

CCREPE: Compositionality Corrected by PERmutation and RENormalization

Emma Schwager, George Weingart, Craig Bielski, Curtis Huttenhower

October 24, 2025

Contents

1	Introduction	1
2	ccrepe	2
2.1	General functionality	2
2.2	Arguments	2
2.3	Output	4
2.4	Usage	4
2.5	Example 1	4
2.6	Example 2	6
2.7	Example 3	8
2.8	Example 4	10
2.9	Example 5	11
3	nc.score	12
3.1	General Functionality	12
3.2	Arguments	12
3.3	Output	13
3.4	Usage	13
3.5	Example 1	13
3.6	Example 2	14
3.7	Example 3	15
4	References	16

1 Introduction

ccrepe is a package for analysis of sparse compositional data. Specifically, it determines the significance of association between features in a composition, using any similarity measure (e.g. Pearson correlation, Spearman correlation, etc.) The CCREPE methodology stands

for Compositionality Corrected by Renormalization and Permutation, as detailed below. The package also provides a novel similarity measure, the N-dimensional checkerboard score (NC-score), particularly appropriate to compositions derived from microbial community sequencing data. This results in p-values and false discovery rate q-values corrected for the effects of compositionality. The package contains two functions `ccrepe` and `nc.score` and is maintained by the Huttenhower lab (ccrepe-users@googlegroups.com).

2 ccrepe

`ccrepe` is the main package function. It calculates compositionality-corrected p-values and q-values for a user-selected similarity measure, operating on either one or two input matrices. If given one matrix, all features (columns) in the matrix are compared to each other using the selected similarity measure. If given two matrices, each feature in the first are compared against all features in the second.

2.1 General functionality

Compositional data induces spurious correlations between features due to the nonindependence of values that must sum to a fixed total. CCREPE abrogates this when determining the significance of a similarity measure for each feature pair using two main steps, permutation/renormalization and bootstrapping. First, given two features to compare, CCREPE generates a null distribution of the similarity expected just due to compositionality by iteratively permuting one feature, renormalizing all samples in the composition to their previous sum, and computing the resulting similarity measures. Second, CCREPE bootstraps over sample subsets in order to assess confidence in the "true" similarity measure. Finally, the two resulting distributions are compared using a pooled-variance Z-test to give a compositionality-corrected p-value. False discovery rate q-values are additionally calculated using the Benjamin-Hochberg-Yekutieli procedure. For greater detail, see [Faust et al. \[2012\]](#) and [Schwager and Colleagues](#).

CCREPE employs several filtering steps before the data are processed. It removes any missing subjects using `na.omit`: in the two dataset case, any subjects missing in *either* dataset will be removed. Any subjects or features which are all zero are removed as well: an all-zero subject cannot be normalized (its sum is 0) and an all-zero feature has standard deviation 0 (in addition to being uninteresting biologically).

2.2 Arguments

- x** First *dataframe* or *matrix* containing relative abundances. Columns are features, rows are samples. Rows should therefore sum to a constant. Row names are used for identification if present.
- y** Second *dataframe* or *matrix* (optional) containing relative abundances. Columns are features, rows are samples. Rows should therefore sum to a constant. If both **x** and **y** are specified, they will be merged by row names. If no row names are specified for either or both datasets, the default is to merge by row number.

CCREPE: Compositionality Corrected by PERmutation and RENormalization

sim.score Similarity measure, such as `cor` or `nc.score`. This can be either an existing R function or user-defined. If the latter, certain properties should be satisfied as detailed below (also see examples). The default similarity measure is Spearman correlation.

A user-defined similarity measure should mimic the interface of `cor`:

1. Take either two *vector* inputs one *matrix* or *dataframe* input.
2. In the case of two inputs, return a single number.
3. In the case of one input, return a matrix in which the (i,j) th entry is the similarity score for column `i` and column `j` in the original matrix.
4. The resulting matrix (in the case of one input) must be symmetric.
5. The inputs must be named `x` and `y`.

sim.score.args An optional list of arguments for the measurement function. When given, they are passed to the `sim.score` function directly. For example, in the case of `cor`, the following would be acceptable:

```
sim.score.args = list(method="spearman", use="complete.obs")
```

min.subj Minimum number (count) of samples that must be non-missing in order to apply the similarity measure. This is to ensure that there are sufficient samples to perform a bootstrap (default: 20).

iterations The number of iterations for both bootstrap and permutation calculations (default: 1000).

subset.cols.x A vector of column indices from `x` to indicate which features to compare

subset.cols.y A vector of column indices from `y` to indicate which features to compare

errthresh If feature has number of zeros greater than $errthresh^{1/n}$, that feature is excluded

verbose If `TRUE`, print periodic progress of the algorithm through the dataset(s), as well as including more detailed debugging output. (default: `FALSE`).

iterations.gap If `verbose=TRUE`, the number of iterations between issuing status messages (default: 100).

distributions Optional output file for detailed log (if given) of all intermediate permutation and renormalization distributions.

compare.within.x A boolean value indicating whether to do comparisons given by taking all subsets of size 2 from `subset.cols.x` or to do comparisons given by taking all possible combinations of `subset.cols.x` and `subset.cols.y`. If `TRUE` but `subset.cols.y=NA`, returns all comparisons involving any features in `subset.cols.x`. This argument is only used when `y=NA`.

CCREPE: Compositionality Corrected by PERmutation and REnormalization

concurrent.output Optional output file to which each comparison will be written as it is calculated.

make.output.table A boolean value indicating whether to include table-formatted output.

2.3 Output

`ccrepe` returns a *list* containing both the calculation results and the parameters used:

sim.score *matrix* of similarity scores for all requested comparisons. The (i,j) th element corresponds to the similarity score of column i (or the i th column of `subset.cols.1`) and column j (or the j th column of `subset.cols.1`) in one dataset, or to the similarity score of column i (or the i th column of `subset.cols.1`) in dataset x and column j (or the j th column of `subset.cols.2`) in dataset y in the case of two datasets.

p.values *matrix* of the corrected p-values for all requested comparisons. The (i,j) th element corresponds to the p-value of the (i,j) th element of `sim.score`.

q.values *matrix* of the Benjamini-Hochberg-Yekutieli corrected p-values. The (i,j) th element corresponds to the p-value of the (i,j) th element of `sim.score`.

z.stat *matrix* of the z-statistics used in generating the p-values for all requested comparisons. The (i,j) th element corresponds to the z-statistic generating the (i,j) th element of `p.values`.

2.4 Usage

```
ccrepe(  
  x = NA,  
  y = NA,  
  sim.score = cor,  
  sim.score.args = list(),  
  min.subj = 20,  
  iterations = 1000,  
  subset.cols.x = NULL,  
  subset.cols.y = NULL,  
  errthresh = 1e-04,  
  verbose = FALSE,  
  iterations.gap = 100,  
  distributions = NA,  
  compare.within.x = TRUE,  
  concurrent.output = NA,  
  make.output.table = FALSE)
```

2.5 Example 1

An example of how to use `ccrepe` with one dataset.

CCREPE: Compositionality Corrected by PERmutation and REnormalization

```
data <- matrix(rlnorm(40,meanlog=0,sdlog=1),nrow=10,ncol=4)
data[,1] = 2*data[,2] + rnorm(10,0,0.01)
data.rowsum <- apply(data,1,sum)
data.norm <- data/data.rowsum
apply(data.norm,1,sum) # The rows sum to 1, so the data are normalized

## [1] 1 1 1 1 1 1 1 1 1 1

test.input <- data.norm

dimnames(test.input) <- list(c(
  "Sample 1", "Sample 2", "Sample 3", "Sample 4", "Sample 5",
  "Sample 6", "Sample 7", "Sample 8", "Sample 9", "Sample 10"),
  c("Feature 1", "Feature 2", "Feature 3", "Feature 4"))

test.output <- ccrepe(x=test.input, iterations=20, min.subj=10)
```

```
par(mfrow=c(1,2))
plot(data[,1],data[,2],xlab="Feature 1",ylab="Feature 2",main="Non-normalized")
plot(data.norm[,1],data.norm[,2],xlab="Feature 1",ylab="Feature 2",
      main="Normalized")
```

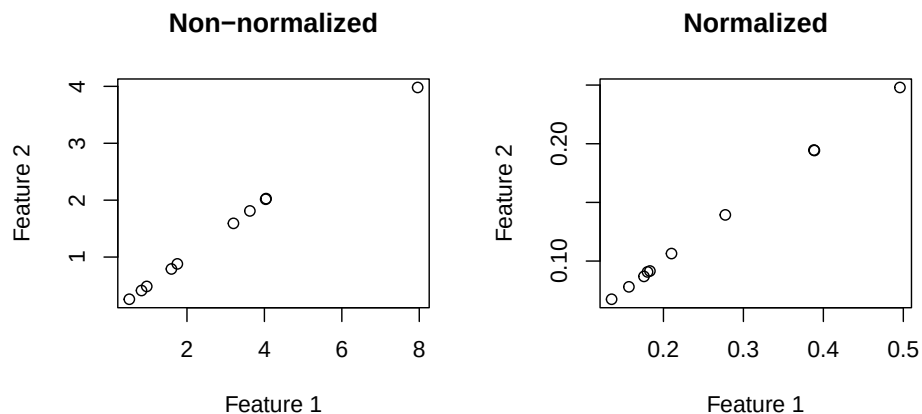


Figure 1: Non-normalized and normalized associations between feature 1 and feature 2. In this case we would expect feature 1 and feature 2 to be associated. In the output we see this by the positive `sim.score` value in the [1,2] element of `test.output$sim.score` and the small `q`-value in the [1,2] element of `test.output$q.values`.

```
test.output

## $p.values
##           Feature 1   Feature 2   Feature 3   Feature 4
## Feature 1          NA 0.0000151584 0.23759345 0.05523248
## Feature 2 0.0000151584          NA 0.06023128 0.13529141
## Feature 3 0.2375934494 0.0602312798          NA 0.33042503
## Feature 4 0.0552324787 0.1352914115 0.33042503          NA
##
## $z.stat
##           Feature 1   Feature 2   Feature 3   Feature 4
## Feature 1          NA  4.326374 -1.1810233  1.9170429
```

CCREPE: Compositionality Corrected by PERmutation and REnormalization

```
## Feature 2  4.326374      NA -1.8790968  1.4935571
## Feature 3 -1.181023 -1.879097      NA -0.9732581
## Feature 4  1.917043  1.493557 -0.9732581      NA
##
## $sim.score
##           Feature 1  Feature 2  Feature 3  Feature 4
## Feature 1      NA  0.99994650 -0.7350114  0.08791348
## Feature 2  0.99994650      NA -0.7359652  0.08933306
## Feature 3 -0.73501141 -0.73596521      NA -0.74004289
## Feature 4  0.08791348  0.08933306 -0.7400429      NA
##
## $q.values
##           Feature 1  Feature 2  Feature 3  Feature 4
## Feature 1      NA  0.0002154592  0.6754236  0.3925331
## Feature 2  0.0002154592      NA  0.2853728  0.4807530
## Feature 3  0.6754235683  0.2853728083      NA  0.7827687
## Feature 4  0.3925331058  0.4807529846  0.7827687      NA
```

2.6 Example 2

An example of how to use `ccrepe` with two datasets.

```
data <- matrix(rlnorm(40,meanlog=0,sdlog=1),nrow=10,ncol=4)
data[,1] = 2*data[,2] + rnorm(10,0,0.01)
data.rowsum <- apply(data,1,sum)
data.norm <- data/data.rowsum
apply(data.norm,1,sum) # The rows sum to 1, so the data are normalized
## [1] 1 1 1 1 1 1 1 1 1 1

test.input <- data.norm

data2 <- matrix(rlnorm(105,meanlog=0,sdlog=1),nrow=15,ncol=7)
aligned.rows <- c(seq(1,4),seq(6,9),11,12) # The datasets dont need
# to have subjects line up exactly
data2[aligned.rows,1] <- 2*data[,3] + rnorm(10,0,0.01)
data2.rowsum <- apply(data2,1,sum)
data2.norm <- data2/data2.rowsum
apply(data2.norm,1,sum) # The rows sum to 1, so the data are normalized
## [1] 1 1 1 1 1 1 1 1 1 1 1 1 1 1 1

test.input.2 <- data2.norm

dimnames(test.input) <- list(paste("Sample",seq(1,10)),paste("Feature",seq(1,4)))
dimnames(test.input.2) <- list(paste("Sample",c(seq(1,4),11,seq(5,8),12,9,10,13,14,15)),paste("Feature",seq(1,4)))

test.output.two.datasets <- ccrepe(x=test.input, y=test.input.2, iterations=20, min.subj=10)

## Warning in preprocess_data(CA): Removing subjects Sample 11, Sample 12, Sample
13, Sample 14, Sample 15 from dataset y because they are not in dataset x.
```

CCREPE: Compositionality Corrected by PERmutation and REnormalization

Please note that we receive a warning because the subjects don't match - only paired observations.

```
par(mfrow=c(1,2))
plot(data2[aligned.rows,1],data[,3],xlab="dataset 2: Feature 1",ylab="dataset 1: Feature 3",main="Non-normalized")
plot(data2.norm[aligned.rows,1],data.norm[,3],xlab="dataset 2: Feature 1",ylab="dataset 1: Feature 3",
      main="Normalized")
```

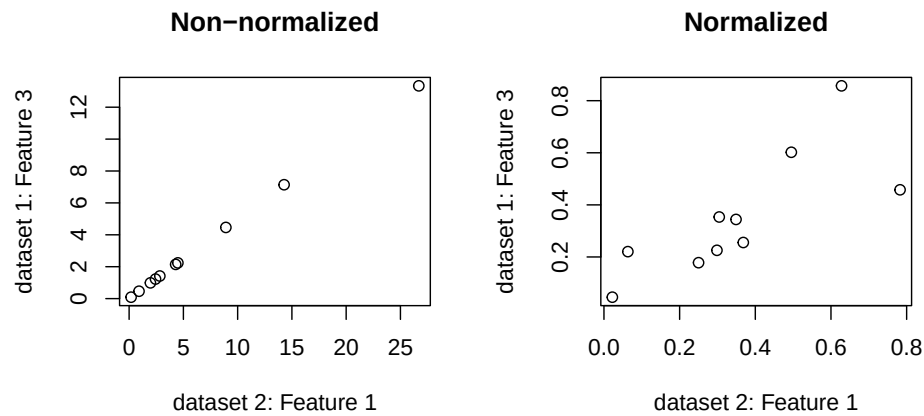


Figure 2: Non-normalized and normalized associations between feature 1 and feature 2. In this case we would expect feature 1 and feature 2 to be associated. In the output we see this by the positive `sim.score` value in the [1,2] element of `test.output$sim.score` and the small `q`-value in the [1,2] element of `test.output$q.values`.

```
test.output.two.datasets

## $p.values
##           Feature 1 Feature 2   Feature 3 Feature 4   Feature 5
## Feature 1 0.0007996296 0.3758273 0.0035227325 0.3389436 0.003725310
## Feature 2 0.0080358268 0.4146369 0.0371076893 0.4009274 0.013985109
## Feature 3 0.0085190218 0.5788777 0.0004169816 0.2110118 0.009252582
## Feature 4 0.4656322519 0.7991743 0.9462954076 0.8778739 0.502736032
##           Feature 6 Feature 7
## Feature 1 0.06912360 0.2414794
## Feature 2 0.06832966 0.3188854
## Feature 3 0.53103872 0.1099252
## Feature 4 0.12465458 0.9655073
##
## $z.stat
##           Feature 1 Feature 2   Feature 3 Feature 4 Feature 5 Feature 6
## Feature 1 -3.352923 -0.8856108 2.91800982 0.9562546 2.900531 1.8176110
## Feature 2 -2.650561 -0.8157609 2.08457727 0.8399660 2.457646 1.8228266
## Feature 3 2.630776 0.5550249 -3.52909612 -1.2507882 -2.602578 -0.6264212
## Feature 4 0.729604 0.2544158 0.06735963 0.1536650 -0.670191 -1.5355264
##           Feature 7
## Feature 1 1.17129715
## Feature 2 0.99675106
## Feature 3 -1.59852932
## Feature 4 0.04324368
##
```

CCREPE: Compositionality Corrected by PERmutation and RENormalization

```
## $sim.score
##           Feature 1 Feature 2 Feature 3 Feature 4 Feature 5 Feature 6
## Feature 1 -0.7225291 -0.3754699  0.6195113  0.2341095  0.6889542  0.6366791
## Feature 2 -0.7219541 -0.3750016  0.6182510  0.2304817  0.6906502  0.6393097
## Feature 3  0.7764977  0.3214216 -0.7840688 -0.4466334 -0.6377653 -0.3655274
## Feature 4  0.0108511  0.1183484  0.1535978  0.2661794 -0.1525060 -0.4490248
##           Feature 7
## Feature 1  0.3199952
## Feature 2  0.3173564
## Feature 3 -0.4951968
## Feature 4  0.2042672
##
## $q.values
##           Feature 1 Feature 2 Feature 3 Feature 4 Feature 5 Feature 6
## Feature 1  0.04376523  2.285526  0.12853719  2.182473  0.1019466  0.6878664
## Feature 2  0.17592637  2.269386  0.45132749  2.309843  0.1913578  0.7479622
## Feature 3  0.15542070  2.640255  0.04564438  1.649867  0.1446889  2.5273694
## Feature 4  2.42713616  3.499225  3.83648363  3.695976  2.5014263  1.0496277
##           Feature 7
## Feature 1  1.762216
## Feature 2  2.181650
## Feature 3  1.002736
## Feature 4  3.774574
```

2.7 Example 3

An example of how to use `ccrepe` with `nc.score` as the similarity score.

```
data <- matrix(rlnorm(40,meanlog=0,sdlog=1),nrow=10,ncol=4)
data[,1] = 2*data[,2] + rnorm(10,0,0.01)
data.rowsum <- apply(data,1,sum)
data.norm <- data/data.rowsum
apply(data.norm,1,sum) # The rows sum to 1, so the data are normalized
## [1] 1 1 1 1 1 1 1 1 1 1

test.input <- data.norm

dimnames(test.input) <- list(paste("Sample",seq(1,10)),paste("Feature",seq(1,4)))

test.output.nc.score <- ccrepe(x=test.input, sim.score=nc.score, iterations=20, min.subj=10)

par(mfrow=c(1,2))
plot(data[,1],data[,2],xlab="Feature 1",ylab="Feature 2",main="Non-normalized")
plot(data.norm[,1],data.norm[,2],xlab="Feature 1",ylab="Feature 2",
      main="Normalized")
```


CCREPE: Compositionality Corrected by PERmutation and REnormalization

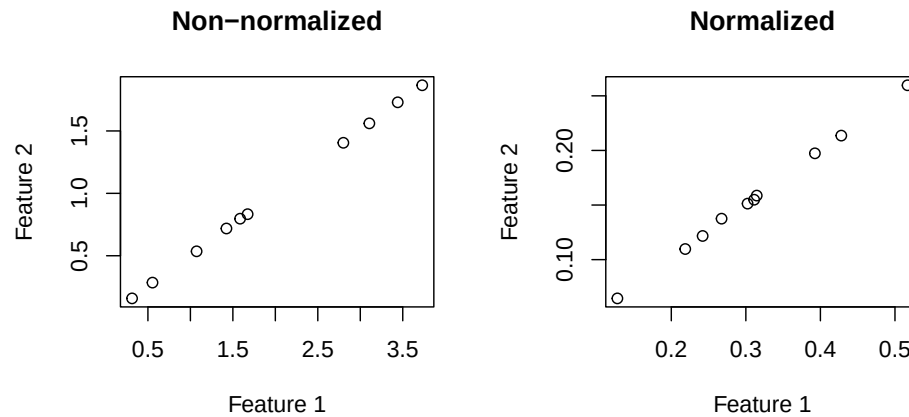


Figure 3: Non-normalized and normalized associations between feature 1 and feature 2. In this case we would expect feature 1 and feature 2 to be associated. In the output we see this by the positive `sim.score` value in the [1,2] element of `test.output$sim.score` and the small `q`-value in the [1,2] element of `test.output$q.values`. In this case, however, the `sim.score` represents the NC-Score between two features rather than the Spearman correlation.

```
test.output.nc.score

## $p.values
##           Feature 1  Feature 2 Feature 3 Feature 4
## Feature 1          NA 0.000105027 0.2393116 0.8989988
## Feature 2 0.000105027          NA 0.4762399 0.9221156
## Feature 3 0.239311626 0.476239904          NA 0.2125140
## Feature 4 0.898998824 0.922115631 0.2125140          NA
##
## $z.stat
##           Feature 1  Feature 2 Feature 3 Feature 4
## Feature 1          NA 3.87867496 1.1767091 -0.12692618
## Feature 2 3.8786750          NA 0.7123632 -0.09776912
## Feature 3 1.1767091 0.71236319          NA -1.24668245
## Feature 4 -0.1269262 -0.09776912 -1.2466825          NA
##
## $sim.score
##           Feature 1 Feature 2 Feature 3 Feature 4
## Feature 1          NA 1.00000 0.21875 -0.59375
## Feature 2 1.00000          NA 0.21875 -0.59375
## Feature 3 0.21875 0.21875          NA -0.65625
## Feature 4 -0.59375 -0.59375 -0.65625          NA
##
## $q.values
##           Feature 1  Feature 2 Feature 3 Feature 4
## Feature 1          NA 0.001492838 1.133847 2.555647
## Feature 2 0.001492838          NA 1.692301 2.184469
## Feature 3 1.133846581 1.692300736          NA 1.510321
## Feature 4 2.555647033 2.184469000 1.510321          NA
```

2.8 Example 4

An example of how to use `ccrepe` with a user-defined `sim.score` function.

```
data <- matrix(rlnorm(40,meanlog=0,sdlog=1),nrow=10,ncol=4)
data[,1] = 2*data[,2] + rnorm(10,0,0.01)
data.rowsum <- apply(data,1,sum)
data.norm <- data/data.rowsum
apply(data.norm,1,sum) # The rows sum to 1, so the data are normalized
## [1] 1 1 1 1 1 1 1 1 1 1

test.input <- data.norm

dimnames(test.input) <- list(paste("Sample",seq(1,10)),paste("Feature",seq(1,4)))

my.test.sim.score <- function(x,y=NA,constant=0.5){
  if(is.vector(x) && is.vector(y)) return(constant)
  if(is.matrix(x) && is.na(y)) return(matrix(rep(constant,ncol(x)^2),ncol=ncol(x)))
  if(is.data.frame(x) && is.na(y)) return(matrix(rep(constant,ncol(x)^2),ncol=ncol(x)))
  else stop('ERROR')
}

test.output.sim.score <- ccrepe(x=test.input, sim.score=my.test.sim.score, iterations=20, min.subj=10, sim)

par(mfrow=c(1,2))
plot(data[,1],data[,2],xlab="Feature 1",ylab="Feature 2",main="Non-normalized")
plot(data.norm[,1],data.norm[,2],xlab="Feature 1",ylab="Feature 2",
      main="Normalized")
```

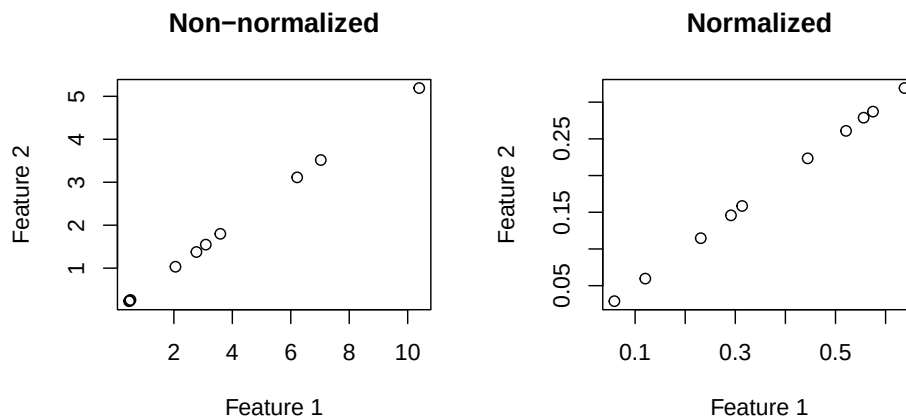


Figure 4: Non-normalized and normalized associations between feature 1 and feature 2. In this case we would expect feature 1 and feature 2 to be associated. Note that the values of `sim.score` are all 0.6 and none of the p-values are very small because of the arbitrary definition of the similarity score.

```
test.output.sim.score
## $p.values
##           Feature 1 Feature 2 Feature 3 Feature 4
## Feature 1         NA        NaN        NaN        NaN
```

CCREPE: Compositionality Corrected by PERmutation and RENormalization

```
## Feature 2      NaN      NA      NaN      NaN
## Feature 3      NaN      NaN      NA      NaN
## Feature 4      NaN      NaN      NaN      NA
##
## $z.stat
##           Feature 1 Feature 2 Feature 3 Feature 4
## Feature 1      NA      NaN      NaN      NaN
## Feature 2      NaN      NA      NaN      NaN
## Feature 3      NaN      NaN      NA      NaN
## Feature 4      NaN      NaN      NaN      NA
##
## $sim.score
##           Feature 1 Feature 2 Feature 3 Feature 4
## Feature 1      NA      0.6      0.6      0.6
## Feature 2      0.6      NA      0.6      0.6
## Feature 3      0.6      0.6      NA      0.6
## Feature 4      0.6      0.6      0.6      NA
##
## $q.values
##           Feature 1 Feature 2 Feature 3 Feature 4
## Feature 1      NA      NaN      NaN      NaN
## Feature 2      NaN      NA      NaN      NaN
## Feature 3      NaN      NaN      NA      NaN
## Feature 4      NaN      NaN      NaN      NA
```

2.9 Example 5

An example of how to use `ccrepe` when specifying column subsets.

```
data <- matrix(rlnorm(40,meanlog=0,sdlog=1),nrow=10,ncol=4)
data.rowsum <- apply(data,1,sum)
data.norm <- data/data.rowsum
apply(data.norm,1,sum) # The rows sum to 1, so the data are normalized
## [1] 1 1 1 1 1 1 1 1 1 1

test.input <- data.norm

dimnames(test.input) <- list(paste("Sample",seq(1,10)),paste("Feature",seq(1,4)))

test.output.1.3 <- ccrepe(x=test.input, iterations=20, min.subj=10, subset.cols.x=c(1,3))
test.output.1 <- ccrepe(x=test.input, iterations=20, min.subj=10, subset.cols.x=c(1), compare.within.x=
test.output.12.3 <- ccrepe(x=test.input, iterations=20, min.subj=10, subset.cols.x=c(1,2),subset.cols.y=c
test.output.1.3$sim.score

##           Feature 1 Feature 3
## Feature 1      NA -0.4991686
## Feature 3 -0.4991686      NA

test.output.1$sim.score

##           Feature 1 Feature 2 Feature 3 Feature 4
## Feature 1      NA -0.3687573 -0.4991686 -0.0847935
```

```
## Feature 2 -0.3687573      NA      NA      NA
## Feature 3 -0.4991686      NA      NA      NA
## Feature 4 -0.0847935      NA      NA      NA

test.output.12.3$sim.score

##           Feature 1 Feature 2 Feature 3 Feature 4
## Feature 1          NA      NA -0.4991686          NA
## Feature 2          NA      NA -0.1180298          NA
## Feature 3 -0.4991686 -0.1180298          NA          NA
## Feature 4          NA      NA          NA          NA
```

3 nc.score

The `nc.score` similarity measure is an N-dimensional extension of the checkerboard score particularly suited to similarity score calculations between compositions derived from ecological relative abundance measurements. In such cases, features typically represent species abundances, and the NC-score discretizes these continuous values into one of N bins before computing a normalized similarity of co-occurrence or co-exclusion. This can be used as a standalone function or with `ccrepe` as above to obtain compositionality-corrected p-values.

3.1 General Functionality

The NC-score is an extension to Diamond's checkerboard score (see [Cody and Diamond \[1975\]](#)) to ordinal data, and simplifies to a calculation of Kendall's τ on binned data instead of ranked data. Let two features in a dataset with n subjects be denoted by

$$\begin{bmatrix} x_1 & x_2 & \dots & x_n \\ y_1 & y_2 & \dots & y_n \end{bmatrix}.$$

The binning function maps from the original data to b numbered bins in $\{1, \dots, b\}$. Let the binning function be denoted by $B(\cdot)$. The co-variation and co-exclusion patterns are the same as concordant and discordant pairs in Kendall's τ . Considering a 2×2 submatrix of the form

$$\begin{bmatrix} B(x_i) & B(x_j) \\ B(y_i) & B(y_j) \end{bmatrix},$$

a co-variation pattern is counted when $(B(x_i) - B(x_j))(B(y_i) - B(y_j)) > 0$ and a co-exclusion pattern, conversely, when $(B(x_i) - B(x_j))(B(y_i) - B(y_j)) < 0$. The NC-score statistic for features x and y is then defined as

$$(\text{number of co-variation patterns}) - (\text{number of co-exclusion patterns}),$$

normalized by the Kendall's τ normalization factor accounting for ties described in [Kendall \[1970\]](#).

3.2 Arguments

- ✕ First numerical *vector*, or single *dataframe* or *matrix*, containing relative abundances. If the latter, columns are features, rows are samples. Rows should therefore sum to a constant.

CCREPE: Compositionality Corrected by PERmutation and RENormalization

y If provided, second numerical *vector* containing relative abundances. If given, **x** must be a *vector* as well.

nbins A non-negative integer of the number of bins to generate (cutoffs will be generated by the `discretize` function from the `infotheo` package).

bin.cutoffs A list of values demarcating the bin cutoffs. The binning is performed using the `findInterval` function.

use An optional character string giving a method for computing covariances in the presence of missing values. This must be (an abbreviation of) one of the strings "everything", "all.obs", "complete.obs", "na.or.complete", or "pairwise.complete.obs".

3.3 Output

nc.score returns either a single number (if called with two vectors) or a *matrix* of all pairwise scores (if called with a *matrix*) of normalized scores. This behaviour is precisely analogous to the `cor` function in R

3.4 Usage

```
nc.score(  
  x,  
  y = NULL,  
  use = "everything",  
  nbins = NULL,  
  bin.cutoffs=NULL)
```

3.5 Example 1

An example of using **nc.score** to get a single similarity score or a matrix.

```
data <- matrix(rlnorm(40,meanlog=0,sdlog=1),nrow=10,ncol=4)  
data.rowsum <- apply(data,1,sum)  
data[,1] = 2*data[,2] + rnorm(10,0,0.01)  
data.norm <- data/data.rowsum  
apply(data.norm,1,sum) # The rows sum to 1, so the data are normalized  
## [1] 0.4225119 1.0653523 0.9198234 0.7336066 1.2048596 1.2686695 0.6536425  
## [8] 1.8421933 1.8407167 1.3301675  
  
test.input <- data.norm  
  
dimnames(test.input) <- list(paste("Sample",seq(1,10)),paste("Feature",seq(1,4)))  
  
test.output.matrix <- nc.score(x=test.input)  
test.output.num <- nc.score(x=test.input[,1],y=test.input[,2])  
  
par(mfrow=c(1, 2))  
plot(data[,1],data[,2],xlab="Feature 1",ylab="Feature 2",main="Non-normalized")
```

```
plot(data.norm[,1],data.norm[,2],xlab="Feature 1",ylab="Feature 2",
     main="Normalized")
```

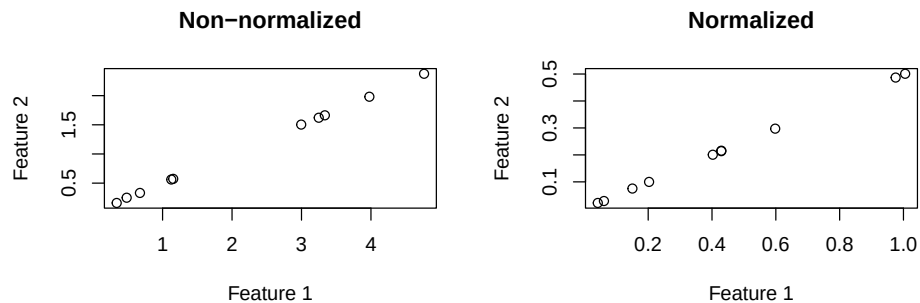


Figure 5: Non-normalized and normalized associations between feature 1 and feature 2 of the second example. Again, we expect to observe a positive association between feature 1 and feature 2. In terms of generalized checkerboard scores, we would expect to see more co-variation patterns than co-exclusion patterns. This is shown by the positive and relatively high value of the [1,2] element of `test.output.matrix` (which is identical to `test.output.num`)

```
test.output.matrix

##           Feature 1 Feature 2 Feature 3 Feature 4
## Feature 1   1.00000   1.00000   0.43750  -0.59375
## Feature 2   1.00000   1.00000   0.43750  -0.59375
## Feature 3   0.43750   0.43750   1.00000  -0.59375
## Feature 4  -0.59375  -0.59375  -0.59375   1.00000

test.output.num

## [1] 1
```

3.6 Example 2

An example of using `nc.score` with an arbitrary bin number.

```
data <- matrix(rlnorm(40,meanlog=0,sdlog=1),nrow=10,ncol=4)
data.rowsum <- apply(data,1,sum)
data[,1] = 2*data[,2] + rnorm(10,0,0.01)
data.norm <- data/data.rowsum
apply(data.norm,1,sum) # The rows sum to 1, so the data are normalized

## [1] 1.9117971 1.2245640 2.1387299 1.8491773 0.6060685 0.9443157 0.9153110
## [8] 0.8535768 0.8486709 1.1937010

test.input <- data.norm

dimnames(test.input) <- list(paste("Sample",seq(1,10)),paste("Feature",seq(1,4)))

test.output <- nc.score(x=test.input,nbins=4)
```

```
par(mfrow=c(1, 2))
plot(data[,1],data[,2],xlab="Feature 1",ylab="Feature 2",main="Non-normalized")
plot(data.norm[,1],data.norm[,2],xlab="Feature 1",ylab="Feature 2",
```

```
main="Normalized")
```

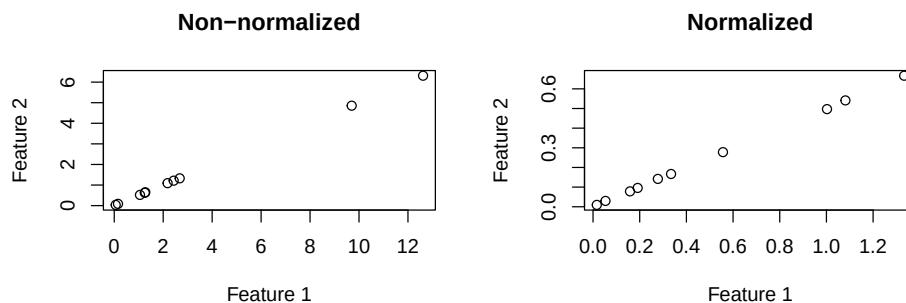


Figure 6: Non-normalized and normalized associations between feature 1 and feature 2 of the second example. Again, we expect to observe a positive association between feature 1 and feature 2. In terms of generalized checkerboard scores, we would expect to see more co-variation patterns than co-exclusion patterns. This is shown by the positive and relatively high value in the [1,2] element of test.output. In this case, the smaller bin number yields a smaller NC-score because of the coarser partitioning of the data.

```
test.output
```

```
##           Feature 1  Feature 2  Feature 3  Feature 4
## Feature 1  1.0000000  1.0000000 -0.77142857 -0.20000000
## Feature 2  1.0000000  1.0000000 -0.77142857 -0.20000000
## Feature 3 -0.7714286 -0.7714286  1.00000000  0.02857143
## Feature 4 -0.2000000 -0.2000000  0.02857143  1.00000000
```

3.7 Example 3

An example of using `nc.score` with user-defined bin edges.

```
data <- matrix(rlnorm(40,meanlog=0,sdlog=1),nrow=10,ncol=4)
data.rowsum <- apply(data,1,sum)
data[,1] = 2*data[,2] + rnorm(10,0,0.01)
data.norm <- data/data.rowsum
apply(data.norm,1,sum) # The rows sum to 1, so the data are normalized

## [1] 0.9160103 0.5917141 1.1612162 0.6204630 1.1466295 0.8921647 1.0956650
## [8] 2.1057495 1.3579113 1.8069848

test.input <- data.norm

dimnames(test.input) <- list(paste("Sample",seq(1,10)),paste("Feature",seq(1,4)))

test.output <- nc.score(x=test.input,bin.cutoffs=c(0.1,0.2,0.3))

par(mfrow=c(1, 2))
plot(data[,1],data[,2],xlab="Feature 1",ylab="Feature 2",main="Non-normalized")
plot(data.norm[,1],data.norm[,2],xlab="Feature 1",ylab="Feature 2",
      main="Normalized")
```

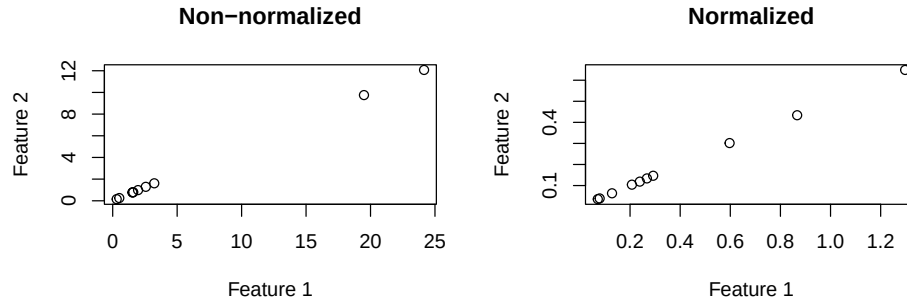


Figure 7: Non-normalized and normalized associations between feature 1 and feature 2 of the second example. Again, we expect to observe a positive association between feature 1 and feature 2. In terms of generalized checkerboard scores, we would expect to see more co-variation patterns than co-exclusion patterns. This is shown by the positive and relatively high value in the [1,2] element of test.output. The bin edges specified here represent almost absent ([0,0.001)), low abundance ([0.001,0.1)), medium abundance ([0.1,0.25)), and high abundance ([0.6,1)).

```
test.output
```

##	Feature 1	Feature 2	Feature 3	Feature 4
## Feature 1	1.00000000	0.97100831	0.02942449	-0.2125119
## Feature 2	0.97100831	1.00000000	0.06060606	-0.2501222
## Feature 3	0.02942449	0.06060606	1.00000000	-0.6878359
## Feature 4	-0.21251186	-0.25012216	-0.68783594	1.00000000

4 References

References

- Martin Leonard Cody and Jared Mason Diamond. *Ecology and evolution of communities*. Harvard University Press, 1975.
- Karoline Faust, J Fah Sathirapongsasuti, Jacques Izard, Nicola Segata, Dirk Gevers, Jeroen Raes, and Curtis Huttenhower. Microbial co-occurrence relationships in the human microbiome. *PLoS computational biology*, 8(7):e1002606, 2012.
- M.G. Kendall. *Rank correlation methods*. Charles Griffin & Co., 1970.
- Emma Schwager and Colleagues. Detecting statistically significant associations between sparse and high dimensional compositional data. In Progress.