

Package ‘tidySummarizedExperiment’

June 5, 2023

Type Package

Title Brings SummarizedExperiment to the Tidyverse

Version 1.11.0

Description

tidySummarizedExperiment is an adapter that abstracts the 'SummarizedExperiment' container in the form of tibble and allows the data manipulation, plotting and nesting using 'tidyverse'

License GPL-3

Depends R (>= 4.1.0), SummarizedExperiment

Imports tibble (>= 3.0.4), dplyr, magrittr, tidyr, ggplot2, rlang,
purrr, lifecycle, methods, plotly, utils, S4Vectors,
tidyselect, ellipsis, vctrs, pillar, stringr, cli, fansi

Suggests BiocStyle, testthat, knitr, markdown

VignetteBuilder knitr

RdMacros lifecycle

Biarch true

biocViews AssayDomain, Infrastructure, RNASeq, DifferentialExpression,
GeneExpression, Normalization, Clustering, QualityControl,
Sequencing, Transcription, Transcriptomics

Encoding UTF-8

LazyData true

RoxygenNote 7.2.3

Roxygen list(markdown = TRUE)

LazyDataCompression xz

URL <https://github.com/stemangiola/tidySummarizedExperiment>

BugReports <https://github.com/stemangiola/tidySummarizedExperiment/issues>

git_url <https://git.bioconductor.org/packages/tidySummarizedExperiment>

git_branch devel

git_last_commit bba72cc

git_last_commit_date 2023-04-25

Date/Publication 2023-06-04
Author Stefano Mangiola [aut, cre]
Maintainer Stefano Mangiola <mangiolastefano@gmail.com>

R topics documented:

as_tibble	2
bind	3
bind_rows	4
count	11
formatting	12
ggplot	13
pasilla	14
plot_ly	15
se	18
tbl_format_header	19
tidy	19
unnest	20
%>%	25
Index	26

as_tibble	<i>Coerce lists, matrices, and more to data frames</i>
-----------	--

Description

[Maturing]

as_tibble() turns a SummarizedExperiment existing object into a so-called tibble, a data frame with class tbl_df.

Arguments

- x A SummarizedExperiment
- ... This parameter includes .subset that can be set to any tidyselect expression. For example .subset = c(sample, type), or .subset = contains("PC").

Value

A tibble

Examples

```
tidySummarizedExperiment::pasilla %>%
  as_tibble()

tidySummarizedExperiment::pasilla %>%
  as_tibble(.subset = -c(condition, type))
```

bind

Efficiently bind multiple data frames by row and column

Description

This is an efficient implementation of the common pattern of `do.call(rbind, dfs)` or `do.call(cbind, dfs)` for binding many data frames into one.

Arguments

<code>...</code>	Data frames to combine. Each argument can either be a data frame, a list that could be a data frame, or a list of data frames. When row-binding, columns are matched by name, and any missing columns will be filled with NA. When column-binding, rows are matched by position, so all data frames must have the same number of rows. To match by value, not position, see mutate-joins.
<code>.id</code>	Data frame identifier. When <code>.id</code> is supplied, a new column of identifiers is created to link each row to its original data frame. The labels are taken from the named arguments to <code>bind_rows()</code>. When a list of data frames is supplied, the labels are taken from the names of the list. If no names are found a numeric sequence is used instead.
<code>add.cell.ids</code>	from SummarizedExperiment 3.0 A character vector of length($x=c(x, y)$). Appends the corresponding values to the start of each objects' cell names.

Details

The output of `bind_rows()` will contain a column if that column appears in any of the inputs.

Value

`bind_rows()` and `bind_cols()` return the same type as the first input, either a data frame, `tbl_df`, or `grouped_df`.

Examples

```
`%>%` <- magrittr::`%>%`
library(tibble)
tt <- tidySummarizedExperiment::pasilla
bind_rows(tt, tt)

num_rows <- nrow(tidySummarizedExperiment::as_tibble(tt))
tt %>% bind_cols(tibble(a=0, num_rows))
```

bind_rows

distinct

Description

`filter()` retains the rows where the conditions you provide a TRUE. Note that, unlike base subsetting with `[],` rows where the condition evaluates to NA are dropped.

`summarise()` creates a new data frame. It will have one (or more) rows for each combination of grouping variables; if there are no grouping variables, the output will have a single row summarising all observations in the input. It will contain one column for each grouping variable and one column for each of the summary statistics that you have specified.

`summarise()` and `summarize()` are synonyms.

`mutate()` adds new variables and preserves existing ones; `transmute()` adds new variables and drops existing ones. New variables overwrite existing variables of the same name. Variables can be removed by setting their value to NULL.

Rename individual variables using `new_name=old_name` syntax.

See [this repository](#) for alternative ways to perform row-wise operations.

`slice()` lets you index rows by their (integer) locations. It allows you to select, remove, and duplicate rows. It is accompanied by a number of helpers for common use cases:

- `slice_head()` and `slice_tail()` select the first or last rows.
- `slice_sample()` randomly selects rows.
- `slice_min()` and `slice_max()` select rows with highest or lowest values of a variable.

If `.data` is a [grouped_df](#), the operation will be performed on each group, so that (e.g.) `slice_head(df, n=5)` will select the first five rows in each group.

Select (and optionally rename) variables in a data frame, using a concise mini-language that makes it easy to refer to variables based on their name (e.g. `a:f` selects all columns from `a` on the left to `f` on the right). You can also use predicate functions like `is.numeric` to select variables based on their properties.

[Superseded] `sample_n()` and `sample_frac()` have been superseded in favour of [slice_sample\(\)](#). While they will not be deprecated in the near future, retirement means that we will only perform critical bug fixes, so we recommend moving to the newer alternative.

These functions were superseded because we realised it was more convenient to have two mutually exclusive arguments to one function, rather than two separate functions. This also made it to clean up a few other smaller design issues with `sample_n()/sample_frac`:

- The connection to `slice()` was not obvious.
- The name of the first argument, `tbl`, is inconsistent with other single table verbs which use `.data`.
- The size argument uses tidy evaluation, which is surprising and undocumented.
- It was easier to remove the deprecated `.env` argument.
- ... was in a suboptimal position.

`pull()` is similar to `$`. It's mostly useful because it looks a little nicer in pipes, it also works with remote data frames, and it can optionally name the output.

Usage

```
bind_rows(..., .id = NULL, add.cell.ids = NULL)
```

```
bind_cols(..., .id = NULL)
```

```
## S3 method for class 'SummarizedExperiment'
filter(.data, ..., .preserve = FALSE)
```

Arguments

<code>...</code>	For use by methods.
<code>.id</code>	Data frame identifier. When <code>.id</code> is supplied, a new column of identifiers is created to link each row to its original data frame. The labels are taken from the named arguments to <code>bind_rows()</code> . When a list of data frames is supplied, the labels are taken from the names of the list. If no names are found a numeric sequence is used instead.
<code>add.cell.ids</code>	from SummarizedExperiment 3.0 A character vector of length($x=c(x, y)$). Appends the corresponding values to the start of each objects' cell names.
<code>.data</code>	A tidySummarizedExperiment object or any data frame
<code>.preserve</code>	when FALSE (the default), the grouping structure is recalculated based on the resulting data, otherwise it is kept as is.
<code>.keep_all</code>	If TRUE, keep all variables in <code>.data</code> . If a combination of ... is not distinct, this keeps the first row of values. (See dplyr)
<code>.drop</code>	When <code>.drop=TRUE</code> , empty groups are dropped. See group_by_drop_default() for what the default value is for this argument.
<code>data</code>	Input data frame.
<code>x</code>	tbls to join. (See dplyr)
<code>y</code>	tbls to join. (See dplyr)
<code>by</code>	A character vector of variables to join by. (See dplyr)
<code>copy</code>	If <code>x</code> and <code>y</code> are not from the same data source, and <code>copy</code> is TRUE, then <code>y</code> will be copied into the same src as <code>x</code> . (See dplyr)
<code>suffix</code>	If there are non-joined duplicate variables in <code>x</code> and <code>y</code> , these suffixes will be added to the output to disambiguate them. Should be a character vector of length 2. (See dplyr)

tbl	A data.frame.
size	<tidy-select> For sample_n(), the number of rows to select. For sample_frac(), the fraction of rows to select. If tbl is grouped, size applies to each group.
replace	Sample with or without replacement?
weight	<tidy-select> Sampling weights. This must evaluate to a vector of non-negative numbers the same length as the input. Weights are automatically standardised to sum to 1.
.env	DEPRECATED.
name	An optional parameter that specifies the column to be used as names for a named vector. Specified in a similar manner as var.

Details

dplyr is not yet smart enough to optimise filtering optimisation on grouped datasets that don't need grouped calculations. For this reason, filtering is often considerably faster on `ungroup()`ed data.

`rowwise()` is used for the results of `do()` when you create list-variables. It is also useful to support arbitrary complex operations that need to be applied to each row.

Currently, rowwise grouping only works with data frames. Its main impact is to allow you to work with list-variables in `summarise()` and `mutate()` without having to use `[[1]]`. This makes `summarise()` on a rowwise tbl effectively equivalent to `plyr::ldply()`.

Slice does not work with relational databases because they have no intrinsic notion of row order. If you want to perform the equivalent operation, use `filter()` and `row_number()`.

Value

A tidySummarizedExperiment object

An object of the same type as `.data`.

- Rows are a subset of the input, but appear in the same order.
- Columns are not modified.
- The number of groups may be reduced (if `.preserve` is not `TRUE`).
- Data frame attributes are preserved.

A [grouped data frame](#), unless the combination of `...` and `add` yields a non empty set of grouping columns, a regular (ungrouped) data frame otherwise.

An object *usually* of the same type as `.data`.

- The rows come from the underlying `group_keys()`.
- The columns are a combination of the grouping keys and the summary expressions that you provide.
- If `x` is grouped by more than one variable, the output will be another [grouped_df](#) with the right-most group removed.
- If `x` is grouped by one variable, or is not grouped, the output will be a [tibble](#).
- Data frame attributes are **not** preserved, because `summarise()` fundamentally creates a new data frame.

An object of the same type as `.data`.

For `mutate()`:

- Rows are not affected.
- Existing columns will be preserved unless explicitly modified.
- New columns will be added to the right of existing columns.
- Columns given value `NULL` will be removed
- Groups will be recomputed if a grouping variable is mutated.
- Data frame attributes are preserved.

For `transmute()`:

- Rows are not affected.
- Apart from grouping variables, existing columns will be removed unless explicitly kept.
- Column order matches order of expressions.
- Groups will be recomputed if a grouping variable is mutated.
- Data frame attributes are preserved.

An object of the same type as `.data`.

- Rows are not affected.
- Column names are changed; column order is preserved
- Data frame attributes are preserved.
- Groups are updated to reflect new names.

A `tbl`

A `tbl`

A `tidySummarizedExperiment` object

A `tidySummarizedExperiment` object

A `tidySummarizedExperiment` object

A `tidySummarizedExperiment` object

An object of the same type as `.data`. The output has the following properties:

- Each row may appear 0, 1, or many times in the output.
- Columns are not modified.
- Groups are not modified.
- Data frame attributes are preserved.

An object of the same type as `.data`. The output has the following properties:

- Rows are not affected.
- Output columns are a subset of input columns, potentially with a different order. Columns will be renamed if `new_name=old_name` form is used.
- Data frame attributes are preserved.
- Groups are maintained; you can't select off grouping variables.

A `tidySummarizedExperiment` object

A vector the same size as `.data`.

Useful filter functions

- `==, >, >=` etc
- `&, |, !, xor()`
- `is.na()`
- `between(), near()`

Grouped tibbles

Because filtering expressions are computed within groups, they may yield different results on grouped tibbles. This will be the case as soon as an aggregating, lagging, or ranking function is involved. Compare this ungrouped filtering:

The former keeps rows with mass greater than the global average whereas the latter keeps rows with mass greater than the gender

average.

Because mutating expressions are computed within groups, they may yield different results on grouped tibbles. This will be the case as soon as an aggregating, lagging, or ranking function is involved. Compare this ungrouped mutate:

Methods

This function is a **generic**, which means that packages can provide implementations (methods) for other classes. See the documentation of individual methods for extra arguments and differences in behaviour.

The following methods are currently available in loaded packages:

These function are **generics**, which means that packages can provide implementations (methods) for other classes. See the documentation of individual methods for extra arguments and differences in behaviour.

Methods available in currently loaded packages:

This function is a **generic**, which means that packages can provide implementations (methods) for other classes. See the documentation of individual methods for extra arguments and differences in behaviour.

The following methods are currently available in loaded packages:

These function are **generics**, which means that packages can provide implementations (methods) for other classes. See the documentation of individual methods for extra arguments and differences in behaviour.

Methods available in currently loaded packages:

This function is a **generic**, which means that packages can provide implementations (methods) for other classes. See the documentation of individual methods for extra arguments and differences in behaviour.

The following methods are currently available in loaded packages:

These function are **generics**, which means that packages can provide implementations (methods) for other classes. See the documentation of individual methods for extra arguments and differences in behaviour.

Methods available in currently loaded packages:

- `slice()`: no methods found.
- `slice_head()`: no methods found.
- `slice_tail()`: no methods found.
- `slice_min()`: no methods found.
- `slice_max()`: no methods found.
- `slice_sample()`: no methods found.

This function is a **generic**, which means that packages can provide implementations (methods) for other classes. See the documentation of individual methods for extra arguments and differences in behaviour.

The following methods are currently available in loaded packages: no methods found.

This function is a **generic**, which means that packages can provide implementations (methods) for other classes. See the documentation of individual methods for extra arguments and differences in behaviour.

The following methods are currently available in loaded packages: no methods found.

Useful functions

- Center: `mean()`, `median()`
- Spread: `sd()`, `IQR()`, `mad()`
- Range: `min()`, `max()`, `quantile()`
- Position: `first()`, `last()`, `nth()`,
- Count: `n()`, `n_distinct()`
- Logical: `any()`, `all()`

Backend variations

The data frame backend supports creating a variable and using it in the same summary. This means that previously created summary variables can be further transformed or combined within the summary, as in `mutate()`. However, it also means that summary variables with the same names as previous variables overwrite them, making those variables unavailable to later summary variables.

This behaviour may not be supported in other backends. To avoid unexpected results, consider using new names for your summary variables, especially when creating multiple summaries.

Useful mutate functions

- `+`, `-`, `log()`, etc., for their usual mathematical meanings
- `lead()`, `lag()`
- `dense_rank()`, `min_rank()`, `percent_rank()`, `row_number()`, `cume_dist()`, `ntile()`
- `cumsum()`, `cummean()`, `cummin()`, `cummax()`, `cumany()`, `cumall()`
- `na_if()`, `coalesce()`
- `if_else()`, `recode()`, `case_when()`

Scoped selection and renaming

Use the three scoped variants ([rename_all\(\)](#), [rename_if\(\)](#), [rename_at\(\)](#)) to renaming a set of variables with a function.

See Also

[filter_all\(\)](#), [filter_if\(\)](#) and [filter_at\(\)](#).

Examples

```
`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  distinct(.sample)

`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  filter(.sample == "untrt1")

# Learn more in ?dplyr_tidy_eval
`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  group_by(.sample)
`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  summarise(mean(counts))
`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  mutate(logcounts=log2(counts))

`%>%` <- magrittr::`%>%`
# tidySummarizedExperiment::pasilla %>%
#
#   rename(cond=condition)
`%>%` <- magrittr::`%>%`
`%>%` <- magrittr::`%>%`

tt <- tidySummarizedExperiment::pasilla
tt %>% left_join(tt %>% distinct(condition) %>% mutate(new_column=1:2))
`%>%` <- magrittr::`%>%`

tt <- tidySummarizedExperiment::pasilla
tt %>% inner_join(tt %>% distinct(condition) %>% mutate(new_column=1:2) %>% slice(1))

`%>%` <- magrittr::`%>%`

tt <- tidySummarizedExperiment::pasilla
tt %>% right_join(tt %>% distinct(condition) %>% mutate(new_column=1:2) %>% slice(1))
```

```

`%>%` <- magrittr::`%>%`

tt <- tidySummarizedExperiment::pasilla
tt %>% full_join(tibble::tibble(condition="treated", dose=10))

`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  slice(1)

`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  select(.sample, .feature, counts)

`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  sample_n(50)
tidySummarizedExperiment::pasilla %>%

  sample_frac(0.1)

`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  pull(feature)

```

count

Count observations by group

Description

`count()` lets you quickly count the unique values of one or more variables: `df %>% count(a, b)` is roughly equivalent to `df %>% group_by(a, b) %>% summarise(n=n())`. `count()` is paired with `tally()`, a lower-level helper that is equivalent to `df %>% summarise(n=n())`. Supply `wt` to perform weighted counts, switching the summary from `n=n()` to `n=sum(wt)`.

`add_count()` and `add_tally()` are equivalents to `count()` and `tally()` but use `mutate()` instead of `summarise()` so that they add a new column with group-wise counts.

Usage

```

count(
  x,
  ...,
  wt = NULL,
  sort = FALSE,

```

```

    name = NULL,
    .drop = group_by_drop_default(x)
  )

```

Arguments

<code>x</code>	A data frame, data frame extension (e.g. a tibble), or a lazy data frame (e.g. from <code>dbplyr</code> or <code>dtplyr</code>).
<code>...</code>	<data-masking> Variables to group by.
<code>wt</code>	<data-masking> Frequency weights. Can be <code>NULL</code> or a variable: • If <code>NULL</code> (the default), counts the number of rows in each group. • If a variable, computes <code>sum(wt)</code> for each group.
<code>sort</code>	If <code>TRUE</code> , will show the largest groups at the top.
<code>name</code>	The name of the new column in the output. If omitted, it will default to <code>n</code> . If there's already a column called <code>n</code> , it will error, and require you to specify the name.
<code>.drop</code>	For <code>count()</code> : if <code>FALSE</code> will include counts for empty groups (i.e. for levels of factors that don't exist in the data). Deprecated in <code>add_count()</code> since it didn't actually affect the output.

Value

An object of the same type as `.data`. `count()` and `add_count()` group transiently, so the output has the same groups as the input.

Examples

```

`%>%` <- magrittr::`%>%`
tidySummarizedExperiment::pasilla %>%

  count(.sample)

```

formatting

Printing tibbles

Description

[Maturing]

One of the main features of the `tbl_df` class is the printing:

- Tibbles only print as many rows and columns as fit on one screen, supplemented by a summary of the remaining rows and columns.
- Tibble reveals the type of each column, which keeps the user informed about whether a variable is, e.g., `<chr>` or `<fct>` (character versus factor).

Printing can be tweaked for a one-off call by calling `print()` explicitly and setting arguments like `n` and `width`. More persistent control is available by setting the options described below.

Arguments

<code>x</code>	Object to format or print.
<code>...</code>	Other arguments passed on to individual methods.
<code>n</code>	Number of rows to show. If <code>NULL</code> , the default, will print all rows if less than option <code>tibble.print_max</code> . Otherwise, will print <code>tibble.print_min</code> rows.
<code>width</code>	Width of text output to generate. This defaults to <code>NULL</code> , which means use <code>getOption("tibble.width")</code> or (if also <code>NULL</code>) <code>getOption("width")</code> ; the latter displays only the columns that fit on one screen. You can also set <code>options(tibble.width = Inf)</code> to override this default and always print all columns.
<code>n_extra</code>	Number of extra columns to print abbreviated information for, if the width is too small for the entire tibble. If <code>NULL</code> , the default, will print information about at most <code>tibble.max_extra_cols</code> extra columns.

Value

Nothing

Package options

The following options are used by the tibble and pillar packages to format and print `tbl_df` objects. Used by the formatting workhorse `trunc_mat()` and, therefore, indirectly, by `print.tbl()`.

- `tibble.print_max`: Row number threshold: Maximum number of rows printed. Set to `Inf` to always print all rows. Default: 20.
- `tibble.print_min`: Number of rows printed if row number threshold is exceeded. Default: 10.
- `tibble.width`: Output width. Default: `NULL` (use width option).
- `tibble.max_extra_cols`: Number of extra columns printed in reduced form. Default: 100.

Examples

```
library(dplyr)
pasilla %>% print()
```

ggplot

Create a new ggplot from a tidySummarizedExperiment object

Description

`ggplot()` initializes a `ggplot` object. It can be used to declare the input data frame for a graphic and to specify the set of plot aesthetics intended to be common throughout all subsequent layers unless specifically overridden.

Arguments

<code>.data</code>	Default dataset to use for plot. If not already a <code>data.frame</code> , will be converted to one by <code>fortify()</code> . If not specified, must be supplied in each layer added to the plot.
<code>mapping</code>	Default list of aesthetic mappings to use for plot. If not specified, must be supplied in each layer added to the plot.
<code>...</code>	Other arguments passed on to methods. Not currently used.
<code>environment</code>	DEPRECATED. Used prior to tidy evaluation.

Details

`ggplot()` is used to construct the initial plot object, and is almost always followed by `+` to add component to the plot. There are three common ways to invoke `ggplot()`:

The first method is recommended if all layers use the same data and the same set of aesthetics, although this method can also be used to add a layer using data from another data frame. See the first example below. The second method specifies the default data frame to use for the plot, but no aesthetics are defined up front. This is useful when one data frame is used predominantly as layers are added, but the aesthetics may vary from one layer to another. The third method initializes a skeleton `ggplot` object which is fleshed out as layers are added. This method is useful when multiple data frames are used to produce different layers, as is often the case in complex graphics.

Value

A `ggplot`

Examples

```
library(ggplot2)

tidySummarizedExperiment::pasilla %>%

  tidySummarizedExperiment::ggplot(aes(sample, counts)) +
  geom_boxplot()
```

<code>pasilla</code>	<i>Read counts of RNA-seq samples of Pasilla knock-down by Brooks et al.</i>
----------------------	--

Description

A `SummarizedExperiment` dataset containing the transcriptome information for *Drosophila Melanogaster*.

Usage

```
data(pasilla)
```

Format

containing 14599 features and 7 biological replicates.

Source

<https://bioconductor.org/packages/release/data/experiment/html/pasilla.html>

plot_ly

Initiate a plotly visualization

Description

This function maps R objects to **plotly.js**, an (MIT licensed) web-based interactive charting library. It provides abstractions for doing common things (e.g. mapping data values to fill colors (via `color`) or creating **animations** (via `frame`)) and sets some different defaults to make the interface feel more 'R-like' (i.e., closer to `plot()` and `ggplot2::qplot()`).

Usage

```
plot_ly(  
  data = data.frame(),  
  ...,  
  type = NULL,  
  name = NULL,  
  color = NULL,  
  colors = NULL,  
  alpha = NULL,  
  stroke = NULL,  
  strokes = NULL,  
  alpha_stroke = 1,  
  size = NULL,  
  sizes = c(10, 100),  
  span = NULL,  
  spans = c(1, 20),  
  symbol = NULL,  
  symbols = NULL,  
  linetype = NULL,  
  linetypes = NULL,  
  split = NULL,  
  frame = NULL,  
  width = NULL,  
  height = NULL,  
  source = "A"  
)
```

Arguments

data	A data frame (optional) or <code>crosstalk::SharedData</code> object.
...	Arguments (i.e., attributes) passed along to the trace type. See <code>schema()</code> for a list of acceptable attributes for a given trace type (by going to <code>traces -> type -> attributes</code>). Note that attributes provided at this level may override other arguments (e.g. <code>plot_ly(x=1:10, y=1:10, color=I("red"), marker=list(color="blue"))</code>).
type	A character string specifying the trace type (e.g. "scatter", "bar", "box", etc). If specified, it <i>always</i> creates a trace, otherwise
name	Values mapped to the trace's name attribute. Since a trace can only have one name, this argument acts very much like <code>split</code> in that it creates one trace for every unique value.
color	Values mapped to relevant 'fill-color' attribute(s) (e.g. <code>fillcolor</code> , <code>marker.color</code> , <code>textfont.color</code> , etc.). The mapping from data values to color codes may be controlled using <code>colors</code> and <code>alpha</code> , or avoided altogether via <code>I()</code> (e.g., <code>color=I("red")</code>). Any color understood by <code>grDevices::col2rgb()</code> may be used in this way.
colors	Either a <code>colorbrewer2.org</code> palette name (e.g. "YlOrRd" or "Blues"), or a vector of colors to interpolate in hexadecimal "#RRGGBB" format, or a color interpolation function like <code>colorRamp()</code> .
alpha	A number between 0 and 1 specifying the alpha channel applied to color. Defaults to 0.5 when mapping to <code>fillcolor</code> and 1 otherwise.
stroke	Similar to <code>color</code> , but values are mapped to relevant 'stroke-color' attribute(s) (e.g., <code>marker.line.color</code> and <code>line.color</code> for filled polygons). If not specified, stroke inherits from <code>color</code> .
strokes	Similar to <code>colors</code> , but controls the stroke mapping.
alpha_stroke	Similar to <code>alpha</code> , but applied to stroke.
size	(Numeric) values mapped to relevant 'fill-size' attribute(s) (e.g., <code>marker.size</code> , <code>textfont.size</code> , and <code>error_x.width</code>). The mapping from data values to symbols may be controlled using <code>sizes</code> , or avoided altogether via <code>I()</code> (e.g., <code>size=I(30)</code>).
sizes	A numeric vector of length 2 used to scale size to pixels.
span	(Numeric) values mapped to relevant 'stroke-size' attribute(s) (e.g., <code>marker.line.width</code> , <code>line.width</code> for filled polygons, and <code>error_x.thickness</code>) The mapping from data values to symbols may be controlled using <code>spans</code> , or avoided altogether via <code>I()</code> (e.g., <code>span=I(30)</code>).
spans	A numeric vector of length 2 used to scale span to pixels.
symbol	(Discrete) values mapped to <code>marker.symbol</code> . The mapping from data values to symbols may be controlled using <code>symbols</code> , or avoided altogether via <code>I()</code> (e.g., <code>symbol=I("pentagon")</code>). Any <code>pch</code> value or <code>symbol name</code> may be used in this way.
symbols	A character vector of <code>pch</code> values or <code>symbol names</code> .
linetype	(Discrete) values mapped to <code>line.dash</code> . The mapping from data values to symbols may be controlled using <code>linetypes</code> , or avoided altogether via <code>I()</code> (e.g., <code>linetype=I("dash")</code>). Any <code>lty</code> (see <code>par</code>) value or <code>dash name</code> may be used in this way.

linetypes	A character vector of lty values or dash names
split	(Discrete) values used to create multiple traces (one trace per value).
frame	(Discrete) values used to create animation frames.
width	Width in pixels (optional, defaults to automatic sizing).
height	Height in pixels (optional, defaults to automatic sizing).
source	a character string of length 1. Match the value of this string with the source argument in <code>event_data()</code> to retrieve the event data corresponding to a specific plot (shiny apps can have multiple plots).

Details

Unless type is specified, this function just initiates a plotly object with 'global' attributes that are passed onto downstream uses of `add_trace()` (or similar). A **formula** must always be used when referencing column name(s) in data (e.g. `plot_ly(mtcars, x=~wt)`). Formulas are optional when supplying values directly, but they do help inform default axis/scale titles (e.g., `plot_ly(x=mtcars$wt)` vs `plot_ly(x=~mtcars$wt)`)

Value

A plotly

Author(s)

Carson Sievert

References

<https://plotly-r.com/overview.html>

See Also

- For initializing a plotly-geo object: `plot_geo()`
- For initializing a plotly-mapbox object: `plot_mapbox()`
- For translating a ggplot2 object to a plotly object: `ggplotly()`
- For modifying any plotly object: `layout()`, `add_trace()`, `style()`
- For linked brushing: `highlight()`
- For arranging multiple plots: `subplot()`, `crosstalk::bscols()`
- For inspecting plotly objects: `plotly_json()`
- For quick, accurate, and searchable plotly.js reference: `schema()`

Examples

```
# Plotly better not run
print("See below examples")

## Not run:
# plot_ly() tries to create a sensible plot based on the information you
# give it. If you don't provide a trace type, plot_ly() will infer one.
plot_ly(economics, x=~pop)
plot_ly(economics, x=~date, y=~pop)
# plot_ly() doesn't require data frame(s), which allows one to take
# advantage of trace type(s) designed specifically for numeric matrices
plot_ly(z=~volcano)
plot_ly(z=~volcano, type="surface")

# plotly has a functional interface: every plotly function takes a plotly
# object as it's first input argument and returns a modified plotly object
add_lines(plot_ly(economics, x=~date, y=~ unemploy / pop))

# To make code more readable, plotly imports the pipe operator from magrittr
economics %>%
  plot_ly(x=~date, y=~ unemploy / pop) %>%
  add_lines()

# Attributes defined via plot_ly() set 'global' attributes that
# are carried onto subsequent traces, but those may be over-written
plot_ly(economics, x=~date, color=I("black")) %>%
  add_lines(y=~uempmed) %>%
  add_lines(y=~psavert, color=I("red"))

# Attributes are documented in the figure reference -> https://plot.ly/r/reference
# You might notice plot_ly() has named arguments that aren't in this figure
# reference. These arguments make it easier to map abstract data values to
# visual attributes.
p <- plot_ly(iris, x=~Sepal.Width, y=~Sepal.Length)
add_markers(p, color=~Petal.Length, size=~Petal.Length)
add_markers(p, color=~Species)
add_markers(p, color=~Species, colors="Set1")
add_markers(p, symbol=~Species)
add_paths(p, linetype=~Species)

## End(Not run)
```

se

*Read counts of RNA-seq samples derived from Pasilla knock-down by
Brooks et al.*

Description

A SummarizedExperiment dataset containing the transcriptome information for *Drosophila Melanogaster*.

Usage

```
data(se)
```

Format

containing 14599 features and 7 biological replicates.

Source

<https://bioconductor.org/packages/release/data/experiment/html/pasilla.html>

tbl_format_header	<i>Format the header of a tibble</i>
-------------------	--------------------------------------

Description**[Experimental]**

For easier customization, the formatting of a tibble is split into three components: header, body, and footer. The `tbl_format_header()` method is responsible for formatting the header of a tibble.

Override this method if you need to change the appearance of the entire header. If you only need to change or extend the components shown in the header, override or extend `tbl_sum()` for your class which is called by the default method.

tidy	<i>tidy for SummarizedExperiment</i>
------	--------------------------------------

Description

DEPRECATED. Not needed any more.

Usage

```
tidy(object)
```

Arguments

object	A SummarizedExperiment object
--------	-------------------------------

Value

A tidySummarizedExperiment object

Examples

```
tidySummarizedExperiment::pasilla %>% tidy()
```

unnest

*unnest***Description**

Given a regular expression with capturing groups, `extract()` turns each group into a new column. If the groups don't match, or the input is NA, the output will be NA.

`pivot_longer()` "lengthens" data, increasing the number of rows and decreasing the number of columns. The inverse transformation is `pivot_wider()`

Learn more in `vignette("pivot")`.

`pivot_wider()` "widens" data, increasing the number of columns and decreasing the number of rows. The inverse transformation is [pivot_longer\(\)](#).

Learn more in `vignette("pivot")`.

Convenience function to paste together multiple columns into one.

Given either a regular expression or a vector of character positions, `separate()` turns a single character column into multiple columns.

Arguments

<code>keep_empty</code>	See <code>tidyr::unnest</code>
<code>ptype</code>	See <code>tidyr::unnest</code>
<code>.drop</code>	See <code>tidyr::unnest</code>
<code>.id</code>	<code>tidyr::unnest</code>
<code>.sep</code>	<code>tidyr::unnest</code>
<code>.preserve</code>	See <code>tidyr::unnest</code>
<code>.data</code>	A tbl. (See <code>tidyr</code>)
<code>.names_sep</code>	See <code>?tidyr::nest</code>
<code>into</code>	Names of new variables to create as character vector. Use NA to omit the variable in the output.
<code>regex</code>	a regular expression used to extract the desired values. There should be one group (defined by <code>()</code>) for each element of <code>into</code> .
<code>convert</code>	If TRUE, will run <code>type.convert()</code> with <code>as.is=TRUE</code> on new columns. This is useful if the component columns are integer, numeric or logical. NB: this will cause string "NA"s to be converted to NAs.
<code>cols</code>	<code><tidy-select></code> Columns to pivot into longer format.
<code>cols_vary</code>	When pivoting <code>cols</code> into longer format, how should the output rows be arranged relative to their original row number? "fastest", the default, keeps individual rows from <code>cols</code> close together in the output. This often produces intuitively ordered output when you have at least one key column from data that is not involved in the pivoting process.

	• "slowest" keeps individual columns from cols close together in the output. This often produces intuitively ordered output when you utilize all of the columns from data in the pivoting process.
names_to	A character vector specifying the new column or columns to create from the information stored in the column names of data specified by cols. • If length 0, or if NULL is supplied, no columns will be created. • If length 1, a single column will be created which will contain the column names specified by cols. • If length >1, multiple columns will be created. In this case, one of names_sep or names_pattern must be supplied to specify how the column names should be split. There are also two additional character values you can take advantage of: – NA will discard the corresponding component of the column name. – ".value" indicates that the corresponding component of the column name defines the name of the output column containing the cell values, overriding values_to entirely.
names_sep, names_pattern	If names_to contains multiple values, these arguments control how the column name is broken up. names_sep takes the same specification as separate(), and can either be a numeric vector (specifying positions to break on), or a single string (specifying a regular expression to split on). names_pattern takes the same specification as extract(), a regular expression containing matching groups (). If these arguments do not give you enough control, use pivot_longer_spec() to create a spec object and process manually as needed.
names_repair	What happens if the output has invalid column names? The default, "check_unique" is to error if the columns are duplicated. Use "minimal" to allow duplicates in the output, or "unique" to de-duplicated by adding numeric suffixes. See vctrs::vec_as_names() for more options.
values_to	A string specifying the name of the column to create from the data stored in cell values. If names_to is a character containing the special .value sentinel, this value will be ignored, and the name of the value column will be derived from part of the existing column names.
values_drop_na	If TRUE, will drop rows that contain only NAs in the value_to column. This effectively converts explicit missing values to implicit missing values, and should generally be used only when missing values in data were created by its structure.
names_transform, values_transform	Optionally, a list of column name-function pairs. Alternatively, a single function can be supplied, which will be applied to all columns. Use these arguments if you need to change the types of specific columns. For example, names_transform = list(week = as.integer) would convert a character variable called week to an integer. If not specified, the type of the columns generated from names_to will be character, and the type of the variables generated from values_to will be the common type of the input columns used to generate them.

names_ptypes, values_ptypes	Optionally, a list of column name-prototype pairs. Alternatively, a single empty prototype can be supplied, which will be applied to all columns. A prototype (or ptype for short) is a zero-length vector (like <code>integer()</code> or <code>numeric()</code>) that defines the type, class, and attributes of a vector. Use these arguments if you want to confirm that the created columns are the types that you expect. Note that if you want to change (instead of confirm) the types of specific columns, you should use <code>names_transform</code> or <code>values_transform</code> instead.
id_cols	<code><tidy-select></code> A set of columns that uniquely identify each observation. Typically used when you have redundant variables, i.e. variables whose values are perfectly correlated with existing variables. Defaults to all columns in data except for the columns specified through <code>names_from</code> and <code>values_from</code>. If a tidyselect expression is supplied, it will be evaluated on data after removing the columns specified through <code>names_from</code> and <code>values_from</code>.
id_expand	Should the values in the <code>id_cols</code> columns be expanded by <code>expand()</code> before pivoting? This results in more rows, the output will contain a complete expansion of all possible values in <code>id_cols</code>. Implicit factor levels that aren't represented in the data will become explicit. Additionally, the row values corresponding to the expanded <code>id_cols</code> will be sorted.
names_from, values_from	<code><tidy-select></code> A pair of arguments describing which column (or columns) to get the name of the output column (<code>names_from</code>), and which column (or columns) to get the cell values from (<code>values_from</code>). If <code>values_from</code> contains multiple values, the value will be added to the front of the output column.
names_sep	If <code>names_from</code> or <code>values_from</code> contains multiple variables, this will be used to join their values together into a single string to use as a column name.
names_prefix	String added to the start of every variable name. This is particularly useful if <code>names_from</code> is a numeric vector and you want to create syntactic variable names.
names_glue	Instead of <code>names_sep</code> and <code>names_prefix</code>, you can supply a glue specification that uses the <code>names_from</code> columns (and special <code>.value</code>) to create custom column names.
names_sort	Should the column names be sorted? If <code>FALSE</code>, the default, column names are ordered by first appearance.
names_vary	When <code>names_from</code> identifies a column (or columns) with multiple unique values, and multiple <code>values_from</code> columns are provided, in what order should the resulting column names be combined? • "fastest" varies <code>names_from</code> values fastest, resulting in a column naming scheme of the form: <code>value1_name1</code>, <code>value1_name2</code>, <code>value2_name1</code>, <code>value2_name2</code>. This is the default. • "slowest" varies <code>names_from</code> values slowest, resulting in a column naming scheme of the form: <code>value1_name1</code>, <code>value2_name1</code>, <code>value1_name2</code>, <code>value2_name2</code>.
names_expand	Should the values in the <code>names_from</code> columns be expanded by <code>expand()</code> before pivoting? This results in more columns, the output will contain column names

	corresponding to a complete expansion of all possible values in <code>names_from</code> . Implicit factor levels that aren't represented in the data will become explicit. Additionally, the column names will be sorted, identical to what <code>names_sort</code> would produce.
<code>values_fill</code>	Optionally, a (scalar) value that specifies what each value should be filled in with when missing. This can be a named list if you want to apply different fill values to different value columns.
<code>values_fn</code>	Optionally, a function applied to the value in each cell in the output. You will typically use this when the combination of <code>id_cols</code> and <code>names_from</code> columns does not uniquely identify an observation. This can be a named list if you want to apply different aggregations to different <code>values_from</code> columns.
<code>unused_fn</code>	Optionally, a function applied to summarize the values from the unused columns (i.e. columns not identified by <code>id_cols</code>, <code>names_from</code>, or <code>values_from</code>). The default drops all unused columns from the result. This can be a named list if you want to apply different aggregations to different unused columns. <code>id_cols</code> must be supplied for <code>unused_fn</code> to be useful, since otherwise all unspecified columns will be considered <code>id_cols</code>. This is similar to grouping by the <code>id_cols</code> then summarizing the unused columns using <code>unused_fn</code>.
<code>data</code>	A data frame.
<code>col</code>	The name of the new column, as a string or symbol. This argument is passed by expression and supports quasiquotation (you can unquote strings and symbols). The name is captured from the expression with <code>rlang::ensym()</code> (note that this kind of interface where symbols do not represent actual objects is now discouraged in the tidyverse; we support it here for backward compatibility).
<code>...</code>	<tidy-select> Columns to unite
<code>na.rm</code>	If TRUE, missing values will be removed prior to uniting each value.
<code>remove</code>	If TRUE, remove input columns from output data frame.
<code>sep</code>	Separator between columns. If character, <code>sep</code> is interpreted as a regular expression. The default value is a regular expression that matches any sequence of non-alphanumeric values. If numeric, <code>sep</code> is interpreted as character positions to split at. Positive values start at 1 at the far-left of the string; negative value start at -1 at the far-right of the string. The length of <code>sep</code> should be one less than <code>into</code>.
<code>extra</code>	If <code>sep</code> is a character vector, this controls what happens when there are too many pieces. There are three valid options: • "warn" (the default): emit a warning and drop extra values. • "drop": drop any extra values without a warning. • "merge": only splits at most <code>length(into)</code> times

fill If `sep` is a character vector, this controls what happens when there are not enough pieces. There are three valid options:

- "warn" (the default): emit a warning and fill from the right
- "right": fill with missing values on the right
- "left": fill with missing values on the left

Details

`pivot_longer()` is an updated approach to `gather()`, designed to be both simpler to use and to handle more use cases. We recommend you use `pivot_longer()` for new code; `gather()` isn't going away but is no longer under active development.

`pivot_wider()` is an updated approach to `spread()`, designed to be both simpler to use and to handle more use cases. We recommend you use `pivot_wider()` for new code; `spread()` isn't going away but is no longer under active development.

Value

A `tidySummarizedExperiment` object or a tibble depending on input

A `tidySummarizedExperiment` object or a tibble depending on input

A `tidySummarizedExperiment` object or a tibble depending on input

A `tidySummarizedExperiment` object or a tibble depending on input

A `tidySummarizedExperiment` object or a tibble depending on input

A `tidySummarizedExperiment` object or a tibble depending on input

See Also

`separate()` to split up by a separator.

`pivot_wider_spec()` to pivot "by hand" with a data frame that defines a pivoting specification.

`separate()`, the complement.

`unite()`, the complement, `extract()` which uses regular expression capturing groups.

Examples

```
tidySummarizedExperiment::pasilla %>%
```

```
  nest(data=-condition) %>%
  unnest(data)
```

```
tidySummarizedExperiment::pasilla %>%
```

```
  nest(data=-condition)
```

```
tidySummarizedExperiment::pasilla %>%
```

```
  extract(type, into="sequencing", regex="([a-z]*)_end", convert=TRUE)
```

```
# See vignette("pivot") for examples and explanation

library(dplyr)
tidySummarizedExperiment::pasilla %>%

  pivot_longer(c(condition, type), names_to="name", values_to="value")

# See vignette("pivot") for examples and explanation

library(dplyr)
tidySummarizedExperiment::pasilla %>%

  pivot_wider(names_from=feature, values_from=counts)

tidySummarizedExperiment::pasilla %>%

  unite("group", c(condition, type))
```

%>%

Pipe operator

Description

See `magrittr::%>%` for details.

Usage

```
lhs %>% rhs
```

Arguments

lhs	A value or the magrittr placeholder.
rhs	A function call using the magrittr semantics.

Value

The result of calling `rhs(lhs)`.

Examples

```
library(magrittr)
1 %>% sum(2)
```

Index

- * **datasets**
 - pasilla, [14](#)
 - se, [18](#)
- * **internal**
 - %>%, [25](#)
 - bind_rows, [4](#)
- * **single table verbs**
 - bind_rows, [4](#)
- +, [9](#)
- ==, [8](#)
- >, [8](#)
- >=, [8](#)
- &, [8](#)
- %>%, [25](#), [25](#)
- add_trace(), [17](#)
- all(), [9](#)
- animation, [15](#)
- any(), [9](#)
- as_tibble, [2](#)
- between(), [8](#)
- bind, [3](#)
- bind_cols (bind_rows), [4](#)
- bind_rows, [4](#)
- case_when(), [9](#)
- coalesce(), [9](#)
- count, [11](#)
- crosstalk::bscols(), [17](#)
- crosstalk::SharedData, [16](#)
- cumall(), [9](#)
- cumany(), [9](#)
- cume_dist(), [9](#)
- cummax(), [9](#)
- cummean(), [9](#)
- cummin(), [9](#)
- cumsum(), [9](#)
- dense_rank(), [9](#)
- distinct (bind_rows), [4](#)
- do(), [6](#)
- event_data(), [17](#)
- expand(), [22](#)
- extract (unnest), [20](#)
- extract(), [21](#), [24](#)
- filter (bind_rows), [4](#)
- filter(), [6](#)
- filter_all(), [10](#)
- filter_at(), [10](#)
- filter_if(), [10](#)
- first(), [9](#)
- formatting, [12](#)
- formula, [17](#)
- fortify(), [14](#)
- full_join (bind_rows), [4](#)
- gather(), [24](#)
- ggplot, [13](#)
- ggplot2::qplot(), [15](#)
- ggplotly(), [17](#)
- grDevices::col2rgb(), [16](#)
- group_by (bind_rows), [4](#)
- group_by_drop_default(), [5](#)
- grouped data frame, [6](#)
- grouped_df, [4](#), [6](#)
- highlight(), [17](#)
- I(), [16](#)
- if_else(), [9](#)
- inner_join (bind_rows), [4](#)
- IQR(), [9](#)
- is.na(), [8](#)
- lag(), [9](#)
- last(), [9](#)
- layout(), [17](#)
- lead(), [9](#)

`left_join(bind_rows)`, 4
`log()`, 9

`mad()`, 9
`max()`, 9
`mean()`, 9
`median()`, 9
`min()`, 9
`min_rank()`, 9
`mutate(bind_rows)`, 4
`mutate()`, 6, 9
`mutate-joins`, 3

`n()`, 9
`n_distinct()`, 9
`na_if()`, 9
`near()`, 8
`nest(unnest)`, 20
`nth()`, 9
`ntile()`, 9

`par`, 16
`pasilla`, 14
`pch`, 16
`percent_rank()`, 9
`pivot_longer(unnest)`, 20
`pivot_longer()`, 20
`pivot_wider(unnest)`, 20
`pivot_wider_spec()`, 24
`plot()`, 15
`plot_geo()`, 17
`plot_ly`, 15
`plot_mapbox()`, 17
`plotly_json()`, 17
`plyr::ldply()`, 6
`pull(bind_rows)`, 4

`quantile()`, 9
`quasiquote`, 23

`recode()`, 9
`rename(bind_rows)`, 4
`rename_all()`, 10
`rename_at()`, 10
`rename_if()`, 10
`right_join(bind_rows)`, 4
`rlang::ensym()`, 23
`row_number()`, 6, 9
`rowwise(bind_rows)`, 4

`sample_frac(bind_rows)`, 4
`sample_n(bind_rows)`, 4
`schema()`, 16, 17
`sd()`, 9
`se`, 18
`select(bind_rows)`, 4
`separate(unnest)`, 20
`separate()`, 21, 24
`slice(bind_rows)`, 4
`slice_sample()`, 4
`spread()`, 24
`style()`, 17
`subplot()`, 17
`summarise(bind_rows)`, 4
`summarise()`, 6

`tbl_format_header`, 19
`tbl_sum()`, 19
`tibble`, 6
`tidy`, 19
`type.convert()`, 20

`ungroup()`, 6
`unite(unnest)`, 20
`unite()`, 24
`unnest`, 20

`vctrs::vec_as_names()`, 21

`xor()`, 8