

Package ‘text2map’

September 15, 2026

Type Package

Title R Tools for Text Matrices, Embeddings, and Networks

Version 0.4.0

Author Dustin Stoltz [aut, cre] (ORCID:
<<https://orcid.org/0000-0002-4774-0765>>),
Marshall Taylor [aut] (ORCID: <<https://orcid.org/0000-0002-7440-0723>>)

Description This is a collection of functions optimized for working with various kinds of text matrices. Focusing on the text matrix as the primary object - represented either as a base R dense matrix or a 'Matrix' package sparse matrix - allows for a consistent and intuitive interface that stays close to the underlying mathematical foundation of computational text analysis. In particular, the package includes functions for working with word embeddings, text networks, and document-term matrices. Methods developed in Stoltz and Taylor (2019) <[doi:10.1007/s42001-019-00048-6](https://doi.org/10.1007/s42001-019-00048-6)>, Taylor and Stoltz (2020) <[doi:10.1007/s42001-020-00075-8](https://doi.org/10.1007/s42001-020-00075-8)>, Taylor and Stoltz (2020) <[doi:10.15195/v7.a23](https://doi.org/10.15195/v7.a23)>, and Stoltz and Taylor (2021) <[doi:10.1016/j.poetic.2021.101567](https://doi.org/10.1016/j.poetic.2021.101567)>.

URL <https://culturalcartography.gitlab.io/text2map>

License MIT + file LICENSE

Encoding UTF-8

Language en-US

BugReports <https://gitlab.com/culturalcartography/text2map/-/issues>

Depends R (>= 4.1.0), Matrix (>= 1.4.2)

Imports parallel, doParallel, foreach, stringi, dplyr, kit, fastmatch, methods, rlang, tibble, rsvd, permute, cli, pillar, tidyselect, text2vec, igraph (>= 1.2.6), RhpcBLASctl

Suggests testthat (>= 3.2.0), covr, spelling, lintr, tm, quanteda, knitr, rmarkdown, ClusterR, RSpectra

Config/testthat/edition 3

Config/Needs/website rmarkdown
Config/roxygen2/version 8.0.0
NeedsCompilation no
Maintainer Dustin Stoltz <dss219@lehhigh.edu>
Repository CRAN
Date/Publication 2026-09-15 11:00:02 UTC

Contents

anchor_lists	3
CMDist	3
CoCA	6
doc_centrality	9
doc_similarity	11
dtm_builder	12
dtm_melter	15
dtm_resampler	16
dtm_stats	17
dtm_stopper	18
find_projection	21
find_rejection	22
find_transformation	22
ft_wv_sample	24
get_anchors	25
get_centroid	26
get_centroids	27
get_direction	28
get_regions	30
get_stoplist	32
jfk_speech	33
meta_shakespeare	34
perm_tester	35
plot.CMDist	37
plot.CoCA	39
print.CMDist	40
print.CoCA	41
rancors_builder	42
rancor_builder	44
seq_builder	46
stoplists	47
test_anchors	49
tiny_gender_tagger	51
vocab_builder	52
weight_ppmi	53

Index	54
--------------	-----------

anchor_lists

A dataset of anchor lists

Description

A dataset containing juxtaposing pairs of English words for 30 semantic relations. These anchors are used with the [get_anchors\(\)](#) function, which can then be used with the [get_direction\(\)](#) function. These have been collected from previously published articles and should be used as a starting point for defining a given relation in a word embedding model.

Format

A data frame with 303 rows and 4 variables.

Variables

Variables:

- pole1. words to be added (or the positive direction)
- pole2. words to be subtracted (or the negative direction)
- relation. the relation to be extracted, 30 relations available
- domain. 6 broader categories within which each relation falls

Source

Curated from previously published semantic direction anchor sets. See: Kozlowski et al. (2019) "The Geometry of Culture"; Garg et al. (2018) "Word Embeddings Quantify 100 Years of Gender and Ethnic Stereotypes"; Bolukbasi et al. (2016) "Man is to Computer Programmer as Woman is to Homemaker".

See Also

[CoCA](#), [get_direction](#), [get_centroid](#), [get_anchors](#)

CMDist

Calculate Concept Mover's Distance

Description

Concept Mover's Distance classifies documents of any length along a continuous measure of engagement with a given concept of interest using word embeddings.

Usage

```
CMDist(
  dtm,
  cw = NULL,
  cv = NULL,
  wv,
  missing = "stop",
  scale = TRUE,
  sens_interval = FALSE,
  alpha = 1,
  n_iters = 20L,
  parallel = FALSE,
  threads = 2L,
  setup_timeout = 120L
)
```

```
cmdist(
  dtm,
  cw = NULL,
  cv = NULL,
  wv,
  missing = "stop",
  scale = TRUE,
  sens_interval = FALSE,
  alpha = 1,
  n_iters = 20L,
  parallel = FALSE,
  threads = 2L,
  setup_timeout = 120L
)
```

Arguments

dtm	Document-term matrix with words as columns. Works with DTMs produced by any popular text analysis package, or using the <code>dtm_builder()</code> function.
cw	Vector with concept word(s) (e.g., <code>c("love", "money")</code> , <code>c("critical thinking")</code>)
cv	Concept vector(s) as output from get_direction() , get_centroid() , or get_regions()
wv	Matrix of word embedding vectors (a.k.a embedding model) with rows as words.
missing	Indicates what action to take if words are not in embeddings. If <code>missing = "stop"</code> (default), the function is stopped and an error messages states which words are missing. If <code>missing = "remove"</code> , output is the same as terms but missing words or rows with missing words are removed. Missing words will be printed as a message.
scale	Logical (default = TRUE) uses <code>scale()</code> on output. This will set zero to the mean of the estimates, and scale by the standard deviation of the estimates. Document estimates will, therefore, be relative to other documents within that specific run, but not necessarily across discrete runs.

<code>sens_interval</code>	logical (default = FALSE), if TRUE several CMDs will be estimate on N resampled DTMs, sensitivity intervals are produced by returning the 2.5 and 97.5 percentiles of estimated CMDs for a given concept word or concept vector.
<code>alpha</code>	If <code>sens_interval</code> = TRUE, a number indicating the proportion of the document length to be resampled for sensitivity intervals. Default is 1 or 100 percent of each documents' length.
<code>n_iters</code>	If <code>sens_interval</code> = TRUE, integer (default = 20L) indicates the number of re-sampled DTMs to produced for sensitivity intervals
<code>parallel</code>	Logical (default = FALSE), whether to parallelize estimate
<code>threads</code>	If <code>parallel</code> = TRUE, an integer indicating attempts to connect to master before failing.
<code>setup_timeout</code>	If <code>parallel</code> = TRUE, maximum number of seconds a worker attempts to connect to master before failing.

Details

`CMDist()` requires three things: a (1) document-term matrix (DTM), a (2) matrix of word embedding vectors, and (3) concept words or concept vectors. The function uses *word counts* from the DTM and *word similarities* from the cosine similarity of their respective word vectors in a word embedding model. The "cost" of transporting all the words in a document to a single vector or a few vectors (denoting a concept of interest) is the measure of engagement, with higher costs indicating less engagement. For intuitiveness the output of `CMDist()` is inverted such that higher numbers will indicate *more engagement* with a concept of interest.

The vector, or vectors, of the concept are specified in several ways. The simplest involves selecting a single word from the word embeddings, the analyst can also specify the concept by indicating a few words. The algorithm then splits the overall flow between each concept word (roughly) depending on which word in the document is nearest. The words need not be in the DTM, but they must be in the word embeddings (the function will either stop or remove words not in the embeddings).

Instead of selecting a word already in the embedding space, the function can also take a vector extracted from the embedding space in the form of a centroid (which averages the vectors of several words) ,a direction (which uses the offset of several juxtaposing words), or a region (which is built by clustering words into k regions). The `get_centroid()`, `get_direction()`, and `get_regions()` functions will extract these.

Value

Returns a data frame (class `CMDist`) with the first column as document ids and each subsequent column as the CMD engagement corresponding to each concept word or concept vector. The upper and lower bound estimates will follow each unique CMD if `sens_interval` = TRUE. The `CMDist` class has `print.CMDist()` and `plot.CMDist()` methods.

Author(s)

Dustin Stoltz and Marshall Taylor

References

- Stoltz, Dustin S., and Marshall A. Taylor. (2019) 'Concept Mover's Distance' *Journal of Computational Social Science* 2(2):293-313. doi:[10.1007/s42001019000486](https://doi.org/10.1007/s42001019000486).
- Taylor, Marshall A., and Dustin S. Stoltz. (2020) 'Integrating semantic directions with concept mover's distance to measure binary concept engagement.' *Journal of Computational Social Science* 1-12. doi:[10.1007/s42001020000758](https://doi.org/10.1007/s42001020000758).
- Taylor, Marshall A., and Dustin S. Stoltz. (2020) 'Concept Class Analysis: A Method for Identifying Cultural Schemas in Texts.' *Sociological Science* 7:544-569. doi:[10.15195/v7.a23](https://doi.org/10.15195/v7.a23).

See Also

[CoCA](#), [get_direction](#), [get_centroid](#), [print.CMDist](#), [plot.CMDist](#)

Examples

```
# load example word embeddings
data(ft_wv_sample)

# load example text
data(jfk_speech)

# minimal preprocessing
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)

# create DTM
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)

# example 1
cm.dists <- CMDist(dtm,
  cw = "space",
  wv = ft_wv_sample
)

# example 2
space <- c("spacecraft", "rocket", "moon")
cen <- get_centroid(anchors = space, wv = ft_wv_sample)

cm.dists <- CMDist(dtm,
  cv = cen,
  wv = ft_wv_sample
)
```

Description

CoCA outputs schematic classes derived from documents' engagement with multiple bi-polar concepts (in a Likert-style fashion). The function requires a (1) DTM of a corpus which can be obtained using any popular text analysis package, or from the `dtm_builder()` function, and (2) semantic directions as output from the `get_direction()`. `CMDist()` works under the hood. Code modified from the `corclass` package.

Usage

```
CoCA(
  dtm,
  wv = NULL,
  directions = NULL,
  filter_sig = TRUE,
  filter_value = 0.05,
  zero_action = c("drop", "ownclass")
)

coca(
  dtm,
  wv = NULL,
  directions = NULL,
  filter_sig = TRUE,
  filter_value = 0.05,
  zero_action = c("drop", "ownclass")
)
```

Arguments

<code>dtm</code>	Document-term matrix with words as columns. Works with DTMs produced by any popular text analysis package, or you can use the <code>dtm_builder()</code> function.
<code>wv</code>	Matrix of word embedding vectors (a.k.a embedding model) with rows as words.
<code>directions</code>	direction vectors output from <code>get_direction()</code>
<code>filter_sig</code>	logical (default = TRUE), sets 'insignificant' ties to 0 to decrease noise and increase stability
<code>filter_value</code>	Minimum significance cutoff. Absolute row correlations below this value will be set to 0
<code>zero_action</code>	If 'drop', CCA drops rows with 0 variance from the analyses (default). If 'own-class', the correlations between 0-variance rows and all other rows is set 0, and the correlations between all pairs of 0-var rows are set to 1

Value

Returns a named list object of class CoCA. List elements include:

- `membership`: document memberships
- `modules`: schematic classes
- `cormat`: correlation matrix

Author(s)

Dustin Stoltz and Marshall Taylor

References

Taylor, Marshall A., and Dustin S. Stoltz. (2020) 'Concept Class Analysis: A Method for Identifying Cultural Schemas in Texts.' *Sociological Science* 7:544-569. doi:10.15195/v7.a23.
Boutyline, Andrei. 'Improving the measurement of shared cultural schemas with correlational class analysis: Theory and method.' *Sociological Science* 4.15 (2017): 353-393. doi:10.15195/v4.a15

See Also

[CMDist](#), [get_direction](#)

Examples

```
# load example word embeddings
data(ft_wv_sample)

# load example text
data(jfk_speech)

# minimal preprocessing
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)

# create DTM
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)

# create semantic directions
gen <- data.frame(
  pole1 = c("woman"),
  pole2 = c("man")
)

die <- data.frame(
  pole1 = c("alive"),
  pole2 = c("die")
)

hot <- data.frame(
  pole1 = c("hot"),
  pole2 = c("icebergs")
)

gen_dir <- get_direction(anchors = gen, wv = ft_wv_sample)
die_dir <- get_direction(anchors = die, wv = ft_wv_sample)
hot_dir <- get_direction(anchors = hot, wv = ft_wv_sample)

sem_dirs <- rbind(gen_dir, die_dir, hot_dir)
```

```

classes <- CoCA(
  dtm = dtm,
  wv = ft_wv_sample,
  directions = sem_dirs,
  filter_sig = TRUE,
  filter_value = 0.1,
  zero_action = "drop"
)

print(classes)

```

doc_centrality

Find a specified document centrality metric

Description

Given a document-term matrix or a document-similarity matrix, this function returns specified text network-based centrality measures. Currently, this includes degree, eigenvector, betweenness, and spanning.

Usage

```
doc_centrality(mat, method, alpha = 1L, two_mode = TRUE)
```

Arguments

mat	Document-term matrix with terms as columns or a document-similarity matrix with documents as rows and columns.
method	Character vector indicating centrality method, including "degree", "eigen", "span", and "between".
alpha	Number (default = 1) indicating the tuning parameter for weighted metrics.
two_mode	Logical (default = TRUE), indicating whether the input matrix is two mode (i.e. a document-term matrix) or one-mode (i.e. document-similarity matrix)

Details

If a document-term matrix is provided, the function obtains the one-mode document-level projection to get the document-similarity matrix using `tcrossprod()`. If a one-mode document-similarity matrix is provided, then this step is skipped. This way document similarities may be obtained using other methods, such as Word-Mover's Distance (see `doc_similarity`). The diagonal is ignored in all calculations.

Document centrality methods include:

- degree: Opsahl's weighted degree centrality with tuning parameter "alpha"
- between: vertex betweenness centrality using Brandes' method

- **eigen**: eigenvector centrality using Freeman's method. When the top two eigenvalues are tied or nearly tied (e.g. a corpus with duplicate or highly symmetric documents), the leading eigenvector is only defined up to an arbitrary rotation within that tied subspace, and results can be unstable under tiny perturbations to the similarity matrix; prefer "degree" or "between" for such corpora.
- **span**: Modified Burt's constraint following Stoltz and Taylor's method, uses a tuning parameter "alpha" and the output is scaled.

Value

A dataframe with two columns

Author(s)

Dustin Stoltz

References

Brandes, Ulrik (2000) 'A faster algorithm for betweenness centrality' *Journal of Mathematical Sociology*. 25(2):163-177 doi:[10.1080/0022250X.2001.9990249](https://doi.org/10.1080/0022250X.2001.9990249).
 Opsahl, Tore, et al. (2010) 'Node centrality in weighted networks: Generalizing degree and shortest paths.' *Social Networks*. 32(3)245:251 doi:[10.1016/j.socnet.2010.03.006](https://doi.org/10.1016/j.socnet.2010.03.006)
 Stoltz, Dustin; Taylor, Marshall (2019) 'Textual Spanning: Finding Discursive Holes in Text Networks' *Socius*. doi:[10.1177/2378023119827674](https://doi.org/10.1177/2378023119827674)

Examples

```
# load example text
data(jfk_speech)

# minimal preprocessing
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)

# create DTM
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)

ddeg <- doc_centrality(dtm, method = "degree")
if (requireNamespace("RSpectra", quietly = TRUE)) {
  deig <- doc_centrality(dtm, method = "eigen")
}
dbet <- doc_centrality(dtm, method = "between")
dspe <- doc_centrality(dtm, method = "span")

# with a document-similarity matrix (dsm)

dsm <- doc_similarity(dtm, method = "cosine")
ddeg <- doc_centrality(dsm, method = "degree", two_mode = FALSE)
```

doc_similarity	<i>Find a similarities between documents</i>
----------------	--

Description

Given a document-term matrix (DTM) this function returns the similarities between documents using a specified method (see details). The result is a square document-by-document similarity matrix (DSM), equivalent to a weighted adjacency matrix in network analysis.

Usage

```
doc_similarity(x, method, wv = NULL)
```

Arguments

x	Document-term matrix with terms as columns.
method	Character vector indicating similarity method, including projection, cosine, wmd, and centroid (see Details).
wv	Matrix of word embedding vectors (a.k.a embedding model) with rows as words. Required for "wmd" and "centroid" similarities.

Details

Document similarity methods include:

- projection: finds the one-mode projection matrix from the two-mode DTM using `tcrossprod()` which measures the shared vocabulary overlap
- cosine: compares row vectors using cosine similarity
- wmd: word mover's distance to compare documents (requires word embedding vectors), using linear-complexity relaxed word mover's distance
- centroid: represents each document as a centroid of their respective vocabulary, then uses cosine similarity to compare centroid vectors (requires word embedding vectors)

Value

A symmetric matrix of document-document similarities.

Author(s)

Dustin Stoltz

Examples

```
# load example word embeddings
data(ft_wv_sample)

# load example text
data(jfk_speech)

# minimal preprocessing
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)

# create DTM
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)

dsm_prj <- doc_similarity(dtm, method = "projection")
dsm_cos <- doc_similarity(dtm, method = "cosine")
dsm_wmd <- doc_similarity(dtm, method = "wmd", wv = ft_wv_sample)
dsm_cen <- doc_similarity(dtm, method = "centroid", wv = ft_wv_sample)
```

dtm_builder

A fast unigram DTM builder

Description

A streamlined function to take raw texts from a column of a data.frame and produce a sparse Document-Term Matrix (of generic class "dgCMatrix").

Usage

```
dtm_builder(
  data,
  text,
  doc_id = NULL,
  vocab = NULL,
  chunk = NULL,
  dense = FALSE,
  omit_empty = FALSE
)
```

Arguments

data	Data.frame with column of texts and column of document ids
text	Name of the column with documents' text
doc_id	Name of the column with documents' unique ids.
vocab	Default is NULL, if a list of terms is provided, the function will return a DTM with terms restricted to this vocabulary. Columns will also be in the same order as the list of terms.

chunk	Default is NULL, if an integer is provided, the function will "re-chunk" the corpus into new documents of a particular length. For example, 100L will divide the corpus into new documents with 100 terms (with the final document likely including slightly less than 100).
dense	The default (FALSE) is to return a matrix of class "dgCMatrix" as DTMs typically have mostly zero cells. This is much more memory efficient. Setting dense to TRUE will return a normal base R matrix.
omit_empty	Logical (default = FALSE) indicating whether to omit rows that have zero tokens.

Details

The function is fast because it has few bells and whistles:

- No weighting schemes other than raw counts
- Tokenizes by the fixed, single whitespace
- Only tokenizes unigrams. No bigrams, trigrams, etc...
- Columns are in the order unique terms are discovered
- No preprocessing during building
- Outputs a basic sparse Matrix or dense matrix

Weighting or stopping terms can be done efficiently after the fact with simple matrix operations, rather than achieved implicitly within the function itself. For example, using the `dtm_stopper()` function. Prior to creating the DTM, texts should have whitespace trimmed, if desired, punctuation removed and terms lowercased.

Like `tidytext`'s DTM functions, `dtm_builder()` is optimized for use in a pipeline, but unlike `tidytext`, it does not build an intermediary tripletlist, so `dtm_builder()` is faster and far more memory efficient.

The function can also chunk the corpus into documents of a given length (default is NULL). If the integer provided is 200L, this will divide the corpus into new documents with 200 terms (with the final document likely including slightly less than 200). If the total terms in the corpus were less than or equal to chunk integer, this would produce a DTM with one document (most will probably not want this).

If the vocabulary is already known, or standardizing vocabulary across several DTMs is desired, a list of terms can be provided to the `vocab` argument. Columns of the DTM will be in the order of the list of terms.

Value

returns a document-term matrix of class "dgCMatrix" or class "matrix"

Author(s)

Dustin Stoltz

Examples

```

library(dplyr)

my_corpus <- data.frame(
  text = c(
    "I hear babies crying I watch them grow",
    "They'll learn much more than I'll ever know",
    "And I think to myself",
    "What a wonderful world",
    "Yes I think to myself",
    "What a wonderful world"
  ),
  line_id = paste0("line", seq_len(6))
)
## some text preprocessing
my_corpus$clean_text <- tolower(gsub("'", "", my_corpus$text))

# example 1 with R 4.1 pipe

dtm <- my_corpus |>
  dtm_builder(clean_text, line_id)

# example 2 without pipe
dtm <- dtm_builder(
  data = my_corpus,
  text = clean_text,
  doc_id = line_id
)

# example 3 with dplyr pipe and mutate

dtm <- my_corpus %>%
  mutate(
    clean_text = gsub("'", "", text),
    clean_text = tolower(clean_text)
  ) %>%
  dtm_builder(clean_text, line_id)

# example 4 with dplyr and chunk of 3 terms
dtm <- my_corpus %>%
  dtm_builder(clean_text,
    line_id,
    chunk = 3L
  )

# example 5 with user defined vocabulary
my.vocab <- c("wonderful", "world", "haiku", "think")

dtm <- dtm_builder(
  data = my_corpus,

```

```
text = clean_text,  
doc_id = line_id,  
vocab = my.vocab  
)
```

`dtm_melter`*Melt a DTM into a triplet data frame*

Description

Converts a DTM into a data frame with three columns: documents, terms, frequency. Each row is a unique document by term frequency. This is akin to reshape2 packages melt function, but works on a sparse matrix. The resulting data frame is also equivalent to the tidytext triplet tibble.

Usage

```
dtm_melter(dtm)
```

Arguments

dtm	Document-term matrix with terms as columns. Works with DTMs produced by any popular text analysis package, or using the dtm_builder() function.
-----	---

Value

returns data frame with three columns: doc_id, term, freq

Author(s)

Dustin Stoltz

Examples

```
data(jfk_speech)  
jfk_speech$sentence <- tolower(jfk_speech$sentence)  
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)  
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)  
  
dtm_melted <- dtm_melter(dtm)  
head(dtm_melted)
```

dtm_resampler

*Resamples an input DTM to generate new DTMs***Description**

Takes any DTM and randomly resamples from each row, creating a new DTM

Usage

```
dtm_resampler(dtm, alpha = NULL, n = NULL)
```

Arguments

dtm	Document-term matrix with terms as columns. Works with DTMs produced by any popular text analysis package, or you can use the <code>dtm_builder()</code> function.
alpha	Number indicating proportion of document lengths, e.g., <code>alpha = 1</code> returns re-sampled rows that are the same lengths as the original DTM.
n	Integer indicating the length of documents to be returned, e.g., <code>n = 100L</code> will bring documents shorter than 100 tokens up to 100, while bringing documents longer than 100 tokens down to 100.

Details

Using the row counts as probabilities, each document's tokens are resampled with replacement up to a certain proportion of the row count (set by `alpha`). This function can be used with iteration to "bootstrap" a DTM without returning to the raw text. It does not iterate, however, so operations can be performed on one DTM at a time without storing multiple DTMs in memory.

If `alpha` is less than 1, then a proportion of each documents' lengths is returned. For example, `alpha = 0.50` will return a resampled DTM where each row has half the tokens of the original DTM. If `alpha = 2`, then each row in the resampled DTM twice the number of tokens of the original DTM. If an integer is provided to `n` then all documents will be resampled to that length. For example, `n = 2000L` will resample each document until they are 2000 tokens long – meaning those shorter than 2000 will be increased in length, while those longer than 2000 will be decreased in length. `alpha` and `n` should not be specified at the same time.

Value

returns a document-term matrix of class "dgCMatrix"

Examples

```
data(jfk_speech)
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)

dtm_resampled <- dtm_resampler(dtm, alpha = 0.5)
```

`dtm_stats`*Gets DTM summary statistics*

Description

`dtm_stats()` provides a summary of corpus-level statistics using any document-term matrix. These include (1) basic information on size (total documents, total unique terms, total tokens), (2) lexical richness, (3) distribution information, (4) central tendency, and (5) character-level information.

Usage

```
dtm_stats(  
  dtm,  
  richness = TRUE,  
  distribution = TRUE,  
  central = TRUE,  
  character = TRUE,  
  simplify = FALSE  
)
```

Arguments

<code>dtm</code>	Document-term matrix with terms as columns. Works with DTMs produced by any popular text analysis package, or you can use the <code>dtm_builder()</code> function.
<code>richness</code>	Logical (default = TRUE), whether to include statistics about lexical richness, i.e. terms that occur once, twice, and three times (hapax, dis, tris), and the total type-token ratio.
<code>distribution</code>	Logical (default = TRUE), whether to include statistics about the distribution, i.e. min, max st. dev, skewness, kurtosis.
<code>central</code>	Logical (default = TRUE), whether to include statistics about the central tendencies i.e. mean and median for types and tokens.
<code>character</code>	Logical (default = TRUE), whether to include statistics about the character lengths of terms, i.e. min, max, mean
<code>simplify</code>	Logical (default = FALSE), whether to return statistics as a data frame where each statistic is a column. Default returns a list of small data frames.

Value

A list of one to five data frames with summary statistics (if `simplify=FALSE`), otherwise a single data frame where each statistic is a column.

Author(s)

Dustin Stoltz

Examples

```
data(jfk_speech)
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)

dtm_stats(dtm)
dtm_stats(dtm, simplify = TRUE)
```

dtm_stopper	<i>Removes terms from a DTM based on rules</i>
-------------	--

Description

dtm_stopper will "stop" terms from the analysis by removing columns in a DTM based on stop rules. Rules include matching terms in a precompiled or custom list, terms meeting an upper or lower document frequency threshold, or terms meeting an upper or lower term frequency threshold.

Usage

```
dtm_stopper(
  dtm,
  stop_list = NULL,
  stop_termfreq = NULL,
  stop_termrank = NULL,
  stop_termprop = NULL,
  stop_docfreq = NULL,
  stop_docprop = NULL,
  stop_hapax = FALSE,
  stop_null = FALSE,
  omit_empty = FALSE,
  dense = FALSE,
  ignore_case = TRUE
)
```

Arguments

dtm	Document-term matrix with terms as columns. Works with DTMs produced by any popular text analysis package, or you can use the dtm_builder function.
stop_list	Vector of terms, from a precompiled stoplist or custom list such as c("never", "gonna", "give").
stop_termfreq	Vector of two numbers indicating the lower and upper threshold for exclusion (see details). Use Inf for max or min, respectively.
stop_termrank	Single integer indicating upper term rank threshold for exclusion (see details).

stop_termprop	Vector of two numbers indicating the lower and upper threshold for exclusion (see details). Use Inf for max or min, respectively.
stop_docfreq	Vector of two numbers indicating the lower and upper threshold for exclusion (see details). Use Inf for max or min, respectively.
stop_docprop	Vector of two numbers indicating the lower and upper threshold for exclusion (see details). Use Inf for max or min, respectively.
stop_hapax	Logical (default = FALSE) indicating whether to remove terms occurring one time (or zero times), a.k.a. hapax legomena
stop_null	Logical (default = FALSE) indicating whether to remove terms that occur zero times in the DTM.
omit_empty	Logical (default = FALSE) indicating whether to omit rows that are empty after stopping any terms.
dense	The default (FALSE) is to return a matrix of class "dgCMatrix". Setting dense to TRUE will return a normal base R dense matrix.
ignore_case	Logical (default = TRUE) indicating whether to ignore capitalization.

Details

Stopping terms by removing their respective columns in the DTM is significantly more efficient than searching raw text with string matching and deletion rules. Behind the scenes, the function relies on the `fastmatch` package to quickly match/not-match terms.

The `stop_list` argument takes a list of terms which are matched and removed from the DTM. If `ignore_case = TRUE` (the default) then word case will be ignored.

The `stop_termfreq` argument provides rules based on a term's occurrences in the DTM as a whole – regardless of its within document frequency. If real numbers between 0 and 1 are provided then terms will be removed by corpus proportion. For example `c(0.01, 0.99)`, terms that are either below 1% of the total tokens or above 99% of the total tokens will be removed. If integers are provided then terms will be removed by total count. For example `c(100, 9000)`, occurring less than 100 or more than 9000 times in the corpus will be removed. This also means that if `c(0, 1)` is provided, then the will only *keep* terms occurring once.

The `stop_termrank` argument provides the upper threshold for a terms' rank in the corpus. For example, 5L will remove the five most frequent terms.

The `stop_docfreq` argument provides rules based on a term's document frequency – i.e. the number of documents within which it occurs, regardless of how many times it occurs. If real numbers between 0 and 1 are provided then terms will be removed by corpus proportion. For example `c(0.01, 0.99)`, terms in more than 99% of all documents or terms that are in less than 1% of all documents. For example `c(100, 9000)`, then words occurring in less than 100 documents or more than 9000 documents will be removed. This means that if `c(0, 1)` is provided, then the function will only *keep* terms occurring in exactly one document, and remove terms in more than one.

The `stop_hapax` argument is a shortcut for removing terms occurring just one time in the corpus – called hapax legomena. Typically, a size-able portion of the corpus tends to be hapax terms, and removing them is a quick solution to reducing the dimensions of a DTM. The DTM must be frequency counts (not relative frequencies).

The `stop_null` argument removes terms that do not occur at all. In other words, there is a column for the term, but the entire column is zero. This can occur for a variety of reasons, such as starting

with a predefined vocabulary (e.g., using `dtm_builder`'s `vocab` argument) or through some cleaning processes.

The `omit_empty` argument will remove documents that are empty

Value

returns a document-term matrix of class "dgCMatrix"

Author(s)

Dustin Stoltz

Examples

```
# create corpus and DTM
my_corpus <- data.frame(
  text = c(
    "I hear babies crying I watch them grow",
    "They'll learn much more than I'll ever know",
    "And I think to myself",
    "What a wonderful world",
    "Yes I think to myself",
    "What a wonderful world"
  ),
  line_id = paste0("line", seq_len(6))
)
## some text preprocessing
my_corpus$clean_text <- tolower(gsub("'", "", my_corpus$text))

dtm <- dtm_builder(
  data = my_corpus,
  text = clean_text,
  doc_id = line_id
)

## example 1 with R 4.1 pipe
dtm_st <- dtm |>
  dtm_stopper(stop_list = c("world", "babies"))

## example 2 without pipe
dtm_st <- dtm_stopper(
  dtm,
  stop_list = c("world", "babies")
)

## example 3 precompiled stoplist
dtm_st <- dtm_stopper(
  dtm,
  stop_list = get_stoplist("snowball2014")
)
```

```
## example 4, stop top 2
dtm_st <- dtm_stopper(
  dtm,
  stop_termrank = 2L
)

## example 5, stop docfreq
dtm_st <- dtm_stopper(
  dtm,
  stop_docfreq = c(2, 5)
)
```

find_projection	<i>Find the 'projection matrix' to a semantic vector</i>
-----------------	--

Description

"Project" each word in a word embedding matrix of D dimension along a vector of D dimensions, extracted from the same embedding space. The vector can be a single word, or a concept vector obtained from [get_centroid\(\)](#), [get_direction\(\)](#), or [get_regions\(\)](#).

Usage

```
find_projection(wv, vec)
```

Arguments

wv	Matrix of word embedding vectors (a.k.a embedding model) with rows as words.
vec	Vector extracted from the embeddings

Details

All the vectors in the matrix A are projected onto the a vector, v , to find the projection matrix, P , defined as:

$$P = \frac{A \cdot v}{v \cdot v} * v$$

Value

A new word embedding matrix, each row of which is parallel to vector.

Examples

```
data(ft_wv_sample)

vec <- ft_wv_sample["space", ]
proj <- find_projection(wv = ft_wv_sample, vec = vec)
```

find_rejection	<i>Find the 'rejection matrix' from a semantic vector</i>
----------------	---

Description

"Reject" each word in a word embedding matrix of D dimension from a vector of D dimensions, extracted from the same embedding space. The vector can be a single word, or a concept vector obtained from `get_centroid()`, `get_direction()`, or `get_regions()`.

Usage

```
find_rejection(wv, vec)
```

Arguments

wv	Matrix of word embedding vectors (a.k.a embedding model) with rows as words.
vec	Vector extracted from the embeddings

Value

A new word embedding matrix, each row of which is rejected from vector.

Examples

```
data(ft_wv_sample)

vec <- ft_wv_sample["space", ]
rej <- find_rejection(wv = ft_wv_sample, vec = vec)
```

find_transformation	<i>Find a specified matrix transformation</i>
---------------------	---

Description

Given a matrix, B , of word embedding vectors (source) with terms as rows, this function finds a transformed matrix following a specified operation. These include: centering (i.e. translation) and normalization (i.e. scaling). In the first, B is centered by subtracting column means. In the second, B is normalized by the L2 norm. Both have been found to improve word embedding representations. The function also finds a transformed matrix that approximately aligns B , with another matrix, A , of word embedding vectors (reference), using Procrustes transformation (see details). Finally, given a term-co-occurrence matrix built on a local corpus, the function can "retrofit" pretrained embeddings to better match the local corpus.

Usage

```
find_transformation(
  wv,
  ref = NULL,
  method = c("align", "norm", "center", "retrofit")
)
```

Arguments

wv	Matrix of word embedding vectors (a.k.a embedding model) with rows as terms (the source matrix to be transformed).
ref	If method = "align", this is the reference matrix toward which the source matrix is to be aligned.
method	Character vector indicating the method to use for the transformation. Current methods include: "align", "norm", "center", and "retrofit" – see details.

Details

Aligning a source matrix of word embedding vectors, B , to a reference matrix, A , has primarily been used as a post-processing step for embeddings trained on longitudinal corpora for diachronic analysis or for cross-lingual embeddings. Aligning preserves internal (cosine) distances, while orient the source embeddings to minimize the sum of squared distances (and is therefore a Least Squares problem). Alignment is accomplished with the following steps:

- translation: centering by column means
- scaling: scale (normalizes) by the L2 Norm
- rotation/reflection: rotates and reflects to minimize sum of squared differences, using singular value decomposition

Alignment is asymmetrical, and only outputs the transformed source matrix, B . Therefore, it is typically recommended to align B to A , and then A to B . However, simply centering and norming A after may be sufficient.

Value

A new word embedding matrix, transformed using the specified method.

References

- Mikel Artetxe, Gorka Labaka, and Eneko Agirre. (2018). 'A robust self-learning method for fully unsupervised cross-lingual mappings of word embeddings.' *In Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics*. 789-798
- Mikel Artetxe, Gorka Labaka, and Eneko Agirre. 2019. 'An effective approach to unsupervised machine translation.' *In Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*. 194-203
- Hamilton, William L., Jure Leskovec, and Dan Jurafsky. (2018). 'Diachronic Word Embeddings Reveal Statistical Laws of Semantic Change.' <https://arxiv.org/abs/1605.09096v6>.
- Lin, Zefeng, Xiaojun Wan, and Zongming Guo. (2019). 'Learning Diachronic Word Embeddings

with Iterative Stable Information Alignment.’ *Natural Language Processing and Chinese Computing*. 749-60. doi:10.1007/9783030322335_58.

Schlechtweg et al. (2019). ‘A Wind of Change: Detecting and Evaluating Lexical Semantic Change across Times and Domains.’ <https://arxiv.org/abs/1906.02979v1>. Shoemark et al. (2019).

‘Room to Glo: A Systematic Comparison of Semantic Change Detection Approaches with Word Embeddings.’ *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing*. 66-76. doi:10.18653/v1/D191007 Borg and Groenen. (1997). *Modern Multidimensional Scaling*. New York: Springer. 340-342

Examples

```
data(ft_wv_sample)

# center and normalize
wv_centered <- find_transformation(ft_wv_sample, method = "center")
wv_normed <- find_transformation(ft_wv_sample, method = "norm")

# align requires a reference matrix (e.g. from a different time period)
# wv_aligned <- find_transformation(wv_source, ref = wv_ref, method = "align")
```

ft_wv_sample	<i>Sample of fastText embeddings</i>
--------------	--------------------------------------

Description

These are a sample of the English fastText embeddings including 770 words matching those used in the jfk_speech. These are intended to be used for example code.

Format

A matrix of 770 rows and 300 columns

Source

Subsample of English fastText word embeddings. Bojanowski et al. (2017) "Enriching Word Vectors with Subword Information." doi:10.1162/tacl_a_00051. Available at <https://fasttext.cc>.

get_anchors	<i>Gets anchor terms from precompiled anchor lists</i>
-------------	--

Description

Produces a data.frame of juxtaposed word pairs used to extract a semantic direction from word embeddings. Can be used as input to [get_direction\(\)](#).

Usage

```
get_anchors(relation)
```

Arguments

relation	String indicating a semantic relation, 30 relations are available in the dataset (see details).
----------	---

Details

Sets of juxtaposed "anchor" pairs are adapted from published work and associated with a particular semantic relation. These should be used as a starting point, not as a "ground truth."

Available relations include:

- activity
- affluence
- age
- attractiveness
- authority
- borders
- concreteness
- cultivation
- dominance
- education
- fairness
- gender
- government
- health
- immigration
- ingroup
- labor
- legality
- morality

- organization
- partisan
- policy
- purity
- safety
- sexuality
- skills
- status
- valence
- whiteness

Value

returns a tibble with two columns

Author(s)

Dustin Stoltz

Examples

```
gen <- get_anchors(relation = "gender")
```

get_centroid	<i>Word embedding semantic centroid extractor</i>
--------------	---

Description

The function outputs an averaged vector from a set of anchor terms' word vectors. This average is roughly equivalent to the intersection of the contexts in which each word is used. This semantic centroid can be used for a variety of ends, and specifically as input to `CMDist()`. `get_centroid()` requires a list of terms, string of terms, data.frame or matrix. In the latter two cases, the first column will be used. The vectors are aggregated using the simple average. Terms can be repeated, and are therefore "weighted" by their counts.

Usage

```
get_centroid(anchors, wv, missing = "stop")
```

Arguments

anchors	List of terms to be averaged
wv	Matrix of word embedding vectors (a.k.a embedding model) with rows as words.
missing	what action to take if terms are not in embeddings. If action = "stop" (default), the function is stopped and an error messages states which terms are missing. If action = "remove", missing terms or rows with missing terms are removed. Missing terms will be printed as a message.

Value

returns a one row matrix

Author(s)

Dustin Stoltz

See Also

[get_centroids\(\)](#)

Examples

```
# load example word embeddings
data(ft_wv_sample)

space1 <- c("spacecraft", "rocket", "moon")

cen1 <- get_centroid(anchors = space1, wv = ft_wv_sample)

space2 <- c("spacecraft rocket moon")
cen2 <- get_centroid(anchors = space2, wv = ft_wv_sample)

identical(cen1, cen2)
```

get_centroids	<i>Compute multiple centroids</i>
---------------	-----------------------------------

Description

Compute multiple centroids in a single call, either from a list of anchor term sets (word-vector mode) or from a feature matrix with a grouping vector (feature-matrix mode). The plural counterpart to [get_centroid\(\)](#).

Usage

```
get_centroids(fv, anchors = NULL, groups = NULL, missing = "stop")
```

Arguments

fv	Matrix or sparse <code>Matrix::dgCMatrix</code> of feature vectors. In word-vector mode this is the word embedding matrix (rows are terms); in feature-matrix mode this is a document-by-feature matrix (rows are documents).
anchors	List of anchor term vectors. Each element is passed to get_centroid() along with fv and missing. If the list is unnamed, integer indices are used as the rownames of the result.

groups	Vector of group labels of the same length as <code>nrow(fv)</code> . A centroid is computed for each level via <code>colMeans()</code> . Character vectors are preferred over factors; empty factor levels are dropped with a message. Rows with a missing (NA) label are excluded from every centroid, with a message.
missing	Character scalar. Either "stop" or "remove". Inherited from <code>get_centroid()</code> and applies only to the anchors (word-vector) branch. See <code>get_centroid()</code> .

Value

Matrix with one row per centroid and `ncol(fv)` columns. Rownames are the names of anchors (or the levels of groups).

Author(s)

Dustin Stoltz

See Also

[get_centroid\(\)](#)

Examples

```
data(ft_wv_sample)

## word-vector mode: named list of anchor sets
anchors <- list(
  space = c("spacecraft", "rocket", "moon"),
  water = c("ocean", "sea")
)
cens <- get_centroids(ft_wv_sample, anchors = anchors)

## feature-matrix mode: group rows of a matrix by a label vector
mat <- ft_wv_sample[1:10, , drop = FALSE]
labs <- rep(c("a", "b"), each = 5)
cens <- get_centroids(mat, groups = labs)
```

get_direction

Word embedding semantic direction extractor

Description

`get_direction()` outputs a vector corresponding to one pole of a "semantic direction" built from sets of antonyms or juxtaposed terms. The output can be used as an input to `CMDist()` and `CoCA()`. Anchors must be a two-column data.frame or a list of length == 2.

Usage

```
get_direction(anchors, wv, method = "paired", missing = "stop", n_dirs = 1L)
```

Arguments

anchors	A data frame or list of juxtaposed 'anchor' terms
wv	Matrix of word embedding vectors (a.k.a embedding model) with rows as terms.
method	Indicates the method used to generate vector offset. Default is 'paired'. See details.
missing	what action to take if terms are not in embeddings. If action = "stop" (default), the function is stopped and an error messages states which terms are missing. If action = "remove", missing terms or rows with missing terms are removed. Missing terms will be printed as a message.
n_dirs	If method = "PCA", an integer indicating how many directions to return. Default = 1L, indicating a single, bipolar, direction. Cannot exceed the number of anchor pairs provided in anchors; errors with a message naming the maximum available if it does.

Details

Semantic directions can be estimated in using a few methods:

- 'paired' (default): each individual term is subtracted from exactly one other paired term. there must be the same number of terms for each side of the direction (although one word may be used more than once).
- 'pooled': terms corresponding to one side of a direction are first averaged, and then these averaged vectors are subtracted. A different number of terms can be used for each side of the direction.
- 'L2': the vector is calculated the same as with 'pooled' but is then divided by the L2 'Euclidean' norm
- 'PCA': vector offsets are calculated for each pair of terms, as with 'paired', and if n_dirs = 1L (the default) then the direction is the first principal component. Users can return more than one direction by increasing the n_dirs parameter.

Value

Returns a matrix with one row per direction (n_dirs rows when method = "PCA"; one row otherwise).

Author(s)

Dustin Stoltz

References

Bolukbasi, Tolga, Kai-Wei Chang, James Zou, Venkatesh Saligrama, Adam Kalai (2016). 'Man Is to Computer Programmer as Woman Is to Homemaker? Debiasing Word Embeddings.' *Advances in Neural Information Processing Systems*, 29. doi:10.48550/arXiv.1606.06121.

Taylor, Marshall A., and Dustin S. Stoltz. (2020) 'Concept Class Analysis: A Method for Identifying Cultural Schemas in Texts.' *Sociological Science* 7:544-569. doi:10.15195/v7.a23.

Taylor, Marshall A., and Dustin S. Stoltz. (2020) 'Integrating semantic directions with concept

mover's distance to measure binary concept engagement.' *Journal of Computational Social Science* 1-12. doi:10.1007/s42001020000758.

Kozlowski, Austin C., Matt Taddy, and James A. Evans. (2019). 'The geometry of culture: Analyzing the meanings of class through word embeddings.' *American Sociological Review* 84(5):905-949. doi:10.1177/0003122419877135.

Arseniev-Koehler, Alina, and Jacob G. Foster. (2020). 'Machine learning as a model for cultural learning: Teaching an algorithm what it means to be fat.' arXiv preprint <https://arxiv.org/abs/2003.12133v2>.

Examples

```
# load example word embeddings
data(ft_wv_sample)

# create anchor list
gen <- data.frame(
  pole1 = c("woman"),
  pole2 = c("man")
)

dir <- get_direction(anchors = gen, wv = ft_wv_sample)

# method = "PCA" requires at least two anchor pairs
gen2 <- data.frame(
  pole1 = c("woman", "stay"),
  pole2 = c("man", "move")
)

dir <- get_direction(
  anchors = gen2, wv = ft_wv_sample,
  method = "PCA", n_dirs = 1L
)

# method = "PCA" can return more than one direction, up to the
# number of anchor pairs provided
dirs <- get_direction(
  anchors = gen2, wv = ft_wv_sample,
  method = "PCA", n_dirs = 2L
)
```

get_regions

Word embedding semantic region extractor

Description

Given a set of word embeddings of d dimensions and v vocabulary, `get_regions()` finds k semantic regions in d dimensions. This, in effect, learns latent topics from an embedding space (a.k.a. topic modeling), which are directly comparable to both terms (with cosine similarity) and documents (with Concept Mover's distance using `CMDist()`).

Usage

```
get_regions(wv, k_regions = 5L, max_iter = 20L, seed = 0)
```

Arguments

wv	Matrix of word embedding vectors (a.k.a embedding model) with rows as words.
k_regions	Integer indicating the k number of regions to return
max_iter	Integer indicating the maximum number of iterations before k-means terminates.
seed	Integer indicating a random seed. Default is 0, which calls 'std::time(NULL)'.

Details

To group words into more encompassing "semantic regions" we use k -means clustering. We choose k -means primarily for its ubiquity and the wide range of available diagnostic tools for k -means cluster.

A word embedding matrix of d dimensions and v vocabulary is "clustered" into k semantic regions which have d dimensions. Each region is represented by a single point defined by the d dimensional vector. The process discretely assigns all word vectors are assigned to a given region so as to minimize some error function, however as the resulting regions are in the same dimensions as the word embeddings, we can measure each terms similarity to each region. This, in effect, is a mixed membership topic model similar to topic modeling by Latent Dirichlet Allocation.

We use the KMeans_arma function from the ClusterR package which uses the Armadillo library.

Value

returns a matrix of class "dgCMatrix" with k rows and d dimensions

Author(s)

Dustin Stoltz

References

Butnaru, Andrei M., and Radu Tudor Ionescu. (2017) 'From image to text classification: A novel approach based on clustering word embeddings.' *Procedia computer science*. 112:1783-1792. [doi:10.1016/j.procs.2017.08.211](https://doi.org/10.1016/j.procs.2017.08.211).

Zhang, Yi, Jie Lu, Feng Liu, Qian Liu, Alan Porter, Hongshu Chen, and Guangquan Zhang. (2018). 'Does Deep Learning Help Topic Extraction? A Kernel K-Means Clustering Method with Word Embedding.' *Journal of Informetrics*. 12(4):1099-1117. [doi:10.1016/j.joi.2018.09.004](https://doi.org/10.1016/j.joi.2018.09.004).

Arseniev-Koehler, Alina and Cochran, Susan D and Mays, Vickie M and Chang, Kai-Wei and Foster, Jacob Gates (2021) 'Integrating topic modeling and word embedding to characterize violent deaths' [doi:10.31235/osf.io/nkyaq](https://doi.org/10.31235/osf.io/nkyaq)

Examples

```
# load example word embeddings
data(ft_wv_sample)

my_regions <- get_regions(
  wv = ft_wv_sample,
  k_regions = 10L,
  max_iter = 10L,
  seed = 01984
)
```

get_stoplist	<i>Gets stoplist from precompiled lists</i>
--------------	---

Description

Provides access to 8 precompiled stoplists, including the most commonly used stoplist from the Snowball stemming package ("snowball2014"), text2map's tiny stoplist ("tiny2020"), a few historically important stop lists. This aims to be a transparent and well-document collection of stoplists. Only includes English language stoplists at the moment.

Usage

```
get_stoplist(source = "tiny2020", language = "en", tidy = FALSE)
```

Arguments

source	Character indicating source, default = "tiny2020"
language	Character (default = "en") indicating language of stopwords by ISO 639-1 code, currently only English is supported.
tidy	logical (default = FALSE), returns a tibble

Details

There is no such thing as a *stopword*! But, there are **tons** of precompiled lists of words that someone thinks we should remove from our texts. (See for example: <https://github.com/igorbrigadir/stopwords>) One of the first stoplists is from C.J. van Rijsbergen's "Information retrieval: theory and practice" (1979) and includes 250 words. text2map's very own stoplist tiny2020 is a lean 34 words.

Below are stoplists available with [get_stoplist](#):

- "tiny2020": Tiny (2020) list of 33 words (Default)
- "snowball2001": Snowball stemming package's (2001) list of 127 words
- "snowball2014": Updated Snowball (2014) list of 175 words
- "van1979": C. J. van Rijsbergen's (1979) list of 250 words

- "fox1990": Christopher Fox's (1990) list of 421 words
- "smart1993": Original SMART (1993) list of 570 words
- "onix2000": ONIX (2000) list of 196 words
- "nltk2009": Python's NLTK (2009) list of 179 words

The Snowball (2014) stoplist is likely the most commonly, it is the default in the `stopwords` package, which is used by `quanteda`, `tidytext` and `tokenizers` packages, followed closely by the Smart (1993) stoplist, the default in the `tm` package. The word counts for SMART (1993) and ONIX (2000) are slightly different than in other places because of duplicate words.

Value

Character vector of words to be stopped, if `tidy = TRUE`, a tibble is returned

Author(s)

Dustin Stoltz

Examples

```
stops <- get_stoplist("snowball2014")
head(stops)

stops_tb <- get_stoplist("snowball2014", tidy = TRUE)
head(stops_tb)
```

jfk_speech

Full Text of JFK's Rice Speech

Description

This is a data frame for the text of JFK's Rice Speech "We choose to go to the moon." Each row is a 10 word string of the speech – roughly a sentence. This is intended to be used for example code.

Format

A data frame with 84 rows and 2 columns

Variables

Variables:

- `sentence_id`. Order and unique ID for the sentence
- `sentence`. The text of a sentence

Source

John F. Kennedy, "We choose to go to the Moon" speech, Rice University, September 12, 1962. Public domain.

meta_shakespeare	<i>Metadata for Shakespeare's First Folio</i>
------------------	---

Description

Metadata related to Shakespeare's First Folio including the IDs to download the plays from Project Gutenberg, and a count of the number of deaths in each play (body count).

Format

A data frame with 37 rows and 8 columns

Variables

Variables:

- short_title.
- gutenberg_title.
- gutenberg_id.
- genre.
- year.
- body_count.
- boas_problem_plays.
- death.

Source

Shakespeare's First Folio plays, obtained via Project Gutenberg (<https://www.gutenberg.org>). Body counts compiled by Stoltz and Taylor (2021) "Integrating Semantic Contrast with Concept Mover's Distance" [doi:10.1016/j.poetic.2021.101567](https://doi.org/10.1016/j.poetic.2021.101567).

perm_tester

*Monte Carlo Permutation Tests for Model P-Values***Description**

`perm_tester()` carries out Monte Carlo permutation tests for model p-values from two-tailed, left-tailed, and/or right-tailed hypothesis testing.

Usage

```
perm_tester(
  data,
  model,
  perm_var = NULL,
  strat_var = NULL,
  statistic,
  perm_n = 1000,
  alternative = "all",
  alpha = 0.05,
  seed = NULL
)
```

Arguments

<code>data</code>	The dataframe from which the model is estimated.
<code>model</code>	The model which will be estimated and re-estimated.
<code>perm_var</code>	The variable in the model that will be permuted. Default is <code>NULL</code> which takes the first Y_n term in the formula of the model
<code>strat_var</code>	Categorical variable for within-stratum permutations. Defaults to <code>NULL</code> .
<code>statistic</code>	The name of the model statistic you want to "grab" after re-running the model with each permutation to compare to the original model statistic. Must be a single string matching one of <code>names(model)</code> (e.g. "coefficients" for an <code>lm</code> object); errors listing the model's actual component names if it doesn't.
<code>perm_n</code>	The total number of permutations. Defaults to 1000.
<code>alternative</code>	The alternative hypothesis. One of "two.sided", "left", "right", and "all" (default). Defaults to "all", which reports the p-value statistics for all three alternative hypotheses.
<code>alpha</code>	Alpha level for the hypothesis test. Defaults to 0.05.
<code>seed</code>	Optional seed for reproducibility of the p-value statistics. Defaults to null.

Details

`perm_tester()` can be used to derive p-values under the randomization model of inference. There are various reasons one might want to do this— with text data, and observational data more generally, this might be because the corpus/sample is not a random sample from a target population.

In such cases, population model p-values might not make much sense since the asymptotically-derived standard errors from which they are constructed themselves do not make sense. We might therefore want to make inferences on the basis of whether or not randomness, as a data-generating mechanism, might reasonably account for a statistic at least as extreme as the one we observed. `perm_tester()` works from this idea.

`perm_tester()` works like this. First, the model (supplied the `model` parameter) is run on the observed data. Second, we take some statistic of interest, which we indicate with the `statistic` parameter, and set it to the side. Third, a variable, `perm_var`, is permuted—meaning the observed values for the rows of data on `perm_var` are randomly reshuffled. Fourth, we estimate the model again, this time with the permuted `perm_var`. Fifth, we get grab that same statistic. We repeat steps two through five a total of `perm_n` times, each time tallying the number of times the statistic from the permutation-derived model is greater than or equal to (for a right-tailed test), less-than or equal to (for a left-tailed test), and/or has an absolute value greater than or equal to (for a two-tailed test) the statistic from the "real" model.

If we divide those tallies by the total number of permutations, then we get randomization-based p-values. This is what `perm_tester()` does. The null hypothesis is that randomness could likely generate the statistic that we observe. The alternative hypothesis is that randomness alone likely can't account for the observed statistic.

We then reject the null hypothesis if the p-value is below a threshold indicated with `alpha`, which, as in population-based inference, is the probability below which we are willing to reject the null hypothesis when it is actually true. So if the p-value is below, say, `alpha = 0.05` and we're performing, a right-tailed test, then fewer than 5% of the statistics derived from the permutation-based models are greater than or equal to our observed statistic. We would then reject the null, as it is unlikely (based on our `alpha` threshold), that randomness as a data-generating mechanism can account for a test statistic at least as large the one we observed.

In most cases, analysts probably cannot expect to perform "exact" permutation tests where every possible permutation is accounted for—i.e., where `perm_n` equals the total number of possible permutations. Instead, we can take random samples of the "population" of permutations. `perm_tester()` does this, and reports the standard errors and $(1 - \alpha)$ confidence intervals for the p-values.

`perm_tester()` can also perform stratified permutation tests, where the observed `perm_var` variables within groups. This can be done by setting the `strat_var` variable to be the grouping variable.

Value

Returns a data frame with the observed statistic (`stat`), the p-values (`P_left`, for left-tailed, `P_right` for right-tailed, and/or `P_two` for two-tailed), and the standard errors and confidence intervals for those p-values, respectively.

Author(s)

Marshall Taylor and Dustin Stoltz

References

Taylor, Marshall A. (2020) 'Visualization Strategies for Regression Estimates with Randomization Inference' *Stata Journal* 20(2):309-335. doi:[10.1177/1536867X20930999](https://doi.org/10.1177/1536867X20930999).

Darlington, Richard B. and Andrew F. Hayes (2016) *Regression analysis and linear models: Concepts, applications, and implementation*. Guilford Publications.

Ernst, Michael D. (2004) 'permutation methods: a basis for exact inference' *Statistical Science* 19(4):676-685. doi:[10.1214/088342304000000396](https://doi.org/10.1214/088342304000000396).

Manly, Bryan F. J. (2007) *Randomization, Bootstrap and Monte Carlo Methods in Biology*. Chapman and Hall/CRC. doi:[10.1201/9781315273075](https://doi.org/10.1201/9781315273075).

See Also

[CMDist](#), [CoCA](#), [get_direction](#)

Examples

```
data(meta_shakespeare)

model <- lm(body_count ~ boas_problem_plays + year + genre, data = meta_shakespeare)

# without stratified permutations, two-sided test
out1 <- perm_tester(
  data = meta_shakespeare,
  model = model,
  statistic = "coefficients",
  perm_n = 40,
  alternative = "two.sided",
  alpha = .01,
  seed = 8675309
)

# with stratified permutations, two-sided test
out2 <- perm_tester(
  data = meta_shakespeare,
  model = model,
  strat_var = "boas_problem_plays",
  statistic = "coefficients",
  perm_n = 40,
  alternative = "two.sided",
  alpha = .01,
  seed = 8675309
)
```

Description

Plots the CMD score for each document, sorted from lowest to highest engagement, with one panel per concept. If `x` was produced with `sens_interval = TRUE`, error bars show each document's sensitivity interval.

Usage

```
## S3 method for class 'CMDist'
plot(x, concept = NULL, ...)
```

Arguments

<code>x</code>	CMDist object returned by <code>CMDist()</code>
<code>concept</code>	Character vector of one or more concept names to plot (default = <code>NULL</code> , which plots every concept found in <code>x</code>). Errors if any name is not one of the concepts in <code>x</code> .
<code>...</code>	Arguments to be passed to methods

Value

Invisibly returns `x`

See Also

[CMDist](#), [print.CMDist](#)

Examples

```
data(ft_wv_sample)
data(jfk_speech)
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)

out <- CMDist(dtm, cw = "space", wv = ft_wv_sample)
plot(out)

# with sensitivity intervals, plotting a single concept
out_sens <- CMDist(dtm, cw = "space", wv = ft_wv_sample,
                  sens_interval = TRUE, n_iters = 5)
plot(out_sens, concept = "space")
```

plot.CoCA

*Plot CoCA***Description**

Plot CoCA

Usage

```
## S3 method for class 'CoCA'
plot(
  x,
  module = NULL,
  min = 0.15,
  max = 1,
  label.cex = 1.2,
  seed = 123,
  main = NULL,
  ...
)
```

Arguments

<code>x</code>	CoCA object returned by <code>CoCA()</code>
<code>module</code>	Integer, which schematic class to plot. Defaults to NULL, but a value is effectively required: omitting it or passing NULL errors, prompting you to specify one. Must be a single whole number between 1 and the number of modules found in <code>x</code> (<code>length(x\$modules)</code>), else errors with a message stating the valid range.
<code>min</code>	edges with absolute weights under this value are not shown (default = 0.15)
<code>max</code>	highest weight to scale the edge widths too (default = 1)
<code>label.cex</code>	label size scaling factor (default = 1.2)
<code>seed</code>	random seed for layout reproducibility (default = 123)
<code>main</code>	title for plot (default = NULL)
<code>...</code>	Arguments to be passed to methods

Value

invisibly returns the igraph object

Examples

```
data(ft_wv_sample)
data(jfk_speech)
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)
```

```

gen <- data.frame(pole1 = c("man", "new", "choose"), pole2 = c("nation", "world", "peace"))
space <- data.frame(pole1 = c("space", "moon", "science"), pole2 = c("nation", "world", "freedom"))
power <- data.frame(pole1 = c("great", "power", "new"), pole2 = c("peace", "freedom", "world"))

gen_dir <- get_direction(anchors = gen, wv = ft_wv_sample)
space_dir <- get_direction(anchors = space, wv = ft_wv_sample)
power_dir <- get_direction(anchors = power, wv = ft_wv_sample)
directions <- rbind(gen_dir, space_dir, power_dir)

classes <- CoCA(dtm, wv = ft_wv_sample, directions = directions)

```

print.CMDist	<i>Print CMDist estimates</i>
--------------	-------------------------------

Description

Print CMDist estimates

Usage

```

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

```

Arguments

x	CMDist object returned by CMDist()
...	Arguments to be passed to methods

Value

Invisibly returns x

See Also

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

Examples

```

data(ft_wv_sample)
data(jfk_speech)
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)

out <- CMDist(dtm, cw = "space", wv = ft_wv_sample)
print(out)

```

print.CoCA	<i>Prints CoCA class information</i>
------------	--------------------------------------

Description

Prints CoCA class information

Usage

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

Arguments

x	CoCA object returned by CoCA()
...	Arguments to be passed to methods

Value

Invisibly returns NULL

Examples

```
data(ft_wv_sample)  
data(jfk_speech)  
jfk_speech$sentence <- tolower(jfk_speech$sentence)  
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)  
dtm <- dtm_builder(jfk_speech, sentence, sentence_id)  
  
gen <- data.frame(pole1 = c("man", "new", "choose"), pole2 = c("nation", "world", "peace"))  
space <- data.frame(pole1 = c("space", "moon", "science"), pole2 = c("nation", "world", "freedom"))  
power <- data.frame(pole1 = c("great", "power", "new"), pole2 = c("peace", "freedom", "world"))  
  
gen_dir <- get_direction(anchors = gen, wv = ft_wv_sample)  
space_dir <- get_direction(anchors = space, wv = ft_wv_sample)  
power_dir <- get_direction(anchors = power, wv = ft_wv_sample)  
directions <- rbind(gen_dir, space_dir, power_dir)  
  
classes <- CoCA(dtm, wv = ft_wv_sample, directions = directions)  
print(classes)
```

rancors_builder

Build Multiple Random Corpora

Description

`rancors_builder()` generates multiple random corpus (rancor) based on a user defined term probabilities and vocabulary. Users can set the number of documents, as well as the mean, standard deviation, minimum, and maximum document lengths (i.e., number of tokens) of the parent normal distribution from which the document lengths are randomly sampled. The output is a list of document-term matrices. To produce a *single* random corpus, use `rancor_builder()` (note the singular).

Usage

```
rancors_builder(
  data,
  vocab,
  probs,
  n_cors,
  n_docs,
  len_mean,
  len_var,
  len_min,
  len_max,
  seed = NULL
)
```

Arguments

<code>data</code>	Data.frame containing vocabulary and probabilities
<code>vocab</code>	Name of the column containing vocabulary
<code>probs</code>	Name of the column containing probabilities
<code>n_cors</code>	Integer indicating the number of corpora to build
<code>n_docs</code>	Integer(s) indicating the number of documents to be returned. If two numbers are provided, number will be randomly sampled within the range for each corpora.
<code>len_mean</code>	Integer(s) indicating the mean of the document lengths in the parent normal sampling distribution. If two numbers are provided, number will be randomly sampled within the range for each corpora.
<code>len_var</code>	Integer(s) indicating the standard deviation of the document lengths in the parent normal sampling distribution. If two numbers are provided, number will be randomly sampled within the range for each corpora.
<code>len_min</code>	Integer(s) indicating the minimum of the document lengths in the parent normal sampling distribution. If two numbers are provided, number will be randomly sampled within the range for each corpora.

len_max	Integer(s) indicating the maximum of the document lengths in the parent normal sampling distribution. If two numbers are provided, number will be randomly sampled within the range for each corpora.
seed	Optional seed for reproducibility. If NULL (default), a seed is derived from the system clock/process id rather than from the global RNG stream, so the caller's .Random.seed is left untouched and repeated unseeded calls still differ from one another. Whichever seed is actually used (supplied or auto-generated) is attached to the output via attr(x, "seed"), so an unseeded run can still be reproduced.

Value

A list of sparse document-term matrices (dgCMatrix), with the seed actually used (supplied or auto-generated) attached as attr(x, "seed").

Author(s)

Dustin Stoltz and Marshall Taylor

Examples

```
# create corpus and DTM
my_corpus <- data.frame(
  text = c(
    "I hear babies crying I watch them grow",
    "They'll learn much more than I'll ever know",
    "And I think to myself",
    "What a wonderful world",
    "Yes I think to myself",
    "What a wonderful world"
  ),
  line_id = paste0("line", seq_len(6))
)
## some text preprocessing
my_corpus$clean_text <- tolower(gsub("'", "", my_corpus$text))

dtm <- dtm_builder(
  data = my_corpus,
  text = clean_text,
  doc_id = line_id
)

# use colSums to get term frequencies
df <- data.frame(
  vocab = colnames(dtm),
  freqs = colSums(dtm)
)
# convert to probabilities
df$probs <- df$freqs / sum(df$freqs)
```

```
# create random DTM
ls_dtms <- df |>
  rancors_builder(vocab,
    probs,
    n_cors = 20,
    n_docs = 100,
    len_mean = c(50, 200),
    len_var = 5,
    len_min = 20,
    len_max = 1000,
    seed = 59801
  )
length(ls_dtms)

# recover the seed used, even when not supplied
attr(ls_dtms, "seed")
```

rancor_builder*Build a Random Corpus*

Description

`rancor_builder()` generates a random corpus (rancor) based on a user defined term probabilities and vocabulary. Users can set the number of documents, as well as the mean, standard deviation, minimum, and maximum document lengths (i.e., number of tokens) of the parent normal distribution from which the document lengths are randomly sampled. The output is a single document-term matrix. To produce multiple random corpora, use `rancors_builder()` (note the plural). Term probabilities/vocabulary can come from a users own corpus, or a pre-compiled frequency list, such as the one derived from the Google Book N-grams corpus

Usage

```
rancor_builder(
  data,
  vocab,
  probs,
  n_docs = 100L,
  len_mean = 500,
  len_var = 10L,
  len_min = 20L,
  len_max = 1000L,
  seed = NULL
)
```

Arguments

<code>data</code>	Data.frame containing vocabulary and probabilities
<code>vocab</code>	Name of the column containing vocabulary
<code>probs</code>	Name of the column containing probabilities
<code>n_docs</code>	Integer indicating the number of documents to be returned
<code>len_mean</code>	Integer indicating the mean of the document lengths in the parent normal sampling distribution
<code>len_var</code>	Integer indicating the standard deviation of the document lengths in the parent normal sampling distribution
<code>len_min</code>	Integer indicating the minimum of the document lengths in the parent normal sampling distribution
<code>len_max</code>	Integer indicating the maximum of the document lengths in the parent normal sampling distribution
<code>seed</code>	Optional seed for reproducibility. If NULL (default), a seed is derived from the system clock/process id rather than from the global RNG stream, so the caller's <code>.Random.seed</code> is left untouched and repeated unseeded calls still differ from one another. Whichever seed is actually used (supplied or auto-generated) is attached to the output via <code>attr(x, "seed")</code> , so an unseeded run can still be reproduced.

Value

A sparse document-term matrix (`dgCMatrix`), with the seed actually used (supplied or auto-generated) attached as `attr(x, "seed")`.

Author(s)

Dustin Stoltz and Marshall Taylor

Examples

```
# create corpus and DTM
my_corpus <- data.frame(
  text = c(
    "I hear babies crying I watch them grow",
    "They'll learn much more than I'll ever know",
    "And I think to myself",
    "What a wonderful world",
    "Yes I think to myself",
    "What a wonderful world"
  ),
  line_id = paste0("line", seq_len(6))
)
## some text preprocessing
my_corpus$clean_text <- tolower(gsub("'", "", my_corpus$text))

dtm <- dtm_builder(
  data = my_corpus,
```

```

    text = clean_text,
    doc_id = line_id
  )

  # use colSums to get term frequencies
  df <- data.frame(
    terms = colnames(dtm),
    freqs = colSums(dtm)
  )
  # convert to probabilities
  df$probs <- df$freqs / sum(df$freqs)

  # create random DTM
  rDTM <- df |>
    rancor_builder(terms, probs)

  # recover the seed used, even when not supplied
  attr(rDTM, "seed")

```

seq_builder

*Represent Documents as Token-Integer Sequences***Description**

First, each token in the vocabulary is mapped to an integer in a lookup dictionary. Next, documents are converted to sequences of integers where each integer is an index of the token from the dictionary.

Usage

```

seq_builder(
  data,
  text,
  doc_id = NULL,
  vocab = NULL,
  maxlen = NULL,
  matrix = TRUE
)

```

Arguments

data	Data.frame with column of texts and column of document ids
text	Name of the column with documents' text
doc_id	Name of the column with documents' unique ids.
vocab	Default is NULL, if a list of terms is provided, the function will return a DTM with terms restricted to this vocabulary. Columns will also be in the same order as the list of terms.

maxlen	Integer indicating the maximum document length. If NULL (default), the length of the longest document is used.
matrix	Logical, TRUE (default) returns a matrix, FALSE a list

Details

Function will return a matrix of integer sequences by default. The columns will be the length of the longest document or maxlen, with shorter documents padded with zeros. The dictionary will be an attribute of the matrix accessed with `attr(seq, "dic")`. If `matrix = FALSE`, the function will return a list of integer sequences. The vocabulary will either be each unique token in the corpus, or a the list of words provided to the `vocab` argument. This kind of text representation is used in [tensorflow](#) and [keras](#).

Value

returns a matrix or list

Author(s)

Dustin Stoltz

Examples

```
data(jfk_speech)
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)

seqs <- seq_builder(jfk_speech, sentence, sentence_id)
dim(seqs)
```

stoplists

A dataset of stoplists

Description

A dataset containing eight English stoplists. Is used with the [get_stoplist\(\)](#) function.

Format

A data frame with 1775 rows and 2 variables.

Details

The stoplists include:

- "tiny2020": Tiny (2020) list of 33 words (Default)
- "snowball2001": Snowball (2001) list of 127 words
- "snowball2014": Updated Snowball (2014) list of 175 words
- "van1979": van Rijsbergen's (1979) list of 250 words
- "fox1990": Christopher Fox's (1990) list of 421 words
- "smart1993": Original SMART (1993) list of 570 words
- "onix2000": ONIX (2000) list of 196 words
- "nltk2009": Python's NLTK (2009) list of 179 words

Tiny 2020, is a very small stop list of the most frequent English conjunctions, articles, prepositions, and demonstratives (N=17). Also includes the 8 forms of the copular verb "to be" and the 8 most frequent personal (singular and plural) pronouns (minus gendered and possessive pronouns).

No contractions are included.

Variables

Variables:

- words. words to be stopped
- source. source of the list

Source

- tiny2020: Stoltz and Taylor (2020)
- snowball2001: Porter (2001) Snowball stemming algorithm
- snowball2014: Porter (2014) updated Snowball stoplist
- van1979: van Rijsbergen (1979) "Information Retrieval"
- fox1990: Fox (1990) "A Stop List for General Text"
- smart1993: Salton and Buckley (1993) SMART retrieval system
- onix2000: ONIX (2000) "Oxford English Dictionary" stoplist
- nltk2009: Bird, Loper and Klein (2009) NLTK

test_anchors

*Evaluate anchor sets in defining semantic relations***Description**

This function evaluates how well an anchor set defines a semantic relations using one of two methods: pairdir (which only evaluates semantic directions) or relco which evaluates semantic directions, semantic centroids and compound concepts). See details.

Usage

```
test_anchors(
  anchors,
  wv,
  non_anchors = NULL,
  method,
  all = FALSE,
  type = c("direction", "centroid", "compound"),
  conf = 0.95,
  dir_method = c("paired", "pooled", "L2", "PCA"),
  n_runs = 100,
  null = 0,
  alpha = 0.5,
  seed = NULL,
  order_non_anchors = FALSE,
  summarize = TRUE
)
```

Arguments

anchors	A data frame or list of 'anchor' terms
wv	Matrix of word embedding vectors (a.k.a embedding model) with rows as terms.
non_anchors	For 'relco', terms that are not anchors (random, unrelated, or distinctive terms).
method	Which metric used to evaluate, 'pairdir' or 'relco'
all	Logical (default FALSE). Whether to evaluate all possible pairwise combinations of two sets of anchors. If FALSE only the input pairs are used in evaluation and anchor sets must be of equal lengths.
type	For 'relco', indicate which kind of relation, "direction", "centroid", "compound"
conf	For 'relco', confidence interval
dir_method	For 'relco' and type = "direction", indicate the method for calculating direction ("paired", "pooled", "L2", "PCA"), See get_direction() for details.
n_runs	For 'relco', number of runs
null	For 'relco', null hypothesis, default is 0.
alpha	For 'relco', term selection weighting (default is 0.5)

seed	For 'relco', set sampling seed
order_non_anchors	Logical (default FALSE). For 'relco', if TRUE the subset of non-anchor terms is fixed across runs.
summarize	Logical (default TRUE). Returns a dataframe with AVERAGE scores for input pairs along with each pairs' contribution. If summarize = FALSE, returns a list with each offset matrix, each contribution, and the average score.

Details

PairDir evaluates how parallel two anchor sets are when used to define a semantic direction. According to Boutyline and Johnston (2023):

"We find that PairDir – a measure of parallelism between the offset vectors (and thus of the internal reliability of the estimated relation) – consistently outperforms other reliability metrics in explaining axis accuracy."

Boutyline and Johnston only consider analyst specified pairs. However, if `all = TRUE`, all pairwise combinations of terms between each set are evaluated. This can allow for unequal sets of anchors, however this increases computational complexity considerably.

Relco (anchor reliability coefficient) evaluates how well individual anchors index a given semantic relation in comparison to a set of non-anchor words. This can be used on semantic directions, semantic relations, or compound concepts. See Taylor et al (2025) for details; see also the `CMDist()` function.

Value

dataframe or list

References

Boutyline, Andrei, and Ethan Johnston. 2023. "Forging Better Axes: Evaluating and Improving the Measurement of Semantic Dimensions in Word Embeddings." [doi:10.31235/osf.io/576h3](https://doi.org/10.31235/osf.io/576h3)

Taylor, Marshall, et al. 2025. "A Simulation-Based Slope Metric for Anchor List Reliability in Word Embedding Spaces." [doi:10.31235/osf.io/sc2ub_v3](https://doi.org/10.31235/osf.io/sc2ub_v3)

Examples

```
# load example word embeddings
data(ft_wv_sample)

df_anchors <- data.frame(
  a = c("rest", "rested", "stay", "stand"),
  z = c("coming", "embarked", "fast", "move")
)

# test pairdir
test_anchors(df_anchors, ft_wv_sample, method = "pairdir")
test_anchors(df_anchors, ft_wv_sample, method = "pairdir", all = TRUE)
```

```
# test relco
non_anchors <- c("writ", "alloys", "ills", "atlas", "saturn", "cape", "unfolds")
## centroid
test_anchors(df_anchors[, 1], ft_wv_sample, method = "relco",
             type = "centroid", non_anchors = non_anchors)
## compound
test_anchors(df_anchors$a, ft_wv_sample, method = "relco",
             type = "compound", non_anchors = non_anchors)
## direction
test_anchors(df_anchors, ft_wv_sample, method = "relco",
             type = "direction", dir_method = "paired",
             non_anchors = non_anchors)
test_anchors(df_anchors, ft_wv_sample, method = "relco",
             type = "direction", dir_method = "pooled",
             non_anchors = non_anchors)
```

tiny_gender_tagger	<i>A very tiny "gender" tagger</i>
--------------------	------------------------------------

Description

Provides a small dictionary which matches common English pronouns and nouns to conventional gender categories ("masculine" or "feminine"). There are 20 words in each category.

Usage

```
tiny_gender_tagger()
```

Value

returns a tibble with two columns

Author(s)

Dustin Stoltz

Examples

```
tags <- tiny_gender_tagger()
head(tags)
```

`vocab_builder`*A fast unigram vocabulary builder*

Description

A streamlined function to take raw texts from a column of a data.frame and produce a list of all the unique tokens. Tokenizes by the fixed, single whitespace, and then extracts the unique tokens. This can be used as input to `dtm_builder()` to standardize the vocabulary (i.e. the columns) across multiple DTMs. Prior to building the vocabulary, texts should have whitespace trimmed, if desired, punctuation removed and terms lowercased.

Usage

```
vocab_builder(data, text)
```

Arguments

<code>data</code>	Data.frame with one column of texts
<code>text</code>	Name of the column with documents' text

Value

returns a list of unique terms in a corpus

Author(s)

Dustin Stoltz

Examples

```
data(jfk_speech)
jfk_speech$sentence <- tolower(jfk_speech$sentence)
jfk_speech$sentence <- gsub("[[:punct:]]+", " ", jfk_speech$sentence)
vocab <- vocab_builder(jfk_speech, sentence)
head(vocab)
```

weight_ppmi

*Positive Pointwise Mutual Information Weighting***Description**

Computes the positive pointwise mutual information (PPMI) weighting of a term co-occurrence matrix (TCM). PPMI is a standard transformation applied prior to building word embeddings via matrix factorization (e.g. SVD): it down-weights chance co-occurrences and clips negative associations to zero, yielding a non-negative matrix amenable to factorization.

Usage

```
weight_ppmi(tcm, smooth = 0)
```

Arguments

tcm	A term co-occurrence matrix; either a dense matrix or a sparse <code>Matrix::dgCMatrix</code> . Row and column names are preserved.
smooth	A non-negative numeric scalar added to every cell of tcm (including zeros) before estimating probabilities. Defaults to 0, yielding standard PPMI. Values greater than zero nudge the whole distribution and reduce the influence of rare co-occurrences.

Details

The estimator uses marginal probabilities (row and column sums) rather than a diagonal-count proxy, following the standard formulation of Levy & Goldberg (2014). Smoothing, if requested, is applied to the raw counts (including zero cells) before estimating probabilities.

Value

A sparse `dgCMatrix` of the same dimensions as tcm containing PPMI weights, with row and column names preserved.

Author(s)

Dustin Stoltz

References

Levy, O., & Goldberg, Y. (2014). Neural word embedding as implicit matrix factorization. *Advances in Neural Information Processing Systems (NeurIPS)*.

Examples

```
tcm <- matrix(c(2, 1, 0, 1, 3, 1, 0, 1, 2), nrow = 3,
              dimnames = list(c("a", "b", "c"), c("a", "b", "c")))
weight_ppmi(tcm)
```

Index

* datasets

- anchor_lists, 3
- ft_wv_sample, 24
- jfk_speech, 33
- meta_shakespeare, 34
- stoplists, 47

anchor_lists, 3

CMDist, 3, 8, 37, 38, 40

cmdist (CMDist), 3

CMDist(), 7, 26, 28, 30, 38, 40

CoCA, 3, 6, 6, 37

coca (CoCA), 6

CoCA(), 28, 39

doc_centrality, 9

doc_similarity, 11

dtm_builder, 12, 20

dtm_builder(), 7

dtm_melter, 15

dtm_resampler, 16

dtm_stats, 17

dtm_stopper, 18

find_projection, 21

find_rejection, 22

find_transformation, 22

ft_wv_sample, 24

get_anchors, 3, 25

get_anchors(), 3

get_centroid, 3, 6, 26

get_centroid(), 4, 5, 21, 22, 27, 28

get_centroids, 27

get_centroids(), 27

get_direction, 3, 6, 8, 28, 37

get_direction(), 3–5, 7, 21, 22, 25

get_regions, 30

get_regions(), 4, 5, 21, 22, 30

get_stoplist, 32, 32

get_stoplist(), 47

jfk_speech, 33

meta_shakespeare, 34

perm_tester, 35

plot.CMDist, 6, 37, 40

plot.CMDist(), 5

plot.CoCA, 39

print.CMDist, 6, 38, 40

print.CMDist(), 5

print.CoCA, 41

rancor_builder, 44

rancors_builder, 42

seq_builder, 46

stoplists, 47

test_anchors, 49

tiny_gender_tagger, 51

vocab_builder, 52

weight_ppmi, 53