

Package ‘amplican’

December 4, 2025

Type Package

Title Automated analysis of CRISPR experiments

Description `amplican` performs alignment of the amplicon reads, normalizes gathered data, calculates multiple statistics (e.g. cut rates, frameshifts) and presents results in form of aggregated reports. Data and statistics can be broken down by experiments, barcodes, user defined groups, guides and amplicons allowing for quick identification of potential problems.

Version 1.33.1

URL <https://github.com/valenlab/amplican>

BugReports <https://github.com/valenlab/amplican/issues>

biocViews ImmunoOncology, Technology, Alignment, qPCR, CRISPR

License GPL-3

LinkingTo Rcpp

Depends R (>= 3.5.0), methods, BiocGenerics (>= 0.22.0), Biostrings (>= 2.44.2), pwalgn

Imports Rcpp, utils (>= 3.4.1), S4Vectors (>= 0.14.3), ShortRead (>= 1.34.0), IRanges (>= 2.10.2), GenomicRanges (>= 1.61.1), Seqinfo, BiocParallel (>= 1.10.1), gtable (>= 0.2.0), gridExtra (>= 2.2.1), ggplot2 (>= 3.3.4), ggthemes (>= 3.4.0), stringr (>= 1.2.0), stats (>= 3.4.1), matrixStats (>= 0.52.2), Matrix (>= 1.2-10), dplyr (>= 0.7.2), data.table (>= 1.10.4-3), rmarkdown (>= 1.6), knitr (>= 1.16), cluster (>= 2.1.4)

RoxygenNote 7.3.3

Suggests testthat, BiocStyle, GenomicAlignments

Collate 'helpers_general.R' 'AlignmentsExperimentSet-class.R' 'RcppExports.R' 'helpers_rmd.R' 'amplicanReport.R' 'helpers_directory.R' 'helpers_warnings.R' 'helpers_filters.R' 'helpers_alignment.R' 'amplicanAlign.R' 'amplican.R' 'amplicanFilter.R' 'amplicanNormalize.R' 'amplicanSummarize.R' 'ggforce_bezier.R' 'helpers_plots.R'

VignetteBuilder knitr

Encoding UTF-8

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

git_branch devel

git_last_commit dd29ebe

git_last_commit_date 2025-11-11

Repository Bioconductor 3.23

Date/Publication 2025-12-04

Author Kornel Labun [aut],
Eivind Valen [cph, cre]

Maintainer Eivind Valen <eivind.valen@gmail.com>

Contents

AlignmentsExperimentSet-class	4
alphabetQuality	7
amplican	8
amplicanAlign	8
amplicanConsensus	11
amplicanFilter	12
amplicanMap	13
amplicanNormalize	14
amplicanOverlap	15
amplicanPipeline	16
amplicanPipelineConservative	19
amplicanReport	23
amplicanSummarize	24
amplican_print_reads	25
assignedCount	26
barcodeData	27
barcodeData<-	27
checkConfigFile	28
checkFileWriteAccess	28
checkPrimers	29
checkTarget	29
cigarsToEvents	30
comb_along	31
cumsumw	31
decode	32
defGR	32
experimentData	33
experimentData<-	34
extractEvents	34
findEOP	35
findLQR	36
findPD	37

flipRanges	38
fwdReads	38
fwdReads<-	39
fwdReadsType	39
fwdReadsType<-	40
geom_bezier	40
getEventInfo	44
getEvents	44
get_left_primer	45
get_right_primer	46
get_seq	46
goodAvgQuality	47
goodBaseQuality	47
is_hdr_strict	48
locate_pr_start	48
lookupAlignment	49
makeAlignment	50
metaplot_deletions	51
metaplot_insertions	52
metaplot_mismatches	53
pairToEvents	54
plot_amplicon	55
plot_cuts	55
plot_deletions	56
plot_height	58
plot_heterogeneity	58
plot_insertions	59
plot_mismatches	60
plot_variants	61
readCounts	63
readCounts<-	64
revComp	65
rveReads	65
rveReads<-	66
rveReadsType	66
rveReadsType<-	67
unassignedCount	67
unassignedData	68
unassignedData<-	69
upperGroups	69
waffle	70
writeAlignments	71

`AlignmentsExperimentSet-class`*An S4 class to represent alignments from multiple experiments*

Description

Class `AlignmentsExperimentSet` holds data from multiple alignments for many experiments. Allows to examine alignments in great detail.

Usage

```
AlignmentsExperimentSet(...)

## S4 method for signature 'AlignmentsExperimentSet'
length(x)

## S4 method for signature 'AlignmentsExperimentSet'
fwdReads(x)

## S4 replacement method for signature 'AlignmentsExperimentSet'
fwdReads(x) <- value

## S4 method for signature 'AlignmentsExperimentSet'
rveReads(x)

## S4 replacement method for signature 'AlignmentsExperimentSet'
rveReads(x) <- value

## S4 method for signature 'AlignmentsExperimentSet'
fwdReadsType(x)

## S4 replacement method for signature 'AlignmentsExperimentSet'
fwdReadsType(x) <- value

## S4 method for signature 'AlignmentsExperimentSet'
rveReadsType(x)

## S4 replacement method for signature 'AlignmentsExperimentSet'
rveReadsType(x) <- value

## S4 method for signature 'AlignmentsExperimentSet'
unassignedData(x)

## S4 replacement method for signature 'AlignmentsExperimentSet'
unassignedData(x) <- value

## S4 method for signature 'AlignmentsExperimentSet'
```

```
readCounts(x)

## S4 replacement method for signature 'AlignmentsExperimentSet'
readCounts(x) <- value

## S4 method for signature 'AlignmentsExperimentSet'
experimentData(x)

## S4 replacement method for signature 'AlignmentsExperimentSet'
experimentData(x) <- value

## S4 method for signature 'AlignmentsExperimentSet'
barcodeData(x)

## S4 replacement method for signature 'AlignmentsExperimentSet'
barcodeData(x) <- value

## S4 method for signature 'AlignmentsExperimentSet'
unassignedCount(x)

## S4 method for signature 'AlignmentsExperimentSet'
assignedCount(x)

## S4 method for signature 'AlignmentsExperimentSet'
names(x)

## S4 method for signature 'AlignmentsExperimentSet'
c(x, ...)

## S4 method for signature 'AlignmentsExperimentSet,numeric,missing,missing'
x[i, j, ..., drop = TRUE]

## S3 method for class 'AlignmentsExperimentSet'
as.list(x, ...)

## S4 method for signature 'AlignmentsExperimentSet'
x$name

## S4 method for signature 'AlignmentsExperimentSet'
writeAlignments(x, file = "", aln_format = "txt")

## S4 method for signature 'AlignmentsExperimentSet'
lookupAlignment(x, ID, read_id = 1)

## S4 method for signature 'AlignmentsExperimentSet'
extractEvents(object, use_parallel = FALSE)
```

Arguments

...	pass any number of AlignmentsExperimentSet objects, make sure experiment IDs can be unique after merging
x, object	(AlignmentsExperimentSet)
value	Represents assignment values for setter methods.
i, j, name, drop	(numeric, missing, character, missing) AlignmentsExperimentSet object can be subsetted using names of the experiments eg. x\$name or x[i], resulting in AlignmentsExperimentSet object that has only one experiment. During this subsetting, values of unassignedData and barcodeData are dropped.
file	(connection or string) Destination file. When empty, defaults to standard output.
aln_format	("txt" or "fasta") Specifies format of the file.
ID	(string) Experiment Identifier
read_id	(numeric) Read Identifier. Reads are sorted by frequency. Defaults to 1, most abundant read.
use_parallel	(boolean) When using extractEvents you can use multicore back-end through register as this is very slow function (despite vectorization).

Value

depending on the function used

Slots

fwdReads, rveReads	(list) Named list where each element is of class PairwiseAlignmentsSingleSubject . Names correspond to the experiment ID. Contains alignments of reads against amplicons.
fwdReadsType, rveReadsType	(list) Named list where each element is of logical vector, so far TRUE corresponds to HDR events. Names correspond to the experiment ID. Contains type of read - HDR/NHEJ.
readCounts	(list) Named list where each element is numeric vector that describes how many reads are compressed into unique representation before alignment in fwdReads and/or rveReads.
unassignedData	(data.frame) Contains reads that failed to be assigned to any of the experiments. Alignment of forward against reverse reads may give hint whether these reads are compromised in any way.
experimentData	(data.frame) Expands on configuration file and provides information about cut rates, frameshifts, PRIMER DIMER detection etc. Each row corresponds to experiment ID.
barcodeData	(data.frame) Information that is gathered on the barcode level is gathered in this data.frame, mainly quality filtering statistics.

View alignments

Write out all alignments in "fasta" or "txt" format.:

```
writeAlignments(x, file = "", aln_format = "txt")
```

Write out human readable alignments for given experiment and read_id.:

```
lookupAlignment(x, ID, read_id = 1)
```

Coercion based on events

Coerce to data.frame compatible with [GRanges](#) .:
 as.data.frame(x)

Examples

```
exampleAlignments <- pwalignment::pairwiseAlignment(
  Biostrings::DNAStringSet(c("ACTGACTG", "CGACGACG")), "ACGTACGTACGT")
new("AlignmentsExperimentSet",
  fwdReads = list(ID_1 = exampleAlignments, ID_2 = exampleAlignments),
  rveReads = list(ID_1 = exampleAlignments, ID_2 = exampleAlignments),
  fwdReadsType = list(ID_1 = c(FALSE, FALSE), ID_2 = c(FALSE, FALSE)),
  rveReadsType = list(ID_1 = c(FALSE, FALSE), ID_2 = c(FALSE, FALSE)),
  readCounts = list(ID_1 = c(2, 20), ID_2 = c(30, 100)),
  unassignedData = NULL,
  experimentData = data.frame(ID = c("ID_1", "ID_2"),
                                Barcode = c("B1", "B1"),
                                whatever = c(50, 100)),
  barcodeData = data.frame(Barcode = "B1", statistic1 = 100))
# Coercion
extractEvents(AlignmentsExperimentSet())
GenomicRanges::GRanges(extractEvents(AlignmentsExperimentSet()))
```

 alphabetQuality

This filters out sequences which have nonstandard nucleotides.

Description

This filters out sequences which have nonstandard nucleotides.

Usage

```
alphabetQuality(reads, batch_size = 1e+07)
```

Arguments

reads (ShortRead object) Loaded reads from fastq.
 batch_size (numeric) How many reads to process at a time.

Value

(boolean) Logical vector with the valid rows as TRUE.

amplican

Automated analysis of CRISPR experiments.

Description

Main goals:

1. Flexible pipeline for analysis of the CRISPR Mi-Seq or Hi-Seq data.
2. Compatible with GRanges and data.table style.
3. Precise quantification of mutation rates.
4. Prepare automatic reports as .Rmd files that are flexible and open for manipulation.
5. Provide specialized plots for deletions, insertions, mismatches, variants, heterogeneity of the reads.

Details

To learn more about amplican, start with the vignettes: `browseVignettes(package = "amplican")`

Author(s)

Maintainer: Eivind Valen <eivind.valen@gmail.com> [copyright holder]

Authors:

- Kornel Labun <kornel.labun@gmail.com>

See Also

Useful links:

- <https://github.com/valenlab/amplican>
- Report bugs at <https://github.com/valenlab/amplican/issues>

amplicanAlign

Align reads to amplicons.

Description

amplicanAlign takes a configuration files, fastq reads and output directory to prepare alignments and summary. It uses global Needleman-Wunsch algorithm with parameters optimized for CRISPR experiment. After alignments, object of `AlignmentsExperimentSet` is returned that allows for coercion into GRanges (plus is for forward and minus for reverse reads). It is also possible to output alignments in other, additional formats.

Usage

```

amplicanAlign(
  config,
  fastq_folder,
  use_parallel = FALSE,
  average_quality = 30,
  min_quality = 20,
  filter_n = FALSE,
  batch_size = 1e+06,
  scoring_matrix = pwalgn::nucleotideSubstitutionMatrix(match = 5, mismatch = -4,
    baseOnly = FALSE, type = "DNA"),
  gap_opening = 25,
  gap_extension = 0,
  fastqfiles = 0.5,
  primer_mismatch = 0,
  donor_mismatch = 3,
  donor_strict = FALSE
)

```

Arguments

config	(string) The path to your configuration file. For example: <code>system.file("extdata", "config.txt", package = "amplican")</code> . Configuration file can contain additional columns, but first 11 columns have to follow the example config specification.
fastq_folder	(string) Path to FASTQ files. If not specified, FASTQ files should be in the same directory as config file.
use_parallel	(boolean) Set to TRUE, if you have registered multicore back-end.
average_quality	(numeric) The FASTQ file have a quality for each nucleotide, depending on sequencing technology there exist many formats. This package uses readFastq to parse the reads. If the average quality of the reads fall below value of average_quality then sequence is filtered. Default is 0.
min_quality	(numeric) Similar as in average_quality, but depicts the minimum quality for ALL nucleotides in given read. If one of nucleotides has quality BELOW min_quality, then the sequence is filtered. Default is 20.
filter_n	(boolean) Whether to filter out reads containing N base.
batch_size	(numeric) How many reads to analyze at a time? Needed for filtering of large fastq files.
scoring_matrix	(matrix) Default is 'NUC44'. Pass desired matrix using nucleotideSubstitutionMatrix .
gap_opening	(numeric) The opening gap score.
gap_extension	(numeric) The gap extension score.
fastqfiles	(numeric) Normally you want to use both FASTQ files. But in some special cases, you may want to use only the forward file, or only the reverse file. Possible options:

- 0** Use both FASTQ files.
- 0.5** Use both FASTQ files, but only for one of the reads (forward or reverse) is required to have primer perfectly matched to sequence - eg. use when reverse reads are trimmed of primers, but forward reads have forward primer in the sequence.
- 1** Use only the forward FASTQ file.
- 2** Use only the reverse FASTQ file.

primer_mismatch

(numeric) Decide how many mismatches are allowed during primer matching of the reads, that groups reads by experiments. When primer_mismatch = 0 no mismatches are allowed, which can increase number of unassigned read.

donor_mismatch

(numeric) How many events of length 1 (mismatches, deletions and insertions of length 1) are allowed when aligning toward the donor template. This parameter is only used when donor template is specified. The higher the parameter the less strict will be algorithm accepting read as HDR. Set to 0 if only perfect alignments to the donor template marked as HDR, unadvised due to error rate of the sequencers.

donor_strict

(logical) Applies more strict algorithm for HDR detection. Only these reads that have all of the donor events will count as HDR. Tolerates 'donor_mismatch' level of noise, but no indels are allowed. Use this when your reads should span over the whole window of the donor events. Might be more time consuming.

Value

(AlignmentsExperimentSet) Check [AlignmentsExperimentSet](#) class for details. You can use [lookupAlignment](#) to examine alignments visually.

See Also

Other analysis steps: [amplicanConsensus\(\)](#), [amplicanFilter\(\)](#), [amplicanMap\(\)](#), [amplicanNormalize\(\)](#), [amplicanOverlap\(\)](#), [amplicanPipeline\(\)](#), [amplicanPipelineConservative\(\)](#), [amplicanReport\(\)](#), [amplicanSummarize\(\)](#)

Examples

```
# path to example config file
config <- system.file("extdata", "config.csv", package = "amplican")
# path to example fastq files
fastq_folder <- system.file("extdata", package = "amplican")
aln <- amplicanAlign(config, fastq_folder)
aln
```

amplicanConsensus	<i>Extract consensus out of forward and reverse events.</i>
-------------------	---

Description

When forward and reverse reads are in agreement on the events (eg. deletion) `amplicanConsensus` will mark forward event as `TRUE` indicating that he represents consensus. In cases where forward and reverse read agree only partially, for example, they share the same start of the deletion, but they have different end `amplicanConsensus` will pick the version of read with higher alignment score, in situation where both of the reads overlap expected cut site, otherwise both events will be rejected and marked `FALSE`. When there are events only on one of the strands they will be rejected.

Usage

```
amplicanConsensus(aln, cfgT, overlaps = "overlaps", promiscuous = TRUE)
```

Arguments

<code>aln</code>	(data.frame) Contains relevant events in GRanges style.
<code>cfgT</code>	(data.frame) Should be table containing at least positions of primers in the amplicons and their identifiers
<code>overlaps</code>	(character) Specifies which metadata column of <code>aln</code> indicates which events are overlapping expected cut site.
<code>promiscuous</code>	(boolean) Allows to relax consensus rules. When <code>TRUE</code> will allow Indels that are not confirmed by the other strand (when both are used).

Details

In situation where you have only forward or only reverse reads don't use this function and assign all `TRUE` to all of your events.

Consensus out of the forward + reverse reads is required for `amplicanSummary`, and `amplicanConsensus` requires `amplicanOverlap`.

Value

(boolean vector) Where `TRUE` means that given event represents consensus out of forward and reverse reads.

See Also

Other analysis steps: [amplicanAlign\(\)](#), [amplicanFilter\(\)](#), [amplicanMap\(\)](#), [amplicanNormalize\(\)](#), [amplicanOverlap\(\)](#), [amplicanPipeline\(\)](#), [amplicanPipelineConservative\(\)](#), [amplicanReport\(\)](#), [amplicanSummarize\(\)](#)

Examples

```
file_path <- system.file("test_data", "test_aln.csv", package = "amplican")
aln <- data.table::fread(file_path)
cfgT <- data.table::fread(
  system.file("test_data", "test_cfg.csv", package = "amplican"))
all(aln$consensus == amplicanConsensus(aln, cfgT))
```

amplicanFilter	<i>Filter Events Overlapping Primers, PRIMER DIMERS and Low Alignment Score Events.</i>
----------------	---

Description

Very often alignments return deletions that are not real deletions, but rather artifact of incomplete reads eg.:

```
ACTGAAAA----- <- this "deletion" should be filtered
ACTG----ACTGACTG
```

We call them Events Overlapping Primers and filter them together with reads that are potentially PRIMER DIMERS. This filter will also remove all events coming from reads with low alignment score - potential Off-targets.

Usage

```
amplicanFilter(aln, cfgT, PRIMER_DIMER)
```

Arguments

aln	(data.frame) Should contain events from alignments in GRanges style with columns eg. seqnames, width, start, end.
cfgT	(data.frame) Needs columns Forward_Primer, ReversePrimer and Amplicon.
PRIMER_DIMER	(numeric) Value specifying buffer for PRIMER DIMER detection. For a given read it will be recognized as PRIMER DIMER when alignment will introduce gap of size bigger than: length of amplicon - (lengths of PRIMERS + PRIMER_DIMER value)

Value

(aln) Reduced by events classified as PRIMER DIMER or overlapping primers.

See Also

[findPD](#) and [findEOP](#)

Other analysis steps: [amplicanAlign\(\)](#), [amplicanConsensus\(\)](#), [amplicanMap\(\)](#), [amplicanNormalize\(\)](#), [amplicanOverlap\(\)](#), [amplicanPipeline\(\)](#), [amplicanPipelineConservative\(\)](#), [amplicanReport\(\)](#), [amplicanSummarize\(\)](#)

Examples

```

file_path <- system.file("extdata", "results", "alignments",
                          "raw_events.csv", package = "amplican")
aln <- data.table::fread(file_path)
cfgT <- data.table::fread(
  system.file("extdata", "results", "config_summary.csv",
              package = "amplican"))
amplicanFilter(aln, cfgT, 30)

```

amplicanMap	<i>Map events to their respective relative coordinates specified with UPPER case.</i>
-------------	---

Description

Translate coordinates of [GRanges](#) events so that they can be relative to the amplicon. As point zero we assume first left sided UPPER case letter in the amplicon. Be weary that events for amplicons without expected cut sites are filtered. Don't use this function, if you don't have expected cut sites specified and don't use any of the metaplots.

Usage

```
amplicanMap(aln, cfgT)
```

Arguments

aln	(data.frame) List of events to map to the relative coordinates.
cfgT	(data.frame) config table

Value

([GRanges](#)) Same as events, but the coordinates are relative to the expected cut sites.

See Also

Other analysis steps: [amplicanAlign\(\)](#), [amplicanConsensus\(\)](#), [amplicanFilter\(\)](#), [amplicanNormalize\(\)](#), [amplicanOverlap\(\)](#), [amplicanPipeline\(\)](#), [amplicanPipelineConservative\(\)](#), [amplicanReport\(\)](#), [amplicanSummarize\(\)](#)

Examples

```

# example config
config <- read.csv(system.file("extdata", "config.csv",
                              package = "amplican"))

# example events
events <- read.csv(system.file("extdata", "results", "alignments",
                              "raw_events.csv", package = "amplican"))

```

```
# make events relative to the UPPER case
amplicanMap(events, config)
```

amplicanNormalize	<i>Remove events that can be found in Controls.</i>
-------------------	---

Description

This function can adjust events for small differences between known annotations (amplicon sequences) and real DNA of the strain that was sequenced. Events from the control are grouped by add and their frequencies are calculated in respect to number of total reads in that groups. In next step events from the control are filtered according to min_freq, all events below are treated as sequencing errors and rejected. Finally, all events that can be found in treatment group that find their exact match (by non skipped columns) in control group are removed. All events from control group are returned back.

Usage

```
amplicanNormalize(
  aln,
  cfgT,
  add = c("guideRNA", "Group"),
  skip = c("counts", "score", "seqnames", "read_id", "strand", "overlaps", "consensus"),
  min_freq = 0.01
)
```

Arguments

aln	(data.frame) Contains events from alignments.
cfgT	(data.frame) Config table with information about experiments.
add	(character vector) Columns from cfgT that should be included in event table for normalization matching. Defaults to c("guideRNA", "Group") , which means that only those events created by the same guideRNA in the same Group will be removed if found in Control.
skip	(character vector) Specifies which columns of aln to skip.
min_freq	(numeric) All events from control group below this frequency will be not included in filtering. Use this to filter out background noise and sequencing errors.

Value

(data.frame) Same as aln, but events are normalized. Events from Control are not changed. Additionally columns from add are added to the data.frame.

See Also

Other analysis steps: `amplicanAlign()`, `amplicanConsensus()`, `amplicanFilter()`, `amplicanMap()`, `amplicanOverlap()`, `amplicanPipeline()`, `amplicanPipelineConservative()`, `amplicanReport()`, `amplicanSummarize()`

Examples

```
aln <- data.frame(seqnames = 1:5, start = 1, end = 2, width = 2,
                  counts = 101:105)
cfgT <- data.frame(ID = 1:5, guideRNA = rep("ACTG", 5),
                  Reads_Filtered = c(2, 2, 3, 3, 4),
                  Group = c("A", "A", "B", "B", "B"),
                  Control = c(TRUE, FALSE, TRUE, FALSE, FALSE))
# all events are same as in the control group, therefore are filtered out
# events from control groups stay
amplicanNormalize(aln, cfgT)
# events that are different from control group are preserved
aln[2, "start"] <- 3
amplicanNormalize(aln, cfgT)
```

<code>amplicanOverlap</code>	<i>Check which events overlap expected cut sites.</i>
------------------------------	---

Description

To determine which deletions, insertions and mismatches (events) are probably created by CRISPR we check whether they overlap expected cut sites. Expected cut sites should be specified in UPPER CASE letters in the amplicon sequences.

Usage

```
amplicanOverlap(aln, cfgT, cut_buffer = 5, relative = FALSE)
```

Arguments

<code>aln</code>	(data.frame) Contains relevant events in GRanges style.
<code>cfgT</code>	(data.frame) Contains amplicon sequences.
<code>cut_buffer</code>	(numeric) Number of bases that should expand 5' and 3' of the specified expected cut sites.
<code>relative</code>	(boolean) Sets whether events are relative to the position of the target site.

Value

(boolean vector) Where TRUE means that given event overlaps cut site.

See Also

Other analysis steps: [amplicanAlign\(\)](#), [amplicanConsensus\(\)](#), [amplicanFilter\(\)](#), [amplicanMap\(\)](#), [amplicanNormalize\(\)](#), [amplicanPipeline\(\)](#), [amplicanPipelineConservative\(\)](#), [amplicanReport\(\)](#), [amplicanSummarize\(\)](#)

Examples

```
file_path <- system.file("test_data", "test_aln.csv", package = "amplican")
aln <- data.table::fread(file_path)
cfgT <- data.table::fread(
  system.file("test_data", "test_cfg.csv", package = "amplican"))
all(aln$overlaps == amplicanOverlap(aln, cfgT))
```

amplicanPipeline

Wraps main package functionality into one function.

Description

amplicanPipeline is convenient wrapper around all functionality of the package with the most robust settings. It will generate all results in the `result_folder` and also knit prepared reports into 'reports' folder.

Usage

```
amplicanPipeline(
  config,
  fastq_folder,
  results_folder,
  knit_reports = TRUE,
  write_alignments_format = "None",
  average_quality = 30,
  min_quality = 0,
  filter_n = FALSE,
  batch_size = 1e+07,
  use_parallel = FALSE,
  scoring_matrix = pwalgn::nucleotideSubstitutionMatrix(match = 5, mismatch = -4,
    baseOnly = FALSE, type = "DNA"),
  gap_opening = 25,
  gap_extension = 0,
  fastqfiles = 0.5,
  primer_mismatch = 2,
  donor_mismatch = 3,
  donor_strict = FALSE,
  PRIMER_DIMER = 30,
  event_filter = TRUE,
  cut_buffer = 5,
```



```

promiscuous_consensus = TRUE,
normalize = c("guideRNA", "Group"),
min_freq = min_freq_default,
continue = TRUE
)

```

Arguments

- config** (string) The path to your configuration file. For example: `system.file("extdata", "config.txt", package = "amplican")`. Configuration file can contain additional columns, but first 11 columns have to follow the example config specification.
- fastq_folder** (string) Path to FASTQ files. If not specified, FASTQ files should be in the same directory as config file.
- results_folder** (string) Where do you want to store results? The package will create files in that folder so make sure you have writing permissions.
- knit_reports** (boolean) whether function should "knit" all reports automatically for you (it is time consuming, be patient), when false reports will be prepared, but not knitted
- write_alignments_format** (character vector) Whether amplicanPipeline should write alignments results to separate files. Alignments are also always saved as `.rds` object of [AlignmentsExperimentSet](#) class. Possible options are:
- "fasta"** outputs alignments in fasta format where header indicates experiment ID, read id and number of reads
 - "txt"** simple format, read information followed by forward read and amplicon sequence followed by reverse read with its amplicon sequence eg.:


```

ID: ID_1 Count: 7
ACTGAAAAA-----
ACTG-----ACTGACTG

-----G-ACTG
ACTGACTGACTG

```
 - "None"** Don't write any alignments to files.
 - c("fasta", "txt")** There are also possible combinations of above formats, pass a vector to get alignments in multiple formats.
- average_quality** (numeric) The FASTQ file have a quality for each nucleotide, depending on sequencing technology there exist many formats. This package uses [readFastq](#) to parse the reads. If the average quality of the reads fall below value of `average_quality` then sequence is filtered. Default is 0.
- min_quality** (numeric) Similar as in `average_quality`, but depicts the minimum quality for ALL nucleotides in given read. If one of nucleotides has quality BELOW `min_quality`, then the sequence is filtered. Default is 20.
- filter_n** (boolean) Whether to filter out reads containing N base.

batch_size	(numeric) How many reads to analyze at a time? Needed for filtering of large fastq files.
use_parallel	(boolean) Set to TRUE, if you have registered multicore back-end.
scoring_matrix	(matrix) Default is 'NUC44'. Pass desired matrix using nucleotideSubstitutionMatrix .
gap_opening	(numeric) The opening gap score.
gap_extension	(numeric) The gap extension score.
fastqfiles	(numeric) Normally you want to use both FASTQ files. But in some special cases, you may want to use only the forward file, or only the reverse file. Possible options: 0 Use both FASTQ files. 0.5 Use both FASTQ files, but only for one of the reads (forward or reverse) is required to have primer perfectly matched to sequence - eg. use when reverse reads are trimmed of primers, but forward reads have forward primer in the sequence. 1 Use only the forward FASTQ file. 2 Use only the reverse FASTQ file.
primer_mismatch	(numeric) Decide how many mismatches are allowed during primer matching of the reads, that groups reads by experiments. When primer_mismatch = 0 no mismatches are allowed, which can increase number of unassigned read.
donor_mismatch	(numeric) How many events of length 1 (mismatches, deletions and insertions of length 1) are allowed when aligning toward the donor template. This parameter is only used when donor template is specified. The higher the parameter the less strict will be algorithm accepting read as HDR. Set to 0 if only perfect alignments to the donor template marked as HDR, unadvised due to error rate of the sequencers.
donor_strict	(logical) Applies more strict algorithm for HDR detection. Only these reads that have all of the donor events will count as HDR. Tolerates 'donor_mismatch' level of noise, but no indels are allowed. Use this when your reads should span over the whole window of the donor events. Might be more time consuming.
PRIMER_DIMER	(numeric) Value specifying buffer for PRIMER DIMER detection. For a given read it will be recognized as PRIMER DIMER when alignment will introduce gap of size bigger than: length of amplicon - (lengths of PRIMERS + PRIMER_DIMER value)
event_filter	(logical) Whether detection of offtarget reads, should be enabled.
cut_buffer	The number of bases by which extend expected cut sites (specified as UPPER case letters in the amplicon) in 5' and 3' directions.
promiscuous_consensus	(boolean) Whether rules of amplicanConsensus should be promiscuous. When promiscuous, we allow indels that have no confirmation on the other strand.
normalize	(character vector) If column 'Control' in config table has all FALSE/0 values then normalization is skipped. Otherwise, normalization is strict, which means events that are found in 'Control' TRUE group will be removed in 'Control' FALSE group. This parameter by default uses columns 'guideRNA' and 'Group'

	to impose additional restrictions on normalized events eg. only events created by the same 'guideRNA' in the same 'Group' will be normalized.
min_freq	(numeric) All events below this frequency are treated as sequencing errors and rejected. This parameter is used during normalization through amplicanNormalize .
continue	(boolean) Default TRUE, decides whether to continue failed ampliCan runs. In case of FALSE, all contents in 'results' folder will be removed.

Value

(invisible) results_folder path

See Also

Other analysis steps: [amplicanAlign\(\)](#), [amplicanConsensus\(\)](#), [amplicanFilter\(\)](#), [amplicanMap\(\)](#), [amplicanNormalize\(\)](#), [amplicanOverlap\(\)](#), [amplicanPipelineConservative\(\)](#), [amplicanReport\(\)](#), [amplicanSummarize\(\)](#)

Examples

```
# path to example config file
config <- system.file("extdata", "config.csv", package = "amplican")
# path to example fastq files
fastq_folder <- system.file("extdata", package = "amplican")
# output folder
results_folder <- tempdir()

#full analysis, not knitting files automatically
amplicanPipeline(config, fastq_folder, results_folder, knit_reports = FALSE)
```

amplicanPipelineConservative

Wraps main package functionality into one function.

Description

amplicanPipelineIndexHopping is identical as amplicanPipeline except that default min_freq threshold is set to 0.15. Setting this threshold higher will decrease risks of inadequate normalization in cases of potential Index Hopping, potentially decreasing precision of true editing rate calling. Index Hopping can be mitigated with use of unique dual indexing pooling combinations. However, in cases when you might expect Index Hopping to occur you should use this function instead of amplicanPipeline.

Usage

```

amplicanPipelineConservative(
  config,
  fastq_folder,
  results_folder,
  knit_reports = TRUE,
  write_alignments_format = "None",
  average_quality = 30,
  min_quality = 0,
  filter_n = FALSE,
  batch_size = 1e+07,
  use_parallel = FALSE,
  scoring_matrix = pwalig::nucleotideSubstitutionMatrix(match = 5, mismatch = -4,
    baseOnly = FALSE, type = "DNA"),
  gap_opening = 25,
  gap_extension = 0,
  fastqfiles = 0.5,
  primer_mismatch = 2,
  donor_mismatch = 3,
  donor_strict = FALSE,
  PRIMER_DIMER = 30,
  event_filter = TRUE,
  cut_buffer = 5,
  promiscuous_consensus = TRUE,
  normalize = c("guideRNA", "Group"),
  min_freq = min_freq_default,
  continue = TRUE
)

```

Arguments

- | | |
|-------------------------|--|
| config | (string) The path to your configuration file. For example: <code>system.file("extdata", "config.txt", package = "amplican")</code> . Configuration file can contain additional columns, but first 11 columns have to follow the example config specification. |
| fastq_folder | (string) Path to FASTQ files. If not specified, FASTQ files should be in the same directory as config file. |
| results_folder | (string) Where do you want to store results? The package will create files in that folder so make sure you have writing permissions. |
| knit_reports | (boolean) whether function should "knit" all reports automatically for you (it is time consuming, be patient), when false reports will be prepared, but not knitted |
| write_alignments_format | (character vector) Whether amplicanPipeline should write alignments results to separate files. Alignments are also always saved as <code>.rds</code> object of AlignmentsExperimentSet class. Possible options are:
"fasta" outputs alignments in fasta format where header indicates experiment ID, read id and number of reads |

"txt" simple format, read information followed by forward read and amplicon sequence followed by reverse read with its amplicon sequence eg.:

```
ID: ID_1 Count: 7
ACTGAAAAA-----
ACTG-----ACTGACTG
```

```
-----G-ACTG
ACTGACTGACTG
```

"None" Don't write any alignments to files.

c("fasta", "txt") There are also possible combinations of above formats, pass a vector to get alignments in multiple formats.

average_quality

(numeric) The FASTQ file have a quality for each nucleotide, depending on sequencing technology there exist many formats. This package uses [readFastq](#) to parse the reads. If the average quality of the reads fall below value of average_quality then sequence is filtered. Default is 0.

min_quality

(numeric) Similar as in average_quality, but depicts the minimum quality for ALL nucleotides in given read. If one of nucleotides has quality BELOW min_quality, then the sequence is filtered. Default is 20.

filter_n

(boolean) Whether to filter out reads containing N base.

batch_size

(numeric) How many reads to analyze at a time? Needed for filtering of large fastq files.

use_parallel

(boolean) Set to TRUE, if you have registered multicore back-end.

scoring_matrix

(matrix) Default is 'NUC44'. Pass desired matrix using [nucleotideSubstitutionMatrix](#).

gap_opening

(numeric) The opening gap score.

gap_extension

(numeric) The gap extension score.

fastqfiles

(numeric) Normally you want to use both FASTQ files. But in some special cases, you may want to use only the forward file, or only the reverse file. Possible options:

0 Use both FASTQ files.

0.5 Use both FASTQ files, but only for one of the reads (forward or reverse) is required to have primer perfectly matched to sequence - eg. use when reverse reads are trimmed of primers, but forward reads have forward primer in the sequence.

1 Use only the forward FASTQ file.

2 Use only the reverse FASTQ file.

primer_mismatch

(numeric) Decide how many mismatches are allowed during primer matching of the reads, that groups reads by experiments. When primer_mismatch = 0 no mismatches are allowed, which can increase number of unassigned read.

donor_mismatch

(numeric) How many events of length 1 (mismatches, deletions and insertions of length 1) are allowed when aligning toward the donor template. This parameter is only used when donor template is specified. The higher the parameter the

	less strict will be algorithm accepting read as HDR. Set to 0 if only perfect alignments to the donor template marked as HDR, unadvised due to error rate of the sequencers.
donor_strict	(logical) Applies more strict algorithm for HDR detection. Only these reads that have all of the donor events will count as HDR. Tolerates 'donor_mismatch' level of noise, but no indels are allowed. Use this when your reads should span over the whole window of the donor events. Might be more time consuming.
PRIMER_DIMER	(numeric) Value specifying buffer for PRIMER DIMER detection. For a given read it will be recognized as PRIMER DIMER when alignment will introduce gap of size bigger than: length of amplicon - (lengths of PRIMERS + PRIMER_DIMER value)
event_filter	(logical) Whether detection of offtarget reads, should be enabled.
cut_buffer	The number of bases by which extend expected cut sites (specified as UPPER case letters in the amplicon) in 5' and 3' directions.
promiscuous_consensus	(boolean) Whether rules of amplicanConsensus should be promiscuous. When promiscuous, we allow indels that have no confirmation on the other strand.
normalize	(character vector) If column 'Control' in config table has all FALSE/0 values then normalization is skipped. Otherwise, normalization is strict, which means events that are found in 'Control' TRUE group will be removed in 'Control' FALSE group. This parameter by default uses columns 'guideRNA' and 'Group' to impose additional restrictions on normalized events eg. only events created by the same 'guideRNA' in the same 'Group' will be normalized.
min_freq	(numeric) All events below this frequency are treated as sequencing errors and rejected. This parameter is used during normalization through amplicanNormalize .
continue	(boolean) Default TRUE, decides whether to continue failed amplican runs. In case of FALSE, all contents in 'results' folder will be removed.

Details

result_folder and also knit prepared reports into 'reports' folder.

Value

(invisible) results_folder path

See Also

Other analysis steps: [amplicanAlign\(\)](#), [amplicanConsensus\(\)](#), [amplicanFilter\(\)](#), [amplicanMap\(\)](#), [amplicanNormalize\(\)](#), [amplicanOverlap\(\)](#), [amplicanPipeline\(\)](#), [amplicanReport\(\)](#), [amplicanSummarize\(\)](#)

amplicanReport	<i>Prepare reports as .Rmd files.</i>
----------------	---------------------------------------

Description

amplicanReport takes a configuration file, fastq reads and output directory to prepare summaries as an editable .Rmd file. You can specify whether you want to make summaries based on ID, Barcode, Group or even guideRNA and Amplicon. This function automatically knits all reports after creation. If you want to postpone knitting and edit reports, use .Rmd templates to create your own version of reports instead of this function.

Usage

```
amplicanReport(
  results_folder,
  levels = c("id", "barcode", "group", "guide", "amplicon", "summary"),
  report_files = c("id_report", "barcode_report", "group_report", "guide_report",
    "amplicon_report", "index"),
  cut_buffer = 5,
  xlab_spacing = 4,
  top = 5,
  knit_reports = TRUE
)
```

Arguments

results_folder	(string) Folder containing results from the amplicanAlign function, do not change names of the files.
levels	(vector) Possible values are: "id", "barcode", "group", "guide", "amplicon", "summary". You can also input more than one value eg. c("id", "barcode") will create two separate reports for each level.
report_files	(vector) You can supply your own names of the files. For each of the levels there has to be one file name. Files are created in current working directory by default.
cut_buffer	(numeric) Default 5. A number of bases that is used around the specified cut site.
xlab_spacing	(numeric) Default is 4. Spacing of the ticks on the x axis of plots.
top	(numeric) Default is 5. How many of the top most frequent unassigned reads to report? It is only relevant when you used forward and reverse reads. We align them to each other as we could not specify correct amplicon.
knit_reports	(boolean) Whether to knit reports automatically.

Value

(string) Path to the folder with results.

See Also

Other analysis steps: `amplicanAlign()`, `amplicanConsensus()`, `amplicanFilter()`, `amplicanMap()`, `amplicanNormalize()`, `amplicanOverlap()`, `amplicanPipeline()`, `amplicanPipelineConservative()`, `amplicanSummarize()`

Examples

```
results_folder <- tempdir()
amplicanReport(results_folder, report_files = file.path(results_folder,
                                                         c("id_report",
                                                           "barcode_report",
                                                           "group_report",
                                                           "guide_report",
                                                           "amplicon_report",
                                                           "index")),
               knit_reports = FALSE)
```

<code>amplicanSummarize</code>	<i>Summarize how many reads have frameshift and how many reads have deletions.</i>
--------------------------------	--

Description

Before using this function make sure events are filtered to represent consensus with `amplicanConsensus`, if you use both forward and reverse reads. If you want to calculate metrics over expected cut site, filter events using `amplicanOverlap`.

Usage

```
amplicanSummarize(aln, cfgT)
```

Arguments

<code>aln</code>	(data.frame) Contains events from the alignments.
<code>cfgT</code>	(data.frame) Config file with the experiments details.

Details

Adds columns to `cfgT`:

HDR Count of reads identified as Homology Directed Repair events.

Reads_Del Count of reads containing at least one deletion.

Reads_In Count of reads containing at least one insertion.

Reads_Edited Count of reads with any edit (insertion, deletion, or HDR).

Reads_Frameshifted Count of reads with a frameshift (net indel length is not a multiple of 3).

Value

(data.frame) As cfgT, but with extra columns.

See Also

Other analysis steps: [amplicanAlign\(\)](#), [amplicanConsensus\(\)](#), [amplicanFilter\(\)](#), [amplicanMap\(\)](#), [amplicanNormalize\(\)](#), [amplicanOverlap\(\)](#), [amplicanPipeline\(\)](#), [amplicanPipelineConservative\(\)](#), [amplicanReport\(\)](#)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                        "events_filtered_shifted_normalized.csv",
                        package = "amplican")
aln <- data.table::fread(file_path)
cfgT <- data.table::fread(
  system.file("extdata", "results", "config_summary.csv",
              package = "amplican"))
amplicanSummarize(aln, cfgT)
```

`amplican_print_reads` *Pretty print forward and reverse reads aligned to each other.*

Description

Usefull and needed for barcode reports.

Usage

```
amplican_print_reads(forward, reverse)
```

Arguments

<code>forward</code>	(character or vector of characters) Forward reads.
<code>reverse</code>	(character or vector of characters) Will be reverse complemented before alignment.

Value

Vector with alignments ready to be printed.

Examples

```
# load example data
unassigned_file <- system.file('extdata', 'results', 'alignments',
                              'unassigned_reads.csv', package = 'amplican')
unassigned <- data.table::setDF(data.table::fread(unassigned_file))
# sort by frequency
unassigned <- unassigned[order(unassigned$BarcodeFrequency,
                              decreasing = TRUE), ]
# print alignment of most frequent unassigned reads
cat(amplican_print_reads(unassigned[1, 'Forward'],
                        unassigned[1, 'Reverse']),
    sep = "\n")
```

assignedCount	<i>Get count of assigned reads.</i>
---------------	-------------------------------------

Description

Get count of assigned reads.

Usage

```
assignedCount(x)
```

Arguments

x (AlignmentsExperimentSet)

Value

(numeric)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                        "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
writeAlignments(aln, file.path(tempdir(), "aln.txt"))
```

barcodeData	<i>Barcode data.</i>
-------------	----------------------

Description

Get barcode data.frame with information on the barcode level.

Usage

```
barcodeData(x)
```

Arguments

x (AlignmentsExperimentSet)

Value

(data.tableOrNULL)

Examples

```
file_path <- system.file("extdata", "results", "alignments",  
                          "AlignmentsExperimentSet.rds", package = "amplican")  
aln <- readRDS(file_path)  
barcodeData(aln)
```

barcodeData<-	<i>Barcode data.</i>
---------------	----------------------

Description

Set barcode data.frame with information on the barcode level.

Usage

```
barcodeData(x) <- value
```

Arguments

x (AlignmentsExperimentSet)
value (data.frame)

Value

(AlignmentsExperimentSet)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
barcodeData(aln) <- barcodeData(aln) #replace with the same values as before
```

checkConfigFile	<i>Pre-process a config file and checks that everything is in order.</i>
-----------------	--

Description

Its takes care of the following: No IDs are duplicated. Every combination of barcode, forward primer and reverse primer is unique. Each barcode has unique forward reads file and reverse read files. Checks that the read files exist with read access.

Usage

```
checkConfigFile(configTable, fastq_folder)
```

Arguments

configTable (data.frame) Config file.
 fastq_folder (string) Path to fastq folder.

Value

(boolean) TRUE, If anything goes wrong stops and prints error.

checkFileWriteAccess	<i>Checks if the given directory exist and can be written to.</i>
----------------------	---

Description

Checks if the given directory exist and can be written to.

Usage

```
checkFileWriteAccess(filePath)
```

Arguments

filePath (string) A string the path to the file.

Value

(invisible) TRUE, Stop if no access.

checkPrimers	<i>Checks if the forward and reverse primer are in the amplicon and where they are located.</i>
--------------	---

Description

Checks if the forward and reverse primer are in the amplicon and where they are located.

Usage

```
checkPrimers(configTable, fastqfiles)
```

Arguments

configTable	(data.frame) A data frame of config file.
fastqfiles	(numeric) Which primers are important.

Value

configTable (data.frame) A data frame of config file with additional fields for start locations of the primers

checkTarget	<i>Checks if the guideRNA is in the amplicon.</i>
-------------	---

Description

Checks if the guideRNA is in the amplicon.

Usage

```
checkTarget(configTable)
```

Arguments

configTable	(data.frame) data frame of config file
-------------	--

Value

(boolean vector) Prints warning when some guides can't be found.

cigarsToEvents	<i>Transform extended CIGAR strings into GRanges.</i>
----------------	---

Description

Transform extended CIGAR strings into [GRanges](#) representation with events of deletions, insertions and mismatches.

Usage

```
cigarsToEvents(  
  cigars,  
  aln_pos_start,  
  query_seq,  
  ref,  
  read_id,  
  mapq,  
  seqnames,  
  strands,  
  counts  
)
```

Arguments

cigars	(character) Extended CIGARS.
aln_pos_start	(integer) Pos of CIGARS.
query_seq	(character) Aligned query sequences.
ref	(character) Reference sequences used for alignment.
read_id	(numeric) Read id for assignment for each of the CIGARS.
mapq	(numeric) Mapping scores.
seqnames	(character) Names of the sequences, potentially ids of the reference sequences.
strands	(character) Strands to assign.
counts	(integer) Vector of cigar counts, if data collapsed.

Value

([GRanges](#)) Same as events.

comb_along	<i>Generate all combinations along string exchanging m characters at a time with dictionary letters.</i>
------------	--

Description

Generate all combinations along string seq swapping m characters at a time with letters defined in dictionary letters. Allows, for instance, to create a list of possible primers with two mismatches.

Usage

```
comb_along(seq, m = 2, letters = c("A", "C", "T", "G"))
```

Arguments

seq	(character) input character to permute
m	(integer) number of elements to permute at each step
letters	(character vector) dictionary source for combinations of elements

Value

(character vector) all unique combinations of permuted string

Examples

```
comb_along("AC")
comb_along("AAA", 1)
comb_along("AAA")
comb_along("AAA", 3)
comb_along("AAAAAAAA")
```

cumsumw	<i>Cumulative sum to calculate shift</i>
---------	--

Description

Cumulative sum to calculate shift

Usage

```
cumsumw(x)
```

Arguments

x	(IRanges)
---	-----------

Value

(numeric vector)

decode	<i>Get codons for given string - translate</i>
--------	--

Description

Get codons for given string - translate

Usage

decode(x)

Arguments

x (string)

Value

(string) codons

defGR	<i>Helper to construct GRanges with additional metadata columns.</i>
-------	--

Description

Helper to construct GRanges with additional metadata columns.

Usage

```
defGR(  
  x,  
  ID,  
  score,  
  strand_info = "+",  
  type = "deletion",  
  originally = "",  
  replacement = ""  
)
```


Arguments

x	(IRanges) names(x) indicating read_id
ID	(string)
score	(numeric) scores from the alignments
strand_info	(string) Either '+', '-'
type	(string)
originally	(string) Base pairs on the amplicon.
replacement	(string) Base pairs on the read.

Value

(GRanges) Object with meta-data

experimentData	<i>Experiment data.</i>
----------------	-------------------------

Description

Get experiment data.frame with information on the experiment level.

Usage

```
experimentData(x)
```

Arguments

x	(AlignmentsExperimentSet)
---	---------------------------

Value

(data.frameOrNULL)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
experimentData(aln)
```

```
experimentData<-      Experiment data.
```

Description

Set experiment data.frame with information on the experiment level.

Usage

```
experimentData(x) <- value
```

Arguments

```
x                (AlignmentsExperimentSet)
value            (data.frame)
```

Value

(AlignmentsExperimentSet)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
experimentData(aln) <- experimentData(aln) # replace with the same values
```

```
extractEvents      Extract AlignmentsExperimentSet events into data.frame.
```

Description

Extracts events (insertions, deletions, mismatches) from alignments into data.frame. Can use multiple cores as process is quite slow. All events are relative towards forward strand. "-" in strand column indicates which events were from reverse reads.

Usage

```
extractEvents(object, use_parallel = FALSE)
```

Arguments

```
object           (AlignmentsExperimentSet)
use_parallel     (boolean) Set to TRUE, if you have registered multicore back-end with register.
```

Value

(data.frame) Compatible with [GRanges](#) style.

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
extractEvents(aln)
```

findEOP

Find Events Overlapping Primers.

Description

Very often alignments return deletions that are not real deletions, but rather artifact of incomplete reads eg.:

```
ACTGAAAAA----- <- this "deletion" should be filtered
ACTG----ACTGACTG
```

Usage

```
findEOP(aln, cfgT)
```

Arguments

aln	(data.frame) Should contain events from alignments in GRanges style with columns eg. seqnames, width, start, end.
cfgT	(data.frame) Needs columns Forward_Primer, ReversePrimer and Amplicon.

Value

(logical vector) where TRUE indicates events that are overlapping primers

See Also

[findPD](#) [findLQR](#)

Other filters: [findLQR\(\)](#), [findPD\(\)](#)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "raw_events.csv", package = "amplican")
aln <- data.table::fread(file_path)
cfgT <- data.table::fread(
  system.file("extdata", "results", "config_summary.csv",
              package = "amplican"))
findEOP(aln, cfgT)
```

findLQR

Find Off-targets and Fragmented alignments from reads.

Description

Will try to detect off-targets and low quality alignments (outliers). It tries k-means clustering on normalized number of events per read and read alignment score. If there are 3 clusters (decided based on silhouette criterion) cluster with high event count and low alignment score will be marked for filtering. When there is less than 1000 scores in aln it will filter nothing.

Usage

```
findLQR(aln)
```

Arguments

aln (data.frame) Should contain events from alignments in GRanges style with columns eg. seqnames, width, start, end, score.

Value

(logical vector) where TRUE indicates events that are potential off-targets or low quality alignments.

See Also

[findPD](#) [findEOP](#)

Other filters: [findEOP\(\)](#), [findPD\(\)](#)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "raw_events.csv", package = "amplican")
aln <- data.table::fread(file_path)
aln <- aln[seqnames == "ID_1"] # for first experiment
findLQR(aln)
```

findPD	<i>Find PRIMER DIMER reads.</i>
--------	---------------------------------

Description

Use to filter reads that are most likely PRIMER DIMERS.

Usage

```
findPD(aln, cfgT, PRIMER_DIMER = 30)
```

Arguments

aln	(data.frame) Should contain events from alignments in GRanges style with columns eg. seqnames, width, start, end.
cfgT	(data.frame) Needs columns Forward_Primer, ReversePrimer and Amplicon.
PRIMER_DIMER	(numeric) Value specifying buffer for PRIMER DIMER detection. For a given read it will be recognized as PRIMER DIMER when alignment will introduce gap of size bigger than: length of amplicon - (lengths of PRIMERS + PRIMER_DIMER value)

Value

(logical) Where TRUE indicates event classified as PRIMER DIMER

See Also

[findEOP](#) [findLQR](#)

Other filters: [findEOP\(\)](#), [findLQR\(\)](#)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "raw_events.csv", package = "amplican")
aln <- data.table::fread(file_path)
cfgT <- data.table::fread(
  system.file("extdata", "results", "config_summary.csv",
              package = "amplican"))
findPD(aln, cfgT)
```

flipRanges	<i>Reverse complement events that have amplicons with direction 1.</i>
------------	--

Description

Reverse complement events that have amplicons with direction 1.

Usage

```
flipRanges(idR, cfgT)
```

Arguments

idR	(data.frame) Loaded events.
cfgT	(data.frame) Loaded configuration file.

Value

(data.frame) Returns input idR, but events for amplicons with direction 1 reverse complemented, "+" and "-" swapped.

fwdReads	<i>Alignments for forward reads.</i>
----------	--------------------------------------

Description

Get alignments for forward reads.

Usage

```
fwdReads(x)
```

Arguments

x	(AlignmentsExperimentSet)
---	---------------------------

Value

(listOrNULL) list with objects of PairwiseAlignmentsSingleSubject

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
fwdReads(aln)
```

fwdReads<-	<i>Alignments for forward reads.</i>
------------	--------------------------------------

Description

Set alignments for forward reads.

Usage

```
fwdReads(x) <- value
```

Arguments

x	(AlignmentsExperimentSet)
value	(list) Named (experiment IDs) list with elements of

Value

(AlignmentsExperimentSet) [PairwiseAlignmentsSingleSubject](#) class.

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
fwdReads(aln) <- fwdReads(aln) # replace with the same values
```

fwdReadsType	<i>Type of forward reads.</i>
--------------	-------------------------------

Description

Get type of forward reads.

Usage

```
fwdReadsType(x)
```

Arguments

x	(AlignmentsExperimentSet)
---	---------------------------

Value

(listOrNULL) list with objects of [PairwiseAlignmentsSingleSubject](#)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
fwdReadsType(aln)
```

fwdReadsType<-	<i>Read type for forward reads.</i>
----------------	-------------------------------------

Description

Set read type for forward reads.

Usage

```
fwdReadsType(x) <- value
```

Arguments

x	(AlignmentsExperimentSet)
value	(list) Named (experiment IDs) list with elements of

Value

(AlignmentsExperimentSet) [PairwiseAlignmentsSingleSubject](#) class.

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
fwdReadsType(aln) <- fwdReadsType(aln) # replace with the same values
```

geom_bezier	<i>Create quadratic or cubic bezier curves [copied from ggforce]</i>
-------------	--

Description

This set of functionality is copied from ggforce package due to dependency issues on Bioconductor and is used internally (not exported) only. This set of geoms makes it possible to connect points creating either quadratic or cubic beziers. bezier works by calculating points along the bezier and connecting these to draw the curve.

Usage

```

stat_bezier(
  mapping = NULL,
  data = NULL,
  geom = "path",
  position = "identity",
  na.rm = FALSE,
  show.legend = NA,
  n = 100,
  inherit.aes = TRUE,
  ...
)

geom_bezier(
  mapping = NULL,
  data = NULL,
  stat = "bezier",
  position = "identity",
  arrow = NULL,
  lineend = "butt",
  na.rm = FALSE,
  show.legend = NA,
  inherit.aes = TRUE,
  n = 100,
  ...
)

```

Arguments

- | | |
|---------|---|
| mapping | Set of aesthetic mappings created by aes() . If specified and <code>inherit.aes = TRUE</code> (the default), it is combined with the default mapping at the top level of the plot. You must supply mapping if there is no plot mapping. |
| data | <p>The data to be displayed in this layer. There are three options:</p> <p>If <code>NULL</code>, the default, the data is inherited from the plot data as specified in the call to ggplot().</p> <p>A <code>data.frame</code>, or other object, will override the plot data. All objects will be fortified to produce a data frame. See fortify() for which variables will be created.</p> <p>A function will be called with a single argument, the plot data. The return value must be a <code>data.frame</code>, and will be used as the layer data. A function can be created from a formula (e.g. <code>~ head(.x, 10)</code>).</p> |
| geom | <p>The geometric object to use to display the data for this layer. When using a <code>stat_*()</code> function to construct a layer, the <code>geom</code> argument can be used to override the default coupling between stats and geoms. The <code>geom</code> argument accepts the following:</p> <ul style="list-style-type: none"> • A <code>Geom</code> ggproto subclass, for example <code>GeomPoint</code>. |

	• A string naming the geom. To give the geom as a string, strip the function name of the <code>geom_</code> prefix. For example, to use <code>geom_point()</code>, give the geom as "point". • For more information and other ways to specify the geom, see the layer geom documentation.
<code>position</code>	A position adjustment to use on the data for this layer. This can be used in various ways, including to prevent overplotting and improving the display. The <code>position</code> argument accepts the following: • The result of calling a position function, such as <code>position_jitter()</code>. This method allows for passing extra arguments to the position. • A string naming the position adjustment. To give the position as a string, strip the function name of the <code>position_</code> prefix. For example, to use <code>position_jitter()</code>, give the position as "jitter". • For more information and other ways to specify the position, see the layer position documentation.
<code>na.rm</code>	If FALSE, the default, missing values are removed with a warning. If TRUE, missing values are silently removed.
<code>show.legend</code>	logical. Should this layer be included in the legends? NA, the default, includes if any aesthetics are mapped. FALSE never includes, and TRUE always includes. It can also be a named logical vector to finely select the aesthetics to display. To include legend keys for all levels, even when no data exists, use TRUE. If NA, all levels are shown in legend, but unobserved levels are omitted.
<code>n</code>	The number of points to create for each segment
<code>inherit.aes</code>	If FALSE, overrides the default aesthetics, rather than combining with them. This is most useful for helper functions that define both data and aesthetics and shouldn't inherit behaviour from the default plot specification, e.g. <code>annotation_borders()</code> .
<code>...</code>	Other arguments passed on to <code>layer()</code>'s <code>params</code> argument. These arguments broadly fall into one of 4 categories below. Notably, further arguments to the <code>position</code> argument, or aesthetics that are required can <i>not</i> be passed through <code>...</code>. Unknown arguments that are not part of the 4 categories below are ignored. • Static aesthetics that are not mapped to a scale, but are at a fixed value and apply to the layer as a whole. For example, <code>colour = "red"</code> or <code>linewidth = 3</code>. The geom's documentation has an Aesthetics section that lists the available options. The 'required' aesthetics cannot be passed on to the <code>params</code>. Please note that while passing unmapped aesthetics as vectors is technically possible, the order and required length is not guaranteed to be parallel to the input data. • When constructing a layer using a <code>stat_*()</code> function, the <code>...</code> argument can be used to pass on parameters to the geom part of the layer. An example of this is <code>stat_density(geom = "area", outline.type = "both")</code>. The geom's documentation lists which parameters it can accept. • Inversely, when constructing a layer using a <code>geom_*()</code> function, the <code>...</code> argument can be used to pass on parameters to the stat part of the layer. An example of this is <code>geom_area(stat = "density", adjust = 0.5)</code>. The stat's documentation lists which parameters it can accept.

	The <code>key_glyph</code> argument of <code>layer()</code> may also be passed on through <code>...</code>. This can be one of the functions described as key glyphs, to change the display of the layer in the legend.
<code>stat</code>	The statistical transformation to use on the data for this layer. When using a <code>geom_*()</code> function to construct a layer, the <code>stat</code> argument can be used to override the default coupling between geoms and stats. The <code>stat</code> argument accepts the following: - A Stat ggproto subclass, for example <code>StatCount</code>. - A string naming the stat. To give the stat as a string, strip the function name of the <code>stat_</code> prefix. For example, to use <code>stat_count()</code>, give the stat as <code>"count"</code>. - For more information and other ways to specify the stat, see the layer stat documentation.
<code>arrow</code>	Arrow specification, as created by <code>grid::arrow()</code> .
<code>lineend</code>	Line end style (round, butt, square).

Details

Input data is understood as a sequence of data points the first being the start point, then followed by one or two control points and then the end point. More than 4 and less than 3 points per group will throw an error.

Aesthetics

`geom_link`, `geom_link2` and `geom_lino` understand the following aesthetics (required aesthetics are in bold):

- **`x`** - **`y`** - `color` - `linewidth` - `linetype` - `alpha` - `lineend`

Computed variables

`x`, **`y`** The interpolated point coordinates

`index` The progression along the interpolation mapped between 0 and 1

Author(s)

Thomas Lin Pedersen

Examples

```
beziers <- data.frame(
  x = c(1, 2, 3, 4, 4, 6, 6),
  y = c(0, 2, 0, 0, 2, 2, 0),
  type = rep(c('cubic', 'quadratic'), c(3, 4)),
  point = c('end', 'control', 'end', 'end', 'control', 'control', 'end')
)
help_lines <- data.frame(
  x = c(1, 3, 4, 6),
  xend = c(2, 2, 4, 6),
```

```

      y = 0,
      yend = 2
    )
    ggplot2::ggplot() + ggplot2::geom_segment(
      ggplot2::aes(x = x, xend = xend, y = y, yend = yend),
      data = help_lines,
      arrow = ggplot2::arrow(length = ggplot2::unit(c(0, 0, 0.5, 0.5), 'cm')),
      colour = 'grey') +
    amplican::geom_bezier(ggplot2::aes(x= x, y = y, group = type, linetype = type),
      data = beziers) +
    ggplot2::geom_point(ggplot2::aes(x = x, y = y, colour = point), data = beziers)

```

getEventInfo	<i>This function takes alignments and gives back the events coordinates.</i>
--------------	--

Description

This function takes alignments and gives back the events coordinates.

Usage

```
getEventInfo(aligned, ID, amplicon_shift, strand_info = "+")
```

Arguments

aligned	(PairwiseAlignmentsSingleSubject)
ID	(string)
amplicon_shift	(numeric vector) Shift events additionally by this value. PairwiseAlignmentsSingleSubject returns truncated alignments.
strand_info	(string) Either '+', '-' or default '*'

Value

(GRanges) Object with meta-data for insertion, deletion, mismatch

getEvents	<i>Transform aligned strings into GRanges representation of events.</i>
-----------	---

Description

Transforms aligned strings into GRanges representation with events of deletions, insertions and mismatches. Subject should come from one amplicon sequence, after alignment to many sequences (patterns).

Usage

```
getEvents(
  pattern,
  subject,
  scores,
  ID = "NA",
  ampl_shift = 1L,
  ampl_start = 1L,
  strand_info = "+"
)
```

Arguments

pattern	(character) Aligned pattern.
subject	(character) Aligned subject.
scores	(integer) Alignment scores of the pattern and subject.
ID	(character) Will be used as seqnames of output GRanges.
ampl_shift	(numeric) Possible shift of the amplicons.
ampl_start	(numeric) Real amplicon starts. pairwiseAlignment clips alignments, therefore to output GRanges relative to the amplicon sequence (subject) ranges have to be shifted.
strand_info	(character) Strands to assign.

Value

([GRanges](#)) Same as events.

get_left_primer	<i>left primer sequence</i>
-----------------	-----------------------------

Description

left primer sequence

Usage

```
get_left_primer(config, id)
```

Arguments

config	(data.frame) config table
id	(vector) a vector of id's

Value

(character) left primer sequence

get_right_primer	<i>right primer sequence</i>
------------------	------------------------------

Description

right primer sequence

Usage

```
get_right_primer(config, id)
```

Arguments

config	(data.frame) config table
id	(vector) a vector of id's

Value

(character) right primer sequence

get_seq	<i>amplicon sequence, reverse complemented when needed</i>
---------	--

Description

amplicon sequence, reverse complemented when needed

Usage

```
get_seq(config, id, column = "Amplicon")
```

Arguments

config	(data.frame) config table
id	(vector) a vector of id's

Value

(character) amplicon sequence, reverse complemented if Direction 1

goodAvgQuality	<i>This filters out sequences which have bad average quality readings.</i>
----------------	--

Description

This filters out sequences which have bad average quality readings.

Usage

```
goodAvgQuality(reads, avg = 30, batch_size = 1e+07)
```

Arguments

reads	(ShortRead object) Loaded reads from fastq.
avg	(numeric) This is what the average score of the quality of sequence should be. For example, if we have a sequence with nucleotides which have quality 70-70-70, the average would be 70. If set the average to 70 or less the sequence will pass. If we set the average to 71 the sequence will not pass.
batch_size	(numeric) How many reads to process at a time.

Value

(boolean) Logical vector with the valid rows as TRUE.

goodBaseQuality	<i>Filters out sequences which have bad base quality readings.</i>
-----------------	--

Description

Filters out sequences which have bad base quality readings.

Usage

```
goodBaseQuality(reads, min = 20, batch_size = 1e+07)
```

Arguments

reads	(ShortRead object) Loaded reads from fastq.
min	(numeric) This is the minimum quality that we accept for every nucleotide. For example, if we have a sequence with nucleotides which have quality 50-50-50-50-10, and we set the minimum to 30, the whole sequence will be a bad sequence.
batch_size	(numeric) How many reads to process at a time.

Value

(boolean) Logical vector with the valid rows as TRUE.

is_hdr_strict	<i>Figure out which reads conform to the HDR using the donor.</i>
---------------	---

Description

This is strict detection as compared to 'is_hdr' which was designed to be less specific and allow for all kinds of donors. This method requires that you have exactly the same events (mismatches, insertions, deletions) as the difference between amplicon and donor sequences. It ignores everything else, so other mismatches and small indels etc. as noise are allowed here for valid HDR.

Usage

```
is_hdr_strict(aln, cfgT, scoring_matrix, gap_opening = 25, gap_extension = 0)
```

Arguments

aln	(data.table) This are events that contain already consensus column, they are also shifted and normalized.
cfgT	(data.table) Config data.table with columns for amplicon and donor.
scoring_matrix	(scoring matrix)
gap_opening	(integer)
gap_extension	(integer)

Value

(aln) same as aln on entry, but readType is updated to TRUE when read is recognized as HDR

locate_pr_start	<i>Find full or partial primer start positions with mismatches.</i>
-----------------	---

Description

This function finds a primer that can be partially truncated at the 5' end. It handles both full matches inside the read and partial matches at the beginning of the read.

Usage

```
locate_pr_start(reads, primer, m = 3, min_overlap = ceiling(nchar(primer)/2))
```

Arguments

reads	A character vector of DNA sequences.
primer	A single character string for the primer sequence.
m	The maximum number of allowed mismatches.
min_overlap	The minimum number of base pairs the primer must overlap with the read. A good value is often around half the primer length.

Value

A numeric vector of the EARLIEST start position for a valid match, with NA for no match.

lookupAlignment	<i>Show alignment in human readable format.</i>
-----------------	---

Description

Prints alignments in blast-like style for human examination.

Usage

```
lookupAlignment(x, ID, read_id = 1)
```

Arguments

x	(AlignmentsExperimentSet)
ID	(string) Experiment Identifier
read_id	(numeric) Read Identifier. Reads are sorted by frequency. Defaults to 1, most abundant read.

Value

(print to view)

Examples

```
# load example object
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
# look at most frequent reads aligned from experiment ID_1
lookupAlignment(aln, "ID_1")
```

makeAlignment	<i>Make alignments helper.</i>
---------------	--------------------------------

Description

Aligning reads to the amplicons for each ID in this barcode, constructing amplicanAlignment. Assume that all IDs here belong to the same barcode.

Usage

```
makeAlignment(
  cfgT,
  average_quality,
  min_quality,
  filter_n,
  batch_size,
  scoring_matrix,
  gap_opening,
  gap_extension,
  fastqfiles,
  primer_mismatch,
  donor_mismatch,
  donor_strict
)
```

Arguments

cfgT	config file as data table
average_quality	(numeric) The FASTQ file have a quality for each nucleotide, depending on sequencing technology there exist many formats. This package uses readFastq to parse the reads. If the average quality of the reads fall below value of average_quality then sequence is filtered. Default is 0.
min_quality	(numeric) Similar as in average_quality, but depicts the minimum quality for ALL nucleotides in given read. If one of nucleotides has quality BELOW min_quality, then the sequence is filtered. Default is 20.
filter_n	(boolean) Whether to filter out reads containing N base.
batch_size	(numeric) How many reads to analyze at a time? Needed for filtering of large fastq files.
scoring_matrix	(matrix) Default is 'NUC44'. Pass desired matrix using nucleotideSubstitutionMatrix .
gap_opening	(numeric) The opening gap score.
gap_extension	(numeric) The gap extension score.
fastqfiles	(numeric) Normally you want to use both FASTQ files. But in some special cases, you may want to use only the forward file, or only the reverse file. Possible options:

- 0** Use both FASTQ files.
- 0.5** Use both FASTQ files, but only for one of the reads (forward or reverse) is required to have primer perfectly matched to sequence - eg. use when reverse reads are trimmed of primers, but forward reads have forward primer in the sequence.
- 1** Use only the forward FASTQ file.
- 2** Use only the reverse FASTQ file.

primer_mismatch

(numeric) Decide how many mismatches are allowed during primer matching of the reads, that groups reads by experiments. When primer_mismatch = 0 no mismatches are allowed, which can increase number of unassigned read.

donor_mismatch

(numeric) How many events of length 1 (mismatches, deletions and insertions of length 1) are allowed when aligning toward the donor template. This parameter is only used when donor template is specified. The higher the parameter the less strict will be algorithm accepting read as HDR. Set to 0 if only perfect alignments to the donor template marked as HDR, unadvised due to error rate of the sequencers.

donor_strict

(logical) Applies more strict algorithm for HDR detection. Only these reads that have all of the donor events will count as HDR. Tolerates 'donor_mismatch' level of noise, but no indels are allowed. Use this when your reads should span over the whole window of the donor events. Might be more time consuming.

Value

amplicanAlignment object for this barcode experiments

metaplot_deletions *MetaPlots deletions using ggplot2.*

Description

This function plots deletions in relation to the amplicons for given selection vector that groups values by given config group. All reads should already be converted to their relative position to their respective amplicon using [amplicanMap](#). Top plot is for the forward reads and bottom plot is for reverse reads.

Usage

```
metaplot_deletions(alnmt, config, group, selection, over = "overlaps")
```

Arguments

alnmt (data.frame) Loaded alignment information from events_filtered_shifted_normalized.csv file.

config (data.frame) Loaded table from config_summary.csv file.

group	(string) Name of the column from the config file to use for grouping. Events are subselected based on this column and values from selection.
selection	(string or vector of strings) Values from config column specified in group argument.
over	(string) Specify which column contains overlaps with expected cut sites generated by amplicanOverlap

Value

(deletions metaplot) ggplot2 object of deletions metaplot

See Also

Other specialized plots: [metaplot_insertions\(\)](#), [metaplot_mismatches\(\)](#), [plot_cuts\(\)](#), [plot_deletions\(\)](#), [plot_heterogeneity\(\)](#), [plot_insertions\(\)](#), [plot_mismatches\(\)](#), [plot_variants\(\)](#)

Examples

```
#example config
config <- read.csv(system.file("extdata", "results", "config_summary.csv",
                             package = "amplican"))

#example alignments results
alignments_file <- system.file("extdata", "results", "alignments",
                              "events_filtered_shifted_normalized.csv",
                              package = "amplican")
alignments <- read.csv(alignments_file)
metaplot_deletions(alignments[alignments$consensus, ],
                  config, "Group", "Betty")
```

metaplot_insertions *MetaPlots insertions using ggplot2.*

Description

This function plots insertions in relation to the amplicons for given selection vector that groups values by given config group. All reads should already be converted to their relative position to their respective amplicon using [amplicanMap](#). Top plot is for the forward reads and bottom plot is for reverse reads.

Usage

```
metaplot_insertions(alnmt, config, group, selection)
```

Arguments

alnmt	(data.frame) Loaded alignment information from alignments_events.csv file.
config	(data.frame) Loaded table from config_summary.csv file.
group	(string) Name of the column from the config file to use for grouping. Events are subselected based on this column and values from selection.
selection	(string or vector of strings) Values from config column specified in group argument.

Value

(insertions metaplot) ggplot2 object of insertions metaplot

See Also

Other specialized plots: [metaplot_deletions\(\)](#), [metaplot_mismatches\(\)](#), [plot_cuts\(\)](#), [plot_deletions\(\)](#), [plot_heterogeneity\(\)](#), [plot_insertions\(\)](#), [plot_mismatches\(\)](#), [plot_variants\(\)](#)

Examples

```
#example config
config <- read.csv(system.file("extdata", "results", "config_summary.csv",
                             package = "amplican"))

#example alignments results
alignments_file <- system.file("extdata", "results", "alignments",
                              "events_filtered_shifted_normalized.csv",
                              package = "amplican")

alignments <- read.csv(alignments_file)
metaplot_insertions(alignments$consensus, [], config,
                   "Group", "Betty")
```

metaplot_mismatches *MetaPlots mismatches using ggplot2.*

Description

Plots mismatches in relation to the amplicons for given selection vector that groups values by given config group. All reads should already be converted to their relative position to their respective amplicon using [amplicanMap](#). Zero position on new coordinates is the most left UPPER case letter of the respective amplicon. This function filters out all alignment events that have amplicons without UPPER case defined. Top plot is for the forward reads and bottom plot is for reverse reads.

Usage

```
metaplot_mismatches(alnmt, config, group, selection)
```

Arguments

alnmt	(data.frame) Loaded alignment information from alignments_events.csv file.
config	(data.frame) Loaded table from config_summary.csv file.
group	(string) Name of the column from the config file to use for grouping. Events are subselected based on this column and values from selection.
selection	(string or vector of strings) Values from config column specified in group argument.

Value

(mismatches metaplot) ggplot2 object of mismatches metaplot

See Also

Other specialized plots: [metaplot_deletions\(\)](#), [metaplot_insertions\(\)](#), [plot_cuts\(\)](#), [plot_deletions\(\)](#), [plot_heterogeneity\(\)](#), [plot_insertions\(\)](#), [plot_mismatches\(\)](#), [plot_variants\(\)](#)

Examples

```
#example config
config <- read.csv(system.file("extdata", "results", "config_summary.csv",
                             package = "amplican"))

#example alignments results
alignments_file <- system.file("extdata", "results", "alignments",
                              "events_filtered_shifted_normalized.csv",
                              package = "amplican")

alignments <- read.csv(alignments_file)
metaplot_mismatches(alignments,
                    config, "Group", "Betty")
```

pairToEvents

Read "pair" format of EMBOSS needle into GRanges as events.

Description

Parse EMBOSS needle (or needleall) "pair" format into GRanges representation with events of deletions, insertions and mismatches. Make sure that each file corresponds to single subject (single amplicon). Assumes that bottom sequence "-bsequence" corresponds to the "subject" and full sequence alignment is returned.

Usage

```
pairToEvents(file, ID = "NA", strand_info = "+")
```

Arguments

file (character) File path.
ID (character) ID of the experiment, will be used as seqnames of the reutner ranges.
strand_info (character) Strand to assign.

Value

([GRanges](#)) Same as events.

plot_amplicon	<i>Plots amplicon sequence using ggplot2.</i>
---------------	---

Description

Plots amplicon sequence using ggplot2.

Usage

```
plot_amplicon(amplicon, from, to)
```

Arguments

amplicon (character) Sequence of the amplicon to plot.
from (number) Minimum on x axis - start of the amplicon
to (number) Maximum on x axis - not necessarily end of the amplicon

Value

(amplicon plot) ggplot2 object of amplicon plot

plot_cuts	<i>Plots cuts using ggplot2.</i>
-----------	----------------------------------

Description

This function plots cuts in relation to the amplicon with distinction for each ID.

Usage

```
plot_cuts(alignments, config, id, cut_buffer = 5, xlab_spacing = 4)
```

Arguments

alignments	(data.frame) Loaded alignment information from alignments_events.csv file.
config	(data.frame) Loaded table from config_summary.csv file.
id	(string or vector of strings) Name of the ID column from config file or name of multiple IDs if it is possible to group them. First amplicon will be used as the basis for plot.
cut_buffer	(numeric) Default is 5, you should specify the same as used in the analysis.
xlab_spacing	(numeric) Spacing of the x axis labels. Default is 4.

Value

(cuts plot) gtable object of cuts plot

See Also

Other specialized plots: [metaplot_deletions\(\)](#), [metaplot_insertions\(\)](#), [metaplot_mismatches\(\)](#), [plot_deletions\(\)](#), [plot_heterogeneity\(\)](#), [plot_insertions\(\)](#), [plot_mismatches\(\)](#), [plot_variants\(\)](#)

Examples

```
#example config
config <- read.csv(system.file("extdata", "results", "config_summary.csv",
                             package = "amplican"))

#example alignments results
alignments_file <- system.file("extdata", "results", "alignments",
                              "events_filtered_shifted_normalized.csv",
                              package = "amplican")
alignments <- read.csv(alignments_file)
plot_cuts(alignments[alignments$consensus & alignments$overlaps, ],
          config, c('ID_1','ID_3'))
```

plot_deletions	<i>Plots deletions using ggplot2.</i>
----------------	---------------------------------------

Description

This function plots deletions in relation to the amplicon, assumes events are relative to the expected cut site. Top plot is for the forward reads, middle one shows amplicon sequence, and bottom plot is for reverse reads.

Usage

```
plot_deletions(
  alignments,
  config,
  id,
  cut_buffer = 5,
  xlab_spacing = 4,
  over = "overlaps"
)
```

Arguments

<code>alignments</code>	(data.frame) Loaded alignment information from alignments.csv file.
<code>config</code>	(data.frame) Loaded table from config_summary.csv file.
<code>id</code>	(string or vector of strings) Name of the ID column from config file or name of multiple IDs if it is possible to group them. First amplicon will be used as the basis for plot.
<code>cut_buffer</code>	(numeric) Default is 5, you should specify the same as used in the analysis.
<code>xlab_spacing</code>	(numeric) Spacing of the x axis labels. Default is 4.
<code>over</code>	(string) Specify which columns contains overlaps with expected cut sites generated by amplicanOverlap

Value

(deletions plot) gtable object of deletions plot

See Also

Other specialized plots: [metaplot_deletions\(\)](#), [metaplot_insertions\(\)](#), [metaplot_mismatches\(\)](#), [plot_cuts\(\)](#), [plot_heterogeneity\(\)](#), [plot_insertions\(\)](#), [plot_mismatches\(\)](#), [plot_variants\(\)](#)

Examples

```
#example config
config <- read.csv(system.file("extdata", "results", "config_summary.csv",
                              package = "amplican"))

#example alignments results
alignments_file <- system.file("extdata", "results", "alignments",
                              "events_filtered_shifted_normalized.csv",
                              package = "amplican")

alignments <- read.csv(alignments_file)
p <- plot_deletions(alignments[alignments$consensus, ],
                    config, c('ID_1', 'ID_3'))
```

plot_height	<i>Get figure height in inches for number of elements on y axis.</i>
-------------	--

Description

Helper function to calculate figure height based on number of elements to plot for automating sizes of figures in knitted reports.

Usage

```
plot_height(x)
```

Arguments

x (numeric) number of elements to fit onto height axis

Value

(numeric) In inches

Examples

```
plot_height(20)
```

plot_heterogeneity	<i>Plots heterogeneity of the reads using ggplot2.</i>
--------------------	--

Description

This function creates stacked barplot explaining reads heterogeneity. It groups reads by user defined levels and measures how unique are reads in this level. Uniqueness of reads is simplified to the bins and colored according to the color gradient. Default color black indicates very high heterogeneity of the reads. The more yellow (default) the more similar are reads and less heterogeneous.

Usage

```
plot_heterogeneity(
  alignments,
  config,
  level = "ID",
  colors = c("#000000", "#F0E442"),
  bins = c(0, 5, seq(10, 100, 10))
)
```

Arguments

alignments	(data.frame) Loaded alignment information from alignments_events.csv file.
config	(data.frame) Loaded table from config_summary.csv file.
level	(string) Name of the column from config file specifying levels to group by.
colors	(html colors vector) Two colours for gradient, eg. c('#000000', '#F0E442').
bins	(numeric vector) Numeric vector from 0 to 100 specifying bins eg. c(0, 5, seq(10, 100, 10)).

Value

(heterogeneity plot) ggplot2 object of heterogeneity plot

See Also

Other specialized plots: [metaplot_deletions\(\)](#), [metaplot_insertions\(\)](#), [metaplot_mismatches\(\)](#), [plot_cuts\(\)](#), [plot_deletions\(\)](#), [plot_insertions\(\)](#), [plot_mismatches\(\)](#), [plot_variants\(\)](#)

Examples

```
#example config
config <- read.csv(system.file("extdata", "results", "config_summary.csv",
                             package = "amplican"))

#example alignments results
alignments_file <- system.file("extdata", "results", "alignments",
                              "events_filtered_shifted_normalized.csv",
                              package = "amplican")

alignments <- read.csv(alignments_file)
plot_heterogeneity(alignments[alignments$consensus, ], config)
```

plot_insertions	<i>Plots insertions using ggplot2.</i>
-----------------	--

Description

This function plots insertions in relation to the amplicon. Top plot is for the forward reads, middle one shows amplicon sequence, and bottom plot is for reverse reads.

Usage

```
plot_insertions(alignments, config, id, cut_buffer = 5, xlab_spacing = 4)
```

Arguments

alignments	(data.frame) Loaded alignment information from alignments_events.csv file.
config	(data.frame) Loaded table from config_summary.csv file.
id	(string or vector of strings) Name of the ID column from config file or name of multiple IDs if it is possible to group them. First amplicon will be used as the basis for plot.
cut_buffer	(numeric) Default is 5, you should specify the same as used in the analysis.
xlab_spacing	(numeric) Spacing of the x axis labels. Default is 4.

Value

(insertions plot) gtable object of insertions plot

See Also

Other specialized plots: [metaplot_deletions\(\)](#), [metaplot_insertions\(\)](#), [metaplot_mismatches\(\)](#), [plot_cuts\(\)](#), [plot_deletions\(\)](#), [plot_heterogeneity\(\)](#), [plot_mismatches\(\)](#), [plot_variants\(\)](#)

Examples

```
#example config
config <- read.csv(system.file("extdata", "results", "config_summary.csv",
                             package = "amplican"))

#example alignments results
alignments_file <- system.file("extdata", "results", "alignments",
                              "events_filtered_shifted_normalized.csv",
                              package = "amplican")

alignments <- read.csv(alignments_file)
p <- plot_insertions(alignments, config, c('ID_1','ID_3'))
```

plot_mismatches	<i>Plots mismatches using ggplot2.</i>
-----------------	--

Description

Plots mismatches in relation to the amplicon, assumes your reads are relative to the respective amplicon sequences predicted cut sites. Top plot is for the forward reads, middle one shows amplicon sequence, and bottom plot is for reverse reads.

Usage

```
plot_mismatches(alignments, config, id, cut_buffer = 5, xlab_spacing = 4)
```

Arguments

alignments	(data.frame) Loaded alignment information from alignments_events.csv file.
config	(data.frame) Loaded table from config_summary.csv file.
id	(string or vector of strings) Name of the ID column from config file or name of multiple IDs, if it is possible to group them. They have to have the same amplicon, amplicons on the reverse strand will be reverse complemented to match forward strand amplicons.
cut_buffer	(numeric) Default is 5, you should specify the same as used in the analysis.
xlab_spacing	(numeric) Spacing of the x axis labels. Default is 4.

Value

(mismatches plot) gtable object of mismatches plot

See Also

Other specialized plots: [metaplot_deletions\(\)](#), [metaplot_insertions\(\)](#), [metaplot_mismatches\(\)](#), [plot_cuts\(\)](#), [plot_deletions\(\)](#), [plot_heterogeneity\(\)](#), [plot_insertions\(\)](#), [plot_variants\(\)](#)

Examples

```
#example config
config <- read.csv(system.file("extdata", "results", "config_summary.csv",
                             package = "amplican"))

#example alignments results
alignments_file <- system.file("extdata", "results", "alignments",
                              "events_filtered_shifted_normalized.csv",
                              package = "amplican")

alignments <- read.csv(alignments_file)
id <- c('ID_1', 'ID_3'); cut_buffer = 5; xlab_spacing = 4;
p <- plot_mismatches(alignments, config, c('ID_1', 'ID_3'))
ggplot2::ggsave("~/test.png", p, width = 25, units = "in")
```

plot_variants

Plots most frequent variants using ggplot2.

Description

This function plots variants in relation to the amplicon. Shows sequences of top mutants without aggregating on deletions, insertions and mismatches.

Usage

```
plot_variants(
  alignments,
  config,
  id,
  cut_buffer = 5,
  top = 10,
  annot = if (amplican::get_seq(config, id, "Donor") == "") "cov" else NA,
  summary_plot = amplican::get_seq(config, id, "Donor") == "",
  frameshift = amplican::get_seq(config, id, "Donor") == ""
)
```

Arguments

alignments	(data.frame) Loaded alignment information from alignments_events.csv file.
config	(data.frame) Loaded table from config_summary.csv file.
id	(string or vector of strings) Name of the ID column from config file or name of multiple IDs if it is possible to group them. First amplicon will be used as the basis for plot. If Donor is available we will try to add the first donor and mark it on the plot.
cut_buffer	(numeric) Default is 5, you should specify the same as used in the analysis.
top	(numeric) Specify number of most frequent reads to plot. By default it is 10. Check plot_heterogeneity to see how many reads will be enough to give good overview of your variants.
annot	("codon" or "cov" or NA) What to display for annotation top plot. When NA will not display anything, also not display total summary. Codon plot is all reading frames for a given window, and the default "cov" is coverage of all indels and mismatches over a given window.
summary_plot	(boolean) Whether small summary plot in the upper right corner should be displayed. Top bar summarizes total reads with frameshift (F), reads with Edits without Frameshift (Edits) and reads without Edits (Match).
frameshift	(boolean) Whether to include Frameshift column in the table.
annot on off	

Details

Top plot shows all six possible frames for given amplicon. Amino acids are colored as follows:

Small nonpolar	G, A, S, T	Orange
Hydrophobic	C, V, I, L, P, F, Y, M, W	Green
Polar	N, Q, H	Magenta
Negatively charged	D, E	Red
Positively charged	K, R	Blue
Other	eg. *, U, +	Grey

Variant plot shows amplicon reference, UPPER letters which were the basis for window selection are highlighted with dashed white box (guideRNA). Black triangles are reflecting insertion points. Dashed letters indicate deletions. Table associated with variant plot represents:

Freq - Frequency of given read in experiment. Variants are ordered by frequency value.

Count - Represents raw count of this variant reads in experiment.

F - Sum of deletion and insertion widths of events overlapping presented window. Green background indicates frameshift.

Value

(variant plot) gtable object of variants plot

See Also

Other specialized plots: [metaplot_deletions\(\)](#), [metaplot_insertions\(\)](#), [metaplot_mismatches\(\)](#), [plot_cuts\(\)](#), [plot_deletions\(\)](#), [plot_heterogeneity\(\)](#), [plot_insertions\(\)](#), [plot_mismatches\(\)](#)

Examples

```
#example config
config <- read.csv(system.file("extdata", "results", "config_summary.csv",
                             package = "amplican"))

#example alignments results
alignments_file <- system.file("extdata", "results", "alignments",
                              "events_filtered_shifted_normalized.csv",
                              package = "amplican")
alignments <- read.csv(alignments_file)

alignments <- alignments[alignments$consensus & alignments$overlaps, ]
p <- plot_variants(alignments[alignments$consensus & alignments$overlaps, ],
                  config, c('ID_1','ID_3'))
# with Donor we dont plot summary and the annot, summary plot and frameshift
p <- plot_variants(alignments[alignments$consensus & alignments$overlaps, ],
                  config, c('ID_5'))
```

readCounts

Alignments for forward reads.

Description

Set alignments for forward reads.

Usage

```
readCounts(x)
```

Arguments

x (AlignmentsExperimentSet)

Value

(listOrNULL)

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
readCounts(aln)
```

readCounts<- *Alignments for forward reads.*

Description

Set alignments for forward reads.

Usage

```
readCounts(x) <- value
```

Arguments

x (AlignmentsExperimentSet)

value (list) Named (experiment IDs) list with elements of

Value

(AlignmentsExperimentSet) [PairwiseAlignmentsSingleSubject](#) class.

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
readCounts(aln) <- readCounts(aln) # replace with the same values
```

revComp	<i>Reverse and complement given string or list of strings</i>
---------	---

Description

Reverse and complement given string or list of strings

Usage

```
revComp(x)
```

Arguments

x (string or vector of strings)

Value

(string or vector of strings) reverse complemented input

rveReads	<i>Alignments for reverse reads.</i>
----------	--------------------------------------

Description

Get alignments for reverse reads.

Usage

```
rveReads(x)
```

Arguments

x (AlignmentsExperimentSet)

Value

(listOrNULL) list with objects of PairwiseAlignmentsSingleSubject

Examples

```
file_path <- system.file("extdata", "results", "alignments",  
                          "AlignmentsExperimentSet.rds", package = "amplican")  
aln <- readRDS(file_path)  
rveReads(aln)
```

rveReads<-	<i>Alignments for forward reads.</i>
------------	--------------------------------------

Description

Set alignments for forward reads.

Usage

```
rveReads(x) <- value
```

Arguments

x	(AlignmentsExperimentSet)
value	(list) Named (experiment IDs) list with elements of

Value

(AlignmentsExperimentSet) [PairwiseAlignmentsSingleSubject](#) class.

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
rveReads(aln) <- rveReads(aln) # replace with the same values
```

rveReadsType	<i>Type of reverse reads.</i>
--------------	-------------------------------

Description

Get type of reverse reads.

Usage

```
rveReadsType(x)
```

Arguments

x	(AlignmentsExperimentSet)
---	---------------------------

Value

(listOrNULL) list with objects of [PairwiseAlignmentsSingleSubject](#)

`rveReadsType<-`

67

Examples

```
file_path <- system.file("extdata", "results", "alignments",  
                          "AlignmentsExperimentSet.rds", package = "amplican")  
aln <- readRDS(file_path)  
rveReadsType(aln)
```

<code>rveReadsType<-</code>	<i>Read type for reverse reads.</i>
--------------------------------	-------------------------------------

Description

Set read type for reverse reads.

Usage

```
rveReadsType(x) <- value
```

Arguments

<code>x</code>	(AlignmentsExperimentSet)
<code>value</code>	(list) Named (experiment IDs) list with elements of

Value

(AlignmentsExperimentSet) [PairwiseAlignmentsSingleSubject](#) class.

Examples

```
file_path <- system.file("extdata", "results", "alignments",  
                          "AlignmentsExperimentSet.rds", package = "amplican")  
aln <- readRDS(file_path)  
rveReadsType(aln) <- rveReadsType(aln) # replace with the same values
```

<code>unassignedCount</code>	<i>Get count of unassigned reads.</i>
------------------------------	---------------------------------------

Description

Get count of unassigned reads.

Usage

```
unassignedCount(x)
```

Arguments

<code>x</code>	(AlignmentsExperimentSet)
----------------	---------------------------

Value

(numeric)

Examples

```
file_path <- system.file("extdata", "results", "alignments",  
                          "AlignmentsExperimentSet.rds", package = "amplican")  
aln <- readRDS(file_path)  
unassignedCount(aln)
```

unassignedData

Unassigned read information.

Description

Get unassigned reads and their characteristics.

Usage

```
unassignedData(x)
```

Arguments

x (AlignmentsExperimentSet)

Value

(data.frameOrNULL)

Examples

```
file_path <- system.file("extdata", "results", "alignments",  
                          "AlignmentsExperimentSet.rds", package = "amplican")  
aln <- readRDS(file_path)  
unassignedData(aln)
```

unassignedData<-	<i>Alignments for forward reads.</i>
------------------	--------------------------------------

Description

Set alignments for forward reads.

Usage

```
unassignedData(x) <- value
```

Arguments

x	(AlignmentsExperimentSet)
value	(list) Named (experiment IDs) list with elements of

Value

(AlignmentsExperimentSet) [PairwiseAlignmentsSingleSubject](#) class.

Examples

```
file_path <- system.file("extdata", "results", "alignments",
                          "AlignmentsExperimentSet.rds", package = "amplican")
aln <- readRDS(file_path)
unassignedData(aln) <- unassignedData(aln) #replace with the same values
```

upperGroups	<i>Detect uppercases as ranges object.</i>
-------------	--

Description

For a given string, detect how many groups of uppercases is inside, where are they, and how long they are.

Usage

```
upperGroups(candidate)
```

Arguments

candidate	(string) A string with the nucleotide sequence.
-----------	---

Details

For example: asdkfaAGASDGAsjaeuradAFDSfsafjjaeiorAuaoeurasjfasdhfashTTSfajeiasjsf
Has 4 groups of uppercases of length 7, 4, 1 and 3.

Value

(IRanges) A IRanges object with uppercases groups for given candidate string

waffle	Create a Waffle Chart using ggplot2
--------	-------------------------------------

Description

This function generates a waffle chart, which is a grid of squares showing parts of a whole. It's a simple replacement for the deprecated 'waffle' package for basic use cases.

Usage

```
waffle(parts, rows = 10, title = NULL, colors = NULL, legend_pos = "bottom")
```

Arguments

parts	A named numeric vector where names are the categories and values are the number of squares to be plotted for each category.
rows	The number of rows in the waffle grid. The number of columns will be calculated based on the total number of squares.
title	A string for the plot title.
colors	A character vector of colors to use for the different categories. The number of colors should match the number of categories in 'parts'.
legend_pos	The position of the legend (e.g., "bottom", "right", "none").

Value

A ggplot object representing the waffle chart.

Examples

```
read_q_per <- c("Good Reads\n100 (80%)" = 80, "Bad Reads\n25 (20%)" = 20)
waffle(
  parts = read_q_per,
  rows = 10,
  title = "Quality of all reads",
  colors = c('#E69F00', '#000000')
)
```

writeAlignments	<i>Write alignments to file.</i>
-----------------	----------------------------------

Description

Saves alignments into txt or fasta file.

Usage

```
writeAlignments(x, file = "", aln_format = "txt")
```

Arguments

x	(AlignmentsExperimentSet)
file	(connection or string) Destination file. When empty, defaults to standard output.
aln_format	("txt" or "fasta") Specifies format of the file.

Value

(invisible)

Examples

```
file_path <- system.file("extdata", "results", "alignments",  
                          "AlignmentsExperimentSet.rds", package = "amplican")  
aln <- readRDS(file_path)  
writeAlignments(aln, file.path(tempdir(), "aln.txt"))
```

Index

* analysis steps

- [amplicanAlign](#), 8
- [amplicanConsensus](#), 11
- [amplicanFilter](#), 12
- [amplicanMap](#), 13
- [amplicanNormalize](#), 14
- [amplicanOverlap](#), 15
- [amplicanPipeline](#), 16
- [amplicanPipelineConservative](#), 19
- [amplicanReport](#), 23
- [amplicanSummarize](#), 24

* filters

- [findEOP](#), 35
- [findLQR](#), 36
- [findPD](#), 37

* internal

- [AlignmentsExperimentSet-class](#), 4
- [alphabetQuality](#), 7
- [assignedCount](#), 26
- [barcodeData](#), 27
- [barcodeData<-](#), 27
- [checkConfigFile](#), 28
- [checkFileWriteAccess](#), 28
- [checkPrimers](#), 29
- [checkTarget](#), 29
- [cumsumw](#), 31
- [decode](#), 32
- [defGR](#), 32
- [experimentData](#), 33
- [experimentData<-](#), 34
- [extractEvents](#), 34
- [flipRanges](#), 38
- [fwdReads](#), 38
- [fwdReads<-](#), 39
- [fwdReadsType](#), 39
- [fwdReadsType<-](#), 40
- [get_left_primer](#), 45
- [get_right_primer](#), 46
- [get_seq](#), 46

- [getEventInfo](#), 44
- [goodAvgQuality](#), 47
- [goodBaseQuality](#), 47
- [lookupAlignment](#), 49
- [makeAlignment](#), 50
- [plot_amplicon](#), 55
- [readCounts](#), 63
- [readCounts<-](#), 64
- [revComp](#), 65
- [rveReads](#), 65
- [rveReads<-](#), 66
- [rveReadsType](#), 66
- [rveReadsType<-](#), 67
- [unassignedCount](#), 67
- [unassignedData](#), 68
- [unassignedData<-](#), 69
- [upperGroups](#), 69
- [writeAlignments](#), 71

* specialized plots

- [metaplot_deletions](#), 51
- [metaplot_insertions](#), 52
- [metaplot_mismatches](#), 53
- [plot_cuts](#), 55
- [plot_deletions](#), 56
- [plot_heterogeneity](#), 58
- [plot_insertions](#), 59
- [plot_mismatches](#), 60
- [plot_variants](#), 61

- [\[,AlignmentsExperimentSet,numeric,missing,missing-method](#)
(AlignmentsExperimentSet-class),
4

- [\\$,AlignmentsExperimentSet-method](#)
(AlignmentsExperimentSet-class),
4

- [aes\(\)](#), 41

- [AlignmentsExperimentSet](#), 8, 10, 17, 20

- [AlignmentsExperimentSet](#)
(AlignmentsExperimentSet-class),
4

- AlignmentsExperimentSet-class, 4
- alphabetQuality, 7
- amplican, 8
- amplican-package (amplican), 8
- amplican_print_reads, 25
- amplicanAlign, 8, 11–13, 15, 16, 19, 22–25
- amplicanConsensus, 10, 11, 12, 13, 15, 16, 18, 19, 22, 24, 25
- amplicanFilter, 10, 11, 12, 13, 15, 16, 19, 22, 24, 25
- amplicanMap, 10–12, 13, 15, 16, 19, 22, 24, 25, 51–53
- amplicanNormalize, 10–13, 14, 16, 19, 22, 24, 25
- amplicanOverlap, 10–13, 15, 15, 19, 22, 24, 25, 52, 57
- amplicanPipeline, 10–13, 15, 16, 16, 22, 24, 25
- amplicanPipelineConservative, 10–13, 15, 16, 19, 19, 24, 25
- amplicanReport, 10–13, 15, 16, 19, 22, 23, 25
- amplicanSummarize, 10–13, 15, 16, 19, 22, 24, 24
- annotation_borders(), 42
- as.list.AlignmentsExperimentSet
(AlignmentsExperimentSet-class), 4
- assignedCount, 26
- assignedCount, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- barcodeData, 27
- barcodeData, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- barcodeData<-, 27
- barcodeData<-, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- c, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- checkConfigFile, 28
- checkFileWriteAccess, 28
- checkPrimers, 29
- checkTarget, 29
- cigarsToEvents, 30
- comb_along, 31
- cumsumw, 31
- decode, 32
- defGR, 32
- experimentData, 33
- experimentData, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- experimentData<-, 34
- experimentData<-, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- extractEvents, 34
- extractEvents, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- findEOP, 12, 35, 36, 37
- findLQR, 35, 36, 37
- findPD, 12, 35, 36, 37
- flipRanges, 38
- fortify(), 41
- fwdReads, 38
- fwdReads, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- fwdReads<-, 39
- fwdReads<-, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- fwdReadsType, 39
- fwdReadsType, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- fwdReadsType<-, 40
- fwdReadsType<-, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class), 4
- geom_bezier, 40
- get_left_primer, 45
- get_right_primer, 46
- get_seq, 46
- getEventInfo, 44
- getEvents, 44
- ggplot(), 41
- goodAvgQuality, 47

- goodBaseQuality, 47
- GRanges, 7, 13, 30, 33, 35, 37, 44, 45, 55
- grid::arrow(), 43
- IRanges, 31, 33, 70
- is_hdr_strict, 48
- key glyphs, 43
- layer geom, 42
- layer position, 42
- layer stat, 43
- layer(), 42, 43
- length, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- locate_pr_start, 48
- lookupAlignment, 10, 49
- lookupAlignment, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- makeAlignment, 50
- metaplot_deletions, 51, 53, 54, 56, 57,
59–61, 63
- metaplot_insertions, 52, 52, 54, 56, 57,
59–61, 63
- metaplot_mismatches, 52, 53, 53, 56, 57,
59–61, 63
- names, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- nucleotideSubstitutionMatrix, 9, 18, 21,
50
- pairToEvents, 54
- pairwiseAlignment, 45
- PairwiseAlignmentsSingleSubject, 6, 39,
40, 44, 64, 66, 67, 69
- plot_amplicon, 55
- plot_cuts, 52–54, 55, 57, 59–61, 63
- plot_deletions, 52–54, 56, 56, 59–61, 63
- plot_height, 58
- plot_heterogeneity, 52–54, 56, 57, 58,
60–63
- plot_insertions, 52–54, 56, 57, 59, 59, 61,
63
- plot_mismatches, 52–54, 56, 57, 59, 60, 60,
63
- plot_variants, 52–54, 56, 57, 59–61, 61
- readCounts, 63
- readCounts, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- readCounts<-, 64
- readCounts<-, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- readFastq, 9, 17, 21, 50
- register, 6, 34
- revComp, 65
- rveReads, 65
- rveReads, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- rveReads<-, 66
- rveReads<-, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- rveReadsType, 66
- rveReadsType, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- rveReadsType<-, 67
- rveReadsType<-, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- stat_bezier (geom_bezier), 40
- unassignedCount, 67
- unassignedCount, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- unassignedData, 68
- unassignedData, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- unassignedData<-, 69
- unassignedData<-, AlignmentsExperimentSet-method
(AlignmentsExperimentSet-class),
4
- upperGroups, 69
- waffle, 70
- writeAlignments, 71

writeAlignments, AlignmentsExperimentSet-method
 (AlignmentsExperimentSet-class),
 4