

A clustering approach for pairwise comparison matrices

Kolos Csaba Ágoston, Sándor Bozóki & László Csató

To cite this article: Kolos Csaba Ágoston, Sándor Bozóki & László Csató (2025) A clustering approach for pairwise comparison matrices, Journal of the Operational Research Society, 76:5, 971-983, DOI: [10.1080/01605682.2024.2406231](https://doi.org/10.1080/01605682.2024.2406231)

To link to this article: <https://doi.org/10.1080/01605682.2024.2406231>



© 2024 The Author(s). Published by Informa UK Limited, trading as Taylor & Francis Group



View supplementary material [↗](#)



Published online: 07 Oct 2024.



Submit your article to this journal [↗](#)



Article views: 506



View related articles [↗](#)



View Crossmark data [↗](#)

RESEARCH ARTICLE

OPEN ACCESS



A clustering approach for pairwise comparison matrices

Kolos Csaba Ágoston^a, Sándor Bozóki^{a,b} and László Csató^{a,b}

^aInstitute of Operations and Decision Sciences, Department of Operations Research and Actuarial Sciences, Corvinus University of Budapest (BCE), Budapest, Hungary; ^bInstitute for Computer Science and Control (SZTAKI), Hungarian Research Network (HUN-REN), Laboratory on Engineering and Management Intelligence, Research Group of Operations Research and Decision Systems, Budapest, Hungary

ABSTRACT

We consider clustering in group decision making where the opinions are given by pairwise comparison matrices. In particular, the k-medoids model is suggested to classify the matrices since it has a linear programming problem formulation that may contain any condition on the properties of the cluster centres. Its objective function depends on the measure of dissimilarity between the matrices but not on the weights derived from them. Our methodology provides a convenient tool for decision support, for instance, it can be used to quantify the reliability of the aggregation. The proposed theoretical framework is applied to a large-scale experimental dataset, on which it is able to automatically detect some mistakes made by the decision-makers, as well as to identify a common source of inconsistency.

ARTICLE HISTORY

Received 21 February 2024
Accepted 14 September 2024

KEYWORDS

Analytic Hierarchy Process (AHP); clustering; decision analysis; large-scale group decision making; pairwise comparison matrix

“There can be many reasons for searching for representative objects. Not only can these objects provide a characterization of the clusters, but they can often be used for further work or research, especially when it is more economical or convenient to use a small set of k objects instead of the large set one started off with.”¹

1. Introduction

The fast development of information technology has enabled large-scale group decision making (LSGDM): in many decision making problems, the possible number of decision makers (DMs) now easily reaches thousands (Garcia-Zamora et al., 2022). However, because the cognitive abilities of humans have not evolved parallel to the exponential growth of data, it is necessary to aggregate the opinions of DMs (Fan et al., 2024; Tang & Liao, 2021). For this purpose, one of the most widely used techniques is *clustering*: allocating the individual judgements into groups such that judgements in the same group (called a cluster) are more similar to each other than to those in other groups (clusters).

The Analytic Hierarchy Process (AHP) (Saaty, 1977, 1980) is a popular multi-criteria decision making (MCDM) methodology, hence, solving group AHP (GAHP) models is almost as old as AHP itself (Aczél & Saaty, 1983). Traditionally, there are two main approaches to aggregating individual preferences and

creating a group consensus: (a) aggregating the individual pairwise comparison matrices (PCMs) and deriving priorities from the common matrix (Aczél & Saaty, 1983); and (b) deriving priorities from the individual PCMs and aggregating these priorities (Basak & Saaty, 1993). Ossadnik et al. (2016) recommend the second option as no other aggregation technique satisfies a comparable number of evaluation criteria. However, this approach is sensitive to extreme opinions, the supposed consensus may not accurately portray the true group preference (Amenta et al., 2020). Furthermore, Duleba and Szádóczki (2022) find that distance-based aggregation outperforms conventional methods for a high number of DMs.

Even though cluster analysis of rankings or preferences is sometimes used in the literature, we think the topic deserves further attention. Albano et al. (2024) apply a clustering approach based on preference relations. However, the preferences in this study are always consistent and not expressed *via* PCMs. A more closely related research field is clustering for fuzzy preference relations and consensus reaching (García-Lapresta & Pérez-Román, 2016; Kamis et al., 2018; Meng et al., 2023; Wu & Xu, 2018; Xu et al., 2015; Zhang et al., 2018). However, none of these papers use (dis)similarity measures proposed for multiplicative PCMs as the basis of clustering.

On the other hand, we do not know of any studies where individual PCMs are clustered in order to

CONTACT László Csató ✉ laszlo.csato@uni-corvinus.hu Institute of Operations and Decision Sciences, Department of Operations Research and Actuarial Sciences, Corvinus University of Budapest (BCE), Budapest, Hungary.

Supplemental data for this article can be accessed online at <https://doi.org/10.1080/01605682.2024.2406231>.

© 2024 The Author(s). Published by Informa UK Limited, trading as Taylor & Francis Group

This is an Open Access article distributed under the terms of the Creative Commons Attribution-NonCommercial-NoDerivatives License (<http://creativecommons.org/licenses/by-nc-nd/4.0/>), which permits non-commercial re-use, distribution, and reproduction in any medium, provided the original work is properly cited, and is not altered, transformed, or built upon in any way. The terms on which this article has been published allow the posting of the Accepted Manuscript in a repository by the author(s) or with their consent.

aggregate them. The current paper aims to fill this research gap, which is our main contribution to the extant literature.

The proposed clustering methodology can be used for several purposes:

- it might provide an alternative aggregation procedure if the number of clusters is restricted to one;
- it gives information on the ability of aggregated preferences to represent the individual preferences;
- it makes it possible to detect some mistakes in the original data, which is far from trivial in LSGDM problems;
- it can help to identify common sources of inconsistency.

In addition, a crucial advantage of our clustering method is the lack of any requirement on the number of missing entries in the PCMs, which can be especially useful in LSGDM problems where the underlying pairwise comparisons are obtained *via* a questionnaire that can be finished at every point (Fan et al., 2024). Naturally, a complete PCM will have a stronger effect on clustering, but this seems to be a fair compensation for the DMs who have devoted more effort and time to reveal their preferences.

The proposed clustering approach guarantees that the cluster centres are PCMs given by at least one DM. This is also an attractive property because a cluster centre is difficult to interpret if it strongly differs from individual opinions in a MCDM problem.

Last but not least, the suggested k -medoids clustering has a linear programming (LP) formulation, which allows for imposing additional restrictions. For example, the inconsistency of the cluster centres can be required to remain below an exogenous threshold.

The theoretical framework is applied to the experimental data of Bozóki et al. (2013).

The paper is structured as follows. Section 2 presents the main concepts related to pairwise comparison matrices and examines some dissimilarity measures suggested for them. Section 3 introduces cluster analysis and shows the LP formulation of the k -medoids clustering. The numerical results are presented in Section 4. Section 5 discusses the incorporation of demographic factors that have been suggested by an anonymous reviewer. Finally, Section 6 offers concluding remarks.

2. Pairwise comparison matrices and their similarity

First, the theoretical background of using pairwise comparison matrices and their measures of similarity—the basis of any clustering algorithm—are discussed. Section 2.1 introduces pairwise comparison matrices, while Section 2.2 overviews

some measures to quantify the (dis)similarity of pairwise comparison matrices.

2.1. Pairwise comparison matrices

Pairwise comparison matrices are often used to determine the cardinal preferences of DMs, who answer questions like “How many times is a criterion more important than another one?” or “How many times is a given alternative better than another one with respect to a fixed criterion?” These pairwise ratios are collected into the $n \times n$ pairwise comparison matrix $A = [a_{ij}]$, where $a_{ij} > 0$ and $a_{ij} = 1/a_{ji}$ hold for all $1 \leq i, j \leq n$. A is called consistent if $a_{ij}a_{jk} = a_{ik}$ is satisfied for all $1 \leq i, j, k \leq n$. Otherwise, when cardinal transitivity is violated, the matrix is called inconsistent.

From a practical point of view, it is natural and reasonable to allow for some degree of inconsistency. Several inconsistency indices have been proposed to measure the contradictions in a pairwise comparison matrix (Bozóki & Rapcsák, 2008; Brunelli, 2018; Csató, 2019a). In particular, Saaty (1977) has suggested the so-called *inconsistency index* CI in his celebrated theory of AHP, which is an affine transformation of the dominant eigenvalue $\lambda_{\max}$ of the pairwise comparison matrix A :

$$CI(A) = \frac{\lambda_{\max} - n}{n - 1}.$$

This is divided by the random index RI_n , the average CI of a large number of randomly generated $n \times n$ pairwise comparison matrices to obtain the *inconsistency ratio* $CR(A) = CI(A)/RI_n$. According to Saaty, CR should remain below 0.1 to accept the matrix as a reasonable representation of consistent preferences.

Analogously, many weighting methods are available to find a positive weight vector $\mathbf{w} = (w_1, w_2, \dots, w_n)$ such that the ratios w_i/w_j give good approximations to the pairwise comparisons a_{ij} (Choo & Wedley, 2004). The *Logarithmic Least Squares Method* (LLSM) (Crawford & Williams, 1985; Csató, 2019b; De Graan, 1980; Fichtner, 1986; Rabinowitz, 1976) is one of the most popular:

$$\begin{aligned} \min \sum_{i,j} \left[\log a_{ij} - \log \left(\frac{w_i}{w_j} \right) \right]^2 \\ \sum_{i=1}^n w_i = 1, \\ w_i > 0, \quad i = 1, 2, \dots, n. \end{aligned}$$

The unique solution to this optimisation problem is given by the geometric means of the entries in the rows:

$$\frac{w_i}{w_j} = \sqrt[n]{\prod_{k=1}^n a_{ik}} \frac{\sqrt[n]{\prod_{k=1}^n a_{jk}}}{\sqrt[n]{\prod_{k=1}^n a_{jk}}}.$$

Therefore, it is often called the *geometric mean* method.

2.2. Quantifying the similarity of pairwise comparison matrices

Several ideas exist in the literature for computing the (dis)similarity of two pairwise comparison matrices A and B . Crawford and Williams (1985, p. 389) essentially introduce the following metric:

$$D_1(A, B) = \sqrt{\left(\sum_{i=1}^n \sum_{j=1}^n (\log(a_{ij}) - \log(b_{ij}))^2 \right)}.$$

The original definition takes the sum only for the entries above the diagonal, but this is equivalent to D_1 except for a constant factor. Obviously, D_1 is symmetric and equals 0 if and only if $A = B$. It also satisfies the triangle inequality (Crawford & Williams, 1985).

Tekile et al. (2023) measure the closeness of two complete pairwise comparison matrices by another formula called Manhattan or L^1 distance:

$$D_2(A, B) = \sum_{i=1}^n \sum_{j=1}^n |\log(a_{ij}) - \log(b_{ij})|.$$

D_2 also satisfies the triangle inequality.

Kuřakowski et al. (2022) use another indicator, the so-called compatibility index, which was originally defined by Saaty (2008):

$$D_{3u}(A, B) = \frac{1}{n^2} \left(\sum_{i=1}^n \sum_{j=1}^n a_{ij} b_{ji} \right).$$

While D_{3u} is symmetric, $D_{3u}(A, B) \geq n^2$. Hence, it is worth normalising D_{3u} as has been done in Ágoston and Csátó (2024):

$$D_3(A, B) = \frac{1}{n^2} \left(\sum_{i=1}^n \sum_{j=1}^n (a_{ij} b_{ji} - 1) \right).$$

D_3 remains symmetric, and $D_3(A, B) = 0$ if and only if $A = B$. However, D_3 is not a distance.

Lemma 2.1 Dissimilarity index D_3 does not satisfy the triangle inequality.

Proof. It is sufficient to give a counterexample. Consider the following three pairwise comparison matrices:

$$A = \begin{bmatrix} 1 & 2 & 1 \\ 1/2 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}, \quad B = \begin{bmatrix} 1 & 3 & 1 \\ 1/3 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}, \quad C = \begin{bmatrix} 1 & 4 & 1 \\ 1/4 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}.$$

It can be checked that $D_3(A, B) = 1/6$, $D_3(B, C) = 1/12$, but $D_3(A, C) = 1/2 > 1/4 = D_3(A, B) + D_3(B, C)$. $\square$

Kuřakowski et al. (2022) consider three other versions of the compatibility index (all of them are modified to ensure $D(A, A) = 0$ and $D(A, B) \geq 0$):

$$D_4(A, B) = \frac{2}{n(n-1)} \left(\sum_{i=1}^{n-1} \sum_{j=i+1}^n \max\{a_{ij} b_{ji}, a_{ji} b_{ij}\} - 1 \right);$$

$$D_5(A, B) = \max\{a_{ij} b_{ji} - 1 : 1 \leq i, j \leq n\};$$

$$D_6(A, B) = -\frac{2}{n(n-1)} \left(\sum_{i=1}^{n-1} \sum_{j=i+1}^n \min\{a_{ij} b_{ji}, a_{ji} b_{ij}\} - 1 \right).$$

Based on D_4 – D_6 , another reasonable index could be

$$D_7(A, B) = -\min\{a_{ij} b_{ji} - 1 : 1 \leq i, j \leq n\}.$$

Lemma 2.2 Indices D_4 and D_5 do not satisfy the triangle inequality.

Proof. The PCMs from the proof of Lemma 2.1 can be used: $D_4(A, B) = 1/6$, $D_4(B, C) = 1/9$ but $D_4(A, C) = 1/3 > 5/18 = D_4(A, B) + D_4(B, C)$, and $D_5(A, B) = 1/2$, $D_5(B, C) = 1/3$ but $D_5(A, C) = 1 > 5/6 = D_5(A, B) + D_5(B, C)$. $\square$

Proposition 2.3 Indices D_6 and D_7 satisfy the triangle inequality, thus, they are distances on the set of pairwise comparison matrices.

Proof. For both D_6 and D_7 , it is sufficient to show that the triangle inequality holds elementwise for any pairwise comparison matrices $A = [a_{ij}]$, $B = [b_{ij}]$, and $C = [c_{ij}]$. $1 = a_{ij} \leq b_{ij} \leq c_{ij}$ can be assumed without loss of generality. Then, by elementary calculus, the following three inequalities hold:

$$1 - \frac{a_{ij}}{b_{ij}} + 1 - \frac{b_{ij}}{c_{ij}} \geq 1 - \frac{a_{ij}}{c_{ij}};$$

$$1 - \frac{a_{ij}}{c_{ij}} + 1 - \frac{b_{ij}}{c_{ij}} \geq 1 - \frac{a_{ij}}{b_{ij}};$$

$$1 - \frac{a_{ij}}{b_{ij}} + 1 - \frac{a_{ij}}{c_{ij}} \geq 1 - \frac{b_{ij}}{c_{ij}}.$$

$\square$

Fichtner (1984) verifies that the popular eigenvector method (Saaty, 1977) can be defined by a metric, too. However, although it satisfies the requirements of a distance, it is discontinuous.

All of the above dissimilarity measures D_1 – D_7 are invariant under the relabeling of the alternatives. Consequently, they are invariant to transposition, that is,

$$D_i(A, B) = D_i(A^\top, B^\top),$$

where $A^\top$ is the transpose of the pairwise comparison matrix A .

Since the input of the k -medoids problem is the dissimilarity matrix of the objects, any of the indices

D_1 – D_7 can be used in the proposed methodology. Our numerical experiment will consider D_1 and D_3 .

3. Cluster analysis and the k -medoids problem

Clustering, the problem of dividing a set of objects into groups (clusters) such that any object is more similar to an arbitrary object in its group than to any object in a distinct group, is one of the most popular methods in exploratory data science. Although no “perfect” clustering algorithm exists, perhaps the most natural solution is k -means clustering, when each object belongs to the cluster with the nearest mean.

However, this leads to an NP-hard non-convex discrete optimisation problem. Therefore, most state-of-the-art statistical and data science software implements a heuristic iterative algorithm that alternates two steps: i) assigning all points to the nearest cluster centre; and ii) recalculating the cluster centres. The operation in step ii) is quite trivial in Euclidean spaces (take the average in every dimension), but it may be challenging if the objects do not belong to a multidimensional Euclidean space (Majstorović et al., 2018). Furthermore, the heuristic iterative algorithm above can converge to a local rather than global optimum. Hence, Ágoston and E.-Nagy (2024) provide a mixed integer linear programming (MILP) formulation to obtain an exact solution for the minimum sum-of-clustering problem that allows for adding arbitrary linear constraints and can be solved up to 150 objects.

An alternative approach to k -means clustering is the k -medoids problem (Kaufman & Rousseeuw, 1986; Schubert & Rousseeuw, 2019). Here, all cluster centres should be objects themselves, allowing for easier interpretability of the cluster centres. This seems to be an important advantage in LSGDM since the DMs may not be willing to accept a solution that has not been suggested by any of them. In addition, k -medoids can be used with an arbitrary dissimilarity measure, while k -means generally require Euclidean distance. Last but not least, minimising the sum of pairwise dissimilarities instead of the sum of squared Euclidean distances makes the k -medoids problem less sensitive to outliers than the k -means problem.

For unknown reasons, the k -medoids model is not widely used and is not implemented in usual statistical software packages. Nonetheless, the problem can be formulated as an LP model, which is described in the following.

Assume that the $m \times m$ matrix $\Delta = [\delta_{ij}]$ contains the (pairwise) dissimilarities between any two of the m objects (individual PCMs). For example, $\delta_{ij} =$

$D_1(A^{(i)}, A^{(j)})$ if measure D_1 is used to determine the dissimilarity of the preferences given by matrices $A^{(i)}$ and $A^{(j)}$ of decision makers i and j . Let M denote the set of integers $\{1, \dots, m\}$.

The k -medoids problem is as follows (see, for instance, Vinod (1969)):

$$\sum_{i=1}^m \sum_{j=1}^m \delta_{ij} x_{ij} \rightarrow \min \quad (1)$$

$$\text{s.t.} \quad \sum_{j=1}^m x_{ij} = 1 \quad \forall i \in M \quad (2)$$

$$x_{ij} \leq y_j \quad \forall i, j \in M \quad (3)$$

$$\sum_{j=1}^m y_j = k \quad (4)$$

$$x_{ij} \geq 0 \quad \forall i, j \in M \quad (5)$$

$$y_j \in \{0, 1\} \quad \forall j \in M \quad (6)$$

In this formulation, the binary variable y_j equals one if object j is a cluster centre (and zero otherwise), while x_{ij} equals one if object i is placed in the cluster whose centre is object j (and zero otherwise). Constraint (4) ensures that there are k different cluster centres. According to (1), each object is classified into exactly one cluster. Finally, due to the constraints (3), object i can be associated with object j ($x_{ij} = 1$) only if object j is a cluster centre, namely, $y_j = 1$.

The LP model (1)–(6) can be solved with standard solvers in small and medium size instances. The above formulation allows any PCM to be a cluster centre. However, further restrictions can be easily added to the LP model, for example, by requiring that the inconsistency ratio CR of each cluster centre is below an exogenous threshold.

Originally, clustering was intended to form groups in Euclidean spaces, but the underlying idea can be applied in other fields. Clustering algorithms appear in network science, where the goal is to classify nodes in a graph such that the nodes in a given cluster are connected more strongly to each other than to nodes in other clusters. Other objects can also be clustered, for example, curves or mortality tables.

To conclude, cluster analysis requires a clustering algorithm and a dissimilarity measure, which is not necessarily a distance. This paper considers the k -medoids problem due to its advantages discussed above.

4. Numerical results

Bozóki et al. (2013) have collected hundreds of pairwise comparison matrices in a controlled experiment where even the questioning order is known for each matrix. However, here we use only the final complete pairwise comparison matrices. The matrices are distinguished by

their size ($n = 4$, $n = 6$, $n = 8$ alternatives) and type (objective, where the students compared maps; subjective, where the students compared summer houses). The sample sizes are reported in Table 1.

4.1. Analysing a sample of objective type

In the case of dataset M4, the “correct” consistent pairwise comparison matrix provided by the ratios of country areas that appear on the map is known:

$$\begin{bmatrix} 1 & 1.691 & 0.282 & 0.770 \\ 0.591 & 1 & 0.167 & 0.455 \\ 3.544 & 5.991 & 1 & 2.725 \\ 1.300 & 2.198 & 0.367 & 1 \end{bmatrix}.$$

For $k = 2$ clusters, the cluster centres coincide for both dissimilarity measures D_1 and D_3 :

$$CC_1^{(1)} = CC_3^{(1)} = \begin{bmatrix} 1 & 1.500 & 0.286 & 0.833 \\ 0.667 & 1 & 0.154 & 0.435 \\ 3.500 & 6.500 & 1 & 2.300 \\ 1.200 & 2.300 & 0.435 & 1 \end{bmatrix},$$

and

$$CC_1^{(2)} = CC_3^{(2)} = \begin{bmatrix} 1 & 1.200 & 0.400 & 0.909 \\ 0.833 & 1 & 0.200 & 0.500 \\ 2.500 & 5.000 & 1 & 2.300 \\ 1.100 & 2.000 & 0.500 & 1 \end{bmatrix}.$$

The two clusters are the same: 47 matrices belong to first and 19 to the second group.

Figure 1 shows boxplots for the inconsistency ratios CR in the two clusters.

We have used multidimensional scaling (MDS) to visualise the pairwise comparison matrices (Kruskal, 1964; Kruskal & Wish, 1978). This technique places each matrix into a lower-dimensional space such that the original distances are preserved to the extent possible. In Figure 2, red dots and blue squares represent the two clusters. A darker mark is associated with a higher inconsistency. The two cluster centres are green and have double sizes. Clearly, the clusters are not determined by the inconsistency of the matrices. The axes are artificial measures in a two-dimensional Euclidean space without

a clear interpretation, as provided by the `cmdscale()` function of R.

The main findings can be summarised as follows:

- The clusters are quite balanced with respect to the number of matrices in them;
- The results are insensitive to the dissimilarity measure (D_1 or D_3) used;
- The two groups are not distinguished by the level of inconsistency, although the variance of CR is somewhat higher for the first group.

The analysis has been repeated with higher numbers of clusters. A strong relationship remains between the resulting groups for different values of k , as well as for the two dissimilarity measures D_1 and D_3 .

In sample M8 with measure D_3 and $k = 2$ clusters, an interesting example has been found that uncovers a possible application of our methodology.

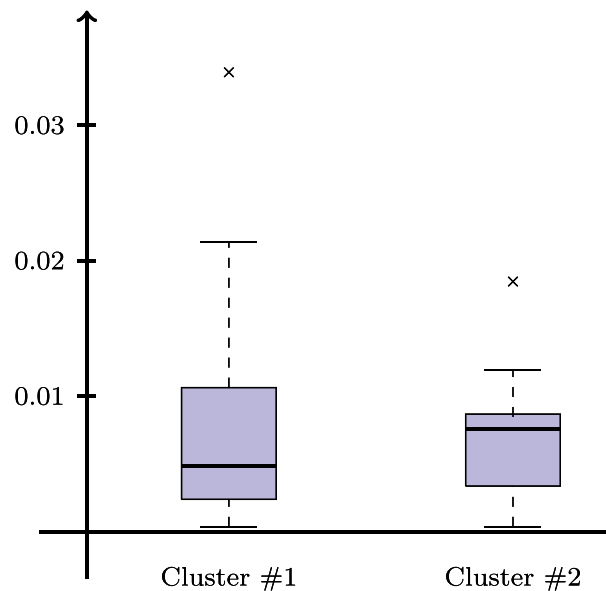


Figure 1. Distribution of inconsistency ratios CR , dataset M4, $k = 2$ clusters.

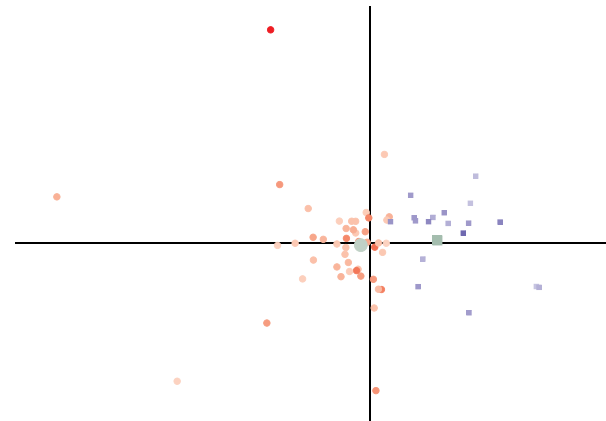


Figure 2. Approximated coordinates of PCMs, dataset M4, measure D_3 .

Notes: The two clusters are represented by red dots and blue squares. A darker mark is associated with a higher inconsistency. The cluster centers are green and have double sizes.

Table 1. Sample sizes in the experimental database.

Label	Type	Size (n)	Number of PCMs (m)
M4	Objective (map)	4	66
M6	Objective (map)	6	77
M8	Objective (map)	8	82
S4	Subjective (summer house)	4	68
S6	Subjective (summer house)	6	77
S8	Subjective (summer house)	8	78

The cluster centres are:

$$CC_3^{(1)} = \begin{bmatrix} 1 & 1.800 & 0.769 & 0.556 & 1.800 & 5.000 & 1.100 & 3.000 \\ 0.556 & 1 & 0.455 & 0.313 & 1.100 & 2.700 & 0.909 & 1.500 \\ 1.300 & 2.200 & 1 & 0.833 & 3.000 & 7.000 & 1.800 & 3.300 \\ 1.800 & 3.200 & 1.200 & 1 & 5.000 & \mathbf{10.100} & 2.200 & 6.000 \\ 0.556 & 0.909 & 0.333 & 0.200 & 1 & 2.300 & 0.625 & 1.400 \\ 0.200 & 0.370 & 0.143 & 0.099 & 0.435 & 1 & 0.250 & 0.769 \\ 0.909 & 1.100 & 0.556 & 0.455 & 1.600 & 4.000 & 1 & 2.200 \\ 0.333 & 0.667 & 0.303 & 0.167 & 0.714 & 1.300 & 0.455 & 1 \end{bmatrix},$$

and

$$CC_3^{(2)} = \begin{bmatrix} 1 & 1.500 & 0.500 & 0.500 & 3.000 & 6.000 & 1.200 & 3.000 \\ 0.667 & 1 & 0.500 & 0.333 & 1.500 & 3.000 & 0.667 & 1.500 \\ 2.000 & 2.000 & 1 & 0.667 & 5.000 & 8.000 & 1.500 & 4.000 \\ 2.000 & 3.000 & 1.500 & 1 & 4.500 & \mathbf{0.100} & 2.000 & 5.000 \\ 0.333 & 0.667 & 0.200 & 0.222 & 1 & 2.000 & 0.667 & 1.200 \\ 0.167 & 0.333 & 0.125 & 10.000 & 0.500 & 1 & 0.286 & 0.500 \\ 0.833 & 1.500 & 0.667 & 0.500 & 1.500 & 3.500 & 1 & 2.500 \\ 0.333 & 0.667 & 0.250 & 0.200 & 0.833 & 2.000 & 0.400 & 1 \end{bmatrix}.$$

The first cluster contains 81 matrices and the second cluster consists of only one. The two cluster centres are similar except for the preference between the fourth and sixth alternatives (countries), highlighted in bold font. Since the pairwise comparisons show the ratio of the area of two countries, the only PCM in the second cluster is almost certain to contain a typo: the student has thought that country 4 is ten times larger than country 6 but mistakenly written the reciprocal 1/10, which is a standard mistake. Unsurprisingly, the associated PCM fundamentally differs from all other matrices, and the presented clustering approach is able to detect the outlier without any other specification. This can be highly advantageous in LSGDM problems.

4.2. Analysing a sample of subjective type

When the students have compared summer houses, no “natural” PCM exists around which the opinions are centred.

Since sample S4 contains four different houses, we have started the analysis with four clusters, supposing that each alternative will be the best in one cluster.

For the measure D_1 , the cluster centres are as follows:

$$\begin{aligned} CC_1^{(1)} &= \begin{bmatrix} 1 & 2.000 & 7.000 & 3.000 \\ 0.500 & 1 & 5.000 & 2.000 \\ 0.143 & 0.200 & 1 & 0.500 \\ 0.333 & 0.500 & 2.000 & 1 \end{bmatrix}, & CC_1^{(2)} &= \begin{bmatrix} 1 & 1.500 & 0.500 & 0.333 \\ 0.667 & 1 & 0.333 & 0.333 \\ 2.000 & 3.000 & 1 & 0.667 \\ 3.000 & 3.000 & 1.500 & 1 \end{bmatrix}, \\ CC_1^{(3)} &= \begin{bmatrix} 1 & 3.000 & 5.000 & 0.400 \\ 0.333 & 1 & 3.000 & 0.143 \\ 0.200 & 0.333 & 1 & 0.111 \\ 2.500 & 7.000 & 9.000 & 1 \end{bmatrix}, & CC_1^{(4)} &= \begin{bmatrix} 1 & 6.000 & 3.000 & 3.000 \\ 0.167 & 1 & 0.200 & 0.200 \\ 0.333 & 5.000 & 1 & 0.500 \\ 0.333 & 5.000 & 2.000 & 1 \end{bmatrix}. \end{aligned}$$

On the other hand, for the measure D_3 :

$$CC_3^{(1)} = \begin{bmatrix} 1 & 2.000 & 7.000 & 3.000 \\ 0.500 & 1 & 5.000 & 2.000 \\ 0.143 & 0.200 & 1 & 0.500 \\ 0.333 & 0.500 & 2.000 & 1 \end{bmatrix}, \quad CC_3^{(2)} = \begin{bmatrix} 1 & 1.500 & 0.333 & 0.200 \\ 0.667 & 1 & 0.333 & 0.200 \\ 3.000 & 3.000 & 1 & 0.500 \\ 5.000 & 5.000 & 2.000 & 1 \end{bmatrix}.$$

Note that the centres of the first clusters (matrices $CC_1^{(1)}$ and $CC_3^{(1)}$) coincide. The cluster sizes are 26, 15, 9, 18 for index D_1 and 27, 9, 12, 20 for index D_3 .

$$CC_3^{(3)} = \begin{bmatrix} 1 & 4.000 & 1.500 & 2.000 \\ 0.250 & 1 & 0.333 & 0.667 \\ 0.667 & 3.000 & 1 & 3.000 \\ 0.500 & 1.500 & 0.333 & 1 \end{bmatrix}, \quad CC_3^{(4)} = \begin{bmatrix} 1 & 5.000 & 3.000 & 1.000 \\ 0.200 & 1 & 0.500 & 0.143 \\ 0.333 & 2.000 & 1 & 0.200 \\ 1.000 & 7.000 & 5.000 & 1 \end{bmatrix}.$$

The contingency table of the two groupings is shown in Table 2. The similarity of the clusters is lower than in the objective task (Section 4.1), but the largest cluster almost coincides.

Table 2. Contingency table for sample S4.

	$\mathcal{C}_3^{(1)}$	$\mathcal{C}_3^{(2)}$	$\mathcal{C}_3^{(3)}$	$\mathcal{C}_3^{(4)}$	Σ
$\mathcal{C}_1^{(1)}$	25	0	1	0	26
$\mathcal{C}_1^{(2)}$	0	9	5	1	15
$\mathcal{C}_1^{(3)}$	2	0	0	7	9
$\mathcal{C}_1^{(4)}$	0	0	6	12	18
Σ	27	9	12	20	68

Four clusters, measures D_1 (row) and D_3 (column).

The inconsistency ratios of the cluster centres are 0.007, 0.008, 0.028, 0.063 in the case of D_1 , while 0.007, 0.009, 0.022, 0.013 in the case of D_3 . Similar to the results in Section 4.1, the clusters are relatively uniform with respect to the level of inconsistency as can be seen in Figure 3. Even though the cluster centres are slightly less inconsistent for D_3 , each centre has an acceptable inconsistency according to the famous 10% rule of thumb. Naturally, this is not guaranteed in other datasets but the LP model of the k -medoids problem (Section 3) might contain any (linear) restriction on the inconsistency of the cluster centres.

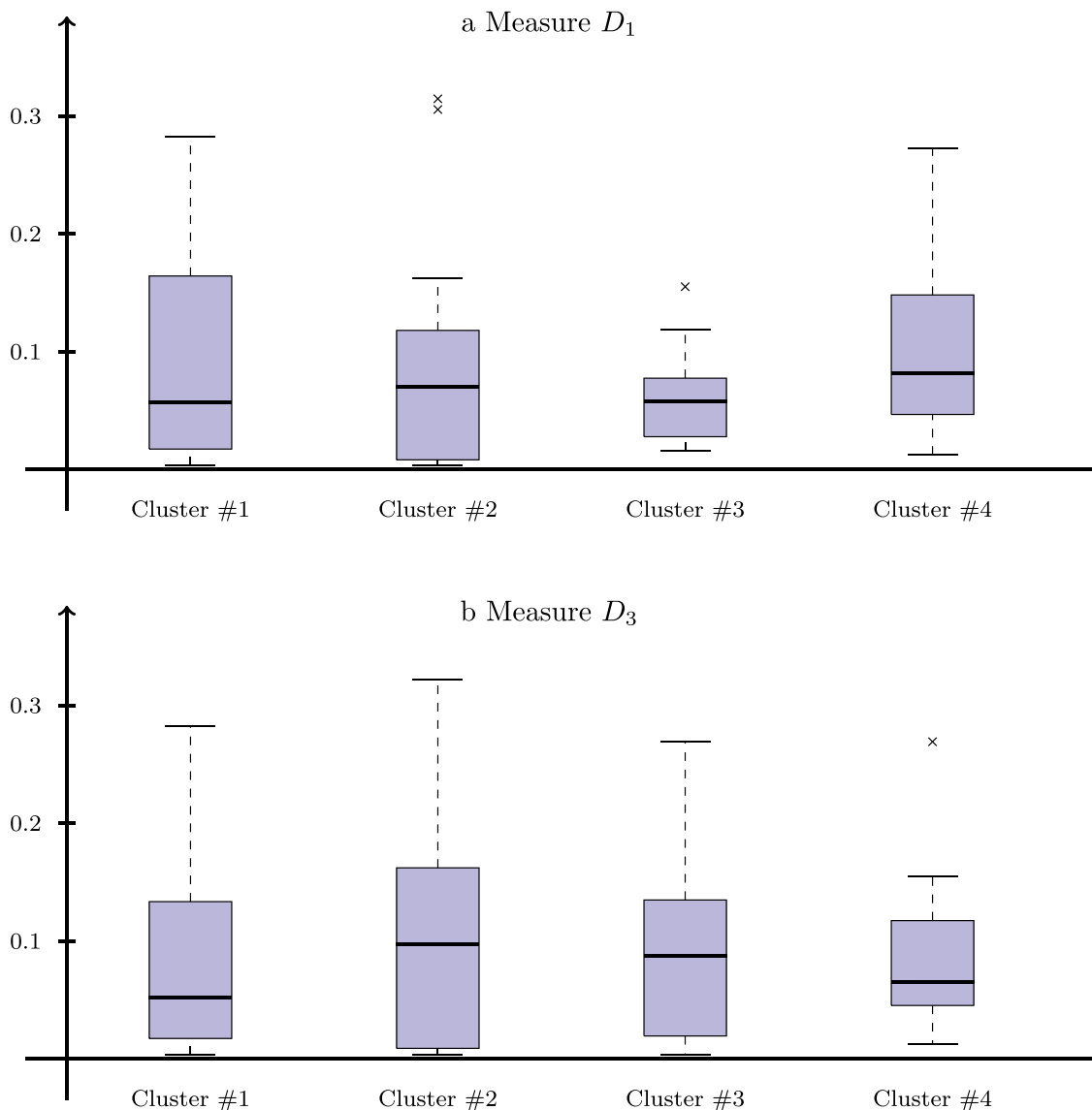


Figure 3. Distribution of inconsistency ratios CR , dataset S4, $k = 4$ clusters.

As the cluster centres are PCMs, the associated priority weights can help the interpretation of the clusters. To that end, we have used the geometric mean method (Section 2.1). For the dissimilarity measure D_1 , they are as follows:

$$\begin{bmatrix} 0.495 \\ 0.291 \\ 0.067 \\ 0.148 \end{bmatrix}, \begin{bmatrix} 0.155 \\ 0.114 \\ 0.310 \\ 0.420 \end{bmatrix}, \begin{bmatrix} 0.260 \\ 0.102 \\ 0.049 \\ 0.589 \end{bmatrix}, \begin{bmatrix} 0.511 \\ 0.054 \\ 0.180 \\ 0.254 \end{bmatrix}.$$

On the other hand, for index D_3 :

$$\begin{bmatrix} 0.495 \\ 0.291 \\ 0.067 \\ 0.148 \end{bmatrix}, \begin{bmatrix} 0.109 \\ 0.089 \\ 0.284 \\ 0.517 \end{bmatrix}, \begin{bmatrix} 0.403 \\ 0.105 \\ 0.339 \\ 0.153 \end{bmatrix}, \begin{bmatrix} 0.368 \\ 0.065 \\ 0.113 \\ 0.455 \end{bmatrix}.$$

Against our conjecture, the second and the third summer houses do not have the highest priority in any cluster; the first and the last alternatives are the best in the two clusters, respectively.

Aggregating the individual pairwise comparison matrices by the geometric means of the entries (Aczél & Saaty, 1983) and applying *LLSM* to the common matrix results in the following priority vector:

$$\begin{bmatrix} 0.410 \\ 0.164 \\ 0.146 \\ 0.279 \end{bmatrix}. \quad (7)$$

The implied ranking $1 \succ 4 \succ 2 \succ 3$ differs from the ranking in any cluster centre, thus, the aggregated ranking does not correspond to the preferences of any group.

Clustering may provide an aggregation procedure if there is only one cluster, whose centre represents “best” the whole set of PCMs. The cluster centres are

$$CC_1 = \begin{bmatrix} 1 & 2.000 & 3.000 & 1.500 \\ 0.500 & 1 & 2.000 & 0.500 \\ 0.333 & 0.500 & 1 & 0.250 \\ 0.667 & 2.000 & 4.000 & 1 \end{bmatrix} \quad \text{and} \quad CC_3 = \begin{bmatrix} 1 & 2.000 & 1.500 & 1.000 \\ 0.500 & 1 & 0.667 & 0.667 \\ 0.667 & 1.500 & 1 & 0.667 \\ 1.000 & 1.500 & 1.500 & 1 \end{bmatrix}$$

for measures D_1 and D_3 , respectively, and the corresponding weight vectors are

$$\begin{bmatrix} 0.381 \\ 0.185 \\ 0.099 \\ 0.334 \end{bmatrix} \quad \text{and} \quad \begin{bmatrix} 0.319 \\ 0.166 \\ 0.219 \\ 0.296 \end{bmatrix},$$

which lead to the rankings $1 \succ 4 \succ 2 \succ 3$ and $1 \succ 4 \succ 3 \succ 2$, respectively. The rankings are identical to the ranking derived from the weights (7) associated with the aggregated matrix, except for a rank reversal at the bottom in the case of dissimilarity measure D_3 . On the other hand, the relative weights of the first and the last alternatives are closer to each other according to both clusters than according to the result obtained by a reasonable aggregation of the individual matrices given in (7).

4.3. The appropriate number of clusters

There is no unique or standard procedure to determine the appropriate number of clusters. Naturally, one might have knowledge about the number of clusters based on expert opinion or the demographic background of the DMs.

Nonetheless, it is a more interesting case if no previous information is available on the expected/ desired number of clusters. A popular choice could be the so-called “elbow” method: the sum of distances from the cluster centres (the objective function of the LP in Section 3) is plotted for all relevant numbers of clusters, and the value for which the decreasing trend flattens is chosen. The corresponding chart for sample S4 is shown in Figure 4, suggesting that $k = 4$ or $k = 5$ clusters is the best option as a further increase in k does not lead to a substantial reduction in the optimal value of the objective function.

A more sophisticated version of the “elbow” method is the gap statistic (Tibshirani et al., 2001). Although it “is designed to be applicable to any clustering method and distance measure” (Tibshirani et al., 2001, p. 411), during the implementation uniformly distributed samples are generated in a multivariate Euclidean space. Therefore, this method does not seem to be suitable for our problem.

Several other approaches are available to obtain the optimal number of clusters but the majority of them are based on the sum of squares. Although the variance could be calculated, we think it has no common meaning for the dissimilarity measures considered.

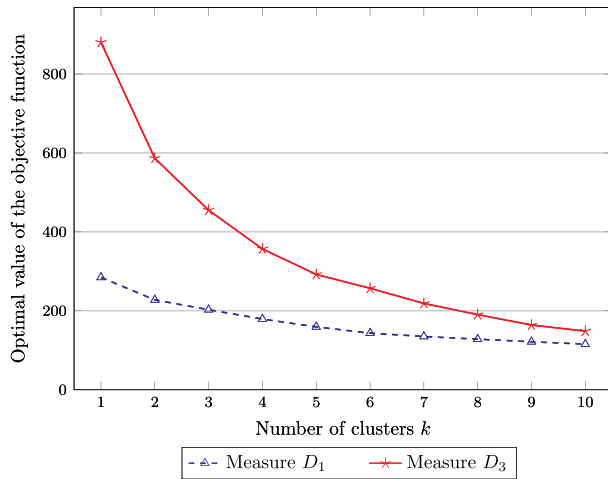


Figure 4. The optimum as a function of the number of clusters, dataset S4.

One exception is the silhouette method (Kaufman & Rousseeuw, 1990), which uses the dissimilarity matrix Δ but does not rely on any specific property of Euclidean spaces. The average silhouette values are given in Table 3. They tend to grow for small values of k , the maximum is reached at 11 (12) clusters in the case of index D_1 (D_3). However, these numbers seem to be unreasonably high.

Another common method is using a hierarchical clustering method and choosing the number of clusters based on dendrograms, which are presented in Figure 5. Note that Figure 5a has been derived by the square of D_1 to make the structure more visible. The dendrogram shows the level of dissimilarity at which the corresponding objects are merged into the same cluster. According to these dendrograms, the number of clusters k should be between three and six in our case.

To conclude, we currently could not recommend any method that immediately gives the number of clusters. The problem is somewhat analogous to the

Table 3. Mean silhouette values, sample S4.

Number of clusters	Measure D_1	Measure D_3
2	0.286	0.467
3	0.285	0.407
4	0.260	0.503
5	0.310	0.510
6	0.293	0.494
7	0.302	0.477
8	0.321	0.488
9	0.337	0.535
10	0.345	0.568
11	0.348	0.556
12	0.325	0.571
13	0.325	0.553
14	0.330	0.545
15	0.327	0.538
16	0.324	0.514
17	0.308	0.512
18	0.310	0.523
19	0.306	0.521
20	0.310	0.518

choice of the dissimilarity measure, where there is no perfect solution, too.

4.4. Typical patterns in pairwise comparison matrices

This section considers another possible application of the proposed clustering approach. In particular, we are not interested in the preferences among the alternatives (i.e. which alternative is the best) but look for typical patterns in the structure of preferences and for common sources of inconsistency. Therefore, the alternatives are relabelled: the most preferred alternative will be the first, the second most preferred the second, and so on. To that end, the alternatives are ranked by the geometric mean method (Section 2.1). For dataset S4, dissimilarity measure D_1 , and four clusters, which is the case presented in Section 4.2, the cluster centres are: and the clusters contain 14, 18, 10, and 26 matrices, respectively.

$$CC_1^{(1)} = \begin{bmatrix} 1 & 2.000 & 5.000 & 5.000 \\ 0.500 & 1 & 3.000 & 3.000 \\ 0.200 & 0.333 & 1 & 1.500 \\ 0.200 & 0.333 & 0.667 & 1 \end{bmatrix},$$

$$CC_1^{(2)} = \begin{bmatrix} 1 & 2.000 & 2.500 & 4.000 \\ 0.500 & 1 & 1.500 & 3.500 \\ 0.400 & 1.667 & 1 & 3.000 \\ 0.250 & 0.286 & 0.333 & 1 \end{bmatrix},$$

$$CC_1^{(3)} = \begin{bmatrix} 1 & 5.000 & 9.000 & 9.000 \\ 0.200 & 1 & 7.000 & 7.000 \\ 0.111 & 0.143 & 1 & 5.000 \\ 0.111 & 0.143 & 0.200 & 1 \end{bmatrix},$$

$$CC_1^{(4)} = \begin{bmatrix} 1 & 3.000 & 5.000 & 7.000 \\ 0.333 & 1 & 3.000 & 6.000 \\ 0.200 & 0.333 & 1 & 4.000 \\ 0.143 & 0.167 & 0.250 & 1 \end{bmatrix},$$

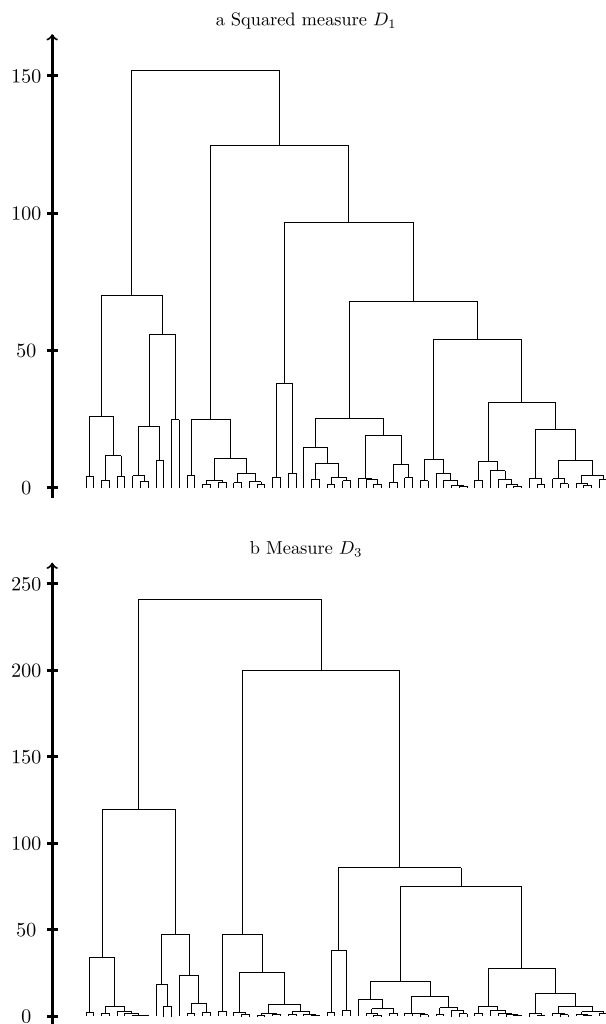


Figure 5. Dendrograms, sample S4.

Figure 6 reports the boxplots of inconsistency ratios in all clusters analogously to Figure 3a. Contrary to the previous results, now there exists a cluster (Cluster #3) that contains highly inconsistent matrices. We have tried to find a pattern in these 10 matrices. Consistency implies $a_{14} = a_{12} \times a_{23} \times a_{34}$. This is definitely not true for $CC_1^{(3)}$, where $a_{41} = 9$ but $a_{12} \times a_{23} \times a_{34} = 175$. The other matrices of Cluster #3 are similar to the cluster centre, including the property that all the six comparisons above the diagonal are higher than one. In other words, while these DMs have clear ordinal preferences, they are not able to appropriately translate them into numbers.

5. Discussion

In the following, two issues are considered: Section 5.1 provides some thoughts on the potential contribution of clustering to the analysis of group decision making problems based on pairwise comparisons, while Section 5.2 discusses how demographic variables can be utilised in the proposed clustering approach.

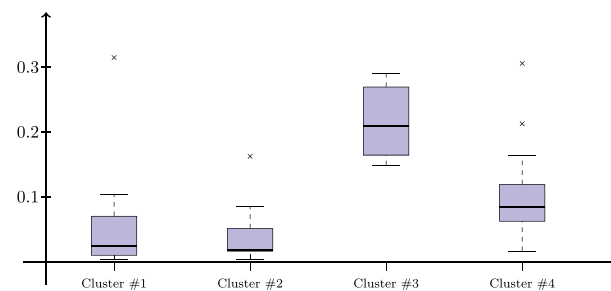


Figure 6. Distribution of inconsistency ratios CR , dataset S4, $k = 4$ clusters, measure D_1 , relabelled pairwise comparison matrices.

5.1. The potential role of cluster analysis in the study of pairwise comparison matrices

Pairwise comparison matrices can be classified in several ways such as by the best alternative, by the ranking of the alternatives, or by their level of inconsistency. However, this procedure usually requires an initial assumption on which the classification is carried out, and the groups are created after a reduction of dimension, which certainly loses some information about the original pairwise comparisons. For instance, inconsistency “measures” the matrices on a one-dimensional scale, hence, many “qualitatively” different matrices might have the same inconsistency. Although this is not necessarily a problem if inconsistency is the focus of the analysis, it can be misleading if dissimilar matrices are grouped together.

On the other hand, cluster analysis is a method of exploratory data analysis that has no underlying idea of the classification, and considers the individual pairwise comparison matrices without any transformation. Naturally, this flexibility makes the result dependent on the measure of similarity (Section 2.2) used. Nonetheless, clustering is able to uncover surprising patterns in a dataset that are difficult to recognise without a particular investigation as Section 4 illustrates:

- A typo (mistakenly reversed comparison) is found in a pairwise comparison matrix (Section 4.1);
- Pairwise comparison matrices can be quite different—that is, belong to different clusters—even if their inconsistencies are close and the best alternatives coincide (Section 4.2);
- A potential source of inconsistency (inappropriate conversion of robust ordinal preferences into cardinal values) is identified (Section 4.4).

Hopefully, our paper will inspire further instructive applications of clustering pairwise comparison matrices.

Table 4. An illustrative contingency table for clusters and gender.

Cluster	Male	Female	Σ
$Cl_1^{(1)}$	7	19	26
$Cl_1^{(2)}$	5	10	15
$Cl_1^{(3)}$	1	8	9
$Cl_1^{(4)}$	3	15	18
Σ	16	52	68

Dataset S4, four clusters, measure D_1 .

5.2. Incorporating demographic factors into the cluster analysis

In many application, further information is available about the DMs, and the clusters might be determined at least partially by demographic variables. The suggested clustering method can help this kind of analysis as well. The dataset studied in Section 4 contains the age category and the gender of the DMs, furthermore, they faced three different filling orders. Table 4 reveals the joint distribution of the four clusters from Section 4.2 and gender. Although the proportion of male and female students is not the same in all clusters, the differences are statistically insignificant. We have found no significant association with the clusters for any other variables.

Naturally, in other applications more information may be available from the DMs, and the cluster analysis can be based on demographic factors. Another possibility could be to define the distance between decision makers by considering both their preferences and others, such as demographics and characteristics; or a bicluster approach where both preferences and demographic values should be similar in any cluster.

6. Conclusions

This paper has provided a new perspective on group decision making, especially large-scale group decision making, by proposing the k -medoids clustering for pairwise comparison matrices. Our method has some advantages over other solutions:

- it is independent of the specification of the weighting method;
- it can handle incomplete data such that the impact of incomplete pairwise comparison matrices, which contain less information, is inherently reduced;
- it is more robust to outliers than the k -means clustering algorithm;
- the cluster centres are guaranteed to be individual pairwise comparison matrices, making them easier to accept by the decision-makers;
- its LP formulation allows for adding various restrictions, for example, regarding the inconsistency of the cluster centres.

The suggested approach has been used to analyse the experimental dataset of Bozóki et al. (2013)

Hopefully, all practitioners dealing with data from large-scale group decision-making may benefit from using the clustering model presented here. Further investigations are especially welcome because, at the moment, there are few results on choosing the dissimilarity measure underlying the k -medoids algorithm or the appropriate number of clusters k .

Note

1. Source: Kaufman and Rousseeuw (1990, p. 71).

Acknowledgements

Four anonymous reviewers gave useful remarks on earlier drafts.

Disclosure statement

No potential conflict of interest was reported by the authors.

Funding

The research was supported by the National Research, Development and Innovation Office under Grants FK 145838 and TKP2021-NKTA-01 NRDI, and the János Bolyai Research Scholarship of the Hungarian Academy of Sciences.

References

- Aczél, J., & Saaty, T. L. (1983). Procedures for synthesizing ratio judgements. *Journal of Mathematical Psychology*, 27(1), 93–102. [https://doi.org/10.1016/0022-2496\(83\)90028-7](https://doi.org/10.1016/0022-2496(83)90028-7)
- Ágoston, K. C., & Csató, L. (2024). A lexicographically optimal completion for pairwise comparison matrices with missing entries. *European Journal of Operational Research*, 314(3), 1078–1086. <https://doi.org/10.1016/j.ejor.2023.10.035>
- Ágoston, K. C., & E.-Nagy, M. (2024). Mixed integer linear programming formulation for K-means clustering problem. *Central European Journal of Operations Research*, 32(1), 11–27. <https://doi.org/10.1007/s10100-023-00881-1>
- Albano, A., García-Lapresta, J. L., Plaia, A., & Sciandra, M. (2024). Clustering alternatives in preference-approvals via novel pseudometrics. *Statistical Methods & Applications*, 33(1), 61–87. <https://doi.org/10.1007/s10260-023-00718-w>
- Amenta, P., Ishizaka, A., Lucadamo, A., Marcarelli, G., & Vyas, V. (2020). Computing a common preference vector in a complex multi-actor and multi-group decision system in Analytic Hierarchy Process context. *Annals of Operations Research*, 284(1), 33–62. <https://doi.org/10.1007/s10479-019-03258-3>
- Basak, I., & Saaty, T. (1993). Group decision making using the analytic hierarchy process. *Mathematical and Computer Modelling*, 17(4–5), 101–109. [https://doi.org/10.1016/0895-7177\(93\)90179-3](https://doi.org/10.1016/0895-7177(93)90179-3)

- Bozóki, S., Dezső, L., Poesz, A., & Temesi, J. (2013). Analysis of pairwise comparison matrices: An empirical research. *Annals of Operations Research*, 211(1), 511–528. <https://doi.org/10.1007/s10479-013-1328-1>
- Bozóki, S., & Rapcsák, T. (2008). On Saaty's and Koczkodaj's inconsistencies of pairwise comparison matrices. *Journal of Global Optimization*, 42(2), 157–175. <https://doi.org/10.1007/s10898-007-9236-z>
- Brunelli, M. (2018). A survey of inconsistency indices for pairwise comparisons. *International Journal of General Systems*, 47(8), 751–771. <https://doi.org/10.1080/03081079.2018.1523156>
- Choo, E. U., & Wedley, W. C. (2004). A common framework for deriving preference values from pairwise comparison matrices. *Computers & Operations Research*, 31(6), 893–908. [https://doi.org/10.1016/S0305-0548\(03\)00042-X](https://doi.org/10.1016/S0305-0548(03)00042-X)
- Crawford, G., & Williams, C. (1985). A note on the analysis of subjective judgment matrices. *Journal of Mathematical Psychology*, 29(4), 387–405. [https://doi.org/10.1016/0022-2496\(85\)90002-1](https://doi.org/10.1016/0022-2496(85)90002-1)
- Csató, L. (2019a). Axiomatizations of inconsistency indices for triads. *Annals of Operations Research*, 280(1–2), 99–110. <https://doi.org/10.1007/s10479-019-03312-0>
- Csató, L. (2019b). A characterization of the Logarithmic Least Squares Method. *European Journal of Operational Research*, 276(1), 212–216. <https://doi.org/10.1016/j.ejor.2018.12.046>
- De Graan, J. G. (1980). *Extensions of the multiple criteria analysis method of T. L. Saaty*. National Institute for Water Supply.
- Duleba, S., & Szádóczi, Z. (2022). Comparing aggregation methods in large-scale group AHP: Time for the shift to distance-based aggregation. *Expert Systems with Applications*, 196, 116667. <https://doi.org/10.1016/j.eswa.2022.116667>
- Fan, S., Liang, H., Pedrycz, W., & Dong, Y. (2024). Integrating incomplete preference estimation and consistency control in consensus reaching. *Information Fusion*, 106, 102268. <https://doi.org/10.1016/j.inffus.2024.102268>
- Fichtner, J. (1984). *Some thoughts about the mathematics of the Analytic Hierarchy Process*. Institut für Angewandte Systemforschung und Operations Research, Universität der Bundeswehr München.
- Fichtner, J. (1986). On deriving priority vectors from matrices of pairwise comparisons. *Socio-Economic Planning Sciences*, 20(6), 341–345. [https://doi.org/10.1016/0038-0121\(86\)90045-5](https://doi.org/10.1016/0038-0121(86)90045-5)
- García-Lapresta, J. L., & Pérez-Román, D. (2016). Consensus-based clustering under hesitant qualitative assessments. *Fuzzy Sets and Systems*, 292, 261–273. <https://doi.org/10.1016/j.fss.2014.05.004>
- García-Zamora, D., Labella, A., Ding, W., Rodríguez, R. M., & Martínez, L. (2022). Large-scale group decision making: A systematic review and a critical analysis. *IEEE/CAA Journal of Automatica Sinica*, 9(6), 949–966. <https://doi.org/10.1109/JAS.2022.105617>
- Kamis, N. H., Chiclana, F., & Levesley, J. (2018). Preference similarity network structural equivalence clustering based consensus group decision making model. *Applied Soft Computing*, 67, 706–720. <https://doi.org/10.1016/j.asoc.2017.11.022>
- Kaufman, L., & Rousseeuw, P. J. (1990). *Finding Groups in Data: An Introduction to Cluster Analysis*. Wiley Series in Probability and Statistics. John Wiley & Sons.
- Kaufman, L., & Rousseeuw, P. J. (1986). Clustering large data sets. In *Pattern Recognition in Practice II*, edited by E. S. Gelsema and L. N. Kanal, North Holland, 425–437. Elsevier.
- Kruskal, J. B. (1964). Multidimensional scaling by optimizing goodness of fit to a nonmetric hypothesis. *Psychometrika*, 29(1), 1–27. <https://doi.org/10.1007/BF02289565>
- Kruskal, J. B., & Wish, M. (1978). *Multidimensional scaling*. Sage Publications.
- Kuřakowski, K., Mazurek, J., & Strada, M. (2022). On the similarity between ranking vectors in the pairwise comparison method. *Journal of the Operational Research Society*, 73(9), 2080–2089. <https://doi.org/10.1080/01605682.2021.1947754>
- Majstorović, S., Sabo, K., Jung, J., & Klarić, M. (2018). Spectral methods for growth curve clustering. *Central European Journal of Operations Research*, 26(3), 715–737. <https://doi.org/10.1007/s10100-017-0515-6>
- Meng, F., Tang, J., & An, Q. (2023). Cooperative game based two-stage consensus adjustment mechanism for large-scale group decision making. *Omega*, 117, 102842. <https://doi.org/10.1016/j.omega.2023.102842>
- Ossadnik, W., Schinke, S., & Kaspar, R. H. (2016). Group aggregation techniques for Analytic Hierarchy Process and Analytic Network Process: A comparative analysis. *Group Decision and Negotiation*, 25(2), 421–457. <https://doi.org/10.1007/s10726-015-9448-4>
- Rabinowitz, G. (1976). Some comments on measuring world influence. *Conflict Management and Peace Science*, 2(1), 49–55. <https://doi.org/10.1177/073889427600200104>
- Saaty, T. L. (1977). A scaling method for priorities in hierarchical structures. *Journal of Mathematical Psychology*, 15(3), 234–281. [https://doi.org/10.1016/0022-2496\(77\)90033-5](https://doi.org/10.1016/0022-2496(77)90033-5)
- Saaty, T. L. (1980). *The Analytic Hierarchy Process: Planning, Priority Setting, Resource Allocation*. McGraw-Hill.
- Saaty, T. L. (2008). “Relative measurement and its generalization in decision making. Why pairwise comparisons are central in mathematics for the measurement of intangible factors. The Analytic Hierarchy/Network Process.” *RACSAM—Revista de la Real Academia de Ciencias Exactas. Físicas y Naturales. Serie A. Matemáticas*, 102(2), 251–318.
- Schubert, E., & Rousseeuw, P. J. (2019). *Faster k-medoids clustering: Improving the PAM, CLARA, and CLARANS algorithms* [Paper presentation]. In *Similarity Search and Applications: 12th International Conference, SISAP 2019, Newark, NJ, USA, October 2–4, 2019, Proceedings*, edited by G. Amato, C. Gennaro, V. Oria, and M. Radovanović, Lecture Notes in Computer Science, Cham, Switzerland, 171–187. Springer.
- Tang, M., & Liao, H. (2021). From conventional group decision making to large-scale group decision making: What are the challenges and how to meet them in big data era? A state-of-the-art survey. *Omega*, 100, 102141. <https://doi.org/10.1016/j.omega.2019.102141>
- Tekile, H. A., Brunelli, M., & Fedrizzi, M. (2023). A numerical comparative study of completion methods for pairwise comparison matrices. *Operations Research Perspectives*, 10, 100272. <https://doi.org/10.1016/j.orp.2023.100272>
- Tibshirani, R., Walther, G., & Hastie, T. (2001). Estimating the number of clusters in a data set via the gap statistic. *Journal of the Royal Statistical Society*

- Series B: Statistical Methodology*, 63(2), 411–423. <https://doi.org/10.1111/1467-9868.00293>
- Vinod, H. D. (1969). Integer programming and the theory of grouping. *Journal of the American Statistical Association*, 64(326), 506–519. <https://doi.org/10.1080/01621459.1969.10500990>
- Wu, Z., & Xu, J. (2018). A consensus model for large-scale group decision making with hesitant fuzzy information and changeable clusters. *Information Fusion*, 41, 217–231. <https://doi.org/10.1016/j.inffus.2017.09.011>
- Xu, X., Zhong, X., Chen, X., & Zhou, Y. (2015). A dynamical consensus method based on exit-delegation mechanism for large group emergency decision making. *Knowledge-Based Systems*, 86, 237–249. <https://doi.org/10.1016/j.knosys.2015.06.006>
- Zhang, H., Dong, Y., & Herrera-Viedma, E. (2018). Consensus building for the heterogeneous large-scale GDM with the individual concerns and satisfactions. *IEEE Transactions on Fuzzy Systems*, 26(2), 884–898. <https://doi.org/10.1109/TFUZZ.2017.2697403>