.179 ⁺	0.101 ⁺	0.112 ⁺	0.099 ⁺	0.132 ⁺	0.124 ⁺	0.202 ⁺
141	0.569 ⁻	0.139 ⁺	0.071 ⁺	/	0.124 ⁺	0.124 ⁻	0.109 ⁺	0.084 ⁺	0.090 ^R	0.111 ⁺	0.109 ^R	0.099 ⁺	0.116 ⁺	0.105 ⁺	0.162 ⁺	0.097 ^R	0.112 ⁺	0.092 ^R	0.115 ⁺	0.107 ⁺	0.179 ⁺
142	0.499 ⁻	0.095 ^R	0.118 ⁺	0.124 ⁺	/	0.091 ⁺	0.084 ⁺	0.113 ^R	0.118 ⁺	0.158 ⁺	0.130 ⁺	0.141 ⁺	0.085 ⁺	0.165 ⁺	0.227 ⁺	0.120 ⁺	0.158 ⁺	0.133 ⁺	0.180 ⁺	0.164 ⁺	0.258 ⁺
143	0.510 ⁻	0.113 ^R	0.104 ^R	0.124 ⁺	0.091 ⁺	/	0.094 ^R	0.118 ^R	0.100 ⁺	0.130 ⁺	0.104 ^R	0.109 ⁺	0.089 ^R	0.159 ⁺	0.196 ⁺	0.116 ^R	0.137 ⁺	0.128 ^R	0.155 ⁺	0.152 ⁺	0.231 ⁺
144	0.521 ⁻	0.109 ^R	0.103 ⁺	0.109 ⁺	0.084 ⁺	0.094 ^R	/	0.073 ^R	0.123 ⁺	0.160 ⁺	0.116 ⁺	0.127 ⁺	0.094 ⁺	0.156 ⁺	0.223 ⁺	0.116 ⁺	0.152 ⁺	0.121 ⁺	0.151 ⁺	0.151 ⁺	0.249 ⁺
145	0.553 ⁻	0.150 ^R	0.089 ^R	0.084 ⁺	0.113 ^R	0.118 ^R	0.073 ^R	/	0.105 ⁺	0.134 ⁺	0.100 ⁺	0.103 ⁺	0.109 ^R	0.137 ⁺	0.202 ⁺	0.105 ⁺	0.153 ^R	0.110 ⁺	0.145 ⁺	0.136 ⁺	0.220 ⁺
146	0.567 ⁻	0.157 ⁻	0.093 ⁻	0.090 ^R	0.118 ⁺	0.100 ⁺	0.123 ⁺	0.105 ⁺	/	0.070 ⁺	0.070 ^R	0.073 ⁺	0.104 ⁻	0.107 ⁺	0.143 ⁺	0.099 ^R	0.130 ^R	0.110 ^R	0.106 ⁺	0.119 ⁺	0.181 ⁺
147	0.594 ⁻	0.187 ⁻	0.122 ⁻	0.111 ⁺	0.158 ⁺	0.130 ⁺	0.160 ⁺	0.134 ⁺	0.070 ⁺	/	0.093 ⁻	0.088 ⁺	0.136 ⁺	0.109 ⁺	0.116 ⁺	0.111 ⁺	0.132 ^R	0.123 ⁻	0.101 ⁺	0.128 ^R	0.154 ⁺
148	0.572 ⁻	0.169 ^R	0.112 ⁻	0.109 ^R	0.130 ⁺	0.104 ^R	0.116 ⁺	0.100 ⁺	0.070 ^R	0.093 ⁺	/	0.062 ⁺	0.112 ⁻	0.118 ⁺	0.159 ⁺	0.106 ⁺	0.144 ^R	0.118 ⁺	0.117 ⁺	0.120 ⁺	0.188 ⁺
149	0.587 ⁻	0.178 ⁻	0.107 ⁻	0.099 ⁺	0.141 ⁺	0.109 ⁻	0.127 ⁺	0.105 ⁺	0.073 ⁺	0.088 ⁺	0.062 ⁻	/	0.122 ⁻	0.121 ⁺	0.149 ⁺	0.103 ⁻	0.132 ^R	0.099 ⁺	0.101 ⁺	0.114 ⁺	0.168 ⁺
150	0.535 ⁻	0.121 ^R	0.107 ⁺	0.116 ⁺	0.085 ⁻	0.089 ^R	0.094 ⁺	0.109 ^R	0.104 ⁺	0.136 ⁺	0.112 ⁺	0.122 ⁺	/	0.143 ⁺	0.200 ⁺	0.099 ⁺	0.123 ⁺	0.099 ⁺	0.156 ⁺	0.134 ⁺	0.122 ⁺
151	0.618 ⁻	0.175 ⁻	0.125 ⁻	0.105 ⁻	0.165 ⁻	0.159 ⁻	0.156 ⁻	0.137 ⁻	0.107 ⁻	0.109 ⁻	0.118 ⁻	0.121 ⁻	0.143 ⁻	/	0.105 ⁺	0.094 ⁻	0.106 ^R	0.124 ⁻	0.105 ⁺	0.082 ^R	0.139 ⁺
152	0.651 ⁻	0.229 ⁻	0.179 ⁻	0.162 ⁻	0.227 ⁻	0.196 ⁻	0.225 ⁻	0.202 ⁻	0.143 ⁻	0.116 ⁻	0.159 ⁻	0.149 ⁻	0.210 ⁻	0.105 ⁻	/	0.156 ⁻	0.138 ⁻	0.183 ⁻	0.120 ⁻	0.138 ⁻	0.108 ^R
153	0.578 ⁻	0.148 ^R	0.101 ⁻	0.097 ^R	0.120 ⁻	0.116 ^R	0.116 ⁻	0.105 ⁻	0.099 ^R	0.111 ⁺	0.106 ⁻	0.103 ⁺	0.099 ⁻	0.094 ⁺	0.156 ⁺	/	0.102 ^R	0.086 ⁺	0.105 ⁺	0.083 ⁺	0.171 ⁺
154	0.585 ⁻	0.146 ⁻	0.112 ⁻	0.112 ^R	0.158 ⁻	0.137 ⁻	0.152 ⁻	0.153 ^R	0.130 ^R	0.132 ^R	0.144 ^R	0.132 ^R	0.122 ⁻	0.106 ^R	0.138 ⁺	0.102 ^R	/	0.101 ^R	0.099 ^R	0.161 ⁺	
155	0.583 ⁻	0.151 ^R	0.099 ⁻	0.092 ^R	0.133 ⁻	0.128 ^R	0.121 ⁻	0.110 ⁻	0.110 ^R	0.132 ⁺	0.118 ⁻	0.115 ⁺	0.099 ⁻	0.124 ⁺	0.183 ⁺	0.086 ⁻	0.101 ^R	/	0.144 ⁺	0.090 ⁺	0.182 ⁺
156	0.625 ⁻	0.199 ⁻	0.132 ⁻	0.115 ⁻	0.180 ⁻	0.155 ⁻	0.170 ⁻	0.145 ⁻	0.106 ⁻	0.101 ⁻	0.117 ⁻	0.101 ⁻	0.156 ⁻	0.105 ⁻	0.120 ⁺	0.105 ⁻	0.121 ^R	0.144 ⁻	/	0.087 ^R	0.116 ⁺
157	0.622 ⁻	0.181 ⁻	0.124 ⁻	0.107 ⁻	0.169 ⁻	0.152 ⁻	0.151 ⁻	0.136 ⁻	0.119 ⁻	0.128 ^R	0.120 ⁻	0.114 ⁻	0.134 ⁻	0.082 ^R	0.138 ⁺	0.083 ⁻	0.099 ^R	0.090 ⁻	0.087 ^R	/	0.137 ⁺
158	0.701 ⁻	0.266 ⁻	0.202 ⁻	0.179 ⁻	0.258 ⁻	0.231 ⁻	0.249 ⁻	0.220 ⁻	0.181 ⁻	0.154 ⁻	0.186 ⁻	0.166 ⁻	0.224 ⁻	0.139 ⁻	0.108 ^R	0.171 ⁻	0.161 ⁻	0.182 ⁻	0.116 ⁻	0.137 ⁻	/