

How to Interpret AHP/ANP Application Results in a Really Meaningful Manner?

Grzegorz Ginda*

Abstract. Final decision recommendations rely heavily on ranking Decision-Making Units (DMUs), often achieved using Saaty's Analytic Hierarchy/Network Process (AHP/ANP). AHP/ANP provides precise overall priority scores which decision-makers commonly treat as definitive for ranking purposes. This reliance means that even minimal numerical differences between DMUs are used to determine the final selection. However, this strict adherence to tiny numerical distinctions – disregarding the actual degree of difference – is problematic. Practically, it risks rejecting DMUs whose performance is only slightly inferior; methodologically, it contradicts the qualitative nature of the input (pairwise comparisons) with the quantitative output. This tension raises the question of achieving an adequate qualitative interpretation of the quantitative rankings. To resolve this, the paper proposes clustering approaches to help decision-makers reliably group and discriminate among similar DMUs. These methods aim to justify more informed choices by avoiding spurious precision. The approaches were tested using two diverse decision cases. The results are promising and indicate that these clustering techniques can be useful under certain specific circumstances.

Keywords: decision analysis, AHP/ANP, results, interpretation, compatibility, qualitative

Mathematics Subject Classification: 68U35

JEL Classification: C44

Submitted: October 1, 2021

Revised: November 18, 2021

© 2021 Author. This is an open-access publication that can be used, distributed, and reproduced in any medium according to the Creative Commons CC-BY 4.0 license. This license requires that the original work was properly cited.

1. INTRODUCTION

The application of the multi-attribute decision analysis (MADA) methodology provides decision makers with recommendations that facilitate actual decision making. The methodology is implemented by means of diverse techniques. The application of

* AGH University of Science and Technology, Faculty of Management, Krakow, Poland, e-mail: gginda@agh.edu.pl

the techniques results in specific outcomes. One of the most popular forms of such outcomes is the ranking of considered alternatives (decision making units – DMUs). This is because a ranking represents an easily interpretable hierarchy of DMUs.

Analytic Hierarchy/Network Process (AHP/ANP) is one of the most popular MADA ranking techniques. It was initially developed by Saaty (1980) in the 1970s as the Analytic Hierarchy Process (AHP), and it was later extended to a more general approach, namely the Analytic Network Process (Saaty, 1996), in the 1990s. Both aforementioned flavors of the technique are based on the application of the same notion of pair-wise comparison. However, they differ in the actual character of fundamental inquiry regarding the differences between the compared DMUs. It is questionable how much two DMUs differ concerning their importance for actually achieving the fundamental analysis goal in AHP, while it is also queried how much they differ according to their influence on achieving that goal in ANP. The flavors also differ in applicable forms of relations between components (and component groups) in a model of a decision making problem.

The application of pair-wise comparison makes AHP/ANP capable of considering both qualitative and quantitative DMU attributes. The 9-point Saaty's scale is applied to provide the necessary means for the qualitative assessment of different DMUs being compared in a pair-wise manner. Subsequent integer numbers from 1 to 9 are utilized to express successive scale levels. A lack of difference in the compared DMUs with regard to their importance/influence corresponds with number 1. The following odd numbers: 3, 5, 7, and 9 express a gradual rise in the assessment of the difference – from a slight difference to an extreme difference. The even numbers: 2, 4, 6, and 8 are used to consider the hesitation regarding which adjacent scale level to choose, e.g., 6 instead of 5 or 7. Note that Saaty's scale levels from 2 to 9 directly conform to the advantage of the first compared DMU. Adequate reciprocal values, i.e., $1/2$, $1/3$, $1/4$, $1/5$, $1/6$, $1/7$, $1/8$, and $1/9$, are applied in the case of a need to express the advantage of the second compared DMU.

It is obvious, therefore, that Saaty's scale application provides the AHP/ANP procedure with qualitative assessments that are expressed by numbers. The numbers are used directly to derive the overall priorities of DMUs and to construct the final hierarchy of DMUs. Unfortunately, despite the generally qualitative character of the input data provided by pair-wise comparisons, the outcomes of AHP/ANP technique application are commonly analyzed in a strictly numerical way. Therefore, this interpretation of the outcomes of AHP/ANP application seems to be incompatible with the actual qualitative nature of the technique.

It ultimately seems that the grouping of DMUs according to their overall priority is particularly well suited to provide suitable means for an adequate and meaningful interpretation of AHP/ANP application results. This is because it would be capable of both distilling sure top DMUs and enriching AHP/ANP use by means of the identification of close DMUs.

Note that there is an AHP-based approach, AHP Sort (Ishizaka et al., 2012), available to group DMUs. The approach also makes use of overall priorities to accomplish this. However, it requires the user to provide information about the subjective limiting profiles of DMUs classes, which are ultimately applied to group DMUs. It nevertheless

seems that, in general, the application of predefined classes is not really necessary. This is because the mere application of overall priorities seems to provide all the necessary means to group DMUs while avoiding the influence of superfluous subjectivity.

The rest of the paper is structured in the following way, therefore. The second section is devoted to the discussion of selected approaches to sorting DMUs. The effects of their sample applications are discussed in the third section. The last section is devoted to some conclusions and final remarks.

2. GROUPING AHP/ANP RESULTS

AHP/ANP provides contextual results as vectors of overall DMU priorities. One-dimensional grouping could be applied, therefore, to divide DMUs that are close enough in terms of their overall priorities into common clusters. There are diverse one-dimensional tools available and some of them are presented in the following subsections.

2.1. Clustering tools

Clustering tools are capable of dividing objects that are described by several attributes into groups – clusters. Individual objects are assigned to the clusters that consist of the most similar objects to them. A concrete metric is applied to express the similarity of the objects. The metric is usually based on a concept of distance between points that represent DMUs in a multi-dimensional space of attributes. The smaller the distance between them, the more similar the corresponding DMUs are. Criteria of diverse kinds are applied when assigning objects to clusters. The criteria may operate on different concepts of cluster similarity provided by appropriate definitions of the distance between clusters. It is also possible to use measures involving the application of statistical concepts, such as extreme distance and the standard deviation of the distance between cluster components, etc. As a result, diverse procedures are available to group objects.

The simple gradient technique is based on the concept of the partition approach. Therefore, its procedure starts with one large cluster that contains all DMUs. It is then divided in a step by-step manner. The actual division deals with a sequence of overall priorities that is gradually partitioned at points of the currently largest difference between the overall priorities. Unfortunately, the stopping criterion depends on somewhat subjective information. A predefined number of distinct clusters (k) or a threshold of absolute difference (ϑ) between overall priorities may be applied in this regard. Note that the simple gradient procedure follows the idea of partitioning a minimum spanning tree (Florek et al., 1951), which is directly expressed by the overall ranking of DMUs.

The Szczotka–Spaeth (Spaeth, 1973; Szczotka, 1972) technique represents another concept: the aggregative approach. Its procedure starts with the initial division of DMUs into n clusters, where n is the number of DMUs. Hence, each initial cluster consists of precisely one DMU. The technique also belongs to optimization approaches, as it uses a goal function as a clustering criterion. The goal function addresses the

minimization of the sum of the average distances between the components of distinct clusters. It looks as follows:

$$Q = \sum_{i=1}^k \frac{2}{n_i(n_i - 1)} \sum_{s,t=1}^{n_i} d_{st} \quad (1)$$

where n_i is the cardinality of the i -th consecutive DMU cluster, and d_{st} is distance between different DMUs, denoted by indices s, t ($s \neq t$), which are components of the i -th consecutive cluster.

Initial clusters are then gradually aggregated to compose more complex clusters in a step by-step manner. The effects of the integration of adjoining DMU clusters on the goal function value (1) is taken into account during each step. Hence, in each step, two adjoining clusters are selected to join. The actual choice of joined clusters corresponds to the smallest increase in the goal function (1) caused by cluster integration. As in the case of the simple gradient technique, the stopping criterion deals with achieving a DMU partition that contains a predefined number of k clusters. However, unlike the simple gradient approach, the technique belongs to aggregation methods. Note that the application of the Szczotka-Spaeth technique involves rather laborious calculations that require software support.

Both above mentioned techniques belong to hierarchical approaches. They are capable, therefore, of providing a cluster hierarchy whose levels define possible DMU partitions without the need to provide any stopping criterion. This is because their procedure stops just after the full cluster hierarchy is derived. Therefore, they are applied in such a way in the paper.

2.2. Other selected approaches

As can be seen from the previous subsection, the applicability of a clustering approach to identifying DMU groups may depend on additional subjective information or relatively complex and time consuming calculations. It seems purposeful, therefore, to try to identify a simple yet reliable one-dimensional clustering approach for AHP/ANP application results.

Let us first try to take advantage of the common 80/20 Pareto principle. It could be readily applied in the case of using AHP/ANP ideals in place of the raw overall DMU priorities. The ideal result from the transformation of overall priorities is that the highest DMU overall priority becomes equal to one. The priorities of other DMUs are then recalculated accordingly. The application of the Pareto rule would allow us to identify the closest DMU, i.e., the ones whose current ideals are contained within a 20% margin from the ideal for the top DMU, i.e., in the $[0.8, 1]$ interval. As a result, all such DMUs would also be regarded as the topmost ones. They would also be excluded from further analysis due to the reduction of the DMUs set to obtain the set of currently active DMUs. Further analysis would involve the step-wise identification of top DMU from the currently active DMUs, as well as a step-wise reduction of the active DMU set until it becomes empty. Hence, Pareto rule-based procedure would

finish without any need for subjective user's intervention. Note that the ideals for the currently active DMU must be updated in the beginning of each step to make the ideal for the top active DMU equal to one.

The Pareto rule exploits the constant threshold value $\vartheta = 0.80$. Note that it could nevertheless matter for the sake of the reliability of active DMU discrimination whether the threshold is constant or not. For example, Opricovic (1988) proposed using a threshold-based rule related to the actual number of DMU set to provide evidence of a significant and necessary advantage of a DMU over another DMU in the VIKOR technique (Opricovic & Tzeng, 2004). The evidence is based on the following ideal advantage threshold:

$$\varsigma = 1/(N - 1) \quad (2)$$

where: N denotes number of DMUs.

The threshold could be used in the same way as a 20% margin threshold in the case of Pareto rule application, with N denoting the number of current active DMUs. Therefore, a certain DMUs would prove to be very close to the topmost one if its current ideal were at least equal to:

$$\vartheta = 1 - \varsigma \quad (3)$$

The application of ϑ causes both rule-based procedures to stop immediately after the set of currently active DMUs becomes empty. However, the subsequent clusters are identified in a sequence from the best down to the worst DMUs. Hence, in the case the worst DMU is too far from the adjoining DMUs, what would imply $N = 1$ during the last procedure step, no further proceeding is actually needed, as it is obvious that the DMU would comprise a distinct cluster.

2.3. Partition validity

To facilitate the comparison of the effects of applying different DMU grouping approaches, additional methods for assessing the obtained DMU partitions should be utilized. There are many partition quality assessment indices available (Kolenda, 2006). The most popular means for partition quality checking is provided, however, by the silhouette coefficient $s(i)$ (Kaufman & Rousseeuw, 1990). t is a type of higher is better index. Its value for the i -th consecutive DMU would be described by the following equation:

$$s(i) = (b_i - a_i) / \max(a_i, b_i), \quad (4)$$

where: a_i denotes average distance between the i -th consecutive DMU and other DMUs from the same cluster; b_i is the smallest average distance of the i -th consecutive DMU from DMUs in other clusters i.e. the average distance from the closest cluster (note that average linkage is applied with this regard in the paper, although any of possible linkage types maybe applied).

Silhouette coefficient values belong to the interval $[-1, +1]$. Its negative values would indicate that the DMU does not fit a given cluster. The better the i -th

consecutive DMU fits a given cluster, the closer the coefficient $s(i)$ value is to +1. Note that the coefficient equals zero for a DMU that comprises a cluster itself. For the sake of simplicity, it is nevertheless assumed in the paper that the existence of single DMU clusters does not influence the quality of a partition at all.

It is generally accepted that $s(i)$ values above 0.71 indicate very high partition quality, and levels above 0.51 signify good quality for the partition from the perspective of the i -th consecutive object. The latter value is therefore ultimately associated in the paper with a notion of acceptable fit for the i -th consecutive DMU.

3. SAMPLE ANALYSIS

3.1. Sample data

The results of two sample AHP/ANP applications are used to present the potential merits and drawbacks of the aforementioned grouping approaches. The applications differ in the number of DMUs, making it possible to examine the effects of different conditions. The first case deals with a rather mean number of DMUs. It pertains to the ranking of sustainability programs of six Brazilian textile industry companies. The programs are denoted by the symbols E1–E6. Overall priorities for the programs are given in Table 1.

Table 1. Overall priorities of sustainability programs in textile industry (Netto et al., 2021)

Sustainability program	E2	E5	E3	E1	E6	E4
Overall priority	0.64	0.52	0.38	0.36	0.28	0.26
Ideal	1	0.813	0.594	0.563	0.438	0.406
Rank	1	2	3	4	5	6

The second case pertains to the results of a recent AHP application for ranking 16 Polish voivodships concerning their biogas technology potential (Ginda & Szyba, 2020). Overall priorities and ideals obtained for the voivodeships are presented in Table 2.

Table 2. The results of AHP application (Ginda & Szyba, 2020)

Rank	DMU	Overall priority	Ideal	Rank	DMU	Overall priority	Ideal
1	B	0.1820	1	2	N	0.1701	0.9346
3	W	0.1648	0.9054	4	L	0.1623	0.8917
5	Z	0.1437	0.7895	6	F	0.1384	0.7604
7	P	0.1019	0.5598	8	T	0.0828	0.4549
9	O	0.0797	0.4379	10	E	0.0681	0.3741
11	C	0.0673	0.3697	12	S	0.0647	0.3554
13	R	0.0624	0.3428	13	G	0.0624	0.3428
15	D	0.0615	0.3379	16	K	0.0524	0.2879

3.2. The case of six DMUs

Let us use information about priority difference ranks in a simple gradient approach. We start from the establishment of a global cluster of all DMUs: {E2 E5 E3 E1 E4 E6}. The largest priority difference is then applied to define the initial partition point. It is clear from Table 1 that the largest priority difference corresponds to the gap between E5 and E3 sustainability programs. We end up, therefore, with two second level clusters in the first step. The first cluster consists of E2 and E5 programs while the second one contains the remaining programs. The application of next priority differences in descending order allows us to complete levels of the cluster hierarchy. The hierarchy is presented in Table 3. We can see that there are 6 final clusters in the bottommost hierarchy level. Three intermediate hierarchy levels correspond to the partitions with two, three, and four distinct clusters, respectively.

Table 3. Cluster level hierarchy for simple gradient technique application

Cluster hierarchy level	Partition
Top	{E2 E5 E3 E1 E6 E4}
sg26	{E2 E5}={E3 E1 E6 E4}
sg36	{E2}={E5}={E3 E1 E6 E4}
sg46	{E2}={E5}={E3 E1}={E6 E4}
Bottom	{E2}={E5}={E3}={E1}={E6}={E4}

Now is the time to make a decision about which partition expresses the best division of sustainability programs. For example, we could prefer having no more than 4 clusters or aim to achieve a similarity measure at a level not higher than an average priority difference (0.76 in this case). It proves, therefore, that we could finally be happy with the sg46 partition given by the fourth cluster level. The partition deals with the division of sustainability programs into four distinct clusters. Two of them consist of a single program (E2 or E5), and the two remaining clusters consist of two programs each: {E3 E1}, and {E6 E5}. Let's take a look at the $s(i)$ values to ensure the quality of the partitions from Table 3.

Individual silhouette coefficient values (see Table 4) obtained for sg26 support the following conclusions:

1. E2 sustainable strategy would fit rather well with a common cluster shared with E5 sustainable strategy. However, a rather low $s(i)$ coefficient value for E5 strategy suggests that it doesn't fit the cluster well. Therefore, it finally turns out that both aforementioned strategies should comprise distinct clusters in their own right.
2. Coefficient values for the remaining strategies suggest that they either fit very well (E6, E4) or fit well (E1, E3) within their common cluster.

We can also see a considerable drop in $s(i)$ values for the worst four strategies after the two top strategies were separated into two distinct clusters (see sg36 partition in Table 3). Such a decrease in the coefficient values suggests that the

E3, E1, E6, and E4 strategies no longer fit the common cluster well. The considerable increase in the coefficient following their division into two separate clusters confirms the superiority of the sg46 partition in the case of the simple gradient approach application.

Table 4. *Silhouette coefficient values $s(i)$ for the partitions – the simple gradient approach application case*

Partition	E2	E5	E3	E1	E6	E4
sg26	0.625	0.4	0.6	0.697	0.778	0.75
sg36	–	–	0.429	0.583	0.722	0.692
sg46	–	–	0.818	0.778	0.778	0.818

Let us now apply the Szczotka–Spaeth technique. The anticipated effects of grouping the adjoining distinct sustainability strategies are tested according to the goal function $Q(1)$, and a new cluster that minimizes the current goal function as much as possible is recorded in each step. Hence, cluster hierarchy emerges in a step-wise manner. The results of consecutive steps are presented in Tables 5–8. Note that the recommended clusters are expressed there by means of boldface. The calculations result in a cluster hierarchy, one which is finally presented in Table 9.

The comparison of cluster hierarchies obtained through the application of the simple gradient approach (Table 3) and the application of the Szczotka–Spaeth technique (Table 9) reveals a difference. The difference is associated with two unique partitions. The first one (ss56) corresponds with the second cluster hierarchy level and consists of one cluster, which comprises two components {E3 E1}, as well as four clusters that contain a single component each: {E2}, {E5}, {E6}, and {E4}. The second partition (ss26) is presented in the fifth cluster hierarchy level and consists of a single component cluster {E2} and a cluster that is composed of the remaining sustainability programs. We finally use silhouette coefficient values again to justify the partition – see Table 10 for details.

Table 5. *The effects of possible joining of adjoining sustainability programs after the initial step of Szczotka–Spaeth technique*

Possible new cluster	{E2 E5}	{E5 E3}	{E3 E1}	{E1 E6}	{E6 E4}
$Q(1)$	0.12	0.14	0.02	0.08	0.02

Table 6. *The effects of possible joining of adjoining sustainability programs after the second step of Szczotka–Spaeth technique*

Possible new cluster	{E2 E5}	{E5 E3 E1}	{E3 E1 E6}	{E6 E4}
$Q(1)$	0.14	0.126	0.087	0.04

Table 7. The effects of possible joining of adjoining sustainability programs after the third step of Szczotka–Spaeth technique

Possible new cluster	{E2 E5}	{E5 E3 E1}	{E3 E1 E6 E4}
Q (1)	0.16	0.127	0.073

Table 8. The effects of possible joining of adjoining sustainability programs after the fourth step of Szczotka–Spaeth technique

Possible new cluster	{E2 E5}	{E5 E3 E1 E6 E4}
Q (1)	0.193	0.124

Table 9. Cluster hierarchy levels – the application of Szczotka–Spaeth technique

Cluster hierarchy level	Partition
Top	{E2 E5 E3 E1 E6 E4}
ss26	{E2}={E5 E3 E1 E6 E4}
ss36=sg36	{E2}={E5}={E3 E1 E6 E4}
ss46=sg46	{E2}={E5}={E3 E1}={E6 E4}
ss56	{E2}={E5}={E3 E1}={E6 E4}
Bottom	{E2}={E5}={E3}={E1}={E6}={E4}

Table 10. Silhouette coefficient values $s(i)$ for the unique partitions according to Szczotka–Spaeth technique

Cluster hierarchy level	E2	E5	E3	E1	E6	E4
ss26	–	–	0.8	0.75	–	–
ss56	–	–0.4	0.635	0.679	0.694	0.671

Hence, it is clear that the use of the partition from the fifth cluster hierarchy level ss56 is inefficient due to the negative coefficient value for the E5 sustainability program, which indicates a total mismatch for the program. It also turns out that the application of a unique partition ss26 would result in lower coefficient values for E3 and E1 sustainability programs than those obtained for partition sg46 in the case of the simple gradient approach – see Tables 3 and 4. It seems, therefore, that it should be rejected in favor of the third cluster hierarchy level partition presented in Table 9, which is identical to the sg46 partition from Table 3. Hence, his final partition recommendation is the same as in the case of using the simple gradient approach.

The use of the Pareto-based rule is grounded in the application of ideals obtained for the sustainability programs (see Table 1). The results of the conducted calculations are illustrated in Table 11. We start from core ideals and the fully active DMU set. The second best DMUs (E5) has an ideal within a 2% margin of the topmost DMU, while the ideals of other DMUs are outside the margin. Hence, the topmost cluster

consists of two top sustainability programs, E2 and E5. We remove them from the set of active DMUs and proceed with the remaining four sustainability programs. The ideal for the best of the remaining programs is then scaled to 1, and the ideals of the other active sustainability programs are recalculated accordingly. The current ideal for the second most active DMU (E1) fits a 20% margin from the current top active DMU (E3). Hence, the second topmost cluster consists of two sustainability programs: E3 and E1. The same conclusion pertains to the bottommost cluster, which consists of the remaining DMUs: E6 and E4. Hence, the application of the Pareto-based rule for clustering results in a unique set of three double sized clusters. Individual silhouette coefficient values for the partition are presented in Table 12.

Table 11. *Pareto rule-based approach illustration ($\vartheta = 0.80$) for the six DMUs case*

Step I		Step II		Step III	
Sustainability program	Ideal	Sustainability program	Ideal	Sustainability program	Ideal
E2	1	—	—	—	—
E5	0.81	—	—	—	—
E3	0.59	E3	1	—	—
E1	0.56	E1	0.94	—	—
E6	0.43	E6	0.73	E6	1
E4	0.40	E4	0.68	E4	0.92

Table 12. *Silhouette coefficient values $s(i)$ for the unique partition $\{E2\ E5\}\ \{E3\ E1\}\ \{E6\ E4\}$*

E2	E5	E3	E1	E6	E4
0.555	0.2	0.8	0.6	0.778	0.818

The contents of Table 12 confirm that the final partition provided by the application of the Pareto rule-based approach is slightly better than some other partitions obtained through the application of different clustering approaches (e.g., the sg26 partition in Table 3) from the perspective of clearly less preferred sustainability programs E3, E1, E6, and E4. It is, nevertheless, unacceptable because the sustainability program E5 does not fit the partition at all.

Let us finally use a VIKOR-like marginal approach. The results of the calculations are illustrated in Table 13. The first step deals with all sustainable programs and their original ideals again. This time, however, the margin from the topmost DMU depends on the current cardinality of the initial active DMU set. The cardinality is equal to $N = 6$. Hence, according to (2) the initial threshold yields: $\vartheta = 0.80$. The application of the threshold results in the same topmost cluster structure as in the case of applying the Pareto-rule $\{E2\ E5\}$. The second step deals with a reduced set of four active DMUs. The current margin from the topmost active DMU is therefore calculated for $N = 4$, therefore. It yields $\vartheta = 0.67$, now. The ideals of all four active DMUs qualify

them for the second highest cluster $\{E3\ E1\ E6\ E4\}$. The procedure stops, therefore, all DMUs have been distributed among clusters. Note that the obtained partition is the same as the inefficient sg26 partition in the second cluster hierarchy level (Table 3).

Table 13. *VIKOR-like rule-based approach illustration*

Step I		Step II	
N	ϑ	N	ϑ
6	0.80	4	0.67
Sustainability program	Ideal	Sustainability program	Ideal
E2	1	E2	–
E5	0.81	E5	–
E3	0.59	E3	1
E1	0.56	E1	0.94
E6	0.43	E6	0.73
E4	0.40	E4	0.68

The best partitions identified by means of the core application of all used approaches are presented in Table 14. Partitions delivered by the Pareto rule and VIKOR-like rule are unacceptable from the E5 sustainability program's point of view. It seems, therefore, that only clustering approaches may be capable of providing partitions in which all the sustainability programs fit well.

Table 14. *The best outcomes for use of applied approaches for the six DMUs case*

Approach	Partition	E2	E5	E3	E1	E6	E4
		$s(i)$					
Simple gradient	{E2}={E5}={E3 E1}={E6 E4}	–	–	0.82	0.78	0.78	0.82
Szczotka–Spaeth							
Pareto rule	$\{E2\ E5\}=\{E3\ E1\}=\{E6+E4\}$	0.56	0.20	0.80	0.60	0.78	0.82
VIKOR-like rule	$\{E2\ E5\}=\{E3\ E1\ E6\ E4\}$	0.63	0.40	0.60	0.70	0.78	0.75

3.3. The case of sixteen DMUs

Cluster hierarchy levels that result from the application of a simple gradient technique in the case of biogas potential analysis are presented in Table 15. A statistical summary of the corresponding silhouette coefficient values $s(i)$ for all meaningful partitions is presented in Table 16.

Note that the partitions that make up the highest levels in the cluster hierarchy (sg2–sg7) correspond to minimal values of silhouette coefficients, which testify that some sustainability programs do not fit the partitions well. These are the eighth and ninth cluster hierarchy levels that define partitions (sg8, sg9) which guaranty a good

fit for all sustainability programs. The advantage in mean and maximum silhouette coefficient values makes sg8 the final recommendation for the partition; however, Note that excellent silhouette coefficient values are obtained in the case of partitions from the twelfth and thirteenth cluster hierarchy levels. These partitions, nevertheless, seem to be unsuitable because they are extremely fragmented.

Table 15. *Cluster hierarchy – the simple gradient technique application case for sixteen DMUs case*

Name	Partition
Top	B+N+W+L+Z+F+P+T+O+E+C+S+G+R+D+K
sg2	B+N+W+L+Z+F P+T+O+E+C+S+G+R+D+K
sg3	B+N+W+L+Z+F P T+O+E+C+S+G+R+D+K
sg4	B+N+W+L Z+F P T+O+E+C+S+G+R+D+K
sg5	B N+W+L Z+F P T+O+E+C+S+G+R+D+K
sg6	B N+W+L Z+F P T+O E+C+S+G+R+D+K
sg7	B N+W+L Z+F P T+O E+C+S+G+R+D K
sg8	B N W+L Z F P T+O E+C+S+G+R+D K
sg9	B N W+L Z F P T O E+C+S+G+R+D K
sg10	B N W+L Z F P T O E+C S+G+R+D K
sg11	B N W L Z F P T O E+C S+G+R+D K
sg12	B N W L Z F P T O E+C S G+R+D K
sg13	B N W L Z F P T O E C S G+R+D K
Bottom	B N W L Z F P T O E C S G+R D K

Table 16. *Silhouette coefficient statistics for unique partitions for simple gradient technique application case for sixteen DMUs case*

Name	Number of clusters	min $s(i)$ (DMU)	mean $s(i)$	max $s(i)$ (DMU)	std.dev. $s(i)$ (DMU)
sg2	2	0.398 (P)	0.793	0.888 (D)	0.126
sg3	3	0.058 (T)	0.634	0.805 (R)	0.229
sg4	4	0.058 (T)	0.661	0.831 (F)	0.210
sg5	5	0.058 (T)	0.665	0.806 (F)	0.221
sg6	6	0.450 (N)	0.712	0.838 (T)	0.118
sg7	7	0.450 (N)	0.722	0.832 (T)	0.096
sg8	9	0.528 (W)	0.707	0.832 (T)	0.085
sg9	10	0.528 (W)	0.670	0.776 (S)	0.074
sg10	11	0.133 (S)	0.668	0.850 (E)	0.224
sg11	12	0.133 (S)	0.689	0.850 (S)	0.251
sg12	13	0.719 (D)	0.776	0.804 (G,R)	0.040
sg13	14	0.719 (D)	0.776	0.804 (G,R)	0.040

Table 17 presents the cluster hierarchy obtained through the application of the Szczotka–Spaeth technique. The technique provides 5 partitions, which also appear in the simple gradient approach to the use-related cluster hierarchy. They include the most notable ones with regard to individual silhouette coefficient values as well. However, the majority of partitions provided by the technique are unique. The unique partitions are expressed in Table 17 by boldface. Silhouette coefficient statistics for them are given in Table 18.

Table 17. Cluster hierarchy – Szczotka–Spaeth technique application case for sixteen DMUs case

Name	Partition
Top	B+N+W+L+Z+F+P+T+O+E+C+S+G+R+D+K
sg2	B+N+W+L+Z+F P+T+O+E+C+S+G+R+D+K
ss3	B+N+W+L Z+F P+T+O+E+C+S+G+R+D+K
ss4	B N+W+L Z+F P+T+O+E+C+S+G+R+D+K
ss5	B N+W+L Z F P+T+O+E+C+S+G+R+D+K
ss6	B N+W+L Z F P T+O+E+C+S+G+R+D+K
ss7	B N+W+L Z F P T O+E+C+S+G+R+D+K
ss8	B N+W+L Z F P T O E+C+S+G+R+D+K
ss9	B N W+L Z F P T O E+C+S+G+R+D+K
ss10	B N W L Z F P T O E+C+S+G+R+D+K
ss11	B N W L Z F P T O E+C+S+G+R+D K
sg11	B N W L Z F P T O E+C S+G+R+D K
sg12	B N W L Z F P T O E+C S G+R+D K
sg13	B N W L Z F P T O E C S G+R+D K
Bottom	B N W L Z F P T O E C S G+R D K

Table 18. Silhouette coefficient statistics for unique partitions for Szczotka–Spaeth technique application case for sixteen DMUs case

Name	Number of clusters	min $s(i)$ (DMU)	mean $s(i)$	max $s(i)$ (DMU)	std.dev. $s(i)$ (DMU)
ss3	3	0.104 (P)	0.723	0.860 (S)	0.188
ss4	4	0.104 (P)	0.731	0.860 (S)	0.196
ss5	5	0.039 (P)	0.712	0.855 (S)	0.223
ss6	6	0.058 (T)	0.645	0.805 (G, R)	0.232
ss7	7	−0.818 (O)	0.505	0.773 (W)	0.430
ss8	8	0.450 (E)	0.651	0.773 (W)	0.127
ss9	9	0.455 (E)	0.642	0.771 (G, R)	0.114
ss10	10	0.455 (E)	0.653	0.771 (G,R)	0.120
ss11	11	0.617 (E)	0.692	0.776 (S)	0.058

It appears that almost all unique partitions derived from the application of the Szczotka–Spaeth technique deal with rather unacceptable silhouette coefficient values for individual DMUs. The only unique partition that provides acceptable values for the coefficient consists of a single cluster {E C S G R D} and 10 clusters made up of distinct DMUs. Coefficient values for individual DMUs seem, nevertheless, to suggest that the partition is slightly worse than the best partitions provided by the gradient approach application (see Table 16).

The sequence of the Pareto rule-based procedure steps is illustrated in Table 19. The boldfaced ideals in the table correspond to the current cluster composition. Note that the application of the procedure finally gives the partition, which consists of 5 clusters: {B N W L}, {Z F}, {P T}, {O E C S}, and {G R D K}. The silhouette coefficient values for individual DMUs, which are presented in the last column of Table 19, show, however, that the partition is inefficient (see the boldfaced entries in the last column). This is because some DMUs generally misfit the partition (T, S, O), while others show a rather poor fit (C, E, P, K).

Table 19. Pareto rule-based approach use illustration ($\vartheta = 0.80$) for the sixteen DMUs case

DMU	I	II	III	IV	V	$s(\hat{i})$
B	1	–	–	–	–	0.603
N	0.934	–	–	–	–	0.713
W	0.905	–	–	–	–	0.649
L	0.891	–	–	–	–	0.529
Z	0.789	1	–	–	–	0.797
F	0.760	0.963	–	–	–	0.831
P	0.559	0.709	1	–	–	0.402
T	0.454	0.576	0.812	–	–	–0.327
O	0.437	0.554	0.782	1	–	–0.027
E	0.374	0.473	0.668	0.854	–	0.375
C	0.369	0.468	0.660	0.844	–	0.309
S	0.355	0.450	0.634	0.811	–	–0.282
G	0.342	0.434	0.612	0.782	1	0.519
R	0.342	0.434	0.612	0.782	1	0.519
D	0.337	0.428	0.603	0.771	0.985	0.570
K	0.287	0.364	0.514	0.657	0.839	0.447

The sequence of steps for the VIKOR-like rule-based application is presented in Table 20. The obtained results suggest a partition that consists of 7 clusters: {B N}, {W L}, {Z F}, {P}, {T O}, {E C S G R D}, and {K}. Although the partition is much better, in terms of the silhouette coefficient, than the partition resulting from the application of the Szczotka–Spaeth technique, it is still ineffective. This is due to an unacceptable coefficient value for the second best DMU and a rather poor value for the top DMU. It is truly a pity, as fairly high silhouette coefficient values are registered for the majority of the remaining DMUs.

The comparison of final results for the sixteen DMUs case is presented in Table 21. It transpires that the application of different approaches results in the identification of partitions that differ in the structure of DMU clusters. The number of clusters ranges from 5 in the case of the Pareto rule, application to 11 in the case of the Szczotka–Spaeth technique. Almost all the approaches proved capable of identifying partitions where more than half of the DMUs form clusters consisting of at least two DMUs. Szczotka–Spaeth is a notable exception in this regard.

Table 20. *VIKOR-like rule approach rule-based use illustration for the sixteen DMUs case*

DMU	I	II	III	IV	V	VI	VII	$s(i)$
N	16	14	12	10	9	7	1	
$\emptyset$	0.067	0.077	0.091	0.111	0.125	0.167	–	
B	1	–	–	–	–	–	–	0.355
N	0.934	–	–	–	–	–	–	–0.450
W	0.905	1	–	–	–	–	–	0.778
L	0.891	0.984	–	–	–	–	–	0.818
Z	0.789	0.872	1	–	–	–	–	0.733
F	0.760	0.839	0.963	–	–	–	–	0.789
P	0.559	0.618	0.709	1	–	–	–	–
T	0.454	0.502	0.576	0.812	1	–	–	0.832
O	0.437	0.483	0.554	0.782	0.962	–	–	0.797
E	0.374	0.413	0.473	0.668	0.822	1	–	0.662
C	0.369	0.408	0.468	0.660	0.812	0.988	–	0.728
S	0.355	0.392	0.450	0.634	0.781	0.950	–	0.776
G	0.342	0.378	0.434	0.612	0.753	0.916	–	0.724
R	0.342	0.378	0.434	0.612	0.753	0.916	–	0.724
D	0.337	0.373	0.428	0.603	0.742	0.903	–	0.618
K	0.287	0.318	0.364	0.514	0.632	0.769	1	–

Table 21. *Final results for the sixteen DMUs case*

Approach (number of clusters)	Partition			
	$\min s(i)$ (DMU)	mean $s(i)$	$\max s(i)$ (DMU)	std.dev. $s(i)$
Simple gradient (9)	{B}={N}={W L}={Z}={F}={P}={T O}={E C S G R D K}			
	0.528 (W)	0.707	0.832 (T)	0.085
Szczotka–Spaeth (11)	{B}={N}={W}={L}={Z}={F}={P}={T}={O}={E C S G R D}={K}			
	0.617 (E)	0.692	0.776 (S)	0.058
Pareto rule (5)	{B N W L}={Z F}={P T}={O E C S}={G R D K}			
	–0.327 (T)	0.414	0.831 (F)	0.336
VIKOR-like rule (7)	{B}={N}={W L}={Z F}={P}={T O}={E C S G R D}={K}			
	–0.450 (N)	0.635	0.854 (G, R)	0.322

4. DISCUSSION

Results of the presented analysis show that traditional clustering approaches seem capable of indicating partitions that provide a good fit for individual DMUs. On the other hand, both proposed rule-based approaches seem to lack such a capability. A closer look at the results of the application of distinct techniques (see Table 14 and Table 21) nevertheless provides some hints regarding possible improvements in the results. For example, there is only a single E5 sustainability program that doesn't fit the final partition suggestion provided by the application of both rule-based approaches for the six DMUs case (see Table 14). Hence, we could try to break its common cluster with the E2 program (VIKOR rule-based partition only) or even merge it with the cluster of worse sustainability programs. The anticipated effects of these actions are presented in Table 22. The effects nonetheless confirm that the corrections result in unsatisfactory outcomes.

Table 22. *Anticipated effects of partition corrections for the six DMUs case*

Approach	Partition	$s(i)$					
		E2	E5	E3	E1	E6	E4
Pareto rule	$\{E2\}=\{E5\ E3\ E1\}=\{E6\ E4\}$	–	–0.20	0.27	0	0.52	0.50
VIKOR-like rule	$\{E2\}=\{E5\}=\{E3\ E1\ E6\ E4\}$	–	–	0.42	0.58	0.72	0.69
	$\{E2\}=\{E5\ E3\ E1\ E6\ E4\}$	–	–0.40	0.63	0.67	0.69	0.67

Let us see if it is possible to improve the results of the application of rule-based approaches in the sixteen DMUs case. The contents of Table 19 suggest that the main problem with the inefficiency of the partition resulting from the application of the Pareto rule clearly pertains to the cluster $\{P\ T\ O\ E\ C\ S\}$. The somewhat poor, although not very bad, silhouette coefficient values for P, E, and C suggest that it could be advantageous to separate them into distinct clusters: $\{P\}$ and $\{E\ C\}$. By the way, conducting the same action regarding the least preferable DMU, namely K, may also help improve individual silhouette coefficient values. Note that negative values for the coefficient in the cases of T and O also suggest that these DMUs could benefit from a common, distinct cluster. On the other hand, the negative silhouette coefficient value for S suggests that it would fit better into a common cluster with slightly better DMUs (C and E) and slightly worse DMUs (G, R, and D). Hence, we could finally obtain the corrected partition, which would consist of 6 clusters: $\{B\ N\ W\ L\}$, $\{Z\ F\}$, $\{P\}$, $\{T\ O\}$, $\{E\ C\ S\ G\ R\ D\}$, and $\{K\}$. The quality of the derived partition is confirmed by the contents of Table 23.

Table 23. *Recommended corrected partition for sixteen DMUs case – Pareto rule approach application*

Partition	min $s(i)$	mean $s(i)$	max $s(i)$	std.dev. $s(i)$
$\{B\ N\ W\ L\}=\{Z\ F\}=\{P\}=\{T\ O\}=\{E\ C\ S\ G\ R\ D\}=\{K\}$	0.529 (L)	0.713	0.832 (T)	0.088

In the case of the VIKOR-like rule application (see Table 20), the presence of only two unsatisfactory values for individual silhouette coefficient values for adjacent top DMUs, B and N, seems to suggest a rather obvious and simple solution for improving the partition. The solution would address the final division of two top DMUs into two separate clusters or the final integration of N into the adjoining cluster {W L}. Note that the latter partition is identical to the sg7 partition from the seventh level of the cluster hierarchy obtained for the simple gradient approach application (see Table 15). Nevertheless, the descriptive statistics for individual silhouette coefficient values for the partition presented in Table 16 show that N fits rather poorly with the common cluster of W and L. The division of the two top DMUs among two distinct clusters is finally recommended. Therefore, the final recommended partition consists of 8 clusters. Note that the corrected partition is only slightly worse than the best partitions derived from the application of both clustering approaches, as its core disadvantage results from a lower silhouette coefficient for W only.

All in all, both clustering approaches and both rule-based approaches proved to be capable of recommending diverse partitions to which DMUs fit well, at least. The diversity of the partitions results in a different number of clusters (ranging from 6 in the case of the Pareto rule to 11 in the case of the Szczotka-Spaeth technique) and differences in silhouette coefficient values for individual DMUs. The results of the calculations conducted show that very good partitions – in terms of both average and maximum silhouette coefficients (values over 0.71) – are provided by the corrected results of the application of the VIKOR-like and Pareto rule-based approaches. Note, however, that the results provided by the use of both clustering approaches and the simple gradient approach, in particular, do not prove to be significantly worse in this regard. Moreover, the partitions provided by correcting the direct VIKOR-like rule and using a simple gradient technique are almost the same, with only two DMUs – Z and F – being treated differently by these methods. Note that balanced partitions are ultimately recommended by the application of all four approaches, as the standard deviation level for individual silhouette coefficient values is relatively low. Almost all the approaches proved to be capable of recommending partitions in which more than half of the DMUs form clusters consisting of at least two DMUs Szczotka-Spaeth is, nevertheless, a notable exception in this regard.

5. CONCLUSIONS

Some possible ways to interpret the results of qualitative AHP/ANP technique applications in a more adequate manner were discussed in the paper. The application of both common hierarchical clustering approaches (the simple gradient approach and the Szczotka-Spaeth technique), as well as simplified approaches based on discrimination rules (the Pareto rule and the VIKOR-like rule), is considered in this regard. Two distinct real case studies, which differed in the number of DMUs, were used to initially assess the suitability of the approaches for an adequate analysis of AHP/ANP results.

It turns out that common clustering approaches seem to be generally capable of providing adequate means for expressing qualitative differences between DMUs, thanks to their reliable division among distinct groups, regardless of how many DMUs are actually considered. Their use, particularly the use of a less tedious and simpler gradient approach, is therefore recommended.

The considered simplified rule-based approaches are generally less tedious but do not seem to directly succeed in recommending partitions in which all DMUs fit well, at least. However, the quality of the corrected results from their application, obtained for a considerable number of DMUs, seems to suggest that they may, nevertheless, prove to be useful in peculiar circumstances, at least. This is because they seem to be capable of providing a reliable basis for the identification of unique competitive partitions that common clustering approaches were not aware of at all. Hence, although they do not seem to directly derive reliable DMUs partitions, they should not be disregarded.

It is obvious that the rather limited scope of the conducted analysis does not allow for the formation of reliable general conclusions. Further research is recommended; therefore, a comprehensive investigation into the effects of using considered approaches for grouping AHP/ANP application results is necessary to facilitate their qualitative interpretation. For instance, simulating the influence of the number of DMUs and other parameters of decision analysis problems on the effects of using the approaches could prove advantageous in this regard. The verification of the suitability of other approaches is also welcome. For example, the application of diverse metrics for cluster separation quality assessment (Bezdek & Pal, 1998; Dunn, 1973; Pakhira et al., 2004; Tibshirani et al., 2001) and various clustering techniques (Ismkhan, 2017; Mankowski & Moshkov, 2021; Szkaliczki, 2016; Wang & Song, 2011). The influence of the use of different priority estimation techniques and assessment scales in AHP/ANP, as well as similarity metrics, may also be considered in this regard. All in all, it is nonetheless hoped that the paper will foster a serious discussion about the adequate way to interpret the results provided by AHP/ANP applications in a meaningful and consistent manner, at least.

REFERENCES

- Bezdek J.C. & Pal N.R. (1998). Some new indexes of cluster validity. *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, **28**(3), pp. 301–315. DOI: <https://doi.org/10.1109/3477.678624>.
- Dunn J.C. (1973). A fuzzy relative of the ISODATA process and its use in detecting compact well-separated clusters. *Journal of Cybernetics*, **3**(3), pp. 32–57. DOI: <https://doi.org/10.1080/01969727308546046>.
- Florek K., Łukaszewicz J., Perkal J., Steinhaus H. & Zubrzycki S. (1951). Taksonomia numeryczna wrocławska. *Przegląd Antropologiczny*, **17**, pp. 193–211.
- Ginda G. & Szyba M. (2020). How different are Polish regions with regard to the utilisation of biogas plant technology application potential? [In:] *ISAHP 2020 Web conference*. URL: https://isahp2020.exordo.com/files/papers/107/final_draft/ggms_isahp2020.pdf [10.12.2020].

- Ishizaka A., Nemery P. & Pearman C. (2012). AHPSort: An AHP based method for sorting problems. *International Journal of Production Research*, **50**(17), pp. 4767–4784. **DOI:** <https://doi.org/10.1080/00207543.2012.657966>.
- Ismkhan H. (2017). A.1D-C: A novel fast automatic heuristic to handle large-scale one-dimensional clustering. *Applied Soft Computing*, **52**, pp. 4767–4784. **DOI:** <https://doi.org/10.1016/j.asoc.2016.09.001>.
- Kaufman L. & Rousseeuw P.J. (1990). *Finding Groups in Data: An Introduction to Cluster Analysis*. Hoboken, NJ: Wiley-Interscience. **DOI:** <https://doi.org/10.1002/9780470316801>.
- Kolenda M. (2006). *Taksonomia numeryczna. Klasyfikacja, porządkowanie i analiza obiektów wielocechowych*. Wrocław: Wydawnictwo Akademii Ekonomicznej im. Os-kara Langego.
- Mankowski M.A. & Moshkov M.J. (2021). One-dimensional clustering. *Studies in Systems, Decision and Control*, **331**, pp. 123–133.
- Netto A.L., Salomon V.A.P., Ortiz-Barros M.A., Florek-Paszkowska A.K., Petrillo A. & De Oliveira O. (2021). Multiple criteria assessment of sustainability programs in the textile industry. *International Transactions in Operational Research*, **28**, pp. 1550–1572. **DOI:** <https://doi.org/10.1111/itor.12871>.
- Opricovic S. (1988). *Multicriteria Optimization of Civil Engineering Systems*. Belgrade: Faculty of Civil Engineering, University of Belgrade.
- Opricovic S. & Tzeng G.-H. (2004). Compromise solution by MCDM methods: A comparative analysis of VIKOR and TOPSIS. *European Journal of Operational Research*, **156**(2), pp. 445–455. **DOI:** [https://doi.org/10.1016/S0377-2217\(03\)00020-1](https://doi.org/10.1016/S0377-2217(03)00020-1).
- Pakhira M.K., Bandyopadhyay S. & Maulik U. (2004). Validity index for crisp and fuzzy clusters. *Pattern Recognition* **37**(3), pp. 487–501. **DOI:** <https://doi.org/10.1016/j.patcog.2003.06.005>.
- Saaty T.L. (1980). *The Analytic Hierarchy Process: Planning, Priority Setting, Resource Allocation*, New York: McGraw-Hill.
- Saaty T.L. (1996). *Decision Making with Dependence and Feedback*. Pittsburgh: RWS.
- Spaeth H. (1973). Algorithm of clustering of one-dimensional data ordered data. *Computing*, **11**(2), pp. 175–177. **DOI:** <https://doi.org/10.1007/BF02252907>.
- Szczotka F.A. (1972). On the method of ordering and clustering objects. *Zastosowania Matematyki*, **XIII**(1), pp. 23–34.
- Szkaliczki T. (2016). Clustering.sc.dp: Optimal clustering with sequential constraint by using dynamic programming. *R Journal*, **8**(1), pp. 318–327. **DOI:** <https://doi.org/10.32614/rj-2016-022>.
- Tibshirani R., Walter 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), pp. 411–423. **DOI:** <https://doi.org/10.1111/1467-9868.00293>.
- Wang H. & Song M. (2011). Ckmeans.1d.dp: Optimal k-means clustering in one dimension by dynamic programming. *R Journal* **3**(2), pp. 29–33. **DOI:** <https://doi.org/10.32614/RJ-2011-015>.