

RESEARCH ARTICLE

Alternative weighting schemes for fine-tuned extended similarity indices

Kenneth López Pérez¹ | Anita Rácz² | Dávid Bajusz³ | Camila Gonzalez¹ |
Károly Héberger²  | Ramón Alain Miranda-Quintana¹

¹Department of Chemistry and Quantum Theory Project, University of Florida, Gainesville, Florida, USA

²Plasma Chemistry Research Group, HUN-REN Research Centre for Natural Sciences, Budapest, Hungary

³Medicinal Chemistry Research Group, HUN-REN Research Centre for Natural Sciences, Budapest, Hungary

Correspondence

Ramón Alain Miranda-Quintana,
Department of Chemistry and Quantum Theory Project, University of Florida,
Gainesville, Florida 32611, USA.
Email: quintana@chem.ufl.edu

Anita Rácz, Plasma Chemistry Research Group, HUN-REN Research Centre for Natural Sciences, Magyar tudósok krt. 2, 1117 Budapest, Hungary.
Email: racz.anita@ttk.hu

Funding information

National Research Development and Innovation Office, Grant/Award Numbers: K134260, FK146063; Hungarian Academy of Sciences: János Bolyai Research Scholarship; National Research Development and Innovation Fund, Grant/Award Number: ÚNKP-23-5 New National Excellence Program; National Institute of General Medical Sciences of the National Institutes of Health, Grant/Award Number: R35GM150620

Abstract

Extended similarity indices (i.e., generalization of pairwise similarity) have recently gained importance because of their simplicity, fast computation, and superiority in tasks like diversity picking. However, they operate with several meta parameters that should be optimized. Earlier, we extended the binary similarity indices to “discrete non-binary” and “continuous” data; now we continue with introducing and comparing multiple weighting functions. As a case study, the similarity of CYP enzyme inhibitors (4016 molecules after curation) was characterized by their extended similarities, based on 2D descriptors, MACCS and Morgan fingerprints. A statistical workflow based on sum of ranking differences (SRD) and analysis of variance (ANOVA) was used for finding the optimal weight function(s). Overall, the best weighting function is the fraction (“frac”), which corresponds to the principle of parsimony. Optimal extended similarity indices were also found, and their differences are revealed across different data sets. We intend this work to be a guideline for users of extended similarity indices regarding the various weighting options available. Source code for the calculations is available at <https://github.com/mqcomplab/MultipleComparisons>.

KEYWORDS

ANOVA, molecular similarity, sum of ranking differences, threshold setting, weighting

Kenneth López Pérez, Anita Rácz, and Dávid Bajusz contributed equally to this work.

This paper is intended for the Conferentia Chemometrica 2023 special issue.

This is an open access article under the terms of the [Creative Commons Attribution](https://creativecommons.org/licenses/by/4.0/) License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited.

© 2024 The Authors. *Journal of Chemometrics* published by John Wiley & Sons Ltd.

1 | INTRODUCTION

The application of molecular fingerprints (binary strings of molecular features with 0/1 denoting absence/presence) and similarity measures has been the backbone of cheminformatics for decades. Ligand similarity searches have constituted a high-throughput alternative for computational drug design and virtual screening,^{1,2} and they have remained tightly integrated even into today's AI-enhanced, structure-based workflows,³ as well as generative models.⁴ After our 2015 work that provided a statistical basis for the decade-long collective habit of preferring the Tanimoto coefficient,⁵ we have thoroughly investigated a large number of alternatives collected from diverse scientific fields by Todeschini et al.,⁶ regarding their use in metabolomics,⁷ molecular modeling,⁸ and even food science.⁹ In certain cases, we could establish a perfect consistency between two or more metrics even with analytical methods.¹⁰

As cheminformatics and related fields move gradually into the domain of big data, for example, by handling molecular datasets in and above the billion regime,^{11,12} a deeply rooted bottleneck of similarity metrics gets revealed. By their original definition, similarity measures are calculated between exactly two entities, resulting in a disadvantageous, quadratic scaling of computational demand if we want to characterize (cluster, etc.) a large dataset by pairwise similarity calculations. To rectify this, we have recently developed and introduced an extension of the mathematical framework of similarity calculations, allowing for calculating a single similarity value for an arbitrarily large group of objects.^{13,14} Without reiterating the whole framework here, we briefly note that, for binary fingerprints, this extension is based on the generalization of the terms

- a : the number of coincident 1's (common *on* bits)
- d : the number of coincident 0's (common *off* bits)
- and $b + c$: the number of 0-1 and 1-0 pairs

to the following, more general terms:

- 1-similarity counters: number of bit positions where 1's occur over a coincidence threshold γ
- 0-similarity counters: number of bit positions where 0's occur over a coincidence threshold γ
- dissimilarity counters: number of bit positions where neither 1's, nor 0's occur over the coincidence threshold γ .

An important feature of the resulting *extended* or *n-ary similarity metrics* is that they realize a true generalization of the well-known pairwise similarities, as they provide the same results in the $n = 2$ case as the “traditional” pairwise definitions. As a further extension of the framework, we have defined *extended many-item*¹⁵ and *extended continuous similarities*¹⁶ to cases of discrete non-binary and continuous data, respectively. Most importantly, the usage of extended similarities displays a much more advantageous, linear scaling of computational demand versus library size, during common tasks such as diversity selection or clustering.¹⁴ Recently, the various flavors of extended similarity metrics have been implemented into several workflows, including data visualization,¹⁷ activity cliff detection,¹⁸ or molecular dynamics trajectory sampling.¹⁹

Besides their many advantages, extended similarities involve the optimization of several parameters that are absent from the traditional pairwise definitions. One such parameter is the definition of the coincidence threshold: when comparing five objects and having a bit position with three co-occurring 1's and two 0's, should we consider that a 1-similarity counter, or a dissimilarity counter? Similarly, should we distinguish the “strength” of similarity counters by assigning a larger weight to a bit position with five co-occurring 1's versus the bit position detailed above (three 1's and two 0's)? While we have thoroughly investigated the first question in our earlier papers, there is much left to be explored regarding the use of different weighting schemes and their effect on the outcome of extended similarity calculations.

Here, we will introduce multiple weighting schemes and compare their usage in a common scenario of differentiating two groups of active versus inactive molecules against the important drug (anti)target, CYP 2C9.²⁰ For continuity, we use the same dataset as in our recent work,¹⁶ but use different molecular representations to cover the binary and continuous extended similarity metrics.

2 | METHODS

2.1 | Extended similarity

The key insight behind the extended similarity indices is that a condensed vector of the whole dataset is sufficient to quantify the similarity of the set. That is, given N molecules represented by vectors with M variables, we will arrange them in a $N \times M$ matrix. Then, we just need to calculate the sum of each of the columns of this matrix, thus generating a vector $\Sigma = [\sigma_1, \sigma_2, \dots, \sigma_M]$. In order to recover the information about the 1-, 0-, and dissimilarity counters (a , d , $b + c$) we need to calculate the indicator $\Delta_k = |2 \cdot \sigma_k - N|$, which quantifies the “agreement” between elements in the k th column (variable), and a coincidence threshold, γ , that indicates up to which point we count the column as contributing to the similarity or dissimilarity of the set. However, even with this classification we still need to distinguish between cases with partial coincidence of elements in a column. For that, we need to define weight functions, that penalize the partial coincidence over the similar columns, f_s , and the dissimilar ones f_d . These functions only need to obey very general conditions, namely,

$$f_s(N) = 1; x > y \Rightarrow f_s(x) > f_s(y) \quad (1)$$

$$f_d(N \bmod 2) = 1; x > y \Rightarrow f_d(x) < f_d(y) \quad (2)$$

Until now, all the extended similarity applications have only used a very particular form of these functions, with a simple linear dependency:

$$f_s(x) = \frac{x}{N}; f_d(x) = 1 - \frac{x - N \bmod 2}{N} \quad (3)$$

Here, we explore the effect of changing these weights in non-linear ways, as summarized in Table 1. (For an example calculation of the eJT similarity index, see Appendix A1).

2.2 | Dataset

A large dataset of cytochrome P450 (CYP) 2C9 ligands from Pubchem Bioassay (AID 1851) was used to compare the different weighting schemes,²¹ to provide continuity with our earlier work.¹⁶ Cytochrome P450 enzymes (CYP) are important mediators of drug metabolism, therefore generally important anti-targets in drug design: consequently, many datapoints of CYP bioactivity have been deposited into public databases, and CYP enzymes are likewise popular targets of QSAR and machine learning studies.²² In total, 12,161 molecules were applied after data curation: 4016 inhibitors with a potency of 10 μ M or better (actives) and 8145 inactive species. MACCS²³ and Morgan fingerprints (radius: 4, length: 1024)²⁴ were generated with RDKit,²⁵ while the Dragon 7 software was used for the calculation of 2D

TABLE 1 Weighting schemes introduced and compared in this work (for reference, we include the original, fractional weighting scheme, labeled “Frac” from here on).

$f_s = \left(\frac{x}{N}\right)^\alpha$ and $f_d = 1 - \left(\frac{x - N \bmod 2}{N}\right)^\alpha$			
Order, α	Notation	Order, α	Notation
1	Frac	$\left(\frac{x}{N}\right)^{\alpha=1}$	$1 - \left(\frac{x - N \bmod 2}{N}\right)^{\alpha=1}$
2	Poly_2	−2	Poly_2
4	Poly_4	−4	Poly_M4
8	Poly_8	−8	Poly_M8
16	Poly_16	−16	Poly_M16

descriptors.²⁶ Highly correlated variables (above 0.997) and constant variables were excluded from the sets.²⁷ The details and descriptions of the different descriptor sets can be found in the Dragon software manual.

2.3 | Sum of ranking differences

Sum of ranking differences, abbreviated as SRD, is simple and intuitive, a non-parametric method that is particularly suitable to compare variables, methods, models, items, and so on (see Bajusz et al.⁵ and references therein). It can be used to numerical and ordinal data alike. Generally, an SRD input matrix is arranged so: variables (alternatives) in the columns and objects (samples, molecules, criteria, etc.) in the rows. SRD scores are calculated as:

$$\text{SRD}_j = \sum_{i=1}^N |\text{rank}(x_{ij}) - \text{rank}(\rho_i)|$$

where $\text{rank}(x_{ij})$ denotes the rank of the i th row of the j th variable, and ρ represents the gold standard. If the ρ is not known beforehand, it can be aggregated as the row mean (median if outliers are present), minimum for error-like quantities, or maximum for correlation- or accuracy-like ones. The variable with the smallest SRD score is the best one. Two validation options are included in the SRD algorithm: randomization test and cross-validation. SRD has been proven to be a multicriteria decision making tool.

3 | RESULTS AND DISCUSSION

First, similarity calculations were carried out for the three separate datasets (2D descriptors, MACCS and Morgan fingerprint) for active and total groups. Our major assumption was that the better similarity metrics can provide higher similarity values for the actives and bigger differences between the similarity values of the "active" and "total" groups. Some restrictions had to be made before the evaluation of the similarity values in each case study. Based on previous assumptions, we have used only the non-weighted version of extended similarity values for the further evaluations. While in the case of the 2D descriptor dataset, the coincidence threshold limit was set to minimum to avoid the combinatorial "explosion" in the further ANOVA evaluations, for the MACCS and Morgan fingerprints, similarity values were calculated for each coincidence threshold limit. SRD was used for the comparison of the similarity metrics in the case of each weighting function for the active group. The results of SRD were channeled into a factorial ANOVA, in which the weighting functions (9) and similarity metrics (16) were used as factors along with the different datasets (3). The workflow of the three datasets will be discussed in the next sections.

In the case of 2D descriptors for the active set, the basic SRD input dataset contained 16 similarity metrics in the columns, and 14 different 2D descriptor sets in the rows. The dimensions of the input dataset were the same for each weighting functions. Thus, the SRD calculations were carried out for the nine different weighting functions and the cross-validated SRD values were united for the further analysis.

In the case of both MACCS and Morgan fingerprint datasets for actives, the similarity metrics were in the columns, and the similarity values for each coincidence threshold limit were in the rows. It resulted in nine rows and 16 columns. Like in the case of the 2D descriptors set, in total nine SRD calculations were carried out with the same settings and further analyzed by ANOVA.

As SRD is in % scale for each case study, we have merged together the results for the factorial ANOVA.

In the factorial ANOVA analysis, our aim was to find the differences and best solutions between the weighting functions for the different dataset alternatives, such as continuous (2D descriptors) and binary vectors (fingerprints). Therefore, at first, we show the comparison of the weighting functions for the SRD results of the three datasets in Figure 1.

The best weighting function depends on the used dataset, with significant differences detected in the factor of the weighting functions (at $\alpha = 0.05$). While MACCS had the lowest (best) SRD values in total, the binary datasets show almost the same pattern. Thus, the positive orders of the weighting functions are optimal for the continuous descriptor sets (especially Poly_8 and Poly_16), while the negative orders and the fraction weighting functions are better for the binary datasets.

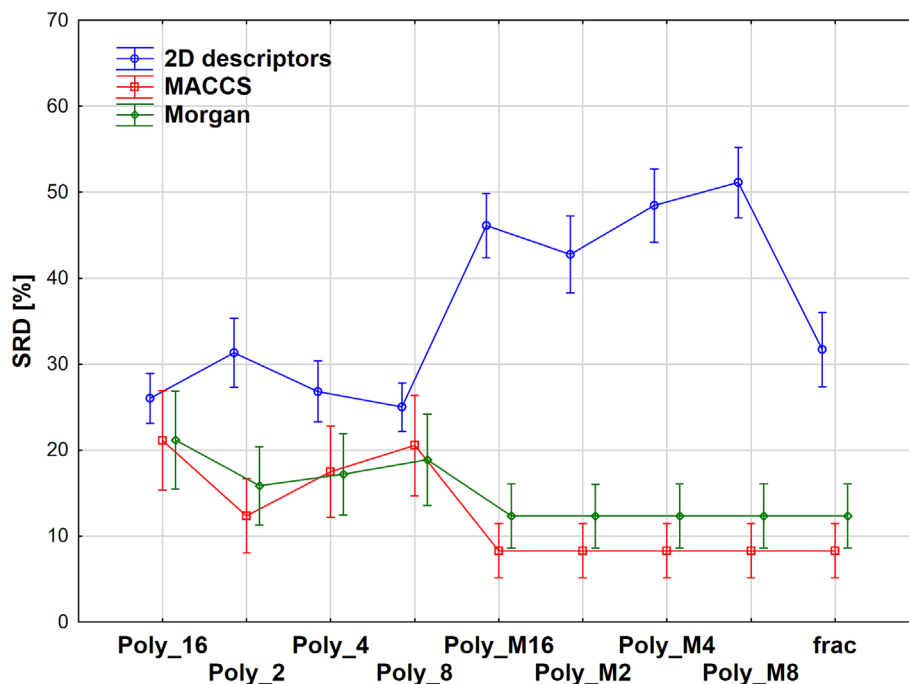


FIGURE 1 SRD values (%) plotted for the different weighting functions (for notations, see Table 1).

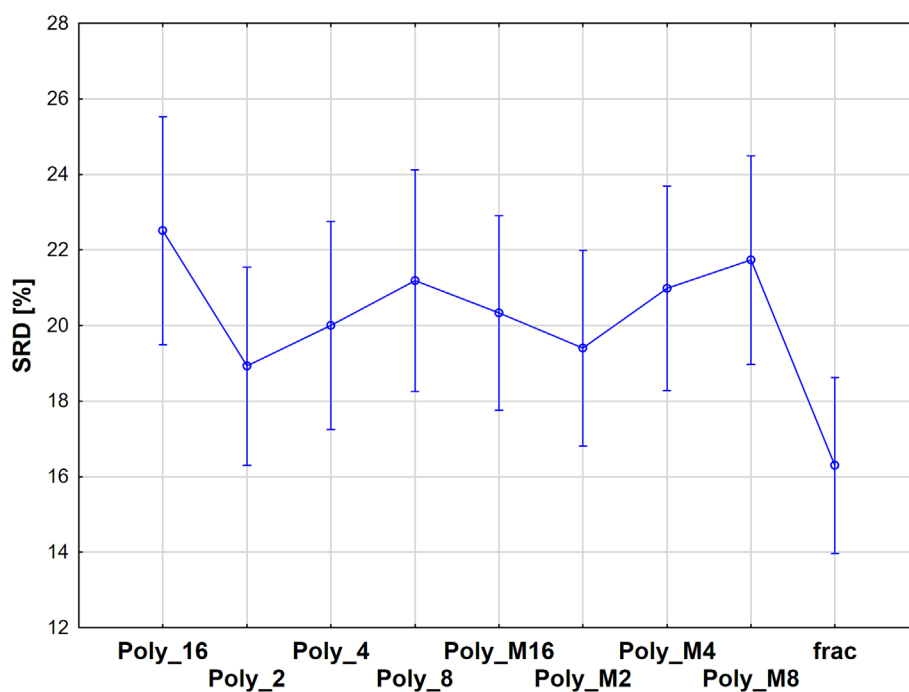


FIGURE 2 SRD values (%) based on the weighting functions (for notations, see Table 1).

If we compare the SRD values for the actives together for the three datasets and compare the weighting functions, we can see the most optimal ones overall in Figure 2. The fractional weighting function looks like the best option, but it is biased toward the binary datasets. Although the confidence intervals (95%) are wider in this case, the weighting functions still have significant impact on the outcome.

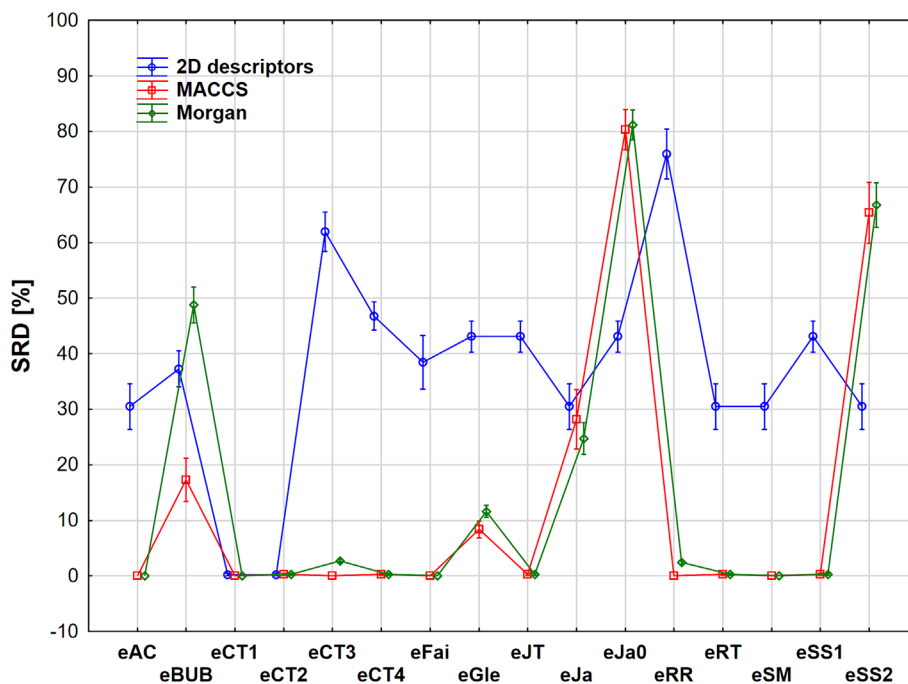


FIGURE 3 SRD values based on the similarity metrics and the dataset types (for abbreviations, see tab. 1 in Miranda-Quintana et al.¹³)

From the similarity metrics point of view, we can safely say, based on Figure 3, that the current metrics are generally worse for continuous sets, than the binary ones. The extended similarity index denoted by eJa0 is close to random or displays even reverse ranking in the case of binary datasets, and interestingly, eBUB produces dramatically different SRD scores for MACCS and Morgan fingerprints. However, the patterns for MACCS and Morgan datasets are still very similar. SRD scores for many extended similarity indices are the same or closely resemble the hypothetical best ranking. These datasets are in concordance with the previous results (cf., fig. 5 in Miranda-Quintana et al.¹³) that eCT1 $\left(s_{eCT1} = \frac{\ln(1 + \sum_{1-s} f_s(\Delta_{n(k)}) C_{n(k)}) + \sum_{0-s} f_s(\Delta_{n(k)}) C_{n(k)}}{\ln(1 + \sum_s C_{n(k)} + \sum_d C_{n(k)})} \right)$ and eCT2 $\left(s_{eCT2} = \frac{\ln(1 + \sum_s f_s(\Delta_{n(k)}) C_{n(k)} + \sum_d f_d(\Delta_{n(k)}) C_{n(k)}) - \ln(1 + \sum_d f_d(\Delta_{n(k)}) C_{n(k)})}{\ln(1 + \sum_s C_{n(k)} + \sum_d C_{n(k)})} \right)$ indices are amongst the best ones for all the three datasets. We can extend this now with eAC, eRT, and eSM, which are still suitable for the two binary case studies and acceptable for the continuous dataset as well. In the case of 2D descriptors, the list can be extended with eJa and eSS2, too.

Next, we have analyzed the occurrences of the weighting functions and similarity indices with regard to the differences between the “active” and “total” group similarities ($\text{Similarity}_{\text{Active}} - \text{Similarity}_{\text{Total}}$) for each dataset. This way we could find out what are the optimal settings to maximize the difference between the similarities of the active and total sets. Figure 4 shows the percentage of occurrences of the weighting functions above the cut-off limit of 0.10, while Figure 5 shows the percentage of occurrences of the similarity metrics above this cut-off limit.

As shown, the top weighting functions in this case were the Poly_8 and Poly_16 variants. These two were the most frequent in all case studies regarding the Active-Total similarity values. Poly_8 and especially Poly_16 can be good options if the aim of the study is to maximize the difference between the active set and the whole database (actives and inactives). On the other hand, the Poly_2 weighting function is also a good choice for continuous descriptors as input data.

We have evaluated the frequencies of the similarity metrics above the 0.10 cut-off value in the same way. Figure 5 shows the distribution of the similarity metrics in the three case studies for the Active-Total values. The result shows that eCT3 $\left(s_{eCT3} = \frac{\ln(1 + \sum_{1-s} f_s(\Delta_{n(k)}) C_{n(k)})}{\ln(1 + \sum_s C_{n(k)} + \sum_d C_{n(k)})} \right)$ or eCT4 $\left(s_{eCT4} = \frac{\ln(1 + \sum_{1-s} f_s(\Delta_{n(k)}) C_{n(k)})}{\ln(1 + \sum_{1-s} C_{n(k)} + \sum_d C_{n(k)})} \right)$ can be an optimal choice for the similarity calculations especially in the case of binary datasets like MACCS or Morgan, where the aim is to determine the differences between the actives and the whole dataset of molecules. On the other hand, the indices denoted by eGle, eJT, eJa, and eSS1 are also good options for continuous variables in the input matrix.

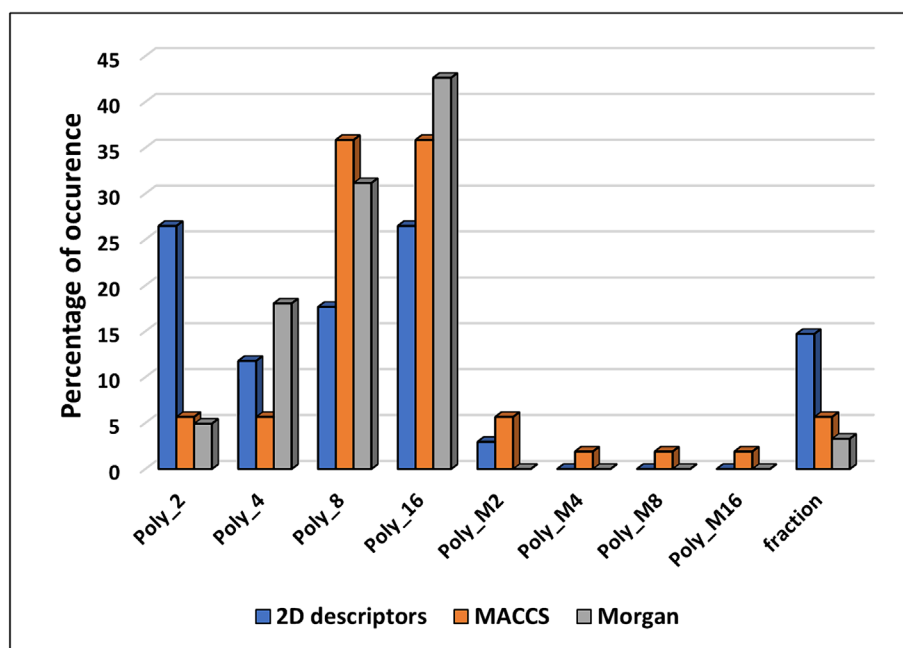


FIGURE 4 Percentage of occurrence of the weighting functions above 0.10 similarity values in the case of Active-Total data (for abbreviations, see Table 1).

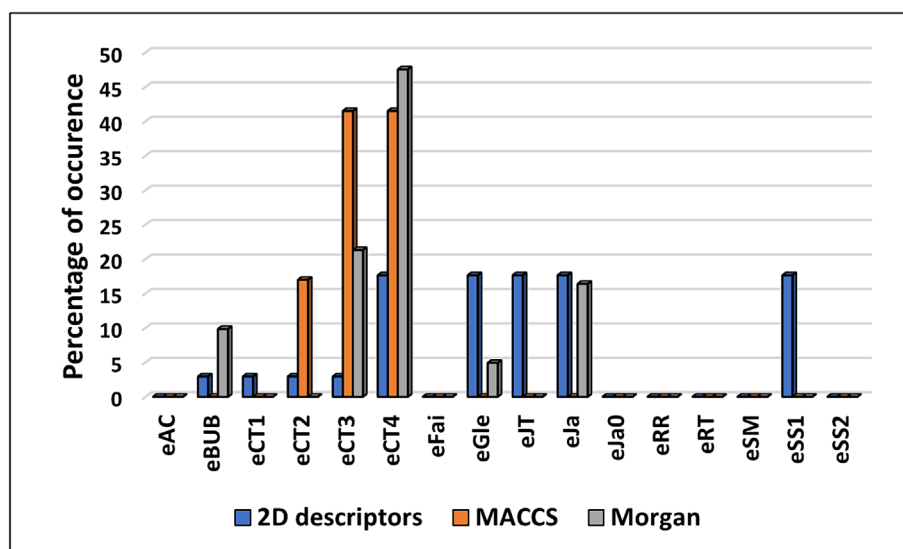


FIGURE 5 Percentage of occurrence of the similarity indices above 0.10 cut-off limit for the Active-Total similarity values of the three case studies.

4 | CONCLUSION

We have compared several weighting functions in the case of extended similarity metrics based on sum of ranking differences (SRD scores) combined with factorial ANOVA. The superiority of positive-power weighting functions was clear for the continuous descriptors, while negative-power alternative (and fraction) are better for binary variables in the input dataset. Interestingly, in the evaluation of Active-Total values (differences in the similarities of the active molecules vs. the whole dataset), positive weighting functions, especially Poly_8 and Poly_16 were among the best for all the case studies. Thus, weights can be optimized based on the aim of the application. From the extended similarity metrics point of view, we have seen that the most optimal ones are input dependent, while eCT1 and eCT2 are good choices for

the similarity calculation of the active sets in any case. On the other hand, eCT3 and eCT4 are very promising for binary case studies in both Actives and Active-Total evaluation, while eJa can be an optimal index for the continuous descriptors in both Actives and Active-Total evaluation. As the used weighting functions and similarity metrics had significant impact on the outcome, selecting the optimal ones in a case specific manner is of outmost importance in the similarity calculations stage of any cheminformatics related research. Source code for the calculations is available at <https://github.com/mqcomplab/MultipleComparisons>.

ACKNOWLEDGEMENTS

The authors received funding from the National Research Development and Innovation Office of Hungary under grant no. K134260 (KH) and FK146063 (DB), as well as the Hungarian Academy of Sciences: János Bolyai Research Scholarship: (AR). The work of AR was supported by the ÚNKP-23-5 New National Excellence Program of the Ministry for Culture and Innovation from the source of the National Research Development and Innovation Fund. RAMQ and KLP thank support from the National Institute of General Medical Sciences of the National Institutes of Health under award number R35GM150620.

CONFLICT OF INTEREST STATEMENT

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

DATA AVAILABILITY STATEMENT

The data that support the findings of this study are available from the corresponding author upon reasonable request.

ORCID

Károly Héberger  <https://orcid.org/0000-0003-0965-939X>

REFERENCES

1. Willett P. Similarity-based approaches to virtual screening. *Biochem Soc Trans.* 2003;31(3):603-606. doi:10.1042/bst0310603
2. Willett P. Similarity-based virtual screening using 2D fingerprints. *Drug Discov Today.* 2006;11(23-24):1046-1053. doi:10.1016/j.drudis.2006.10.005
3. Tropsha A, Popov KI, Wellnitz J, Maxfield T. Hit discovery using docking ENriched by GEnerative modeling (HIDDEN GEM): a novel computational workflow for accelerated virtual screening of ultra-large chemical libraries. *Molec Informatics.* 2024;43(1):e202300207. doi:10.1002/minf.202300207
4. Tibo A, He J, Janet JP, Nittinger E, Engkvist O. Exhaustive local chemical space exploration using a transformer model. 2023. doi:10.26434/chemrxiv-2023-v25xb
5. Bajusz D, Rácz A, Héberger K. Why is Tanimoto index an appropriate choice for fingerprint-based similarity calculations? *J Chem.* 2015; 7(1):20. doi:10.1186/s13321-015-0069-3
6. Todeschini R, Consonni V, Xiang H, Holliday J, Buscema M, Willett P. Similarity coefficients for binary chemoinformatics data: overview and extended comparison using simulated and real data sets. *J Chem Inf Model.* 2012;52(11):2884-2901. doi:10.1021/ci300261r
7. Rácz A, Andrić F, Bajusz D, Héberger K. Binary similarity measures for fingerprint analysis of qualitative metabolomic profiles. *Metabolomics.* 2018;14(3):29. doi:10.1007/s11306-018-1327-y
8. Rácz A, Bajusz D, Héberger K. Life beyond the Tanimoto coefficient: similarity measures for interaction fingerprints. *J Chem.* 2018; 10(1):48. doi:10.1186/s13321-018-0302-y
9. Gere A, Bajusz D, Biró B, Rácz A. Discrimination ability of assessors in check-all-that-apply tests: method and product development. *Foods.* 2021;10(5):1123. doi:10.3390/foods10051123
10. Miranda-Quintana RA, Bajusz D, Rácz A, Héberger K. Differential consistency analysis: which similarity measures can be applied in drug discovery? *Mol Inf.* 2021;40(7):2060017. doi:10.1002/minf.202060017
11. Bajusz D, Keserű GM. Maximizing the integration of virtual and experimental screening in hit discovery. *Expert Opin Drug Discovery.* 2022;17(6):629-640. doi:10.1080/17460441.2022.2085685
12. da Silva MMP, Guedes IA, Custódio FL, Krempser E, Dardenne LE. Deep learning strategies for enhanced molecular docking and virtual screening. 2023. doi:10.26434/chemrxiv-2023-zfv87-v2
13. Miranda-Quintana RA, Bajusz D, Rácz A, Héberger K. Extended similarity indices: the benefits of comparing more than two objects simultaneously. Part 1: theory and characteristics†. *J Chem.* 2021;13(1):32. doi:10.1186/s13321-021-00505-3
14. Miranda-Quintana RA, Rácz A, Bajusz D, Héberger K. Extended similarity indices: the benefits of comparing more than two objects simultaneously. Part 2: speed, consistency, diversity selection. *J Chem.* 2021;13(1):33. doi:10.1186/s13321-021-00504-4
15. Bajusz D, Miranda-Quintana RA, Rácz A, Héberger K. Extended many-item similarity indices for sets of nucleotide and protein sequences. *Comput Struct Biotechnol J.* 2021;19:3628-3639. doi:10.1016/j.csbj.2021.06.021

16. Rácz A, Dunn TB, Bajusz D, Kim TD, Miranda-Quintana RA, Héberger K. Extended continuous similarity indices: theory and application for QSAR descriptor selection. *J Comput Aided Mol des.* 2022;36(3):157-173. doi:10.1007/s10822-022-00444-7
17. López-Pérez K, López-López E, Medina-Franco JL, Miranda-Quintana RA. Sampling and mapping chemical space with extended similarity indices. *Molecules.* 2023;28(17):6333. doi:10.3390/molecules28176333
18. Dunn TB, López-López E, Kim TD, Medina-Franco JL, Miranda-Quintana RA. Exploring activity landscapes with extended similarity: is Tanimoto enough? *Molec Informatics.* 2023;42(7):2300056. doi:10.1002/minf.202300056
19. Rácz A, Mihalovits LM, Bajusz D, Héberger K, Miranda-Quintana RA. Molecular dynamics simulations and diversity selection by extended continuous similarity indices. *J Chem Inf Model.* 2022;62(14):3415-3425. doi:10.1021/acs.jcim.2c00433
20. Rácz A, Keserű GM. Large-scale evaluation of cytochrome P450 2C9 mediated drug interaction potential with machine learning-based consensus modeling. *J Comput Aided Mol des.* 2020;34(8):831-839. doi:10.1007/s10822-020-00308-y
21. National Center for Biotechnology Information. *PubChem database.* Source=NCGC, AID=1851.
22. Rácz A, Bajusz D, Miranda-Quintana RA, Héberger K. Machine learning models for classification tasks related to drug safety. *Mol Divers.* 2021;25(3):1409-1424. doi:10.1007/s11030-021-10239-x
23. Bajusz D, Rácz A, Héberger K. Chemical data formats, fingerprints, and other molecular descriptions for database analysis and searching. In: Chackalamannil S, Rotella DP, Ward SE, eds. *Comprehensive medicinal chemistry III.* Elsevier; 2017:329-378. doi:10.1016/B978-0-12-409547-2.12345-5
24. Rogers D, Hahn M. Extended-connectivity fingerprints. *J Chem Inf Model.* 2010;50(5):742-754. doi:10.1021/ci100050t
25. RDKit: open-source cheminformatics software. Accessed May 2, 2016. <http://rdkit.org/>
26. Mauri A, Consonni V, Pavan M, Todeschini R. Dragon software: an easy approach to molecular descriptor calculations. *MATCH Commun Math Comput Chem.* 2006;56:237-248. Dragon 7.0, Kode Cheminformatics. Dragon 70, Kode Cheminformatics.
27. Rácz A, Bajusz D, Héberger K. Interrelation limits in molecular descriptor preselection for QSAR/QSPR. *Mol Inf.* 2019;38(8-9):1800154. doi:10.1002/minf.201800154

How to cite this article: López Pérez K, Rácz A, Bajusz D, Gonzalez C, Héberger K, Miranda-Quintana RA. Alternative weighting schemes for fine-tuned extended similarity indices. *Journal of Chemometrics.* 2024;e3558. doi:10.1002/cem.3558

APPENDIX A

A.1 | Sample calculation of extended similarity indices

Let us start from a set of $N = 10$ fingerprints, each with $M = 6$ variables:

$$\begin{bmatrix} 1 & 1 & 1 & 1 & 0 & 0 \\ 0 & 1 & 1 & 1 & 0 & 1 \\ 0 & 1 & 0 & 0 & 1 & 1 \\ 1 & 0 & 0 & 0 & 1 & 0 \\ 1 & 1 & 1 & 0 & 1 & 0 \\ 0 & 1 & 1 & 0 & 0 & 1 \\ 0 & 0 & 1 & 1 & 1 & 1 \\ 0 & 1 & 0 & 0 & 1 & 0 \\ 1 & 1 & 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 0 & 0 & 1 \end{bmatrix} \quad (4)$$

The vector of column sums is then:

$$\Sigma = [4 \ 8 \ 5 \ 3 \ 6 \ 6] \quad (5)$$

The next step is to calculate the corresponding Δ_k indicators:

$$\begin{aligned}\Delta_1 &= |2 \cdot \sigma_1 - 10| = |2 \cdot 4 - 10| = 2 \\ \Delta_2 &= |2 \cdot \sigma_2 - 10| = |2 \cdot 8 - 10| = 6 \\ \Delta_3 &= |2 \cdot \sigma_3 - 10| = |2 \cdot 5 - 10| = 0 \\ \Delta_4 &= |2 \cdot \sigma_4 - 10| = |2 \cdot 3 - 10| = 4 \\ \Delta_5 &= |2 \cdot \sigma_5 - 10| = |2 \cdot 6 - 10| = 2 \\ \Delta_6 &= |2 \cdot \sigma_6 - 10| = |2 \cdot 6 - 10| = 2\end{aligned}\quad (6)$$

The coincidence threshold, γ , must be bounded by the following limits: $N \bmod 2 \leq \gamma \leq N - 1$. The mod function is necessary to distinguish between the cases in which we have and odd or even number of molecules to be compared, and reflects the maximum possible difference of different characters in any given column of the fingerprint matrix. For this particular example, we will choose the lower bound, that is: $\gamma = 10 \bmod 2 = 0$. Following this, the different columns are classified as follows:

$$\begin{aligned}\Delta_1 &= 2 > \gamma \rightarrow \text{sim} \\ \Delta_2 &= 6 > \gamma \rightarrow \text{sim} \\ \Delta_3 &= 0 = \gamma \rightarrow \text{dis} \\ \Delta_4 &= 4 > \gamma \rightarrow \text{sim} \\ \Delta_5 &= 2 > \gamma \rightarrow \text{sim} \\ \Delta_6 &= 2 > \gamma \rightarrow \text{sim}\end{aligned}\quad (7)$$

For more detail on the classification of the similarity columns:

$$\begin{aligned}\Delta_1 &= 2 \cdot 4 - 10 < 0 \rightarrow 0 - \text{sim} \\ \Delta_2 &= 2 \cdot 8 - 10 > 0 \rightarrow 1 - \text{sim} \\ \Delta_4 &= 2 \cdot 3 - 10 < 0 \rightarrow 0 - \text{sim} \\ \Delta_5 &= 2 \cdot 6 - 10 > 0 \rightarrow 1 - \text{sim} \\ \Delta_6 &= 2 \cdot 6 - 10 > 0 \rightarrow 1 - \text{sim}\end{aligned}\quad (8)$$

The final step is to calculate the weight functions, which for this example we will just take as the simple linear fractions shown in Equation (3). (Once again, note that the use of the mod function in the difference function is to guarantee that the maximum possible mismatch between fingerprints is assigned the same value independently of if we are considering odd or even numbers of molecules.) The weight values are

$$\begin{aligned}f_s(\Delta_1) &= \frac{2}{10} \\ f_s(\Delta_2) &= \frac{6}{10} \\ f_d(\Delta_3) &= 1 - \frac{0-0}{10} = 1 \\ f_s(\Delta_4) &= \frac{4}{10} \\ f_s(\Delta_5) &= \frac{2}{10} \\ f_s(\Delta_6) &= \frac{2}{10}\end{aligned}\quad (9)$$

With this, we have all the ingredients needed to calculate any of the extended similarity indices. For instance, the eJT index value is

$$eJT = \frac{\frac{6}{10} + \frac{2}{10} + \frac{2}{10}}{3 + 1} = 0.25 \quad (10)$$