

Package ‘immunarch’

October 15, 2025

Type Package

Title Multi-Modal Immune Repertoire Analytics for Immunotherapy and Vaccine Design in R

Version 0.10.3

Contact support@immunomind.com

Description A comprehensive analytics framework for building reproducible pipelines on T-cell and B-cell immune receptor repertoire data. Delivers multi-modal immune profiling (bulk, single-cell, CITE-seq/AbSeq, spatial, immunogenicity data), feature engineering (ML-ready feature tables and matrices), and biomarker discovery workflows (cohort comparisons, longitudinal tracking, repertoire similarity, enrichment). Provides a user-friendly interface to widely used AIRR methods — clonality/diversity, V(D)J usage, similarity, annotation, tracking, and many more. Think Scanpy or Seurat, but for AIRR data, a.k.a. Adaptive Immune Receptor Repertoire, VDJ-seq, RepSeq, or VDJ sequencing data. A successor to our previously published ``tcR'' R package (Nazarov 2015).

License Apache License (>= 2.0)

URL <https://immunomind.github.io/docs/>,
<https://github.com/immunomind/immunarch/>,
<https://immunarch.com/>

BugReports <https://github.com/immunomind/immunarch/issues>

Depends R (>= 4.2.0), ggplot2 (>= 3.1.0), immundata (>= 0.0.5), patchwork

Imports dplyr, dtplyr (>= 1.0.0), data.table (>= 1.12.6), cli, pheatmap (>= 1.0.12), reshape2 (>= 1.4.2), circlize, airr, Rcpp (>= 1.0), magrittr, methods, scales, rlang, plyr, stringdist, readr, stringr, tibble, tidyselect, tidyr, ape, doParallel, rlist, utils, glue, checkmate, duckplyr (>= 1.1.2), dbplyr, lifecycle, purrr, stats, vctrs, ggthemes, ggsci

LinkingTo Rcpp

Suggests knitr (>= 1.8), roxygen2 (>= 3.0.0), testthat (>= 3.0.0),
 pkgdown (>= 0.1.0), assertthat, rmarkdown, factoextra (>= 1.0.4), fpc, ggpubr (>= 0.2), ggraph, ggseqlogo, igraph,
 phangorn, ggalluvial (>= 0.10.0), UpSetR (>= 1.4.0), ggrepel
 (>= 0.8.0), shiny (>= 1.4.0), shinythemes, quarto, MASS (>= 7.3), Rtsne (>= 0.15)

VignetteBuilder knitr, quarto

Encoding UTF-8

RoxygenNote 7.3.3

LazyData true

LazyDataCompression xz

Config/Needs/website rmarkdown

Config/testthat/edition 3

NeedsCompilation yes

Author Vadim I. Nazarov [aut, cre] (ORCID:
<https://orcid.org/0000-0003-3659-2709>),
 Vasily O. Tsvetkov [aut],
 Aleksandr A. Popov [aut],
 Ivan Balashov [aut]

Maintainer Vadim I. Nazarov <support@immunomind.com>

Repository CRAN

Date/Publication 2025-10-14 23:30:02 UTC

Contents

.quant_column_choice	4
aa_properties	5
aa_table	5
add_class	5
airr_clonality	6
airr_diversity	8
airr_public	12
airr_stats	14
annotate_clonality	17
apply_symm	19
bcrdata	19
bunch_translate	20
check_distribution	21
coding	22
dbAnnotate	23
dbLoad	24
entropy	25
fixVis	26
geneUsage	27

geneUsageAnalysis	28
gene_segments	29
gene_stats	30
getKmers	30
get_immunarch_news	31
group_from_metadata	32
has_class	32
immdata	33
immunarch_v1_updates	33
immunr_data_format	34
immunr_hclust	34
immunr_pca	36
inc_overlap	37
list_immunarch_news	38
matrixdiagcopy	39
public_matrix	39
pubRep	40
pubRepApply	41
pubRepFilter	42
pubRepStatistics	43
repAlignLineage	43
repClonalFamily	45
repClonality	46
repDiversity	48
repExplore	51
repFilter	52
repGermline	53
repLoad	54
repOverlap	57
repOverlapAnalysis	59
repSample	60
repSave	62
repSomaticHypermutation	63
select_barcodes	64
select_clusters	65
seqCluster	66
seqDist	67
set_pb	69
spectratype	69
split_to_kmers	70
switch_type	71
top	72
trackClonotypes	73
vis	74
vis.clonal_family	76
vis.clonal_family_tree	77
vis.immunr_chao1	78
vis.immunr_clonal_prop	80

vis.immunr_dynamics 82

vis.immunr_exp_vol 83

vis.immunr_gene_usage 85

vis.immunr_hclust 86

vis.immunr_inc_overlap 87

vis.immunr_kmeans 88

vis.immunr_kmer_table 89

vis.immunr_mds 90

vis.immunr_ov_matrix 91

vis.immunr_public_repertoire 92

vis.immunr_public_statistics 93

vis.step_failure_ignored 94

vis_bar 95

vis_box 96

vis_circos 98

vis_heatmap 99

vis_heatmap2 100

vis_hist 102

vis_immunr_kmer_profile_main 104

vis_public_clonotypes 105

vis_public_frequencies 106

vis_textlogo 107

Index 109

.quant_column_choice	Get a column's name using the input alias
----------------------	---

Description

Get a column's name using the input alias

Usage

.quant_column_choice(x)

Arguments

x Character vector of length 1.

Value

A string with the column name.

Developer Examples

immunarch:::quant_column_choice("count") immunarch:::quant_column_choice("freq")

aa_properties	<i>Tables with amino acid properties</i>
---------------	--

Description

Tables with amino acid properties

aa_table	<i>Amino acid / codon table</i>
----------	---------------------------------

Description

Amino acid / codon table

Usage

AA_TABLE

Format

An object of class table of length 65.

add_class	<i>Add a new class attribute</i>
-----------	----------------------------------

Description

Add a new class attribute

Usage

add_class(.obj, .class)

Arguments

- .obj R object.
- .class String with the desired class name.

Value

Input object with additional class .class.

Developer Examples

tmp <- "abc" class(tmp) tmp <- immunarch:::add_class(tmp, "new_class") class(tmp)

Description

[Experimental]

A family of functions to quantify **receptor overabundance** per repertoire. Helps in deciphering the structure and partition the repertoire.

Available functions:

Supported methods are the following.

`airr_clonality_line` - build ranked abundance lines: for each repertoire, take the top limit receptors by count and attach repertoire metadata. Useful for per-repertoire rank-abundance plots.

`airr_clonality_rank` - aggregate clonal space by **rank bins**. Receptors are ordered by proportion within each repertoire; each receptor is assigned to the smallest threshold in bins that contains its rank.

`airr_clonality_prop` - aggregate clonal space by **proportion bins**. Each receptor is assigned to a named bin according to its proportion (e.g., Hyperexpanded $\geq 1e-2$, Large $\geq 1e-3$, ...). Thresholds are matched in descending order; unmatched receptors fall into "Ultra-rare".

Usage

```
airr_clonality_line(
  idata,
  limit = 1e+05,
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)

airr_clonality_rank(
  idata,
  bins = c(10, 30, 100, 300, 1000, 10000, 1e+05),
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)

airr_clonality_prop(
  idata,
  bins = c(Hyperexpanded = 0.01, Large = 0.001, Medium = 1e-04, Small = 1e-05, Rare =
    1e-06),
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)
```

Arguments

<code>idata</code>	An ImmunData object.
<code>limit</code>	Positive integer ≥ 10 : maximum number of top receptors to keep per repertoire (default 100000).
<code>autojoin</code>	Logical. If TRUE, join repertoire metadata by the schema repertoire id. Change the default behaviour by calling <code>options(immunarch.autojoin = FALSE)</code> .
<code>format</code>	String. One of "long" ("long" tibble with <code>imd_repertoire_id</code> , facet columns, and value; useful for visualizations) or "wide" (wide/unmelted table of features, with each row corresponding to a specific repertoire / pair of repertoires; useful for Machine Learning).
<code>bins</code>	A named numeric vector of thresholds (e.g., <code>c(Hyperexpanded = 1e-2, Large = 1e-3, ...)</code>). Names become bin labels and must be non-empty. Internally sorted in descending order.

Value

`airr_clonality_line`:

A tibble with columns:

- `repertoire_id` - repertoire identifier
- `index` - rank within repertoire (1 = most abundant)
- `count` - receptor count used for ranking
- plus any repertoire metadata columns carried from `idata$repertoires`

`airr_clonality_rank`:

A tibble with

- `repertoire_id`
- `clonal_rank_bin` - the rank threshold (e.g., 10, 100, ...)
- `occupied_prop` - sum of proportion within the bin
- plus repertoire metadata columns from `idata$repertoires`

`airr_clonality_prop`:

A tibble with

- `repertoire_id`
- `clonal_prop_bin` - factor-like label from `names(bins)` or "Ultra-rare"
- `occupied_prop` - sum of proportion within the bin
- plus repertoire metadata columns from `idata$repertoires`

See Also

- Per-repertoire summaries: [annotate_clonality](#)
- Data container: [immundata::ImmunData](#)

Examples

```
# Limit the number of threads used by the underlying DB for this session.
# Change this only if you know what you're doing (e.g., multi-user machines, shared CI/servers).
db_exec("SET threads TO 1")

# Load data
## Not run:
immdata <- get_test_idata() |> agg_repertoires("Therapy")

## End(Not run)

#
# airr_clonality_line
#
## Not run:
top_line <- airr_clonality_line(immdata, limit = 1000)

## End(Not run)

#
# airr_clonality_rank
#
## Not run:
rank_stat <- airr_clonality_rank(immdata, bins = c(10, 100))

## End(Not run)

#
# airr_clonality_prop
#
## Not run:
prop_stat <- airr_clonality_prop(immdata)

## End(Not run)
```

airr_diversity

Diversity - estimating the heterogeneity of immune repertoires

Description

[Experimental]

A family of functions to quantify **receptor diversity** per repertoire. A characteristic of a whole repertoire.

Available functions:

Supported methods are the following.

`airr_diversity_dxx` - **coverage diversity**: minimal number of top receptors needed to reach perc% of clonal space (by proportion). Great for spotting dominance/overexpansion and for quick, interpretable dashboards (e.g., D50 = receptors to cover half of the repertoire).

`airr_diversity_chao1` - Chao1 estimator is a nonparameteric asymptotic estimator of species richness (number of species in a population). One of the most used methods for estimating immune repertoire diversity.

`airr_diversity_shannon` - Shannon entropy (base 2) per repertoire computed from proportion. Ideal when you want a single evenness-aware diversity score; pair with Pielou/Hill for samples with very different richness.

`airr_diversity_pielou` - Pielou's evenness $H / \log_2(S)$ with richness S . Best when you need a **size-normalized** evenness score that's comparable across repertoires with different receptor counts.

`airr_diversity_index` - convenience alias for Hill number with $q = 1$ ($\exp(\text{Shannon})$) using natural log). A solid **default single metric** that's relatively robust to rare-count noise and easy to compare across samples.

`airr_diversity_hill` - Hill numbers ("true diversity") for orders $q \in \{0, 1, 2, \dots\}$: $q=0$ richness, $q=1$ $\exp(\text{Shannon})$, $q>1$ emphasizes abundant receptors. Perfect when you want a **diversity profile** that tunes sensitivity to rare vs. abundant clonotypes.

Usage

```
airr_diversity_dxx(
  idata,
  perc = 50,
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)
```

```
airr_diversity_chao1(
  idata,
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)
```

```
airr_diversity_shannon(
  idata,
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)
```

```
airr_diversity_pielou(
  idata,
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)
```

```
airr_diversity_index(
  idata,
```

```

    autojoin = getOption("immundata.autojoin", TRUE),
    format = c("long", "wide")
  )

  airr_diversity_hill(
    idata,
    q = 0:5,
    autojoin = getOption("immundata.autojoin", TRUE),
    format = c("long", "wide")
  )

```

Arguments

<code>idata</code>	An ImmunData object.
<code>perc</code>	A number or numeric vector in $(0, 100]$ (default 50), e.g. 50 for D50, 20 for D20.
<code>autojoin</code>	Logical. If TRUE, join repertoire metadata by the schema repertoire id. Change the default behaviour by calling <code>options(immunarch.autojoin = FALSE)</code> .
<code>format</code>	String. One of "long" ("long" tibble with <code>imd_repertoire_id</code> , facet columns, and value; useful for visualizations) or "wide" (wide/unmelted table of features, with each row corresponding to a specific repertoire / pair of repertoires; useful for Machine Learning).
<code>q</code>	A scalar or vector of non-negative orders. Defaults to <code>0:5</code> .

Value

`airr_diversity_dxx:`

A tibble with:

- `imd_repertoire_id`
- `perc`
- `dxx` - minimal count of top receptors to reach `perc%`
- plus repertoire metadata from `idata$repertoires`

`airr_diversity_chao1:`

A tibble with:

- `imd_repertoire_id`
- Estimator - number of species
- SD - standard deviation for the estimator value
- `Conf.95.lo` - CI 0.025
- `Conf.95.hi` - CI 0.975
- plus repertoire metadata from `idata$repertoires`

`airr_diversity_shannon:`

A tibble with:

- `imd_repertoire_id`

- shannon - entropy in bits

airr_diversity_pielou:

A tibble with:

- imd_repertoire_id
- shannon
- n_receptors
- pielou - evenness in $[0, 1]$ (NA if $S \leq 1$)

airr_diversity_index:

A tibble with:

- imd_repertoire_id
- $q = 1$
- hill_number
- plus repertoire metadata from `idata$repertoires`

airr_diversity_hill:

A tibble with:

- imd_repertoire_id
- q - Hill order
- hill_number - true diversity of order q
- plus repertoire metadata from `idata$repertoires`

See Also

[immunodata::ImmunData](#)

Examples

```
# Limit the number of threads used by the underlying DB for this session.
# Change this only if you know what you're doing (e.g., multi-user machines, shared CI/servers).
db_exec("SET threads TO 1")
# Load data
## Not run:
immdata <- get_test_idata() |> agg_repertoires("Therapy")

## End(Not run)

#
# airr_diversity_dxx
#
## Not run:
d50 <- airr_diversity_dxx(immdata, perc = 50)
d_multi <- airr_diversity_dxx(immdata, perc = c(20, 50, 80))

## End(Not run)

#
```

```

# airr_diversity_chao1
#
## Not run:
chao <- airr_diversity_chao1(immdata)

## End(Not run)

#
# airr_diversity_shannon
#
## Not run:
sh <- airr_diversity_shannon(immdata)

## End(Not run)

#
# airr_diversity_pielou
#
## Not run:
pj <- airr_diversity_pielou(immdata)

## End(Not run)

#
# airr_diversity_index
#
## Not run:
idx <- airr_diversity_index(immdata)

## End(Not run)

#
# airr_diversity_hill
#
## Not run:
hill <- airr_diversity_hill(immdata, q = c(0, 1, 2))

## End(Not run)

```

airr_public

Public indices - pairwise repertoire overlap

Description

[Experimental]

A family of functions to quantify **public or shared receptors** between repertoire.

Available functions:

Supported methods are the following.

`airr_public_intersection` - number of **shared receptors** between each pair of repertoires (intersection size). Handy for quick overlap heatmaps, QC of replicate similarity, or spotting donor-shared "public" clonotypes.

`airr_public_jaccard` - **Jaccard similarity** of receptor sets between repertoires ($A \cap B / A \cup B$). Best when comparing cohorts with different sizes to get a scale-invariant overlap score.

Usage

```
airr_public_intersection(
  idata,
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)

airr_public_jaccard(
  idata,
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)
```

Arguments

<code>idata</code>	An ImmunData object.
<code>autojoin</code>	Logical. If TRUE, join repertoire metadata by the schema repertoire id. Change the default behaviour by calling <code>options(immunarch.autojoin = FALSE)</code> .
<code>format</code>	String. One of "long" ("long" tibble with <code>imd_repertoire_id</code> , facet columns, and value; useful for visualizations) or "wide" (wide/unmelted table of features, with each row corresponding to a specific repertoire / pair of repertoires; useful for Machine Learning).

Value

`airr_public_intersection`:

A **symmetric numeric matrix** where rows/columns are `repertoire_id` and each cell is the count of shared unique receptors. The diagonal contains per-repertoire richness (total unique receptors). Row/column names are repertoire IDs.

`airr_public_jaccard`:

A **symmetric numeric matrix** where rows/columns are `repertoire_id` and each cell is the Jaccard similarity in $[0, 1]$. The diagonal is 1. Row/column names are repertoire IDs.

See Also

[immundata::ImmunData](#)

Examples

```
# Limit the number of threads used by the underlying DB for this session.
# Change this only if you know what you're doing (e.g., multi-user machines, shared CI/servers).
db_exec("SET threads TO 1")
# Load data
immdata <- get_test_idata() |> agg_repertoires("Therapy")

#
# airr_public_intersection
#
## Not run:
m_pub <- airr_public_intersection(immdata)

## End(Not run)

#
# airr_public_jaccard
#
## Not run:
m_jac <- airr_public_jaccard(immdata)

## End(Not run)
```

airr_stats

Compute key immune repertoire statistics

Description

[Experimental]

A family of functions that extract **core descriptive statistics** from an ImmunData object.

Available functions:

Supported methods are the following.

airr_stats_chains — count V(D)J *chains* per repertoire (optionally split by locus). Quickly gauges capture depth per repertoire and, when split by locus, reveals TRA/TRB/IGH balance. Use it for QC, library-size checks, and to spot locus-specific dropouts or over-representation.

airr_stats_lengths — count the number of sequence lengths per repertoire. Summarizes the CDR3 length distribution, a sensitive QC fingerprint of repertoire prep and selection. Helpful for detecting primer/UMI biases, comparing cohorts, and deriving length-based features for models.

airr_stats_genes - count V(D)J gene segments per repertoire, optionally split by locus and using either receptor counts or barcode/UMI counts as the measure. Profiles V/D/J gene usage to characterize repertoire composition and germline biases, with optional locus split. Useful for cohort comparisons, flagging clonal expansions, and producing ML-ready features for repertoire-level ML tasks.

Usage

```

airr_stats_chains(
  idata,
  locus_col = NA,
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)

airr_stats_lengths(
  idata,
  seq_col = "cdr3_aa",
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)

airr_stats_genes(
  idata,
  gene_col = "v_call",
  level = c("receptor", "barcode"),
  by = c(NA, "locus"),
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)

```

Arguments

<code>idata</code>	An ImmunData object.
<code>locus_col</code>	Column in <code>idata\$annotations</code> that stores the locus (e.g. "locus"). If NULL or missing, the result is not split by locus.
<code>autojoin</code>	Logical. If TRUE, join repertoire metadata by the schema repertoire id. Change the default behaviour by calling <code>options(immunarch.autojoin = FALSE)</code> .
<code>format</code>	String. One of "long" ("long" tibble with <code>imd_repertoire_id</code> , facet columns, and value; useful for visualizations) or "wide" (wide/unmelted table of features, with each row corresponding to a specific repertoire / pair of repertoires; useful for Machine Learning).
<code>seq_col</code>	Character vector with names of the columns containing sequences.
<code>gene_col</code>	A single column name in <code>idata\$annotations</code> with gene segment calls (e.g., "v_call", "d_call", "j_call", "c_call"). Default is "v_call".
<code>level</code>	One of "receptor" or "barcode". If "receptor" (default), the function counts unique receptors (one per receptor ID) that carry a given gene segment. If "barcode", the function sums counts (e.g., cells/UMIs) per gene segment using the column defined by <code>immundata::imd_schema("count")</code> .
<code>by</code>	Either NULL (no split) or "locus". When "locus", the result is further split by the locus column if present (as given by <code>immundata::imd_schema("locus")</code>); otherwise a warning is emitted and the split is ignored.

Value

airr_stats_chains **Returns a tibble with columns::**

- repertoire_id – repertoire identifier
- locus – TRA, TRB, IGH, ... (present only if locus_col is supplied)
- n_chains – number of chains

airr_stats_lengths **Returns a tibble with columns::**

- repertoire_id – repertoire identifier
- seq_len – lengths of sequences
- n – number of receptors

airr_stats_genes **A tibble with columns::**

- repertoire_id - repertoire identifier
- *(optional)* locus - TRA, TRB, IGH, ... (present only when by = "locus" and the locus column exists)
- <gene_col> - the gene segment value (e.g., V gene)
- n - the measure:
 - if level = "receptor": number of receptors carrying the gene segment
 - if level = "barcode": sum of counts across receptors for the segment

See Also

[immundata::ImmunData](#)

Examples

```
# Limit the number of threads used by the underlying DB for this session.
# Change this only if you know what you're doing (e.g., multi-user machines, shared CI/servers).
db_exec("SET threads TO 2")

# Load data
## Not run:
immdata <- get_test_idata() |> agg_repertoires("Therapy")

## End(Not run)

#
# airr_stats_chains
#

## Not run:
airr_stats_chains(immdata)

## End(Not run)

#
# airr_stats_lengths
#
```



```
## Not run:
airr_stats_lengths(immdata)

## End(Not run)

#
# airr_stats_genes
#

## Not run:
# V gene usage by receptor count
airr_stats_genes(immdata, gene_col = "v_call", level = "receptor")

# V gene usage by summed cell/UMI counts (if a count column is present)
airr_stats_genes(immdata, gene_col = "v_call", level = "barcode")

# Split by locus (TRA/TRB/... if locus column exists)
airr_stats_genes(immdata, gene_col = "v_call", level = "receptor", by = "locus")

## End(Not run)
```

annotate_clonality *Annotate clonality - per-receptor labels for overabundance*

Description

[Experimental]

A small family of helpers that **add clonality labels to each receptor** in an [immundata::ImmunData](#) object.

Available functions:

- `annotate_clonality_rank()` - label by **rank bins** within each repertoire.
- `annotate_clonality_prop()` - label by **proportion bins** (named thresholds).

`annotate_clonality_rank()` - for each repertoire, receptors are ordered by within-repertoire abundance (proportion) and assigned a **rank bin** label.

`annotate_clonality_prop()` - label each receptor by **proportion bin** using named thresholds (matched in descending order; else "Ultra-rare").

Usage

```
annotate_clonality_rank(
  idata,
  bins = c(10, 30, 100, 300, 1000, 10000, 1e+05),
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
```

```

)

annotate_clonality_prop(
  idata,
  bins = c(Hyperexpanded = 0.01, Large = 0.001, Medium = 1e-04, Small = 1e-05, Rare =
    1e-06),
  autojoin = getOption("immundata.autojoin", TRUE),
  format = c("long", "wide")
)

```

Arguments

<code>idata</code>	An immundata::ImmunData object.
<code>bins</code>	A named numeric vector of thresholds (e.g., <code>c(Hyperexpanded = 1e-2, Large = 1e-3, ...)</code>). Names become bin labels and must be non-empty. Internally sorted in descending order.
<code>autojoin</code>	Logical. If TRUE, join repertoire metadata by the schema repertoire id. Change the default behaviour by calling <code>options(immunarch.autojoin = FALSE)</code> .
<code>format</code>	String. One of "long" ("long" tibble with <code>imd_repertoire_id</code> , facet columns, and value; useful for visualizations) or "wide" (wide/unmelted table of features, with each row corresponding to a specific repertoire / pair of repertoires; useful for Machine Learning).

Value

An [immundata::ImmunData](#) whose `$annotations` gains:

- `clonal_rank_bin` - integer-like label with the applied rank threshold (outside all thresholds -> NA).

An [immundata::ImmunData](#) whose `$annotations` gains:

- `clonal_prop_bin` - label from `names(bins)` or "Ultra-rare".

See Also

- Per-repertoire summaries: [airr_clonality](#)
- Data container: [immundata::ImmunData](#)

Examples

```

## Not run:
idata <- get_test_idata() |> agg_repertoires("Therapy")
idata_rank <- annotate_clonality_rank(idata)
idata_prop <- annotate_clonality_prop(idata)

## End(Not run)

```

apply_symm	<i>Apply function to each pair of data frames from a list.</i>
------------	--

Description

Apply the given function to every pair in the given datalist. Function either symmetrical (i.e. $\text{fun}(x,y) == \text{fun}(y,x)$) or assymmetrical (i.e. $\text{fun}(x,y) != \text{fun}(y,x)$).

Usage

```
apply_symm(.datalist, .fun, ..., .diag = NA, .verbose = TRUE)
```

```
apply_asymm(.datalist, .fun, ..., .diag = NA, .verbose = TRUE)
```

Arguments

.datalist	List with some data.frames.
.fun	Function to apply, which return basic class value.
...	Arguments passed to .fun.
.diag	Either NA for NA or something else != NULL for .fun(x,x).
.verbose	if TRUE then output a progress bar.

Value

Matrix with values $M(i,j) = \text{fun}(\text{datalist}(i), \text{datalist}(j))$

Examples

```
data(immdata)
apply_symm(immdata$data, function(x, y) {
  nrow(x) + nrow(y)
})
```

bcrdata	<i>BCR dataset</i>
---------	--------------------

Description

A dataset with BCR data for testing and exemplary purposes.

Usage

```
bcrdata
```

Format

A list of two elements. The first element ("data") is a list of 1 element named "full_clones" that contains immune repertoire data frame. The second element ("meta") is empty metadata table.

data List of immune repertoire data frames.

meta Metadata ...

bunch_translate	<i>Nucleotide to amino acid sequence translation</i>
-----------------	--

Description

Nucleotide to amino acid sequence translation

Usage

```
bunch_translate(.seq, .two.way = TRUE, .ignore.n = FALSE)
```

Arguments

<code>.seq</code>	Vector or list of strings.
<code>.two.way</code>	Logical. If TRUE (default) then translate from the both ends (like MIXCR).
<code>.ignore.n</code>	Logical. If FALSE (default) then return NA for sequences that have N, else parse triplets with N as ~

Value

Character vector of translated input sequences.

Examples

```
data(immdata)
head(bunch_translate(immdata$data[[1]]$CDR3.nt))
```

check_distribution	<i>Check and normalise distributions</i>
--------------------	--

Description

Check if the given `.data` is a distribution and normalise it if necessary with an optional Laplace correction.

Usage

```
check_distribution(  
  .data,  
  .do.norm = NA,  
  .laplace = 1,  
  .na.val = 0,  
  .warn.zero = FALSE,  
  .warn.sum = TRUE  
)
```

Arguments

<code>.data</code>	Numeric vector of values.
<code>.do.norm</code>	One of the three values - NA, TRUE or FALSE. If NA then checks for distribution ($\text{sum}(.data) == 1$) and normalises if needed with the given laplace correction value. if TRUE then does the normalisation and laplace correction. If FALSE then doesn't do either normalisation or laplace correction.
<code>.laplace</code>	Value for the laplace correction.
<code>.na.val</code>	Replace all NAs with this value.
<code>.warn.zero</code>	if TRUE then the function checks if in the resulted vector (after normalisation) are any zeros, and prints a warning message if there are some.
<code>.warn.sum</code>	if TRUE then the function checks if the sum of resulted vector (after normalisation) is equal to one, and prints a warning message if not.

Value

Numeric vector.

Developer Examples

```
immunarch:::check_distribution(c(1, 2, 3)) immunarch:::check_distribution(c(1, 2, 3), TRUE) im-  
munarch:::check_distribution(c(1, 2, 3), FALSE)
```

coding

*Filter out coding and non-coding clonotype sequences***Description**

Filter out clonotypes with non-coding, coding, in-frame or out-of-frame CDR3 sequences:

`coding()` - remove all non-coding sequences (i.e., remove all sequences with stop codons and frame shifts);

`noncoding()` - remove all coding sequences (i.e., leave sequences with stop codons and frame shifts only);

`inframes()` - remove all out-of-frame sequences (i.e., remove all sequences with frame shifts);

`outofframes()` - remove all in-frame sequences (i.e., leave sequences with frame shifts only).

Note: the function will remove all clonotypes sequences with NAs in the CDR3 amino acid column.

Usage

```
coding(.data)
```

```
noncoding(.data)
```

```
inframes(.data)
```

```
outofframes(.data)
```

Arguments

`.data` The data to be processed. Can be [data.frame](#), [data.table::data.table](#), or a list of these objects.

Every object must have columns in the immunarch compatible format. [immunarch_data_format](#)

Competent users may provide advanced data representations: DBI database connections, Apache Spark DataFrame from "copy_to" or a list of these objects. They are supported with the same limitations as basic objects.

Note: each connection must represent a separate repertoire.

Value

Filtered data frame.

Examples

```
data(immdata)
immdata_cod <- coding(immdata$data)
immdata_cod1 <- coding(immdata$data[[1]])
```

dbAnnotate	<i>Annotate clonotypes in immune repertoires using clonotype databases (e.g., VDJDB, McPAS)</i>
------------	---

Description

[Deprecated]

Annotate clonotypes by matching them to known condition-associated immune receptors in a database. Before using this function, you must download or load the relevant database files. For more information, see the [online tutorial](#).

Usage

```
dbAnnotate(.data, .db, .data.col, .db.col)
```

Arguments

<code>.data</code>	The data to process. It can be a data.frame, a data.table::data.table, or a list of these objects. Every object must have columns in the immunarch compatible format. immunarch_data_format Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects. Note: each connection must represent a separate repertoire.
<code>.db</code>	A data frame or a data table with an immune receptor database. See dbLoad on how to load databases into R.
<code>.data.col</code>	Character vector. Vector of columns in the input repertoires to use for clonotype search. E.g., "CDR3.aa" or c("CDR3.aa", "V.name").
<code>.db.col</code>	Character vector. Vector of columns in the database to use for clonotype search. The order must match the order of ".data.col". E.g., if ".data.col" is c("CDR3.aa", "V.name"), then ".db.col" must have the exact order of columns. i.e., the first column must correspond to CDR3 amino acid sequences, and the second column must correspond to V gene segment names.

Value

Data frame with input sequences and counts or proportions for each of the input repertoire.

Examples

```
data(immdata)

#' # Example file path
file_path <- paste0(system.file(package = "immunarch"), "/extdata/db/vdjdb.example.txt")
```

```
# Load the database with human-only TRB-only receptors for all known antigens
db <- dbLoad(file_path, "vdjdb", "HomoSapiens", "TRB")

res <- dbAnnotate(immdata$data, db, "CDR3.aa", "cdr3")
res
```

dbLoad	<i>Load clonotype databases such as VDJDB and McPAS into the R workspace</i>
--------	--

Description

[Deprecated]

The function automatically detects the database format and loads it into R. Additionally, the function provides a general query interface to databases that allows filtering by species, chain types (i.e., locus) and pathology (i.e., antigen species).

Currently we support three popular databases:

VDJDB - <https://github.com/antigenomics/vdjdb-db>

McPAS-TCR - <https://friedmanlab.weizmann.ac.il/McPAS-TCR/>

TBAdb from PIRD - <https://db.cngb.org/pird/>

Usage

```
dbLoad(.path, .db, .species = NA, .chain = NA, .pathology = NA)
```

Arguments

.path	Character. A path to the database file, e.g., "/Users/researcher/Downloads/McPAS-TCR.csv".
.db	Character. A database type: either "vdjdb", "vdjdb-search", "mcpas" or "tbadb". "vdjdb" for VDJDB; "vdjdb-search" for search table obtained from the web interface of VDJDB; "mcpas" for McPAS-TCR; "tbadb" for PIRD TBAdb.
.species	Character. A string or a vector of strings specifying which species need to be in the database, e.g., "HomoSapiens". Pass NA (by default) to load all available species.
.chain	Character. A string or a vector of strings specifying which chains need to be in the database, e.g., "TRB". Pass NA (by default) to load all available chains.
.pathology	Character. A string or a vector of strings specifying which disease, virus, bacteria or any condition needs to be in the database, e.g., "CMV". Pass NA (by default) to load all available conditions.

Value

Data frame with the input database records.

Examples

```
# Example file path
file_path <- paste0(system.file(package = "immunarch"), "/extdata/db/vdjdb.example.txt")

# Load the database with human-only TRB-only receptors for all known antigens
db <- dbLoad(file_path, "vdjdb", "HomoSapiens", "TRB")
db
```

entropy

*Information measures***Description****[Deprecated]**

Compute information-based estimates and distances.

Usage

```
entropy(.data, .base = 2, .norm = FALSE, .do.norm = NA, .laplace = 1e-12)

kl_div(.alpha, .beta, .base = 2, .do.norm = NA, .laplace = 1e-12)

js_div(.alpha, .beta, .base = 2, .do.norm = NA, .laplace = 1e-12, .norm.entropy = FALSE)

cross_entropy(.alpha, .beta, .base = 2, .do.norm = NA,
              .laplace = 1e-12, .norm.entropy = FALSE)
```

Arguments

.data	Numeric vector. Any distribution.
.base	Numeric. A base of logarithm.
.norm	Logical. If TRUE then normalises the entropy by the maximal value of the entropy.
.do.norm	If TRUE then normalises the input distributions to make them sum up to 1.
.laplace	Numeric. A value for the laplace correction.
.alpha	Numeric vector. A distribution of some random value.
.beta	Numeric vector. A distribution of some random value.
.norm.entropy	Logical. If TRUE then normalises the resulting value by the average entropy of input distributions.

Value

A numeric value.

Examples

```
P <- abs(rnorm(10))
Q <- abs(rnorm(10))
entropy(P)
kl_div(P, Q)
js_div(P, Q)
cross_entropy(P, Q)
```

fixVis

Manipulate ggplot plots and create publication-ready plots

Description

[Deprecated]

The fixVis is a built-in software tool for the manipulation of plots, such as adjusting title text font and size, axes, and more. It is a powerful tool designed to produce publication-ready plots with minimal amount of coding.

Usage

```
fixVis(.plot = NA)
```

Arguments

.plot A ggplot2 plot.

Value

No return value because it is an application.

Examples

```
if (interactive()) {
  # Compute gene usage, visualise it and tweak via fixVis
  data(immdata) # load test data
  gu <- geneUsage(immdata$data)
  p <- vis(gu)
  fixVis(p)
}
```

geneUsage

Main function for estimation of V-gene and J-gene statistics

Description

[Deprecated]

An utility function to analyse the immune receptor gene usage (IGHD, IGHJ, IDHV, IGIJ, IGKJ, IGKV, IGLJ, IGLV, TRAJ, TRAV, TRBD, etc.) and statistics. For gene details run `gene_stats()`.

Usage

```
geneUsage(
  .data,
  .gene = c("hs.trbv", "HomoSapiens.TRBJ", "macmul.IGHV"),
  .quant = c(NA, "count"),
  .ambig = c("inc", "exc", "maj"),
  .type = c("segment", "allele", "family"),
  .norm = FALSE
)
```

Arguments

- | | |
|---------------------|---|
| <code>.data</code> | <p>The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects.</p> <p>Every object must have columns in the immunarch compatible format. immunarch_data_format</p> <p>Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects.</p> <p>Note: each connection must represent a separate repertoire.</p> |
| <code>.gene</code> | <p>A character vector of length one with the name of the gene you want to analyse of the specific species. If you provide a vector of different length, only the first element will be used. The string should also contain the species of interest, for example, valid ".gene" arguments are "hs.trbv", "HomoSapiens.TRBJ" or "macmul.IGHV". For details run <code>gene_stats()</code>.</p> |
| <code>.quant</code> | <p>Selects the column with data to evaluate. Pass NA if you want to compute gene statistics at the clonotype level without re-weighting. Pass "count" to use the "Clones" column to weight genes by abundance of their corresponding clonotypes.</p> |
| <code>.ambig</code> | <p>An option to handle ambiguous gene assignments, e.g., "TRAV1,TRAV2".</p> <ul style="list-style-type: none"> • Pass "inc" to include all possible gene segments, so "TRAV1,TRAV2" is counted as a different gene segment. • Pass "exc" to exclude all ambiguous gene assignments, so "TRAV1,TRAV2" is excluded from the resultant gene table. |

We recommend to turn it on by passing "inc" (turned on by default). You can exclude data for the cases where there is no clear match for gene, include it for every supplied gene, or pick only first from the set. Set it to "exc", "inc" or "maj", respectively.

.type Set the type of data to evaluate: "segment", "allele", or "family".
 .norm If TRUE then return proportions of genes. If FALSE then return counts of genes.

Value

A data frame with rows corresponding to gene segments and columns corresponding to the input samples.

Examples

```
data(immdata)
gu <- geneUsage(immdata$data)
vis(gu)
```

geneUsageAnalysis	<i>Post-analysis of V-gene and J-gene statistics: PCA, clustering, etc.</i>
-------------------	---

Description

[Deprecated]

The `geneUsageAnalysis()` function deploys several data analysis methods, including PCA, multidimensional scaling, Jensen-Shannon divergence, k-means, hierarchical clustering, DBscan, and different correlation coefficients.

Usage

```
geneUsageAnalysis(
  .data,
  .method = c("js+hclust", "pca+kmeans", "anova", "js+pca+kmeans"),
  .base = 2,
  .norm.entropy = FALSE,
  .cor = c("pearson", "kendall", "spearman"),
  .do.norm = TRUE,
  .laplace = 1e-12,
  .verbose = TRUE,
  .k = 2,
  .eps = 0.01,
  .perp = 1,
  .theta = 0.1
)
```

Arguments

<code>.data</code>	The <code>geneUsageAnalysis()</code> function runs on the output from <code>geneUsage()</code> .
<code>.method</code>	A string that defines the type of analysis to perform. Can be "pca", "mds", "js", "kmeans", "hclust", "dbscan" or "cor" if you want to calculate correlation coefficient. In the latter case you have to provide <code>.cor</code> argument.
<code>.base</code>	A numerical value that defines the logarithm base for Jensen-Shannon divergence.
<code>.norm.entropy</code>	A logical value. Set TRUE to normalise your data if you haven't done it already.
<code>.cor</code>	A string that defines the correlation coefficient for analysis. Can be "pearson", "kendall" or "spearman".
<code>.do.norm</code>	A logical value. If TRUE it forces Laplace smoothing, if NA it checks if smoothing is necessary, if FALSE does nothing.
<code>.laplace</code>	The numeric value, which is used as a pseudocount for Laplace smoothing.
<code>.verbose</code>	A logical value.
<code>.k</code>	The number of clusters to create, passed as <code>k</code> to <code>hcut</code> or as <code>centers</code> to <code>kmeans</code> .
<code>.eps</code>	A numerical value, DBscan epsilon parameter, see <code>immunr_dbscan()</code> .
<code>.perp</code>	A numerical value, t-SNE perplexity, see <code>immunr_tsne()</code> .
<code>.theta</code>	A numerical value, t-SNE theta parameter, see <code>immunr_tsne()</code> .

Value

Depends on the last element in the `.method` string. See `immunr_tsne` for more info.

Examples

```
data(immdata)
gu <- geneUsage(immdata$data, .norm = TRUE)
geneUsageAnalysis(gu, "js+hclust", .verbose = FALSE) %>% vis()
```

gene_segments

Gene segments table

Description

Gene segments table

gene_stats	WIP
Description	
WIP	
Usage	
gene_stats()	
Value	
gene_stats returns all segment gene statistics	
Examples	
gene_stats() get_genes("hs.trbv", "segment")	
getKmers	Calculate the k-mer statistics of immune repertoires

Description

[Deprecated]

Usage

getKmers(.data, .k, .col = c("aa", "nt"), .coding = TRUE)

Arguments

- | | |
|---------|--|
| .data | The data to be processed. Can be data.frame , data.table::data.table , or a list of these objects.
Every object must have columns in the immunarch compatible format. immunarch_data_format
Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects.
Note: each connection must represent a separate repertoire. |
| .k | Integer. Length of k-mers. |
| .col | Character. Which column to use, pass "aa" (by default) for CDR3 amino acid sequence, pass "nt" for CDR3 nucleotide sequences. |
| .coding | Logical. If TRUE (by default) then removes all non-coding sequences from input data first. |

Value

Data frame with two columns (k-mers and their counts).

Examples

```
data(immdata)
kmers <- getKmers(immdata$data[[1]], 5)
kmers %>% vis()
```

get_immunarch_news	<i>Get the Latest immunarch Update</i>
--------------------	--

Description

Retrieves an update message for immunarch.

Usage

```
get_immunarch_news(datepoint = "latest")
```

Arguments

datepoint	A string specifying the update date. Use "latest" for the most recent update or supply a valid date key (e.g., "Apr 2025").
-----------	---

Details

If datepoint is set to "latest", the function returns the most recent update. Otherwise, specify the update date key (e.g., "Apr 2025") to retrieve that particular update. If no matching update is found, a warning is issued along with available update keys.

Value

A character string with the update details or a warning if the key is not found.

See Also

[list_immunarch_news\(\)](#)

group_from_metadata	<i>Get a character vector of samples' groups from the input metadata file</i>
---------------------	---

Description

Get a character vector of samples' groups from the input metadata file

Usage

```
group_from_metadata(.by, .metadata, .sep = "; ")
```

Arguments

.by	Character vector. Specify a column or columns in the input metadata to group by.
.metadata	Metadata object.
.sep	Character vector. Defines a separator between groups if more than one group passed in .by.

Value

Character vector with group names.

Developer Examples

```
immunarch:::group_from_metadata("Status", data.frame(Status = c("A", "A", "B", "B", "C")))
```

has_class	<i>Check for the specific class</i>
-----------	-------------------------------------

Description

A function to check if an input object has a specific class name.

Usage

```
has_class(.data, .class)
```

Arguments

.data	Any R object.
.class	Character vector. Specifies a class name to check against.

Value

Logical value.

Developer Examples

```
tmp <- "abc" immunarch:::has_class(tmp, "new_class") tmp <- immunarch:::add_class(tmp, "new_class")
immunarch:::has_class(tmp, "new_class")
```

immdata	<i>Single chain immune repertoire dataset</i>
---------	---

Description

A dataset with single chain TCR data for testing and exemplary purposes.

Usage

```
immdata
```

Format

A list of two elements. The first element ("data") is a list with data frames with clonotype tables. The second element ("meta") is a metadata table.

- data** List of immune repertoire data frames.
- meta** Metadata ...

immunarch_v1_updates	<i>Get a list of package updates</i>
----------------------	--------------------------------------

Description

Get a list of package updates

Usage

```
immunarch_v1_updates
```

Format

An object of class list of length 1.

immunr_data_format

*Specification of the data format used by immunarch dataframes***Description**

- "Clones" - number of barcodes (events, UMIs) or reads;
- "Proportion" - proportion of barcodes (events, UMIs) or reads;
- "CDR3.nt" - CDR3 nucleotide sequence;
- "CDR3.aa" - CDR3 amino acid sequence;
- "V.name" - names of aligned Variable gene segments;
- "D.name" - names of aligned Diversity gene segments or NA;
- "J.name" - names of aligned Joining gene segments;
- "V.end" - last positions of aligned V gene segments (1-based);
- "D.start" - positions of D'5 end of aligned D gene segments (1-based);
- "D.end" - positions of D'3 end of aligned D gene segments (1-based);
- "J.start" - first positions of aligned J gene segments (1-based);
- "VJ.ins" - number of inserted nucleotides (N-nucleotides) at V-J junction (-1 for receptors with VDJ recombination);
- "VD.ins" - number of inserted nucleotides (N-nucleotides) at V-D junction (-1 for receptors with VJ recombination);
- "DJ.ins" - number of inserted nucleotides (N-nucleotides) at D-J junction (-1 for receptors with VJ recombination);
- "Sequence" - full nucleotide sequence.

immunr_hclust

*Clustering of objects or distance matrices***Description****[Deprecated]**

Clusters the data with one of the following methods:

- immunr_hclust clusters the data using the hierarchical clustering from [hcut](#);
- immunr_kmeans clusters the data using the K-means algorithm from [kmeans](#);
- immunr_dbscan clusters the data using the DBSCAN algorithm from [dbscan](#).

Usage

```
immunr_hclust(.data, .k = 2, .k.max = nrow(.data) - 1, .method = "complete", .dist = TRUE)

immunr_kmeans(.data, .k = 2, .k.max = as.integer(sqrt(nrow(.data))) + 1,
  .method = c("silhouette", "gap_stat"))

immunr_dbscan(.data, .eps, .dist = TRUE)
```

Arguments

<code>.data</code>	Matrix or data frame with features, distance matrix or output from repOverlapAnalysis or geneUsageAnalysis functions.
<code>.k</code>	The number of clusters to create, defined as <code>k</code> to hcut or as centers to kmeans .
<code>.k.max</code>	Limits the maximum number of clusters. It is passed as <code>k.max</code> to factoextra::fviz_nbclust for <code>immunr_hclust</code> and <code>immunr_kmeans</code> .
<code>.method</code>	Passed to factoextra::hcut or as factoextra::fviz_nbclust . In case of factoextra::hcut the agglomeration method is going to be used (argument <code>hc_method</code>). In case of factoextra::fviz_nbclust it is the method to be used for estimating the optimal number of clusters (argument <code>method</code>).
<code>.dist</code>	If <code>TRUE</code> then <code>".data"</code> is expected to be a distance matrix. If <code>FALSE</code> then the euclidean distance is computed for the input objects.
<code>.eps</code>	Local radius for expanding clusters, minimal distance between points to expand clusters. Passed as <code>eps</code> to dbscan .

Value

`immunr_hclust` - list with two elements. The first element is an output from [factoextra::hcut](#). The second element is an output from [factoextra::fviz_nbclust](#)

`immunr_kmeans` - list with three elements. The first element is an output from [kmeans](#). The second element is an output from [factoextra::fviz_nbclust](#). The third element is the input dataset `.data`.

`immunr_dbscan` - list with two elements. The first element is an output from [fpc::dbscan](#). The second element is the input dataset `.data`.

Examples

```
data(immdata)
gu <- geneUsage(immdata$data, .norm = TRUE)
immunr_hclust(t(as.matrix(gu[, -1])), .dist = FALSE)

gu[is.na(gu)] <- 0
immunr_kmeans(t(as.matrix(gu[, -1])))
```

immunr_pca

*Dimensionality reduction***Description****[Deprecated]**

Collects a set of principal variables, reducing the number of not important variables to analyse. Dimensionality reduction makes data analysis algorithms work faster and sometimes more accurate, since it also reduces noise in the data. Currently available methods are:

- immunr_pca performs PCA (Principal Component Analysis) using [prcomp](#);
- immunr_mds performs MDS (Multi-Dimensional Scaling) using isoMODS from MASS package.
- immunr_tsne performs tSNE (t-Distributed Stochastic Neighbour Embedding) using Rtsne Rtsne package.

Usage

```
immunr_pca(.data, .scale = default_scale_fun, .raw = TRUE, .orig = FALSE, .dist = FALSE)
```

```
immunr_mds(.data, .scale = default_scale_fun, .raw = TRUE, .orig = FALSE, .dist = TRUE)
```

```
immunr_tsne(.data, .perp = 1, .dist = TRUE, ...)
```

Arguments

<code>.data</code>	A matrix or a data frame with features, distance matrix or output from repOverlapAnalysis or geneUsageAnalysis functions.
<code>.scale</code>	A function to apply to your data before passing it to any of dimensionality reduction algorithms. There is no scaling by default.
<code>.raw</code>	If TRUE then returns the non-processed output from dimensionality reduction algorithms. Pass FALSE if you want to visualise results.
<code>.orig</code>	If TRUE then returns the original result from algorithms. Pass FALSE if you want to visualise results.
<code>.dist</code>	If TRUE then assumes that ".data" is a distance matrix.
<code>.perp</code>	The perplexity parameter for Rtsne. Specifies the number of neighbors each data point must have in the resulting plot.
<code>...</code>	Other parameters passed to Rtsne.

Value

immunr_pca - an output from [prcomp](#).

immunr_mds - an output from isoMDS.

immunr_tsne - an output from Rtsne.

See Also

[vis.immunr_pca](#) for visualisations.

Examples

```
data(immdata)
gu <- geneUsage(immdata$data)
gu[is.na(gu)] <- 0
gu <- t(as.matrix(gu[, -1]))
immunr_pca(gu)
immunr_mds(dist(gu))
immunr_tsne(dist(gu))
```

inc_overlap

Incremental counting of repertoire similarity

Description**[Deprecated]**

For reference please look up <https://www.pnas.org/content/111/16/5980> (Fig. 4).

Usage

```
inc_overlap(
  .data,
  .fun,
  .step = 1000,
  .n.steps = 10,
  .downsample = FALSE,
  .bootstrap = NA,
  .verbose.inc = TRUE,
  ...
)
```

Arguments

.data	The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects. Every object must have columns in the immunarch compatible format. immunarch_data_format Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects. Note: each connection must represent a separate repertoire.
.fun	Function to compute overlaps. e.g., <code>morisita_index</code>.

<code>.step</code>	Either an integer or a numeric vector. In the first case, the integer defines the step of incremental overlap. In the second case, the vector encodes all repertoire sampling depths.
<code>.n.steps</code>	Integer. Number of steps if <code>.step</code> is a single integer. Skipped if <code>".step"</code> is a numeric vector.
<code>.downsample</code>	If TRUE then performs downsampling to N clonotypes at each step instead of choosing the top N clonotypes.
<code>.bootstrap</code>	Set NA to turn off any bootstrapping, set a number to perform bootstrapping with this number of tries.
<code>.verbose.inc</code>	Logical. If TRUE then shows the output from the computation process.
<code>...</code>	Other arguments passed to <code>.fun</code> .

Value

List with overlap matrices.

Examples

```
## Not run:
data(immdata)
ov <- repOverlap(immdata$data, "inc+overlap", .step = 100, .verbose.inc = FALSE, .verbose = FALSE)
vis(ov)

## End(Not run)
```

list_immunarch_news *List Available immunarch Updates*

Description

Returns the list of available update keys for immunarch v1.

Usage

```
list_immunarch_news()
```

Value

A character vector containing all the date keys for the available updates.

See Also

[get_immunarch_news\(\)](#)

matrixdiagcopy	<i>Copy the upper matrix triangle to the lower one</i>
----------------	--

Description

Copy the upper matrix triangle to the lower one

Usage

```
matrixdiagcopy(.mat)
```

Arguments

.mat Matrix.

Value

Matrix with its upper tri part copied to the lower tri part.

Developer Examples

```
mat <- matrix(0, 3, 3) mat[1, 3] <- 1 mat <- immunarch::matrixdiagcopy(mat) mat
```

public_matrix	<i>Get a matrix with public clonotype frequencies</i>
---------------	---

Description

[Deprecated]

Usage

```
public_matrix(.data)
```

Arguments

.data Public repertoire, an output from [pubRep](#).

Value

Matrix with per-sample clonotype counts / proportions only.

Examples

```
data(immdata)
immdata$data <- lapply(immdata$data, head, 2000)
pr <- pubRep(immdata$data, .verbose = FALSE)
pr.mat <- public_matrix(pr)
dim(pr.mat)
head(pr.mat)
```

pubRep

Create a repertoire of public clonotypes

Description

[Deprecated]

Usage

```
pubRep(
  .data,
  .col = "aa+v",
  .quant = c("count", "prop"),
  .coding = TRUE,
  .min.samples = 1,
  .max.samples = NA,
  .verbose = TRUE
)
```

Arguments

- | | |
|---------|---|
| .data | <p>The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects.</p> <p>Every object must have columns in the immunarch compatible format. immunarch_data_format</p> <p>Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects.</p> <p>Note: each connection must represent a separate repertoire.</p> |
| .col | <p>A string that specifies the column(s) to be processed. Outputs one of the following strings, separated by the plus sign: "nt" for nucleotide sequences, "aa" for amino acid sequences, "v" for V gene segments, "j" for J gene segments. E.g., pass "aa+v" to compute overlaps on CDR3 amino acid sequences paired with V gene segments, i.e., in this case a unique clonotype is a pair of CDR3 amino acid and V gene segment.</p> |
| .quant | <p>A string that specifies the column to be processed. Set "count" to see public clonotype sharing with the number of clones, set "prop" to see proportions.</p> |
| .coding | <p>Logical. If TRUE then preprocesses the data to filter out non-coding sequences.</p> |

<code>.min.samples</code>	Integer. A minimal number of samples a clonotype must have to be included in the public repertoire table.
<code>.max.samples</code>	Integer. A maximal number of samples a clonotype must have to be included in the public repertoire table. Set NA (by default) to have the maximal amount of samples.
<code>.verbose</code>	Logical. If TRUE then outputs the progress.

Value

Data table with columns for:

- Clonotypes (e.g., CDR3 sequence, or two columns for CDR3 sequence and V gene)
- Incidence of clonotypes
- Per-sample proportions or counts

Examples

```
# Subset the data to make the example faster to run
immdata$data <- lapply(immdata$data, head, 2000)
pr <- pubRep(immdata$data, .verbose = FALSE)
vis(pr, "clonotypes", 1, 2)
```

pubRepApply

Apply transformations to public repertoires

Description

[Deprecated]

Usage

```
pubRepApply(.pr1, .pr2, .fun = function(x) log10(x[1])/log10(x[2]))
```

Arguments

<code>.pr1</code>	First public repertoire.
<code>.pr2</code>	Second public repertoire.
<code>.fun</code>	A function to apply to pairs of frequencies of same clonotypes from "pr1" and "pr2". By default - $\log(X) / \log(Y)$ where X, Y - frequencies of the same clonotype, found in both public repertoires.

Value

Work in progress.

Examples

```
data(immdata)
immdata$data <- lapply(immdata$data, head, 2000)
pr <- pubRep(immdata$data, .verbose = FALSE)
pr1 <- pubRepFilter(pr, immdata$meta, .by = c(Status = "MS"))
pr2 <- pubRepFilter(pr, immdata$meta, .by = c(Status = "C"))
prapp <- pubRepApply(pr1, pr2)
head(prapp)
```

pubRepFilter

*Filter out clonotypes from public repertoires***Description****[Deprecated]**

Filter our clonotypes with low incidence in a specific group.

Usage

```
pubRepFilter(.pr, .meta, .by, .min.samples = 1)
```

Arguments

.pr	Public repertoires, an output from pubRep .
.meta	Metadata file.
.by	Named character vector. Names of the group to filter by.
.min.samples	Integer. Filters out clonotypes with the number of samples below than this number.

Value

Data frame with filtered clonotypes.

Examples

```
data(immdata)
immdata$data <- lapply(immdata$data, head, 2000)
pr <- pubRep(immdata$data, .verbose = FALSE)
pr1 <- pubRepFilter(pr, immdata$meta, .by = c(Status = "MS"))
head(pr1)
```

pubRepStatistics	<i>Statistics of number of public clonotypes for each possible combinations of repertoires</i>
------------------	--

Description**[Deprecated]****Usage**

```
pubRepStatistics(.data, .by = NA, .meta = NA)
```

Arguments

.data	Public repertoire, an output from the pubRep function.
.by	Work in Progress.
.meta	Work in Progress.

Value

Data frame with incidence statistics per sample.

Examples

```
data(immdata)
immdata$data <- lapply(immdata$data, head, 2000)
pr <- pubRep(immdata$data, .verbose = FALSE)
pubRepStatistics(pr) %>% vis()
```

repAlignLineage	<i>Aligns all sequences including germline within each clonal lineage within each cluster</i>
-----------------	---

Description**[Deprecated]**

This function aligns all sequences (including germline) that belong to one clonal lineage and one cluster. After clustering and building the clonal lineage and germline, the next step is to analyze the degree of mutation and maturity of each clonal lineage. This allows for finding high mature cells and cells with a large number of offspring. The phylogenetic analysis will find mutations that increase the affinity of BCR. Making alignment of the sequence is the first step towards sequence analysis including BCR.

Usage

```
repAlignLineage(.data, .min_lineage_sequences, .prepare_threads, .align_threads, .nofail)
```

Arguments

- `.data` The data to be processed. Can be [data.frame](#), [data.table::data.table](#) or a list of these objects.
- `.min_lineage_sequences` If number of sequences in the same clonal lineage and the same cluster (not including germline) is lower than this threshold, this group of sequences will be filtered out from the dataframe; so only large enough lineages will be included.
- `.prepare_threads` Number of threads to prepare results table. Please note that high number can cause heavy memory usage!
- `.align_threads` Number of threads for lineage alignment.
It must have columns in the immunarch compatible format [immunarch_data_format](#), and also must contain 'Cluster' column, which is added by `seqCluster()` function, and 'Germline.sequence' column, which is added by `repGermline()` function.
- `.nofail` Will return NA instead of stopping if Clustal W is not installed. Used to avoid raising errors in examples on computers where Clustal W is not installed.

Value

Dataframe or list of dataframes (if input is a list with multiple samples). The dataframe has these columns:

- Cluster: cluster name
- Germline: germline sequence
- Alignment: DNABin object with alignment
- Sequences: nested dataframe containing all sequences for this combination of cluster and germline; it has columns
 - Sequence, CDR1.nt, CDR2.nt, CDR3.nt, FR1.nt, FR2.nt, FR3.nt, FR4.nt, V.allele, J.allele, V.aa, J.aa: all values taken from the input dataframe
 - Clone.ID: taken from the input dataframe, or created (filled with row numbers) if missing
 - Clones: taken from the input dataframe, or created (filled with '1' values) if missing

Examples

```
data(bcrdata)
bcr_data <- bcrdata$data

bcr_data %>%
  seqCluster(seqDist(bcr_data), .fixed_threshold = 3) %>%
  repGermline(.threads = 1) %>%
  repAlignLineage(.min_lineage_sequences = 2, .align_threads = 2, .nofail = TRUE)
```

repClonalFamily	<i>Builds a phylogenetic tree using the sequences of a clonal lineage</i>
-----------------	---

Description

[Deprecated]

This function uses the PHYLIP package to make phylogenetic analysis. For making trees it uses maximum parsimony methods.

Usage

```
repClonalFamily(.data, .vis_groups, .threads, .nofail)
```

Arguments

<code>.data</code>	The data to be processed, output of repAlignLineage() function.
<code>.vis_groups</code>	Groups for visualization, used to annotate specific clones on chart and display them in different colors. This is a named list, where names are for the chart legend, and list items are clone IDs that belong to the groups. It's not necessary to assign groups to all clonotypes; unassigned ones will be displayed on the chart as "Clonotype" category. It's also possible to assign multiple clonotypes to the same group by providing nested lists or vectors of clone IDs instead of single clone IDs. Example: <code>.vis_groups = list(A = 817, B = 201, C = list(303, 42))</code>
<code>.threads</code>	Number of threads to use.
<code>.nofail</code>	Returns NA instead of stopping if PHYLIP is not installed. Used to avoid raising errors in examples on computers where PHYLIP is not installed.

Value

Dataframe or list of dataframes (if input is a list with multiple samples). The dataframe has these columns:

- Cluster: cluster name
- Germline.Input: germline sequence, like it was in the input; not aligned
- Germline.Output: germline sequence, parsed from PHYLIP dnapars function output; it contains difference of germline from the common ancestor; "." characters mean matching letters
- Common.Ancestor: common ancestor sequence, parsed from PHYLIP dnapars function output
- Trunk.Length: mean trunk length, representing the distance between the most recent common ancestor and germline sequence as a measure of the maturity of a lineage
- Tree: output tree in "phylo" format, loaded from by PHYLIP dnapars function output
- TreeStats: nested dataframe containing data about tree nodes, needed for visualization
- Sequences: nested dataframe containing all sequences for this combination of cluster and germline; it contains regions from original sequences, saved for repSomaticHypermutation() calculation, and also data needed for visualizations

Examples

```
data(bcrdata)
bcr_data <- bcrdata$data

bcr_data %>%
  seqCluster(seqDist(bcr_data), .fixed_threshold = 3) %>%
  repGermline(.threads = 1) %>%
  repAlignLineage(.min_lineage_sequences = 2, .align_threads = 2, .nofail = TRUE) %>%
  repClonalFamily(.threads = 1, .nofail = TRUE)
```

 repClonality

Clonality analysis of immune repertoires

Description

[Deprecated]

repClonality function encompasses several methods to measure clonal proportions in a given repertoire.

Usage

```
repClonality(
  .data,
  .method = c("clonal.prop", "homeo", "top", "rare"),
  .perc = 10,
  .clone.types = c(Rare = 1e-05, Small = 1e-04, Medium = 0.001, Large = 0.01,
    Hyperexpanded = 1),
  .head = c(10, 100, 1000, 3000, 10000, 30000, 1e+05),
  .bound = c(1, 3, 10, 30, 100)
)
```

Arguments

.data	The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects. Every object must have columns in the immunarch compatible format. immunarch_data_format Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects. Note: each connection must represent a separate repertoire.
.method	A String with one of the following options: "clonal.prop", "homeo", "top" or "rare". Set "clonal.prop" to compute clonal proportions or in other words percentage of clonotypes required to occupy specified by .perc percent of the total immune repertoire.

Set "homeo" to analyse relative abundance (also known as clonal space homeostasis), which is defined as the proportion of repertoire occupied by clonal groups with specific abundances..

Set "top" to estimate relative abundance for the groups of top clonotypes in repertoire, e.g., ten most abundant clonotypes. Use ".head" to define index intervals, such as 10, 100 and so on.

Set "rare" to estimate relative abundance for the groups of rare clonotypes with low counts. Use ".bound" to define the threshold of clonotype groups.

.perc	A single numerical value ranging from 0 to 100.
.clone.types	A named numerical vector with the threshold of the half-closed intervals that mark off clonal groups.
.head	A numerical vector with ranges of the top clonotypes.
.bound	A numerical vector with ranges of abundance for the rare clonotypes in the dataset.

Details

Clonal proportion assessment is a different approach to estimate repertoire diversity. When visualised, it allows for thorough examination of immune repertoire structure and composition.

In its core this type of analysis is similar to the relative species abundance concept in ecology. Relative abundance is the percent composition of an organism of a particular kind relative to the total number of organisms in the area.

A stacked barplot of relative clonotype abundances can be therefore viewed as a non-parametric approach to comparing their underlying distributions.

Value

If input data is a single immune repertoire, then the function returns a numeric vector with clonality statistics.

Otherwise, it returns a numeric matrix with clonality statistics for all input repertoires.

See Also

[repDiversity](#)

Examples

```
# Load the data
data(immdata)

imm_pr <- repClonality(immdata$data, .method = "clonal.prop")
vis(imm_pr)

imm_top <- repClonality(immdata$data, .method = "top", .head = c(10, 100, 1000, 3000, 10000))
vis(imm_top)

imm_rare <- repClonality(immdata$data, .method = "rare")
vis(imm_rare)
```

```
imm_hom <- repClonality(immdata$data, .method = "homeo")
vis(imm_hom)
```

repDiversity

The main function for immune repertoire diversity estimation

Description

[Deprecated]

This is a utility function to estimate the diversity of species or objects in the given distribution.

Note: functions will check if .data is a distribution of a random variable (sum == 1) or not. To force normalisation and / or to prevent this, set .do.norm to TRUE (do normalisation) or FALSE (don't do normalisation), respectively.

Usage

```
repDiversity(
  .data,
  .method = "chao1",
  .col = "aa",
  .max.q = 6,
  .min.q = 1,
  .q = 5,
  .step = NA,
  .quantile = c(0.025, 0.975),
  .extrapolation = NA,
  .perc = 50,
  .norm = TRUE,
  .verbose = TRUE,
  .do.norm = NA,
  .laplace = 0
)
```

Arguments

- | | |
|---------|---|
| .data | <p>The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects.</p> <p>Every object must have columns in the immunarch compatible format. immunarch_data_format</p> <p>Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects.</p> <p>Note: each connection must represent a separate repertoire.</p> |
| .method | <p>Picks a method used for estimation out of a following list: chao1, hill, div, gini.simp, inv.simp, gini, raref, d50, dxx.</p> |

<code>.col</code>	A string that specifies the column(s) to be processed. Pass one of the following strings, separated by the plus sign: "nt" for nucleotide sequences, "aa" for amino acid sequences, "v" for V gene segments, "j" for J gene segments. E.g., pass "aa+v" to compute diversity estimations on CDR3 amino acid sequences paired with V gene segments, i.e., in this case a unique clonotype is a pair of CDR3 amino acid and V gene segment. Clonal counts of equal clonotypes will be summed up.
<code>.max.q</code>	The max hill number to calculate (default: 5).
<code>.min.q</code>	Function calculates several hill numbers. Set the min (default: 1).
<code>.q</code>	q-parameter for the Diversity index.
<code>.step</code>	Rarefaction step's size.
<code>.quantile</code>	Numeric vector with quantiles for confidence intervals.
<code>.extrapolation</code>	An integer. An upper limit for the number of clones to extrapolate to. Pass 0 (zero) to turn extrapolation subroutines off.
<code>.perc</code>	Set the percent to dXX index measurement.
<code>.norm</code>	Normalises rarefaction curves.
<code>.verbose</code>	If TRUE then outputs progress.
<code>.do.norm</code>	One of the three values - NA, TRUE or FALSE. If NA then checks for distribution ($\text{sum}(.data) == 1$) and normalises if needed with the given laplace correction value. if TRUE then does normalisation and laplace correction. If FALSE then doesn't do neither normalisation nor laplace correction.
<code>.laplace</code>	A numeric value, which is used as a pseudocount for Laplace smoothing.

Details

- True diversity, or the effective number of types, refers to the number of equally-abundant types needed for the average proportional abundance of the types to equal that observed in the dataset of interest where all types may not be equally abundant.
- Inverse Simpson index is the effective number of types that is obtained when the weighted arithmetic mean is used to quantify average proportional abundance of types in the dataset of interest.
- The Gini coefficient measures the inequality among values of a frequency distribution (for example levels of income). A Gini coefficient of zero expresses perfect equality, where all values are the same (for example, where everyone has the same income). A Gini coefficient of one (or 100 percents) expresses maximal inequality among values (for example where only one person has all the income).
- The Gini-Simpson index is the probability of interspecific encounter, i.e., probability that two entities represent different types.
- Chao1 estimator is a nonparameteric asymptotic estimator of species richness (number of species in a population).
- Rarefaction is a technique to assess species richness from the results of sampling through extrapolation.
- Hill numbers are a mathematically unified family of diversity indices (differing among themselves only by an exponent q).

- d50 is a recently developed immune diversity estimate. It calculates the minimum number of distinct clonotypes amounting to greater than or equal to 50 percent of a total of sequencing reads obtained following amplification and sequencing
- dXX is a similar to d50 index where XX corresponds to desirable percent of total sequencing reads.

Value

div, gini, gini.simp, inv.simp, raref return numeric vector of length 1 with value.

chao1 returns 4 values: estimated number of species, standart deviation of this number and two 95% confidence intervals for the species number.

hill returns a vector of specified length .max.q - .min.q

For most methods, if input data is a single immune repertoire, then the function returns a numeric vector with diversity statistics.

Otherwise, it returns a numeric matrix with diversity statistics for all input repertoires.

For Chao1 the function returns a matrix with diversity estimations.

For rarefaction the function returns either a matrix with diversity estimatinos on different step of the simulaiton process or a list with such matrices.

See Also

[repOverlap](#), [entropy](#), [repClonality](#) Rarefaction wiki [https://en.wikipedia.org/wiki/Rarefaction_\(ecology\)](https://en.wikipedia.org/wiki/Rarefaction_(ecology)) Hill numbers paper <https://www.uvm.edu/~ngotelli/manuscriptpdfs/ChaoHill.pdf> Diversity wiki https://en.wikipedia.org/wiki/Measurement_of_biodiversity

Examples

```
data(immdata)

# Make data smaller for testing purposes
immdata$data <- top(immdata$data, 4000)

# chao1
repDiversity(.data = immdata$data, .method = "chao1") %>% vis()

# Hill numbers
repDiversity(
  .data = immdata$data, .method = "hill", .max.q = 6,
  .min.q = 1, .do.norm = NA, .laplace = 0
) %>% vis()

# diversity
repDiversity(.data = immdata$data, .method = "div", .q = 5, .do.norm = NA, .laplace = 0) %>%
  vis()

# Gini-Simpson
repDiversity(.data = immdata$data, .method = "gini.simp", .q = 5, .do.norm = NA, .laplace = 0) %>%
  vis()
```

```
# inverse Simpson
repDiversity(.data = immdata$data, .method = "inv.simp", .do.norm = NA, .laplace = 0) %>% vis()

# Gini coefficient
repDiversity(.data = immdata$data, .method = "gini", .do.norm = NA, .laplace = 0)

# d50
repDiversity(.data = immdata$data, .method = "d50") %>% vis()
```

repExplore	<i>Main function for exploratory data analysis: compute the distribution of lengths, clones, etc.</i>
------------	---

Description

[Deprecated]

The repExplore function calculates the basic statistics of repertoire: the number of unique immune receptor clonotypes, their relative abundances, and sequence length distribution across the input dataset.

Usage

```
repExplore(
  .data,
  .method = c("volume", "count", "len", "clones"),
  .col = c("nt", "aa"),
  .coding = TRUE
)
```

Arguments

.data	The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects. Every object must have columns in the immunarch compatible format. immunarch_data_format Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects. Note: each connection must represent a separate repertoire.
.method	A string that specifies the method of analysis. It can be either "volume", "count", "len" or "clones". When .method is set to "volume" the repExplore calculates the number of unique clonotypes in the input data. When .method is set to "count" the repExplore calculates the distribution of clonotype abundances, i.e., how frequent receptors with different abundances are.

	When <code>.method</code> is set to "len" the <code>repExplore</code> calculates the distribution of CDR3 sequence lengths.
	When <code>.method</code> is set to "clones" the <code>repExplore</code> returns the number of clones (i.e., cells) per input repertoire.
<code>.col</code>	A string that specifies the column to be processed. Pass "nt" for nucleotide sequence or "aa" for amino acid sequence.
<code>.coding</code>	If TRUE, then only coding sequences will be analysed.

Value

If input data is a single immune repertoire, then the function returns a numeric vector with exploratory analysis statistics.

Otherwise, it returns a numeric matrix with exploratory analysis statistics for all input repertoires.

See Also

[vis.immunr_exp_vol](#)

Examples

```
data(immdata)

# Calculate statistics and generate a visual output with vis()
repExplore(immdata$data, .method = "volume") %>% vis()

repExplore(immdata$data, .method = "count") %>% vis()

repExplore(immdata$data, .method = "len") %>% vis()
```

repFilter

Main function for data filtering

Description

[Deprecated]

Usage

```
repFilter(
  .data,
  .method = "by.clonotype",
  .query = list(CDR3.aa = exclude("partial", "out_of_frame")),
  .match = "exact"
)
```

Arguments

<code>.data</code>	The data to be processed. Must be the list of 2 elements: a data table and a metadata table.
<code>.method</code>	Method of filtering. Implemented methods: <code>by.meta</code> , <code>by.repertoire</code> (<code>by.rep</code>), <code>by.clonotype</code> (<code>by.cl</code>) Default value: <code>'by.clonotype'</code> .
<code>.query</code>	Filtering query. It's a named list of filters that will be applied to data. Possible values for names in this list are dependent on filter methods: <code>by.meta</code>: filters by metadata. Names in the named list are metadata column headers. <code>by.repertoire</code>: filters by the number of clonotypes or total number of clones in sample. Possible names in the named list are <code>"n_clonotypes"</code> and <code>"n_clones"</code>. <code>by.clonotype</code>: filters by data in all samples. Names in the named list are data column headers. Elements of the named list for each of the filters are filtering options. Possible values for filtering options: <code>include("STR1", "STR2", ...)</code>: keeps only rows with matching values. Available for methods: <code>"by.meta"</code>, <code>"by.clonotype"</code>. <code>exclude("STR1", "STR2", ...)</code>: removes rows with matching values. Available for methods: <code>"by.meta"</code>, <code>"by.clonotype"</code>. <code>lessthan(value)</code>: keeps rows/samples with numeric values less than specified. Available for methods: <code>"by.meta"</code>, <code>"by.repertoire"</code>, <code>"by.clonotype"</code>. <code>morethan(value)</code>: keeps rows/samples with numeric values more than specified. Available for methods: <code>"by.meta"</code>, <code>"by.repertoire"</code>, <code>"by.clonotype"</code>. <code>interval(from, to)</code>: keeps rows/samples with numeric values that fits in this interval. <code>from</code> is inclusive, <code>to</code> is exclusive. Available for methods: <code>"by.meta"</code>, <code>"by.repertoire"</code>, <code>"by.clonotype"</code>. Default value: <code>'list(CDR3.aa = exclude("partial", "out_of_frame"))'</code>.
<code>.match</code>	Matching method for <code>"include"</code> and <code>"exclude"</code> options in query. Possible values: <code>exact</code>: matches only the exact specified string; <code>startswith</code>: matches all strings starting with the specified substring; <code>substring</code>: matches all strings containing the specified substring. Default value: <code>'exact'</code>.

repGermline

*Creates germlines for clonal lineages***Description****[Deprecated]**

This function creates germlines for clonal lineages. B cell clonal lineage represents a set of B cells that presumably have a common origin (arising from the same VDJ rearrangement event) and a common ancestor. Each clonal lineage has its own germline sequence that represents the ancestral sequence for each BCR in clonal lineage. In other words, germline sequence is a sequence of B-cells immediately after VDJ recombination, before B-cell maturation and hypermutation process. Germline sequence is useful for assessing the degree of mutation and maturity of the repertoire.

Usage

```
repGermline(.data, .species, .min_nuc_outside_cdr3, .threads)
```

Arguments

<code>.data</code>	The data to be processed. Can be data.frame , data.table::data.table or a list of these objects. It must have columns in the immunarch compatible format immunarch_data_format .
<code>.species</code>	Species from which the data was acquired. Available options: "HomoSapiens" (default), "MusMusculus", "BosTaurus", "CamelusDromedarius", "CanisLupusFamiliaris", "DanioRerio", "MacacaMulatta", "MusMusculusDomesticus", "MusMusculusCastaneus", "MusMusculusMolossinus", "MusMusculusMusculus", "MusSpretus", "OncorhynchusMykiss", "OrnithorhynchusAnatinus", "OryctolagusCuniculus", "RattusNorvegicus", "SusScrofa".
<code>.min_nuc_outside_cdr3</code>	This parameter sets how many nucleotides should have V or J chain outside of CDR3 to be considered good for further alignment.
<code>.threads</code>	Number of threads to use.

Value

Data with added columns:

- Sequence (FR1+CDR1+FR2+CDR2+FR3+CDR3+FR4 in nucleotides; the column will be replaced if exists)
- V.allele, J.allele (chosen alleles of V and J genes),
- V.aa, J.aa (V and J sequences from original clonotype, outside CDR3, converted to amino acids)
- Germline.sequence (combined germline nucleotide sequence)

Examples

```
data(bcrdata)

bcrdata$data %>%
  top(5) %>%
  repGermline()
```

 repLoad

Load immune repertoire files into the R workspace

Description**[Deprecated]**

The repLoad function loads repertoire files into R workspace in the immunarch format where you can immediately use them for the analysis. repLoad automatically detects the right format for your files, so all you need is simply provide the path to your files.

See "Details" for more information on supported formats. See "Examples" for diving right into it.

Usage

```
repLoad(.path, .mode = "paired", .coding = TRUE, ...)
```

Arguments

<code>.path</code>	A character string specifying the path to the input data. Input data can be one of the following: • a single repertoire file. In this case repLoad returns an R data.frame; • a vector of paths to repertoire files. Same as in the case with no metadata file presented in the next section below; • a path to the folder with repertoire files and, if available, metadata file "metadata.txt". If the metadata file is presented, then the repLoad returns a list with two elements "data" and "meta". "data" is another list with repertoire R data.frames. "meta" is a data frame with the metadata. If the metadata file "metadata.txt" is not presented, then the repLoad creates a dummy metadata file with sample names and returns a list with two elements "data" and "meta". If input data has multiple chains or cell types stored in the same file (for example, like in 10xGenomics repertoire files), such repertoire files will be splitted to different R data frames with only one type of chain and cell presented. The metadata file will have additional columns specifying cell and chain types for different samples.
<code>.mode</code>	Either "single" for single chain data or "paired" for paired chain data. Currently "single" works for every format, and "paired" works only for 10X Genomics data. By default, 10X Genomics data will be loaded as paired chain data, and other files will be loaded as single chain data.
<code>.coding</code>	A logical value. Set TRUE to get coding-only clonotypes (by default). Set FALSE to get all clonotypes.
<code>...</code>	Extra arguments for parsing functions

Details

The metadata has to be a tab delimited file with first column named "Sample". It can have any number of additional columns with arbitrary names. The first column should contain base names of files without extensions in your folder. Example:

Sample	Sex	Age	Status
immunoseq_1	M	1	C
immunoseq_2	M	2	C
immunoseq_3	FALSE	3	A

Currently, Immunarch support the following formats:

- "immunoseq" - ImmunoSEQ of any version. <http://www.adaptivebiotech.com/immunoseq>
- "mitcr" - MiTCR. <https://github.com/milaboratory/mitcr>
- "mixcr" - MiXCR (the "all" files) of any version. <https://github.com/milaboratory/mixcr>

- "migec" - MiGEC. <http://migec.readthedocs.io/en/latest/>
- "migmap" - For parsing IgBLAST results postprocessed with MigMap. <https://github.com/mikessh/migmap>
- "tcr" - tcR, our previous package. <https://imminfo.github.io/tcr/>
- "vdjtools" - VDJtools of any version. <http://vdjtools-doc.readthedocs.io/en/latest/>
- "imgt" - IMGT HighV-QUEST. <http://www.imgt.org/HighV-QUEST/>
- "airr" - adaptive immune receptor repertoire (AIRR) data format. <http://docs.airr-community.org/en/latest/datarep/overview.html>
- "10x" - 10XGenomics clonotype annotations tables. <https://support.10xgenomics.com/single-cell-vdj/software/pipelines/latest/output/annotation>
- "archer" - ArcherDX clonotype tables. <https://archerdx.com/>

Value

A list with two named elements:

- "data" is a list of input samples;
- "meta" is a data frame with sample metadata.

See Also

[immunr_data_format](#) for immunarch data format; [repSave](#) for file saving; [repOverlap](#), [geneUsage](#) and [repDiversity](#) for starting with immune repertoires basic statistics.

Examples

```
# To load the data from a single file (note that you don't need to specify the data format):
file_path <- paste0(system.file(package = "immunarch"), "/extdata/io/Sample1.tsv.gz")
immdata <- repLoad(file_path)

# Suppose you have a following structure in your folder:
# >_ ls
# immunoseq1.txt
# immunoseq2.txt
# immunoseq3.txt
# metadata.txt

# To load the whole folder with every file in it type:
file_path <- paste0(system.file(package = "immunarch"), "/extdata/io/")
immdata <- repLoad(file_path)
print(names(immdata))

# We recommend creating a metadata file named "metadata.txt" in the folder.

# In that case, when you load your data you will see:
# > immdata <- repLoad("path/to/your/folder/")
# > names(immdata)
# [1] "data" "meta"

# If you do not have "metadata.txt", you will see the same output,
# but your metadata will be almost empty:
```



```
# > immdata <- repLoad("path/to/your/folder/")
# > names(immdata)
# [1] "data" "meta"
```

repOverlap

Main function for public clonotype statistics calculations

Description

[Deprecated]

The repOverlap function is designed to analyse the overlap between two or more repertoires. It contains a number of methods to compare immune receptor sequences that are shared between individuals.

Usage

```
repOverlap(
  .data,
  .method = c("public", "overlap", "jaccard", "tversky", "cosine", "morisita",
    "inc+public", "inc+morisita"),
  .col = "aa",
  .a = 0.5,
  .b = 0.5,
  .verbose = TRUE,
  .step = 1000,
  .n.steps = 10,
  .downsample = FALSE,
  .bootstrap = NA,
  .verbose.inc = NA,
  .force.matrix = FALSE
)
```

Arguments

- | | |
|---------|---|
| .data | <p>The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects.</p> <p>Every object must have columns in the immunarch compatible format. immunarch_data_format</p> <p>Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects.</p> <p>Note: each connection must represent a separate repertoire.</p> |
| .method | <p>A string that specifies the method of analysis or a combination of methods. The repOverlap function supports following basic methods: "public", "overlap", "jaccard", "tversky", "cosine", "morisita". If vector of multiple methods is given for this parameter, the first method will be used.</p> |

<code>.col</code>	A string that specifies the column(s) to be processed. Pass one of the following strings, separated by the plus sign: "nt" for nucleotide sequences, "aa" for amino acid sequences, "v" for V gene segments, "j" for J gene segments. E.g., pass "aa+v" to compute overlaps on CDR3 amino acid sequences paired with V gene segments, i.e., in this case a unique clonotype is a pair of CDR3 amino acid and V gene segment. Clonal counts of equal clonotypes will be summed up.
<code>.a, .b</code>	Alpha and beta parameters for Tversky Index. Default values give the Jaccard index measure.
<code>.verbose</code>	if TRUE then output the progress.
<code>.step</code>	Either an integer or a numeric vector. In the first case, the integer defines the step of incremental overlap. In the second case, the vector encodes all repertoire sampling depths.
<code>.n.steps</code>	Skipped if ".step" is a numeric vector.
<code>.downsample</code>	If TRUE then performs downsampling to N clonotypes at each step instead of choosing the top N clonotypes in incremental overlaps. Change nothing of you are using conventional methods.
<code>.bootstrap</code>	Set NA to turn off any bootstrapping, set a number to perform bootstrapping with this number of tries.
<code>.verbose.inc</code>	Logical. If TRUE then shows output from the computation process.
<code>.force.matrix</code>	Logical. If TRUE then always forces the matrix output even in case of two input repertoires.

Details

"public" and "shared" are synonyms that exist for the convenience of researchers.

The "overlap" coefficient is a similarity measure that measures the overlap between two finite sets.

The "jaccard" index is conceptually a percentage of how many objects two sets have in common out of how many objects they have total.

The "tversky" index is an asymmetric similarity measure on sets that compares a variant to a prototype.

The "cosine" index is a measure of similarity between two non-zero vectors of an inner product space that measures the cosine of the angle between them.

The "morisita" index measures how many times it is more likely to randomly select two sampled points from the same quadrat (the dataset is covered by a regular grid of changing size) then it would be in the case of a random distribution generated from a Poisson process. Duplicate objects are merged with their counts are summed up.

Value

In most cases the return value is a matrix with overlap values for each pair of repertoires.

If only two repertoires were provided, return value is single numeric value.

If one of the incremental method is chosen, return list of overlap matrix.

See Also[inc_overlap](#), [vis](#)**Examples**

```
data(immdata)

# Make data smaller for testing purposes
immdata$data <- top(immdata$data, 4000)

ov <- repOverlap(immdata$data, .verbose = FALSE)
vis(ov)

ov <- repOverlap(immdata$data, "jaccard", .verbose = FALSE)
vis(ov, "heatmap2")
```

repOverlapAnalysis	<i>Post-analysis of public clonotype statistics: PCA, clustering, etc.</i>
--------------------	--

Description**[Deprecated]**

The [repOverlapAnalysis\(\)](#) function contains advanced data analysis methods. You can use several clustering and dimensionality reduction techniques in order to investigate further the difference between repertoires provided.

To cluster a subset of similar data with [repOverlapAnalysis\(\)](#) you can perform hierarchical clustering, k-means or dbscan ('hclust', 'kmeans', 'dbscan' respectively).

To reduce dimensions, for example, to select features for subsequent analysis, you can execute the multidimensional scaling or t-sne algorithms ('mds' and 'tsne' respectively).

Usage

```
repOverlapAnalysis(
  .data,
  .method = ("hclust"),
  .scale = default_scale_fun,
  .raw = TRUE,
  .perp = 1,
  .theta = 0.1,
  .eps = 0.01,
  .k = 2
)
```

Arguments

<code>.data</code>	Any distance matrix between pairs of repertoires. You can also pass your output from repOverlap() .
<code>.method</code>	A string that defines the type of analysis to perform.
<code>.scale</code>	A function to scale the data before passing it to the MDS algorithm.
<code>.raw</code>	A logical value. Set TRUE if you want to receive raw output of clustering or dimensionality reduction function of choice. Set FALSE if you want to receive processed output that can be subjected to visualisation with vis() function.
<code>.perp</code>	A numerical value, t-SNE parameter, see immunr_tsne() .
<code>.theta</code>	A numerical value, t-SNE parameter, see immunr_tsne() .
<code>.eps</code>	A numerical value, DBscan epsilon parameter, see immunr_dbscan() .
<code>.k</code>	The number of clusters to create, passed as k to hcut or as centers to kmeans .

Value

Depends on the last element in the `.method` string. See [immunr_tsne](#) for more info.

Examples

```
data(immdata)
ov <- repOverlap(immdata$data)
repOverlapAnalysis(ov, "mds+hclust") %>% vis()
```

repSample

Downsampling and resampling of immune repertoires

Description**[Deprecated]**

Sample (downsample) repertoires using different approaches.

Usage

```
repSample(
  .data,
  .method = c("downsample", "resample", "sample"),
  .n = NA,
  .prob = TRUE
)
```

Arguments

<code>.data</code>	The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects. Every object must have columns in the immunarch compatible format. immunarch_data_format Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects. Note: each connection must represent a separate repertoire.
<code>.method</code>	Character. Name of a sampling method. See "Details" for more details. Default value is "downsample" that downsamples the repertoires to the number of clones (i.e., reads / UMIs) that the smallest repertoire has, if user doesn't set any value to the ".n" argument.
<code>.n</code>	Integer. Number of clones / clonotypes / reads / UMIs to choose, depending on the method. Set NA to sample repertoires to the size of the smallest repertoire in the ".data".
<code>.prob</code>	Logical. If TRUE then samples the clonotypes with probability weights equal to their number of clones. Used only if ".method" is "sample".

Details

If `.method` is "downsample" then `repSample` chooses `.n` clones (not clonotypes!) from the input repertoires without any probabilistic simulation, but exactly computing each choosed clones. Such approach is more consistent and biologically pleasant than an output from the function if `.method` is "resample".

If `.method` is "resample" then `repSample` uses multinomial distribution to compute the number of occurrences for each cloneset. then it removes zero-number clonotypes and return the resulting data frame. Probabilities for `rmultinom` for each cloneset is a percentage of this cloneset in the "Proportion" column. It's a some sort of simulation of how clonotypes are chosen from the organisms.

if `.method` is "sample" then `repSample` chooses `.n` clonotypes (not clones!) randomly. Depending on the `.prob` argument, the function chooses clonotypes either according to their size (if `.prob` is TRUE, by default), or each clonotype has an equal chance to be choosed (if `.prob` is FALSE). Note that sampling is done without replacing.

Value

Subsampled immune repertoire or a list of subsampled immune repertoires.

See Also

[rmultinom](#), [clonal_proportion](#)

Examples

```
data(imdata)
# Downsampling to 1000 clones (not clonotypes!)
tmp <- repSample(imdata$data[[1]], .n = 1000)
```

```

sum(tmp$Clones)

# Downsampling to 1000 clonotypes
tmp <- repSample(immdata$data[[1]], "sample", .n = 1000)
nrow(tmp)

# Downsampling to the smallest repertoire by clones (not clonotypes!)
tmp <- repSample(immdata$data[c(1, 2)])
sum(tmp[[1]]$Clones)
sum(tmp[[2]]$Clones)

# Downsampling to the smallest repertoire by clonotypes
tmp <- repSample(immdata$data[c(1, 2)], "sample")
nrow(tmp[[1]]$Clones)
nrow(tmp[[2]]$Clones)

```

repSave

Save immune repertoires to the disk

Description

[Deprecated]

The repSave function is designed to save your data to the disk in desirable format. Currently supports "immunarch" and "vdjtools" file formats.

Usage

```
repSave(.data, .path, .format = c("immunarch", "vdjtools"), .compress = TRUE)
```

Arguments

.data	An R dataframe, a list of R dataframes or a list with data and meta where first element is a list of dataframes and the latter is a dataframe with metadata.
.path	A string with the path to the output directory. It should include file name if a single dataframe is provided to .data argument.
.format	A string with desirable format specification. Current options are "immunarch" and "vdjtools".
.compress	A boolean value. Defines whether the output will be compressed or not.

Details

It is not necessary to create directories beforehand. If the provided directory does not exist it will be created automatically.

Value

No return value.

Examples

```
## Not run:
data(immdata)
# Reduce data to save time on examples
immdata$data <- map(immdata$data, ~ .x %>% head(10))
dirpath <- tempdir()
# Save the list of repertoires
repSave(immdata, dirpath)
# Load it and check if it is the same
new_immdata <- repLoad(dirpath)
# sum(immdata$data[[1]] != new_immdata$data[[1]], na.rm = TRUE)
# sum(immdata$data[[2]] != new_immdata$data[[2]], na.rm = TRUE)
# sum(immdata$meta != new_immdata$meta, na.rm = TRUE)

## End(Not run)
```

repSomaticHypermutation

Calculates number of mutations against the germline for each clonotype

Description**[Deprecated]**

This function aligns V and J genes from the germline in each cluster with corresponding genes in each clonotype, saves the alignments for purpose of visualization, and calculates number of mutations for each clonotype.

Usage

```
repSomaticHypermutation(.data, .threads, .nofail)
```

Arguments

.data	The data to be processed: an output of repClonalFamily(); variants with one sample and list of samples are both supported.
.threads	Number of threads to use.
.nofail	Will return NA instead of stopping if Clustal W is not installed. Used to avoid raising errors in examples on computers where Clustal W is not installed.

Value

Dataframe or list of dataframes (if input is a list with multiple samples). The dataframe has all the columns from repClonalFamily() output dataframe, with Sequence column unnested: the resulting dataframe has one line per clonotype. Clone.ID column contains original IDs for clonotypes, and can be used as dataframe key. New columns are added:

- `Germline.Alignment.V`: contains V gene alignment of current clonotype with the germline
- `Germline.Alignment.J`: contains J gene alignment of current clonotype with the germline
- `Substitutions`: contains number of substitutions in the alignment (summary for V and J)
- `Insertions`: contains number of insertions in the clonotype relative to germline (summary for V and J)
- `Deletions`: contains number of deletions in the clonotype relative to germline (summary for V and J)
- `Mutations`: contains total number of mutations in the alignment (summary for V and J)

Examples

```
data(bcrdata)
bcr_data <- bcrdata$data

bcr_data %>%
  seqCluster(seqDist(bcr_data), .fixed_threshold = 3) %>%
  repGermline(.threads = 1) %>%
  repAlignLineage(.min_lineage_sequences = 2, .align_threads = 2, .nofail = TRUE) %>%
  repClonalFamily(.threads = 1, .nofail = TRUE) %>%
  repSomaticHypermutation(.threads = 1, .nofail = TRUE)
```

select_barcodes

Select specific clonotypes using barcodes from single-cell metadata

Description

[Deprecated]

Subsets the input immune repertoire by barcodes. Creates a vector of barcodes to subset or a vector cluster IDs and corresponding barcodes to get a list of immune repertoires corresponding to cluster IDs. Columns with clonotype counts and proportions are changed accordingly to the filtered barcodes.

Usage

```
select_barcodes(.data, .barcodes, .force.list = FALSE)
```

Arguments

<code>.data</code>	The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects. Every object must have columns in the immunarch compatible format. immunarch_data_format Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects. Note: each connection must represent a separate repertoire.
--------------------	---

- .barcodes Either a character vector with barcodes or a named character/factor vector with barcodes as names and cluster IDs a vector elements. The output of Seurat's Idents function works.
- .force.list Logical. If TRUE then always returns a list, even if the result is one data frame.

Value

An immune repertoire (if ".barcodes" is a barcode vector) or a list of immune repertoires (if ".barcodes" is named vector or an output from Seurat::Idents()). Each element is an immune repertoire with clonotype barcodes corresponding to the input barcodes. The output list names are cluster names in the ".barcode" argument (Seurat::Idents() case only).

See Also

[select_clusters](#)

Examples

```
## Not run:
data(immdata)
# Create a fake single-cell data
df <- immdata$data[[1]]
df$Barcode <- "AAAAACCCCC"
df$Barcode[51:nrow(df)] <- "GGGGCCCCC"
barcodes <- "AAAAACCCCC"
df <- select_barcodes(df, barcodes)
nrow(df)

## End(Not run)
```

select_clusters	<i>Split the immune repertoire data to clusters from single-cell barcodes</i>
-----------------	---

Description

[Deprecated]

Given the vector of barcodes from Seurat, splits the input repertoires to separate subsets following the barcodes' assigned IDs. Useful in case you want to split immune repertoires by patients or clusters.

Usage

```
select_clusters(.data, .clusters, .field = "Cluster")
```

Arguments

<code>.data</code>	List of two elements "data" and "meta", with "data" being a list of immune repertoires, and "meta" being a metadata table.
<code>.clusters</code>	Factor vector with barcodes as vector names and cluster IDs as vector elements. The output of the Seurat Idents function works.
<code>.field</code>	A string specifying the name of the field in the input metadata. New immune repertoire subsets will have cluster IDs in this field.

Value

A list with two elements "data" and "meta" with updated immune repertoire tables and metadata.

See Also

[select_barcodes](#)

Examples

```
## Not run:
library(Seurat)
Idents(pbmc_small)
new_cluster_ids <- c("A", "B", "C")
new_cluster_ids <- levels(pbmc_small)
new_cluster_ids
pbmc_small <- RenameIdents(pbmc_small, new_cluster_ids)

## End(Not run)
```

seqCluster

Function for assigning clusters based on sequences similarity

Description**[Deprecated]**

Graph clustering based on distances between sequences

Usage

```
seqCluster(.data, .dist, .perc_similarity, .nt_similarity, .fixed_threshold)
```

Arguments

<code>.data</code>	The data which was used to calculate <code>.dist</code> object. Can be data.frame , data.table::data.table , or a list of these objects. Every object must have columns in the immunarch compatible format immunarch_data_format
<code>.dist</code>	List of distance objects produced with seqDist function.

`.perc_similarity` Numeric value between 0 and 1 specifying the maximum acceptable weight of an edge in a graph. This threshold depends on the length of sequences.

`.nt_similarity` Numeric between 0-sequence length specifying the threshold of allowing a 1 in n nucleotides mismatch in sequences.

`.fixed_threshold` Numeric specifying the threshold on the maximum weight of an edge in a graph.

Value

Immdata data format object. Same as `.data`, but with extra 'Cluster' column with clusters assigned.

Examples

```
## Not run:
data(immdata)
# In this example, we will use only 2 samples with 500 clonotypes in each for time saving
input_data <- lapply(immdata$data[1:2], head, 500)
dist_result <- seqDist(input_data)
cluster_result <- seqCluster(input_data, dist_result, .fixed_threshold = 1)

## End(Not run)
```

seqDist

Function for computing distance for sequences

Description

[Deprecated]

Computing sequential distances between clonotypes from two repertoires:

Usage

```
seqDist(.data, .col = 'CDR3.nt', .method = 'hamming',
        .group_by = c("V.name", "J.name"), .group_by_seqLength = TRUE, .trim_genes = TRUE, ...)
```

Arguments

`.data` The data to be processed. Can be [data.frame](#), [data.table::data.table](#), or a list of these objects.
Every object must have columns in the immunarch compatible format [immunarch_data_format](#)

`.col` A string that specifies the column name to be processed. The default value is 'CDR3.nt'.

`.method` Character value or user-defined function.

`.group_by` Character vector of column names to group sequence by. The default value is `c("V.first", "J.first")`. Columns "V.first" and "J.first" containing first genes without allele suffixes are calculated automatically from "V.name" and "J.name" if absent in the data. Pass NA for no grouping options.

`.group_by_seqLength` If TRUE - adds grouping by sequence length of `.col` argument

`.trim_genes` If TRUE - use only general gene values (e.g. "IGHV1-18") of `.group_by` columns for clustering; if FALSE - can cause very small clusters in case of high resolution genotyping

... Extra arguments for user-defined function.

The default value is 'hamming' for Hamming distance which counts the number of character substitutions that turns b into a. If a and b have different number of characters the distance is Inf.

Other possible values are:

'lv' for Levenshtein distance which counts the number of deletions, insertions and substitutions necessary to turn b into a.

'lcs' for longest common substring is defined as the longest string can be obtained by pairing characters from a and b while keeping the order of characters intact.

In case of user-defined function, it should take x and y parameters as input and return `dist` object.

Value

Named list of list with `dist` objects for given repertoires for each combination of `.group_by` variable(s) and/or sequence length of `.col`.

Examples

```
## Not run:
data(immdata)
# Reducing data to save time on examples
immdata$data <- map(immdata$data, ~ .x %>% head(10))
# Computing hamming distance for the first two repertoires in `immdata`
seqDist(immdata$data[1:2])

# Here we define a custom distance function
# that will count the difference in number of characters in sequences.

f <- function(x, y) {
  res <- matrix(nrow = length(x), ncol = length(y))
  for (i in 1:length(x)) {
    res[i, ] <- abs(nchar(x[i]) - nchar(y))
  }
  dimnames(res) <- list(x, y)
  return(as.dist(res))
}

seqDist(immdata$data[1:2], .method = f, .group_by_seqLength = FALSE)
```

```
## End(Not run)
```

set_pb	<i>Set and update progress bars</i>
--------	-------------------------------------

Description

Set and update progress bars

Usage

```
set_pb(.max)

add_pb(.pb, .value = 1)
```

Arguments

.max	Integer. Maximal value of the progress bar.
.pb	Progress bar object from set_pb.
.value	Numeric. Value to add to the progress bar at each step.

Value

An updated progress bar.

Developer Examples

```
pb <- immunarch:::set_pb(100) immunarch:::add_pb(pb, 25) immunarch:::add_pb(pb, 25) immunarch:::add_pb(pb, 25) immunarch:::add_pb(pb, 25) immunarch:::add_pb(pb, 25) close(pb)
```

spectratype	<i>Immune repertoire spectratyping</i>
-------------	--

Description

[Deprecated]

Usage

```
spectratype(.data, .quant = c("id", "count"), .col = "nt")
```

Arguments

<code>.data</code>	The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects. Every object must have columns in the immunarch compatible format. immunarch_data_format Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects. Note: each connection must represent a separate repertoire.
<code>.quant</code>	Select the column with clonal counts to evaluate. Set to "id" to count every clonotype once. Set to "count" to take into the account number of clones per clonotype.
<code>.col</code>	A string that specifies the column(s) to be processed. The output is one of the following strings, separated by the plus sign: "nt" for nucleotide sequences, "aa" for amino acid sequences, "v" for V gene segments, "j" for J gene segments. E.g., pass "aa+v" for spectratyping on CDR3 amino acid sequences paired with V gene segments, i.e., in this case a unique clonotype is a pair of CDR3 amino acid and V gene segment. Clonal counts of equal clonotypes will be summed up.

Value

Data frame with distributions of clonotypes per CDR3 length.

Examples

```
# Load the data
data(immdata)
sp <- spectratype(immdata$data[[1]], .col = "aa+v")
vis(sp)
```

split_to_kmers

Analysis immune repertoire kmer statistics: sequence profiles, etc.

Description

[Deprecated]

Usage

```
split_to_kmers(.data, .k)
```

```
kmer_profile(.data, .method = c("freq", "prob", "wei", "self"), .remove.stop = TRUE)
```

Arguments

<code>.data</code>	Character vector or the output from <code>getKmers</code> .
<code>.k</code>	Integer. Size of k-mers.
<code>.method</code>	Character vector of length one. If "freq" then returns a position frequency matrix (PFM) - a matrix with occurrences of each amino acid in each position. If "prob" then returns a position probability matrix (PPM) - a matrix with probabilities of occurrences of each amino acid in each position. This is a traditional representation of sequence motifs. If "wei" then returns a position weight matrix (PWM) - a matrix with log likelihoods of PPM elements. If "self" then returns a matrix with self-information of elements in PWM. For more information see https://en.wikipedia.org/wiki/Position_weight_matrix .
<code>.remove.stop</code>	Logical. If TRUE (by default) remove stop codons.

Value

`split_to_kmers` - Data frame with two columns (k-mers and their counts).

`kmer_profile` - a matrix with per-position amino acid statistics.

Examples

```
data(immdata)
kmers <- getKmers(immdata$data[[1]], 5)
kmer_profile(kmers) %>% vis()
```

switch_type	<i>Return a column's name</i>
-------------	-------------------------------

Description

Return a column's name

Usage

```
switch_type(type)

process_col_argument(.col)
```

Arguments

<code>type</code>	Character. Specifies the column to choose: "nt" chooses the CDR3 nucleotide column, "aa" chooses the CDR3 amino acid column, "v" chooses the V gene segment column, "j" chooses the J gene segment column.
<code>.col</code>	A string that specifies the column(s) to be processed. Select one of the following strings, separated by the plus sign: "nt" for nucleotide sequences, "aa" for amino acid sequences, "v" for V gene segments, "j" for J gene segments.

Value

A column's name.

Developer Examples

```
immunarch:::switch_type("nuc") immunarch:::switch_type("v")
```

top	<i>Get the N most abundant clonotypes</i>
-----	---

Description

Get the N most abundant clonotypes

Usage

```
top(.data, .n = 10)
```

Arguments

.data	The data to be processed. Can be data.frame, data.table::data.table, or a list of these objects. Every object must have columns in the immunarch compatible format. immunarch_data_format Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects. Note: each connection must represent a separate repertoire.
.n	Numeric. Number of the most abundant clonotypes to return.

Value

Data frame with the .n most abundant clonotypes only.

Examples

```
data(immdata)
top(immdata$data)
top(immdata$data[[1]])
```

trackClonotypes	<i>Track clonotypes across time and data points</i>
-----------------	---

Description

[Deprecated]

Tracks the temporal dynamics of clonotypes in repertoires. For example, tracking across multiple time points after vaccination.

Note: duplicated clonotypes are merged and their counts are summed up.

Usage

```
trackClonotypes(.data, .which = list(1, 15), .col = "aa", .norm = TRUE)
```

Arguments

- | | |
|---------------------|---|
| <code>.data</code> | <p>The data to process. It can be a data.frame, a data.table::data.table, or a list of these objects.</p> <p>Every object must have columns in the immunarch compatible format. immunarch_data_format</p> <p>Competent users may provide advanced data representations: DBI database connections, or a list of these objects. They are supported with the same limitations as basic objects.</p> <p>Note: each connection must represent a separate repertoire.</p> |
| <code>.which</code> | <p>An argument that regulates which clonotypes to choose for tracking. There are three options for this argument:</p> <ol style="list-style-type: none"> 1. passes a list with two elements <code>list(X, Y)</code>, where X is the name or the index of a target repertoire from ".data", and Y is the number of the most abundant clonotypes to take from X. 2. passes a character vector of sequences to take from all data frames; 3. passes a data frame (data table, database) with one or more columns - first for sequences, and other for gene segments (if applicable). <p>See the "Examples" below with examples for each option.</p> |
| <code>.col</code> | <p>A character vector of length 1. Specifies an identifier for a column, from which the function chooses clonotype sequences. Specify "nt" for nucleotide sequences, "aa" for amino acid sequences, "aa+v" for amino acid sequences and Variable genes, "nt+j" for nucleotide sequences with Joining genes, or any combination of the above. Used only if ".which" has option 1) or option 2).</p> |
| <code>.norm</code> | <p>Logical. If TRUE then uses Proportion instead of the number of Clones per clonotype to store in the function output.</p> |

Value

Data frame with input sequences and counts or proportions for each of the input repertoire.

Examples

```
# Load an example data that comes with immunarch
data(immdata)

# Make the data smaller in order to speed up the examples
immdata$data <- immdata$data[c(1, 2, 3, 7, 8, 9)]
immdata$meta <- immdata$meta[c(1, 2, 3, 7, 8, 9), ]

# Option 1
# Choose the first 10 amino acid clonotype sequences
# from the first repertoire to track
tc <- trackClonotypes(immdata$data, list(1, 10), .col = "aa")
# Choose the first 20 nucleotide clonotype sequences
# and their V genes from the "MS1" repertoire to track
tc <- trackClonotypes(immdata$data, list("MS1", 20), .col = "nt+v")

# Option 2
# Choose clonotypes with amino acid sequences "CASRGLITDTQYF" or "CSASRGSPNEQYF"
tc <- trackClonotypes(immdata$data, c("CASRGLITDTQYF", "CSASRGSPNEQYF"), .col = "aa")

# Option 3
# Choose the first 10 clonotypes from the first repertoire
# with amino acid sequences and V segments
target <- immdata$data[[1]] %>%
  select(CDR3.aa, V.name) %>%
  head(10)
tc <- trackClonotypes(immdata$data, target)

# Visualise the output regardless of the chosen option
# There are three ways to visualise it, regulated by the .plot argument
vis(tc, .plot = "smooth")
vis(tc, .plot = "area")
vis(tc, .plot = "line")

# Visualising timepoints
# First, we create an additional column in the metadata with randomly chosen timepoints:
immdata$meta$Timepoint <- sample(1:length(immdata$data))
immdata$meta
# Next, we create a vector with samples in the right order,
# according to the "Timepoint" column (from smallest to greatest):
sample_order <- order(immdata$meta$Timepoint)
# Sanity check: timepoints are following the right order:
immdata$meta$Timepoint[sample_order]
# Samples, sorted by the timepoints:
immdata$meta$Sample[sample_order]
# And finally, we visualise the data:
vis(tc, .order = sample_order)
```

Description

[Deprecated]

Output from every function in immunarch can be visualised with a single function - `vis`. The `vis` automatically detects the type of the data and draws a proper visualisation. For example, output from the `repOverlap` function will be identified as repertoire overlap values and respective visualisation will be chosen without any additional arguments. See "Details" for the list of available visualisations.

Usage

```
vis(.data, ...)
```

Arguments

<code>.data</code>	Pass the output from any immunarch analysis tool to <code>vis()</code> .
<code>...</code>	Any other arguments, see the "Details" section for specific visualisation functions.

Details

List of available visualisations for different kinds of data.

Basic analysis:

- Exploratory analysis results (from [repExplore](#)) - see [vis.immunr_exp_vol](#);
- Clonality statistics (from [repClonality](#)) - see [vis.immunr_homeo](#).

Overlaps and public clonotypes:

- Overlaps (from [repOverlap](#)) using heatmaps, circos plots, polar area plots - see [vis.immunr_ov_matrix](#);
- Overlap clustering (from [repOverlapAnalysis](#)) - see [vis.immunr_hclust](#);
- Repertoire incremental overlaps (from [repOverlap](#)) - see [vis.immunr_inc_overlap](#);
- Public repertoire abundance (from [pubRep](#)) - see [vis.immunr_public_repertoire](#).

Gene usage:

- Gene usage statistics (from [geneUsage](#)) using bar plots, box plots - see [vis.immunr_gene_usage](#);
- Gene usage distances (from [geneUsageAnalysis](#)) using heatmaps, circos plots, polar area plots - see [vis.immunr_ov_matrix](#);
- Gene usage clustering (from [geneUsageAnalysis](#)) - see [vis.immunr_hclust](#).

Diversity estimation:

- Diversity estimations (from [repDiversity](#)) - see [vis.immunr_chao1](#).

BCR analysis:

- Clonal tree (from [repClonalFamily](#)) - see [vis.clonal_family](#) and [vis.clonal_family_tree](#).

Advanced analysis:

- Repertoire dynamics (from [trackClonotypes](#)) - see [vis.immunr_dynamics](#);
- Sequence logo plots of amino acid distributions (from [kmer_profile](#)) - see [vis_seqlogo](#);
- Kmers distributions (from [getKmers](#)) - see [vis.immunr_kmer_table](#);
- Mutation networks (from [mutationNetwork](#)) - Work In Progress on [vis.immunr_mutation_network](#);
- CDR3 amino acid properties, e.g., biophysical (from [cdrProp](#)) - Work In Progress on [vis.immunr_cdr_prop](#).

Additionally, we provide a wrapper functions for visualisations of common data types:

- Any data frames or matrices using heatmaps - see [vis_heatmap](#) and [vis_heatmap2](#);
- Any data frames or matrices using circos plots - see [vis_circos](#).

Value

A ggplot2, pheatmap or circlize object.

See Also

[fixVis](#) for precise manipulation of plots.

Examples

```
## Not run:
# Load the test data
data(immdata)

# Compute and visualise:
ov <- repOverlap(immdata$data)
vis(ov)

gu <- geneUsage(immdata$data)
vis(gu)

dv <- repDiversity(immdata$data)
vis(dv)

## End(Not run)
```

vis.clonal_family	<i>Visualise clonal family tree: wrapper for calling on the entire repClonalFamily output</i>
-------------------	---

Description

[Deprecated]

Usage

```
## S3 method for class 'clonal_family'
vis(.data, ...)
```

Arguments

.data Clonal families from 1 or multiple samples: [repClonalFamily\(\)](#) output.
 ... Not used here.

Value

A ggraph object.

Examples

```
## Not run:
data(bcrdata)
bcr_data <- bcrdata$data

clonal_family <- bcr_data %>%
  seqCluster(seqDist(bcr_data), .fixed_threshold = 3) %>%
  repGermline(.threads = 1) %>%
  repAlignLineage(.min_lineage_sequences = 2, .align_threads = 2, .nofail = TRUE) %>%
  repClonalFamily(.threads = 1, .nofail = TRUE) %>%
  vis()

## End(Not run)
```

vis.clonal_family_tree

Visualise clonal family tree

Description

[Deprecated]

Usage

```
## S3 method for class 'clonal_family_tree'
vis(.data, ...)
```

Arguments

.data Single clonal family tree data from 1 cluster: 1 element from TreeStats column
 from [repClonalFamily\(\)](#) output.
 ... Not used here.

Value

A ggraph object.

Examples

```
## Not run:
data(bcrdata)
bcr_data <- bcrdata$data

clonal_family <- bcr_data %>%
  seqCluster(seqDist(bcr_data), .fixed_threshold = 3) %>%
  repGermline(.threads = 1) %>%
  repAlignLineage(.min_lineage_sequences = 2, .align_threads = 2, .nofail = TRUE) %>%
  repClonalFamily(.threads = 1, .nofail = TRUE)

# This condition can be omitted; it prevents the example from crashing
# when ClustalW or PHYLIP are not installed
if (!("step_failure_ignored" %in% class(clonal_family))) {
  vis(clonal_family[["full_clones"]][["TreeStats"]][[2]])
}

## End(Not run)
```

vis.immunr_chao1

Visualise diversity.

Description

[Deprecated]

An utility function to visualise the output from [repDiversity\(\)](#).

Usage

```
## S3 method for class 'immunr_chao1'
vis(
  .data,
  .by = NA,
  .meta = NA,
  .errorbars = c(0.025, 0.975),
  .errorbars.off = FALSE,
  .points = TRUE,
  .test = TRUE,
  .signif.label.size = 3.5,
  ...
)
```

Arguments

.data Output from [repDiversity\(\)](#).

<code>.by</code>	Pass NA if you want to plot samples without grouping. You can pass a character vector with one or several column names from ".meta" to group your data before plotting. In this case you should provide ".meta". You can pass a character vector that exactly matches the number of samples in your data, each value should correspond to a sample's property. It will be used to group data based on the values provided. Note that in this case you should pass NA to ".meta".
<code>.meta</code>	A metadata object. An R dataframe with sample names and their properties, such as age, serostatus or hla.
<code>.errorbars</code>	A numeric vector of length two with quantiles for error bars on sectors. Disabled if ".errorbars.off" is TRUE.
<code>.errorbars.off</code>	If TRUE then plot CI bars for distances between each group. Disabled if no group passed to the ".by" argument.
<code>.points</code>	A logical value defining whether points will be visualised or not.
<code>.test</code>	A logical vector whether statistical tests should be applied. See "Details" for more information.
<code>.signif.label.size</code>	An integer value defining the size of text for p-value.
<code>...</code>	Not used here.

Details

If data is grouped, then statistical tests for comparing means of groups will be performed, unless `.test = FALSE` is supplied. In case there are only two groups, the Wilcoxon rank sum test (https://en.wikipedia.org/wiki/Wilcoxon_signed-rank_test) is performed (R function `wilcox.test()` with an argument `exact = FALSE`) for testing if there is a difference in mean rank values between two groups. In case there more than two groups, the Kruskal-Wallis test (https://en.wikipedia.org/wiki/Kruskal%E2%80%93Wallis_test) is performed (R function `kruskal.test()`), that is equivalent to ANOVA for ranks and it tests whether samples from different groups originated from the same distribution. A significant Kruskal-Wallis test indicates that at least one sample stochastically dominates one other sample. Adjusted for multiple comparisons P-values are plotted on the top of groups. P-value adjusting is done using the Holm method (https://en.wikipedia.org/wiki/Holm%E2%80%93Bonferroni_method) (also known as Holm-Bonferroni correction). You can execute the command `?p.adjust` in the R console to see more.

Value

A ggplot2 object.

See Also

[repDiversity vis](#)

Examples

```
## Not run:
data(immdata)
```

```
dv <- repDiversity(immdata$data, "chao1")
vis(dv)

## End(Not run)
```

```
vis.immunr_clonal_prop
```

Visualise results of the clonality analysis

Description

[Deprecated]

An utility function to visualise the output from `repClonality()`.

Usage

```
## S3 method for class 'immunr_clonal_prop'
vis(
  .data,
  .by = NA,
  .meta = NA,
  .errorbars = c(0.025, 0.975),
  .errorbars.off = FALSE,
  .points = TRUE,
  .test = TRUE,
  .signif.label.size = 3.5,
  ...
)
```

Arguments

<code>.data</code>	Output from <code>repClonality()</code> .
<code>.by</code>	Pass NA if you want to plot samples without grouping. You can pass a character vector with one or several column names from <code>".meta"</code> to group your data before plotting. In this case you should provide <code>".meta"</code>. You can pass a character vector that exactly matches the number of samples in your data, each value should correspond to a sample's property. It will be used to group data based on the values provided. Note that in this case you should pass NA to <code>".meta"</code>.
<code>.meta</code>	A metadata object. An R dataframe with sample names and their properties, such as age, serostatus or hla.
<code>.errorbars</code>	A numeric vector of length two with quantiles for error bars on sectors. Disabled if <code>".errorbars.off"</code> is TRUE.
<code>.errorbars.off</code>	If TRUE then plot CI bars for distances between each group. Disabled if no group passed to the <code>".by"</code> argument.

<code>.points</code>	A logical value defining whether points will be visualised or not.
<code>.test</code>	A logical vector whether statistical tests should be applied. See "Details" for more information.
<code>.signif.label.size</code>	An integer value defining the size of text for p-value.
<code>...</code>	Not used here.

Details

If data is grouped, then statistical tests for comparing means of groups will be performed, unless `.test = FALSE` is supplied. In case there are only two groups, the Wilcoxon rank sum test (https://en.wikipedia.org/wiki/Wilcoxon_signed-rank_test) is performed (R function `wilcox.test()` with an argument `exact = FALSE`) for testing if there is a difference in mean rank values between two groups. In case there more than two groups, the Kruskal-Wallis test (https://en.wikipedia.org/wiki/Kruskal%E2%80%93Wallis_test) is performed (R function `kruskal.test()`), that is equivalent to ANOVA for ranks and it tests whether samples from different groups originated from the same distribution. A significant Kruskal-Wallis test indicates that at least one sample stochastically dominates one other sample. Adjusted for multiple comparisons P-values are plotted on the top of groups. P-value adjusting is done using the Holm method (https://en.wikipedia.org/wiki/Holm%E2%80%93Bonferroni_method) (also known as Holm-Bonferroni correction). You can execute the command `?p.adjust` in the R console to see more.

Value

A ggplot2 object.

See Also

[repClonality vis](#)

Examples

```
## Not run:
data(immdata)
clp <- repClonality(immdata$data, "clonal.prop")
vis(clp)

hom <- repClonality(immdata$data, "homeo")
# Remove p values and points from the plot
vis(hom, .by = "Status", .meta = immdata$meta, .test = FALSE, .points = FALSE)

## End(Not run)
```

vis.immunr_dynamics *Visualise clonotype dynamics*

Description**[Deprecated]****Usage**

```
## S3 method for class 'immunr_dynamics'
vis(.data, .plot = c("smooth", "area", "line"), .order = NA, .log = FALSE, ...)
```

Arguments

.data	Output from the trackClonotypes function.
.plot	Character. Either "smooth", "area" or "line". Each specifies a type of plot for visualisation of clonotype dynamics.
.order	Numeric or character vector. Specifies the order to samples, e.g., it used for ordering samples by timepoints. Either See "Examples" below for more details.
.log	Logical. If TRUE then use log-scale for the frequency axis.
...	Not used here.

Value

A ggplot2 object.

Examples

```
## Not run:
# Load an example data that comes with immunarch
data(immdata)

# Make the data smaller in order to speed up the examples
immdata$data <- immdata$data[c(1, 2, 3, 7, 8, 9)]
immdata$meta <- immdata$meta[c(1, 2, 3, 7, 8, 9), ]

# Option 1
# Choose the first 10 amino acid clonotype sequences
# from the first repertoire to track
tc <- trackClonotypes(immdata$data, list(1, 10), .col = "aa")
# Choose the first 20 nucleotide clonotype sequences
# and their V genes from the "MS1" repertoire to track
tc <- trackClonotypes(immdata$data, list("MS1", 20), .col = "nt+v")

# Option 2
# Choose clonotypes with amino acid sequences "CASRGLITDTQYF" or "CSASRGSPNEQYF"
tc <- trackClonotypes(immdata$data, c("CASRGLITDTQYF", "CSASRGSPNEQYF"), .col = "aa")
```

```

# Option 3
# Choose the first 10 clonotypes from the first repertoire
# with amino acid sequences and V segments
target <- immdata$data[[1]] %>%
  select(CDR3.aa, V.name) %>%
  head(10)
tc <- trackClonotypes(immdata$data, target)

# Visualise the output regardless of the chosen option
# There are three ways to visualise it, regulated by the .plot argument
vis(tc, .plot = "smooth")
vis(tc, .plot = "area")
vis(tc, .plot = "line")

# Visualising timepoints
# First, we create an additional column in the metadata with randomly chosen timepoints:
immdata$meta$Timepoint <- sample(1:length(immdata$data))
immdata$meta
# Next, we create a vector with samples in the right order,
# according to the "Timepoint" column (from smallest to greatest):
sample_order <- order(immdata$meta$Timepoint)
# Sanity check: timepoints are following the right order:
immdata$meta$Timepoint[sample_order]
# Samples, sorted by the timepoints:
immdata$meta$Sample[sample_order]
# And finally, we visualise the data:
vis(tc, .order = sample_order)

## End(Not run)

```

vis.immunr_exp_vol	<i>Visualise results of the exploratory analysis</i>
--------------------	--

Description

[Deprecated]

An utility function to visualise the output from [repExplore\(\)](#).

Usage

```

## S3 method for class 'immunr_exp_vol'
vis(
  .data,
  .by = NA,
  .meta = NA,
  .errorbars = c(0.025, 0.975),
  .errorbars.off = FALSE,
  .points = TRUE,
  .test = TRUE,

```

```

    .signif.label.size = 3.5,
    ...
)

```

Arguments

<code>.data</code>	Output from <code>repExplore()</code> .
<code>.by</code>	Pass NA if you want to plot samples without grouping. You can pass a character vector with one or several column names from ".meta" to group your data before plotting. In this case you should provide ".meta". You can pass a character vector that exactly matches the number of samples in your data, each value should correspond to a sample's property. It will be used to group data based on the values provided. Note that in this case you should pass NA to ".meta".
<code>.meta</code>	A metadata object. An R dataframe with sample names and their properties, such as age, serostatus or hla.
<code>.errorbars</code>	A numeric vector of length two with quantiles for error bars on sectors. Disabled if ".errorbars.off" is TRUE.
<code>.errorbars.off</code>	If TRUE then plot CI bars for distances between each group. Disabled if no group passed to the ".by" argument.
<code>.points</code>	A logical value defining whether points will be visualised or not.
<code>.test</code>	A logical vector whether statistical tests should be applied. See "Details" for more information.
<code>.signif.label.size</code>	An integer value defining the size of text for p-value.
<code>...</code>	Not used here.

Details

If data is grouped, then statistical tests for comparing means of groups will be performed, unless `.test = FALSE` is supplied. In case there are only two groups, the Wilcoxon rank sum test (https://en.wikipedia.org/wiki/Wilcoxon_signed-rank_test) is performed (R function `wilcox.test()` with an argument `exact = FALSE`) for testing if there is a difference in mean rank values between two groups. In case there more than two groups, the Kruskal-Wallis test (https://en.wikipedia.org/wiki/Kruskal%E2%80%93Wallis_test) is performed (R function `kruskal.test()`), that is equivalent to ANOVA for ranks and it tests whether samples from different groups originated from the same distribution. A significant Kruskal-Wallis test indicates that at least one sample stochastically dominates one other sample. Adjusted for multiple comparisons P-values are plotted on the top of groups. P-value adjusting is done using the Holm method (https://en.wikipedia.org/wiki/Holm%E2%80%93Bonferroni_method) (also known as Holm-Bonferroni correction). You can execute the command `?p.adjust` in the R console to see more.

Value

A `ggplot2` object.

See Also[repExplore vis](#)**Examples**

```
## Not run:
data(immdata)
repExplore(immdata$data, "volume") %>% vis()
repExplore(immdata$data, "count") %>% vis()
repExplore(immdata$data, "len") %>% vis()
repExplore(immdata$data, "clones") %>% vis()

## End(Not run)
```

vis.immunr_gene_usage *Histograms and boxplots (general case / gene usage)*

Description**[Deprecated]**

Visualise distributions of genes using heatmaps or other plots.

Usage

```
## S3 method for class 'immunr_gene_usage'
vis(.data, .plot = c("hist", "box", "heatmap", "heatmap2", "circos"), ...)
```

Arguments

.data	Output from the geneUsage function.
.plot	String specifying the plot type: • "hist" for histograms using vis_hist; • "heatmap" for heatmaps using vis_heatmap; • "heatmap2" for heatmaps using vis_heatmap2; • "circos" for circos plots using vis_circos.
...	Other arguments passed to corresponding functions depending on the plot type: • "hist" - passes arguments to vis_hist; • "box" - passes arguments to vis_box; • "heatmap" - passes arguments to vis_heatmap; • "heatmap2" - passes arguments to vis_heatmap2 and heatmap from the "pheatmap" package; • "circos" - passes arguments to vis_circos and circlize::chordDiagram from the "circlize" package.

Value

A ggplot2 object, pheatmap or circlize object.

See Also

[geneUsage](#)

Examples

```
## Not run:
data(immdata)

gu <- geneUsage(immdata$data[[1]])
vis(gu)

gu <- geneUsage(immdata$data)
vis(gu, .by = "Status", .meta = immdata$meta)
vis(gu, "box", .by = "Status", .meta = immdata$meta)

## End(Not run)
```

vis.immunr_hclust

Visualisation of hierarchical clustering

Description**[Deprecated]**

Visualisation of the results of hierarchical clustering. For other clustering visualisations see [vis.immunr_kmeans](#).

Usage

```
## S3 method for class 'immunr_hclust'
vis(.data, .rect = FALSE, .plot = c("clust", "best"), ...)
```

Arguments

.data	Clustering results from repOverlapAnalysis or geneUsageAnalysis .
.rect	Passed to factoextra::fviz_dend - whether to add a rectangle around groups.
.plot	A character vector of length one or two specifying which plots to visualise. If "clust" then plot only the clustering. If "best" then plot the number of optimal clusters. If both then plot both.
...	Not used here.

Value

Ggplot2 objects inside the patchwork container.

See Also

[vis](#), [repOverlapAnalysis](#), [geneUsageAnalysis](#)

Examples

```
## Not run:
data(immdata)
ov <- repOverlap(immdata$data)
repOverlapAnalysis(ov, "mds+hclust") %>% vis()

## End(Not run)
```

```
vis.immunr_inc_overlap
```

Visualise incremental overlaps

Description

[Deprecated]

Usage

```
## S3 method for class 'immunr_inc_overlap'
vis(.data, .target = 1, .grid = FALSE, .ncol = 2, ...)
```

Arguments

.data	Output from the repOverlap function that uses "top" methods.
.target	Index of a repertoire to plot. Omitted if .grid is TRUE.
.grid	Logical. If TRUE then plot all similarities in a grid.
.ncol	Numeric. Number of columns in the resulting grid.
...	Not used here.

Value

A ggplot2 object.

See Also

[repOverlap](#)

Examples

```
## Not run:
data(immdata)
tmp <- repOverlap(immdata$data[1:4], "inc+overlap", .verbose.inc = FALSE, .verbose = FALSE)
vis(tmp, .target = 1)
vis(tmp, .grid = TRUE)

## End(Not run)
```

vis.immunr_kmeans

Visualisation of K-means and DBSCAN clustering

Description

[Deprecated]

Visualisation of the results of K-means and DBSCAN clustering. For hierarchical clustering visualisations see [vis.immunr_hclust](#).

Usage

```
## S3 method for class 'immunr_kmeans'
vis(
  .data,
  .point = TRUE,
  .text = TRUE,
  .ellipse = TRUE,
  .point.size = 2,
  .text.size = 10,
  .plot = c("clust", "best"),
  ...
)
```

Arguments

.data	Clustering results from repOverlapAnalysis or geneUsageAnalysis .
.point	If TRUE then plot sample points. Passed to factoextra::fviz_cluster .
.text	If TRUE then plot text labels. Passed to factoextra::fviz_cluster .
.ellipse	If TRUE then plot ellipses around all samples. Passed to "ellipse" from factoextra::fviz_cluster .
.point.size	Size of points, passed to "pointsize" from factoextra::fviz_cluster .
.text.size	Size of text labels, passed to labelsize from factoextra::fviz_cluster .
.plot	A character vector of length one or two specifying which plots to visualise. If "clust" then plot only the clustering. If "best" then plot the number of optimal clusters. If both then plot both.
...	Not used here.

Value

Ggplot2 objects inside the pathwork container.

See Also

[vis](#), [repOverlapAnalysis](#), [geneUsageAnalysis](#)

Examples

```
## Not run:
data(immdata)
ov <- repOverlap(immdata$data)
repOverlapAnalysis(ov, "mds+kmeans") %>% vis()

## End(Not run)
```

vis.immunr_kmer_table *Most frequent kmers visualisation.*

Description**[Deprecated]**

Plot a distribution (bar plot) of the most frequent kmers in a data.

Usage

```
## S3 method for class 'immunr_kmer_table'
vis(
  .data,
  .head = 100,
  .position = c("stack", "dodge", "fill"),
  .log = FALSE,
  ...
)
```

Arguments

.data	Data frame with two columns "Kmers" and "Count" or a list with such data frames. See Examples.
.head	Number of the most frequent kmers to choose for plotting from each data frame.
.position	Character vector of length 1. Position of bars for each kmers. Value for the ggplot2 argument position.
.log	Logical. If TRUE then plot log-scaled plots.
...	Not used here.

Value

A ggplot2 object.

See Also

get.kmers

Examples

```
## Not run:
# Load necessary data and package.
data(immdata)
# Get 5-mers.
imm.km <- getKmers(immdata$data[[1]], 5)
# Plots for kmer proportions in each data frame in immdata.
p1 <- vis(imm.km, .position = "stack")
p2 <- vis(imm.km, .position = "fill")
p1 + p2

## End(Not run)
```

vis.immunr_mds

PCA / MDS / tSNE visualisation (mainly overlap / gene usage)

Description

[Deprecated]

Usage

```
## S3 method for class 'immunr_mds'
vis(
  .data,
  .by = NA,
  .meta = NA,
  .point = TRUE,
  .text = TRUE,
  .ellipse = TRUE,
  .point.size = 2,
  .text.size = 4,
  ...
)
```

Arguments

.data Output from analysis functions such as [geneUsageAnalysis](#) or [immunr_pca](#), [immunr_mds](#) or [immunr_tsne](#).

.by	Pass NA if you want to plot samples without grouping. You can pass a character vector with one or several column names from ".meta" to group your data before plotting. In this case you should provide ".meta". You can pass a character vector that exactly matches the number of samples in your data, each value should correspond to a sample's property. It will be used to group data based on the values provided. Note that in this case you should pass NA to ".meta".
.meta	A metadata object. An R dataframe with sample names and their properties, such as age, serostatus or hla.
.point	Logical. If TRUE then plot points corresponding to objects.
.text	Logical. If TRUE then plot sample names.
.ellipse	Logical. If TRUE then plot ellipses around clusters of grouped samples.
.point.size	Numeric. A size of points to plot.
.text.size	Numeric. A size of sample names' labels.
...	Not used here.

Details

Other visualisation methods:

- PCA - [vis.immunr_pca](#)
- MDS - [vis.immunr_mds](#)
- tSNE - [vis.immunr_tsne](#)

Value

A ggplot2 object.

Examples

```
## Not run:
data(immdata)
ov <- repOverlap(immdata$data)
repOverlapAnalysis(ov, "mds") %>% vis()

## End(Not run)
```

vis.immunr_ov_matrix *Repertoire overlap and gene usage visualisations*

Description

[Deprecated]

Visualises matrices with overlap values or gene usage distances among samples. For details see the links below.

Usage

```
## S3 method for class 'immunr_ov_matrix'
vis(.data, .plot = c("heatmap", "heatmap2", "circos"), ...)
```

Arguments

<code>.data</code>	Output from repOverlap or geneUsageAnalysis .
<code>.plot</code>	A string specifying the plot type: • "heatmap" for heatmaps using vis_heatmap; • "heatmap2" for heatmaps using vis_heatmap2; • "circos" for circos plots using vis_circos;
<code>...</code>	Other arguments are passed through to the underlying plotting function: • "heatmap" - passes arguments to vis_heatmap; • "heatmap2" - passes arguments to vis_heatmap2 and heatmap from the "pheatmap" package; • "circos" - passes arguments to vis_circos and circlize::chordDiagram from the "circlize" package;

Value

A ggplot2, pheatmap or circlize object.

Examples

```
## Not run:
data(immdata)
ov <- repOverlap(immdata$data)
vis(ov)
vis(ov, "heatmap")
vis(ov, "heatmap2")
vis(ov, "circos")

## End(Not run)
```

```
vis.immunr_public_repertoire
      Public repertoire visualisation
```

Description

[Deprecated]

Usage

```
## S3 method for class 'immunr_public_repertoire'
vis(.data, .plot = c("freq", "clonotypes"), ...)
```

Arguments

- | | |
|-------|--|
| .data | Public repertoire, an output from pubRep . |
| .plot | A string specifying the plot type: <ul style="list-style-type: none"> • "freq" for visualisation of the distribution of occurrences of clonotypes and their frequencies using vis_public_frequencies. • "clonotypes" for visualisation of public clonotype frequency correlations between pairs of samples using vis_public_clonotypes |
| ... | Further arguments passed vis_public_frequencies or vis_public_clonotypes , depending on the ".plot" argument. |

Value

A ggplot2 object.

Examples

```
## Not run:
data(immdata)
immdata$data <- lapply(immdata$data, head, 300)
pr <- pubRep(immdata$data, .verbose = FALSE)
vis(pr, "freq")
vis(pr, "freq", .type = "none")

vis(pr, "clonotypes", 1, 2)

## End(Not run)
```

vis.immunr_public_statistics

Visualise sharing of clonotypes among samples

Description**[Deprecated]**

Visualise public clonotype frequencies.

Usage

```
## S3 method for class 'immunr_public_statistics'
vis(.data, ...)
```

Arguments

- | | |
|-------|---|
| .data | Public repertoire - an output from the pubRep function. |
| ... | Other arguments passed directly to UpSetR::upset . |

Value

A ggplot2 object.

Examples

```
## Not run:
data(immdata)
immdata$data <- lapply(immdata$data, head, 2000)
pr <- pubRep(immdata$data, .verbose = FALSE)
pubRepStatistics(pr) %>% vis()

## End(Not run)
```

vis.step_failure_ignored

Handler for .nofail argument of pipeline steps that prevents examples from crashing on computers where certain dependencies are not installed

Description

[Deprecated]

Usage

```
## S3 method for class 'step_failure_ignored'
vis(.data, ...)
```

Arguments

.data	Not used here.
...	Not used here.

Value

An empty object with "step_failure_ignored" class.

vis_bar

*Bar plots***Description****[Deprecated]****Usage**

```
vis_bar(
  .data,
  .by = NA,
  .meta = NA,
  .errorbars = c(0.025, 0.975),
  .errorbars.off = FALSE,
  .stack = FALSE,
  .points = TRUE,
  .test = TRUE,
  .signif.label.size = 3.5,
  .errorbar.width = 0.2,
  .defgroupby = "Sample",
  .grouping.var = "Group",
  .labs = c("X", "Y"),
  .title = "Barplot (.title argument)",
  .subtitle = "Subtitle (.subtitle argument)",
  .legend = NA,
  .leg.title = "Legend (.leg.title argument)",
  .legend.pos = "right",
  .rotate_x = 90
)
```

Arguments

.data	Data to visualise.
.by	Pass NA if you want to plot samples without grouping. You can pass a character vector with one or several column names from ".meta" to group your data before plotting. In this case you should provide ".meta". You can pass a character vector that exactly matches the number of samples in your data, each value should correspond to a sample's property. It will be used to group data based on the values provided. Note that in this case you should pass NA to ".meta".
.meta	A metadata object. An R dataframe with sample names and their properties, such as age, serostatus or hla.
.errorbars	A numeric vector of length two with quantiles for error bars on sectors. Disabled if ".errorbars.off" is TRUE.

<code>.errorbars.off</code>	If TRUE then plot CI bars for distances between each group. Disabled if no group passed to the ".by" argument.
<code>.stack</code>	If TRUE and <code>.errorbars.off</code> is TRUE then plot stacked bar plots for each Group or Sample
<code>.points</code>	A logical value defining whether points will be visualised or not.
<code>.test</code>	A logical vector whether statistical tests should be applied. See "Details" for more information.
<code>.signif.label.size</code>	An integer value defining the size of text for p-value.
<code>.errorbar.width</code>	Numeric. Width for error bars.
<code>.defgroupby</code>	A name for the column with sample names.
<code>.grouping.var</code>	A name for the column to group by.
<code>.labs</code>	A character vector of length two specifying names for x-axis and y-axis.
<code>.title</code>	The text for the plot's title.
<code>.subtitle</code>	The text for the plot's subtitle.
<code>.legend</code>	If TRUE then displays a legend, otherwise removes legend from the plot.
<code>.leg.title</code>	The text for the plots's legend. Provide NULL to remove the legend's title completely.
<code>.legend.pos</code>	Positions of the legend: either "top", "bottom", "left" or "right".
<code>.rotate_x</code>	How much the x tick text should be rotated? In angles.

Value

A ggplot2 object.

Examples

```
## Not run:
vis_bar(data.frame(Sample = c("A", "B", "C"), Value = c(1, 2, 3)))

## End(Not run)
```

vis_box

Flexible box-plots for visualisation of distributions

Description

[Deprecated]

Visualisation of distributions using ggplot2-based boxplots.

Usage

```
vis_box(
  .data,
  .by = NA,
  .meta = NA,
  .melt = TRUE,
  .points = TRUE,
  .test = TRUE,
  .signif.label.size = 3.5,
  .defgroupby = "Sample",
  .grouping.var = "Group",
  .labs = c("X", "Y"),
  .title = "Boxplot (.title argument)",
  .subtitle = "Subtitle (.subtitle argument)",
  .legend = NA,
  .leg.title = "Legend (.leg.title argument)",
  .legend.pos = "right"
)
```

Arguments

<code>.data</code>	Input matrix or data frame.
<code>.by</code>	Pass NA if you want to plot samples without grouping. You can pass a character vector with one or several column names from ".meta" to group your data before plotting. In this case you should provide ".meta". You can pass a character vector that exactly matches the number of samples in your data, each value should correspond to a sample's property. It will be used to group data based on the values provided. Note that in this case you should pass NA to ".meta".
<code>.meta</code>	A metadata object. An R dataframe with sample names and their properties, such as age, serostatus or hla.
<code>.melt</code>	If TRUE then apply reshape2::melt to the ".data" before plotting. In this case ".data" is supposed to be a data frame with the first character column reserved for names of genes and other numeric columns reserved to counts or frequencies of genes. Each numeric column should be associated with a specific repertoire sample.
<code>.points</code>	A logical value defining whether points will be visualised or not.
<code>.test</code>	A logical vector whether statistical tests should be applied. See "Details" for more information.
<code>.signif.label.size</code>	An integer value defining the size of text for p-value.
<code>.defgroupby</code>	A name for the column with sample names.
<code>.grouping.var</code>	A name for the column to group by.
<code>.labs</code>	Character vector of length two with names for x-axis and y-axis, respectively.
<code>.title</code>	The text for the title of the plot.

.subtitle	The The text for the plot's subtitle.
.legend	If TRUE then displays a legend, otherwise removes legend from the plot.
.leg.title	The The text for the plots's legend. Provide NULL to remove the legend's title completely.
.legend.pos	Positions of the legend: either "top", "bottom", "left" or "right".

Value

A ggplot2 object.

See Also

[vis.immunr_gene_usage](#), [geneUsage](#)

Examples

```
## Not run:
vis_box(data.frame(Sample = sample(c("A", "B", "C"), 100, TRUE), Value = rnorm(100)), .melt = FALSE)

## End(Not run)
```

vis_circos

Visualisation of matrices using circos plots

Description**[Deprecated]**

Visualise matrices with the [circlize::chordDiagram](#) function from the circlize package.

Usage

```
vis_circos(.data, .title = NULL, ...)
```

Arguments

.data	Input matrix.
.title	The The text for the title of the plot.
...	Other arguments passed to circlize::chordDiagram from the 'circlize' package.

Value

A circlize object.

See Also

[vis](#), [repOverlap](#).

Examples

```
## Not run:
data(immdata)
ov <- repOverlap(immdata$data)
vis(ov, .plot = "circos")

## End(Not run)
```

vis_heatmap	<i>Visualisation of matrices and data frames using ggplot2-based heatmaps</i>
-------------	---

Description**[Deprecated]**

Fast and easy visualisations of matrices or data frames with functions based on the ggplot2 package.

Usage

```
vis_heatmap(
  .data,
  .text = TRUE,
  .scientific = FALSE,
  .signif.digits = 2,
  .text.size = 4,
  .axis.text.size = NULL,
  .labs = c("Sample", "Sample"),
  .title = "Overlap",
  .leg.title = "Overlap values",
  .legend = TRUE,
  .na.value = NA,
  .transpose = FALSE,
  ...
)
```

Arguments

<code>.data</code>	Input object: a matrix or a data frame. If matrix: column names and row names (if presented) will be used as names for labs. If data frame: the first column will be used for row names and removed from the data. Other columns will be used for values in the heatmap.
<code>.text</code>	If TRUE then plots values in the heatmap cells. If FALSE does not plot values, just plot coloured cells instead.
<code>.scientific</code>	If TRUE then uses the scientific notation for numbers (e.g., "2.0e+2").
<code>.signif.digits</code>	Number of significant digits to display on plot.

<code>.text.size</code>	Size of text in the cells of heatmap.
<code>.axis.text.size</code>	Size of text on the axis labels.
<code>.labs</code>	A character vector of length two with names for x-axis and y-axis, respectively.
<code>.title</code>	The The text for the plot's title.
<code>.leg.title</code>	The The text for the plots's legend. Provide NULL to remove the legend's title completely.
<code>.legend</code>	If TRUE then displays a legend, otherwise removes the legend from the plot.
<code>.na.value</code>	Replace NA values with this value. By default they remain NA.
<code>.transpose</code>	Logical. If TRUE then switch rows and columns.
<code>...</code>	Other passed arguments.

Value

A ggplot2 object.

See Also

[vis](#), [repOverlap](#).

Examples

```
## Not run:
data(immdata)
ov <- repOverlap(immdata$data)
vis_heatmap(ov)
gu <- geneUsage(immdata$data, "hs.trbj")
vis_heatmap(gu)

## End(Not run)
```

vis_heatmap2

Visualisation of matrices using pheatmap-based heatmaps

Description**[Deprecated]**

Visualise matrices with the functions based on the [pheatmap](#) package with minimum amount of arguments.

Usage

```
vis_heatmap2(  
  .data,  
  .meta = NA,  
  .by = NA,  
  .title = NA,  
  .color = colorRampPalette(c("#67001f", "#d6604d", "#f7f7f7", "#4393c3",  
    "#053061"))(1024),  
  ...  
)
```

Arguments

<code>.data</code>	Input matrix. Column names and row names (if presented) will be used as names for labs.
<code>.meta</code>	A metadata object. An R dataframe with sample names and their properties, such as age, serostatus or hla.
<code>.by</code>	Set NA if you want to plot samples without grouping.
<code>.title</code>	The text for the plot's title (same as the "main" argument in pheatmap).
<code>.color</code>	A vector specifying the colors (same as the "color" argument in pheatmap). Pass NA to use the default pheatmap colors.
<code>...</code>	Other arguments for the pheatmap function.

Value

A pheatmap object.

See Also

[vis](#), [repOverlap](#)

Examples

```
## Not run:  
data(immdata)  
ov <- repOverlap(immdata$data)  
vis_heatmap2(ov)  
  
## End(Not run)
```

vis_hist

*Visualisation of distributions using histograms***Description****[Deprecated]**

Visualisation of distributions using ggplot2-based histograms.

Usage

```
vis_hist(
  .data,
  .by = NA,
  .meta = NA,
  .title = "Gene usage",
  .ncol = NA,
  .points = TRUE,
  .test = TRUE,
  .coord.flip = FALSE,
  .grid = FALSE,
  .labs = c("Gene", NA),
  .melt = TRUE,
  .legend = NA,
  .add.layer = NULL,
  ...
)
```

Arguments

<code>.data</code>	Input matrix or data frame.
<code>.by</code>	Pass NA if you want to plot samples without grouping. You can pass a character vector with one or several column names from ".meta" to group your data before plotting. In this case you should provide ".meta". You can pass a character vector that exactly matches the number of samples in your data, each value should correspond to a sample's property. It will be used to group data based on the values provided. Note that in this case you should pass NA to ".meta".
<code>.meta</code>	A metadata object. An R dataframe with sample names and their properties, such as age, serostatus or hla.
<code>.title</code>	The text for the title of the plot.
<code>.ncol</code>	A number of columns to display. Provide NA (by default) if you want the function to automatically detect the optimal number of columns.
<code>.points</code>	A logical value defining whether points will be visualised or not.
<code>.test</code>	A logical vector whether statistical tests should be applied. See "Details" for more information.

<code>.coord.flip</code>	If TRUE then swap x- and y-axes.
<code>.grid</code>	If TRUE then plot separate visualisations for each sample.
<code>.labs</code>	A character vector of length two with names for x-axis and y-axis, respectively.
<code>.melt</code>	If TRUE then apply reshape2::melt to the ".data" before plotting. In this case ".data" is supposed to be a data frame with the first character column reserved for names of genes and other numeric columns reserved to counts or frequencies of genes. Each numeric column should be associated with a specific repertoire sample.
<code>.legend</code>	If TRUE then plots the legend. If FALSE removes the legend from the plot. If NA automatically detects the best way to display legend.
<code>.add.layer</code>	Additional ggplot2 layers, that added to each plot in the output plot or grid of plots.
<code>...</code>	Is not used here.

Details

If data is grouped, then statistical tests for comparing means of groups will be performed, unless `.test = FALSE` is supplied. In case there are only two groups, the Wilcoxon rank sum test (https://en.wikipedia.org/wiki/Wilcoxon_signed-rank_test) is performed (R function [wilcox.test\(\)](#) with an argument `exact = FALSE`) for testing if there is a difference in mean rank values between two groups. In case there more than two groups, the Kruskal-Wallis test (https://en.wikipedia.org/wiki/Kruskal%E2%80%93Wallis_test) is performed (R function [kruskal.test\(\)](#)), that is equivalent to ANOVA for ranks and it tests whether samples from different groups originated from the same distribution. A significant Kruskal-Wallis test indicates that at least one sample stochastically dominates one other sample. Adjusted for multiple comparisons P-values are plotted on the top of groups. P-value adjusting is done using the Holm method (https://en.wikipedia.org/wiki/Holm%E2%80%93Bonferroni_method) (also known as Holm-Bonferroni correction). You can execute the command `?p.adjust` in the R console to see more.

Value

A ggplot2 object.

See Also

[vis.immunr_gene_usage](#), [geneUsage](#)

Examples

```
## Not run:
data(immdata)
imm_gu <- geneUsage(immdata$data[[1]])
vis(imm_gu,
    .plot = "hist", .add.layer =
      theme(axis.text.x = element_text(angle = 75, vjust = 1))
)
imm_gu <- geneUsage(immdata$data[1:4])
vis(imm_gu,
```

```

    .plot = "hist", .grid = TRUE, .add.layer =
      theme(axis.text.x = element_text(angle = 75, vjust = 1))
  )

  ## End(Not run)

```

```

vis_immunr_kmer_profile_main
  Visualise kmer profiles

```

Description

#' [Deprecated]

Usage

```
vis_immunr_kmer_profile_main(.data, .plot, ...)
```

Arguments

.data	Kmer data, an output from kmer_profile .
.plot	String specifying the plot type: "seqlogo" for traditional sequence logo plots using vis_seqlogo; "textlogo" for modified approach to sequence logo plots via text labels using vis_textlogo;
...	Other arguments passed to vis_textlogo or vis_seqlogo , depending on the ".plot" argument.

Value

A ggplot2 object.

Examples

```

## Not run:
data(immdata)
getKmers(immdata$data[[1]], 5) %>%
  kmer_profile() %>%
  vis("seqlogo")

## End(Not run)

```

vis_public_clonotypes *Visualisation of public clonotypes*

Description**[Deprecated]**

Visualise correlation of public clonotype frequencies in pairs of repertoires.

Usage

```
vis_public_clonotypes(
  .data,
  .x.rep = NA,
  .y.rep = NA,
  .title = NA,
  .ncol = 3,
  .point.size.modif = 1,
  .cut.axes = TRUE,
  .density = TRUE,
  .lm = TRUE,
  .radj.size = 3.5
)
```

Arguments

<code>.data</code>	Public repertoire data - an output from the pubRep function.
<code>.x.rep</code>	Either indices of samples or character vector of sample names for the x-axis. Must be of the same length as ".y.rep".
<code>.y.rep</code>	Either indices of samples or character vector of sample names for the y-axis. Must be of the same length as ".x.rep".
<code>.title</code>	The text for the title of the plot.
<code>.ncol</code>	An integer number of columns to print in the grid of pairs of repertoires.
<code>.point.size.modif</code>	An integer value that is a modifier of the point size. The larger the number, the larger the points.
<code>.cut.axes</code>	If TRUE then axes limits become shorter.
<code>.density</code>	If TRUE then displays density plot for distributions of clonotypes for each sample. If FALSE then removes density plot from the visualisation.
<code>.lm</code>	If TRUE then fit a linear model and displays an R adjusted coefficient that shows how similar samples are in terms of shared clonotypes.
<code>.radj.size</code>	An integer value, that defines the size of the The text for the R adjusted coefficient.

Value

A ggplot2 object.

See Also

[pubRep](#), [vis.immunr_public_repertoire](#)

Examples

```
## Not run:
data(immdata)
pr <- pubRep(immdata$data, .verbose = FALSE)
vis(pr, "clonotypes", 1, 2)

## End(Not run)
```

vis_public_frequencies

Public repertoire visualisation

Description**[Deprecated]**

Visualise public clonotype frequencies.

Usage

```
vis_public_frequencies(
  .data,
  .by = NA,
  .meta = NA,
  .type = c("boxplot", "none", "mean")
)
```

Arguments

.data	Public repertoire - an output from the pubRep function.
.by	Pass NA if you want to plot samples without grouping. You can pass a character vector with one or several column names from ".meta" to group your data before plotting. In this case you should provide ".meta". You can pass a character vector that exactly matches the number of samples in your data, each value should correspond to a sample's property. It will be used to group data based on the values provided. Note that in this case you should pass NA to ".meta".
.meta	A metadata object. An R dataframe with sample names and their properties, such as age, serostatus or hla.
.type	Character. Either "boxplot" for plotting distributions of frequencies, "none" for plotting everything, or "mean" for plotting average values only.

Value

A ggplot2 object.

Examples

```
## Not run:
data(immdata)
immdata$data <- lapply(immdata$data, head, 500)
pr <- pubRep(immdata$data, .verbose = FALSE)
vis(pr, "freq", .type = "boxplot")
vis(pr, "freq", .type = "none")
vis(pr, "freq", .type = "mean")
vis(pr, "freq", .by = "Status", .meta = immdata$meta)

## End(Not run)
```

vis_textlogo

Sequence logo plots for amino acid profiles.

Description**[Deprecated]**

Plot sequence logo plots for visualising of amino acid motif sequences / profiles.

vis_textlogo plots sequences in a text format - each letter has the same height. Useful when there are no big differences between occurrences of amino acids in the motif.

vis_seqlogo is a traditional sequence logo plots. Useful when there are one or two amino acids with clear differences in their occurrences.

Usage

```
vis_textlogo(.data, .replace.zero.with.na = TRUE, .width = 0.1, ...)
```

```
vis_seqlogo(.data, .scheme = "chemistry", ...)
```

Arguments

.data	Output from the kmer.profile function.
.replace.zero.with.na	if TRUE then replace all zeros with NAs, therefore letters with zero frequency wont appear at the plot.
.width	Width for jitter, i.e., how much points will scatter around the vertical line. Pass 0 (zero) to plot points on the straight vertical line for each position.
...	Not used here.
.scheme	Character. An argument passed to geom_logo from ggseqlogo package specifying how to colour symbols.

Value

A ggplot2 object.

See Also

[getKmers](#), [kmer_profile](#)

Examples

```
## Not run:
data(immdata)
kmers <- getKmers(immdata$data[[1]], 5)
ppm <- kmer_profile(kmers, "prob")
vis(ppm, .plot = "text")
vis(ppm, .plot = "seq")

d <- kmer_profile(c("CASLL", "CASSQ", "CASGL"))
vis_textlogo(d)
vis_seqlogo(d)

## End(Not run)
```

Index

- * **Clonality**
 - airr_clonality, [6](#)
 - annotate_clonality, [17](#)
- * **Diversity**
 - airr_diversity, [8](#)
- * **Key AIRR statistics**
 - airr_stats, [14](#)
- * **Public indices**
 - airr_public, [12](#)
- * **align_lineage**
 - repAlignLineage, [43](#)
- * **annotation**
 - dbAnnotate, [23](#)
 - dbLoad, [24](#)
- * **clonality**
 - repClonality, [46](#)
 - vis.immunr_clonal_prop, [80](#)
- * **datasets**
 - aa_table, [5](#)
 - bcrdata, [19](#)
 - immdata, [33](#)
 - immunarch_v1_updates, [33](#)
- * **data**
 - aa_properties, [5](#)
 - aa_table, [5](#)
 - bcrdata, [19](#)
 - gene_segments, [29](#)
 - gene_stats, [30](#)
 - immdata, [33](#)
 - immunr_data_format, [34](#)
- * **distance**
 - seqDist, [67](#)
- * **diversity**
 - repDiversity, [48](#)
 - vis.immunr_chao1, [78](#)
- * **dynamics**
 - trackClonotypes, [73](#)
 - vis.immunr_dynamics, [82](#)
- * **explore**
 - repExplore, [51](#)
 - vis.immunr_exp_vol, [83](#)
- * **filters**
 - repFilter, [52](#)
- * **fixvis**
 - fixVis, [26](#)
- * **gene_usage**
 - geneUsage, [27](#)
 - geneUsageAnalysis, [28](#)
 - vis.immunr_gene_usage, [85](#)
- * **germline**
 - repGermline, [53](#)
- * **io**
 - repLoad, [54](#)
 - repSave, [62](#)
- * **k-mers**
 - getKmers, [30](#)
 - split_to_kmers, [70](#)
- * **kmers**
 - vis.immunr_kmer_table, [89](#)
 - vis.immunr_kmer_profile_main, [104](#)
 - vis_textlogo, [107](#)
- * **migration_utility**
 - get_immunarch_news, [31](#)
 - immunarch_v1_updates, [33](#)
 - list_immunarch_news, [38](#)
- * **overlap**
 - inc_overlap, [37](#)
 - repOverlap, [57](#)
 - repOverlapAnalysis, [59](#)
 - vis.immunr_inc_overlap, [87](#)
 - vis.immunr_ov_matrix, [91](#)
- * **phylip**
 - repClonalFamily, [45](#)
 - vis.clonal_family, [76](#)
 - vis.clonal_family_tree, [77](#)
- * **post_analysis**
 - immunr_hclust, [34](#)
 - immunr_pca, [36](#)

- vis.immunr_hclust, 86
- vis.immunr_kmeans, 88
- vis.immunr_mds, 90
- * **preprocessing**
 - bunch_translate, 20
 - coding, 22
 - repSample, 60
 - top, 72
- * **pubrep**
 - public_matrix, 39
 - pubRep, 40
 - pubRepApply, 41
 - pubRepFilter, 42
 - pubRepStatistics, 43
 - vis.immunr_public_repertoire, 92
 - vis.immunr_public_statistics, 93
 - vis_public_clonotypes, 105
 - vis_public_frequencies, 106
- * **seq_cluster**
 - seqCluster, 66
- * **single_cell**
 - select_barcodes, 64
 - select_clusters, 65
- * **somatic_hypermutation**
 - repSomaticHypermutation, 63
- * **utility_private**
 - .quant_column_choice, 4
 - add_class, 5
 - check_distribution, 21
 - group_from_metadata, 32
 - has_class, 32
 - matrixdiagcopy, 39
 - set_pb, 69
 - switch_type, 71
- * **utility_public**
 - apply_symm, 19
 - entropy, 25
- * **vis**
 - spectratype, 69
 - vis_bar, 95
 - vis_box, 96
 - vis_circos, 98
 - vis_heatmap, 99
 - vis_heatmap2, 100
 - vis_hist, 102
- .quant_column_choice, 4
- AA_PROP (aa_properties), 5
- aa_prop (aa_properties), 5
- aa_properties, 5
- AA_TABLE (aa_table), 5
- aa_table, 5
- AA_TABLE_REVERSED (aa_table), 5
- add_class, 5
- add_pb (set_pb), 69
- airr_clonality, 6, 18
- airr_clonality_line (airr_clonality), 6
- airr_clonality_prop (airr_clonality), 6
- airr_clonality_rank (airr_clonality), 6
- airr_diversity, 8
- airr_diversity_chao1 (airr_diversity), 8
- airr_diversity_dxx (airr_diversity), 8
- airr_diversity_hill (airr_diversity), 8
- airr_diversity_index (airr_diversity), 8
- airr_diversity_pielou (airr_diversity), 8
- airr_diversity_shannon (airr_diversity), 8
- airr_public, 12
- airr_public_intersection (airr_public), 12
- airr_public_jaccard (airr_public), 12
- airr_stats, 14
- airr_stats_chains (airr_stats), 14
- airr_stats_genes (airr_stats), 14
- airr_stats_lengths (airr_stats), 14
- annotate_clonality, 7, 17
- annotate_clonality_prop (annotate_clonality), 17
- annotate_clonality_rank (annotate_clonality), 17
- apply_asymm (apply_symm), 19
- apply_symm, 19
- ATCHLEY (aa_properties), 5
- atchley (aa_properties), 5
- bcrdata, 19
- bunch_translate, 20
- chao1 (repDiversity), 48
- check_distribution, 21
- circize::chordDiagram, 85, 92, 98
- clonal.prop (repClonality), 46
- clonal_proportion, 61
- clonal_proportion (repClonality), 46
- clonal_space_homeostasis (repClonality), 46
- clonality (repClonality), 46

- coding, 22
- cross_entropy (entropy), 25
- data.frame, 22, 23, 27, 30, 37, 40, 44, 46, 48, 51, 54, 55, 57, 61, 64, 66, 67, 70, 72, 73
- data.table::data.table, 22, 23, 27, 30, 37, 40, 44, 46, 48, 51, 54, 57, 61, 64, 66, 67, 70, 72, 73
- dbAnnotate, 23
- dbLoad, 23, 24
- dbscan, 34, 35
- dist, 68
- diversity_eco (repDiversity), 48
- entropy, 25, 50
- exclude (repFilter), 52
- factoextra::fviz_cluster, 88
- factoextra::fviz_dend, 86
- factoextra::fviz_nbclust, 35
- factoextra::hcut, 35
- fixVis, 26, 76
- fpc::dbscan, 35
- GENE_SEGMENTS (gene_segments), 29
- gene_segments, 29
- gene_stats, 30
- genes (gene_segments), 29
- geneUsage, 27, 56, 75, 85, 86, 98, 103
- geneUsage(), 29
- geneUsageAnalysis, 28, 35, 36, 75, 86–90, 92
- geneUsageAnalysis(), 28, 29
- get.kmers (getKmers), 30
- get_aliases (geneUsage), 27
- get_genes (geneUsage), 27
- get_immunarch_news, 31
- get_immunarch_news(), 38
- getKmers, 30, 76, 108
- gini_coef (repDiversity), 48
- gini_simpson (repDiversity), 48
- group_from_metadata, 32
- has_class, 32
- hcut, 29, 34, 35, 60
- heatmap, 85, 92
- hill_numbers (repDiversity), 48
- immdata, 33
- immunarch_data_format, 22, 23, 27, 30, 37, 40, 44, 46, 48, 51, 54, 57, 61, 64, 66, 67, 70, 72, 73
- immunarch_data_format (immunr_data_format), 34
- immunarch_v1_updates, 33
- immdata::ImmunData, 7, 11, 13, 16–18
- immunr_data_format, 34, 56
- immunr_dbscan (immunr_hclust), 34
- immunr_dbscan(), 29, 60
- immunr_hclust, 34
- immunr_kmeans (immunr_hclust), 34
- immunr_mds, 90
- immunr_mds (immunr_pca), 36
- immunr_pca, 36, 90
- immunr_tsne, 29, 60, 90
- immunr_tsne (immunr_pca), 36
- immunr_tsne(), 29, 60
- inc_overlap, 37, 59
- include (repFilter), 52
- inframes (coding), 22
- interval (repFilter), 52
- inverse_simpson (repDiversity), 48
- js_div (entropy), 25
- KIDERA (aa_properties), 5
- kidera (aa_properties), 5
- kl_div (entropy), 25
- kmeans, 29, 34, 35, 60
- kmer_profile, 76, 104, 108
- kmer_profile (split_to_kmers), 70
- kruskal.test(), 79, 81, 84, 103
- lessthan (repFilter), 52
- list_immunarch_news, 38
- list_immunarch_news(), 31
- makeKmerTable (getKmers), 30
- matrixdiagcopy, 39
- morethan (repFilter), 52
- noncoding (coding), 22
- outofframes (coding), 22
- pheatmap, 100, 101
- prcomp, 36
- process_col_argument (switch_type), 71
- properties (aa_properties), 5

public_matrix, 39
 publicRepertoire (pubRep), 40
 publicRepertoireApply (pubRepApply), 41
 publicRepertoireFilter (pubRepFilter), 42
 pubRep, 39, 40, 42, 43, 75, 93, 105, 106
 pubRepApply, 41
 pubRepFilter, 42
 pubRepStatistics, 43

 rare_proportion (repClonality), 46
 rarefaction (repDiversity), 48
 repAlignLineage, 43
 repClonalFamily, 45, 75
 repClonalFamily(), 77
 repClonality, 46, 50, 75, 81
 repClonality(), 80
 repDiversity, 47, 48, 56, 75, 79
 repDiversity(), 78
 repExplore, 51, 75, 85
 repExplore(), 83, 84
 repFilter, 52
 repGermline, 53
 repLoad, 54
 repOverlap, 50, 56, 57, 75, 87, 92, 98, 100, 101
 repOverlap(), 60
 repOverlapAnalysis, 35, 36, 59, 75, 86–89
 repOverlapAnalysis(), 59
 repSample, 60
 repSave, 56, 62
 repSomaticHypermutation, 63
 reshape2::melt, 97, 103
 rmultinom, 61

 segments (gene_segments), 29
 select_barcodes, 64, 66
 select_clusters, 65, 65
 seqCluster, 66
 seqDist, 66, 67
 set_pb, 69
 spectratype, 69
 split_to_kmers, 70
 switch_type, 71

 top, 72
 top_proportion (repClonality), 46
 trackClonotypes, 73, 76, 82
 translate_bunch (bunch_translate), 20

 UpSetR::upset, 93

 vis, 59, 74, 79, 81, 85, 87, 89, 98, 100, 101
 vis(), 60
 vis.clonal_family, 75, 76
 vis.clonal_family_tree, 75, 77
 vis.immunr_chao1, 75, 78
 vis.immunr_clonal_prop, 80
 vis.immunr_dbscan (vis.immunr_kmeans), 88
 vis.immunr_div (vis.immunr_chao1), 78
 vis.immunr_dxx (vis.immunr_chao1), 78
 vis.immunr_dynamics, 76, 82
 vis.immunr_exp_clones (vis.immunr_exp_vol), 83
 vis.immunr_exp_count (vis.immunr_exp_vol), 83
 vis.immunr_exp_len (vis.immunr_exp_vol), 83
 vis.immunr_exp_vol, 52, 75, 83
 vis.immunr_gene_usage, 75, 85, 98, 103
 vis.immunr_ginisimp (vis.immunr_chao1), 78
 vis.immunr_gu_matrix (vis.immunr_ov_matrix), 91
 vis.immunr_hclust, 75, 86, 88
 vis.immunr_hill (vis.immunr_chao1), 78
 vis.immunr_homeo, 75
 vis.immunr_homeo (vis.immunr_clonal_prop), 80
 vis.immunr_inc_overlap, 75, 87
 vis.immunr_invsimp (vis.immunr_chao1), 78
 vis.immunr_kmeans, 86, 88
 vis.immunr_kmer_table, 76, 89
 vis.immunr_mds, 90, 91
 vis.immunr_ov_matrix, 75, 91
 vis.immunr_pca, 37, 91
 vis.immunr_pca (vis.immunr_mds), 90
 vis.immunr_public_repertoire, 75, 92, 106
 vis.immunr_public_statistics, 93
 vis.immunr_rarefaction (vis.immunr_chao1), 78
 vis.immunr_tail_prop (vis.immunr_clonal_prop), 80
 vis.immunr_top_prop (vis.immunr_clonal_prop), 80
 vis.immunr_tsne, 91

`vis.immunr_tsne (vis.immunr_mds)`, 90
`vis.step_failure_ignored`, 94
`vis_bar`, 95
`vis_box`, 85, 96
`vis_circos`, 76, 85, 92, 98
`vis_heatmap`, 76, 85, 92, 99
`vis_heatmap2`, 76, 85, 92, 100
`vis_hist`, 85, 102
`vis_immunr_kmer_profile_main`, 104
`vis_public_clonotypes`, 93, 105
`vis_public_frequencies`, 93, 106
`vis_seqlogo`, 76, 104
`vis_seqlogo (vis_textlogo)`, 107
`vis_textlogo`, 104, 107

`wilcox.test()`, 79, 81, 84, 103