

Package ‘ClustBlock’

September 5, 2026

Title Clustering of Datasets

Version 6.1.0

Maintainer Fabien Llobell <fabienllobellresearch@gmail.com>

Description Hierarchical and partitioning algorithms to cluster blocks of variables. The partitioning algorithm includes an option called noise cluster to set aside atypical blocks of variables. Different thresholds per cluster can be sets. The CLUSTATIS method (for quantitative blocks) (Llobell, Cariou, Vigneau, Labenne & Qannari (2020) <[doi:10.1016/j.foodqual.2018.05.013](https://doi.org/10.1016/j.foodqual.2018.05.013)>, Llobell, Vigneau & Qannari (2019) <[doi:10.1016/j.foodqual.2019.02.017](https://doi.org/10.1016/j.foodqual.2019.02.017)>) and the CLUS-CATA method (for Check-All-That-Apply data) (Llobell, Cariou, Vigneau, Labenne & Qannari (2019) <[doi:10.1016/j.foodqual.2018.09.006](https://doi.org/10.1016/j.foodqual.2018.09.006)>, Llobell, Giacalone, Labenne & Qannari (2019) <[doi:10.1016/j.foodqual.2019.05.017](https://doi.org/10.1016/j.foodqual.2019.05.017)>) are the core of this package. The CATATIS methods allows to compute some indices and tests to control the quality of CATA data (Llobell, Bonnet & Giacalone (2024) <[doi:10.1111/joss.12941](https://doi.org/10.1111/joss.12941)>). Multivariate analysis and clustering of subjects for quantitative multiblock data, CATA, RATA, Free Sorting and JAR experiments are available. Clustering of observations (products in sensory analysis) in multi-block context (notably with ClusMB strategy) is also included (Llobell & Giacalone (2025) <[doi:10.1111/joss.70024](https://doi.org/10.1111/joss.70024)>). Performing clustering based on CATA and liking at the same time is possible thanks to cluscata_liking function (Vigneau, Cariou, Giacalone, Berget & Llobell (2022) <[doi:10.1016/j.foodqual.2021.104358](https://doi.org/10.1016/j.foodqual.2021.104358)>). Clustering of variables (quantitative, qualitative or mixed) can be done thanks to the MixCluStatis() function. Clustering on JAR + Liking can be achieved thanks to preprocess_JAR_liking function.

Depends R (>= 3.5)

License MIT + file LICENSE

Encoding UTF-8

LazyData true

Imports FactoMineR

Suggests ClustVarLV

Date 2026-09-05

RoxygenNote 7.3.3

NeedsCompilation no

Author Fabien Llobell [aut, cre] (SensElevation),
Evelyne Vigneau [ctb] (Oniris),

Veronique Cariou [ctb] (Oniris),
El Mostafa Qannari [ctb] (Oniris)

Repository CRAN

Date/Publication 2026-09-05 15:40:27 UTC

Contents

ClustBlock-package	3
catatis	4
catatis_jar	6
catatis_rata	8
cata_ryebread	9
change_cata_format	10
change_cata_format2	11
cheese	12
choc	12
cluscata	13
cluscata_jar	15
cluscata_kmeans	18
cluscata_kmeans_jar	19
cluscata_liking	21
cluscata_rata	23
ClusMB	25
clustatis	27
clustatis_FreeSort	29
clustatis_FreeSort_kmeans	32
clustatis_kmeans	33
clustRowsOnStatisAxes	35
combinCATALiking	37
consistency_cata	38
consistency_cata_panel	39
croissant	40
fish	40
indicesClusters	41
liking_ryebread	42
mixclustatis	43
plot.catatis	45
plot.cluscata	46
plot.cluscata_liking	47
plot.clusRows	48
plot.clustatis	49
plot.statis	50
preprocess_FreeSort	51
preprocess_JAR	52
preprocess_JAR_liking	53
print.catatis	55
print.cluscata	55

print.cluscata_liking	56
print.clusRows	56
print.clustatis	57
print.statis	57
RATAchoc	58
simil_groups_cata	58
smoo	59
statis	60
statis_FreeSort	61
straw	63
summary.catatis	63
summary.cluscata	64
summary.cluscata_liking	65
summary.clusRows	65
summary.clustatis	66
summary.statis	67
Index	68

ClustBlock-package	<i>Clustering of Datasets</i>
--------------------	-------------------------------

Description

Hierarchical and partitioning algorithms of blocks of variables. The CLUSTATIS method and the CLUSCATA method are the core of this package. The CATATIS methods allows to compute some indices and tests to control the quality of CATA data. Multivariate analysis and clustering of subjects for quantitative multiblock data, CATA, RATA, Free Sorting and JAR experiments are available. Clustering of rows in multi-block context (notably with ClusMB strategy) is also included. Performing clustering based on CATA and liking at the same time is possible thanks to cluscata_liking function. Clustering of variables (quantitative, qualitative or mixed) can be done thanks to the Mix-CluStatis function. Clustering on JAR + Liking can be achieved thanks to preprocess_JAR_liking function.

Details

Package:	ClustBlock
Type:	Package
Version:	6.1.0
First version Date:	2019-03-06
Last version Date:	2026-09-05

Author(s)

Fabien Llobell, Evelyne Vigneau, Veronique Cariou, El Mostafa Qannari

Maintainer: <fabienllobellresearch@gmail.com>

References

- Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2020). Analysis and clustering of multiblock datasets by means of the STATIS and CLUSTATIS methods. Application to sensometrics. *Food Quality and Preference*, 79, 103520.
- Llobell, F., Vigneau, E., & Qannari, E. M. (2019). Clustering datasets by means of CLUSTATIS with identification of atypical datasets. Application to sensometrics. *Food Quality and Preference*, 75, 97-104.
- Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2019). A new approach for the analysis of data and the clustering of subjects in a CATA experiment. *Food quality and preference*, 72, 31-39.
- Llobell, F., Giacalone, D., Labenne, A., & Qannari, E. M. (2019). Assessment of the agreement and cluster analysis of the respondents in a CATA experiment. *Food Quality and Preference*, 77, 184-190.
- Llobell, F., & Qannari, E. M. (2020). CLUSTATIS: Cluster analysis of blocks of variables. *Electronic Journal of Applied Statistical Analysis*, 13(2), 436-453.
- Llobell, F. (2020). Classification de tableaux de donnees, applications en analyse sensorielle (Doctoral dissertation, Nantes, Ecole nationale veterinaire).
- Vigneau, E., Cariou, V., Giacalone, D., Berget, I., & Llobell, F. (2022). Combining hedonic information and CATA description.
- Llobell, F., Bonnet, L., & Giacalone, D. (2024). Assessment of panel performance in CATA and RATA experiment. *Journal of Sensory Studies*, 39(4), e12941.
- Llobell, F., & Giacalone, D. (2025). Two Methods for Clustering Products in a Sensory Study: STATIS and ClusMB. *Journal of Sensory Studies*, 40(1), e70024.

catatis

Perform the CATATIS method on different blocks from a CATA experiment

Description

CATATIS method. Additional outputs are also computed. Non-binary data are accepted and weights can be tested.

Usage

```
catatis(Data,nblo,NameBlocks=NULL, NameVar=NULL, Graph=TRUE, Graph_weights=TRUE,
Test_weights=FALSE, nperm=100)
```

Arguments

Data	data frame or matrix where the blocks of binary variables are merged horizontally. If you have a different format, see change_cata_format
nblo	integer. Number of blocks (subjects).
NameBlocks	string vector. Name of each block (subject). Length must be equal to the number of blocks. If NULL, the names are S1,...Sm. Default: NULL
NameVar	string vector. Name of each variable (attribute, the same names for each subject). Length must be equal to the number of attributes. If NULL, the colnames of the first block are taken. Default: NULL
Graph	logical. Show the graphical representation? Default: TRUE
Graph_weights	logical. Should the barplot of the weights be plotted? Default: TRUE
Test_weights	logical. Should the the weights be tested? Default: FALSE
nperm	integer. Number of permutation for the weight tests. Default: 100

Value

a list with:

- S: the S matrix: a matrix with the similarity coefficient among the subjects
- compromise: a matrix which is the compromise of the subjects (akin to a weighted average)
- weights: the weights associated with the subjects to build the compromise
- weights_tests: the weights tests results
- lambda: the first eigenvalue of the S matrix
- overall error: the error for the CATATIS criterion
- error_by_sub: the error by subject (CATATIS criterion)
- error_by_prod: the error by product (CATATIS criterion)
- s_with_compromise: the similarity coefficient of each subject with the compromise
- homogeneity: homogeneity of the subjects (in percentage)
- CA: the results of correspondence analysis performed on the compromise dataset
- eigenvalues: the eigenvalues associated to the correspondence analysis
- inertia: the percentage of total variance explained by each axis of the CA
- scalefactors: the scaling factors of each subject
- nb_1: the number of 1 in each block, i.e. the number of checked attributes by subject.
- param: parameters called

References

Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2019). A new approach for the analysis of data and the clustering of subjects in a CATA experiment. *Food Quality and Preference*, 72, 31-39.

Bonnet, L., Ferney, T., Riedel, T., Qannari, E.M., Llobell, F. (September 14, 2022) .Using CATA for sensory profiling: assessment of the panel performance. Eurosense, Turku, Finland.

See Also

[plot.catatis](#), [summary.catatis](#), [cluscata](#), [change_cata_format](#), [change_cata_format2](#)

Examples

```
data(straw)
res.cat=catatis(straw, nblo=114)
summary(res.cat)
plot(res.cat)

#Vertical format with sessions
data("fish")
chang=change_cata_format2(fish, nprod= 6, nattrib= 27, nsub = 12, nsess= 3)
res.cat2=catatis(Data= chang$Datafinal, nblo = 12, NameBlocks = chang$NameSub, Test_weights=TRUE)

#Vertical format without sessions
Data=fish[1:66,2:30]
chang2=change_cata_format2(Data, nprod= 6, nattrib= 27, nsub = 11, nsess= 1)
res.cat3=catatis(Data= chang2$Datafinal, nblo = 11, NameBlocks = chang2$NameSub)
```

catatis_jar

Perform the CATATIS method on Just About Right data.

Description

CATATIS method adapted to JAR data.

Usage

```
catatis_jar(Data, nprod, nsub, levelsJAR=3, beta=0.1, Graph=TRUE, Graph_weights=TRUE,
Test_weights=FALSE, nperm=100)
```

Arguments

Data	data frame where the first column is the Assessors, the second is the products and all other columns the JAR attributes with numbers (1 to 3 or 1 to 5, see levelsJAR)
nprod	integer. Number of products.
nsub	integer. Number of subjects.
levelsJAR	integer. 3 or 5 levels. If 5, the data will be transformed in 3 levels.
beta	numerical. Parameter for agreement between JAR and other answers. Between 0 and 0.5.
Graph	logical. Show the graphical representation? Default: TRUE
Graph_weights	logical. Should the barplot of the weights be plotted? Default: TRUE
Test_weights	logical. Should the the weights be tested? Default: FALSE
nperm	integer. Number of permutation for the weight tests. Default: 100

Value

a list with:

- S: the S matrix: a matrix with the similarity coefficient among the subjects
- compromise: a matrix which is the compromise of the subjects (akin to a weighted average)
- weights: the weights associated with the subjects to build the compromise
- weights_tests: the weights tests results
- lambda: the first eigenvalue of the S matrix
- overall error: the error for the CATATIS criterion
- error_by_sub: the error by subject (CATATIS criterion)
- error_by_prod: the error by product (CATATIS criterion)
- s_with_compromise: the similarity coefficient of each subject with the compromise
- homogeneity: homogeneity of the subjects (in percentage)
- CA: the results of correspondance analysis performed on the compromise dataset
- eigenvalues: the eigenvalues associated to the correspondance analysis
- inertia: the percentage of total variance explained by each axis of the CA
- scalefactors: the scaling factors of each subject
- nb_1: Can be ignored
- param: parameters called

References

Llobell, F., Vigneau, E. & Qannari, E. M. ((September 14, 2022). Multivariate data analysis and clustering of subjects in a Just about right task. Eurosense, Turku, Finland.

See Also

[catatis](#), [plot.catatis](#), [summary.catatis](#), [cluscata_jar](#), [preprocess_JAR](#), [cluscata_kmeans_jar](#)

Examples

```
data(cheese)
res.cat=catatis_jar(Data=cheese, nprod=8, nsub=72, levelsJAR=5)
summary(res.cat)
#plot(res.cat)
```

catatis_rata	<i>Perform the CATATIS method on different blocks from a RATA experiment</i>
--------------	--

Description

CATATIS method for RATA data. Additional outputs are also computed. Non-binary data are accepted and weights can be tested.

Usage

```
catatis_rata(Data,nblo,NameBlocks=NULL, NameVar=NULL, Graph=TRUE, Graph_weights=TRUE,
  Test_weights=FALSE, nperm=100)
```

Arguments

Data	data frame or matrix where the blocks of variables are merged horizontally. If you have a different format, see change_cata_format
nblo	integer. Number of blocks (subjects).
NameBlocks	string vector. Name of each block (subject). Length must be equal to the number of blocks. If NULL, the names are S1,...Sm. Default: NULL
NameVar	string vector. Name of each variable (attribute, the same names for each subject). Length must be equal to the number of attributes. If NULL, the colnames of the first block are taken. Default: NULL
Graph	logical. Show the graphical representation? Default: TRUE
Graph_weights	logical. Should the barplot of the weights be plotted? Default: TRUE
Test_weights	logical. Should the the weights be tested? Default: FALSE
nperm	integer. Number of permutation for the weight tests. Default: 100

Value

a list with:

- S: the S matrix: a matrix with the similarity coefficient among the subjects
- compromise: a matrix which is the compromise of the subjects (akin to a weighted average)
- weights: the weights associated with the subjects to build the compromise
- weights_tests: the weights tests results
- lambda: the first eigenvalue of the S matrix
- overall error: the error for the CATATIS criterion
- error_by_sub: the error by subject (CATATIS criterion)
- error_by_prod: the error by product (CATATIS criterion)
- s_with_compromise: the similarity coefficient of each subject with the compromise
- homogeneity: homogeneity of the subjects (in percentage)

- CA: the results of correspondence analysis performed on the compromise dataset
- eigenvalues: the eigenvalues associated to the correspondence analysis
- inertia: the percentage of total variance explained by each axis of the CA
- scalefactors: the scaling factors of each subject
- param: parameters called

References

Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2019). A new approach for the analysis of data and the clustering of subjects in a CATA experiment. *Food Quality and Preference*, 72, 31-39.

Bonnet, L., Ferney, T., Riedel, T., Qannari, E.M., Llobell, F. (September 14, 2022) .Using CATA for sensory profiling: assessment of the panel performance. Eurosense, Turku, Finland.

Bonnet, L., Llobell, F., Qannari, E.M. (Pangborn 2023). Assessment of the panel performance in a RATA experiment.

See Also

[catatis](#), [plot.catatis](#), [summary.catatis](#), [change_cata_format](#), [change_cata_format2](#)

Examples

```
#RATA data with session
data(RATAchoc)
chang2=change_cata_format2(RATAchoc, nprod= 12, nattr= 13, nsub = 9, nsess= 3)
res.cat4=catatis_rata(Data= chang2$Datafinal, nblo = 9, NameBlocks =  chang2$NameSub)
summary(res.cat4)

#RATA data without session
Data=RATAchoc[1:108,2:16]
chang2=change_cata_format2(Data, nprod= 12, nattr= 13, nsub = 9, nsess = 1)
res.cat5=catatis_rata(Data= chang2$Datafinal, nblo = 9, NameBlocks =  chang2$NameSub)
summary(res.cat5)
graphics.off()
```

cata_ryebread

cata_ryebread data

Description

cata_ryebread data

Usage

```
data(cata_ryebread)
```

Format

CATA data. A data frame with 6 rows (the number of ryebread) and 1848 columns (the number of consumers (132) * the number of attributes (14)). For each consumer, each attribute and each product, there is 1 if the attribute has been checked by the consumer for the product, and 0 if not.

References

Giactalone, D. (2018). Product Performance Optimization. In Ares, G., & Varela, P. (Eds.) Methods in Consumer Research, Volume I (Chapter 7. pp. 159-185), Elsevier.

Examples

```
data(cata_ryebread)
```

change_cata_format	<i>Change format of CATA datasets to perform CATATIS or CLUSCATA function</i>
--------------------	---

Description

CATATIS and CLUSCATA operate on data where the blocksvariables are merged horizontally. If you have a different format, you can use this function to change the format. Format=1 is for data merged vertically with the dataset of the first subject, then the second,... with products in same order Format=2 is for data merged vertically with the dataset for the first product, then the second... with subjects in same order

Unlike change_cata_format2, you don't need to specify products and subjects, just make sure they are in the right order.

Usage

```
change_cata_format(Data, nprod, nattr, nsub, format=1, NameProds=NULL, NameAttr=NULL)
```

Arguments

Data	data frame or matrix. Correspond to your data
nprod	integer. Number of products
nattr	integer. Number of attributes
nsub	integer. Number of subjects.
format	integer (1 or 2). See the description
NameProds	string vector with the names of the products (length must be nprod)
NameAttr	string vector with the names of attributes (length must be nattr)

Value

The arranged data for CATATIS and CLUSCATA function

See Also

[catatis](#), [cluscata](#), [change_cata_format2](#)

change_cata_format2 *Change format of CATA datasets to perform the package functions*

Description

CATATIS and CLUSCATA operate on data where the blocks of variables are merged horizontally. If you have a vertical format, you can use this function to change the format. The first column must contain the sessions, the second the subjects, the third the products and the others the attributes. If you don't have sessions, then the first column must contain the subjects and the second the products. Unlike change_cata_format function, you can enter data with sessions and/or mixed data in terms of products/subjects. However, you have to set columns to indicate this beforehand.

Usage

```
change_cata_format2(Data, nprod, nattr, nsub, nsess)
```

Arguments

Data	data frame or matrix. Correspond to your data
nprod	integer. Number of products
nattr	integer. Number of attributes
nsub	integer. Number of subjects.
nsess	integer. Number of sessions

Value

The arranged data for CATATIS and CLUSCATA function and the subjects names in the correct order.

See Also

[catatis](#), [cluscata](#), [change_cata_format](#)

Examples

```
#Vertical format with sessions
data("fish")
chang=change_cata_format2(fish, nprod= 6, nattr= 27, nsub = 12, nsess= 3)
res.cat2=catatis(Data= chang$Datafinal, nblo = 12, NameBlocks = chang$NameSub)

#Vertical format without sessions
Data=fish[1:66,2:30]
chang2=change_cata_format2(Data, nprod= 6, nattr= 27, nsub = 11, nsess= 1)
res.cat3=catatis(Data= chang2$Datafinal, nblo = 11, NameBlocks = chang2$NameSub)
res.clu3=cluscata(Data= chang2$Datafinal, nblo = 11, NameBlocks = chang2$NameSub)
```

cheese	<i>cheese Just About Right data</i>
--------	-------------------------------------

Description

cheese Just About Right data

Usage

```
data(cheese)
```

Format

JAR data. A data frame with Assessors, Products and JAR attributes. 8 products, 9 attributes and 72 subjects.

References

Luc, A., Lê, S., Philippe, M., Qannari, E. M., & Vigneau, E. (2022). Free JAR experiment: Data analysis and comparison with JAR task. *Food Quality and Preference*, 98, 104453.

Examples

```
data(cheese)
```

choc	<i>chocolates data</i>
------	------------------------

Description

chocolates data

Usage

```
data(choc)
```

Format

Free sorting data. A data frame with 14 rows (the chocolates) and 25 columns (the subjects). The numbers indicate the groups to which the products (rows) are assigned.

References

Courcoux, P., Qannari, E. M., Taylor, Y., Buck, D., & Greenhoff, K. (2012). Taxonomic free sorting. *Food Quality and Preference*, 23(1), 30-35.

Examples

```
data(choc)
```

cluscata

*Perform a cluster analysis of subjects from a CATA experiment***Description**

Clustering of subjects (blocks) from a CATA experiment. Each cluster of blocks is associated with a compromise computed by the CATATIS method. The hierarchical clustering is followed by a partitioning algorithm (consolidation). Non-binary data are accepted.

Usage

```
cluscata(Data, nblo, NameBlocks=NULL, NameVar=NULL, Noise_cluster=FALSE,
         Unique_threshold=TRUE, Itermax=30, Graph_dend=TRUE, Graph_bar=TRUE,
         printlevel=FALSE, gpmx=min(6, nblo-2), rhoparam=NULL,
         Testonlyoneclust=FALSE, alpha=0.05, nperm=50, Warnings=FALSE)
```

Arguments

Data	data frame or matrix where the blocks of variables (attributes) are merged horizontally. If you have a different format, see change_cata_format
nblo	numerical. Number of blocks (subjects).
NameBlocks	string vector. Name of each block (subject). Length must be equal to the number of blocks. If NULL, the names are S1,...Sm. Default: NULL
NameVar	string vector. Name of each variable (attribute, the same names for each subject). Length must be equal to the number of attributes. If NULL, the colnames of the first block are taken. Default: NULL
Noise_cluster	logical. Should a noise cluster be computed? Default: FALSE
Unique_threshold	logical. Use same rho for every cluster? Default: TRUE
Itermax	numerical. Maximum of iteration for the partitioning algorithm. Default:30
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging step of the hierarchical algorithm be plotted? Default: TRUE
printlevel	logical. Print the number of remaining levels during the hierarchical clustering algorithm? Default: FALSE
gpmx	logical. What is maximum number of clusters to consider? Default: min(6, nblo-2)
rhoparam	numerical or vector. What is the threshold for the noise cluster? Between 0 and 1, high value can imply lot of blocks set aside. If NULL, automatic threshold is computed. Can be different for each group (in this case, provide a vector)
Testonlyoneclust	logical. Test if there is more than one cluster? Default: FALSE

alpha	numerical between 0 and 1. What is the threshold to test if there is more than one cluster? Default: 0.05
nperm	numerical. How many permutations are required to test if there is more than one cluster? Default: 50
Warnings	logical. Display warnings about the fact that none of the subjects in some clusters checked an attribute or product? Default: FALSE

Value

Each partitionK contains a list for each number of clusters of the partition, K=1 to gpmx with:

- group: the clustering partition after consolidation. If Noise_cluster=TRUE, some subjects could be in the noise cluster ("K+1")
- rho: the threshold for the noise cluster
- homogeneity: homogeneity index (
- s_with_compromise: similarity coefficient of each subject with its cluster compromise
- weights: weight associated with each subject in its cluster
- compromise: the compromise of each cluster
- CA: list. the correspondance analysis results on each cluster compromise (coordinates, contributions...)
- inertia: percentage of total variance explained by each axis of the CA for each cluster
- s_all_cluster: the similarity coefficient between each subject and each cluster compromise
- criterion: the CLUSCATA criterion error
- param: parameters called
- type: parameter passed to other functions

There is also at the end of the list:

- dend: The CLUSCATA dendrogram
- cutree_k: the partition obtained by cutting the dendrogram in K clusters (before consolidation).
- overall_homogeneity_ng: percentage of overall homogeneity by number of clusters before consolidation (and after if there is no noise cluster)
- diff_crit_ng: variation of criterion when a merging is done before consolidation (and after if there is no noise cluster)
- test_one_cluster: decision and pvalue to know if there is more than one cluster
- param: parameters called
- type: parameter passed to other functions

References

- Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2019). A new approach for the analysis of data and the clustering of subjects in a CATA experiment. *Food Quality and Preference*, 72, 31-39.
- Llobell, F., Giacalone, D., Labenne, A., Qannari, E.M. (2019). Assessment of the agreement and cluster analysis of the respondents in a CATA experiment. *Food Quality and Preference*, 77, 184-190.

See Also

[plot.cluscata](#), [summary.cluscata](#), [catatis](#), [cluscata_kmeans](#), [change_cata_format](#), [change_cata_format2](#)

Examples

```
data(straw)
#with 40 subjects
res=cluscata(Data=straw[,1:(16*40)], nblo=40)
#plot(res, ngroups=3, Graph_dend=FALSE)
summary(res, ngroups=3)
#With noise cluster
res2=cluscata(Data=straw[,1:(16*40)], nblo=40, Noise_cluster=TRUE,
Graph_dend=FALSE, Graph_bar=FALSE)
#With noise cluster and defined rho threshold
#(high threshold for this example, you can put low threshold
#(ex: 0.2 or 0.3) to avoid set aside lot of respondents)
res3=cluscata(Data=straw[,1:(16*40)], nblo=40, Noise_cluster=TRUE,
Graph_dend=FALSE, Graph_bar=FALSE, rhoparam=0.6)
#different Noise cluster thresholds
res3=cluscata(Data=straw[,1:(16*40)], nblo=40, Noise_cluster=TRUE,
Graph_dend=FALSE, Graph_bar=FALSE, Unique_threshold= FALSE,
rhoparam=c(0.6, 0.5,0.4))
#with all subjects
res=cluscata(Data=straw, nblo=114, printlevel=TRUE)

#Vertical format
data("fish")
Data=fish[1:66,2:30]
chang2=change_cata_format2(Data, nprod= 6, nattr= 27, nsub = 11, nsess= 1)
res3=cluscata(Data= chang2$Datafinal, nblo = 11, NameBlocks = chang2$NameSub)
```

cluscata_jar

Perform a cluster analysis of subjects in a JAR experiment.

Description

Hierarchical clustering of subjects from a JAR experiment. Each cluster of subjects is associated with a compromise computed by the CATATIS method. The hierarchical clustering is followed by a partitioning algorithm (consolidation).

Usage

```
cluscata_jar(Data, nprod, nsub, levelsJAR=3, beta=0.1, Noise_cluster=FALSE,
Unique_threshold=TRUE, Itermax=30, Graph_dend=TRUE, Graph_bar=TRUE,
printlevel=FALSE, gpmx=min(6, nsub-2), rhoparam=NULL,
Testonlyoneclust=FALSE, alpha=0.05, nperm=50, Warnings=FALSE)
```

Arguments

Data	data frame where the first column is the Assessors, the second is the products and all other columns the JAR attributes with numbers (1 to 3 or 1 to 5, see levelsJAR)
nprod	integer. Number of products.
nsub	integer. Number of subjects.
levelsJAR	integer. 3 or 5 levels. If 5, the data will be transformed in 3 levels.
beta	numerical. Parameter for agreement between JAR and other answers. Between 0 and 0.5.
Noise_cluster	logical. Should a noise cluster be computed? Default: FALSE
Unique_threshold	logical. Use same rho for every cluster? Default: TRUE
Itermax	numerical. Maximum of iteration for the partitioning algorithm. Default:30
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging step of the hierarchical algorithm be plotted? Default: TRUE
printlevel	logical. Print the number of remaining levels during the hierarchical clustering algorithm? Default: FALSE
gpmx	logical. What is maximum number of clusters to consider? Default: min(6, nblo-2)
rhoparam	numerical or vector. What is the threshold for the noise cluster? Between 0 and 1, high value can imply lot of blocks set aside. If NULL, automatic threshold is computed. Can be different for each group (in this case, provide a vector)
Testonlyoneclust	logical. Test if there is more than one cluster? Default: FALSE
alpha	numerical between 0 and 1. What is the threshold to test if there is more than one cluster? Default: 0.05
nperm	numerical. How many permutations are required to test if there is more than one cluster? Default: 50
Warnings	logical. Display warnings about the fact that none of the subjects in some clusters checked an attribute or product? Default: FALSE

Value

Each partitionK contains a list for each number of clusters of the partition, K=1 to gpmx with:

- group: the clustering partition after consolidation. If Noise_cluster=TRUE, some subjects could be in the noise cluster ("K+1")
- rho: the threshold(s) for the noise cluster
- homogeneity: homogeneity index (
- s_with_compromise: similarity coefficient of each subject with its cluster compromise

- weights: weight associated with each subject in its cluster
- compromise: the compromise of each cluster
- CA: list. the correspondance analysis results on each cluster compromise (coordinates, contributions...)
- inertia: percentage of total variance explained by each axis of the CA for each cluster
- s_all_cluster: the similarity coefficient between each subject and each cluster compromise
- criterion: the CLUSCATA criterion error
- param: parameters called
- type: parameter passed to other functions

There is also at the end of the list:

- dend: The CLUSCATA dendrogram
- cutree_k: the partition obtained by cutting the dendrogram in K clusters (before consolidation).
- overall_homogeneity_ng: percentage of overall homogeneity by number of clusters before consolidation (and after if there is no noise cluster)
- diff_crit_ng: variation of criterion when a merging is done before consolidation (and after if there is no noise cluster)
- test_one_cluster: decision and pvalue to know if there is more than one cluster
- param: parameters called
- type: parameter passed to other functions

References

Llobell, F., Vigneau, E. & Qannari, E. M. ((September 14, 2022). Multivariate data analysis and clustering of subjects in a Just about right task. Eurosense, Turku, Finland.

See Also

[plot.cluscata](#), [summary.cluscata](#), [catatis_jar](#), [preprocess_JAR](#), [cluscata_kmeans_jar](#)

Examples

```
data(cheese)
res=cluscata_jar(Data=cheese, nprod=8, nsub=72, levelsJAR=5)
#plot(res, ngroups=4, Graph_dend=FALSE)
summary(res, ngroups=4)
```

cluscata_kmeans	<i>Compute the CLUSCATA partitioning algorithm on different blocks from a CATA experiment</i>
-----------------	---

Description

Partitioning of binary Blocks from a CATA experiment. Each cluster is associated with a compromise computed by the CATATIS method. Can be performed using a multi-start strategy or initial partition provided by the user. Moreover, a noise cluster can be set up.

Usage

```
cluscata_kmeans(Data,nblo, clust, nstart=100, rho=0, NameBlocks=NULL, NameVar=NULL,
                Itermax=30, Graph_groups=TRUE, print_attempt=FALSE, Warnings=FALSE)
```

Arguments

Data	data frame or matrix where the blocks of binary variables are merged horizontally. If you have a different format, see change_cata_format
nblo	numerical. Number of blocks (subjects).
clust	numerical vector or integer. Initial partition or number of starting partitions if integer. If numerical vector, the numbers must be 1,2,3,...,number of clusters
nstart	numerical. Number of starting partitions. Default: 100
rho	numerical or vector between 0 and 1. Threshold for the noise cluster. Default:0. If you want a different threshold for each cluster, you can provide a vector.
NameBlocks	string vector. Name of each block. Length must be equal to the number of blocks. If NULL, the names are S1,...,Sm. Default: NULL
NameVar	string vector. Name of each variable (attribute, the same names for each subject). Length must be equal to the number of attributes. If NULL, the colnames of the first block are taken. Default: NULL
Itermax	numerical. Maximum of iterations by partitioning algorithm. Default: 30
Graph_groups	logical. Should each cluster compromise graphical representation be plotted? Default: TRUE
print_attempt	logical. Print the number of remaining attempts in multi-start case? Default: FALSE
Warnings	logical. Display warnings about the fact that none of the subjects in some clusters checked an attribute or product? Default: FALSE

Value

a list with:

- group: the clustering partition. If $\rho > 0$, some subjects could be in the noise cluster ("K+1")
- rho: the threshold for the noise cluster

- homogeneity: percentage of homogeneity of the subjects in each cluster and the overall homogeneity
- s_with_compromise: Similarity coefficient of each subject with its cluster compromise
- weights: weight associated with each subject in its cluster
- compromise: The compromise of each cluster
- CA: The correspondance analysis results on each cluster compromise (coordinates, contributions...)
- inertia: percentage of total variance explained by each axis of the CA for each cluster
- s_all_cluster: the similarity coefficient between each subject and each cluster compromise
- param: parameters called
- criterion: the CLUSCATA criterion error
- type: parameter passed to other functions

References

Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2019). A new approach for the analysis of data and the clustering of subjects in a CATA experiment. *Food Quality and Preference*, 72, 31-39.

Llobell, F., Giacalone, D., Labenne, A., Qannari, E.M. (2019). Assessment of the agreement and cluster analysis of the respondents in a CATA experiment. *Food Quality and Preference*, 77, 184-190.

See Also

[plot.cluscata](#), [summary.cluscata](#), [catatis](#), [cluscata](#), [change_cata_format](#)

Examples

```
data(straw)
cl_km=cluscata_kmeans(Data=straw[,1:(16*40)], nblo=40, clust=3)
#plot(cl_km, Graph_groups=FALSE, Graph_weights = TRUE)
summary(cl_km)
```

cluscata_kmeans_jar	<i>Perform a cluster analysis of subjects in a JAR experiment</i>
---------------------	---

Description

Partitioning of subject from a JAR experiment. Each cluster is associated with a compromise computed by the CATATIS method. Moreover, a noise cluster can be set up.

Usage

```
cluscata_kmeans_jar(Data, nprod, nsub, levelsJAR=3, beta=0.1, clust, nstart=100, rho=0,
Itermax=30, Graph_groups=TRUE, print_attempt=FALSE, Warnings=FALSE)
```

Arguments

Data	data frame where the first column is the Assessors, the second is the products and all other columns the JAR attributes with numbers (1 to 3 or 1 to 5, see levelsJAR)
nprod	integer. Number of products.
nsub	integer. Number of subjects.
levelsJAR	integer. 3 or 5 levels. If 5, the data will be transformed in 3 levels.
beta	numerical. Parameter for agreement between JAR and other answers. Between 0 and 0.5.
clust	numerical vector or integer. Initial partition or number of starting partitions if integer. If numerical vector, the numbers must be 1,2,3,...,number of clusters
nstart	numerical. Number of starting partitions. Default: 100
rho	numerical or vector between 0 and 1. Threshold for the noise cluster. Default:0. If you want a different threshold for each cluster, you can provide a vector.
Itermax	numerical. Maximum of iterations by partitioning algorithm. Default: 30
Graph_groups	logical. Should each cluster compromise graphical representation be plotted? Default: TRUE
print_attempt	logical. Print the number of remaining attempts in multi-start case? Default: FALSE
Warnings	logical. Display warnings about the fact that none of the subjects in some clusters checked an attribute or product? Default: FALSE

Value

a list with:

- group: the clustering partition. If $\rho > 0$, some subjects could be in the noise cluster ("K+1")
- rho: the threshold(s) for the noise cluster
- homogeneity: percentage of homogeneity of the subjects in each cluster and the overall homogeneity
- s_with_compromise: Similarity coefficient of each subject with its cluster compromise
- weights: weight associated with each subject in its cluster
- compromise: The compromise of each cluster
- CA: The correspondance analysis results on each cluster compromise (coordinates, contributions...)
- inertia: percentage of total variance explained by each axis of the CA for each cluster
- s_all_cluster: the similarity coefficient between each subject and each cluster compromise
- param: parameters called
- criterion: the CLUSCATA criterion error
- type: parameter passed to other functions

References

Llobell, F., Vigneau, E. & Qannari, E. M. ((September 14, 2022). Multivariate data analysis and clustering of subjects in a Just about right task. Eurosense, Turku, Finland.

See Also

[plot.cluscata](#), [summary.cluscata](#), [catatis_jar](#), [preprocess_JAR](#), [cluscata_jar](#)

Examples

```
data(cheese)
res=cluscata_kmeans_jar(Data=cheese, nprod=8, nsub=72, levelsJAR=5, clust=4)
#plot(res)
summary(res)
```

cluscata_liking	<i>Perform a cluster analysis of subjects on CATA/liking combination</i>
-----------------	--

Description

Clustering of subjects (blocks) from combination of CATA and liking experiments.

Usage

```
cluscata_liking(Data, nblo, NameBlocks=NULL, NameVar=NULL, Itermax=30,
               Graph_dend=TRUE, Graph_bar=TRUE,
               printlevel=FALSE, gpmax=min(6, nblo-1))
```

Arguments

Data	data frame or matrix where the blocks of variables (attributes) are merged horizontally thanks to combinCATALiking function
nblo	numerical. Number of blocks (subjects).
NameBlocks	string vector. Name of each block (subject). Length must be equal to the number of blocks. If NULL, the names are S1,...Sm. Default: NULL
NameVar	string vector. Name of each variable (attribute, the same names for each subject). Length must be equal to the number of attributes. If NULL, the colnames of the first block are taken. Default: NULL
Itermax	numerical. Maximum of iteration for the partitioning algorithm. Default:30
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging step of the hierarchical algorithm be plotted? Default: TRUE

printlevel	logical. Print the number of remaining levels during the hierarchical clustering algorithm? Default: FALSE
gpmx	logical. What is maximum number of clusters to consider? Default: min(6, nblo-2)

Value

Each partitionK contains a list for each number of clusters of the partition, K=1 to gpmx with:

- group: the clustering partition after consolidation.
- compromise: the compromise of each cluster
- dist_all_cluster: the distance between each subject and each cluster compromise
- criterion: the CLUSCATA-liking criterion error
- param: parameters called
- type: parameter passed to other functions

There is also at the end of the list:

- dend: The CLUSCATA dendrogram
- cutree_k: the partition obtained by cutting the dendrogram in K clusters (before consolidation).
- diff_crit_ng: variation of criterion when a merging is done before consolidation (and after)
- param: parameters called
- type: parameter passed to other functions

References

Vigneau, E., Cariou, V., Giacalone, D., Berget, I., & Llobell, F. (2022). Combining hedonic information and CATA description for consumer segmentation. *Food Quality and Preference*, 95, 104358.

See Also

[plot.cluscata_liking](#), [summary.cluscata_liking](#), [combinCATALiking](#)

Examples

```
data(cata_ryebread)
data(liking_ryebread)
cataliking=combinCATALiking(cata_ryebread, liking_ryebread)

#with only 40 subjects
resclustcatal=cluscata_liking(Data=cataliking[,1:(40*14)], nblo=40, gpmx=5)
plot(resclustcatal, cata_ryebread[,1:(40*14)], liking_ryebread[,1:40])
```

cluscata_rata	<i>Perform a cluster analysis of subjects from a RATA experiment</i>
---------------	--

Description

Hierarchical clustering of subjects (blocks) from a RATA experiment. Each cluster of blocks is associated with a compromise computed by the CATATIS method. The hierarchical clustering is followed by a partitioning algorithm (consolidation).

Usage

```
cluscata_rata(Data, nblo, NameBlocks=NULL, NameVar=NULL, Noise_cluster=FALSE,
  Unique_threshold =TRUE, Itermax=30, Graph_dend=TRUE,
  Graph_bar=TRUE, printlevel=FALSE,
  gpmax=min(6, nblo-2), rhoparam=NULL, Testonlyoneclust=FALSE, alpha=0.05,
  nperm=50, Warnings=FALSE)
```

Arguments

Data	data frame or matrix where the blocks of binary variables are merged horizontally. If you have a different format, see change_cata_format
nblo	numerical. Number of blocks (subjects).
NameBlocks	string vector. Name of each block (subject). Length must be equal to the number of blocks. If NULL, the names are S1,...Sm. Default: NULL
NameVar	string vector. Name of each variable (attribute, the same names for each subject). Length must be equal to the number of attributes. If NULL, the colnames of the first block are taken. Default: NULL
Noise_cluster	logical. Should a noise cluster be computed? Default: FALSE
Unique_threshold	logical. Use same rho for every cluster? Default: TRUE
Itermax	numerical. Maximum of iteration for the partitioning algorithm. Default:30
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging step of the hierarchical algorithm be plotted? Default: TRUE
printlevel	logical. Print the number of remaining levels during the hierarchical clustering algorithm? Default: FALSE
gpmax	logical. What is maximum number of clusters to consider? Default: min(6, nblo-2)
rhoparam	numerical or vector. What is the threshold for the noise cluster? Between 0 and 1, high value can imply lot of blocks set aside. If NULL, automatic threshold is computed. Can be different for each group (in this case, provide a vector)
Testonlyoneclust	logical. Test if there is more than one cluster? Default: FALSE

alpha	numerical between 0 and 1. What is the threshold to test if there is more than one cluster? Default: 0.05
nperm	numerical. How many permutations are required to test if there is more than one cluster? Default: 50
Warnings	logical. Display warnings about the fact that none of the subjects in some clusters checked an attribute or product? Default: FALSE

Value

Each partitionK contains a list for each number of clusters of the partition, K=1 to gpmax with:

- group: the clustering partition after consolidation. If Noise_cluster=TRUE, some subjects could be in the noise cluster ("K+1")
- rho: the threshold(s) for the noise cluster
- homogeneity: homogeneity index (
- s_with_compromise: similarity coefficient of each subject with its cluster compromise
- weights: weight associated with each subject in its cluster
- compromise: the compromise of each cluster
- CA: list. the correspondance analysis results on each cluster compromise (coordinates, contributions...)
- inertia: percentage of total variance explained by each axis of the CA for each cluster
- s_all_cluster: the similarity coefficient between each subject and each cluster compromise
- criterion: the CLUSCATA criterion error
- param: parameters called
- type: parameter passed to other functions

There is also at the end of the list:

- dend: The CLUSCATA dendrogram
- cutree_k: the partition obtained by cutting the dendrogram in K clusters (before consolidation).
- overall_homogeneity_ng: percentage of overall homogeneity by number of clusters before consolidation (and after if there is no noise cluster)
- diff_crit_ng: variation of criterion when a merging is done before consolidation (and after if there is no noise cluster)
- test_one_cluster: decision and pvalue to know if there is more than one cluster
- param: parameters called
- type: parameter passed to other functions

References

Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2019). A new approach for the analysis of data and the clustering of subjects in a CATA experiment. Food Quality and Preference, 72, 31-39.

Llobell, F., Giacalone, D., Labenne, A., Qannari, E.M. (2019). Assessment of the agreement and cluster analysis of the respondents in a CATA experiment. Food Quality and Preference, 77, 184-190.

Llobell, F., Jaeger, S.R. (September 11, 2024). Consumer segmentation based on sensory product characterisations elicited by RATA questions? Eurosense conference, Dublin, Ireland.

See Also

[plot.cluscata](#), [summary.cluscata](#), [catatis_rata](#), [change_cata_format](#), [change_cata_format2](#)

Examples

```
#RATA data without session
data(RATAchoc)
Data=RATAchoc[1:108,2:16]
chang2=change_cata_format2(Data, nprod= 12, nattr= 13, nsub = 9, nsess = 1)
res.clus=cluscata_rata(Data= chang2$Datafinal, nblo = 9, NameBlocks = chang2$NameSub)
summary(res.clus)
plot(res.clus)
```

ClusMB	<i>Perform a cluster analysis of rows in a Multi-block context with the ClusMB method</i>
--------	---

Description

Clustering of rows (products in sensory analysis) in a Multi-block context. The hierarchical clustering is followed by a partitioning algorithm (consolidation).

Usage

```
ClusMB(Data, Blocks, NameBlocks=NULL, scale=FALSE, center=TRUE,
nclust=NULL, gpmx=6)
```

Arguments

- | | |
|------------|--|
| Data | data frame or matrix. Correspond to all the blocks of variables merged horizontally |
| Blocks | numerical vector. The number of variables of each block. The sum must be equal to the number of columns of Data. |
| NameBlocks | string vector. Name of each block. Length must be equal to the length of Blocks vector. If NULL, the names are B1,...Bm. Default: NULL |

scale	logical. Should the data variables be scaled? Default: FALSE
center	logical. Should the data variables be centered? Default: TRUE. Please set to FALSE for a CATA experiment
nclust	numerical. Number of clusters to consider. If NULL, the Hartigan index advice is taken.
gpmx	logical. What is maximum number of clusters to consider? Default: min(6, number of blocks -2)

Value

- group: the clustering partition after consolidation.
- nbgH: Advised number of clusters per Hartigan index
- nbgCH: Advised number of clusters per Calinski-Harabasz index
- cutree_k: the partition obtained by cutting the dendrogram in K clusters (before consolidation).
- dend: The ClusMB dendrogram
- param: parameters called
- type: parameter passed to other functions

References

Llobell, F., & Giacalone, D. (2025). Two Methods for Clustering Products in a Sensory Study: STATIS and ClusMB. *Journal of Sensory Studies*, 40(1), e70024.

Llobell, F., Qannari, E.M. (June 10, 2022). Cluster analysis in a multi-bloc setting. SMTDA, Athens, Greece.

Llobell, F., Giacalone, D., Qannari, E. M. (Pangborn 2021). Cluster Analysis of products in CATA experiments.

See Also

[indicesClusters](#), [summary.clusRows](#), [clustRowsOnStatisAxes](#)

Examples

```
#####projective mapping####
library(ClusBlock)
data(smoo)
res1=ClusMB(smoo, rep(2,24))
summary(res1)
indicesClusters(smoo, rep(2,24), res1$group)

####CATA####
data(fish)
Data=fish[1:66,2:30]
chang2=change_cata_format2(Data, nprod= 6, nattr= 27, nsub = 11, nsess= 1)
res2=ClusMB(Data= chang2$Datafinal, Blocks= rep(27, 11), center=FALSE)
indicesClusters(Data= chang2$Datafinal, Blocks= rep(27, 11), cut = res2$group, center=FALSE)
```

```
graphics.off()
```

clustatis

Perform a cluster analysis of blocks of quantitative variables

Description

Hierarchical clustering of quantitative Blocks followed by a partitioning algorithm (consolidation). Each cluster of blocks is associated with a compromise computed by the STATIS method. Moreover, a noise cluster can be set up.

Usage

```
clustatis(Data,Blocks,NameBlocks=NULL,Noise_cluster=FALSE,
  Unique_threshold=TRUE,scale=FALSE,
  Itermax=30, Graph_dend=TRUE, Graph_bar=TRUE,
  printlevel=FALSE, gpmax=min(6, length(Blocks)-2), rhoparam=NULL,
  Testonlyoneclust=FALSE, alpha=0.05, nperm=50)
```

Arguments

Data	data frame or matrix. Correspond to all the blocks of variables merged horizontally
Blocks	numerical vector. The number of variables of each block. The sum must be equal to the number of columns of Data
NameBlocks	string vector. Name of each block. Length must be equal to the length of Blocks vector. If NULL, the names are B1,...Bm. Default: NULL
Noise_cluster	logical. Should a noise cluster be computed? Default: FALSE
Unique_threshold	logical. Use same rho for every cluster? Default: TRUE
scale	logical. Should the data variables be scaled? Default: FALSE
Itermax	numerical. Maximum of iteration for the partitioning algorithm. Default: 30
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging step of the hierarchical algorithm be plotted? Default: TRUE
printlevel	logical. Print the number of remaining levels during the hierarchical clustering algorithm? Default: FALSE
gpmax	logical. What is maximum number of clusters to consider? Default: min(6, number of blocks -2)
rhoparam	numerical or vector. What is the threshold for the noise cluster? Between 0 and 1, high value can imply lot of blocks set aside. If NULL, automatic threshold is computed. Can be different for each group (in this case, provide a vector)

Testonlyoneclust

logical. Test if there is more than one cluster? Default: FALSE

alpha numerical between 0 and 1. What is the threshold to test if there is more than one cluster? Default: 0.05

nperm numerical. How many permutations are required to test if there is more than one cluster? Default: 50

Value

Each partitionK contains a list for each number of clusters of the partition, K=1 to gpmx with:

- group: the clustering partition of datasets after consolidation. If Noise_cluster=TRUE, some blocks could be in the noise cluster ("K+1")
- rho: the threshold(s) for the noise cluster (computed or input parameter)
- homogeneity: homogeneity index (
- rv_with_compromise: RV coefficient of each block with its cluster compromise
- weights: weight associated with each block in its cluster
- comp_RV: RV coefficient between the compromises associated with the various clusters
- compromise: the W compromise of each cluster
- coord: the coordinates of objects of each cluster
- inertia: percentage of total variance explained by each axis for each cluster
- rv_all_cluster: the RV coefficient between each block and each cluster compromise
- criterion: the CLUSTATIS criterion error
- param: parameters called in the consolidation
- type: parameter passed to other functions

There is also at the end of the list:

- dend: The CLUSTATIS dendrogram
- cutree_k: the partition obtained by cutting the dendrogram for K clusters (before consolidation).
- overall_homogeneity_ng: percentage of overall homogeneity by number of clusters before consolidation (and after if there is no noise cluster)
- diff_crit_ng: variation of criterion when a merging is done before consolidation (and after if there is no noise cluster)
- test_one_cluster: decision and pvalue to know if there is more than one cluster
- param: parameters called
- type: parameter passed to other functions

References

Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2018). Analysis and clustering of multiblock datasets by means of the STATIS and CLUSTATIS methods. Application to sensometrics. *Food Quality and Preference*, in Press.

Llobell, F., Vigneau, E., Qannari, E. M. (2019). Clustering datasets by means of CLUSTATIS with identification of atypical datasets. Application to sensometrics. *Food Quality and Preference*, 75, 97-104.

Llobell, F., & Qannari, E. M. (2020). CLUSTATIS: Cluster analysis of blocks of variables. *Electronic Journal of Applied Statistical Analysis*, 13(2).

See Also

[plot.clustatis](#), [summary.clustatis](#), [clustatis_kmeans](#), [statist](#)

Examples

```
data(smoo)
NameBlocks=paste0("S",1:24)
cl=clustatis(Data=smoo,Blocks=rep(2,24),NameBlocks = NameBlocks)
#plot(cl, ngroups=3, Graph_dend=FALSE)
summary(cl)
#with noise cluster
cl2=clustatis(Data=smoo,Blocks=rep(2,24),NameBlocks = NameBlocks,
Noise_cluster=TRUE, Graph_dend=FALSE, Graph_bar=FALSE)
#with noise cluster and defined rho threshold
cl3=clustatis(Data=smoo,Blocks=rep(2,24),NameBlocks = NameBlocks,
Noise_cluster=TRUE, Graph_dend=FALSE, Graph_bar=FALSE, rhoparam=0.5)
#different Noise cluster thresholds
cl4=clustatis(Data=smoo,Blocks=rep(2,24),NameBlocks = NameBlocks,
Noise_cluster=TRUE, Graph_dend=FALSE, Graph_bar=FALSE, Unique_threshold= FALSE,
rhoparam=c(0.6, 0.5,0.4))

graphics.off()
```

clustatis_FreeSort	<i>Perform a cluster analysis of free sorting data</i>
--------------------	--

Description

Hierarchical clustering of free sorting data followed by a partitioning algorithm (consolidation). Each cluster of blocks is associated with a compromise computed by the STATIS method. Moreover, a noise cluster can be set up.

Usage

```
clustatis_FreeSort(Data, NameSub=NULL, Noise_cluster=FALSE,
Unique_threshold = TRUE, Itermax=30,
Graph_dend=TRUE, Graph_bar=TRUE, printlevel=FALSE,
gpmx=min(6, ncol(Data)-1),rhoparam=NULL,
Testonlyoneclust=FALSE, alpha=0.05, nperm=50)
```

Arguments

Data	data frame or matrix. Corresponds to all variables that contain subjects results. Each column corresponds to a subject and gives the groups to which the products (rows) are assigned
NameSub	string vector. Name of each subject. Length must be equal to the number of column of the Data. If NULL, the names are S1,...Sm. Default: NULL
Noise_cluster	logical. Should a noise cluster be computed? Default: FALSE
Unique_threshold	logical. Use same rho for every cluster? Default: TRUE
Itermax	numerical. Maximum of iteration for the partitioning algorithm. Default: 30
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging be plotted? Default: FALSE
printlevel	logical. Print the number of remaining levels during the hierarchical clustering algorithm? Default: FALSE
gpmx	logical. What is maximum number of clusters to consider? Default: min(6, number of subjects -1)
rhoparam	numerical or vector. What is the threshold for the noise cluster? Between 0 and 1, high value can imply lot of blocks set aside. If NULL, automatic threshold is computed. Can be different for each group (in this case, provide a vector)
Testonlyoneclust	logical. Test if there is more than one cluster? Default: FALSE
alpha	numerical between 0 and 1. What is the threshold to test if there is more than one cluster? Default: 0.05
nperm	numerical. How many permutations are required to test if there is more than one cluster? Default: 50

Value

Each partitionK contains a list for each number of clusters of the partition, K=1 to gpmx with:

- group: the clustering partition of subjects after consolidation. If Noise_cluster=TRUE, some subjects could be in the noise cluster ("K+1")
- rho: the threshold(s) for the noise cluster
- homogeneity: homogeneity index (
- rv_with_compromise: RV coefficient of each block with its cluster compromise
- weights: weight associated with each subject in its cluster
- comp_RV: RV coefficient between the compromises associated with the various clusters
- compromise: the W compromise of each cluster
- coord: the coordinates of objects of each cluster
- inertia: percentage of total variance explained by each axis for each cluster
- rv_all_cluster: the RV coefficient between each subject and each cluster compromise

- criterion: the CLUSTATIS criterion error
- param: parameters called in the consolidation
- type: parameter passed to other functions

There is also at the end of the list:

- dend: The CLUSTATIS dendrogram
- cutree_k: the partition obtained by cutting the dendrogram for K clusters (before consolidation).
- overall_homogeneity_ng: percentage of overall homogeneity by number of clusters before consolidation (and after if there is no noise cluster)
- diff_crit_ng: variation of criterion when a merging is done before consolidation (and after if there is no noise cluster)
- test_one_cluster: decision and pvalue to know if there is more than one cluster
- param: parameters called
- type: parameter passed to other functions

References

Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2018). Analysis and clustering of multiblock datasets by means of the STATIS and CLUSTATIS methods. Application to sensometrics. Food Quality and Preference, in Press.

Llobell, F., Vigneau, E., Qannari, E. M. (2019). Clustering datasets by means of CLUSTATIS with identification of atypical datasets. Application to sensometrics. Food Quality and Preference, 75, 97-104.

See Also

[clustatis](#), [preprocess_FreeSort](#), [summary.clustatis](#), [plot.clustatis](#)

Examples

```
data(choc)
res.clu=clustatis_FreeSort(choc)
plot(res.clu, Graph_dend=FALSE)
summary(res.clu)
graphics.off()
```

clustatis_FreeSort_kmeans

Compute the CLUSTATIS partitioning algorithm on free sorting data

Description

partitioning algorithm for Free Sorting data. Each cluster is associated with a compromise computed by the STATIS method. Moreover, a noise cluster can be set up.

Usage

```
clustatis_FreeSort_kmeans(Data, NameSub=NULL, clust, nstart=100, rho=0, Itermax=30,
  Graph_groups=TRUE, Graph_weights=FALSE, print_attempt=FALSE)
```

Arguments

Data	data frame or matrix. Corresponds to all variables that contain subjects results. Each column corresponds to a subject and gives the groups to which the products (rows) are assigned
NameSub	string vector. Name of each subject. Length must be equal to the number of column of the Data. If NULL, the names are S1,...Sm. Default: NULL
clust	numerical vector or integer. Initial partition or number of starting partitions if integer. If numerical vector, the numbers must be 1,2,3,...,number of clusters
nstart	integer. Number of starting partitions. Default: 100
rho	numerical or vector between 0 and 1. Threshold for the noise cluster. Default:0. If you want a different threshold for each cluster, you can provide a vector.
Itermax	numerical. Maximum of iterations by partitioning algorithm. Default: 30
Graph_groups	logical. Should each cluster compromise be plotted? Default: TRUE
Graph_weights	logical. Should the barplot of the weights in each cluster be plotted? Default: FALSE
print_attempt	logical. Print the number of remaining attempts in the multi-start case? Default: FALSE

Value

a list with:

- group: the clustering partition. If $\rho > 0$, some subjects could be in the noise cluster ("K+1")
- rho: the threshold(s) for the noise cluster
- homogeneity: percentage of homogeneity of the subjects in each cluster and the overall homogeneity
- rv_with_compromise: RV coefficient of each subject with its cluster compromise
- weights: weight associated with each subject in its cluster

- `comp_RV`: RV coefficient between the compromises associated with the various clusters
- `compromise`: the W compromise of each cluster
- `coord`: the coordinates of objects of each cluster
- `inertia`: percentage of total variance explained by each axis for each cluster
- `rv_all_cluster`: the RV coefficient between each subject and each cluster compromise
- `criterion`: the CLUSTATIS criterion error
- `param`: parameters called
- `type`: parameter passed to other functions

References

Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2018). Analysis and clustering of multiblock datasets by means of the STATIS and CLUSTATIS methods. Application to sensometrics. Food Quality and Preference, in Press.

Llobell, F., Vigneau, E., Qannari, E. M. (2019). Clustering datasets by means of CLUSTATIS with identification of atypical datasets. Application to sensometrics. Food Quality and Preference, 75, 97-104.

See Also

[clustatis_FreeSort](#), [preprocess_FreeSort](#), [summary.clustatis](#), [plot.clustatis](#)

Examples

```
data(choc)
res.clu=clustatis_FreeSort_kmeans(choc, clust=2)
plot(res.clu, Graph_groups=FALSE, Graph_weights=TRUE)
summary(res.clu)
```

<code>clustatis_kmeans</code>	<i>Compute the CLUSTATIS partitioning algorithm on different blocks of quantitative variables</i>
-------------------------------	---

Description

Partitioning algorithm for quantitative variables. Each cluster is associated with a compromise computed by the STATIS method. Can be performed using a multi-start strategy or initial partition provided by the user. Moreover, a noise cluster can be set up.

Usage

```
clustatis_kmeans(Data, Blocks, clust, nstart=100, rho=0, NameBlocks=NULL,
Itermax=30, Graph_groups=TRUE, Graph_weights=FALSE,
scale=FALSE, print_attempt=FALSE)
```

Arguments

Data	data frame or matrix. Correspond to all the blocks of variables merged horizontally
Blocks	numerical vector. The number of variables of each block. The sum must be equal to the number of columns of Data
clust	numerical vector or integer. Initial partition or number of starting partitions if integer. If numerical vector, the numbers must be 1,2,3,...,number of clusters
nstart	integer. Number of starting partitions. Default: 100
rho	numerical or vector between 0 and 1. Threshold for the noise cluster. Default:0. If you want a different threshold for each cluster, you can provide a vector.
NameBlocks	string vector. Name of each block. Length must be equal to the length of Blocks vector. If NULL, the names are B1,...Bm. Default: NULL
Itermax	numerical. Maximum of iterations by partitioning algorithm. Default: 30
Graph_groups	logical. Should each cluster compromise be plotted? Default: TRUE
Graph_weights	logical. Should the barplot of the weights in each cluster be plotted? Default: FALSE
scale	logical. Should the data variables be scaled? Default: FALSE
print_attempt	logical. Print the number of remaining attempts in the multi-start case? Default: FALSE

Value

a list with:

- group: the clustering partition. If $\rho > 0$, some blocks could be in the noise cluster ("K+1")
- rho: the threshold(s) for the noise cluster
- homogeneity: percentage of homogeneity of the blocks in each cluster and the overall homogeneity
- rv_with_compromise: RV coefficient of each block with its cluster compromise
- weights: weight associated with each block in its cluster
- comp_RV: RV coefficient between the compromises associated with the various clusters
- compromise: the W compromise of each cluster
- coord: the coordinates of objects of each cluster
- inertia: percentage of total variance explained by each axis for each cluster
- rv_all_cluster: the RV coefficient between each block and each cluster compromise
- criterion: the CLUSTATIS criterion error
- param: parameters called
- type: parameter passed to other functions

References

- Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2018). Analysis and clustering of multiblock datasets by means of the STATIS and CLUSTATIS methods. Application to sensorimetrics. Food Quality and Preference, in Press.
- Llobell, F., Vigneau, E., Qannari, E. M. (2019). Clustering datasets by means of CLUSTATIS with identification of atypical datasets. Application to sensorimetrics. Food Quality and Preference, 75, 97-104.

See Also

[plot.clustatis](#), [clustatis](#), [summary.clustatis](#), [statis](#)

Examples

```
data(smoo)
NameBlocks=paste0("S",1:24)
#with multi-start
cl_km=clustatis_kmeans(Data=smoo,Blocks=rep(2,24),NameBlocks = NameBlocks, clust=3)
#with an initial partition
cl=clustatis(Data=smoo,Blocks=rep(2,24),NameBlocks = NameBlocks,
Graph_dend=FALSE)
partition=cl$cutree_k$partition3
cl_km2=clustatis_kmeans(Data=smoo,Blocks=rep(2,24),NameBlocks = NameBlocks,
clust=partition, Graph_weights=FALSE, Graph_groups=FALSE)
graphics.off()
```

`clustRowsOnStatisAxes` *Perform a cluster analysis of rows in a Multi-block context with clustering on STATIS axes*

Description

Clustering of rows (products in sensory analysis) in a Multi-block context. The STATIS method is followed by a hierarchical algorithm.

Usage

```
clustRowsOnStatisAxes(Data, Blocks, NameBlocks=NULL, scale=FALSE,
nclust=NULL, gpmx=6, ncomp=5)
```

Arguments

Data	data frame or matrix. Correspond to all the blocks of variables merged horizontally
Blocks	numerical vector. The number of variables of each block. The sum must be equal to the number of columns of Data.

NameBlocks	string vector. Name of each block. Length must be equal to the length of Blocks vector. If NULL, the names are B1,...Bm. Default: NULL
scale	logical. Should the data variables be scaled? Default: FALSE
nclust	numerical. Number of clusters to consider. If NULL, the Hartigan index advice is taken.
gpmx	logical. What is maximum number of clusters to consider? min(6, number of blocks -2)
ncomp	numerical. Number of axes to consider. Default:5

Value

- group: the clustering partition.
- nbgH: Advised number of clusters per Hartigan index
- nbgCH: Advised number of clusters per Calinski-Harabasz index
- cutree_k: the partition obtained by cutting the dendrogram in K clusters
- dend: The dendrogram
- param: parameters called
- type: parameter passed to other functions

References

Llobell, F., & Giacalone, D. (2025). Two Methods for Clustering Products in a Sensory Study: STATIS and ClusMB. *Journal of Sensory Studies*, 40(1), e70024.

See Also

[indicesClusters](#), [summary.clusRows](#), [ClusMB](#)

Examples

```
#####projective mapping####
library(ClustBlock)
data(smoo)
res1=clustRowsOnStatisAxes(smoo, rep(2,24))
summary(res1)
indicesClusters(smoo, rep(2,24), res1$group)

####CATA####
data(fish)
Data=fish[1:66,2:30]
chang2=change_cata_format2(Data, nprod= 6, nattr= 27, nsub = 11, nsess= 1)
res2=clustRowsOnStatisAxes(Data= chang2$Datafinal, Blocks= rep(27, 11))
indicesClusters(Data= chang2$Datafinal, Blocks= rep(27, 11),cut = res2$group, center=FALSE)
```

combinCATALiking	<i>Combination of CATA and liking data for CLUSCATA-liking</i>
------------------	--

Description

For CLUSCATA-liking, this preprocessing is needed.

Usage

```
combinCATALiking(cata, liking, center=TRUE, scale=FALSE)
```

Arguments

cata	data frame or matrix where the blocks of binary variables are merged horizontally. If you have a different format, see change_cata_format
liking	data frame or matrix where the products are in rows and the assessors in columns
center	Centering of consumer liking. Default: TRUE
scale	Scaling of consumer liking. Default: FALSE

Value

Combined data

References

Vigneau, E., Cariou, V., Giacalone, D., Berget, I., & Llobell, F. (2022). Combining hedonic information and CATA description for consumer segmentation. *Food Quality and Preference*, 95, 104358.

See Also

[cluscata_liking](#)

Examples

```
data(cata_ryebread)
data(liking_ryebread)
cataliking=combinCATALiking(cata_ryebread, liking_ryebread)
```

consistency_cata	<i>Test the consistency of each attribute in a CATA experiment</i>
------------------	--

Description

Permutation test on the agreement between subjects for each attribute in a CATA experiment

Usage

```
consistency_cata(Data,nblo, nperm=100, alpha=0.05, printAttrTest=FALSE)
```

Arguments

Data	data frame or matrix. Correspond to all the blocks of variables merged horizontally
nblo	numerical. Number of blocks (subjects).
nperm	numerical. How many permutations are required? Default: 100
alpha	numerical between 0 and 1. What is the threshold? Default: 0.05
printAttrTest	logical. Print the number of remaining attributes to be tested? Default: FALSE

Value

a list with:

- consist: the consistent attributes
- no_consist: the inconsistent attributes
- pval: pvalue for each test

References

Llobell, F., Giacalone, D., Labenne, A., Qannari, E.M. (2019). Assessment of the agreement and cluster analysis of the respondents in a CATA experiment. Food Quality and Preference, 77, 184-190.

See Also

[consistency_cata_panel](#), [change_cata_format](#), [change_cata_format2](#)

Examples

```
data(straw)
#with only 40 subjects
consistency_cata(Data=straw[,1:(16*40)], nblo=40)
#with all subjects
consistency_cata(Data=straw, nblo=114, printAttrTest=TRUE)
```

`consistency_cata_panel`*Test the consistency of the panel in a CATA experiment*

Description

Permutation test on the agreement between subjects in a CATA experiment

Usage

```
consistency_cata_panel(Data, nblo, nperm=100, alpha=0.05)
```

Arguments

<code>Data</code>	data frame or matrix. Correspond to all the blocks of variables merged horizontally
<code>nblo</code>	numerical. Number of blocks (subjects).
<code>nperm</code>	numerical. How many permutations are required? Default: 100
<code>alpha</code>	numerical between 0 and 1. What is the threshold? Default: 0.05

Value

a list with:

- `answer`: the answer of the test
- `pval`: pvalue of the test
- `dis`: distance between the homogeneity and the median of the permutations

References

Llobell, F., Giacalone, D., Labenne, A., Qannari, E.M. (2019). Assessment of the agreement and cluster analysis of the respondents in a CATA experiment. *Food Quality and Preference*, 77, 184-190.

Bonnet, L., Ferney, T., Riedel, T., Qannari, E.M., Llobell, F. (September 14, 2022) .Using CATA for sensory profiling: assessment of the panel performance. Eurosense, Turku, Finland.

See Also

[consistency_cata](#), [change_cata_format](#), [change_cata_format2](#)

Examples

```
data(straw)
#with all subjects
consistency_cata_panel(Data=straw, nblo=114)
```

croissant	<i>Croissant JAR and liking data</i>
-----------	--------------------------------------

Description

Croissant JAR and liking data

Usage

```
data(croissant)
```

Format

Just-About-Right (JAR) and liking data. A data frame with one row per subject-product combination. The first column identifies the subjects, the second column identifies the products, the following columns contain the JAR evaluations on a 5-point scale, and the last column contains the overall liking scores.

References

Llobell, F. & Guksch, T. (2026). Beyond penalty analysis: Joint clustering of consumers using JAR and liking data. EuroSense, Oslo, Norway.

Examples

```
data(croissant)
```

fish	<i>fish data</i>
------	------------------

Description

fish data

Usage

```
data(fish)
```

Format

CATA data with sessions. A data frame with the sessions, the panelists, the products and CATA attributes.

References

Bonnet, L., Ferney, T., Riedel, T., Qannari, E.M., Llobell, F. (September 14, 2022) .Using CATA for sensory profiling: assessment of the panel performance. Eurosense, Turku, Finland.

Examples

```
data(fish)
```

indicesClusters	<i>Compute the indices to evaluate the quality of the cluster partition in multi-block context</i>
-----------------	--

Description

Compute the II index to evaluate the agreement between each block and the global partition (in sensory: agreement between each subject and the global partition)

Compute the JI index to evaluate if each block has a partition (in sensory: if each subject made a partition of products)

Usage

```
indicesClusters(Data, Blocks, cut, NameBlocks=NULL, center=TRUE, scale=FALSE)
```

Arguments

Data	data frame or matrix. Correspond to all the blocks of variables merged horizontally
Blocks	numerical vector. The number of variables of each block. The sum must be equal to the number of columns of Data.
cut	numerical vector. The partition of the cluster analysis.
NameBlocks	string vector. Name of each block. Length must be equal to the length of Blocks vector. If NULL, the names are B1,...Bm. Default: NULL
center	logical. Should the data variables be centered? Default: TRUE. Please set to FALSE for a CATA experiment
scale	logical. Should the data variables be scaled? Default: FALSE

Value

- II: the II indices
- jI: the jI indices

References

Llobell, F., & Giacalone, D. (2025). Two Methods for Clustering Products in a Sensory Study: STATIS and ClusMB. *Journal of Sensory Studies*, 40(1), e70024.

Llobell, F., Qannari, E.M. (June 10, 2022). Cluster analysis in a multi-bloc setting. SMTDA, Athens, Greece.

Llobell, F., Giacalone, D., Qannari, E. M. (Pangborn 2021). Cluster Analysis of products in CATA experiments.

See Also

[clustRowsOnStatisAxes](#), [ClusMB](#)

Examples

```
#####projective mapping####
library(ClustBlock)
data(smoo)
res1=ClusMB(smoo, rep(2,24))
summary(res1)
indicesClusters(smoo, rep(2,24), res1$group)

####CATA####
data(fish)
Data=fish[1:66,2:30]
chang2=change_cata_format2(Data, nprod= 6, nattr= 27, nsub = 11, nsess= 1)
res2=ClusMB(Data= chang2$Datafinal, Blocks= rep(27, 11), center=FALSE)
indicesClusters(Data= chang2$Datafinal, Blocks= rep(27, 11),cut = res2$group, center=FALSE)
```

liking_ryebread

liking_ryebread data

Description

liking_ryebread data

Usage

```
data(liking_ryebread)
```

Format

Liking data. A data frame with 6 rows (the number of ryebread) and 132 columns (the number of consumers)

References

Giactalone, D. (2018). Product Performance Optimization. In Ares, G., & Varela, P. (Eds.) Methods in Consumer Research, Volume I (Chapter 7. pp. 159-185), Elsevier.

Examples

```
data(liking_ryebread)
```

mixclustatis

Perform a cluster analysis of variables

Description

Perform cluster analysis of variables in context of quantitative, qualitative or mixed datasets with MixCluStatis.

Usage

```
mixclustatis(Data, quanti=NULL,quali=NULL,Noise_cluster=FALSE,
  Itermax=20, printlevel=FALSE, Graph_dend=TRUE, Graph_bar=TRUE,
  gpmax=min(6, length(quali)+length(quanti)-1), rhoparam = NULL)
```

Arguments

Data	data frame or matrix. Correspond to all the data (variables are columns)
quanti	numerical vector. The number of the columns containing quantitative variables.
quali	numerical vector. The number of the columns containing qualitative variables.
Noise_cluster	logical. Should a noise cluster be computed? Default: FALSE
Itermax	numerical. Maximum of iteration for the partitioning algorithm. Default: 30
printlevel	logical. Print the number of remaining levels during the hierarchical clustering algorithm? Default: FALSE
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging step of the hierarchical algorithm be plotted? Default: TRUE
gpmax	logical. What is maximum number of clusters to consider? Default: min(6, number of variables -2)
rhoparam	numerical or vector. What is the threshold for the noise cluster? Between 0 and 1, high value can imply lot of variables set aside. If NULL, automatic threshold is computed. Can be different for each group (in this case, provide a vector)

Value

Each partitionK contains a list for each number of clusters of the partition, K=1 to gpmax with:

- group: the clustering partition of variables after consolidation. If Noise_cluster=TRUE, some variables could be in the noise cluster ("K+1")
- rho: the threshold(s) for the noise cluster (computed or input parameter)
- homogeneity: homogeneity index (
- rv_with_compromise: RV coefficient of each variable with its cluster compromise
- weights: weight associated with each variable in its cluster

- `comp_RV`: RV coefficient between the compromises associated with the various clusters
- `compromise`: the W compromise of each cluster
- `coord`: the coordinates of objects of each cluster
- `inertia`: percentage of total variance explained by each axis for each cluster
- `rv_all_cluster`: the RV coefficient between each variable and each cluster compromise
- `criterion`: the CLUSTATIS criterion error
- `param`: parameters called in the consolidation
- `type`: parameter passed to other functions

There is also at the end of the list:

- `dend`: The CLUSTATIS dendrogram
- `cutree_k`: the partition obtained by cutting the dendrogram for K clusters (before consolidation).
- `overall_homogeneity_ng`: percentage of overall homogeneity by number of clusters before consolidation (and after if there is no noise cluster)
- `diff_crit_ng`: variation of criterion when a merging is done before consolidation (and after if there is no noise cluster)
- `param`: parameters called
- `type`: parameter passed to other functions

References

Paper submitted: Llobell, F., Abdi, H., Eslami, A. (2026). Clustering of categorical and mixed data variables around latent variables. Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2018). Analysis and clustering of multiblock datasets by means of the STATIS and CLUSTATIS methods. Application to sensometrics. Food Quality and Preference, in Press.

Llobell, F., Vigneau, E., Qannari, E. M. (2019). Clustering datasets by means of CLUSTATIS with identification of atypical datasets. Application to sensometrics. Food Quality and Preference, 75, 97-104.

Llobell, F., & Qannari, E. M. (2020). CLUSTATIS: Cluster analysis of blocks of variables. Electronic Journal of Applied Statistical Analysis, 13(2).

See Also

[clustatis](#), [plot.clustatis](#), [summary.clustatis](#)

Examples

```
data("wine", package = "FactoMineR")
res=mixclustatis(wine, quanti = 3:29, quali = 1:2)
summary(res)
plot(res, Graph_groups = FALSE, Graph_weights = TRUE)
```

plot.catatis	<i>Displays the CATATIS graphs</i>
--------------	------------------------------------

Description

This function plots the CATATIS map and CATATIS weights

Usage

```
## S3 method for class 'catatis'
plot(x, Graph=TRUE, Graph_weights=TRUE, Graph_eig=TRUE,
     axes=c(1,2), tit="CATATIS", cex=1, col.obj="blue", col.attr="red", ...)
```

Arguments

x	object of class 'catatis'
Graph	logical. Show the graphical representation? Default: TRUE
Graph_weights	logical. Should the barplot of the weights be plotted? Default: TRUE
Graph_eig	logical. Should the barplot of the eigenvalues be plotted? Only with Graph=TRUE. Default: TRUE
axes	numerical vector (length 2). Axes to be plotted
tit	string. Title for the graphical representation. Default: 'CATATIS'
cex	numerical. Numeric character expansion factor; multiplied by par("cex") yields the final character size. NULL and NA are equivalent to 1.0.
col.obj	numerical or string. Color for the objects points. Default: "blue"
col.attr	numerical or string. Color for the attributes points. Default: "red"
...	further arguments passed to or from other methods

Value

the CATATIS map

See Also

[catatis](#)

Examples

```
data(straw)
res.cat=catatis(straw, nblo=114)
plot(res.cat, Graph_weights=FALSE, axes=c(1,3))
```

plot.cluscata	<i>Displays the CLUSCATA graphs</i>
---------------	-------------------------------------

Description

This function plots dendrogram, variation of the merging criterion, weights and CATATIS map of each cluster

Usage

```
## S3 method for class 'cluscata'
plot(x, ngroups=NULL, Graph_groups=TRUE, Graph_dend=TRUE,
     Graph_bar=FALSE, Graph_weights=FALSE, axes=c(1,2), cex=1,
     col.obj="blue", col.attr="red", ...)
```

Arguments

x	object of class 'cluscata'.
ngroups	number of groups to consider. Ignored for cluscata_kmeans results. Default: recommended number of clusters
Graph_groups	logical. Should each cluster compromise graphical representation be plotted? Default: TRUE
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging step of the hierarchical algorithm be plotted? Also available after consolidation if Noise_cluster=FALSE. Default: FALSE
Graph_weights	logical. Should the barplot of the weights in each cluster be plotted? Default: FALSE
axes	numerical vector (length 2). Axes to be plotted. Default: c(1,2)
cex	numerical. Numeric character expansion factor; multiplied by par("cex") yields the final character size. NULL and NA are equivalent to 1.0.
col.obj	numerical or string. Color for the objects points. Default: "blue"
col.attr	numerical or string. Color for the attributes points. Default: "red"
...	further arguments passed to or from other methods

Value

the CLUSCATA graphs

See Also

[cluscata](#), [cluscata_kmeans](#)

Examples

```
data(straw)
res=cluscata(Data=straw[,1:(16*40)], nblo=40)
plot(res, ngroups=3, Graph_dend=FALSE)
plot(res, ngroups=3, Graph_dend=FALSE, Graph_bar=FALSE, Graph_weights=FALSE, axes=c(1,3))
```

plot.cluscata_liking *Displays the CLUSCATA-liking graphs*

Description

This function plots dendrogram, variation of the merging criterion, weights and CATATIS map of each cluster

Usage

```
## S3 method for class 'cluscata_liking'
plot(x, DataCata, Dataliking, ngroups=NULL,
     center=TRUE, scale=FALSE,
     sep_graphs=TRUE, Graph_dend=TRUE, Graph_bar=FALSE,
     cex=1, xlimfreq=c(0,0.6), ylimchangelik=c(-2, 2), ...)
```

Arguments

x	object of class 'cluscata_liking'.
DataCata	data frame or matrix where the blocks of binary variables are merged horizontally. If you have a different format, see change_cata_format
Dataliking	data frame or matrix where the products are in rows and the assessors in columns
ngroups	number of groups to consider. Default: recommended number of clusters
center	Centering of consumer liking. Default: TRUE
scale	Scaling of consumer liking. Default: FALSE
sep_graphs	logical. Should cata and liking data analyzed separately first?
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging step of the hierarchical algorithm be plotted? Also available after consolidation if Noise_cluster=FALSE. Default: FALSE
cex	numerical. Numeric character expansion factor; multiplied by par("cex") yields the final character size. NULL and NA are equivalent to 1.0.
xlimfreq	vector of 2 values. Used for the graph of the frequency of citations
ylimchangelik	vector of 2 values. Used for the graphs of change in liking
...	further arguments passed to or from other methods

Value

the CLUSCATA-liking graphs

See Also

[cluscata_liking](#), [summary.cluscata_liking](#)

Examples

```
data(cata_ryebread)
data(liking_ryebread)
cataliking=combinCATALiking(cata_ryebread, liking_ryebread)

#with only 40 subjects
resclustcatal=cluscata_liking(Data=cataliking[,1:(40*14)], nblo=40, gpmax=5)
plot(resclustcatal, cata_ryebread[,1:(40*14)], liking_ryebread[,1:40],
xlimfreq=c(0,0.6), ylimchangelik=c(-2, 2))
```

plot.clusRows

Displays the ClusMB and clustRowsOnstatisAxes graphs

Description

This function plots the dendrogram of ClusMB or clustRowsOnstatisAxes

Usage

```
## S3 method for class 'clusRows'
plot(x, ...)
```

Arguments

x object of class 'clusRows'

... further arguments passed to or from other methods

Value

the dendrogram

See Also

[ClusMB](#), [clustRowsOnStatisAxes](#)

Examples

```
##'
####projective mapping####
library(ClustBlock)
data(smoo)
res1=ClusMB(smoo, rep(2,24))
plot(res1)

graphics.off()
```

plot.clustatis	<i>Displays the CLUSTATIS graphs</i>
----------------	--------------------------------------

Description

This function plots dendrogram, variation of the merging criterion, weights and STATIS map of each cluster

Usage

```
## S3 method for class 'clustatis'
plot(x, ngroups=NULL, Graph_groups=TRUE, Graph_dend=TRUE,
     Graph_bar=FALSE, Graph_weights=FALSE, axes=c(1,2), col=NULL, cex=1, font=1, ...)
```

Arguments

x	object of class 'clustatis'.
ngroups	number of groups to consider. Ignored for clustatis_kmeans results. Default: recommended number of clusters
Graph_groups	logical. Should each cluster compromise graphical representation be plotted? Default: TRUE
Graph_dend	logical. Should the dendrogram be plotted? Default: TRUE
Graph_bar	logical. Should the barplot of the difference of the criterion and the barplot of the overall homogeneity at each merging step of the hierarchical algorithm be plotted? Also available after consolidation if Noise_cluster=FALSE. Default: FALSE
Graph_weights	logical. Should the barplot of the weights in each cluster be plotted? Default: FALSE
axes	numerical vector (length 2). Axes to be plotted. Default: c(1,2)
col	vector. Color for each object. Default: rainbow(nrow(Data))
cex	numerical. Numeric character expansion factor; multiplied by par("cex") yields the final character size. NULL and NA are equivalent to 1.0.
font	numerical. Integer specifying font to use for text. 1=plain, 2=bold, 3=italic, 4=bold italic, 5=symbol. Default: 1
...	further arguments passed to or from other methods

Value

the CLUSTATIS graphs

See Also

[clustatis](#), [clustatis_kmeans](#)

Examples

```
data(smoo)
NameBlocks=paste0("S",1:24)
cl=clustatis(Data=smoo,Blocks=rep(2,24),NameBlocks = NameBlocks)
plot(cl, ngroups=3, Graph_dend=FALSE)
plot(cl, ngroups=3, Graph_dend=FALSE, axes=c(1,3))
graphics.off()
```

plot.stat	<i>Display the STATIS charts</i>
-----------	----------------------------------

Description

This function plots the STATIS map and STATIS weights

Usage

```
## S3 method for class 'stat'
plot(x, axes=c(1,2), Graph_obj=TRUE,
Graph_weights=TRUE, Graph_eig=TRUE, tit="STATIS", col=NULL, cex=1, font=1,
xlim=NULL, ylim=NULL, ...)
```

Arguments

x	object of class 'stat'
axes	numerical vector (length 2). Axes to be plotted. Default: c(1,2)
Graph_obj	logical. Should the compromise graphical representation be plotted? Default: TRUE
Graph_weights	logical. Should the barplot of the weights be plotted? Default: TRUE
Graph_eig	logical. Should the barplot of the eigenvalues be plotted? Only with Graph_obj=TRUE. Default: TRUE
tit	string. Title for the objects graphical representation. Default: 'STATIS'
col	vector. Color for each object. If NULL, col=rainbow(nrow(Data)). Default: NULL
cex	numerical. Numeric character expansion factor; multiplied by par("cex") yields the final character size. NULL and NA are equivalent to 1.0.

font	numerical. Integer specifying font to use for text. 1=plain, 2=bold, 3=italic, 4=bold italic, 5=symbol. Default: 1
xlim	numerical vector (length 2). Minimum and maximum for x coordinates.
ylim	numerical vector (length 2). Minimum and maximum for y coordinates.
...	further arguments passed to or from other methods

Value

the STATIS graphs

See Also

[statis](#)

Examples

```
data(smoo)
NameBlocks=paste0("S",1:24)
st=statis(Data=smoo,Blocks=rep(2,24),NameBlocks = NameBlocks)
plot(st, axes=c(1,3), Graph_weights=FALSE)
```

```
preprocess_FreeSort
```

Preprocessing for Free Sorting Data

Description

For Free Sorting Data, this preprocessing is needed.

Usage

```
preprocess_FreeSort(Data, NameSub=NULL)
```

Arguments

Data	data frame or matrix. Corresponds to all variables that contain subjects results. Each column corresponds to a subject and gives the groups to which the products (rows) are assigned
NameSub	string vector. Name of each subject. Length must be equal to the number of column of the Data. If NULL, the names are S1,...Sm. Default: NULL

Value

A list with:

- new_Data: the Data transformed
- Blocks: the number of groups for each subject
- NameBlocks: the name of each subject

References

Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2018). Analysis and clustering of multiblock datasets by means of the STATIS and CLUSTATIS methods. Application to sensometrics. Food Quality and Preference, in Press.

See Also

[clustatis](#), [clustatis_FreeSort](#)

Examples

```
data(choc)
prepro=preprocess_FreeSort(choc)
```

```
preprocess_JAR
```

Preprocessing for Just About Right Data

Description

For JAR data, this preprocessing is needed.

Usage

```
preprocess_JAR(Data, nprod, nsub, levelsJAR=3, beta=0.1)
```

Arguments

Data	data frame where the first column is the Assessors, the second is the products and all other columns the JAR attributes with numbers (1 to 3 or 1 to 5, see levelsJAR)
nprod	integer. Number of products.
nsub	integer. Number of subjects.
levelsJAR	integer. 3 or 5 levels. If 5, the data will be transformed in 3 levels.
beta	numerical. Parameter for agreement between JAR and other answers. Between 0 and 0.5.

Value

A list with:

- Datafinal: the Data transformed
- NameSub: the name of each subject in the right order

References

Llobell, F., Vigneau, E. & Qannari, E. M. (September 14, 2022). Multivariate data analysis and clustering of subjects in a Just about right task. Eurosense, Turku, Finland.

See Also

[catatis_jar](#), [cluscata_jar](#), [cluscata_kmeans_jar](#)

Examples

```
data(cheese)
prepro=preprocess_JAR(cheese, nprod=8, nsub=72, levelsJAR=5)
```

preprocess_JAR_liking *Preprocessing JAR and Liking Data for CLUSCATA-liking*

Description

Preprocesses Just-About-Right (JAR) and liking data for use with CLUSCATA-liking. JAR responses are converted into binary data, where 1 indicates a JAR response and 0 indicates a non-JAR response. The function also reshapes the JAR and liking data into the formats required by CLUSCATA-liking.

Usage

```
preprocess_JAR_liking(Data, nprod, nsub, liking_col,
                      levelsJAR = 3, scale = FALSE)
```

Arguments

Data	A data frame where the first column contains the subjects, the second column contains the products, and the remaining columns contain the JAR attributes and the liking variable.
nprod	Integer. Number of products.
nsub	Integer. Number of subjects.
liking_col	Character. Name of the column containing the liking scores.
levelsJAR	Integer. Number of levels of the JAR scale. Must be either 3 or 5. For a 3-level scale, level 2 is considered JAR. For a 5-level scale, level 3 is considered JAR. All other levels are considered non-JAR.
scale	Logical. Should the liking data be scaled when combining JAR and liking data? Default is FALSE.

Value

A list containing:

- Datafinal: combined JAR and liking data ready for CLUSCATA-liking.
- CATA: binary JAR data in CLUSCATA format, with 1 = JAR and 0 = non-JAR.
- liking: liking matrix with products in rows and subjects in columns.
- JAR_binary: binary JAR data before conversion to CLUSCATA format.
- NameSub: subject names in the order used in the analysis.
- NameProd: product names in the order used in the analysis.

References

Llobell, F. & Guksch, T. (2026). Beyond penalty analysis: Joint clustering of consumers using JAR and liking data. EuroSense, Oslo, Norway.

See Also

[cluscata_liking](#)

Examples

```
data(croissant)

# Use a subset of 40 subjects for a faster example
subjects <- unique(croissant[[1]])[1:40]
croissant40 <- croissant[croissant[[1]] %in% subjects, ]

prepro <- preprocess_JAR_liking(
  Data = croissant40,
  nprod = 6,
  nsub = 40,
  liking_col = "OVL_Overall Liking",
  levelsJAR = 5,
  scale = FALSE
)

CATA <- prepro$CATA
liking <- prepro$liking
Data <- prepro$Datafinal

res <- cluscata_liking(
  Data,
  nblo = 40,
  NameBlocks = prepro$NameSub,
  printlevel = FALSE
)

summary(res)

plot(
```

```
    res,  
    prepro$CATA,  
    prepro$liking,  
    scale = FALSE  
  )
```

print.catatis	<i>Print the CATATIS results</i>
---------------	----------------------------------

Description

Print the CATATIS results

Usage

```
## S3 method for class 'catatis'  
print(x, ...)
```

Arguments

x	object of class 'catatis'
...	further arguments passed to or from other methods

See Also

[catatis](#)

print.cluscata	<i>Print the CLUSCATA results</i>
----------------	-----------------------------------

Description

Print the CLUSCATA results

Usage

```
## S3 method for class 'cluscata'  
print(x, ...)
```

Arguments

x	object of class 'cluscata'
...	further arguments passed to or from other methods

See Also

[cluscata](#), [cluscata_kmeans](#)

`print.cluscata_liking` *Print the CLUSCATA-liking results*

Description

Print the CLUSCATA-liking results

Usage

```
## S3 method for class 'cluscata_liking'  
print(x, ...)
```

Arguments

<code>x</code>	object of class 'cluscata_liking'
<code>...</code>	further arguments passed to or from other methods

See Also

[cluscata_liking](#)

`print.clusRows` *Print the ClusMB or clustering on STATIS axes results*

Description

Print the ClusMB or clustering on STATIS axes results

Usage

```
## S3 method for class 'clusRows'  
print(x, ...)
```

Arguments

<code>x</code>	object of class 'clusRows'
<code>...</code>	further arguments passed to or from other methods

See Also

[ClusMB](#), [clustRowsOnStatisAxes](#)

print.clustatis	<i>Print the CLUSTATIS results</i>
-----------------	------------------------------------

Description

Print the CLUSTATIS results

Usage

```
## S3 method for class 'clustatis'  
print(x, ...)
```

Arguments

x	object of class 'clustatis'
...	further arguments passed to or from other methods

See Also

[clustatis](#), [clustatis_kmeans](#)

print.statist	<i>Print the STATIS results</i>
---------------	---------------------------------

Description

Print the STATIS results

Usage

```
## S3 method for class 'statist'  
print(x, ...)
```

Arguments

x	object of class 'statist'
...	further arguments passed to or from other methods

See Also

[statist](#)

RATAchoc	<i>RATA data on chocolates</i>
Description	
RATA data on chocolates	
Usage	
data(RATAchoc)	
Format	
RATA data with sessions. A data frame with 3 sessions, 9 panelists, 12 products and 27 RATA attributes.	
References	
Pangborn 2023	
Examples	
data(RATAchoc)	
simil_groups_cata	<i>Testing the difference in perception between two predetermined groups of subjects in a CATA experiment</i>

Description	
Test adapted to CATA data to determine whether two predetermined groups of subjects have a different perception or not. For example, men and women.	
Usage	
simil_groups_cata(Data, groups, one=1, two=2, nperm=50, Graph=TRUE, alpha= 0.05, printl=FALSE)	
Arguments	
Data	data frame or matrix. Correspond to all the blocks of variables merged horizontally
groups	categorical vector. The groups of each subject . The length must be the number of subjects.
one	string. Name of the group 1 in groups vector.

two	string. Name of the group 2 in groups vector.
nperm	numerical. How many permutations are required? Default: 50
Graph	logical. Should the CATATIS graph of each group be plotted? Default: TRUE
alpha	numerical between 0 and 1. What is the threshold of the test? Default: 0.05
printl	logical. Print the number of remaining permutations during the algorithm? Default: FALSE

Value

a list with:

- decision: the decision of the test
- pval: pvalue of the test

References

Llobell, F., Giacalone, D., Jaeger, S.R. & Qannari, E. M. (2021). CATA data: Are there differences in perception? JSM conference.

Llobell, F., Giacalone, D., Jaeger, S.R. & Qannari, E. M. (2021). CATA data: Are there differences in perception? AgroStat conference.

Examples

```
data(straw)
groups=sample(1:2, 114, replace=TRUE)
simil_groups_cata(straw, groups, one=1, two=2)
```

smoo

smoothies data

Description

smoothies data

Usage

```
data(smoo)
```

Format

Projective mapping (or Napping) data. A data frame with 8 rows (the number of smoothies) and 48 columns (the number of consumers * 2). For each consumer, we have the coordinates of the products on the sheet of paper.

References

Francois Husson, Sebastien Le and Marine Cadoret (2017). SensoMineR: Sensory Data Analysis. R package version 1.23. <https://CRAN.R-project.org/package=SensoMineR>

Examples

```
data(smoo)
```

statis	<i>Performs the STATIS method on different blocks of quantitative variables</i>
--------	---

Description

STATIS method on quantitative blocks. SUPplementary outputs are also computed

Usage

```
statis(Data,Blocks,NameBlocks=NULL,Graph_obj=TRUE, Graph_weights=TRUE, scale=FALSE)
```

Arguments

Data	data frame or matrix. Correspond to all the blocks of variables merged horizontally
Blocks	numerical vector. The number of variables of each block. The sum must be equal to the number of columns of Data
NameBlocks	string vector. Name of each block. Length must be equal to the length of Blocks vector. If NULL, the names are B1,...Bm. Default: NULL
Graph_obj	logical. Show the graphical representation of the objects? Default: TRUE
Graph_weights	logical. Should the barplot of the weights be plotted? Default: TRUE
scale	logical. Should the data variables be scaled? Default: FALSE

Value

a list with:

- RV: the RV matrix: a matrix with the RV coefficient between blocks of variables
- compromise: a matrix which is the compromise of the blocks (akin to a weighted average)
- weights: the weights associated with the blocks to build the compromise
- lambda: the first eigenvalue of the RV matrix
- overall error : the error for the STATIS criterion
- error_by_conf: the error by configuration (STATIS criterion)
- rv_with_compromise: the RV coefficient of each block with the compromise
- homogeneity: homogeneity of the blocks (in percentage)

- coord: the coordinates of each object
- eigenvalues: the eigenvalues of the svd decomposition
- inertia: the percentage of total variance explained by each axis
- error_by_obj: the error by object (STATIS criterion)
- scalefactors: the scaling factors of each block
- proj_config: the projection of each object of each configuration on the axes: presentation by configuration
- proj_objects: the projection of each object of each configuration on the axes: presentation by object

References

- Lavit, C., Escoufier, Y., Sabatier, R., Traissac, P. (1994). The act (statis method). Computational 462 Statistics & Data Analysis, 18 (1), 97-119.
- Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2018). Analysis and clustering of multiblock datasets by means of the STATIS and CLUSTATIS methods. Application to sensometrics. Food Quality and Preference, in Press.

See Also

[plot.statis](#), [clustatis](#)

Examples

```
data(smoo)
NameBlocks=paste0("S",1:24)
st=statis(Data=smoo, Blocks=rep(2,24),NameBlocks = NameBlocks)
#plot(st, axes=c(1,3))
summary(st)
#with variables scaling
st2=statis(Data=smoo, Blocks=rep(2,24),NameBlocks = NameBlocks, Graph_weights=FALSE, scale=TRUE)
```

statis_FreeSort

Performs the STATIS method on Free Sorting data

Description

STATIS method on Free Sorting data. A lot of supplementary informations are also computed

Usage

```
statis_FreeSort(Data, NameSub=NULL, Graph_obj=TRUE, Graph_weights=TRUE)
```

Arguments

Data	data frame or matrix. Corresponds to all variables that contain subjects results. Each column corresponds to a subject and gives the groups to which the products (rows) are assigned
NameSub	string vector. Name of each subject. Length must be equal to the number of column of the Data. If NULL, the names are S1,...Sm. Default: NULL
Graph_obj	logical. Show the graphical representation of the objects? Default: TRUE
Graph_weights	logical. Should the barplot of the weights be plotted? Default: TRUE

Value

a list with:

a list with:

- RV: the RV matrix: a matrix with the RV coefficient between subjects
- compromise: a matrix which is the compromise of the subjects (akin to a weighted average)
- weights: the weights associated with the subjects to build the compromise
- lambda: the first eigenvalue of the RV matrix
- overall error : the error for the STATIS criterion
- error_by_conf: the error by configuration (STATIS criterion)
- rv_with_compromise: the RV coefficient of each subject with the compromise
- homogeneity: homogeneity of the subjects (in percentage)
- coord: the coordinates of each object
- eigenvalues: the eigenvalues of the svd decomposition
- inertia: the percentage of total variance explained by each axis
- error_by_obj: the error by object (STATIS criterion)
- scalefactors: the scaling factors of each subject
- proj_config: the projection of each object of each subject on the axes: presentation by subject
- proj_objects: the projection of each object of each subject on the axes: presentation by object

References

- Lavit, C., Escoufier, Y., Sabatier, R., Traissac, P. (1994). The act (statis method). Computational 462 Statistics & Data Analysis, 18 (1), 97-119.
- Llobell, F., Cariou, V., Vigneau, E., Labenne, A., & Qannari, E. M. (2018). Analysis and clustering of multiblock datasets by means of the STATIS and CLUSTATIS methods. Application to sensometrics. Food Quality and Preference, in Press.

See Also

[preprocess_FreeSort](#), [clustatis_FreeSort](#)

Examples

```
data(choc)
res.sta=statis_FreeSort(choc)
```

straw	<i>strawberries data</i>
-------	--------------------------

Description

strawberries data

Usage

```
data(straw)
```

Format

CATA data. A data frame with 6 rows (the number of strawberries) and 1824 columns (the number of consumers (114) * the number of attributes (16)). For each consumer, each attribute and each product, there is 1 if the attribute has been checked by the consumer for the product, and 0 if not.

References

Ares, G., & Jaeger, S. R. (2013). Check-all-that-apply questions: Influence of attribute order on sensory product characterization. *Food Quality and Preference*, 28(1), 141-153.

Examples

```
data(straw)
```

summary.catatis	<i>Show the CATATIS results</i>
-----------------	---------------------------------

Description

This function shows the CATATIS results

Usage

```
## S3 method for class 'catatis'
summary(object, ...)
```

Arguments

object	object of class 'catatis'.
...	further arguments passed to or from other methods

Value

a list with:

- homogeneity: homogeneity of the subjects (in percentage)
- weights: the weights associated with the subjects to build the compromise
- eigenvalues: the eigenvalues associated to the correspondance analysis
- inertia: the percentage of total variance explained by each axis of the CA

See Also

[catatis](#)

summary.cluscata

Show the CLUSCATA results

Description

This function shows the cluscata results

Usage

```
## S3 method for class 'cluscata'
summary(object, ngroups=NULL, ...)
```

Arguments

object	object of class 'cluscata'.
ngroups	number of groups to consider. Ignored for cluscata_kmeans results. Default: recommended number of clusters
...	further arguments passed to or from other methods

Value

the CLUSCATA principal results

a list with:

- group: the clustering partition
- homogeneity: homogeneity index (
- weights: weight associated with each subject in its cluster
- rho: the threshold for the noise cluster
- test_one_cluster: decision and pvalue to know if there is more than one cluster

See Also

[cluscata](#), [cluscata_kmeans](#)

`summary.cluscata_liking`*Show the CLUSCATA-liking results*

Description

This function shows the cluscata_liking results

Usage

```
## S3 method for class 'cluscata_liking'  
summary(object, ngroups=NULL, ...)
```

Arguments

object	object of class 'cluscata_liking'.
ngroups	number of groups to consider. Default: recommended number of clusters
...	further arguments passed to or from other methods

Value

the CLUSCATA-liking principal results

a list with:

- group: the clustering partition
- groupsvector: the groups in a vector

See Also

[cluscata_liking](#), [plot.cluscata_liking](#)

`summary.clusRows`*Show the ClusMB or clustering on STATIS axes results*

Description

This function shows the ClusMB or clustering on STATIS axes results

Usage

```
## S3 method for class 'clusRows'  
summary(object, ...)
```

Arguments

object object of class 'clusRows'.
 ... further arguments passed to or from other methods

Value

a list with:

- groups: clustering partition
- nbClustRetained: the number of clusters retained
- nbgH: Advised number of clusters per Hartigan index
- nbgCH: Advised number of clusters per Calinski-Harabasz index

See Also

[ClusMB](#), [clustRowsOnStatisAxes](#)

summary.clustatis	<i>Show the CLUSTATIS results</i>
-------------------	-----------------------------------

Description

This function shows the clustatis results

Usage

```
## S3 method for class 'clustatis'
summary(object, ngroups=NULL, ...)
```

Arguments

object object of class 'clustatis'.
 ngroups number of groups to consider. Ignored for clustatis_kmeans results. Default: recommended number of clusters
 ... further arguments passed to or from other methods

Value

the CLUSTATIS principal results

a list with:

- group: the clustering partition
- homogeneity: homogeneity index (
- weights: weight associated with each block in its cluster
- rho: the threshold for the noise cluster
- test_one_cluster: decision and pvalue to know if there is more than one cluster

See Also[clustat](#), [clustat_kmeans](#)

summary.stat	Show the STATIS results
--------------	-------------------------

Description

This function shows the STATIS results

Usage

```
## S3 method for class 'stat'  
summary(object, ...)
```

Arguments

object	object of class 'stat'.
...	further arguments passed to or from other methods

Value

a list with:

- homogeneity: homogeneity of the blocks (in percentage)
- weights: the weights associated with the blocks to build the compromise
- eigenvalues: the eigenvalues of the svd decomposition
- inertia: the percentage of total variance explained by each axis

See Also[stat](#)

Index

- * **Blocks**
 - ClustBlock-package, 3
- * **CATA**
 - catatis, 4
 - change_cata_format, 10
 - change_cata_format2, 11
 - cluscata, 13
 - cluscata_kmeans, 18
 - cluscata_liking, 21
 - ClusMB, 25
 - ClustBlock-package, 3
 - clustRowsOnStatisAxes, 35
 - combinCATALiking, 37
 - consistency_cata, 38
 - consistency_cata_panel, 39
 - indicesClusters, 41
 - plot.catatis, 45
 - plot.cluscata, 46
 - plot.cluscata_liking, 47
 - plot.clusRows, 48
 - print.catatis, 55
 - print.cluscata, 55
 - print.cluscata_liking, 56
 - print.clusRows, 56
 - simil_groups_cata, 58
 - summary.catatis, 63
 - summary.cluscata, 64
 - summary.cluscata_liking, 65
 - summary.clusRows, 65
- * **Categorical**
 - ClustBlock-package, 3
- * **Check-All-That-Apply**
 - ClustBlock-package, 3
- * **Cluster analysis**
 - ClustBlock-package, 3
- * **Clustering**
 - ClustBlock-package, 3
- * **Free Sorting**
 - ClustBlock-package, 3
- * **FreeSorting**
 - clustatis_FreeSort, 29
 - clustatis_FreeSort_kmeans, 32
 - preprocess_FreeSort, 51
 - statis_FreeSort, 61
- * **Hedonic**
 - ClustBlock-package, 3
- * **JAR**
 - catatis_jar, 6
 - cluscata_jar, 15
 - cluscata_kmeans_jar, 19
 - ClustBlock-package, 3
 - preprocess_JAR, 52
 - preprocess_JAR_liking, 53
- * **Just About Right**
 - ClustBlock-package, 3
- * **Liking**
 - ClustBlock-package, 3
- * **Multi-block**
 - ClustBlock-package, 3
- * **Noise Cluster**
 - ClustBlock-package, 3
- * **Profiling**
 - ClustBlock-package, 3
- * **Projective mapping**
 - ClustBlock-package, 3
- * **Qualitative**
 - ClustBlock-package, 3
- * **RATA**
 - catatis_rata, 8
 - change_cata_format2, 11
 - cluscata_rata, 23
 - ClusMB, 25
 - ClustBlock-package, 3
 - clustRowsOnStatisAxes, 35
 - consistency_cata, 38
 - consistency_cata_panel, 39
 - indicesClusters, 41
 - plot.catatis, 45

- plot.clusRows, 48
- print.catatis, 55
- print.clusRows, 56
- summary.catatis, 63
- summary.clusRows, 65
- * **Rate-All-That-Apply**
 - ClustBlock-package, 3
- * **Sensory**
 - ClustBlock-package, 3
- * **Variables**
 - ClustBlock-package, 3
- * **clustering**
 - preprocess_JAR_liking, 53
- * **datasets**
 - cata_ryebread, 9
 - cheese, 12
 - choc, 12
 - croissant, 40
 - fish, 40
 - liking_ryebread, 42
 - RATAchoc, 58
 - smoo, 59
 - straw, 63
- * **liking**
 - cluscata_liking, 21
 - combinCATALiking, 37
 - plot.cluscata_liking, 47
 - preprocess_JAR_liking, 53
 - print.cluscata_liking, 56
 - summary.cluscata_liking, 65
- * **mixed**
 - ClustBlock-package, 3
- * **quantitative**
 - ClusMB, 25
 - clustatis, 27
 - clustatis_kmeans, 33
 - clustRowsOnStatisAxes, 35
 - indicesClusters, 41
 - plot.clusRows, 48
 - plot.clustatis, 49
 - plot.statis, 50
 - print.clusRows, 56
 - print.clustatis, 57
 - print.statis, 57
 - statis, 60
 - summary.clusRows, 65
 - summary.clustatis, 66
 - summary.statis, 67
- * **variables**
 - mixclustatis, 43
- cata_ryebread, 9
- catatis, 4, 7, 9, 11, 15, 19, 45, 55, 64
- catatis_jar, 6, 17, 21, 53
- catatis_rata, 8, 25
- change_cata_format, 5, 6, 8, 9, 10, 11, 13, 15, 18, 19, 23, 25, 37–39, 47
- change_cata_format2, 6, 9, 11, 11, 15, 25, 38, 39
- cheese, 12
- choc, 12
- cluscata, 6, 11, 13, 19, 46, 55, 64
- cluscata_jar, 7, 15, 21, 53
- cluscata_kmeans, 15, 18, 46, 55, 64
- cluscata_kmeans_jar, 7, 17, 19, 53
- cluscata_liking, 21, 37, 48, 54, 56, 65
- cluscata_rata, 23
- ClusMB, 25, 36, 42, 48, 56, 66
- clustatis, 27, 31, 35, 44, 50, 52, 57, 61, 67
- clustatis_FreeSort, 29, 33, 52, 62
- clustatis_FreeSort_kmeans, 32
- clustatis_kmeans, 29, 33, 50, 57, 67
- ClustBlock (ClustBlock-package), 3
- ClustBlock-package, 3
- clustRowsOnStatisAxes, 26, 35, 42, 48, 56, 66
- combinCATALiking, 21, 22, 37
- consistency_cata, 38, 39
- consistency_cata_panel, 38, 39
- croissant, 40
- fish, 40
- indicesClusters, 26, 36, 41
- liking_ryebread, 42
- mixclustatis, 43
- plot.catatis, 6, 7, 9, 45
- plot.cluscata, 15, 17, 19, 21, 25, 46
- plot.cluscata_liking, 22, 47, 65
- plot.clusRows, 48
- plot.clustatis, 29, 31, 33, 35, 44, 49
- plot.statis, 50, 61
- preprocess_FreeSort, 31, 33, 51, 62
- preprocess_JAR, 7, 17, 21, 52
- preprocess_JAR_liking, 53

print.catatis, [55](#)
print.cluscata, [55](#)
print.cluscata_liking, [56](#)
print.clusRows, [56](#)
print.clustatis, [57](#)
print.statis, [57](#)

RATAchoc, [58](#)

simil_groups_cata, [58](#)
smoo, [59](#)
statis, [29](#), [35](#), [51](#), [57](#), [60](#), [67](#)
statis_FreeSort, [61](#)
straw, [63](#)
summary.catatis, [6](#), [7](#), [9](#), [63](#)
summary.cluscata, [15](#), [17](#), [19](#), [21](#), [25](#), [64](#)
summary.cluscata_liking, [22](#), [48](#), [65](#)
summary.clusRows, [26](#), [36](#), [65](#)
summary.clustatis, [29](#), [31](#), [33](#), [35](#), [44](#), [66](#)
summary.statis, [67](#)