

Package ‘GMTM’

September 22, 2026

Title Gaussian Mixture Topic Models

Version 0.2.0

Maintainer Kohei Watanabe <watanabe.kohei@gmail.com>

Description Gaussian mixture models (GMM) and k-means for topic analysis of dense document vectors. The underlying clustering functions rely on the Armadillo library (Sanderson & Curtin, 2017) <[doi:10.1109/ICSPCS.2017.8270510](https://doi.org/10.1109/ICSPCS.2017.8270510)>.

License Apache License (>= 2.0)

Encoding UTF-8

Depends R (>= 3.5.0)

Imports quanteda (>= 4.0.0), proxyC, wordvector (>= 0.6.4), Rcpp,
RcppArmadillo

Suggests testthat, spelling, rmarkdown, knitr, withr

LinkingTo Rcpp, RcppArmadillo (>= 0.7.600.1.0)

Language en-US

Config/roxygen2/version 8.1.0

NeedsCompilation yes

Author Kohei Watanabe [aut, cre, cph] (ORCID:
<<https://orcid.org/0000-0001-6519-5265>>),
Sanderson Conrad [ctb, cph] (C++ code for GMM and k-means),
Curtin Ryan [ctb, cph] (C++ code for GMM and k-means)

Repository CRAN

Date/Publication 2026-09-22 07:50:19 UTC

Contents

as.seedwords	2
probability.textmodel_gmm	3
terms	3
textmodel_gmm	4
textmodel_kmeans	5
topics	6

Index**8**

as.seedwords	<i>Convert a dictionary to a seed word matrix</i>
--------------	---

Description

Convert a dictionary to a seed word matrix

Usage

```
as.seedwords(x, model, residual = 0, levels = 1)
```

Arguments

x	a quanteda::dictionary of seed words.
model	a wordvector::textmodel_word2vec object.
residual	the number of unseeded topics.
levels	integers specifying the levels of entries in a hierarchical dictionary.

Details

Unseeded topics are labeled "other", but it can be changed via options("GMTM.residual.name").

Value

Returns a seed word matrix.

Examples

```
library(quanteda)
library(wordvector)
options(wordvector_threads = 2)

corp <- head(wordvector::data_corpus_news2014, 1000)
toks <- tokens(corp, remove_punct = TRUE,
               remove_symbols = TRUE, remove_numbers = TRUE) %>%
  tokens_remove(stopwords("en"), min_nchar = 2)
wov <- textmodel_word2vec(toks, dim = 50)

dict <- dictionary(list(eco = "econom*", sec = "securit*", spo = "sport*",
                       pol = "politi*", cri = "crime"))
seed <- as.seedwords(dict, wov, residual = 2)
```

probability.textmodel_gmm

Extract the probabilities of topics

Description

Extract the probabilities of topics

Usage

```
## S3 method for class 'textmodel_gmm'
probability(x, group = FALSE, ...)
```

Arguments

x	a fitted model.
group	if TRUE, aggregate the probability of topics by the original document doc_id. Ignored if x is a textmodel_kmeans object.
...	not used.

Details

The original doc_id is inherited from [quanteda::dfm](#) or [quanteda::tokens](#) and saved in x\$dovars\$docid_ as factor.

Value

Returns the probabilities of topics as a matrix.

terms

Extract words for topics from documents

Description

Identify distinctive words for each topic by applying TF-IDF weights to the original [quanteda::dfm](#).

Usage

```
terms(x, data, n = 10, ...)
```

Arguments

x	a fitted model or a factor from <code>GMTM::topics()</code> .
data	a quanteda::dfm or quanteda::tokens from which words are extracted for each topic.
n	the number of topic words.
...	passed to functions.

Details

To identify distinctive words for topics, original documents must be provided along with a fitted model because the information about individual words are lost in document vectors. The documents in data is grouped by topic and weighted by TF-IDF to select the most distinctive words for each topic. This technique is commonly known as c-TF-IDF.

Value

Returns a character matrix with the most distinctive words for each topic.

textmodel_gmm	<i>Topic analysis using Gaussian mixture models</i>
---------------	---

Description

Perform topic analysis of document vectors using Gaussian mixture models.

Usage

```
textmodel_gmm(
  x,
  k = 10,
  model = NULL,
  seeds = NULL,
  omit = NULL,
  verbose = quanteda_options("verbose"),
  ...
)
```

Arguments

x	a wordvector::textmodel_doc2vec or a dense matrix of document vectors in the rows.
k	the number of topics to identify.
model	a fitted model from which initial centroids are extracted.
seeds	a matrix created using as.seedwords .
omit	indices of singular values of x to be zero. See the details.
verbose	print the progress if TRUE.
...	passed to the underlying function.

Details

Users can change the number of threads for the parallel computing via `options(GMTM.threads)` or `OMP_THREAD_LIMIT` in the environmental variable. To reproduce results, set `options(GMTM.threads = 1)` and call `set.seed()` immediately before `textmodel_gmm()` or `textmodel_kmeans()`.

On MacOS, only one thread is used regardless of `GMTM.threads` because CRAN's toolchain for the platform does not support OpenMP.

`omit` is used to reduce the noise in the `x` by applying `base::svd` before clustering. If it is not `NULL`, singular values corresponding to `omit` are set to zero, removing their variance in `x`. See Chan et al. (2020) [doi:10.1080/19312458.2020.1812555](https://doi.org/10.1080/19312458.2020.1812555) for the methodology.

The number of iterations in k-means (`iter_km`) and expectation maximization (`iter_em`) stages can be set via `...`

Value

Returns a fitted `textmodel_gmm` object.

Examples

```
library(quanteda)
library(wordvector)
options(wordvector_threads = 2)

corp <- head(wordvector::data_corpus_news2014, 1000)
toks <- tokens(corp, remove_punct = TRUE,
               remove_symbols = TRUE, remove_numbers = TRUE) %>%
  tokens_remove(stopwords("en"), min_nchar = 2)
wov <- textmodel_word2vec(toks, dim = 50)
dov <- as.textmodel_doc2vec(dfm(toks), wov)

gmm <- textmodel_gmm(dov, k = 10)
table(topics(gmm))
```

textmodel_kmeans	<i>Topic analysis using k-means</i>
------------------	-------------------------------------

Description

Perform topic analysis of document vectors using k-means.

Usage

```
textmodel_kmeans(
  x,
  k = 10,
  model = NULL,
  seeds = NULL,
  verbose = quanteda_options("verbose"),
  ...
)
```

Arguments

x	a wordvector::textmodel_doc2vec or a dense matrix of document vectors in the rows.
k	the number of topics to identify.
model	a fitted model from which initial centroids are extracted.
seeds	a matrix created using as.seedwords .
verbose	print the progress if TRUE.
...	passed to the underlying function.

Value

Returns a fitted `textmodel_kmeans` object.

Examples

```
library(quantda)
library(wordvector)
options(wordvector_threads = 2)

corp <- head(wordvector::data_corpus_news2014, 1000)
toks <- tokens(corp, remove_punct = TRUE,
               remove_symbols = TRUE, remove_numbers = TRUE) %>%
  tokens_remove(stopwords("en"), min_nchar = 2)
wov <- textmodel_word2vec(toks, dim = 50)
dov <- as.textmodel_doc2vec(dfm(toks), wov)

km <- textmodel_kmeans(dov, k = 10)
table(topics(km))
```

topics	<i>Extract the topics of documents</i>
--------	--

Description

Extract the topics of documents

Usage

```
topics(x, group = FALSE, ...)
```

Arguments

x	a fitted model.
group	if TRUE, aggregate the probability of topics by the original document <code>doc_id</code> . Ignored if x is a <code>textmodel_kmeans</code> object.
...	not used.

Details

The original doc_id is inherited from [quanteda::dfm](#) or [quanteda::tokens](#) and saved in `x$dovars$docid_` as factor.

Value

Returns predicted topics as a vector.

Index

`as.seedwords`, [2](#), [4](#), [6](#)

`probability.textmodel_gmm`, [3](#)

`quanteda::dfm`, [3](#), [7](#)

`quanteda::dictionary`, [2](#)

`quanteda::tokens`, [3](#), [7](#)

`terms`, [3](#)

`textmodel_gmm`, [4](#)

`textmodel_kmeans`, [5](#)

`topics`, [6](#)

`wordvector::textmodel_doc2vec`, [4](#), [6](#)

`wordvector::textmodel_word2vec`, [2](#)