



TECHNICAL UNIVERSITY OF CLUJ-NAPOCA

ACTA TECHNICA NAPOCENSIS

Series: Applied Mathematics, Mechanics, and Engineering

Vol. 69, Issue II, June, 2026

ADAPTABILITY ANALYSIS OF APPARENT FRONT RANKING

Mihai NEGHINĂ

Abstract: In the context of multi-objective optimization through genetic algorithms, the ranking of individuals is a delicate and consequential operation. Introduced in 2020, the method of Apparent Front Ranking has shown potential, being successfully included in a GA. The current paper offers a deeper analysis of the flexibility of AFR, with insights about the evolution of ranks as the AF template power is modified, showing excellent Kendall and Pearson correlations with the reference hierarchical methods in uniform datasets of up to 10 dimensions. By adjusting the AF template power, the best Kendall rank correlation is above 0.79 with all reference methods, while the best Pearson correlation is above 0.95 with Saaty's AHP and above 0.75 with Köppen's average FPD.

Key words: Ranking Methods, Pareto Optimality, Multi-Objective Optimization Methods.

1. INTRODUCTION

Optimization problems, especially the ones involving multiple partially conflicting objectives, are inherently challenging due to the need to balance trade-offs. Although rarely having a clear-cut, objectively best solution from every point of view, early methods nevertheless imposed a decision-making structure through defining a combined cost function and attempting to optimize it a solution representing a local/global extremum under a set of constraints.

Genetic Algorithms (GAs) have embraced the diversity of incomparable solutions [1], most notably through the introduction of Pareto-optimality [2]. Rather than reducing the problem to a single function as with earlier methods, those approaches generate sets of non-dominated solutions that represent the Pareto front [3]. Numerous creative strategies have been explored for selecting solutions of the Pareto front [4], including dimensionality reduction through principal component analysis [5], prioritizing diversity over dominance among candidates [6],[7],[8],[9],[10] and redefining Pareto dominance using fuzzy logic [11],[12],[13],[14], all of which aim to establish a ranking system for individuals. All these methods offer a more nuanced comparison of individuals, but cannot avoid the inconvenience of pairwise

comparisons, the most computationally demanding step in ranking.

The Pareto front, although composed from discrete non-dominated solutions, could be accurately represented for most problems as a smooth hypersurface in the continuous space of objectives. The Apparent Front (AF), introduced in [15], attempts to leverage this observation and turn the problem of identifying individual solutions of the Pareto set into estimating the Pareto front in its entirety. The AF is also defined as a smooth hypersurface, constructed from a set of support vectors, that tries to approximate the completely unknown Pareto front and to replace it during further computations of distance or population fitness, effectively establishing the Apparent Front Ranking (AFR) [15].

A version of AFR has been successfully integrated in a genetic algorithm, (C)AFZGA [16], that was proven comparable to other GAs, while maintaining a lower computational cost. (C)AFZGA was also included as a viable alternative in the 2024 review paper considering Optimization Frameworks for Multi-Objective Evolutionary Algorithms [17].

This paper continues the investigation on the flexibility and the potential of AFR by exploring in more detail various ways of generating AFs in datasets of up to ten dimensions and comparing the resulting ranking to Saaty's Analytic

Hierarchy Process (AHP) [18] and Köppen et al. Fuzzy Pareto Domination (FPD) [12].

The remainder of this paper is organized as follows: Section 2 presents the datasets as sets of vectors in the space of objectives, without concern about the population individuals that have produced them. Section 3 reviews AFR, describing various strategies of implementation at each step of the method. Section 4 presents the experimental results of AFR against FPD and AHP, using Kendall [19] and Pearson [20] as correlation metrics, and, finally, section 5 offers concluding remarks and perspectives.

2. DATASETS

For the sake of simplicity, the first dataset is a small set of solutions in two-dimensional space of objectives, needing maximization, as represented in Fig.1. This set, called "Dataset#1" for the remainder of the paper, is considered during the explanation and analysis of Apparent Front Ranking (AFR) variants.

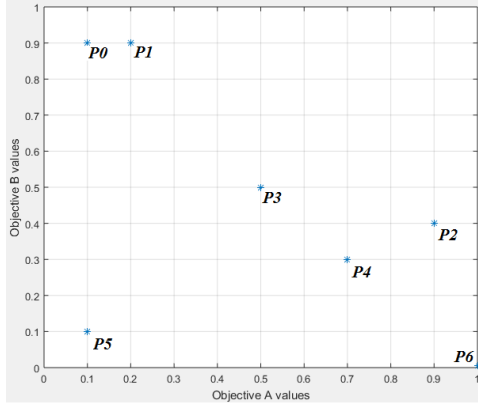


Fig. 1. Dataset #1 vectors in the space of objectives $P_0=[0.1,0.9]^T$ $P_1=[0.2,0.9]^T$ $P_2=[0.9,0.4]^T$ $P_3=[0.5,0.5]^T$ $P_4=[0.7,0.3]^T$ $P_5=[0.1,0.1]^T$ $P_6=[1.0,0.005]^T$

Nine larger Datasets containing randomly generated vectors are identified by numbers #2 to #10, the identification number corresponding to the number of dimensions in the space of objectives. Vectors from these sets obey eq. (1), that does not correspond to any Apparent Front (AF) from the available templates.

Dataset ID (Dimensions)	No. of vectors	Dataset ID (Dimensions)	No. of vectors
Dataset#1(2D)	7	Dataset#6(6D)	230
Dataset#2(2D)	212	Dataset#7(7D)	198
Dataset#3(3D)	279	Dataset#8(8D)	286
Dataset#4(4D)	303	Dataset#9(9D)	249
Dataset#5(5D)	269	Dataset#10(10D)	255

Fig. 2 and 3 show Datasets #2 and #3, including the magenta-colored non-dominated support vectors and one computed AF for each.

$$(-1)^N \cdot \prod_{k=1}^N (p_k - 2) \geq 1, \quad (1)$$

where $0 \leq p_k \leq 2$ for all k

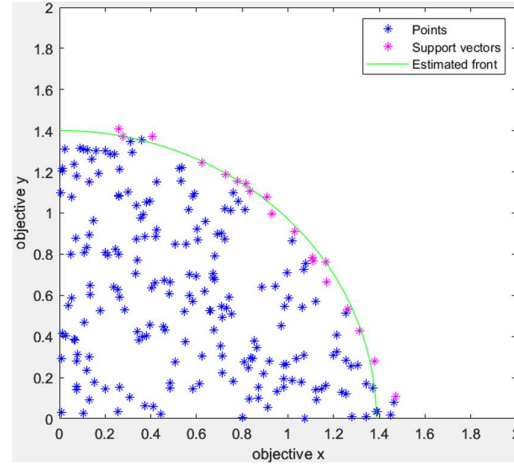


Fig. 2. Dataset #2, vectors, support vectors and an estimated elliptical Apparent Front

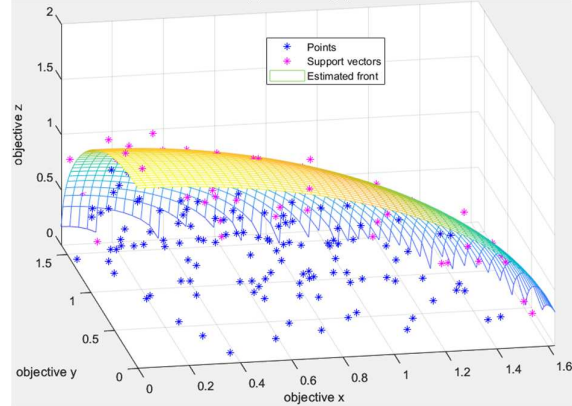


Fig. 3. Dataset #3, vectors, support vectors and an estimated ellipsoidal Apparent Front

All sets only have points in the first quadrant, which is an arbitrary constraint that helps with the visual representations and explanations. However, this constraint does not affect the generality of the conclusions, since any finite dataset can be turned into a first quadrant set through translation of the origin or a more general affine transformation.

3. APPARENT FRONT RANKING (AFR)

The AFR algorithm has four main steps, but each can be implemented in a variety of ways without changing the main idea of the method, emphasizing its flexibility:

1. Choice of the AF template
2. Selection of support vectors
3. Estimation of the AF
4. Computation of distances and ranking
- 5.

3.1 Choice of the AF template

The AF template is a family of parametric functions, whose parameters need to be optimized at step 3. While the current paper focuses on the template described in eq. (2), other templates can be considered.

$$\sum_{k=1}^N a_k \cdot x_k^n - 1 = 0, \quad (2)$$

with $n \in \mathbb{R}_+^*$; $a_k > 0, \forall k = 1, \dots, N$

These families of functions should ideally offer the properties expected from a Pareto front, i.e. no point on these hypersurfaces would be dominated by any other point on the curves/surfaces in the first quadrant. The template in eq. (2) adheres to these constraints. Examples of functions from this 2D template $x^n + 4 \cdot y^n - 1 = 0$ are shown in Fig. 4, all passing through the point of coordinates [1,0]. The other quadrants apart from quadrant I, which is the space of objectives, are grayed out, but their representation helps with visualizing the bigger picture regarding the function.

As shown in Fig. 4, the power n may be a real positive value, although only integer powers generate values outside quadrant I, and only even integer powers create generalized ellipses. For $n = 1$, the vectors that lie along the quadrant symmetry axis would fare better than vectors near the edges. Large values of n , would be advantageous for points near the edges of the quadrant. Although it is possible to choose $n < 1$, with the effect of over-emphasizing vectors that are good in all objectives but not necessarily exceptional in any, the meaning of an estimated Pareto front would be partially distorted due to the loss of the function concavity.

3.2 Selection of support vectors

Ideally, the support vectors should be the non-dominated vectors of the current population. For Dataset#1, the vectors on the Pareto front are marked with magenta in Fig. 5a.

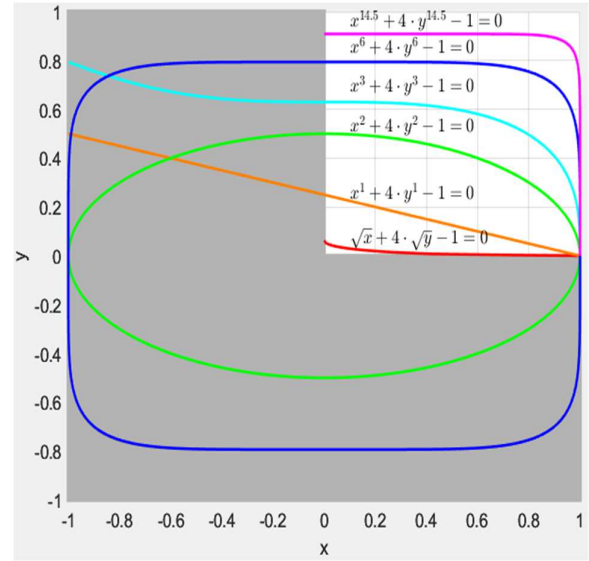


Fig. 4. Examples of AFs for various powers n

However, finding all non-dominated vectors runs into two conceptual problems: first, for a sufficiently large number of dimensions, almost all vectors become non-dominated, and second, the pairwise comparisons needed to assert non-dominance defeat the purpose of AFR as a computationally efficient ranking method.

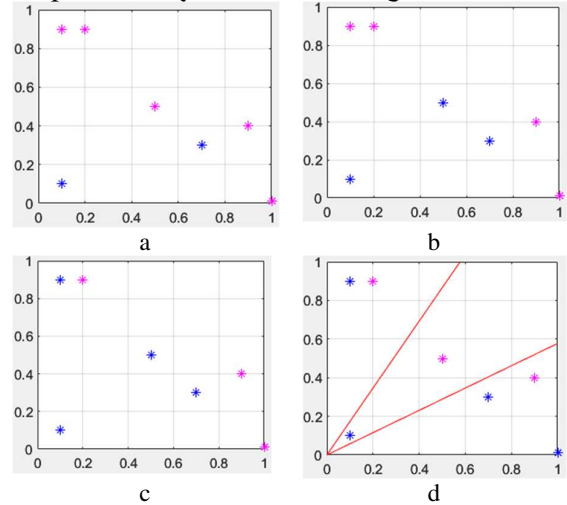


Fig. 5 Methods for generating support vectors:

- a. Pairwise comparisons (non-dominance);
- b. Best two vectors on each dimension;
- c. Best average on each pair of dimensions;
- d. Best sum on each partition ($\theta=\pi/6$).

As the template function may require only a handful of support vectors, at least $N + 1$ in the case of the chosen template from eq. (2), various alternative methods of selecting enough support vectors (without pairwise comparisons) may be considered, as shown in Fig. 5b,c,d:

- Selecting best-fitted vectors for each dimension, regardless of their evaluation along other dimensions. Fig. 5b shows the selection of two vectors for each dimension.
- Selecting one or more vectors with best average values for each pair of dimensions. As long as any vector is chosen only once, there would be a selection for $N \cdot (N - 1)/2$ vectors. It should be noted that the complexity introduced by such a procedure would still be significantly lower than any pairwise comparison between all the vectors. Although this method is less relevant for 2D spaces, Fig. 5c shows the selection of the best 3 vectors according to this criterion.
- Partitioning the quadrant into a large enough number of regions and then selecting one or multiple vectors from each region (e.g. the vectors with the highest sum of the values). Fig. 5d shows the partition into 3 regions, two along each axis by an angle $\theta = \pi/6$, the third being the remaining center, and the selection of the best vector in each region.
- When using our AFR algorithm in an iterative way (as with multi-objective GAs), the AF of the previous generation could be used to identify the most advanced individuals of the current generation as part of the support vectors.

Each of these support vectors selection methods has its drawbacks but all of them are more computationally efficient than the pairwise comparisons used in classical Pareto approaches. For each such method of vectors selection, there need to be mechanisms in place ensuring that a sufficiently large number of support vectors are selected, as well as ensuring at least some diversity. For the template considered, the minimum number of necessary support vectors is one more than the number of dimensions or $N + 1$, but more support vectors could help in better estimating the AF.

3.3 Estimation of the AF

The estimation of the curve that best fits the support vectors can be done through a number of statistical techniques, such as least-square minimization. This paper considers a pre-selected power n and uses a fast least square

fitting technique, adapted from the Fitzgibbon et al. [21] to apply to general ellipsoids described mathematically through eq. (2). Fig. 6 shows the estimated AF of Dataset #1 for $n = 1, 2, 3$, using support vectors from Fig. 5a.

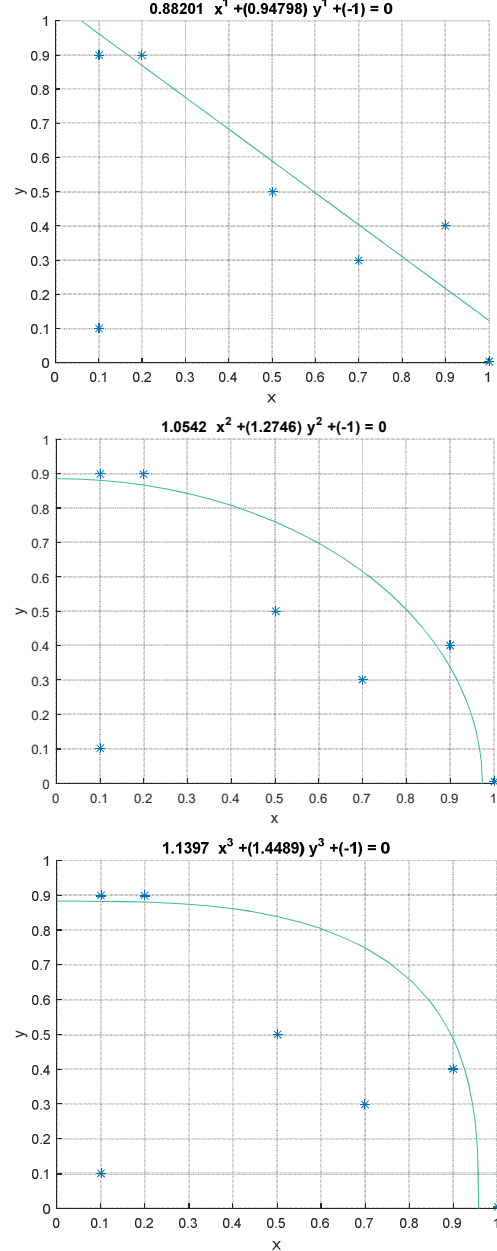


Fig. 6. AFs for Dataset #1, for powers $n = 1, 2, 3$

In case the objectives are not equally important, the function could be fitted on weighted virtual support vectors, such as replacing P_0 from Dataset #1 with its weighted version P_{0w} that considers weights w_x and w_y on the two dimensions as in eq.(3):

$$P_{0w} = [0.1 \cdot w_x, \quad 0.9 \cdot w_y]^T \quad (3)$$

The weights w_x and w_y need to be positive but are otherwise unconstrained and can be adjusted separately. Without loss of generality, last weight could be fixed (e.g. $w_y = 1$) or the weights could be required to sum up to N . These additional constraints may be more related to the interpretation and positioning of the weighted virtual vectors in the space of objectives, rather than the mathematical requirements.

The unweighted 2D vectors could be seen as having $w_x = 1$ and $w_y = 1$. If the weights are $w_x > 1, w_y = 1$, the expectation on objective x is higher and therefore the vectors on the Pareto front with high x value (e.g. P6) would score more poorly compared to vectors with high y value (e.g. P0) because it would be easier for them to surpass the AF built upon the weighted support vectors. Conversely, if $w_x < 1, w_y = 1$, the AF is skewed so that vectors like P6 would score better compared to vectors like P0.

3.4 Computation of proximity to the AF

The proximity to the AF for each vector, $d(P)$, would simply be the evaluation of the determined function with computed coefficients a_k , as shown in eq. (4):

$$d(P) = \sum a_k \cdot p_k^n - 1 \quad (4)$$

Computing distances is equivalent to finding the hypersurface (curve) similar to the AF that goes through the considered vector, equivalent to adjusting the free term a_0 , as shown in eq. (5) and Fig. 7.

$$\sum a_k \cdot x_k^n - a_0 = 0 \quad (5)$$

For points beyond the AF, $a_0 > 1 \Rightarrow d > 0$; for points lying on the AF, $a_0 = 1 \Rightarrow d = 0$; for points below the AF, $a_0 < 1 \Rightarrow d < 0$; for $[0,0]^T$, $a_0 = 0 \Rightarrow d = -1$.

4. RESULTS

4.1 Algorithms for Comparison

For comparison, two ranking methods from the literature are considered: Köppen's algorithm called Fuzzy Pareto Dominance (FPD) [12] and Saaty's Analytic Hierarchy Process (AHP) [18].

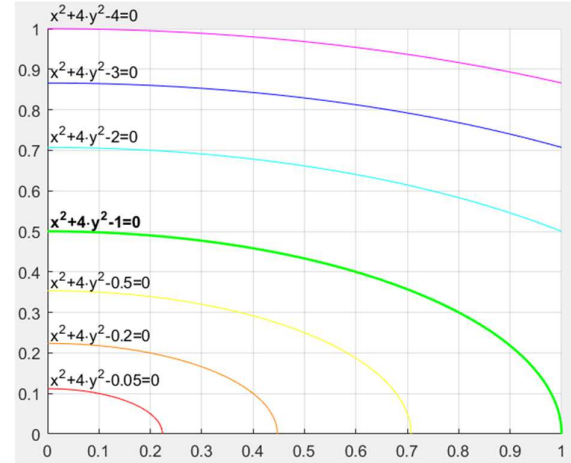


Fig. 7. Set of curves similar to the AF, with various values for the free term a_0

AHP generates the criteria weights and the matrix of option scores which are turned into rankings. Since no supposition can be made on the importance of each criterion in this abstract setting, all weights are considered equal. The matrix of option scores is computed from the pairwise comparison matrices along each dimension. Elements of the pairwise comparison matrices $m_{a,b}^{(k)}$ represent the evaluation of vectors a and b with respect to criterion k and satisfy $m_{a,b}^{(k)} \cdot m_{b,a}^{(k)} = 1$. The versions considered in this study differ in the way elements $m_{a,b}^{(k)}$ are computed, shown through eq. (6) for **AHP v1** and eq. (7) for **AHP v2**:

$$m_{a,b}^{(k)} = \frac{P_a(k)}{P_b(k)} \quad (6)$$

$$m_{a,b}^{(k)} = \left(8 \cdot \frac{|P_a(k) - P_b(k)|}{MAX(k) - MIN(k)} + 1 \right)^{s_k} \quad (7)$$

where $s_k = \text{sign}(P_a(k) - P_b(k))$ and $MAX(k)$ and $MIN(k)$ represent the maximum and minimum value along objective k .

For maximizing objectives, FPD estimates the asymmetric fuzzy domination of vector a over another vector b according to eq. (8):

$$\mu_{ab} = \frac{\prod_i \min(P_a(i), P_b(i))}{\prod_i P_b(i)} \leq 1 \quad (8)$$

The two versions of FPD ranking considered in this study are:

FPDmin, ranking according to the minimum dominance of each vector over all others, and

FPDave, ranking according to the average dominance of each vector over all others.

4.2 Performance Metrics

In order to evaluate the similarity of the rankings, the Kendall [19] and Pearson [20] similarity measures are considered.

The Kendall similarity measure can be briefly described as follows: any set $S = \{s_1, s_2, \dots, s_N\}$ of N objects ordered according to their ranks $R = \{r_1, r_2, \dots, r_N\}$, $r_i \neq r_j$ for $i \neq j$, can be decomposed as a set of $0.5 \cdot N \cdot (N - 1)$ ordered pairs $O_{S,R} = \{[s_i, s_j] | r_i < r_j\}$. The symmetric difference between two sets of ordered pairs $O_{S,R1}$ and $O_{S,R2}$, $\Delta(O_{S,R1}, O_{S,R2})$, is the number of different pairs between these two sets. The Kendall correlation coefficient τ is then computed as eq. (9):

$$\tau = \frac{0.5 \cdot N \cdot (N - 1) - \Delta(O_{S,R1}, O_{S,R2})}{0.5 \cdot N \cdot (N - 1)} \quad (9)$$

The Kendall correlation coefficient τ is between 0 (no correlation) and 1 (perfect correlation, i.e. identical ranking of objects)

Pearson linear correlation coefficient can be computed on ranks or on the values leading to those ranks, with eq. (10):

$$\begin{aligned} \rho_{X,Y} &= \frac{\text{cov}(X,Y)}{\sigma_X \cdot \sigma_Y} \\ &= \frac{\sum_{i=1}^N (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^N (x_i - \bar{x})^2} \cdot \sqrt{\sum_{i=1}^N (y_i - \bar{y})^2}} \end{aligned} \quad (10)$$

The values considered when computing the correlation are the AHP values for both versions, the FPD average and minimum dominance values and AFR distances to the Apparent Front.

4.3 Results for Dataset #1

Table 1 presents the computed values for Dataset #1, using the two versions of FPD and AHP, as well as three versions of AFR considering the template in eq. (2) with powers $n = 1, 2$ and 3. The numbers in parentheses are the ranks of the vectors according to each method, with rank 1 denoting the best vector.

Table 1. Computed values and ranks for each vector according to AHP, FPD and AFR

	AHP v1	AHP v2	FPD min	FPD ave
P₀	2.229 (3)	2.558 (3)	0.100 (5)	0.436 (5)
P₁	2.429 (2)	2.648 (2)	0.200 (4)	0.586 (4)
P₂	2.701 (1)	2.519 (4)	0.444 (2)	0.798 (1)
P₃	2.127 (4)	1.506 (6)	0.500 (1)	0.697 (2)

P₄	2.076 (5)	1.559 (5)	0.333 (3)	0.650 (3)
P₅	0.425 (7)	0.429 (7)	0.027 (6)	0.197 (6)
P₆	2.011 (6)	2.779 (1)	0.005 (7)	0.157 (7)
	AFR $n = 1$	AFR $n = 2$	AFR $n = 3$	
P₀	-0.0586 (3)	0.0430 (4)	0.0574 (3)	
P₁	0.0296 (2)	0.0746 (1)	0.0653 (2)	
P₂	0.1730 (1)	0.0578 (2)	-0.0764 (4)	
P₃	-0.0850 (4)	-0.4178 (6)	-0.6764 (6)	
P₄	-0.0982 (5)	-0.3687 (5)	-0.5700 (5)	
P₅	-0.8170 (7)	-0.9767 (7)	-0.9974 (7)	
P₆	-0.1133 (6)	0.0542 (3)	0.1397 (1)	

On such a small set, it might be a coincidence that the ranks of AHP v1 and AFR with $n = 1$ are the same and AHP v2 with AFR with $n = 3$. In order to have a broader view, Fig. 8 shows the evolution of ranks on Dataset #1 when varying n from 0.01 to 4.00.

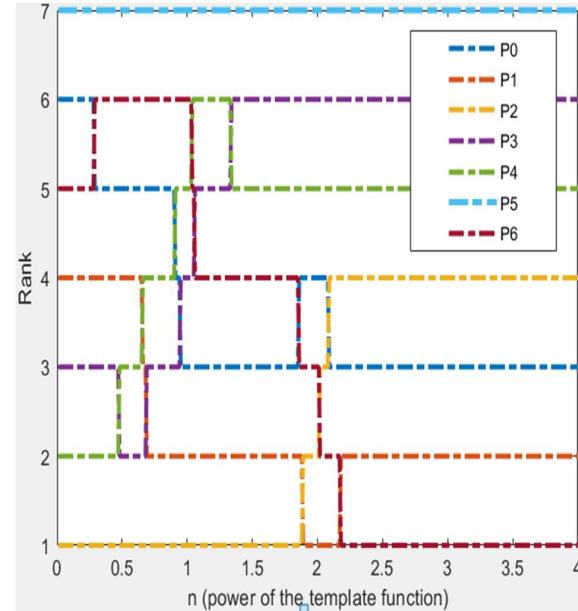


Fig. 8. Evolution of ranks on Dataset #1

It can be noticed that for low values of n , P6 starts with a low rank, as expected by AHP v1 and FPD, but as n increases and the template becomes more curved P6 moves up the ranks, ending up as the best vector in AHP v2. The rank of P0 is always lower than the rank of P1, but both eventually overtake P2 for sufficiently curved AFs. To understand the evolution of the ranks, it might be useful to also observe the evolution of values when varying the power of the template function, as displayed in Fig. 9.

The ranks of AHP v1 and AFR n1 are identical and the correlation on the values that lead to those ranks is not quite perfect (0.999) but sufficiently close to 1, which might indicate that AHP v1 and AFR n1 could produce similar

rankings even on larger sets. Another broader way of presenting the information is to include experiments with a lot more values for the power of the AF (n varying from 0.01 to 4) and to display the evolutions of Kendall and Pearson correlations, as shown in Fig. 10. From these evolutions we can see that AFR comes closest to the FPD variants for very low values of n , to AHP v1 for $n = 1$ and to AHP v2 for larger values of n .

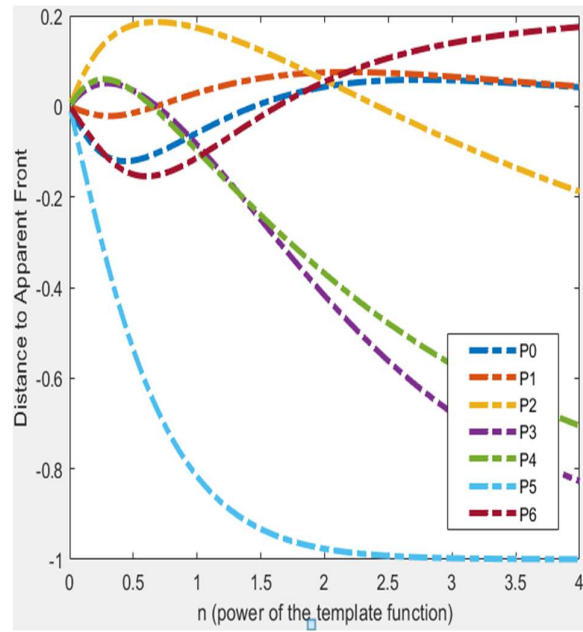


Fig. 9. Evolution of values on Dataset #1

4.4 Results on test sets of up to 10 dimensions

Although Fig. 9 may be useful for a limited number of vectors in the set, such as in Dataset #1, for hundreds of vectors as in Datasets #2 to #10 such images become overloaded with information and therefore more confusing than helpful, as shown in Fig. 11.

For these datasets, the most useful synthesis comes from explicit tables with the highest correlation values for Kendall and Pearson similarity (Tables 2 to 10).

Each table presents the best correlations of AFR against both AHP versions and both FPD versions, also noting the value of the power n for which that correlation is achieved. Coincidentally, the table numerals correspond to the Dataset identification numbers.

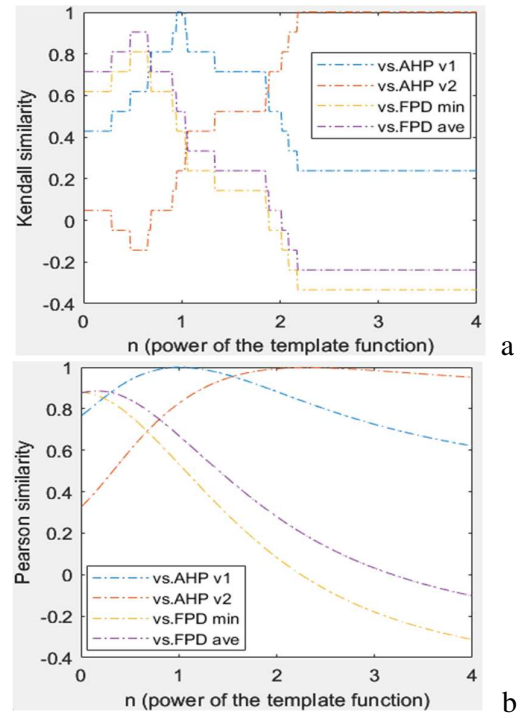


Fig. 10. Evolution of similarity on Dataset #1:
a. Kendall rank correlation;
b. Pearson value correlation.

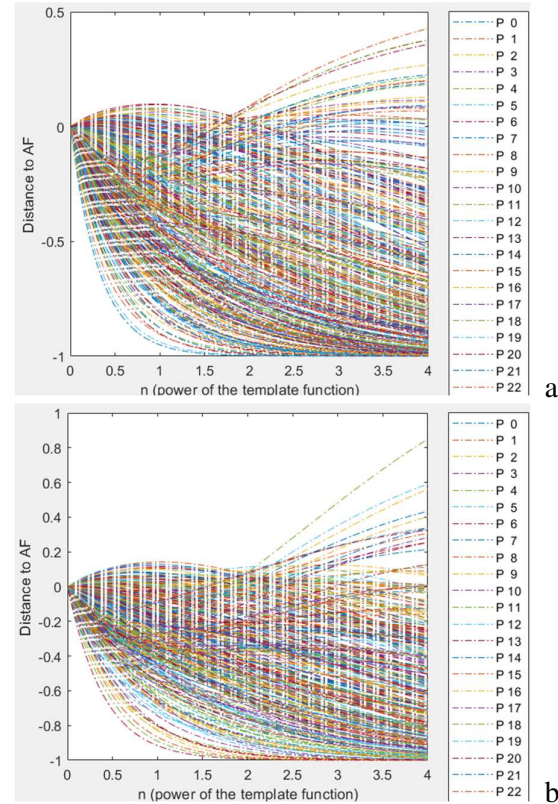


Fig. 11. Evolution of values on:
a. Dataset #2;
b. Dataset #3.

Table 2. Best correlations for Dataset #2

Dataset #2	Kendall on ranks	Pearson on ranks	Pearson on values
vs. AHP v1	0.9599 n = 1.05	0.9974 n = 1.00	0.9980 n = 1.02
vs. AHP v2	0.9757 n = 2.27	0.9990 n = 2.27	0.9990 n = 2.30
vs. FPD min	0.7776 n = 0.08	0.9332 n = 0.08	0.8592 n = 0.30
vs. FPD ave	0.8371 n = 0.09	0.9599 n = 0.11	0.9375 n = 0.21

Table 3. Best correlations for Dataset #3

Dataset #3	Kendall on ranks	Pearson on ranks	Pearson on values
vs. AHP v1	0.9681 n = 1.03	0.9983 n = 1.03	0.9987 n = 1.02
vs. AHP v2	0.9740 n = 2.16	0.9989 n = 2.17	0.9990 n = 2.17
vs. FPD min	0.8286 n = 0.01	0.9610 n = 0.01	0.8009 n = 0.28
vs. FPD ave	0.8533 n = 0.01	0.9692 n = 0.01	0.9011 n = 0.15

Table 4. Best correlations for Dataset #4

Dataset #4	Kendall on ranks	Pearson on ranks	Pearson on values
vs. AHP v1	0.8739 n = 0.96	0.9777 n = 1.00	0.9818 n = 0.98
vs. AHP v2	0.8925 n = 2.23	0.9837 n = 2.28	0.9874 n = 2.24
vs. FPD min	0.7845 n = 0.07	0.9344 n = 0.06	0.7222 n = 0.35
vs. FPD ave	0.8018 n = 0.07	0.9431 n = 0.06	0.8503 n = 0.27

Table 5. Best correlations for Dataset #5

Dataset #5	Kendall on ranks	Pearson on ranks	Pearson on values
vs. AHP v1	0.9227 n = 1.01	0.9909 n = 1.01	0.9940 n = 1.01
vs. AHP v2	0.9399 n = 2.23	0.9943 n = 2.23	0.9953 n = 2.23
vs. FPD min	0.8597 n = 0.01	0.9725 n = 0.01	0.7375 n = 0.18
vs. FPD ave	0.8732 n = 0.01	0.9768 n = 0.01	0.8575 n = 0.11

Table 6. Best correlations for Dataset #6

Dataset #6	Kendall on ranks	Pearson on ranks	Pearson on values
vs. AHP v1	0.8765 n = 1.02	0.9781 n = 1.02	0.9826 n = 1.01
vs. AHP v2	0.8670 n = 2.09	0.9735 n = 2.09	0.9833 n = 2.11
vs. FPD min	0.7637 n = 0.12	0.9277 n = 0.09	0.7089 n = 0.21

vs. FPD ave	0.8016 n = 0.07	0.9492 n = 0.09	0.8436 n = 0.17
-------------	--------------------	--------------------	--------------------

Table 7. Best correlations for Dataset #7

Dataset #7	Kendall on ranks	Pearson on ranks	Pearson on values
vs. AHP v1	0.9344 n = 1.00	0.9933 n = 1.00	0.9957 n = 1.00
vs. AHP v2	0.9204 n = 2.16	0.9901 n = 2.16	0.9933 n = 2.14
vs. FPD min	0.8335 n = 0.01	0.9617 n = 0.01	0.6526 n = 0.24
vs. FPD ave	0.8583 n = 0.01	0.9722 n = 0.01	0.8178 n = 0.16

Table 8. Best correlations for Dataset #8

Dataset #8	Kendall on ranks	Pearson on ranks	Pearson on values
vs. AHP v1	0.8863 n = 0.97	0.9822 n = 0.97	0.9857 n = 0.99
vs. AHP v2	0.8814 n = 2.15	0.9802 n = 2.15	0.9863 n = 2.17
vs. FPD min	0.8040 n = 0.02	0.9481 n = 0.02	0.6379 n = 0.24
vs. FPD ave	0.8351 n = 0.01	0.9623 n = 0.01	0.7946 n = 0.14

Table 9. Best correlations for Dataset #9

Dataset #9	Kendall on ranks	Pearson on ranks	Pearson on values
vs. AHP v1	0.7952 n = 0.99	0.9409 n = 1.00	0.9510 n = 1.00
vs. AHP v2	0.8000 n = 2.14	0.9438 n = 2.26	0.9567 n = 2.21
vs. FPD min	0.7845 n = 0.01	0.9390 n = 0.03	0.5791 n = 0.26
vs. FPD ave	0.8379 n = 0.01	0.9654 n = 0.01	0.7887 n = 0.19

Table 10. Best correlations for Dataset #10

Dataset #10	Kendall on ranks	Pearson on ranks	Pearson on values
vs. AHP v1	0.8763 n = 1.00	0.9780 n = 0.99	0.9833 n = 1.01
vs. AHP v2	0.8631 n = 2.21	0.9727 n = 2.20	0.9821 n = 2.21
vs. FPD min	0.8565 n = 0.01	0.9717 n = 0.01	0.5926 n = 0.21
vs. FPD ave	0.8832 n = 0.01	0.9806 n = 0.01	0.7593 n = 0.14

The summary of the strongest correlation values and the powers for which these are obtained is shown in Table 11.

Table 11. Summary of the best correlation values

Vs.	Kendall on ranks	Pearson on ranks	Pearson on values
AHP v1	0.79 ... 1	0.94 .. 1	0.95 .. 1
	Obtained for $n = 1 \pm 0.05$		
AHP v2	0.79 ... 1	0.94 .. 1	0.95 .. 1
	Obtained for $n = 2.2 \pm 0.11$		
FPD min	0.76 ... 0.90 $n < 0.1$	0.92 .. 0.98 $n < 0.1$	0.57 ... 0.85 $n < 0.35$
FPD ave	0.76 ... 0.90 $n < 0.1$	0.92 .. 0.98 $n < 0.1$	0.75 ... 0.94 $n < 0.35$

Taking both the Kendall correlation on ranks and Pearson correlation on values into account, it appears that AFR can produce results similar to AHP by maintaining a similar clustering of decision values, while still showing comparable results to FPD, albeit with different relationships between values. The underlying principles of AFR seem to align more closely with AHP than with FPD, while bridging both approaches.

5. CONCLUSIONS AND PERSPECTIVES

Apparent Front Ranking is thus proven to be an efficient method of ranking multidimensional vectors, more computationally efficient than the well-established methods discussed, but also flexible enough to be able to produce similar results to either of them, depending on one parameter. Comparing ranks, AHP v1 seems to be closest for powers around $n = 1$, AHP v2 is closest for powers around $n = 2.2$ and FPD is closest for powers $n < 0.1$. Considering ranking values, AFR seems more aligned with the principles of AHP, showing a greater correlation which can be interpreted as generating similar values in order to generate similar ranks.

The deeper analysis of the flexibility of AFR includes insights about the evolution of ranks and values, as the AF template power is modified from almost 0 up to the value 4. On Dataset#1, where there are fewer vectors to keep track of, it can be explicitly seen that points with relatively high values for all objectives (but not exceptional in any) fare better and have are more favorable ranks with lower template powers, while points exceptional in only one objective are scoring better with higher template powers. These observations can be generalized to any number of dimensions, albeit more difficult to follow for a large number of individuals,

showing the particularities of the chosen template function.

The efficiency is mainly driven by the avoidance of computing pairwise comparisons between all vectors, both when selecting support vectors and when generating the ranking values. The method only needs the go through the set of vectors twice, not consider all combinations of two vectors. The flexibility of the method is given by the template function. The template can be adjusted to the studied problem and shaped based on any external information available. Within the template, the parameters can be adjusted to best fit the set of support vectors, resulting in rankings similar to AHP, FPD or neither. The flexibility of the AF could prove valuable in genetic algorithms that require dynamic adaptation of their selection process.

A version of AFR has already been successfully included in a genetic algorithm called (C)AFZGA [16], with encouraging results compared to NSGAI [22] and SPEA2 [23]. However, the full potential and flexibility of Apparent Fronts have not yet been explored, especially considering that other template function families can be also used and analyzed.

6. REFERENCES

- [1] Katoch S. et al., *A review on genetic algorithms: past, present, and future*. Multimedia Tools and Applications 80(5), 8091–8126, 2021. [DOI](#)
- [2] Pareto V., *Cours d'economie politique*. Vol. I/II., F. Rouge, 1896-7.
- [3] Zitzler E., *Evolutionary Algorithms for Multiobjective Optimization: Methods and Applications*. Swiss Federal Institute of Technology Zurich, Switzerland, 1999. ISBN-13: 9783826568312
- [4] Deb K., *Multi-objective optimization using evolutionary algorithms*. Wiley, UK, 2001. ISBN-13: 9780471873396
- [5] Deb K., Saxena D., *On finding Pareto-optimal solutions through dimensionality reduction for certain large-dimensional multi-objective optimization problems*. Kangal report 2005011, Indian Institute of Technology, 2005. [SemanticsScholarLink](#)

- [6] Laumanns M., Thiele L., Deb K., Zitzler E. *Combining convergence and diversity in evolutionary multiobjective optimization*. *Evol. Comp.* 10.3, 263–282, 2002. [DOI](#)
- [7] Adra S., Fleming P. *Diversity management in evolutionary many-objective optim.* *Evol. Comp.* 15.2, 183–195, 2011. [DOI](#)
- [8] Jiang S., Yang S., *Convergence versus diversity in multiobjective optimization*, *PPSN XIV* 9921, pp. 984–993, 2016. [DOI](#)
- [9] Rostami S., Neri F., *A Fast Hypervolume Driven Selection Mechanism for Many-Objective Optimisation Problems*. *Swarm and Evol. Comp.* 34, 50–67, 2017. [DOI](#)
- [10] Liu Z.Z. et al., *A Many-Objective Evolutionary Algorithm with Angle-Based Selection and Shift-Based Density Estimation*. *Information Sciences* 509, 400–419, 2020. [SemanticsScholarLink](#)
- [11] Farina M., Amato P. *Fuzzy optimality and evolutionary multiobjective optimization*. *Evolutionary Multi-Criterion Optimization EMO2003*, 58–72, 2003. [DOI](#)
- [12] Köppen M. et al. *Fuzzy-Pareto-Dominance and its Application in Evolutionary Multi-objective Optimization*. *Evolutionary Multi-Criterion Optimization*, LN in *Comp. Science* 3410, 399–412, 2005. [DOI](#)
- [13] Silva R.C., Yamakami A. *Definition of Fuzzy Pareto-Optimality by Using Possibility Theory*. *IFSA/EUSFLAT*, 1234–1239, 2009. [SemanticsScholarLink](#)
- [14] Farina M., Amato P. *On the optimal solution definition for many-criteria optimization problems*. *NAFIPS-FLINT*, 233–238, 2002. [DOI](#)
- [15] Neghină M., Vințan L., *Apparent Front Ranking: novel population ranking method for genetic multi-objective algorithms*. *Proc. of the Romanian Academy A*, 21 (2), 179–186, 2020. [AcadRoLink](#)
- [16] Neghină M. et al., *A competitive new multi-objective optimization genetic algorithm based on Apparent Front Ranking*, *Engineering Applications of Artificial Intelligence*, Vol. 132(107870), 2024. [DOI](#)
- [17] Pătrăușanu A. et al., *A Systematic Review of Multi-Objective Evolutionary Algorithms Optimization Frameworks Processes* 12(5):869, 2024. [DOI](#)
- [18] Saaty T.L. *The Analytical Hierarchy Process*. McGraw-Hill, New York, 1980. ISBN-13: 9780070543713
- [19] Kendall M.G.. *Rank Correlation Methods*. Hafner Publishing Co., New York, 1955. ISBN-13: 9780852641996
- [20] Pearson K. *Note on regression and inheritance in the case of two parents*. *Proc. Royal Soc. London* 58, 1895. [DOI](#)
- [21] Fitzgibbon A. et al., *Direct least square fitting of ellipses*. *Pattern and Analysis Machine Intel.* 21.5, 476–480, 1999. [DOI](#)
- [22] Deb K. et al., *A fast and elitist multiobjective genetic algorithm: NSGA-II*. *Evol. Comp.* 6.2, 182–197, 2002. [DOI](#)
- [23] Zitzler E. et al., *SPEA2: Improving the strength pareto evolutionary algorithm for multiobjective optimization*., *EUROGEN*, 95–100, 2001. [DOI](#)

Analiza adaptabilității metodei de ordonare cu front aparent

Abstract: În contextul problemelor de optimizare multi-obiectiv cu ajutorul algoritmilor genetici, ordonarea indivizilor este o operațiune delicată și importantă. Introdusă în 2020, metoda de ordonare cu front aparent a arătat potențial, fiind inclusă cu succes într-un algoritm genetic. Această lucrare oferă o analiză mai amănunțită a flexibilității metodei AFR, cu informații despre evoluția rangurilor și a valorilor atunci când puterea funcției șablon a frontului aparent este modificată, prezentând corelații Kendall și Pearson excelente față de metodele ierarhice de referință pentru seturi de date cu până la 10 dimensiuni. Prin ajustarea puterii funcției șablon, cea mai bună corelație Kendall este mai mare de 0.79 față de toate metodele de referință, în timp ce cea mai bună corelație Pearson este peste 0.95 față de metoda AHP a lui Saaty și peste 0.75 față de metoda FPD medie a lui Köppen.

NEGHINĂ Mihai, Lecturer at the Department of Computer Science and Electrical Engineering, Faculty of Engineering, University “Lucian Blaga” of Sibiu, Romania
e-mail: mihai.neghina@ulbsibiu.ro, phone: +40 740 027 369.