

Correlation of Divergency: c_δ . Being Different in a Similar Way or Not

Johan F. Hoorn^{1,2,3,4}

¹ School of Design, The Hong Kong Polytechnic University, Hung Hom, Hong Kong SAR

² Dept. of Computing, The Hong Kong Polytechnic University, Hung Hom, Hong Kong SAR

³ Research Institute for Quantum Technology (RIQT), The Hong Kong Polytechnic University, Hung Hom, Hong Kong SAR

⁴ Dept. of Communication Science, Vrije University, Amsterdam, Netherlands

18 October 2025

Abstract

This paper introduces the correlation-of-divergency coefficient, c_δ , a custom statistical measure designed to quantify the similarity of internal divergence patterns between two groups of values. Unlike conventional correlation coefficients such as Pearson or Spearman, which assess the association between paired values, c_δ evaluates whether the way values differ within one group is mirrored in another. The method involves calculating, for each value, its divergence from all other values in its group, and then comparing these patterns across the two groups (e.g., human vs machine intelligence). The coefficient is normalised by the average root mean square divergence within each group, ensuring scale invariance. Potential applications of c_δ span quantum physics, where it can compare the spread of measurement outcomes between quantum systems, as well as fields such as genetics, ecology, psychometrics, manufacturing, machine learning, and social network analysis. The measure is particularly useful for benchmarking, clustering validation, and assessing the similarity of variability structures. While c_δ is not bounded between -1 and 1 and may be sensitive to outliers (but so is PMCC), it offers a new perspective for analysing internal variability and divergence. The article discusses the mathematical formulation, potential adaptations for complex data, and the interpretative considerations relevant to this alternative approach.

Keywords: correlation of divergency; variability comparison; dispersion patterns; quantum systems; statistical measure

1. Introduction

On many occasions, the comparison of internal patterns of divergence or variability between two sets of quantum states, observables, or measurement outcomes is highly desired. One may want to have a measure that could be used to compare the spread or variability of properties (such as energy levels, spin states, or other observables) within two ensembles of quantum states. For example, it may be interesting to determine whether two different quantum systems exhibit similar patterns of state dispersion or coherence. In studies of entanglement, one might want to compare the divergence patterns of measurement outcomes between entangled and non-entangled pairs. Developing a measure of divergency could help quantify whether the way outcomes differ within

one subsystem is mirrored in another, potentially providing insight into the structure of quantum correlations.

When analysing the effects of noise or decoherence on quantum systems, a correlation of divergence could be used to compare how the variability of quantum observables changes under different environmental conditions or noise models. In quantum information as well, such a measure might be applied to compare the internal variability of information content or uncertainty between two quantum channels, states, or protocols. Therefore, I have set out to devise a measure that could assist in benchmarking quantum simulators by comparing the spread of simulated outcomes to those of theoretical or experimental reference systems.

2. The c_δ coefficient

Correlation of divergency, c_δ , is a custom equation I developed to measure the association between patterns of divergence within two groups of values (x and y). It compares how each value in x differs from all other values in x , and does the same for y , then assesses whether these patterns of divergence are related.

To calculate the numerator (signal), for each x_i , sum the squared differences between x_i and all other x_j ($j \neq i$), and do the same for y_i . Take the square root of each sum (for magnitude), multiply the results for x_i and y_i , and sum over all i .

For the denominator (noise), calculate the average root mean square difference in x (across all pairs), and the same for y , then multiply these averages.

Next, the details for both numerator and denominator are worked out: to define the sums, for each x_i and each y_i :

$$D_{x,i} = \sqrt{\sum_{j \neq i}^n (x_i - x_j)^2}$$

$$D_{y,i} = \sqrt{\sum_{j \neq i}^n (y_i - y_j)^2}$$

In calculating the numerator (signal), multiply the D s for each i , then sum over all i :

$$\sum_{i=1}^n D_{x,i} \cdot D_{y,i} = \sum_{i=1}^n \left(\sqrt{\sum_{j \neq i}^n (x_i - x_j)^2} \cdot \sqrt{\sum_{j \neq i}^n (y_i - y_j)^2} \right)$$

For the denominator (noise), find the average root mean square difference for x and y :

$$\bar{D}_x = \frac{1}{n} \sum_{i=1}^n \sqrt{\frac{1}{n-1} \sum_{j \neq i} (x_i - x_j)^2}$$

$$\bar{D}_y = \frac{1}{n} \sum_{i=1}^n \sqrt{\frac{1}{n-1} \sum_{j \neq i} (y_i - y_j)^2}$$

Then put everything into the full equation:

$$c_\delta = \frac{\sum_{i=1}^n \left[\sqrt{\sum_{j \neq i} (x_i - x_j)^2} \cdot \sqrt{\sum_{j \neq i} (y_i - y_j)^2} \right]}{\left(\frac{1}{n} \sum_{i=1}^n \sqrt{\frac{1}{n-1} \sum_{j \neq i} (x_i - x_j)^2} \right) \left(\frac{1}{n} \sum_{i=1}^n \sqrt{\frac{1}{n-1} \sum_{j \neq i} (y_i - y_j)^2} \right)}$$

In the numerator, the correlation of divergency, c_δ , measures for each data point the magnitude of divergence from all other points in its group (using squared differences for robustness), for both x and y , multiplies these, and sums over all points. The denominator normalises by the average root mean square divergence in each group, making the measure scale-invariant and comparable across datasets.

As another option, one could also write absolute values instead. In certain cases, absolute differences (Gini mean difference) may be preferred, replacing squared differences and square roots with absolute values:

$$c_\delta = \frac{\sum_{i=1}^n \left[\sum_{j \neq i} |x_i - x_j| \cdot \sum_{j \neq i} |y_i - y_j| \right]}{\left(\frac{1}{n} \sum_{i=1}^n \frac{1}{n-1} \sum_{j \neq i} |x_i - x_j| \right) \left(\frac{1}{n} \sum_{i=1}^n \frac{1}{n-1} \sum_{j \neq i} |y_i - y_j| \right)}$$

A high value of the c_δ coefficient indicates that the patterns of divergence within x and y are similar, such that when a value in x is distant from others, the corresponding value in y is likewise distant (being different in the same way). Conversely, a low value suggests that the patterns of divergence are dissimilar or unrelated (being different in a different, perhaps unique, way).

3. Relation to other work

The c_δ coefficient differs from Pearson (1904) and Spearman (1904) correlations in several fundamental ways. Pearson's correlation measures the linear association between two variables, focusing on whether increases in one variable correspond to increases or decreases in the other. It is sensitive to the actual values and assumes interval data with a linear relationship. Spearman's correlation, on the other hand, assesses the monotonic relationship between two variables by

ranking the data and comparing the ranks, making it robust to non-linear associations and suitable for ordinal data.

The c_δ coefficient does not measure direct association between paired values. Instead, it compares the internal patterns of divergence or variability within each group. For each value in x and y , it calculates how much that value differs from all other values in its own group, then assesses whether these patterns of divergence are similar between the two groups. This means c_δ is concerned with whether the way values are spread out or dispersed within x mirrors the dispersion within y , rather than whether high values in x correspond to high values in y .

Unlike Pearson and Spearman, which are bounded between -1 and 1 and have clear interpretations regarding strength and direction of association, c_δ may not be bounded in this way, its value ranges between 0 and ∞ . Its scale and interpretation are different, as it is normalised by the average root mean square divergence within each group, making it scale-invariant but not directly comparable to standard correlation coefficients.

Pearson and Spearman focus on relationships between paired observations, while c_δ focuses on relationships between patterns of internal variability. This makes c_δ suitable for comparing the structure of dispersion, divergence, or variability between two datasets, rather than their direct association.

In other words, c_δ is a measure of similarity in the way two groups of values are internally divergent, whereas Pearson and Spearman measure association between paired values. The mathematical approach, interpretation, and applications of c_δ are distinct from those of the standard correlation coefficients.

4. High, low, zero: how far is far?

Magnitude of c_δ is indicative of (dis)similarity in divergence patterns. In general, for the c_δ coefficient, the further away from zero, the stronger the relationship between the divergence patterns of the two groups. Note that a negative c_δ is not possible with the standard squared difference divergence formula. Also note that the interpretation of its strength is not exactly the same as for standard correlation coefficients, and some caution is needed.

In principle, high c_δ means a strong similarity in divergence patterns, *irrespective of the research units that produced those patterns*. When a value in set X is far from others in set X , the corresponding value in Y (e.g., paired-samples, within-subjects) is also far from others in Y (similar divergence structure).

A completely diametrically opposed divergence structure will not be recognised by c_δ as ‘different,’ because c_δ does not produce negative values and because the *pattern* of divergence in sets X and Y will be the same, although coming from different sources. A completely inverse relationship in divergence patterns exemplifies a theoretical limitation of c_δ that practically may hardly occur. Let us pursue an example, supposing that a researcher obtains four groups of values:

Group X	2, 4, 6, 8	
Group Y	10, 12, 14, 16	(similar divergence pattern)
Group Y'	8, 6, 4, 2	(opposite divergence pattern will not be recognised as such)
Group Y''	11, 11, 13, 15	(random divergence pattern)

The values in Group X and Group Y are evenly spaced. The way each value differs from the others is the same in both groups: they show similar divergence, albeit at different scales, and so the c_δ will be high. Group X and Group Y' are a mirror image of each other. Where X increases, Y' decreases. Yet, opposite divergence will not be visible in the same high c_δ . Group Y'' does not follow the same pattern as X . Random divergence will yield a c_δ that moves closer to the zero.

Near zero c_δ suggests little to no relationship in divergence patterns: the divergencies differ. Be aware that if one group has no internal divergence (all values identical within set X and/or Y), c_δ remains undefined because there is no divergence to compare, which technically is not zero although the equation will render $0/0$. Cases in which c_δ becomes undefined are when at least one group is constant, all values being the same, irrespective of measurement scale. Thus, our anchoring point $c_\delta = 0$ *under the condition* that not all data points are identical within set X and/or Y . Although this technically is contradictory, $c_\delta = 0$ provides us with an abstract anchoring point that can be approximated in an actual sample. Closer to that abstract zero, then, means that at least one group shows near-absent divergence.

The value of c_δ will *never* be zero or close to it when X and Y contain the same set of numbers, even if those numbers are very spread out (i.e., have large divergence). Divergence vectors, or permutations thereof, will be identical. This is our second anchoring point. For data set X , we know that the highest possible c_δ similarity score that we can be certain about comes when $X = Y$, data being identical, rendering a sample-related c_{δ_max} for self-similarity. With that, we know the upper bound of c_δ for that particular data sample and we can rescale the observed value of c_δ proportionally between 0 and c_{δ_max} (i.e., $c_{\delta_observed} / c_{\delta_max}$).

Suppose that $X = \{2, 1, 1, 1, 1\}$ and $Y = \{2, 3, 4, 6, 9\}$. To obtain c_{δ_max} , calculate c_δ for $X = X$ and then for $Y = Y$. For $X = X$, $c_\delta = 8 / 1.44 \approx 5.56$ and for $Y = Y$, $c_\delta = 308.09 / 52.30 \approx 5.89$. Thus, the maximum c_δ that can be achieved in these data is $c_{\delta_max} = 5.89$, our sample-related found upper bound, which we resize to $c_{\delta_max} = 5.89 = 100\%$ by multiplying with ~ 16.98 .

For $X = \{2, 1, 1, 1, 1\}$ and $Y = \{2, 3, 4, 6, 9\}$, $c_\delta = 46.50 / 9.156 \approx 5.08$, proportionally rescaled by multiplying with $\sim 16.98 \approx 86\% \approx .86$ similar divergency structure found between sets X and Y . Instead of a theoretical range $[0-\infty]$, our empirically found range is $[0-1]$.

Suppose that $X = \{2, 3, 4, 6, 9\}$ and $Y = \{21, 32, 43, 65, 98\}$. For $X = X$, $c_\delta = 308.09 / 58.19 \approx 5.29$ and for $Y = Y$, $c_\delta = 37269.88 / 7035.13 \approx 5.30$, resized to $c_{\delta_max} = 5.30 = 100\%$ by multiplying with ~ 18.87 . For X compared to Y , $c_\delta = 3388.37 / 639.06 \approx 5.30 \approx c_{\delta_max} \approx 1$, concluding for near-identity of divergency patterns.

Rescaling c_δ between 0 and c_{δ_max} may be a good practice for comparing across different data sets or for interpretability. This makes c_δ analogous to a similar-divergency index bounded between 0 and 1 for any given sample. Our framework of interpretation would look like this:

- At least one group is constant, the c_δ value is undefined ($0/0$), because there is no divergence to compare
- One group has near-zero divergence, the c_δ value comes near zero, so there is little to no relationship in divergence
- Sets X and Y contain same numbers (spread out), a large c_δ value indicates high similarity in divergence pattern
- Set $X = Y$ (identical order), c_{δ_max} is the highest possible for that sample
- Sets X and Y are unrelated, both with some divergence. The c_δ value is intermediate, some similarity in divergence is observed

Nonetheless, one should use one's understanding of the data to interpret what a large or small c_δ means in context. For example, in some fields, even a moderate c_δ might be meaningful if variability is generally low (e.g., speed of light in near-vacuum).

Since the c_δ coefficient is not bounded on the high end (unlike Pearson's r , which is always between -1 and 1), we saw that interpreting what counts as a high or low value requires a different approach, such as proportional rescaling between 0 and c_{δ_max} . Another way would be empirical benchmarking through simulation or bootstrapping. For instance, one could generate Null distributions by calculating c_δ for many pairs of random or permuted datasets that have no relationship (cf. Group Y''). This gives us a distribution of c_δ values expected by chance. Then compare the observed values, and see where the observed c_δ falls within this distribution. If it is much higher (or lower) than most values from the null distribution, it is high (or low) in a more meaningful sense. Additionally, the analyst can report that the obtained c_δ is, for example, in the top 5% of values expected by chance (e.g., 99th percentile = "very high similarity"), which is somewhat analogous to a p -value.

Relative comparison would be another approach. If within one's domain, the analyst has multiple datasets or repeated experiments, c_δ values can be compared across them. The highest c_δ indicates the most similar divergence patterns, and the lowest c_δ indicates the most dissimilar patterns. The analyst also can rank c_δ values to see which pairs are most or least similar in their divergence structure.

5. Applications

The correlation of divergency c_δ can be employed to assess the similarity of deviating structure between two datasets. However, data may involve complex numbers and probabilistic distributions (quantum data in particular often do), so the c_δ coefficient may need to be adapted to handle these appropriately, for example, by considering magnitudes or expectation values. This is work to be taken up in the future.

Interpretation should be made with care, as always, as systems can exhibit non-classical correlations and behaviours that differ from classical datasets (cf. quantum). Nevertheless, while this measure is not standard in physics, it could provide a novel way to quantify and compare patterns of divergence, variability, or dissimilarity between quantum systems, states, or measurement outcomes. Its meaningfulness would depend on the specific context and the nature of the data being analysed, which is beyond the reach of this here first attempt.

Beyond quantum physics, the c_δ coefficient may be used to determine whether the dispersion of gene expression levels in two distinct tissues, or the spread of test scores in two separate cohorts, is comparable. In addition, correlation of divergency is applicable to clustering validation, as it enables the comparison of the internal cohesion or dispersion of clusters identified through clustering analyses.

In the context of evolutionary biology, where x and y represent measures of genetic divergence in two species, this equation facilitates the evaluation of whether patterns of divergence are analogous. For instance, it can be used to assess whether the same pairs of individuals (e.g., mother-child) exhibit similar degrees of divergence in both species (e.g., humans vs apes). Similarly, in ecological studies, the measure may be applied to compare patterns of trait divergence between populations or communities.

Within psychometrics, if x and y denote scores from different tests or experimental conditions, this measure can be used to determine whether inter-individual differences are consistent across assessments. In manufacturing quality control, it enables the comparison of the variability of measurements obtained from different machines or production batches.

In the field of machine learning, correlation of divergency may serve as a feature to quantify the similarity of internal variance structures between two sets of features or samples, thereby supporting advanced data analysis and model development. Furthermore, in social network analysis, where x and y represent distances or dissimilarities in two distinct networks, this equation provides a means to compare the overall structure of relationships between the networks.

There are several limitations and considerations associated with the c_δ measure. Firstly, it is not a standard correlation coefficient, and as such, its scale and interpretation may differ from those of Pearson or Spearman correlations. Although the denominator serves to normalise for scale, the measure itself may not be bounded between -1 and 1. Just like Pearson's r , the measure is sensitive to outliers (Hoorn & Ho, 2025), particularly when squared differences are employed. Both issues need further investigation.

Hence, the correlation-of-divergency equation can be used to quantify and compare the internal patterns of divergence, such as spread, variability, or dissimilarity, between two groups of values in any context where the analyst deems such comparison meaningful.

6. Conclusions/Discussion

This paper introduces a new statistical measure, the correlation-of-divergency coefficient (c_δ), which is designed to quantify the similarity of internal divergence patterns between two groups of values. Unlike traditional correlation coefficients such as Pearson or Spearman, which assess direct association between paired values, c_δ compares the internal variability structures of two datasets. The measure is original and addresses a gap in statistical methodology, as it focuses on comparing divergence patterns rather than direct associations. I have outlined a wide range of possible applications, including quantum physics, genetics, machine learning, and social network analysis. A mathematical formulation is presented, and both squared and absolute difference variants are discussed.

There are, however, several caveats that should be considered. Firstly, the c_δ coefficient is not naturally bounded between -1 and 1, which makes interpretation less intuitive for practitioners accustomed to conventional measures. I proposed to rescale c_δ between zero and c_{δ_max} for interpretability, but this approach is sample-dependent and does not provide a universal scale. The maximum value is not theoretically fixed but is empirically determined for each dataset, which may lead to inconsistent comparisons across studies.

The use of squared differences in the calculation makes c_δ highly sensitive to outliers, potentially exaggerating divergence patterns due to a single extreme value. This remains an issue, especially in real-world data where outliers are common. As of yet, I did not discuss alternatives, such as trimmed means or winsorisation, nor is there any analysis yet of how the measure behaves under contaminated data.

The measure becomes undefined when one group has no internal divergence, that is, when all values are identical. It is a practical issue that could arise, particularly in small samples or highly homogeneous data. The interpretation of near-zero c_δ is somewhat ambiguous. It means little to no relationship in divergence, but there is no guidance yet on thresholds or statistical significance.

Another conceptual limitation is that the c_δ measure cannot take negative values, even when divergence patterns are perfectly opposed, such as when one group increases while the other decreases. This means c_δ cannot distinguish between similar and inverse divergence structures, resulting in a loss of information. In cases where the divergence pattern is mirrored, c_δ will still be high, failing to capture the directionality of divergence.

Be aware that as is, the equations use summation and square root notation that could be clarified, as the order of operations is not always explicit. There is a potential for double counting in the calculation of divergence, since each pairwise difference is considered for each i and j not equal to i . Depending on implementation, this could lead to double counting or inconsistent normalisation.

Comparability across datasets is another concern. Since c_{δ_max} is calculated for each dataset, comparisons across datasets with different sizes or distributions may be misleading. There is no theoretical justification for this scaling factor, and it may not be invariant under transformations or permutations. Unlike z -scores or other normalised measures, c_δ does not have a standardised distribution under the null hypothesis, which makes statistical inference difficult.

The paper does not provide p -values, confidence intervals, or a formal method for hypothesis testing. While empirical benchmarking via simulation is suggested, there is no analytical null distribution, making it difficult for now to assess the significance of observed c_δ values. As said, adaptation is needed for complex or probabilistic data, such as quantum states, but I did not provide a concrete method or examples yet, which limits immediate applicability in fields where such data are common.

There is some overlap with existing measures. The absolute difference variant of c_δ is related to the Gini mean difference, but I did not discuss this connection or potential redundancy. I offered little discussion of how c_δ compares to other measures of dispersion similarity, such as the coefficient of variation or entropy-based measures.

This paper provides only a few small-scale examples and does not demonstrate the measure on real or simulated data. There is no assessment of behaviour under different scenarios, such as skewed distributions or multimodal data, which limits practical understanding. All of this is future work.

In sum, the correlation-of-divergency coefficient (c_δ) may be an interesting and potentially useful measure for comparing internal divergence patterns between datasets. However, its practical utility is limited by several caveats not resolved yet. It is not bounded, which makes interpretation and comparison difficult. It is sensitive to outliers, but so is PMCC, and undefined for zero-divergency cases. It cannot distinguish between similar and inverse divergence patterns. There is no standardisation or statistical inference framework. The mathematical formulation could be clarified, and more empirical examples are needed. Further development is required to address these issues, including robustification, standardisation, and formal statistical inference. Until then, c_δ should be used with caution and only as a supplementary measure alongside established statistics.

Acknowledgements

This study was supported by the Research Grants Council (grant number T43-518/24-N) under the University Grants Committee, Hong Kong Special Administrative Region Government.

References

- Hoorn, J. F., & Ho, J. K. W. (2025). A test statistic, h^* , for outlier analysis. *arXiv:stat.ME*, 2508.06792v1, 1-35. Available from <https://arxiv.org/abs/2508.06792>
- Pearson, K. (1904). *On the theory of contingency and its relation to association and normal correlation*. London: Cambridge University.
- Spearman, C. (1904). The proof and measurement of association between two things. *The American Journal of Psychology*, 15(1), 72-101. doi: 10.2307/1412159