

Package ‘SQMtools’

August 27, 2026

Title Analyze Results Generated by the 'SqueezeMeta' Pipeline

Version 1.8.1

Description 'SqueezeMeta' is a versatile pipeline for the automated analysis of metagenomics/metatranscriptomics data (<<https://github.com/jtamames/SqueezeMeta>>). This package provides functions loading 'SqueezeMeta' results into R, filtering them based on different criteria, and visualizing the results using basic plots. The 'SqueezeMeta' project (and any subsets of it generated by the different filtering functions) is parsed into a single object, whose different components (e.g. tables with the taxonomic or functional composition across samples, contig/gene abundance profiles) can be easily analyzed using other R packages such as 'vegan' or 'DESeq2'. The methods in this package are further described in Puente-Sánchez et al., (2020) <[doi:10.1186/s12859-020-03703-2](https://doi.org/10.1186/s12859-020-03703-2)>.

Author Fernando Puente-Sánchez [aut, cre],
Natalia García-García [aut]

Maintainer Fernando Puente-Sánchez <fernando.puente.sanchez@slu.se>

Depends R (>= 3.5.0)

LazyDataCompression xz

Imports zip, reshape2, ggplot2, pathview, data.table, Biostrings

Suggests vegan, microeco, phyloseq

License GPL-3

Encoding UTF-8

LazyData true

BugReports <https://github.com/jtamames/SqueezeMeta/issues>

URL <https://github.com/jtamames/SqueezeMeta>

Config/roxygen2/version 8.1.0

NeedsCompilation no

Repository CRAN

Date/Publication 2026-08-27 08:10:08 UTC

Contents

CheckMProkaryote	3
combineSQM	3
combineSQMlite	5
create_bin	6
exportBins	6
exportContigs	7
exportKrona	7
exportORFs	8
exportPathway	8
exportTable	10
find_redundant_contigs	11
Hadza	12
loadSQM	13
loadSQMlite	16
MGKOs	19
MGOs	20
mostAbundant	20
mostVariable	22
plotBars	22
plotBins	24
plotFunctions	26
plotHeatmap	27
plotTaxonomy	28
RecA	30
remove_contigs_from_bin	31
renameSamples	32
rowMaxs	32
rowMins	33
seqvec2fasta	33
SQM_to_microeco	34
SQM_to_phyloseq	35
subsetBins	37
subsetContigs	39
subsetFun	40
subsetORFs	42
subsetRand	43
subsetSamples	44
subsetTax	45
summary.SQM	46
summary.SQMbunch	47
summary.SQMlite	47
USiCGs	48

Index

CheckMProkaryote*CheckM reference markers for Prokaryotes*

Description

List of Universal Single Copy Genes for Bacteria and Archaea.

Usage

data(CheckMProkaryote)

Format

List containing vectors of PFAMs, each vector corresponding to a different set of collocated markers.

Source

[doi:10.1101/gr.186072.114](https://doi.org/10.1101/gr.186072.114)

References

Parks, Imelfort, Skennerton, Hugenholtz & Tyson (2015). CheckM: assessing the quality of microbial genomes recovered from isolates, single cells, and metagenomes *Genome Res.* **25**:1043-1055. ([doi:10.1101/gr.186072.114](https://doi.org/10.1101/gr.186072.114)).

See Also

[USiCGs](#), [MGOGs](#) and [MGKOs](#) for an alternative set of single copy genes, and for examples on how to generate copy numbers.

combineSQM*Combine several SQM objects*

Description

Combine an arbitrary number of SQM objects into a single SQM object (if the input objects contain the same samples, i.e. they come from the same SqueezeMeta run) or a single SQMbunch object. For combining results from `sqm_reads.pl` or `sqm_longreads.pl` please check [combineSQMlite](#). The parameters below (other than ...) will take only effect if the input objects contain the same samples. Otherwise the input objects will be taken as they are, with no recalculation of taxonomy, function or rescaling,

Usage

```
combineSQM(
  ...,
  tax_source = "orfs",
  trusted_functions_only = FALSE,
  ignore_unclassified_functions = FALSE,
  rescale_tpm = TRUE,
  rescale_copy_number = TRUE,
  recalculate_bin_stats = TRUE
)
```

Arguments

...	an arbitrary number of SQM objects. Alternatively, a single list containing an arbitrary number of SQM objects.
tax_source	character, source data used for the taxonomy tables present in SQM\$taxa, either "orfs", "contigs", "bins" (GTDB bin taxonomy if available, SQM bin taxonomy otherwise), "bins_gtdb" (GTDB bin taxonomy) or "bins_sqm" (SQM bin taxonomy). Default "orfs". If the objects being combined contain a subset of taxa or bins, we recommend adjusting this parameter.
trusted_functions_only	logical. If TRUE, only highly trusted functional annotations (best hit + best average) will be considered when generating aggregated function tables. If FALSE, best hit annotations will be used (default FALSE).
ignore_unclassified_functions	logical. If FALSE, ORFs with no functional classification will be aggregated together into an "Unclassified" category. If TRUE, they will be ignored (default FALSE).
rescale_tpm	logical. If TRUE, TPMs for KEGGs, COGs, and PFAMs will be recalculated (so that the TPMs in the subset actually add up to 1 million). Otherwise, per-function TPMs will be calculated by aggregating the TPMs of the ORFs annotated with that function, and will thus keep the scaling present in the parent object (default TRUE).
rescale_copy_number	logical. If TRUE, copy numbers will be recalculated using the median single-copy gene coverages in the subset. Otherwise, single-copy gene coverages will be taken from the parent object. By default it is set to TRUE, which means that the returned copy numbers will represent the average copy number per function <i>in the genomes of the selected bins or contigs</i> . If any SQM objects that are being combined contain a functional subset rather than a contig/bins subset, this parameter should be set to FALSE.
recalculate_bin_stats	logical. If TRUE, bin abundance, quality and taxonomy are recalculated based on the contigs present in the subsetted object (default TRUE).

Value

A SQM or SQMbunch object

See Also

[subsetFun](#), [subsetTax](#), [combineSQMLite](#)

Examples

```
data(Hadza)
# Select Carbohydrate metabolism ORFs in Bacteroidota,
# and Amino acid metabolism ORFs in Proteobacteria
bact = subsetTax(Hadza, "phylum", "Bacteroidota")
bact.carb = subsetFun(bact, "Carbohydrate metabolism")
baci = subsetTax(Hadza, "phylum", "Bacillota")
baci.amins = subsetFun(baci, "Amino acid metabolism")
bact.carb_proteo.amins = combineSQM(bact.carb, baci.amins, rescale_copy_number=FALSE)
```

combineSQMLite

Combine several SQM or SQMLite objects

Description

Combine an arbitrary number of SQM or SQMLite objects into a single SQMLite object. This function accepts objects originating from different projects (i.e. different SqueezeMeta runs).

Usage

```
combineSQMLite(...)
```

Arguments

... an arbitrary number of SQM or SQMLite objects. Alternatively, a single list containing an arbitrary number of SQMLite objects.

Value

A SQMLite object

See Also

[combineSQM](#)

Examples

```
## Not run:
data(Hadza)
# Load data coming from a different run
other = loadSQMLite("/path/to/other/project/tables") # e.g. if the project was run using sqm_reads
# (We could also use loadSQM to load the data as long as the data comes from a SqueezeMeta run)
combined = combineSQMLite(Hadza, other)
# Now we can plot together the samples from Hadza and the second project
plotTaxonomy(combined, 'family')
```

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

create_bin	<i>Create a bin from a vector of contigs</i>
------------	--

Description

Create a bin from a vector of contigs

Usage

```
create_bin(SQM, bin, contigs, delete_overlapping_bins = FALSE)
```

Arguments

SQM	A SQM object.
bin	character. Name of the bin to be created.
contigs	character. Vector with the names of the contigs that will be included in the new bin.
delete_overlapping_bins	logical. If TRUE, bins that originally contained any of the provided contigs will be removed from the object. Default FALSE.

Value

SQM object with the new binning information, including recalculated bin statistics if possible.

See Also

[find_redundant_contigs](#), [remove_contigs_from_bin](#)

exportBins	<i>Export the bins of a SQM object</i>
------------	--

Description

Export the bins of a SQM object

Usage

```
exportBins(SQM, output_dir = "")
```

Arguments

SQM	A SQM object.
output_dir	Existing output directory to which to write the bins.

exportContigs	<i>Export the contigs of a SQM object</i>
---------------	---

Description

Export the contigs of a SQM object

Usage

```
exportContigs(SQM, output_name = "")
```

Arguments

SQM	A SQM object.
output_name	A connection, or a character string naming the file to print to. If "" (the default), sequences will be printed to the standard output connection.

exportKrona	<i>Export the taxonomy of a SQM object into a Krona Chart</i>
-------------	---

Description

Generate a krona chart containing the full taxonomy from a SQM object.

Usage

```
exportKrona(SQM, output_name = NA)
```

Arguments

SQM	A SQM, SQMbunch or SQMlite object.
output_name	character. Name of the output file containing the Krona charts in html format (default "<project_name>.krona.html").

Details

Original code was kindly provided by Giuseppe D'Auria (dauria_giu@gva.es).

Value

No return value, but a krona chart is produced in the current working directory.

See Also

[plotTaxonomy](#) for plotting the most abundant taxa of a SQM object.

Examples

```
data(Hadza)
# Check that kronatools is present.
ecode = system('ktImportText', ignore.stdout = TRUE, ignore.stderr = TRUE)
# If so, run.
if(ecode==0) { exportKrona(Hadza, output_name = file.path(tempdir(), "krona.html")) }
```

exportORFs	<i>Export the ORFs of a SQM object</i>
------------	--

Description

Export the ORFs of a SQM object

Usage

```
exportORFs(SQM, output_name = "")
```

Arguments

SQM	A SQM object.
output_name	A connection, or a character string naming the file to print to. If "" (the default), sequences will be printed to the standard output connection.

exportPathway	<i>Export the functions of a SQM object into KEGG pathway maps</i>
---------------	--

Description

This function is a wrapper for the pathview package (Luo *et al.*, 2017. *Nucleic acids research*, 45:W501-W508). It will generate annotated KEGG pathway maps showing which reactions are present in the different samples. It will also generate legends with the color scales for each sample in separate png files.

Usage

```
exportPathway(
  SQM,
  pathway_id,
  count = "copy_number",
  samples = NULL,
  split_samples = FALSE,
  sample_colors = NULL,
  log_scale = FALSE,
```

```

    fold_change_groups = NULL,
    fold_change_colors = NULL,
    max_scale_value = NULL,
    color_bins = 10,
    rescale_percent = FALSE,
    output_dir = ".",
    output_suffix = "pathview"
)

```

Arguments

SQM	A SQM, SQMbunch or SQMlite object.
pathway_id	character. The five-number KEGG pathway identifier. A list of all pathway identifiers can be found in https://www.genome.jp/kegg/pathway.html .
count	character. Either "abund" for raw abundances, "percent" for percentages, "bases" for raw base counts, "tpm" for TPM normalized values or "copy_number" for copy numbers (default "copy_number"). Note that a given count type might not be available in this object (e.g. TPM or copy number in SQMlite objects originating from a SQM reads project).
samples	character. An optional vector with the names of the samples to export. If absent, all samples will be exported (default NULL).
split_samples	logical. Generate a different output file for each sample (default FALSE).
sample_colors	character. An optional vector with the plotting colors for each sample (default NULL).
log_scale	logical. Use a base 10 logarithmic transformation for the color scale. Will have no effect if fold_change_groups is provided (default FALSE).
fold_change_groups	list. An optional list containing two vectors of samples. If provided, the function will generate a single plot displaying the log2 fold-change between the median abundances of both groups of samples ($\log(\text{second group} / \text{first group})$) (default NULL).
fold_change_colors	character. An optional vector with the plotting colors of both groups in the fold-change plot. Will be ignored if fold_change_group is not provided.
max_scale_value	numeric. Maximum value to include in the color scale. By default it is the maximum value in the selected samples (if plotting abundances in samples) or the maximum absolute log2 fold-change (if plotting fold changes) (default NULL).
color_bins	numeric. Number of bins used to generate the gradient in the color scale (default 10).
rescale_percent	logical. Calculate percent counts over the number of reads in the input object, instead of over the total number of reads in the original project (default FALSE).
output_dir	character. Directory in which to write the output files (default ".").
output_suffix	character. Suffix to be added to the output files (default "pathview").

Value

A ggplot if split_samples = FALSE and the **ggpattern** package is installed, otherwise nothing. Additionally, Pathview figures will be written in the directory specified by output_dir.

See Also

[plotFunctions](#) for plotting the most functions taxa of a SQM object.

Examples

```
data(Hadza)

exportPathway(Hadza, "00910", count = 'copy_number',
              output_dir = tempdir(),
              output_suffix = "nitrogen_metabolism",
              sample_colors = c("red", "blue"))
exportPathway(Hadza, "00250", count = 'tpm',
              output_dir = tempdir(),
              output_suffix = "ala_asp_glu_metabolism_FoldChange",
              fold_change_groups = list(c("H1"), c("H12")), max_scale_value=2)
```

exportTable	<i>Export results in tabular format</i>
-------------	---

Description

This function is a wrapper for R’s write.table function.

Usage

```
exportTable(table, output_name)
```

Arguments

table	vector, matrix or data.frame. The table to be written.
output_name	Either a character string naming a file or a connection open for writing. “” indicates output to the console.

Examples

```
data(Hadza)
Hadza.iron = subsetFun(Hadza, "iron")
# Write the taxonomic distribution at the genus level of all the genes related to iron.
exportTable(Hadza.iron$taxa$genus$percent, file.path(tempdir(), "Hadza.ironGenes.genus.tsv"))
# Now write the distribution of the different iron-related COGs
# (Clusters of Orthologous Groups) across samples.
exportTable(Hadza.iron$functions$COG$tpm, file.path(tempdir(), "Hadza.ironGenes.COG.tsv"))
```

```
# Now write all the information contained in the ORF table.
exportTable(Hadza.iron$orfs$table, file.path(tempdir(), "Hadza.ironGenes.orfTable.tsv"))
```

```
find_redundant_contigs
```

Find redundant contigs within a bin

Description

Find contigs with overlapping marker genes in a given bin, and suggest contigs to be removed in order to reduce contamination without affecting completeness. Note that this can give a quick idea of the contigs that are sources of contamination within a bin, but is not a replacement for proper bin refinement with other tools such as anvi'o.

Usage

```
find_redundant_contigs(SQM, bin, minimum_overlap_for_removal = 1)
```

Arguments

SQM	A SQM object.
bin	character. Name of the bin to be created.
minimum_overlap_for_removal	numeric. Fraction of marker genes in the contigs present in another contig needed to suggest it for removal. If set to 1 (default), contigs will only suggested for removal if their markers fully overlap with those in another contig (and thus completeness will not change after removing them). Smaller values will result in more contigs being suggested for removal, which will further reduce contamination at the expense of some completeness.

Value

A character vector with the contigs deemed to be redundant. A heatmap showing how marker genes overlap over different contigs will also be produced.

See Also

[create_bin](#), [remove_contigs_from_bin](#)

Examples

```
data(Hadza)
bin_name = "Hadza2merged.concoct.28.fa.contigs"
# Get redundant contigs that could be removed from our bin
candidates_for_removal = find_redundant_contigs(Hadza, bin_name)
# We can now remove them from the bin
Hadza.new.1 = remove_contigs_from_bin(Hadza, bin_name, candidates_for_removal)
```

```
# Or we can create a new bin out of them
# which will also remove them from the original bin
Hadza.new.2 = create_bin(Hadza, "new_bin_name", candidates_for_removal)
```

Hadza

Hadza hunter-gatherer gut metagenomes

Description

Subset of two bins (and the associated contigs and genes) generated by running SqueezeMeta on two gut metagenomic samples obtained from two hunter-gatherers of the Hadza ethnic group.

Usage

```
data(Hadza)
```

Format

A SQM object; see [loadSQM](#).

Source

[SRR1927149](#), [SRR1929485](#).

References

Rampelli *et al.*, 2015. Metagenome Sequencing of the Hadza Hunter-Gatherer Gut Microbiota. *Curr. biol.* **25**:1682-93 ([doi:10.1016/j.cub.2015.04.055](#)).

Examples

```
data(Hadza)
plotTaxonomy(Hadza, "genus", rescale=TRUE)
plotFunctions(Hadza, "COG")
```

loadSQM

*Load a SqueezeMeta project into R***Description**

This function takes the path to a project directory generated by **SqueezeMeta** (whose name is specified in the `-p` parameter of the `SqueezeMeta.pl` script) and parses the results into a SQM object. Alternatively, it can load the project data from a zip file produced by `sqm2zip.py`.

Usage

```
loadSQM(
  project_path,
  tax_mode = "prokfilter",
  tax_source = "contigs",
  trusted_functions_only = FALSE,
  single_copy_genes = "MGOGs",
  load_sequences = TRUE,
  engine = "data.table"
)
```

Arguments

project_path	character, a vector of project directories generated by SqueezeMeta, and/or zip files generated by <code>sqm2zip.py</code> .
tax_mode	character, which taxonomic classification should be loaded? SqueezeMeta applies the identity thresholds described in Luo <i>et al.</i> , 2014 (doi:10.1093/nar/gku169). Use <code>allfilter</code> for applying the minimum identity threshold to all taxa, <code>prokfilter</code> for applying the threshold to Bacteria and Archaea, but not to Eukaryotes, and <code>nofilter</code> for applying no thresholds at all (default <code>prokfilter</code>).
tax_source	character, source data used for the taxonomy tables present in <code>SQM\$taxa</code> , either <code>"orfs"</code> , <code>"contigs"</code> , <code>"bins"</code> (GTDB bin taxonomy if available, SQM bin taxonomy otherwise), <code>"bins_gtdb"</code> (GTDB bin taxonomy) or <code>"bins_sqm"</code> (SQM bin taxonomy). Default <code>"contigs"</code> .
trusted_functions_only	logical. If TRUE, only highly trusted functional annotations (best hit + best average) will be considered when generating aggregated function tables. If FALSE, best hit annotations will be used (default FALSE). Will only have an effect if <code>project_path</code> is not a zip file, and <code>project_path/results/tables</code> is not already present.
single_copy_genes	character, source of single copy genes for copy number normalization, either <code>"RecA"</code> (COG0468, RecA/RadA), <code>"MGOGs"</code> (COGs for 10 single copy and housekeeping genes, Salazar, G <i>et al.</i> 2019), <code>"MGKOs"</code> (KOs for 10 single copy and housekeeping genes, Salazar, G <i>et al.</i> , 2019) or <code>"USiCGs"</code> (KOs for 15 single copy genes, Carr <i>et al.</i> , 2013. Table S1). For <code>"MGOGs"</code> , <code>"MGKOs"</code> and <code>"USiCGs"</code> ,

	the median coverage of a set of single copy genes will be used for normalization. Default "MGOGs".
load_sequences	logical. If TRUE, contig and orf sequences will be loaded in the SQM object. Setting it to FALSE will reduce memory usage. Default TRUE.
engine	character. Engine used to load the ORFs and contigs tables. Either "data.frame" or "data.table" (significantly faster if your project is large). Default "data.table".

Value

SQM object containing the parsed project. If more than one path is provided in `project_path` this function will return a SQMbunch object instead. The structure of this object is similar to that of a SQMlite object (see [loadSQMlite](#)) but with an extra entry named `projects` that contains one SQM object for input project. SQM and SQMbunch objects will otherwise behave similarly when used with the `subset` and `plot` functions from this package.

Prerequisites

Run **SqueezeMeta**! An example call for running it would be:

```
/path/to/SqueezeMeta/scripts/SqueezeMeta.pl
-m coassembly -f fastq_dir -s samples_file -p project_dir
```

The SQM object structure

The SQM object is a nested list which contains the following information:

lvl1	lvl2	lvl3	type	rows/names	columns	data
\$orfs	\$table		<i>dataframe</i>	orfs	misc. data	misc. data
	\$abund		<i>numeric matrix</i>	orfs	samples	abundances
	\$bases		<i>numeric matrix</i>	orfs	samples	abundances
	\$cov		<i>numeric matrix</i>	orfs	samples	coverages
	\$cpm		<i>numeric matrix</i>	orfs	samples	covs. / 10 ⁶
	\$tpm		<i>numeric matrix</i>	orfs	samples	tpm
	\$seqs		<i>character vector</i>	orfs	(n/a)	sequences
	\$tax		<i>character matrix</i>	orfs	tax. ranks	taxonomy
	\$tax16S		<i>character vector</i>	orfs	(n/a)	16S rRNA ta
	\$tax_abund		See SQM\$taxa			
\$contigs	\$markers		<i>list</i>	orfs	(n/a)	CheckM1 m
	\$table		<i>dataframe</i>	contigs	misc. data	misc. data
	\$abund		<i>numeric matrix</i>	contigs	samples	abundances
	\$bases		<i>numeric matrix</i>	contigs	samples	abundances
	\$cov		<i>numeric matrix</i>	contigs	samples	coverages
	\$cpm		<i>numeric matrix</i>	contigs	samples	covs. / 10 ⁶
	\$tpm		<i>numeric matrix</i>	contigs	samples	tpm
	\$seqs		<i>character vector</i>	contigs	(n/a)	sequences
	\$tax		<i>character matrix</i>	contigs	tax. ranks	taxonomies
	\$tax_abund		See SQM\$taxa			
\$bins	\$bins		<i>character matrix</i>	contigs	bin. methods	bins
	\$table		<i>dataframe</i>	bins	misc. data	misc. data

		\$length	<i>numeric vector</i>	bins	(n/a)	length
		\$abund	<i>numeric matrix</i>	bins	samples	abundances
		\$percent	<i>numeric matrix</i>	bins	samples	abundances
		\$bases	<i>numeric matrix</i>	bins	samples	abundances
		\$cov	<i>numeric matrix</i>	bins	samples	coverages
		\$cpm	<i>numeric matrix</i>	bins	samples	covs. / 10 ⁶
		\$tax	<i>character matrix</i>	bins	tax. ranks	taxonomy
		\$tax_abund	See SQM\$taxa			
		\$tax_gtdb	<i>character matrix</i>	bins	tax. ranks	GTDB taxon
		\$tax_abund_gtdb	See SQM\$taxa			
\$taxa		\$superkingdom	<i>numeric matrix</i>	superkingdoms	samples	abundances
		\$percent	<i>numeric matrix</i>	superkingdoms	samples	percentages
		\$phylum	<i>numeric matrix</i>	phyla	samples	abundances
		\$percent	<i>numeric matrix</i>	phyla	samples	percentages
		\$class	<i>numeric matrix</i>	classes	samples	abundances
		\$percent	<i>numeric matrix</i>	classes	samples	percentages
		\$order	<i>numeric matrix</i>	orders	samples	abundances
		\$percent	<i>numeric matrix</i>	orders	samples	percentages
		\$family	<i>numeric matrix</i>	families	samples	abundances
		\$percent	<i>numeric matrix</i>	families	samples	percentages
		\$genus	<i>numeric matrix</i>	genera	samples	abundances
		\$percent	<i>numeric matrix</i>	genera	samples	percentages
		\$species	<i>numeric matrix</i>	species	samples	abundances
		\$percent	<i>numeric matrix</i>	species	samples	percentages
\$functions	\$KEGG	\$abund	<i>numeric matrix</i>	KEGG ids	samples	abundances
		\$bases	<i>numeric matrix</i>	KEGG ids	samples	abundances
		\$cov	<i>numeric matrix</i>	KEGG ids	samples	coverages
		\$cpm	<i>numeric matrix</i>	KEGG ids	samples	covs. / 10 ⁶
		\$tpm	<i>numeric matrix</i>	KEGG ids	samples	tpm
		\$copy_number	<i>numeric matrix</i>	KEGG ids	samples	avg. copies
	\$COG	\$abund	<i>numeric matrix</i>	COG ids	samples	abundances
		\$bases	<i>numeric matrix</i>	COG ids	samples	abundances
		\$cov	<i>numeric matrix</i>	COG ids	samples	coverages
		\$cpm	<i>numeric matrix</i>	COG ids	samples	covs. / 10 ⁶
		\$tpm	<i>numeric matrix</i>	COG ids	samples	tpm
		\$copy_number	<i>numeric matrix</i>	COG ids	samples	avg. copies
	\$PFAM	\$abund	<i>numeric matrix</i>	PFAM ids	samples	abundances
		\$bases	<i>numeric matrix</i>	PFAM ids	samples	abundances
		\$cov	<i>numeric matrix</i>	PFAM ids	samples	coverages
		\$cpm	<i>numeric matrix</i>	PFAM ids	samples	covs. / 10 ⁶
		\$tpm	<i>numeric matrix</i>	PFAM ids	samples	tpm
		\$copy_number	<i>numeric matrix</i>	PFAM ids	samples	avg. copies
\$total_reads			<i>numeric vector</i>	samples	(n/a)	total reads
\$misc	\$project_name		<i>character vector</i>	(empty)	(n/a)	project name
	\$samples		<i>character vector</i>	(empty)	(n/a)	samples
	\$tax_names_long	\$superkingdom	<i>character vector</i>	short names	(n/a)	full names
		\$phylum	<i>character vector</i>	short names	(n/a)	full names
		\$class	<i>character vector</i>	short names	(n/a)	full names

	\$order	<i>character vector</i>	short names	(n/a)	full names
	\$family	<i>character vector</i>	short names	(n/a)	full names
	\$genus	<i>character vector</i>	short names	(n/a)	full names
	\$species	<i>character vector</i>	short names	(n/a)	full names
\$tax_names_short		<i>character vector</i>	full names	(n/a)	short names
\$KEGG_names		<i>character vector</i>	KEGG ids	(n/a)	KEGG names
\$KEGG_paths		<i>character vector</i>	KEGG ids	(n/a)	KEGG hierarchies
\$COG_names		<i>character vector</i>	COG ids	(n/a)	COG names
\$COG_paths		<i>character vector</i>	COG ids	(n/a)	COG hierarchies
\$ext_annot_sources		<i>character vector</i>	COG ids	(n/a)	external data

If external databases for functional classification were provided to SqueezeMeta via the `-extdb` argument, the corresponding abundance (reads and bases), coverages, tpm and copy number profiles will be present in `SQM$functions` (e.g. results for the CAZy database would be present in `SQM$functions$CAZy`). Additionally, the extended names of the features present in the external database will be present in `SQM$misc` (e.g. `SQM$misc$CAZy_names`).

Examples

```
## Not run:
## (outside R)
## Run SqueezeMeta on the test data.
/path/to/SqueezeMeta/scripts/SqueezeMeta.pl -p Hadza -f raw -m coassembly -s test.samples
## Now go into R.
library(SQMtools)
Hadza = loadSQM("Hadza") # Where Hadza is the path to the SqueezeMeta output directory.

## End(Not run)

data(Hadza) # We will illustrate the structure of the SQM object on the test data
# Which are the ten most abundant KEGG IDs in our data?
topKEGG = names(sort(rowSums(Hadza$functions$KEGG$tpm), decreasing=TRUE))[1:11]
topKEGG = topKEGG[topKEGG!="Unclassified"]
# Which functions do those KEGG IDs represent?
Hadza$misc$KEGG_names[topKEGG]
# What is the relative abundance of the Negativicutes class across samples?
Hadza$taxa$class$percent["Negativicutes",]
# Which information is stored in the orf, contig and bin tables?
colnames(Hadza$orfs$table)
colnames(Hadza$contigs$table)
colnames(Hadza$bins$table)
# What is the GC content distribution of my metagenome?
boxplot(Hadza$contigs$table[, "GC perc"]) # Not weighted by contig length or abundance!
```

loadSQMLite

*Load tables generated by sqm2tables.py, sqmreads2tables.py or
combine-sqm-tables.py into R.*

Description

This function takes the path to the output directory generated by `sqm2tables.py`, `sqmreads2tables.py` or `combine-sqm-tables.py` a SQMLite object. The SQMLite object will contain taxonomic and functional profiles, but no detailed information on ORFs, contigs or bins. However, it will also have a much smaller memory footprint. A SQMLite object can be used for plotting and exporting, but it can not be subsetted.

Usage

```
loadSQMLite(tables_path, tax_mode = "allfilter")
```

Arguments

tables_path character, tables directory generated by `sqm2table.py`, `sqmreads2tables.py` or `combine-sqm-tables.py`.

tax_mode character, which taxonomic classification should be loaded? SqueezeMeta applies the identity thresholds described in Luo *et al.*, 2014 ([doi:10.1093/nar/gku169](https://doi.org/10.1093/nar/gku169)). Use `allfilter` for applying the minimum identity threshold to all taxa (default), `prokfilter` for applying the threshold to Bacteria and Archaea, but not to Eukaryotes, and `nofilter` for applying no thresholds at all.

Value

SQMLite object containing the parsed tables.

The SQMLite object structure

The SQMLite object is a nested list which contains the following information:

lvl1	lvl2	lvl3	type	rows/names	columns	data
\$taxa	\$superkingdom	\$abund	numeric matrix	superkingdoms	samples	abundances
		\$percent	numeric matrix	superkingdoms	samples	percentages
	\$phylum	\$abund	numeric matrix	phyla	samples	abundances
		\$percent	numeric matrix	phyla	samples	percentages
	\$class	\$abund	numeric matrix	classes	samples	abundances
		\$percent	numeric matrix	classes	samples	percentages
	\$order	\$abund	numeric matrix	orders	samples	abundances
		\$percent	numeric matrix	orders	samples	percentages
	\$family	\$abund	numeric matrix	families	samples	abundances
		\$percent	numeric matrix	families	samples	percentages
	\$genus	\$abund	numeric matrix	genera	samples	abundances
		\$percent	numeric matrix	genera	samples	percentages
	\$species	\$abund	numeric matrix	species	samples	abundances
		\$percent	numeric matrix	species	samples	percentages
\$functions	\$KEGG	\$abund	numeric matrix	KEGG ids	samples	abundances (read)
		\$bases	numeric matrix	KEGG ids	samples	abundances (base)
		\$tpm	numeric matrix	KEGG ids	samples	tpm
		\$copy_number	numeric matrix	KEGG ids	samples	avg. copies

	\$COG	\$abund	numeric matrix	COG ids	samples	abundances (reads)	
		\$bases	numeric matrix	COG ids	samples	abundances (bases)	
		\$tpm	numeric matrix	COG ids	samples	tpm	
	\$PFAM	\$copy_number	numeric matrix	COG ids	samples	avg. copies	
		\$abund	numeric matrix	PFAM ids	samples	abundances (reads)	
		\$bases	numeric matrix	PFAM ids	samples	abundances (bases)	
		\$tpm	numeric matrix	PFAM ids	samples	tpm	
		\$copy_number	numeric matrix	PFAM ids	samples	avg. copies	
	\$total_reads		numeric vector	samples	(n/a)	total reads	
	\$misc	\$project_name		character vector	(empty)	(n/a)	project name
		\$samples		character vector	(empty)	(n/a)	samples
		\$tax_names_long	\$superkingdom	character vector	short names	(n/a)	full names
\$phylum			character vector	short names	(n/a)	full names	
\$class			character vector	short names	(n/a)	full names	
\$order			character vector	short names	(n/a)	full names	
\$family			character vector	short names	(n/a)	full names	
\$genus			character vector	short names	(n/a)	full names	
\$species			character vector	short names	(n/a)	full names	
\$tax_names_short			character vector	full names	(n/a)	short names	
\$KEGG_names		character vector	KEGG ids	(n/a)	KEGG names		
\$KEGG_paths		character vector	KEGG ids	(n/a)	KEGG hierarchy		
\$COG_names		character vector	COG ids	(n/a)	COG names		
\$COG_paths		character vector	COG ids	(n/a)	COG hierarchy		
\$ext_annot_sources			character vector	(empty)	(n/a)	external database	

If external databases for functional classification were provided to SqueezeMeta or SqueezeMeta_reads via the `-extdb` argument, the corresponding abundance, tpm and copy number profiles will be present in `SQM$functions` (e.g. results for the CAZy database would be present in `SQM$functions$CAZy`). Additionally, the extended names of the features present in the external database will be present in `SQM$misc` (e.g. `SQM$misc$CAZy_names`). Note that results generated by SqueezeMeta_reads will contain only read abundances, but not bases, tpm or copy number estimations.

See Also

[plotBars](#) and [plotFunctions](#) will plot the most abundant taxa and functions in a SQMLite object. [exportKrona](#) will generate Krona charts reporting the taxonomy in a SQMLite object.

Examples

```
## Not run:
## (outside R)
## Run SqueezeMeta on the test data.
/path/to/SqueezeMeta/scripts/SqueezeMeta.pl -p Hadza -f raw -m coassembly -s test.samples
## Generate the tabular outputs!
/path/to/SqueezeMeta/utis/sqm2tables.py Hadza Hadza/results/tables
## Now go into R.
library(SQMtools)
Hadza = loadSQMLite("Hadza/results/tables")
```

```
# Where Hadza is the path to the SqueezeMeta output directory.
# Note that this is not the whole SQM project, just the directory containing the tables.
# It would also work with tables generated by sqmreads2tables.py, or combine-sqm-tables.py
plotTaxonomy(Hadza)
plotFunctions(Hadza)
exportKrona(Hadza, 'myKronaTest.html')

## End(Not run)
```

MGKOs	<i>Single Copy Phylogenetic Marker Genes from Sunagawa's group (KOs)</i>
-------	--

Description

Lists of Single Copy Phylogenetic Marker Genes. These are useful for transforming coverages or tpms into copy numbers. This is an alternative way of normalizing data in order to be able to compare functional profiles in samples with different sequencing depths.

Usage

```
data(MGKOs)
```

Format

Character vector with the KEGG identifiers for 10 Single Copy Phylogenetic Marker Genes.

Source

[doi:10.1016/j.cell.2019.10.014](https://doi.org/10.1016/j.cell.2019.10.014),

References

Salazar, G *et al.* (2019). Gene Expression Changes and Community Turnover Differentially Shape the Global Ocean Metatranscriptome *Cell* **179**:1068-1083. ([doi:10.1016/j.cell.2019.10.014](https://doi.org/10.1016/j.cell.2019.10.014)).

See Also

[MGOGs](#) for an equivalent list using OGs instead of KOs; [USiCGs](#) for an alternative set of single copy genes, and for examples on how to generate copy numbers.

MGOGs	<i>Single Copy Phylogenetic Marker Genes from Sunagawa's group (OGs)</i>
-------	--

Description

Lists of Single Copy Phylogenetic Marker Genes. These are useful for transforming coverages or tpms into copy numbers. This is an alternative way of normalizing data in order to be able to compare functional profiles in samples with different sequencing depths.

Usage

```
data(MGOGs)
```

Format

Character vector with the COG identifiers for 10 Single Copy Phylogenetic Marker Genes.

Source

[doi:10.1016/j.cell.2019.10.014](https://doi.org/10.1016/j.cell.2019.10.014)

References

Salazar, G *et al.* (2019). Gene Expression Changes and Community Turnover Differentially Shape the Global Ocean Metatranscriptome *Cell* **179**:1068-1083. ([doi:10.1016/j.cell.2019.10.014](https://doi.org/10.1016/j.cell.2019.10.014)).

See Also

[MGKOs](#) for an equivalent list using KOs instead of OGs; [USiCGs](#) for an alternative set of single copy genes, and for examples on how to generate copy numbers.

mostAbundant	<i>Get the N most abundant rows (or columns) from a numeric table</i>
--------------	---

Description

Return a subset of an input matrix or data frame, containing only the N most abundant rows (or columns), sorted. Alternatively, a custom set of rows can be returned.

Usage

```
mostAbundant(
  data,
  N = 10,
  items = NULL,
  extra_items = NULL,
  ignore = NULL,
  others = FALSE,
  rescale = FALSE,
  bycol = FALSE
)
```

Arguments

<code>data</code>	numeric matrix or data frame
<code>N</code>	integer Number of rows to return (default 10).
<code>items</code>	character vector. Custom row names to return. If provided, it will override <code>N</code> and <code>extra_items</code> (default <code>NULL</code>).
<code>extra_items</code>	character vector. Extra row names to return on top of the <code>N</code> most abundant (default <code>NULL</code>)
<code>ignore</code>	character. Custom row names to drop before abundance calculation.
<code>others</code>	logical. If <code>TRUE</code> , an extra row will be returned containing the aggregated abundances of the elements not selected with <code>N</code> or <code>items</code> (default <code>FALSE</code>).
<code>rescale</code>	logical. Scale result to percentages column-wise (default <code>FALSE</code>).
<code>bycol</code>	logical. Operate on columns instead of rows (default <code>FALSE</code>).

Value

A matrix or data frame (same as input) with the selected rows (or columns).

Examples

```
data(Hadza)
Hadza.carb = subsetFun(Hadza, "Carbohydrate metabolism")
# Which are the 20 most abundant KEGG functions in the ORFs related to carbohydrate metabolism?
topCarb = mostAbundant(Hadza.carb$functions$KEGG$tpm, N=20)
# Now print them with nice names.
rownames(topCarb) = paste(rownames(topCarb),
                          Hadza.carb$misc$KEGG_names[rownames(topCarb)], sep="; ")

topCarb
# We can pass this to any R function.
heatmap(topCarb)
# But for convenience we provide wrappers for plotting ggplot2 heatmaps and barplots.
plotHeatmap(topCarb, label_y="TPM")
plotBars(topCarb, label_y="TPM")
```

mostVariable	<i>Get the N most variable rows (or columns) from a numeric table</i>
--------------	---

Description

Return a subset of an input matrix or data frame, containing only the N most variable rows (or columns), sorted. Variability is calculated as the Coefficient of Variation (sd/mean).

Usage

```
mostVariable(data, N = 10, bycol = FALSE)
```

Arguments

data	numeric matrix or data frame
N	integer Number of rows to return (default 10).
bycol	logical. Operate on columns instead of rows (default FALSE).

Value

A matrix or data frame (same as input) with the selected rows or columns.

Examples

```
data(Hadza)
Hadza.carb = subsetFun(Hadza, "Carbohydrate metabolism")
# Which are the 20 most variable KEGG functions in the ORFs related to carbohydrate metabolism?
topCarb = mostVariable(Hadza.carb$functions$KEGG$tpm, N=20)
# Now print them with nice names
rownames(topCarb) = paste(rownames(topCarb),
                          Hadza.carb$misc$KEGG_names[rownames(topCarb)], sep="; ")
topCarb
# We can pass this to any R function
heatmap(topCarb)
# But for convenience we provide wrappers for plotting ggplot2 heatmaps and barplots
plotHeatmap(topCarb, label_y="TPM")
plotBars(topCarb, label_y="TPM")
```

plotBars	<i>Plot a barplot using ggplot2</i>
----------	-------------------------------------

Description

Plot a ggplot2 barplot from a matrix or data frame. The data should be in tabular format (e.g. features in rows and samples in columns).

Usage

```
plotBars(
  data,
  label_x = "Samples",
  label_y = "Abundances",
  label_fill = "Features",
  color = NULL,
  base_size = 11,
  max_scale_value = NULL,
  metadata_groups = NULL
)
```

Arguments

<code>data</code>	Numeric matrix or data frame.
<code>label_x</code>	character Label for the x axis (default "Samples").
<code>label_y</code>	character Label for the y axis (default "Abundances").
<code>label_fill</code>	character Label for color categories (default "Features").
<code>color</code>	Vector with custom colors for the different features. If empty, the default ggplot2 palette will be used (default NULL).
<code>base_size</code>	numeric. Base font size (default 11).
<code>max_scale_value</code>	numeric. Maximum value to include in the y axis. By default it is handled automatically by ggplot2 (default NULL).
<code>metadata_groups</code>	list. Split the plot into groups defined by the user: <code>list('G1' = c('sample1', 'sample2'), 'G2' = c('sample3', 'sample4'))</code> default NULL).

Value

a ggplot2 plot object.

See Also

[plotTaxonomy](#) for plotting the most abundant taxa of a SQM object; [plotHeatmap](#) for plotting a heatmap with arbitrary data; [mostAbundant](#) for selecting the most abundant rows in a dataframe or matrix.

Examples

```
data(Hadza)
sk = Hadza$taxa$superkingdom$abund
plotBars(sk, label_y = "Raw reads", label_fill = "Superkingdom")
```

plotBins

*Barplot of the most abundant bins in a SQM object***Description**

This function selects the most abundant bins across all samples in a SQM object and represents their abundances in a barplot. It can also plot their aggregated abundances at higher taxonomic ranks. Alternatively, a custom set of bins/taxa can be represented.

Usage

```
plotBins(
  SQM,
  rank = "bin",
  count = "percent",
  N = 15,
  tax_source = "bins",
  bins = NULL,
  tax = NULL,
  others = TRUE,
  samples = NULL,
  ignore_unmapped = FALSE,
  ignore_nobin = FALSE,
  no_partial_classifications = FALSE,
  rescale = FALSE,
  color = NULL,
  base_size = 11,
  max_scale_value = NULL,
  metadata_groups = NULL
)
```

Arguments

SQM	A SQM object.
rank	Taxonomic rank to plot (default bin).
count	character. Either "abund" for raw abundances, "percent" for percentages, "cov" for coverages, or "cpm" for coverages per million reads (default "percent").
N	integer Plot the N most abundant bins (default 15).
tax_source	character. Source of taxonomic annotations, can be "bins" (GTDB bin taxonomy if available, SQM bin taxonomy otherwise), "bins_gtdb" (GTDB bin taxonomy) or "bins_sqm" (SQM bin taxonomy) (default: "bins").
bins	character. Custom bins/taxa to plot. If provided, it will override N (default NULL).
tax	character. Custom bins/taxa to plot. If provided, it will override N and bins (default NULL).

others	logical. Collapse the abundances of least abundant bins, and include the result in the plot (default TRUE).
samples	character. Character vector with the names of the samples to include in the plot. Can also be used to plot the samples in a custom order. If not provided, all samples will be plotted (default NULL).
ignore_unmapped	logical. Don't include unmapped reads in the plot (default FALSE).
ignore_nobin	logical. Don't include reads which are not in a bin in the plot (default FALSE).
no_partial_classifications	logical. Treat reads not fully classified at the requested level (e.g. "Unclassified Bacteroidota" at the class level or below) as fully unclassified. This takes effect before ignore_unclassified, so if both are TRUE the plot will only contain fully classified contigs (default FALSE).
rescale	logical. Re-scale results to percentages (default FALSE).
color	Vector with custom colors for the different features. If empty, we will use our own hand-picked palette if N<=15, and the default ggplot2 palette otherwise (default NULL).
base_size	numeric. Base font size (default 11).
max_scale_value	numeric. Maximum value to include in the y axis. By default it is handled automatically by ggplot2 (default NULL).
metadata_groups	list. Split the plot into groups defined by the user: list('G1' = c('sample1', 'sample2'), 'G2' = c('sample3', 'sample4')) default NULL).

Value

a ggplot2 plot object.

See Also

[plotTaxonomy](#) for plotting the most abundant taxa of a SQM object; [plotBars](#) and [plotHeatmap](#) for plotting barplots or heatmaps with arbitrary data.

Examples

```
data(Hadza)
# Bins distribution.
plotBins(Hadza)
# Aggregated bin taxonomy
plotBins(Hadza, rank = 'order')
```

plotFunctions

Heatmap of the most abundant functions in a SQM object

Description

This function selects the most abundant functions across all samples in a SQM object and represents their abundances in a heatmap. Alternatively, a custom set of functions can be represented.

Usage

```
plotFunctions(
  SQM,
  fun_level = "KEGG",
  count = "copy_number",
  N = 25,
  fun = NULL,
  samples = NULL,
  display_function_names = TRUE,
  ignore_unmapped = TRUE,
  ignore_unclassified = TRUE,
  gradient_col = c("ghostwhite", "dodgerblue4"),
  rescale_percent = FALSE,
  base_size = 11,
  metadata_groups = NULL
)
```

Arguments

SQM	A SQM, SQMbunch or SQMlite object.
fun_level	character. Either "KEGG", "COG", "PFAM" or any other custom database used for annotation (default "KEGG").
count	character. Either "abund" for raw abundances, "percent" for percentages, "bases" for raw base counts, "cpm" for coverages per million reads, "tpm" for TPM normalized values or "copy_number" for copy numbers (default "copy_number"). Note that a given count type might not be available in this object (e.g. TPM or copy number in SQMlite objects originating from a SQM reads project).
N	integer Plot the N most abundant functions (default 25).
fun	character. Custom functions to plot. If provided, it will override N (default NULL).
samples	character. Character vector with the names of the samples to include in the plot. Can also be used to plot the samples in a custom order. If not provided, all samples will be plotted (default NULL).
display_function_names	logical. Plot function names alongside function IDs, if available (default TRUE).

ignore_unmapped logical. Don't include unmapped reads in the plot (default TRUE).
ignore_unclassified logical. Don't include unclassified ORFs in the plot (default TRUE).
gradient_col A vector of two colors representing the low and high ends of the color gradient (default `c("ghostwhite", "dodgerblue4")`).
rescale_percent logical. Calculate percent counts over the number of reads in the input object, instead of over the total number of reads in the original project (default FALSE).
base_size numeric. Base font size (default 11).
metadata_groups list. Split the plot into groups defined by the user: `list('G1' = c('sample1', sample2'), 'G2' = c('sample3', 'sample4'))` default NULL).

Value

a ggplot2 plot object.

See Also

[plotTaxonomy](#) for plotting the most abundant taxa of a SQM object; [plotBars](#) and [plotHeatmap](#) for plotting barplots or heatmaps with arbitrary data.

Examples

```
data(Hadza)
plotFunctions(Hadza)
```

plotHeatmap	<i>Plot a heatmap using ggplot2</i>
-------------	-------------------------------------

Description

Plot a ggplot2 heatmap from a matrix or data frame. The data should be in tabular format (e.g. features in rows and samples in columns).

Usage

```
plotHeatmap(
  data,
  label_x = "Samples",
  label_y = "Features",
  label_fill = "Abundance",
  gradient_col = c("ghostwhite", "dodgerblue4"),
  base_size = 11,
  metadata_groups = NULL
)
```

Arguments

<code>data</code>	numeric matrix or data frame.
<code>label_x</code>	character Label for the x axis (default "Samples").
<code>label_y</code>	character Label for the y axis (default "Features").
<code>label_fill</code>	character Label for color scale (default "Abundance").
<code>gradient_col</code>	A vector of two colors representing the low and high ends of the color gradient (default <code>c("ghostwhite", "dodgerblue4")</code>).
<code>base_size</code>	numeric. Base font size (default 11).
<code>metadata_groups</code>	list. Split the plot into groups defined by the user: <code>list('G1' = c('sample1', sample2'), 'G2' = c('sample3', 'sample4'))</code> default NULL).

Value

A ggplot2 plot object.

See Also

[plotFunctions](#) for plotting the top functional categories of a SQM object; [plotBars](#) for plotting a barplot with arbitrary data; [mostAbundant](#) for selecting the most abundant rows in a dataframe or matrix.

Examples

```
data(Hadza)
topPFAM = mostAbundant(Hadza$functions$PFAM$tpm)
topPFAM = topPFAM[rownames(topPFAM) != "Unclassified",] # Take out the Unclassified ORFs.
plotHeatmap(topPFAM, label_x = "Samples", label_y = "PFAMs", label_fill = "TPM")
```

plotTaxonomy

Barplot of the most abundant taxa in a SQM object

Description

This function selects the most abundant taxa across all samples in a SQM object and represents their abundances in a barplot. Alternatively, a custom set of taxa can be represented.

Usage

```
plotTaxonomy(
  SQM,
  rank = "phylum",
  count = "percent",
  N = 15,
  tax_source = NULL,
  tax = NULL,
```

```

    others = TRUE,
    samples = NULL,
    nocds = "treat_separately",
    ignore_unmapped = FALSE,
    ignore_unclassified = FALSE,
    ignore_nobin = FALSE,
    no_partial_classifications = FALSE,
    rescale = FALSE,
    color = NULL,
    base_size = 11,
    max_scale_value = NULL,
    metadata_groups = NULL
)

```

Arguments

SQM	A SQM, SQMbunch or a SQMLite object.
rank	Taxonomic rank to plot (default phylum).
count	character. Either "percent" for percentages, or "abund" for raw abundances (default "percent").
N	integer Plot the N most abundant taxa (default 15).
tax_source	character. Source of taxonomic annotations, can be "orfs", "contigs", "bins" (GTDB bin taxonomy if available, SQM bin taxonomy otherwise), "bins_gtdb" (GTDB bin taxonomy) or "bins_sqm" (SQM bin taxonomy) (default: use the 'tax_source' from the input object).
tax	character. Custom taxa to plot. If provided, it will override N (default NULL).
others	logical. Collapse the abundances of least abundant taxa, and include the result in the plot (default TRUE).
samples	character. Character vector with the names of the samples to include in the plot. Can also be used to plot the samples in a custom order. If not provided, all samples will be plotted (default NULL).
nocds	character. Either "treat_separately" to treat reads annotated as No CDS separately, "treat_as_unclassified" to treat them as Unclassified or "ignore" to ignore them in the plot (default "treat_separately").
ignore_unmapped	logical. Don't include unmapped reads in the plot (default FALSE).
ignore_unclassified	logical. Don't include unclassified reads in the plot (default FALSE).
ignore_nobin	logical. Ignore reads not mapping to any bin when tax_source is bins, bins_gtdb or bins_sqm (default FALSE).
no_partial_classifications	logical. Treat reads not fully classified at the requested level (e.g. "Unclassified Bacteroidota" at the class level or below) as fully unclassified. This takes effect before ignore_unclassified, so if both are TRUE the plot will only contain fully classified contigs (default FALSE).

rescale	logical. Re-scale results to percentages (default FALSE).
color	Vector with custom colors for the different features. If empty, we will use our own hand-picked palette if N<=15, and the default ggplot2 palette otherwise (default NULL).
base_size	numeric. Base font size (default 11).
max_scale_value	numeric. Maximum value to include in the y axis. By default it is handled automatically by ggplot2 (default NULL).
metadata_groups	list. Split the plot into groups defined by the user: list('G1' = c('sample1', sample2'), 'G2' = c('sample3', 'sample4')) default NULL).

Value

a ggplot2 plot object.

See Also

[plotFunctions](#) for plotting the most abundant functions of a SQM object; [plotBars](#) and [plotHeatmap](#) for plotting barplots or heatmaps with arbitrary data.

Examples

```
data(Hadza)
Hadza.amin = subsetFun(Hadza, "Amino acid metabolism")
# Taxonomic distribution of amino acid metabolism ORFs at the family level.
plotTaxonomy(Hadza.amin, "family")
```

RecA	<i>RecA/RadA recombinase</i>
------	------------------------------

Description

The recombination protein RecA/RadA is essential for the repair and maintenance of DNA, and has homologs in every bacteria and archaea. By dividing the coverage of functions by the coverage of RecA, abundances can be transformed into copy numbers, which can be used to compare functional profiles in samples with different sequencing depths. RecA-derived copy numbers are available in the SQM object (SQM\$functions\$<annotation_type>\$copy_number).

Usage

```
data(RecA)
```

Format

Character vector with the COG identifier for RecA/RadA.

Examples

```
data(Hadza)
data(RecA)
### Let's calculate the average copy number of each function in our samples.
# We do it for COG annotations here, but we could also do it for KEGG or PFAMs.
COG.coverage = Hadza$functions$COG$cov
COG.copynumber = t(t(COG.coverage) / COG.coverage[RecA,]) # Sample-wise division by RecA coverage.
```

```
remove_contigs_from_bin
```

Remove contigs from a given bin

Description

Remove contigs from a given bin

Usage

```
remove_contigs_from_bin(SQM, bin, contigs)
```

Arguments

SQM	A SQM object.
bin	character. Name of the bin from which the contigs will be removed.
contigs	character. Vector with the names of the contigs that will be removed from the new bin.

Value

SQM object with the new binning information, including recalculated bin statistics if possible.

See Also

[find_redundant_contigs](#), [create_bin](#)

renameSamples	<i>Change sample names</i>
---------------	----------------------------

Description

Change the sample names of a SQM or SQMLite object

Usage

```
renameSamples(SQM, new_sample_names)
```

Arguments

SQM	SQM or SQMLite object
new_sample_names	character. New sample names

Value

SQM or SQMLite object with the new sample names

rowMaxs	<i>Return a vector with the row-wise maxima of a matrix or dataframe.</i>
---------	---

Description

Return a vector with the row-wise maxima of a matrix or dataframe.

Usage

```
rowMaxs(table)
```

Arguments

table	matrix or dataframe.
-------	----------------------

Value

a vector with the row-wise maxima.

rowMins	<i>Return a vector with the row-wise minima of a matrix or dataframe.</i>
---------	---

Description

Return a vector with the row-wise minima of a matrix or dataframe.

Usage

```
rowMins(table)
```

Arguments

table	matrix or dataframe.
-------	----------------------

Value

a vector with the row-wise minima.

seqvec2fasta	<i>Print a named vector of sequences as a fasta-formatted string</i>
--------------	--

Description

Print a named vector of sequences as a fasta-formatted string

Usage

```
seqvec2fasta(seqvec, output_name = "")
```

Arguments

seqvec	vector. The vector to be written as a fasta string.
output_name	A connection, or a character string naming the file to print to. If "" (the default), sequences will be printed to the standard output connection.

Examples

```
data(Hadza)
seqvec2fasta(Hadza$orfs$seqs[1:10])
```

SQM_to_microeco	<i>Convert a SQM object into a microtable object from the microeco package</i>
-----------------	--

Description

This function will convert the selected features from a SQM object into an object of the `microtable` class from the `microeco` package. When possible, it will also include the taxonomy of the included features (for functional classifications, the taxonomy table will instead include the description of each feature ID). Optionally, it accepts a meta table that will be passed as provided to `microtable$new`.

Usage

```
SQM_to_microeco(
  SQM,
  features = "genus",
  count = "abund",
  md = NULL,
  nocds = "treat_separately",
  no_partial_classifications = FALSE,
  ignore_unclassified = FALSE,
  ignore_unmapped = FALSE,
  bin_tax_source = "SQM",
  include_seqs = FALSE
)
```

Arguments

SQM	A SQM, SQMbunch or SQMlite object.
features	character. Either "orfs", "contigs", "bins", any taxonomic rank included in SQM\$taxa or any functional classification included in SQM\$functions (default "tax"). Note that a given feature type might not be available in this objects (e.g. "contigs" in SQMlite objects originating from a SQM reads project).
count	character. Either "abund" for raw abundances, "percent" for percentages, "bases" for raw base counts, "cov" for coverages, "cpm" for coverages per million reads, "tpm" for TPM normalized values or "copy_number" for copy numbers (default "abund"). Note that a given count type might not available in this object (e.g. TPM or copy number in SQMlite objects originating from a SQM reads project).
md	data.frame. A optional data.frame containing metadata for the samples in the SQM object.
nocds	character. Either "treat_separately" to treat features annotated as No CDS separately, "treat_as_unclassified" to treat them as Unclassified or "ignore" to ignore them in the output (default "treat_separately").

<code>no_partial_classifications</code>	logical. When <code>features</code> is a taxonomic rank, treat features not fully classified at the requested level (e.g. "Unclassified bacteroidota" at the class level or below) as fully unclassified. This takes effect before <code>ignore_unclassified</code> , so if both are TRUE the plot will only contain features that were fully classified at the requested level (default FALSE).
<code>ignore_unclassified</code>	logical. When <code>features</code> is a taxonomic rank or functional category, don't include unclassified reads in the output (default FALSE).
<code>ignore_unmapped</code>	logical. Don't include unmapped reads in the output (default FALSE).
<code>bin_tax_source</code>	character. Source of taxonomy when <code>features</code> = "bins", either "SQM" or "gtdb" (default "gtdb").
<code>include_seqs</code>	logical. Whether to include sequences or not if creating a microtable from contigs (default FALSE).

Value

A [microtable](#).

See Also

[SQM_to_phyloseq](#) for exporting a SQM/SQMLite/SQM object as a phyloseq object.

<code>SQM_to_phyloseq</code>	<i>Convert a SQM object into a phyloseq object from the phyloseq package</i>
------------------------------	--

Description

This function will convert the selected features from a SQM object into a phyloseq object from the [phyloseq](#) package. When possible, it will also include the taxonomy of the included features (for functional classifications, the taxonomy table will instead include the description of each feature ID). Optionally, it accepts a meta table that will be passed as provided to the phyloseq object constructor.

Usage

```
SQM_to_phyloseq(
  SQM,
  features = "genus",
  count = "abund",
  md = NULL,
  nocds = "treat_separately",
  no_partial_classifications = FALSE,
  ignore_unclassified = FALSE,
```

```

    ignore_unmapped = FALSE,
    bin_tax_source = "SQM",
    include_seqs = FALSE
  )

```

Arguments

<code>SQM</code>	A SQM, SQMbunch or SQMLite object.
<code>features</code>	character. Either "orfs", "contigs", "bins", any taxonomic rank included in <code>SQM\$taxa</code> or any functional classification included in <code>SQM\$functions</code> (default "tax"). Note that a given feature type might not be available in this objects (e.g. "contigs" in SQMLite objects originating from a SQM reads project).
<code>count</code>	character. Either "abund" for raw abundances, "percent" for percentages, "bases" for raw base counts, "cov" for coverages, "cpm" for coverages per million reads, "tpm" for TPM normalized values or "copy_number" for copy numbers (default "abund"). Note that a given count type might not available in this object (e.g. TPM or copy number in SQMLite objects originating from a SQM reads project).
<code>md</code>	data.frame. A optional data.frame containing metadata for the samples in the SQM object.
<code>nocds</code>	character. Either "treat_separately" to treat features annotated as No CDS separately, "treat_as_unclassified" to treat them as Unclassified or "ignore" to ignore them in the output (default "treat_separately").
<code>no_partial_classifications</code>	logical. When features is a taxonomic rank, treat features not fully classified at the requested level (e.g. "Unclassified bacteroidota" at the class level or below) as fully unclassified. This takes effect before <code>ignore_unclassified</code> , so if both are TRUE the plot will only contain features that were fully classified at the requested level (default FALSE).
<code>ignore_unclassified</code>	logical. When features is a taxonomic rank or functional category, don't include unclassified reads in the output (default FALSE).
<code>ignore_unmapped</code>	logical. Don't include unmapped reads in the output (default FALSE).
<code>bin_tax_source</code>	character. Source of taxonomy when features = "bins", either "SQM" or "gtdb" (default "gtdb").
<code>include_seqs</code>	logical. Whether to include sequences or not if creating a microtable from ORFs or contigs (default FALSE).

Value

A phyloseq object.

See Also

[SQM_to_microeco](#) for exporting a SQM/SQMLite/SQM object as a microtable object.

subsetBins	Create a SQM object containing only the requested bins, and the contigs and ORFs contained in them.
------------	---

Description

Create a SQM object containing only the requested bins, and the contigs and ORFs contained in them.

Usage

```
subsetBins(
  SQM,
  bins = NULL,
  rank = NULL,
  tax = NULL,
  min_completeness = NULL,
  max_contamination = NULL,
  tax_source = "bins",
  trusted_functions_only = FALSE,
  ignore_unclassified_functions = FALSE,
  rescale_tpm = TRUE,
  rescale_copy_number = TRUE,
  allow_empty = FALSE
)
```

Arguments

SQM	SQM object to be subsetted.
bins	character. Vector of bins to be selected. If provided, will override rank, tax, min_completeness and max_contamination.
rank	character. The taxonomic rank from which to select the desired taxa (superkingdom, phylum, class, order, family, genus, species)
tax	character. A taxon or vector of taxa to be selected.
min_completeness	numeric. Discard bins with completeness lower than this value (default NULL).
max_contamination	numeric. Discard bins with contamination higher than this value (default NULL).
tax_source	character, source data used for taxonomic subsetting (if rank and tax are provided) and for the aggregate taxonomy tables present in SQM\$taxa, either "orfs", "contigs", "bins" (GTDB bin taxonomy if available, SQM bin taxonomy otherwise), "bins_gtdb" (GTDB bin taxonomy) or "bins_sqm" (SQM bin taxonomy). If using bins_gtdb, note that GTDB taxonomy may differ from the NCBI taxonomy used throughout the rest of SqueezeMeta. Default "bins".

trusted_functions_only	logical. If TRUE, only highly trusted functional annotations (best hit + best average) will be considered when generating aggregated function tables. If FALSE, best hit annotations will be used (default FALSE).
ignore_unclassified_functions	logical. If FALSE, ORFs with no functional classification will be aggregated together into an "Unclassified" category. If TRUE, they will be ignored (default FALSE).
rescale_tpm	logical. If TRUE, TPMs for KEGGs, COGs, and PFAMs will be recalculated (so that the TPMs in the subset actually add up to 1 million). Otherwise, per-function TPMs will be calculated by aggregating the TPMs of the ORFs annotated with that function, and will thus keep the scaling present in the parent object. By default it is set to TRUE, which means that the returned TPMs will be scaled <i>by million of reads of the selected bins</i> .
rescale_copy_number	logical. If TRUE, copy numbers will be recalculated using the median single-copy gene coverages in the subset. Otherwise, single-copy gene coverages will be taken from the parent object. By default it is set to TRUE, which means that the returned copy numbers for each function will represent the average copy number of that function <i>per genome of the selected taxon</i> .
allow_empty	(internal use only).

Value

SQM object containing only the requested bins.

See Also

[subsetContigs](#), [subsetORFs](#)

Examples

```
data(Hadza)
# Which are the most complete bins?
topBinNames = rownames(Hadza$bins$table)[order(Hadza$bins$table[, "Completeness"],
                                                decreasing=TRUE)][1:2]

# Subset with the most complete bin.
topBin = subsetBins(Hadza, topBinNames[1])

# Subset with all the bins over 90% completeness
over90 = subsetBins(Hadza, min_completeness = 90)

# Subset with bins from the Phascolarctobacterium genus using SqueezeMeta's taxonomy
phasco = subsetBins(Hadza, tax_source = "bins_sqm", rank = "genus", tax = "Phascolarctobacterium")

# Subset with bins from the Bacteroidota phylum using GTDB taxonomy
bact = subsetBins(Hadza, tax_source = "bins_gtdb", rank = "phylum", tax = "Bacteroidota")
```

subsetContigs	<i>Select contigs</i>
---------------	-----------------------

Description

Create a SQM object containing only the requested contigs, the ORFs contained in them and the bins that contain them.

Usage

```
subsetContigs(
  SQM,
  contigs,
  tax_source = "contigs",
  trusted_functions_only = FALSE,
  ignore_unclassified_functions = FALSE,
  rescale_tpm = FALSE,
  rescale_copy_number = FALSE,
  recalculate_bin_stats = TRUE,
  allow_empty = FALSE
)
```

Arguments

SQM	SQM object to be subsetted.
contigs	character. Vector of contigs to be selected.
tax_source	character, source data used for the taxonomy tables present in SQM\$taxa, either "bins" (GTDB bin taxonomy if available, SQM bin taxonomy otherwise), "bins_gtdb" (GTDB bin taxonomy) or "bins_sqm" (SQM bin taxonomy). Default "contigs".
trusted_functions_only	logical. If TRUE, only highly trusted functional annotations (best hit + best average) will be considered when generating aggregated function tables. If FALSE, best hit annotations will be used (default FALSE).
ignore_unclassified_functions	logical. If FALSE, ORFs with no functional classification will be aggregated together into an "Unclassified" category. If TRUE, they will be ignored (default FALSE).
rescale_tpm	logical. If TRUE, TPMs for KEGGs, COGs, and PFAMs will be recalculated (so that the TPMs in the subset actually add up to 1 million). Otherwise, per-function TPMs will be calculated by aggregating the TPMs of the ORFs annotated with that function, and will thus keep the scaling present in the parent object (default FALSE).
rescale_copy_number	logical. If TRUE, copy numbers will be recalculated using the median single-copy gene coverages in the subset. Otherwise, single-copy gene coverages will

be taken from the parent object. By default it is set to FALSE, which means that the returned copy numbers for each function will represent the average copy number of that function per genome in the parent object.

recalculate_bin_stats

logical. If TRUE, bin abundance, quality and taxonomy are recalculated based on the contigs present in the subsetted object (default TRUE).

allow_empty (internal use only).

Value

SQM object containing only the selected contigs.

See Also

[subsetORFs](#)

Examples

```
data(Hadza)
# Which contigs have a GC content below 40?
lowGCcontigNames = rownames(Hadza$contigs$table[Hadza$contigs$table[, "GC perc"] < 40,])
lowGCcontigs = subsetContigs(Hadza, lowGCcontigNames)
hist(lowGCcontigs$contigs$table[, "GC perc"])
```

subsetFun

Filter results by function

Description

Create a SQM or SQMbunch object containing only the ORFs with a given function, and the contigs and bins that contain them.

Usage

```
subsetFun(
  SQM,
  fun,
  columns = NULL,
  ignore_case = TRUE,
  fixed = FALSE,
  trusted_functions_only = FALSE,
  ignore_unclassified_functions = FALSE,
  rescale_tpm = FALSE,
  rescale_copy_number = FALSE,
  recalculate_bin_stats = FALSE,
  allow_empty = FALSE
)
```

Arguments

SQM	SQM or SQMbunch object to be subsetted.
fun	character. Pattern to search for in the different functional classifications.
columns	character. Restrict the search to the provided column names from <code>SQM\$orfs\$table</code> . If not provided the search will be performed in all the columns containing functional information (default NULL).
ignore_case	logical. Make pattern matching case-insensitive (default TRUE).
fixed	logical. If TRUE, pattern is a string to be matched as is. If FALSE the pattern is treated as a regular expression (default FALSE).
trusted_functions_only	logical. If TRUE, only highly trusted functional annotations (best hit + best average) will be considered when generating aggregated function tables. If FALSE, best hit annotations will be used (default FALSE).
ignore_unclassified_functions	logical. If FALSE, ORFs with no functional classification will be aggregated together into an "Unclassified" category. If TRUE, they will be ignored (default FALSE).
rescale_tpm	logical. If TRUE, TPMs for KEGGs, COGs, and PFAMs will be recalculated (so that the TPMs in the subset actually add up to 1 million). Otherwise, per-function TPMs will be calculated by aggregating the TPMs of the ORFs annotated with that function, and will thus keep the scaling present in the parent object (default FALSE).
rescale_copy_number	logical. If TRUE, copy numbers will be recalculated using the median single-copy gene coverages in the subset. Otherwise, single-copy gene coverages will be taken from the parent object. By default it is set to FALSE, which means that the returned copy numbers for each function will represent the average copy number of that function per genome in the parent object.
recalculate_bin_stats	logical. If TRUE, bin abundance, quality and taxonomy are recalculated based on the contigs present in the subsetted object (default FALSE).
allow_empty	(internal use only).

Value

SQM or SQMbunch object containing only the requested function.

See Also

[subsetTax](#), [subsetORFs](#), [subsetSamples](#), [combineSQM](#). The most abundant items of a particular table contained in a SQM object can be selected with [mostAbundant](#).

Examples

```
data(Hadza)
Hadza.iron = subsetFun(Hadza, "iron")
```

```
Hadza.carb = subsetFun(Hadza, "Carbohydrate metabolism")
# Search for multiple patterns using regular expressions
Hadza.twoKOs = subsetFun(Hadza, "K00812|K00813", fixed=FALSE)
```

subsetORFs

Select ORFs

Description

Create a SQM object containing only the requested ORFs, and the contigs and bins that contain them. Internally, all the other subset functions in this package end up calling subsetORFs to do the work for them.

Usage

```
subsetORFs(
  SQM,
  orfs,
  tax_source = "orfs",
  trusted_functions_only = FALSE,
  ignore_unclassified_functions = FALSE,
  rescale_tpm = FALSE,
  rescale_copy_number = FALSE,
  recalculate_bin_stats = TRUE,
  contigs_override = NULL,
  allow_empty = FALSE
)
```

Arguments

SQM	SQM object to be subsetted.
orfs	character. Vector of ORFs to be selected.
tax_source	character, source data used for the taxonomy tables present in SQM\$taxa, either "orfs", "contigs", "bins" (GTDB bin taxonomy if available, SQM bin taxonomy otherwise), "bins_gtdb" (GTDB bin taxonomy) or "bins_sqm" (SQM bin taxonomy). Default "orfs".
trusted_functions_only	logical. If TRUE, only highly trusted functional annotations (best hit + best average) will be considered when generating aggregated function tables. If FALSE, best hit annotations will be used (default FALSE).
ignore_unclassified_functions	logical. If FALSE, ORFs with no functional classification will be aggregated together into an "Unclassified" category. If TRUE, they will be ignored (default FALSE).

rescale_tpm	logical. If TRUE, TPMs for KEGGs, COGs, and PFAMs will be recalculated (so that the TPMs in the subset actually add up to 1 million). Otherwise, per-function TPMs will be calculated by aggregating the TPMs of the ORFs annotated with that function, and will thus keep the scaling present in the parent object (default FALSE).
rescale_copy_number	logical. If TRUE, copy numbers will be recalculated using the median single-copy gene coverages in the subset. Otherwise, single-copy gene coverages will be taken from the parent object. By default it is set to FALSE, which means that the returned copy numbers for each function will represent the average copy number of that function per genome in the parent object.
recalculate_bin_stats	logical. If TRUE, bin abundance, quality and taxonomy are recalculated based on the contigs present in the subsetted object (default TRUE).
contigs_override	character. Optional vector of contigs to be included in the subsetted object.
allow_empty	(internal use only).

Value

SQM object containing the requested ORFs.

A note on contig/bins subsetting

While this function selects the contigs and bins that contain the desired orfs, it DOES NOT recalculate contig abundance and statistics based on the selected ORFs only. This means that the abundances presented in tables such as `SQM$contig$abund` will still refer to the complete contigs, regardless of whether only a fraction of their ORFs are actually present in the returned SQM object. This is also true for the statistics presented in `SQM$contigs$table`. Bin statistics may be recalculated if `rescale_copy_number` is set to TRUE, but recalculation will be based on contigs, not ORFs.

Examples

```
data(Hadza)
# Select the 100 most abundant ORFs in our dataset.
mostAbundantORFnames = names(sort(rowSums(Hadza$orfs$tpm), decreasing=TRUE))[1:100]
mostAbundantORFs = subsetORFs(Hadza, mostAbundantORFnames)
```

subsetRand

Select random ORFs

Description

Create a random subset of a SQM object.

Usage

```
subsetRand(SQM, N)
```

Arguments

SQM SQM object to be subsetted.
 N numeric. number of random ORFs to select.

Value

SQM object containing a random subset of ORFs.

See Also

[subsetORFs](#)

subsetSamples	<i>Filter results by sample</i>
---------------	---------------------------------

Description

Create a SQM or SQMlite object containing only samples specified by the user, and the ORFs, contigs, bins, taxa and functions present in those samples.

Usage

```
subsetSamples(SQM, samples, remove_missing = TRUE, new_sample_names = NULL)
```

Arguments

SQM SQM or SQMlite object to be subsetted.
 samples character. Samples to be included in the subset.
 remove_missing bool. If TRUE, ORFs, contigs, bins, taxa and functions absent from the selected samples will be removed from the subsetted object (default TRUE).
 new_sample_names character. New sample names to be included in the subset (default NULL).

Value

SQM or SQMlite object containing only the requested samples.

See Also

[subsetTax](#), [subsetFun](#), [subsetORFs](#), [combineSQM](#). The most abundant items of a particular table contained in a SQM object can be selected with [mostAbundant](#).

subsetTax	<i>Filter results by taxonomy</i>
-----------	-----------------------------------

Description

Create a SQM or SQMbunch object containing only the contigs/bins with a given consensus taxonomy, as well as the ORFs contained in them.

Usage

```
subsetTax(
  SQM,
  rank,
  tax,
  tax_source = NULL,
  trusted_functions_only = FALSE,
  ignore_unclassified_functions = FALSE,
  rescale_tpm = TRUE,
  rescale_copy_number = TRUE,
  recalculate_bin_stats = FALSE,
  allow_empty = FALSE
)
```

Arguments

SQM	SQM object to be subsetted.
rank	character. The taxonomic rank from which to select the desired taxa (superkingdom, phylum, class, order, family, genus, species)
tax	character. A taxon or vector of taxa to be selected.
tax_source	character, source data used for feature selection, and to generate the taxonomy tables present in SQM\$taxa, either "orfs", "contigs", "bins" (GTDB bin taxonomy if available, SQM bin taxonomy otherwise), "bins_gtdb" (GTDB bin taxonomy) or "bins_sqm" (SQM bin taxonomy). When "bins", "bins_gtdb" or "bins_sqm", this function will select the bins from the desired taxa, otherwise it will select the contigs from the desired taxa. If using "bins_gtdb", note that GTDB taxonomy may differ from the NCBI taxonomy used throughout the rest of SqueezeMeta. Default "contigs", unless the project was created with the '–onlybins' flag, where it will be "bins_gtdb" if GTDB taxonomy is available for the bins.
trusted_functions_only	logical. If TRUE, only highly trusted functional annotations (best hit + best average) will be considered when generating aggregated function tables. If FALSE, best hit annotations will be used (default FALSE).
ignore_unclassified_functions	logical. If FALSE, ORFs with no functional classification will be aggregated together into an "Unclassified" category. If TRUE, they will be ignored (default FALSE).

rescale_tpm	logical. If TRUE, TPMs for KEGGs, COGs, and PFAMs will be recalculated (so that the TPMs in the subset actually add up to 1 million). Otherwise, per-function TPMs will be calculated by aggregating the TPMs of the ORFs annotated with that function, and will thus keep the scaling present in the parent object. By default it is set to TRUE, which means that the returned TPMs will be scaled <i>by million of reads of the selected taxon</i> .
rescale_copy_number	logical. If TRUE, copy numbers will be recalculated using the median single-copy gene coverages in the subset. Otherwise, single-copy gene coverages will be taken from the parent object. By default it is set to TRUE, which means that the returned copy numbers for each function will represent the average copy number of that function <i>per genome of the selected taxon</i> .
recalculate_bin_stats	logical. If TRUE, bin abundance, quality and taxonomy are recalculated based on the contigs present in the subsetted object (default TRUE).
allow_empty	(internal use only).

Value

SQM or SQMbunch object containing only the requested taxon.

See Also

[subsetFun](#), [subsetContigs](#), [subsetSamples](#), [combineSQM](#). The most abundant items of a particular table contained in a SQM object can be selected with [mostAbundant](#).

Examples

```
data(Hadza)
Hadza.Prevotella = subsetTax(Hadza, "genus", "Prevotella")
Hadza.Bacteroidota = subsetTax(Hadza, "phylum", "Bacteroidota")
```

summary.SQM

summary method for class SQM

Description

Computes different statistics of the data contained in the SQM object.

Usage

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

Arguments

object	SQM object to be summarized.
...	Additional parameters (ignored).

Value

A list of summary statistics.

summary.SQMbunch	<i>summary method for class SQMbunch</i>
------------------	--

Description

Computes different statistics of the data contained in the SQMbunch object.

Usage

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

Arguments

- object SQMbunch object to be summarized.
- ... Additional parameters (ignored).

Value

A list of summary statistics.

summary.SQMLite	<i>summary method for class SQMLite</i>
-----------------	---

Description

Computes different statistics of the data contained in the SQMLite object.

Usage

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

Arguments

- object SQMLite object to be summarized.
- ... Additional parameters (ignored).

Value

A list of summary statistics.

USiCGs

*Universal Single-Copy Genes***Description**

Lists of Universal Single Copy Genes for Bacteria and Archaea. These are useful for transforming coverages or tpms into copy numbers. This is an alternative way of normalizing data in order to be able to compare functional profiles in samples with different sequencing depths.

Usage

```
data(USiCGs)
```

Format

Character vector with the KEGG identifiers for 15 Universal Single Copy Genes.

Source

[doi:10.1371/journal.pcbi.1003292](https://doi.org/10.1371/journal.pcbi.1003292), Table S1

References

Carr, Shen-Orr & Borenstein (2013). Reconstructing the Genomic Content of Microbiome Taxa through Shotgun Metagenomic Deconvolution *PLoS Comput. Biol.* **9**:e1003292. ([doi:10.1371/journal.pcbi.1003292](https://doi.org/10.1371/journal.pcbi.1003292)).

See Also

[MGOGs](#) and [MGKOs](#) for an alternative set of single copy genes, and for examples on how to generate copy numbers.

Examples

```
data(Hadza)
data(USiCGs)
### Let's look at the Universal Single Copy Gene distribution in our samples.
KEGG.tpm = Hadza$functions$KEGG$tpm
all(USiCGs %in% rownames(KEGG.tpm)) # Are all the USiCGs present in our dataset?
# Plot a boxplot of USiCGs tpms and calculate median USiCGs tpm.
# This looks weird in the test dataset because it contains only a small subset of the metagenomes.
# In a set of complete metagenomes USiCGs should have fairly similar TPM averages
# and low dispersion across samples.
boxplot(t(KEGG.tpm[USiCGs,]), names=USiCGs, ylab="TPM", col="slateblue2")

### Now let's calculate the average copy numbers of each function.
# We do it for KEGG annotations here, but we could also do it for COGs or PFAMs.
USiCGs.cov = apply(Hadza$functions$KEGG$cov[USiCGs,], 2, median)
# Sample-wise division by the median USiCG coverage.
KEGG.copynumber = t(t(Hadza$functions$KEGG$cov) / USiCGs.cov)
```

Index

* datasets

- CheckMProkaryote, [3](#)
- Hadza, [12](#)
- MGKOs, [19](#)
- MGOGs, [20](#)
- RecA, [30](#)
- USiCGs, [48](#)

CheckMProkaryote, [3](#)

combineSQM, [3](#), [5](#), [41](#), [44](#), [46](#)

combineSQMlite, [3](#), [5](#), [5](#)

create_bin, [6](#), [11](#), [31](#)

exportBins, [6](#)

exportContigs, [7](#)

exportKrona, [7](#), [18](#)

exportORFs, [8](#)

exportPathway, [8](#)

exportTable, [10](#)

find_redundant_contigs, [6](#), [11](#), [31](#)

Hadza, [12](#)

loadSQM, [12](#), [13](#)

loadSQMlite, [14](#), [16](#)

MGKOs, [3](#), [19](#), [20](#), [48](#)

MGOGs, [3](#), [19](#), [20](#), [48](#)

microtable, [34](#), [35](#)

mostAbundant, [20](#), [23](#), [28](#), [41](#), [44](#), [46](#)

mostVariable, [22](#)

plotBars, [18](#), [22](#), [25](#), [27](#), [28](#), [30](#)

plotBins, [24](#)

plotFunctions, [10](#), [18](#), [26](#), [28](#), [30](#)

plotHeatmap, [23](#), [25](#), [27](#), [27](#), [30](#)

plotTaxonomy, [7](#), [23](#), [25](#), [27](#), [28](#)

RecA, [30](#)

remove_contigs_from_bin, [6](#), [11](#), [31](#)

renameSamples, [32](#)

rowMaxs, [32](#)

rowMins, [33](#)

seqvec2fasta, [33](#)

SQM_to_microeco, [34](#), [36](#)

SQM_to_phyloseq, [35](#), [35](#)

subsetBins, [37](#)

subsetContigs, [38](#), [39](#), [46](#)

subsetFun, [5](#), [40](#), [44](#), [46](#)

subsetORFs, [38](#), [40](#), [41](#), [42](#), [44](#)

subsetRand, [43](#)

subsetSamples, [41](#), [44](#), [46](#)

subsetTax, [5](#), [41](#), [44](#), [45](#)

summary.SQM, [46](#)

summary.SQMbunch, [47](#)

summary.SQMlite, [47](#)

USiCGs, [3](#), [19](#), [20](#), [48](#)