

Copy number calling and SNV classification using targeted short read sequencing

Markus Riester¹

¹Novartis Biomedical Research, Cambridge, MA

October 7, 2025

Abstract

PureCN [1] is a purity and ploidy aware copy number caller for cancer samples inspired by the *ABSOLUTE* algorithm [2]. It was designed for hybrid capture sequencing data, especially with medium-sized targeted gene panels without matching normal samples in mind (matched whole-exome data is of course supported).

It can be used to supplement existing normalization and segmentation algorithms, i.e. the software can start from BAM files, from target-level coverage data, from copy number log2-ratios or from already segmented data. After the correct purity and ploidy solution was identified, *PureCN* will accurately classify variants as germline vs. somatic or clonal vs. sub-clonal.

PureCN was further designed to integrate well with industry standard pipelines [3], but it is straightforward to generate input data from other pipelines.

Package

PureCN 2.15.4

Contents

1	Introduction	3
2	Basic input files	3
2.1	VCF	3
2.2	Target information	4
2.3	Coverage data	5
2.4	Third-party coverage tools	5
2.5	Third-party segmentation tools	6
2.6	Example data	6
3	Library-specific coverage bias	6
4	Pool of normals	8
4.1	Coverage normalization.	8

Copy number calling and SNV classification using targeted short read sequencing

4.2	Artifact filtering	8
4.2.1	VCF	9
4.3	Artifact filtering without a pool of normals	9
5	Recommended run	10
6	Output	11
6.1	Plots	11
6.2	Data structures	17
6.2.1	Prediction of somatic status and cellular fraction	17
6.2.2	Amplifications and deletions	18
6.2.3	Amplifications in low purity samples	20
6.2.4	Find genomic regions in LOH	22
7	Curation	23
8	Cell lines	24
9	Maximizing the number of heterozygous SNPs	25
10	Advanced usage	25
10.1	Custom normalization and segmentation	25
10.1.1	Custom segmentation	25
10.1.2	Custom normalization	26
10.2	COSMIC annotation	27
10.3	ExAC and gnomAD annotation	28
10.4	Mutation burden	28
10.5	Chromosomal Instability	28
10.6	Detect cross-sample contamination	29
10.7	Power to detect somatic mutations	30
11	Limitations	33
12	Support	33
12.1	Checklist	33
12.2	FAQ	34

1 Introduction

This tutorial will demonstrate on a toy example how we recommend running *PureCN* on targeted sequencing data. To estimate tumor purity, we jointly utilize both target-level¹ coverage data and allelic fractions of single nucleotide variants (SNVs), inside - and optionally outside - the targeted regions. Knowledge of purity will in turn allow us to accurately (i) infer integer copy number and (ii) classify variants (somatic vs. germline, mono-clonal vs. sub-clonal, heterozygous vs. homozygous etc.).

This requires 3 basic input files:

1. A VCF file containing germline SNPs and somatic mutations. Somatic status is not required in case the variant caller was run without matching normal sample.
2. The tumor BAM file.
3. At least one BAM file from a normal control sample, either matched or process-matched.

In addition, we need to know a little bit more about the assay. This is the annoying step since here the user needs to provide some information. Most importantly, we need to know the positions of all targets. Then we need to correct for GC-bias, for which we need GC-content for each target. Optionally, if gene-level calls are wanted, we also need for each target a gene symbol. We may also observe subtle differences in coverage of tumor compared to normal due to varying proliferation rates and we can provide replication timing data to check and correct for such a bias. To obtain best results, we can finally use a pool of normal samples to automatically learn more about the assay and its biases and common artifacts.

The next sections will show how to do all this with *PureCN* alone or with the help of *GATK* and/or existing copy number pipelines.

All the steps described in the following are available in easy to use command line scripts described in a separate vignette.

2 Basic input files

2.1 VCF

Germline SNPs and somatic mutations are expected in a single VCF file. At the bare minimum, this VCF should contain read depths of reference and alt alleles in an AD format field and a DB info flag for membership in germline databases².

Without DB flag, variant ids starting with `rs` are assumed to be in dbSNP. Population allele frequencies are expected in a POP_AF info field. These frequencies are currently only used to infer membership in germline databases when the DB flag is missing; in future versions they will be used calculate somatic prior probabilities more accurately.

If a matched normal is available, then somatic status information is currently expected in a SOMATIC info flag in the VCF. The *VariantAnnotation* package provides examples how to add info fields to a VCF in case the used variant caller does not add this flag. If the VCF contains a BQ format field containing base quality scores, *PureCN* can remove low quality calls.

¹The captured genomic regions, e.g. exons.

²The name of the flag is customizable in `runAbsoluteCN`

Copy number calling and SNV classification using targeted short read sequencing

VCF files generated by *MuTect* [4] should work well and in general require no post-processing. *PureCN* can handle *MuTect* VCF files generated in both tumor-only and matched normal mode. Beta support for *MuTect 2* and *FreeBayes* VCFs is available.

As alternative to VCF input, read counts of common SNPs are also supported in case classification of somatic variants is of no interest and those genotyped SNPs are already generated in an upstream pipeline. The function `readAllelicCountsFile` can be used to parse a *GATK4 CollectAllelicCounts* file into a VCF supported by *PureCN*³:

```
@HD VN:1.6
@SQ SN:1 LN:248956422
@SQ SN:2 LN:242193529
...
@RG ID:GATK SM:Sample1
CONTIG POSITION REF_COUNT ALT_COUNT REF_NUCLEOTIDE ALT_NUCLEOTIDE
1 114515871 177 189 G A
1 150044293 119 157 T G
```

```
ac.file <- system.file("extdata", "example_allelic_counts.tsv",
  package="PureCN")
vcf.ac <- readAllelicCountsFile(ac.file)

## INFO [2025-10-07 19:30:12] Found 20 variants in VCF file.
## INFO [2025-10-07 19:30:12] 20 (100.0%) variants annotated as likely germline (DB INFO flag).
## WARN [2025-10-07 19:30:12] Did not find base quality scores, will use global error rate of 0.0010 instead
```

2.2 Target information

For the default segmentation function provided by *PureCN*, the algorithm first needs to calculate log2-ratios of tumor vs. normal control coverage. To do this, we need to know the locations of the captured genomic regions (targets). These are provided by the manufacturer of your capture kit⁴. Please double check that the genome version of the target file matches the reference. Usually the manufacturer provides two files: the baits file containing the coordinates of the actual capture baits, and the target file containing the coordinates of the actual regions we want to capture. We recommend to use the baits file (and recognize the confusing nomenclature that we follow due to convention in established tools).

Default parameters assume that these targets do NOT include a "padding" to include flanking regions. *PureCN* will automatically include variants in the 50bp flanking regions if the variant caller was either run without interval file or with interval padding (See section 12.2).

PureCN will attempt to optimize the targets for copy number calling (similar to [5]):

- Large targets are split to obtain a higher resolution
- Targets in regions of low mappability are dropped
- Optionally, accessible regions in-between the target (off-target) regions are included so that available coverage information in on- and off-target reads can be used by the segmentation function. In the following, we will use *intervals* when something applies to both on-target and off-target regions and *targets* when it only applies to on-target.

³Only the RG header line is required for extracting the sample id

⁴While *PureCN* can use a pool of normal samples to learn which intervals are reliable and which not, it is highly recommended to provide the correct intervals. Garbage in, garbage out.

Copy number calling and SNV classification using targeted short read sequencing

It further annotates intervals by GC-content (how coverage is normalized is described later in Section 3).

PureCN provides the `preprocessIntervals` function:

```
reference.file <- system.file("extdata", "ex2_reference.fa",
  package = "PureCN", mustWork = TRUE)
bed.file <- system.file("extdata", "ex2_intervals.bed",
  package = "PureCN", mustWork = TRUE)
mappability.file <- system.file("extdata", "ex2_mappability.bigWig",
  package = "PureCN", mustWork = TRUE)

intervals <- import(bed.file)
mappability <- import(mappability.file)

preprocessIntervals(intervals, reference.file,
  mappability = mappability, output.file = "ex2_gc_file.txt")

## WARN [2025-10-07 19:30:13] Found small target regions (< 100bp). Will resize them.
## INFO [2025-10-07 19:30:13] Splitting 5 large targets to an average width of 400.
## WARN [2025-10-07 19:30:13] No repriming scores provided.
## INFO [2025-10-07 19:30:13] Calculating GC-content...
```

A command line script described in a separate vignette provides convenient access to this function and also attempts to annotate the targets with gene symbols using the `annotateTargets` function.

2.3 Coverage data

The `calculateBamCoverageByInterval` function can be used to generate the required coverage data from BAM files. All we need to do is providing the desired intervals (as generated by `preprocessIntervals`):

```
bam.file <- system.file("extdata", "ex1.bam", package="PureCN",
  mustWork = TRUE)
interval.file <- system.file("extdata", "ex1_intervals.txt",
  package = "PureCN", mustWork = TRUE)

calculateBamCoverageByInterval(bam.file = bam.file,
  interval.file = interval.file, output.file = "ex1_coverage.txt")
```

2.4 Third-party coverage tools

Calculating coverage from BAM files is a common task and your pipeline might already provide this information. As alternative to `calculateBamCoverageByInterval`, *PureCN* currently supports coverage files generated by *GATK3 DepthOfCoverage*, *GATK4 CollectFragmentCounts* and by *CNVkit*. By providing files with standard file extension, *PureCN* will automatically detect the correct format and all following steps are the same for all tools. You will, however, still need the interval file generated in Section 2.2 and the third-party tool must use the exact same intervals. See also FAQ Section 12.2 for recommended settings for *GATK3 DepthOfCoverage*.

2.5 Third-party segmentation tools

PureCN integrates well with existing copy number pipelines. Instead of coverage data, the user then needs to provide either already segmented data or a wrapper function. This is described in Section 10.1.

2.6 Example data

We now load a few example files that we will use throughout this tutorial:

```
library(PureCN)

normal.coverage.file <- system.file("extdata", "example_normal.txt.gz",
  package = "PureCN")
normal2.coverage.file <- system.file("extdata", "example_normal2.txt.gz",
  package = "PureCN")
normal.coverage.files <- c(normal.coverage.file, normal2.coverage.file)
tumor.coverage.file <- system.file("extdata", "example_tumor.txt.gz",
  package = "PureCN")
seg.file <- system.file("extdata", "example_seg.txt",
  package = "PureCN")
vcf.file <- system.file("extdata", "example.vcf.gz", package = "PureCN")
interval.file <- system.file("extdata", "example_intervals.txt",
  package = "PureCN")
```

3 Library-specific coverage bias

In coverage normalization, we distinguish between assay- and library-specific biases. Assay-specific biases, for example due to probe density, probe capture efficiency and read mappability, are best removed with a pool of normal samples (Section 4.1, [5]). In other words, by examining the coverage of particular intervals in a pool of normals, we can estimate how well this assay captures these intervals and will then adjust the tumor coverage accordingly.

Other biases are library-specific, meaning a patient sample captured in different libraries may display dramatically different coverage profiles across libraries. Data from great sequencing centers usually show relatively small technical variance nowadays, but some biases are not completely avoidable. The most important library-specific bias is due to GC-content, i.e. regions of high AT- or GC-content are not always captured with exactly the same efficiency in tumor and normals.

We usually also observe that early replicating regions have a slightly higher coverage than late replicating regions [6, 7]. Since there is often a significant difference in proliferation rates of tumor and normal, the pool of normals might also not completely adjust for this small bias.

A good panel of normal samples will remove both assay- and library-specific biases. We recommend sequencing the normals in multiple batches so that the coverage normalization has access to any expected day-to-day assay variance. This will greatly increase the normalization's capability in removing library-specific biases.

Sometimes it can be beneficial to remove library-specific biases explicitly, especially with smaller panels of normals that might show significant differences to tumors (FFPE vs blood, high vs low coverage, etc.).

Copy number calling and SNV classification using targeted short read sequencing

In the following example, we correct the raw coverage of all samples, tumor and normal, for the two major sources of library-specific coverage biases mentioned above (Figure 1). For GC-normalization, we use a 2-step loess normalization [8]. For the replication timing bias, a linear model of log-transformed coverage and provided replication timing score is used.

```
correctCoverageBias(normal.coverage.file, interval.file,  
  output.file = "example_normal_loess.txt.gz", plot.bias = TRUE)
```

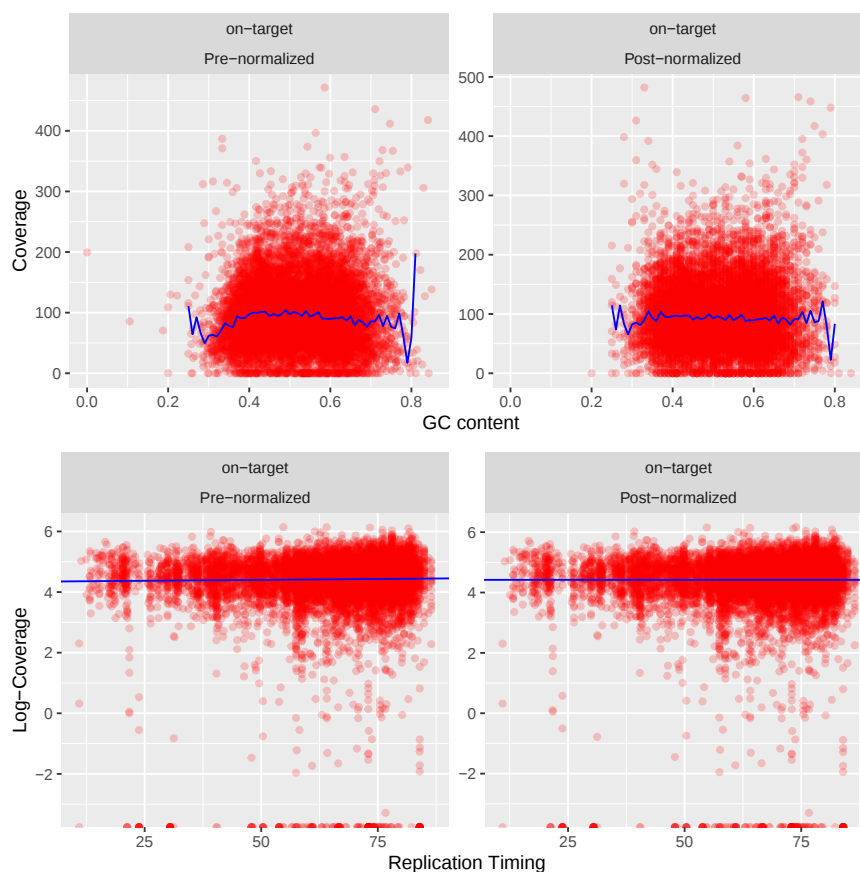


Figure 1: Coverage before and after normalization for GC-content and replication timing. This plot shows coverage as a function of on- and off-target GC-content and replication timing before and after normalization. Each dot is an interval. The example files are already GC-normalized; real data will show more dramatic differences.

The example coverage files are already GC-normalized. We provide a convenient command line script for generating normalized coverage data from BAM files or from *GATK* coverage files (see Quick vignette).

4 Pool of normals

4.1 Coverage normalization

For calculating copy number log2-ratios of tumor vs. normal, *PureCN* requires coverage from a process-matched normal sample. Using a normal that was sequenced using a similar, but not identical assay, rarely works. Differently covered genomic regions simply result in too many log2-ratio outliers. This section describes how to optimally normalize coverage against a pool of normals.

The `createNormalDatabase` function builds a database of coverage files (a command line script providing this functionality is described in a separate vignette):

```
normalDB <- createNormalDatabase(normal.coverage.files)

## INFO [2025-10-07 19:30:22] 576 on-target bins with low coverage in all samples.
## WARN [2025-10-07 19:30:22] You are likely not using the correct baits file!
## WARN [2025-10-07 19:30:22] Allosome coverage missing, cannot determine sex.
## WARN [2025-10-07 19:30:22] Allosome coverage missing, cannot determine sex.
## INFO [2025-10-07 19:30:22] Processing on-target regions...
## INFO [2025-10-07 19:30:23] Removing 930 intervals with low coverage in normalDB.
## INFO [2025-10-07 19:30:23] Removing 1 intervals with zero coverage in more than 3% of normalDB.

# serialize, so that we need to do this only once for each assay
saveRDS(normalDB, file = "normalDB.rds")
```

When using a GC-normalization in tumors, please make sure that all provided coverage files here were also GC-normalized prior to building the database (Section 3). Internally, `createNormalDatabase` determines the sex of the samples and trains a PCA that is later used for denoising tumor coverage using Tangent normalization [9]:

```
normalDB <- readRDS("normalDB.rds")
pool <- calculateTangentNormal(tumor.coverage.file, normalDB)
```

This `createNormalDatabase` function further automatically estimates copy number log2-ratio standard deviations. Assuming that all normal samples are in general diploid, a high variance in log2-ratio is indicative of an interval with either common germline alterations or frequent artifacts; high or low copy number log2-ratios in these intervals are unlikely measuring somatic copy number events. The segmentation function can use this information to skip over such noisy regions.

4.2 Artifact filtering

It is important to remove as many artifacts as possible, because low ploidy solutions are typically punished more by artifacts than high ploidy solutions. High ploidy solutions are complex and usually find ways of explaining artifacts reasonably well. The following steps in this section are optional, but recommended since they will reduce the number of samples requiring manual curation, especially when matching normal samples are not available.

Copy number calling and SNV classification using targeted short read sequencing

4.2.1 VCF

We recommend running *MuTect* with a pool of normal samples to filter common sequencing errors and alignment artifacts from the VCF. *MuTect* requires a single VCF containing all normal samples, for example generated by the *GATK3 CombineVariants* tool (see Section 12.2). For *Mutect 2* in *GATK4*, please follow the somatic best practices including the generation of a *GenomicsDB* containing variants from the normal samples.

It is highly recommended to provide *PureCN* these variant pools of normals as well; it will help the software correcting non-reference read mapping biases. This is described in the `setMappingBiasVcf` documentation. To reduce memory usage, the normal panel VCF can be reduced to contain only variants present in 2-3 or more samples. *GenomicsDBs* from *GATK4* can be provided as is, but require the installation of the *genomicsdb* R package from <https://github.com/nalinigans/GenomicsDB-R>.

Because these pools can become huge with increasing number of normal samples, we pre-compute the mapping bias:

```
# speed-up future runtimes by pre-calculating variant mapping biases
normal.panel.vcf.file <- system.file("extdata", "normalpanel.vcf.gz",
                                     package = "PureCN")

bias <- calculateMappingBiasVcf(normal.panel.vcf.file, genome = "h19")

## INFO [2025-10-07 19:30:25] Processing variants 1 to 50000...
## INFO [2025-10-07 19:30:26] Clustering beta binomial fits...
## INFO [2025-10-07 19:30:26] Assigning (6/20) variants a clustered beta binomial fit.

saveRDS(bias, "mapping_bias.rds")
mapping.bias.file <- "mapping_bias.rds"
```

This function will first fit beta-binomial distributions of allelic fractions from heterozygous variants (between fractions of 0.1 and 0.9) observed in sufficient number of samples; by default `min.normals.betafit` is set to 7. Under the assumption that there are only a limited number of error modes resulting in different distributions, we next cluster those position-specific distributions into `num.betafit.clusters` (default 9). Next, all variants observed heterozygous in less than `min.normals.betafit` but equal or more than `min.normals.assign.betafit` (default 3) are assigned the best fitting beta-distribution cluster.

For variants with at least `min.normals` (default 1) but less than `min.normals.assign.betafit`, the empirical Bayes estimate described in the original *PureCN* paper is used where the reference and alternate counts of an average high quality SNP is added to the observed counts, thus forcing the estimate to the average when position-specific information is limited.

Note that variants with low median coverage (`min.median.coverage.betafit`, default 5) are not attempted to fit to beta-binomial distributions and instead the empirical Bayes method is used.

4.3 Artifact filtering without a pool of normals

By default, *PureCN* will exclude targets with coverage below 15X from segmentation (with a pool of normals, targets are filtered based on the coverage and variance in normal database only). For variants in the provided VCF, the same 15X cutoff is applied. *MuTect* applies more sophisticated artifact tests and flags suspicious variants. If *MuTect* was run in matched normal mode, then both potential artifacts and germline variants are rejected, that means

Copy number calling and SNV classification using targeted short read sequencing

we cannot just filter by the PASS/REJECT *MuTect* flags. The `filterVcfMuTect` function optionally reads the *MuTect* 1.1.7 stats file and will keep germline variants, while removing potential artifacts. Without the stats file, *PureCN* will use only the filters based on read depths as defined in `filterVcfBasic`. Both functions are automatically called by *PureCN*, but can be easily modified and replaced if necessary.

We can also use a BED file to blacklist regions expected to be problematic, for example the simple repeats track from the UCSC:

```
# Instead of using a pool of normals to find low quality regions,
# we use suitable BED files, for example from the UCSC genome browser.

# We do not download these in this vignette to avoid build failures
# due to internet connectivity problems.
downloadFromUCSC <- FALSE
if (downloadFromUCSC) {
  library(rtracklayer)
  mySession <- browserSession("UCSC")
  genome(mySession) <- "hg19"
  simpleRepeats <- track(ucscTableQuery(mySession,
    track = "Simple Repeats", table = "simpleRepeat"))
  export(simpleRepeats, "hg19_simpleRepeats.bed")
}

snp.blacklist <- "hg19_simpleRepeats.bed"
```

5 Recommended run

Finally, we can run *PureCN* with all that information:

```
ret <- runAbsoluteCN(normal.coverage.file = pool,
  tumor.coverage.file = tumor.coverage.file, vcf.file = vcf.file,
  genome = "hg19", sampleid = "Sample1",
  interval.file = interval.file, normalDB = normalDB,
  # args.setMappingBiasVcf=list(
  #   mapping.bias.file = mapping.bias.file
  # ),
  args.filterVcf=list(
    # snp.blacklist = snp.blacklist,
    # stats.file = mutect.stats.file
  ),
  args.filterIntervals = list(
    min.total.counts = 50
  ),
  post.optimize = FALSE, plot.cnv = FALSE, verbose = FALSE)

## WARN [2025-10-07 19:30:27] Allosome coverage missing, cannot determine sex.
## WARN [2025-10-07 19:30:27] Allosome coverage missing, cannot determine sex.
```

The `normal.coverage.file` argument points to a coverage file obtained from either a matched or a process-matched normal sample, but can be also a small pool of best normals (Section 4.1).

Copy number calling and SNV classification using targeted short read sequencing

The `normalDB` argument (Section 4.1) provides a pool of normal samples and for example allows the segmentation function to skip targets with low coverage or common germline deletions in the pool of normals. If available, a VCF containing all variants from the normal samples should be provided via `args.setMappingBiasVcf` to correct read mapping biases. The files specified in `args.filterVcf` help [PureCN](#) filtering SNVs more efficiently for artifacts as described in Sections 4.2 and 4.3. The `snp.blacklist` is only necessary if neither a matched normal nor a large pool of normals is available. We set `min.total.counts` to a very low value to make this toy example use all simulated intervals.

The `post.optimize` flag will increase the runtime by about a factor of 2-5, but might return slightly more accurate purity estimates. For high quality whole-exome data, this is typically not necessary for copy number calling (but might be for variant classification, see Section 6.2.1). For smaller targeted panels, the runtime increase is typically marginal and `post.optimize` should be always set to `TRUE`.

The `plot.cnv` argument allows the segmentation function to generate

additional plots if set to `TRUE`. Finally, `verbose` outputs important and helpful information about all the steps performed and is therefore set to `TRUE` by default.

6 Output

6.1 Plots

We now create a few output files:

```
file.rds <- "Sample1_PureCN.rds"
saveRDS(ret, file = file.rds)
pdf("Sample1_PureCN.pdf", width = 10, height = 11)
plotAbs(ret, type = "all")
dev.off()

## pdf
## 2
```

The RDS file now contains the serialized return object of the `runAbsoluteCN` call. The PDF contains helpful plots for all local minima, sorted by likelihood. The first plot in the generated PDF is displayed in Figure 2 and shows the purity and ploidy local optima, sorted by final likelihood score after fitting both copy number and allelic fractions.

```
plotAbs(ret, type = "overview")
```

We now look at the main plots of the maximum likelihood solution in more detail.

```
plotAbs(ret, 1, type = "hist")
```

Figure 3 displays a histogram of tumor vs. normal copy number log2-ratios for the maximum likelihood solution (number 1 in Figure 2). The height of a bar in this plot is proportional to the fraction of the genome falling into the particular log2-ratio copy number range. The vertical dotted lines and numbers visualize the, for the given purity/ploidy combination, expected log2-ratios for all integer copy numbers from 0 to 7. It can be seen that most of the log2-ratios of the maximum likelihood solution align well to expected values for copy numbers of 0, 1, 2 and 4.

Copy number calling and SNV classification using targeted short read sequencing

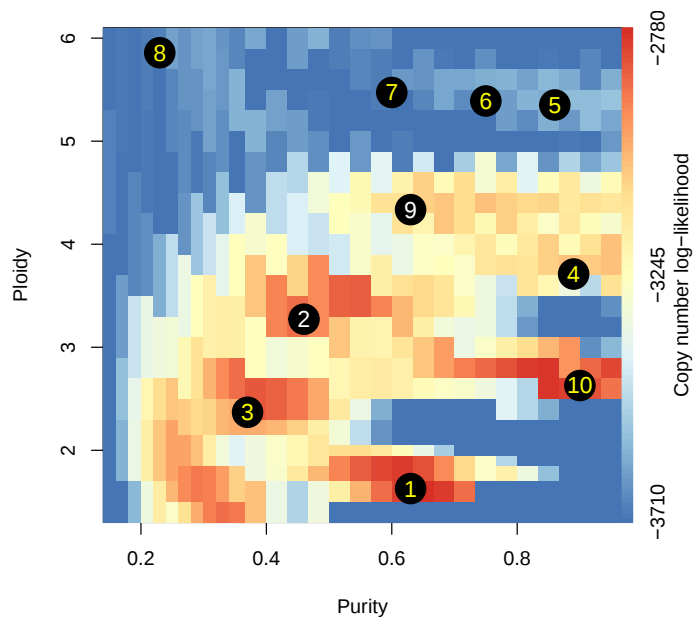


Figure 2: Overview. The colors visualize the copy number fitting score from low (blue) to high (red). The numbers indicate the ranks of the local optima. Yellow fonts indicate that the corresponding solutions were flagged, which does not necessarily mean the solutions are wrong. The correct solution (number 1) of this toy example was flagged due to large amount of LOH.

```
plotAbs(ret, 1, type = "BAF")
```

Germline variant data are informative for calculating integer copy number because unbalanced maternal and paternal chromosome numbers in the tumor portion of the sample lead to unbalanced germline allelic fractions. Figure 4 shows the allelic fractions of predicted germline SNPs. The goodness of fit (GoF) is provided on an arbitrary scale in which 100% corresponds to a perfect fit and 0% to the worst possible fit. The latter is defined as a fit in which allelic fractions on average differ by 0.2 from their expected fractions. Note that this does not take purity into account and low purity samples are expected to have a better fit. In the middle panel, the corresponding copy number log₂-ratios are shown. The lower panel displays the calculated integer copy numbers, corrected for purity and ploidy. We can zoom into particular chromosomes (Figure 5).

```
plotAbs(ret, 1, type = "BAF", chr = "chr19")
```

```
plotAbs(ret, 1, type = "AF")
```

Finally, Figure 6 provides more insight into how well the variants fit the expected values.

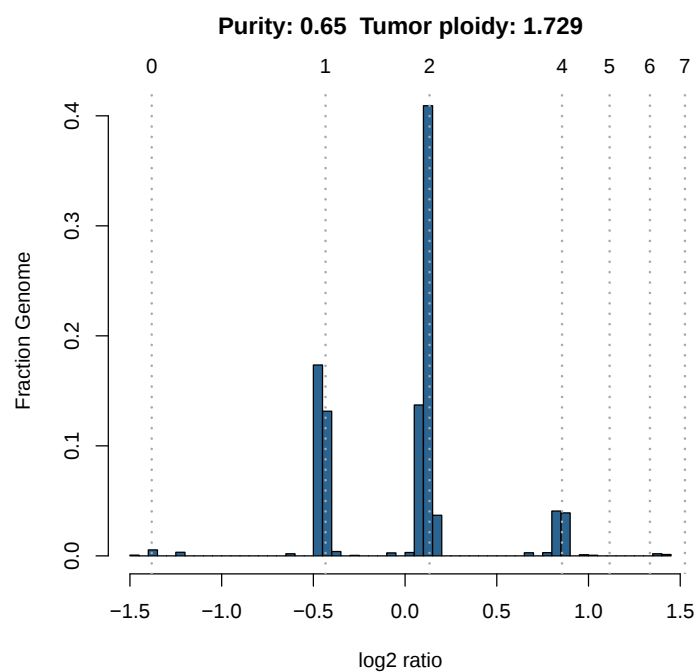


Figure 3: Log-ratio histogram.

Copy number calling and SNV classification using targeted short read sequencing

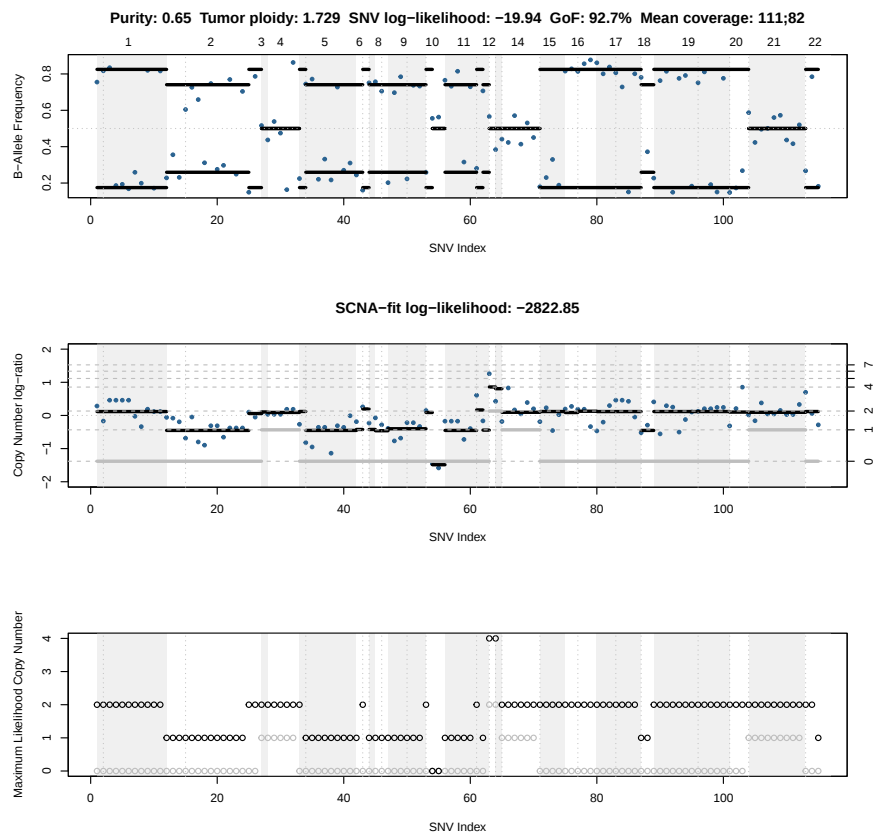


Figure 4: B-allele frequency plot. Each dot is a (predicted) germline SNP. The first panel shows the allelic fractions as provided in the VCF file. The alternating grey and white background colors visualize odd and even chromosome numbers, respectively. The black lines visualize the expected (not the average!) allelic fractions in the segment. These are calculated using the estimated purity and the total and minor segment copy numbers. These are visualized in black and grey, respectively, in the second and third panel. The second panel shows the copy number log₂-ratios, the third panel the integer copy numbers.

Copy number calling and SNV classification using targeted short read sequencing

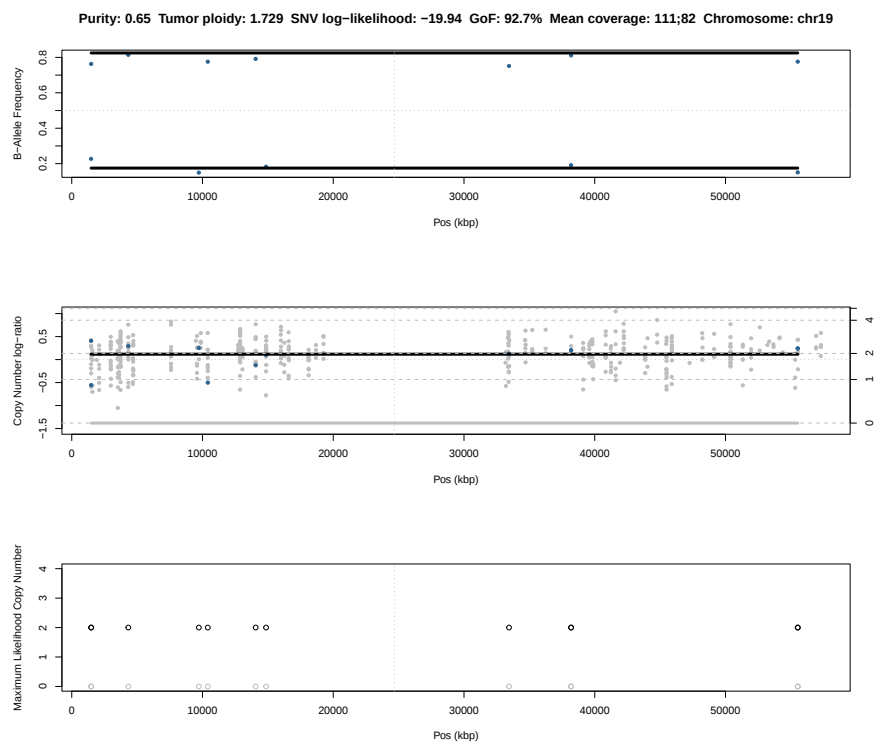


Figure 5: Chromosome plot. Similar to Figure 4, but zoomed into a particular chromosome. The grey dots in the middle panel visualize copy number log₂-ratios of targets without heterozygous SNPs, which are omitted from the previous genome-wide plot. The x-axis now indicates genomic coordinates in kbps.

Copy number calling and SNV classification using targeted short read sequencing

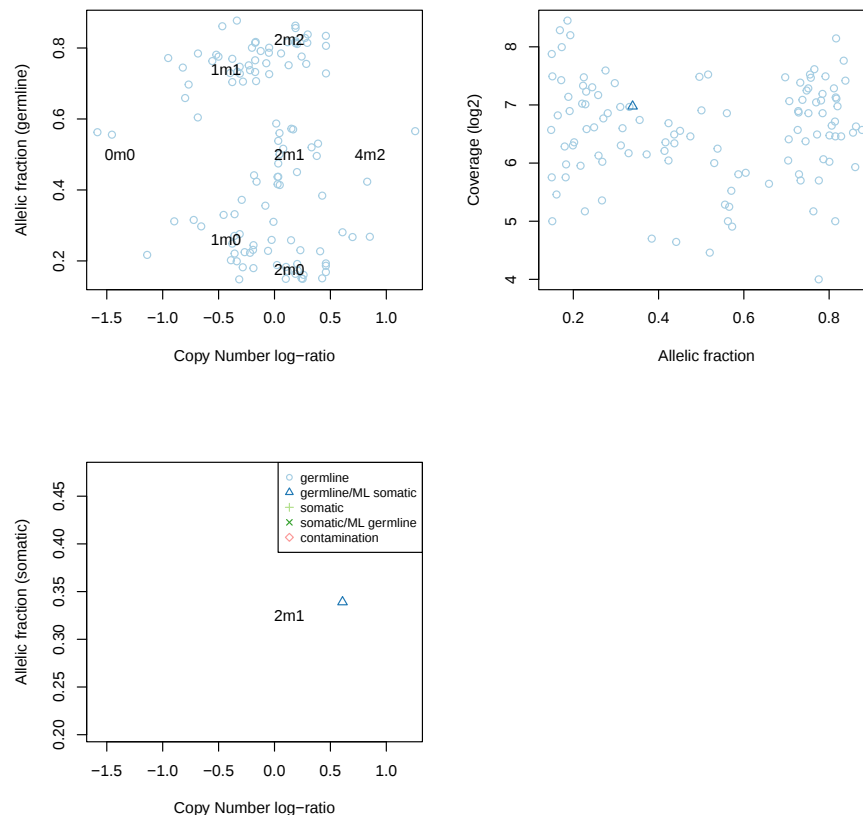


Figure 6: Allele fraction plots. Each dot is again a (predicted) germline SNP. The shapes visualize the different SNV groups based on prior and posterior probabilities. The labels show the expected values for all called states; 2m1 would be diploid, heterozygous, 2m2 diploid, homozygous. The relationship of allelic fraction and coverage is shown in the top right panel. This plot normally also shows somatic mutations in two additional panels, with the left panel showing the same plot as for germline SNPs and the bottom right panel a histogram of cellular fraction of predicted somatic mutations. This toy example contains only germline SNPs however.

6.2 Data structures

The R data file (`file.rds`) contains gene-level copy number calls, SNV status and LOH calls. The purity/ploidy combinations are sorted by likelihood and stored in `ret$results`.

```
names(ret)

## [1] "candidates" "results"    "input"      "version"
```

We provide convenient functions to extract information from this data structure and show their usage in the next sections. We recommend using these functions instead of accessing the data directly since data structures might change in future versions.

6.2.1 Prediction of somatic status and cellular fraction

To understand allelic fractions of particular SNVs, we must know the (i) somatic status, the (ii) tumor purity, the (iii) local copy number, as well as the (iv) number of chromosomes harboring the mutations or SNPs. One of *PureCN* main functions is to find the most likely combination of these four values. We further assign posterior probabilities to all possible combinations or states. Availability of matched normals reduces the search space by already providing somatic status.

The `predictSomatic` function provides access to these probabilities. For predicted somatic mutations, this function also provides cellular fraction estimates [10], i.e. the fraction of tumor cells with mutation. Fractions significantly below 1 indicate sub-clonality:

```
head(predictSomatic(ret), 3)
```

##	chr	start	end	ID	REF	ALT	SOMATIC.M0
## 1	chr1	114515871	114515871	chr1114515871xxx	G	A	4.046059e-102
## 2	chr1	150044293	150044293	chr1150044293xxx	T	G	8.631774e-92
## 3	chr1	158449835	158449835	chr1158449835xxx	A	G	1.248635e-147
##		SOMATIC.M1	SOMATIC.M2	SOMATIC.M3	SOMATIC.M4	SOMATIC.M5	SOMATIC.M6
## 1		8.420377e-39	2.107876e-07	5.330398e-310		0	0
## 2		4.835490e-39	2.452234e-10	4.617747e-312		0	0
## 3		2.177662e-62	5.337189e-15	5.337599e-314		0	0
##		SOMATIC.M7	GERMLINE.M0	GERMLINE.M1	GERMLINE.M2	GERMLINE.M3	
## 1		0	1.607523e-69	6.946405e-16	0.9999998	7.018564e-308	
## 2		0	3.623805e-64	1.081334e-18	1.0000000	3.778064e-306	
## 3		0	1.172374e-105	5.100431e-30	1.0000000	1.157132e-305	
##		GERMLINE.M4	GERMLINE.M5	GERMLINE.M6	GERMLINE.M7	GERMLINE.CONTHIGH	
## 1		0	0	0	0	5.799092e-42	
## 2		0	0	0	0	1.891829e-20	
## 3		0	0	0	0	1.087505e-25	
##		GERMLINE.CONTLOW	GERMLINE.HOMOZYGOUS	ML.SOMATIC	POSTERIOR.SOMATIC	ML.M	
## 1		8.219408e-289		0	FALSE	2.107876e-07	2
## 2		2.345594e-243		0	FALSE	2.452234e-10	2
## 3		0.000000e+00		0	FALSE	5.337189e-15	2
##	ML.C	ML.M.SEGMENT	M.SEGMENT.POSTERIOR	M.SEGMENT.FLAGGED	ML.AR	AR	
## 1	2	0	1	FALSE	0.825	0.755183	
## 2	2	0	1	FALSE	0.825	0.817078	
## 3	2	0	1	FALSE	0.825	0.834266	
##	AR.ADJUSTED	MAPPING.BIAS	ML.LOH	CN.SUBCLONAL	CELLFRACTION		
## 1	0.7775993	0.9711724	TRUE	FALSE	NA		

Copy number calling and SNV classification using targeted short read sequencing

```
## 2 0.8413316 0.9711724 TRUE FALSE NA
## 3 0.8590298 0.9711724 TRUE FALSE NA
## CELLFRACTION.95.LOWER CELLFRACTION.95.UPPER ALLELIC.IMBALANCE FLAGGED
## 1 NA NA -17.46605 FALSE
## 2 NA NA -20.15393 FALSE
## 3 NA NA -33.34446 FALSE
## log.ratio depth prior.somatic prior.contamination on.target seg.id
## 1 0.2842136 184 9.9e-05 0.01 1 1
## 2 -0.1686186 138 9.9e-05 0.01 1 1
## 3 0.4596841 217 9.9e-05 0.01 1 1
## pon.count gene.symbol
## 1 0 HIPK1
## 2 0 VPS45
## 3 0 <NA>
```

The output columns are explained in Table 1.

To annotate the input VCF file with these values:

```
vcf <- predictSomatic(ret, return.vcf = TRUE)
writeVcf(vcf, file = "Sample1_PureCN.vcf")
```

For optimal classification results:

- Set `post.optimize = TRUE` since small inaccuracies in purity can decrease the classification performance significantly
- Provide `args.setMappingBias` a pool of normal VCF to obtain position-specific mapping bias information
- Exclude variants in regions of low mappability
- Use a somatic posterior probability cutoff of 0.8 and 0.2 for somatic and germline variants, respectively. This appears to be a good compromise of call rate and accuracy. If the beta-binomial model was selected in the `model` argument of `runAbsoluteCN` with small panel of normals, these cutoffs might need to be relaxed to get acceptable call rates.
- Add a `Cosmic.CNT` info field to the VCF or provide a COSMIC VCF in `runAbsoluteCN` (see Section 10.2).

Note that the posterior probabilities assume that the purity and ploidy combination is correct. Before classifying variants, it is thus recommended to manually curate flagged samples.

6.2.2 Amplifications and deletions

To call amplifications, we recommend using a cutoff of 6 for focal amplifications and a cutoff of 7 otherwise. For homozygous deletions, a cutoff of 0.5 is useful to allow some heterogeneity in copy number.

For samples that failed *PureCN* calling we recommended using common log2-ratio cutoffs to call amplifications, for example 0.9.

This strategy is implemented in the `callAlterations` function:

Copy number calling and SNV classification using targeted short read sequencing

Table 1: `predictSomatic` output columns.

Column name	Description
<code>chr, start, end</code>	Variant coordinates
<code>ID</code>	The variant ID as provided in the VCF
<code>REF, ALT</code>	The reference and alt alleles
<code>SOMATIC.M*</code>	Posterior probabilities for all somatic states. M0 to M7 are multiplicity values, i.e. the number of chromosomes harboring the mutation (e.g. 1 heterozygous, 2 homozygous if copy number C is 2). <code>SOMATIC.M0</code> represents a sub-clonal state (somatic mutations by definition have a multiplicity larger than 0).
<code>GERMLINE.M*</code>	Posterior probabilities for all heterozygous germline states
<code>GERMLINE.CONTHIGH</code>	Posterior probability for contamination. This state corresponds to homozygous germline SNPs that were not filtered out because reference alleles from another individual were sequenced, resulting in allelic fractions smaller than 1.
<code>GERMLINE.CONTLOW</code>	Posterior probability for contamination. This state corresponds to non-reference alleles only present in the contamination.
<code>GERMLINE.HOMOZYGOUS</code>	Posterior probability that SNP is homozygous in normal. Requires the <code>model.homozygous</code> option in <code>runAbsoluteCN</code> . See Section 8.
<code>ML.SOMATIC</code>	TRUE if the maximum likelihood state is a somatic state
<code>POSTERIOR.SOMATIC</code>	The posterior probability that the variant is somatic (sum of all somatic state posterior probabilities)
<code>ML.M</code>	Maximum likelihood multiplicity
<code>ML.C</code>	Maximum likelihood integer copy number
<code>ML.M.SEGMENT</code>	Maximum likelihood minor segment copy number
<code>M.SEGMENT.POSTERIOR</code>	Posterior probability of <code>ML.M.SEGMENT</code>
<code>M.SEGMENT.FLAGGED</code>	Segment flag indicating <code>ML.M.SEGMENT</code> is unreliable, either due to low posterior probability (< 0.5) or few variants ($< \text{min.variants.segment}$). Indels are always flagged.
<code>ML.AR</code>	Expected allelic fraction of the maximum likelihood state
<code>AR</code>	Observed allelic fraction (as provided in VCF)
<code>AR.ADJUSTED</code>	Observed allelic fraction adjusted for mapping bias
<code>ML.LOH</code>	TRUE if variant is most likely in LOH
<code>CN.SUBCLONAL</code>	TRUE if copy number segment is sub-clonal
<code>CELLFRACTION</code>	Fraction of tumor cells harboring the somatic mutation [10]
<code>CELLFRACTION.95.LOWER</code>	Lower 95% of confidence interval
<code>CELLFRACTION.95.UPPER</code>	Upper 95% of confidence interval
<code>ALLELIC.IMBALANCE</code>	Log-likelihood that variant is in allelic balance. Requires position-specific mapping bias information to be reliable.
<code>FLAGGED</code>	Flag indicating that call is unreliable (currently only due to high mapping bias and high pool of normal counts)
<code>log.ratio</code>	The copy number log2-ratio (tumor vs. normal) for this variant
<code>depth</code>	The total sequencing depth at this position
<code>prior.somatic</code>	Prior probability of variant being somatic
<code>prior.contamination</code>	Prior probability that variant is contamination from another individual
<code>on.target</code>	1 for variants within intervals, 2 for variants in flanking regions, 0 for off-target variants
<code>seg.id</code>	Segment id
<code>pon.count</code>	Number of hits in the <code>mapping.bias.file</code>
<code>gene.symbol</code>	Gene symbol if available

Copy number calling and SNV classification using targeted short read sequencing

```
gene.calls <- callAlterations(ret)
head(gene.calls)
```

##	chr	start	end	C	C.flagged	seg.mean	seg.id	
##	EIF2A	chr3	150270143	150301699	6	FALSE	1.3859	5
##	MIR548H2	chr3	151531951	151538241	6	FALSE	1.3859	5
##	AADAC	chr3	151542451	151545961	6	FALSE	1.3859	5
##	GNMB	chr7	23286477	23313844	7	FALSE	1.4462	19
##	SH2D4B	chr10	82329910	82403838	0	FALSE	-1.4863	26
##	SLC35G1	chr10	95653791	95661248	0	FALSE	-1.2007	28

##	number.targets	gene.mean	gene.min	gene.max	focal	breakpoints	
##	EIF2A	13	1.5643189	0.7544062	2.1392167	TRUE	0
##	MIR548H2	3	0.8261594	0.4682628	1.2000253	TRUE	0
##	AADAC	2	0.7851630	0.6864438	0.8838822	TRUE	1
##	GNMB	11	1.4462179	0.9397072	1.8200775	TRUE	0
##	SH2D4B	6	-1.4862885	-1.5836964	-1.4113313	FALSE	0
##	SLC35G1	3	-1.5817134	-1.8066816	-1.4493561	FALSE	0

##	type	num.snps	M	M.flagged	loh	
##	EIF2A	AMPLIFICATION	0	NA	NA	NA
##	MIR548H2	AMPLIFICATION	0	NA	NA	NA
##	AADAC	AMPLIFICATION	0	NA	NA	NA
##	GNMB	AMPLIFICATION	0	NA	NA	NA
##	SH2D4B	DELETION	2	0	TRUE	TRUE
##	SLC35G1	DELETION	0	NA	NA	NA

It is also often useful to filter the list further by known biology, for example to exclude non-focal amplifications of tumor suppressor genes. The Sanger Cancer Gene Census [11] for example provides such a list.

The output columns of `callAlterations` are explained in Table 2.

6.2.3 Amplifications in low purity samples

The default `callAlterations` is not suited for samples of very low purity below 10% such as frequently observed in cfDNA samples. To call high level amplifications in samples between 4 to 10%, we recommend a simple z-score approach implemented in `callAmplificationsInLowPurity`. This will calculate gene-level log2 copy number ratio standard deviations in all normal samples that are then used to obtain gene-level cutoffs for amplifications in tumor samples. The output of this function is nearly identical to `callAlterations` with additional columns explained in Table 3.

```
gene.calls.zscore <- callAmplificationsInLowPurity(ret, normalDB)
## WARN [2025-10-07 19:31:46] Extensive noise in tumor compared to normals.
head(gene.calls.zscore)
```

##	chr	start	end	C	C.flagged	number.targets	gene.mean
##	EIF2A	chr3	150270143	150301699	NA	FALSE	13 1.5643189
##	MIR548H2	chr3	151531951	151538241	NA	FALSE	3 0.8261594
##	FTLP10	chr4	69056959	69078196	NA	FALSE	3 0.5852635
##	FGF5	chr4	81187979	81207827	NA	FALSE	3 1.0400133
##	SPP1	chr4	88898048	88904049	NA	FALSE	5 0.7200921
##	PPM1K	chr4	89183747	89199736	NA	FALSE	6 0.8253634

Copy number calling and SNV classification using targeted short read sequencing

Table 2: `callAlterations` output columns.

Column name	Description
<code>chr, start, end</code>	Gene coordinates
<code>C</code>	Segment integer copy number. See also <code>seg.id</code> .
<code>C.flagged</code>	Flagged when segment interval weights are low. Requires <code>normalDB</code> . In case multiple segments overlap, <code>C.flagged</code> will be <code>TRUE</code> when any segment is flagged.
<code>seg.mean</code>	Segment mean of copy number log2-ratios (not adjusted for purity/ploidy). See also <code>seg.id</code> .
<code>seg.id</code>	Segment id. In case multiple segments overlap, <code>seg.id</code> will point to the smallest overlapping segment and the column <code>breakpoints</code> will be non-zero.
<code>number.targets</code>	Number of targets annotated with this symbol
<code>gene.*</code>	Gene copy number log2-ratios (not adjusted for purity/ploidy). <code>gene.mean</code> is weighted by interval weights when these are available.
<code>focal</code>	<code>TRUE</code> for focal amplification, as defined by the <code>fun.focal</code> argument in <code>runAbsoluteCN</code>
<code>breakpoints</code>	Number of breakpoints between <code>start</code> and <code>end</code>
<code>type</code>	Amplification or deletion
<code>num.snps</code>	Number of SNPs in this segment informative for LOH detection (requires VCF)
<code>M</code>	minor copy number of segment (requires VCF)
<code>M.flagged</code>	flag indicating that <code>M</code> is unreliable (requires VCF)
<code>loh</code>	<code>TRUE</code> if segment is in LOH, <code>FALSE</code> if not and <code>NA</code> if number of SNPs is insufficient (requires VCF)

```
##          gene.min gene.max breakpoints p.value percentile.genome
## EIF2A      0.7544062 2.139217          0         0          100.00000
## MIR548H2  0.4682628 1.200025          0         0           95.05495
## FTLP10    0.4369763 1.262028          0         0           90.76923
## FGF5      0.8330433 1.249963          0         0           98.35165
## SPP1      0.4282286 1.061844          0         0           92.30769
## PPM1K     0.1361852 1.456415          0         0           94.94505
##          percentile.chromosome          type
## EIF2A                    100.00000 AMPLIFICATION
## MIR548H2                  95.65217 AMPLIFICATION
## FTLP10                    77.50000 AMPLIFICATION
## FGF5                      97.50000 AMPLIFICATION
## SPP1                      82.50000 AMPLIFICATION
## PPM1K                     92.50000 AMPLIFICATION

# filter to known amplifications
known.calls.zscore <- gene.calls.zscore[ gene.calls.zscore$gene.symbol %in%
c("AR", "BRAF", "CCND1", "CCNE1", "CDK4", "CDK6", "EGFR", "ERBB2", "FGFR1",
"FGFR2", "KIT", "KRAS", "MCL1", "MDM2", "MDM4", "MET", "MYC", "PDGFRA",
"PIK3CA", "RAF1"),]
```

Copy number calling and SNV classification using targeted short read sequencing

Since this algorithm is intended to rescue known high-level amplifications in very low tumor purities, we recommend using a liberal p-value cutoff, but restricting the search to a limited set of genes.

Table 3: `callAmplificationsInLowPurity` output columns not already explained in Table 2.

Column name	Description
<code>p.value</code>	P-value of observing such a gene-level log2-ratio in normal samples
<code>percentile.genome</code>	Percentile of gene-level log2-ratio in the tumor sample. Useful for distinguishing focal amplifications from broad gains.
<code>percentile.chromosome</code>	Same as <code>percentile.genome</code> , but for the chromosome

6.2.4 Find genomic regions in LOH

The `gene.calls` data.frame described above provides gene-level LOH information. To find the corresponding genomic regions in LOH, we can use the `callLOH` function:

```
loh <- callLOH(ret)
head(loh)
```

##	chr	start	end	arm	C	M		type	seg.mean
## 1	chr1	114515871	121535434	p	2	0	WHOLE ARM	COPY-NEUTRAL LOH	0.1173894
## 2	chr1	124535434	247419499	q	2	0	WHOLE ARM	COPY-NEUTRAL LOH	0.1173894
## 3	chr2	10262881	92326171	p	1	0		WHOLE ARM LOH	-0.4517019
## 4	chr2	95326171	231775678	q	1	0		WHOLE ARM LOH	-0.4517019
## 5	chr2	236403331	239039169	q	2	0		COPY-NEUTRAL LOH	0.0579000
## 6	chr3	11888043	90504854	p	2	1			0.1056000
##	num.mark	num.snps	M.flagged	maf.expected	maf.estimated				
## 1	886	11	FALSE	0.1750000	0.1767647				
## 2	886	11	FALSE	0.1750000	0.1767647				
## 3	657	13	FALSE	0.2592593	0.2751720				
## 4	657	13	FALSE	0.2592593	0.2751720				
## 5	86	2	TRUE	0.1750000	0.1719800				
## 6	403	1	TRUE	0.5000000	0.4685382				

The output columns are explained in Table 4.

Table 4: `callLOH` output columns.

Column name	Description
<code>chr, start, end</code>	Segment coordinates
<code>arm</code>	Chromosome arm
<code>C</code>	Segment integer copy number
<code>M</code>	Minor integer copy number ($M + N = C, M \leq N$)
<code>type</code>	LOH type if $M = 0$
<code>seg.mean</code>	Segment mean of copy number log2-ratios (not adjusted for purity/ploidy)
<code>num.mark</code>	Number of intervals in this segment (following DNACopy naming convention)
<code>num.snps</code>	Number of SNPs in this segment informative for LOH detection
<code>M.flagged</code>	flag indicating that <code>M</code> is unreliable (due to small number of variants defined by <code>min.variants.segment</code> or because of ambiguous <code>M</code> call)
<code>maf.expected</code>	Expected minor allele frequency of heterozygous SNPs in this segment
<code>maf.observed</code>	Median of observed minor allele frequency of heterozygous SNPs in this segment

7 Curation

For prediction of variant status (germline vs. somatic, sub-clonal vs. clonal, homozygous vs. heterozygous), it is important that both purity and ploidy are correct. We provide functionality for curating results:

```
createCurationFile(file.rds)
```

This will generate a CSV file in which the correct purity and ploidy values can be manually entered. It also contains a column "Curated", which should be set to `TRUE`, otherwise the file will be overwritten when re-run.

Then in R, the correct solution (closest to the combination in the CSV file) can be loaded with the `readCurationFile` function:

```
ret <- readCurationFile(file.rds)
```

This function has various handy features, but most importantly it will re-order the local optima so that the curated purity and ploidy combination is ranked first. This means `plotAbs(ret,1,type="hist")` would show the plot for the curated purity/ploidy combination, for example.

The default curation file will list the maximum likelihood solution:

```
read.csv("Sample1_PureCN.csv")

## Sampleid Purity Ploidy Sex Contamination Flagged Failed Curated
## 1 Sample1 0.65 1.728527 ? 0 TRUE FALSE FALSE
##
## Comment
## 1 EXCESSIVE LOSSES;EXCESSIVE LOH
```

Copy number calling and SNV classification using targeted short read sequencing

PureCN currently only flags samples with warnings, it does not mark any samples as failed. The `Failed` column in the curation file can be used to manually flag samples for exclusion in downstream analyses. See Table 5 for an explanation of all flags.

Table 5: `createCurationFile` flags.

Flag	Description
EXCESSIVE LOH	> 50% of genome in LOH and ploidy < 2.6
EXCESSIVE LOSSES	$\geq 1\%$ of genome deleted
HIGH AT- OR GC-DROPOUT	High GC-bias exceeding cutoff in <code>max.dropout</code>
HIGH PURITY	(when <code>model.homozygous=FALSE</code>). For very high purity samples, it is recommended to set <code>model.homozygous=TRUE</code> . See Section 8.
LOW PURITY	Purity < 30%
LOW BOOTSTRAP VALUE	<code>bootstrapResults</code> identified multiple plausible solutions
NOISY LOG-RATIO	Log-ratio standard deviation > <code>max.logr.sdev</code>
NOISY SEGMENTATION	More than <code>max.segments</code>
NON-ABERRANT	$\geq 99\%$ of genome has identical copy number and $\geq 0.5\%$ has second most common state
POLYGENOMIC	$\geq 0.75 \times \text{max.non.clonal}$ fraction of the genome in sub-clonal state
POOR GOF	GoF < <code>min.gof</code>
POTENTIAL SAMPLE CONTAMINATION	Significant portion of variants present in germline databases are potentially cross-contaminated
RARE KARYOTYPE	Ploidy < 1.5 or > 4.5

8 Cell lines

Default parameters assume some normal contamination. In 100% pure samples without matching normal samples such as cell lines, we cannot distinguish homozygous SNPs from LOH by looking at single allelic fractions alone. It is thus necessary to keep homozygous variants and include this uncertainty in the likelihood model. This is done by setting the `runAbsoluteCN` argument `model.homozygous=TRUE`. If matched normals are available, then variants homozygous in normal are automatically removed since they are uninformative.

For technical reasons, the maximum purity *PureCN* currently models is 0.99. We recommend setting `test.purity=seq(0.9,0.99,by=0.01)` in `runAbsoluteCN` for cell lines.

Please note that in order to detect homozygous deletions in 100% pure samples, you will need to provide a `normalDB` in `runAbsoluteCN` to filter low quality targets efficiently (Section 5).

9 Maximizing the number of heterozygous SNPs

It is possible to use SNPs in off-target reads in the variant fitting step by running *MuTect* without interval file and then setting the `filterVcfBasic` argument `remove.off.target.snvs` to `FALSE`. We recommend a large pool of normals in this case and then generating SNP blacklists as described in Sections 4.2 and 4.3. Remember to also run all the normals in *MuTect* without interval file.

An often better alternative to including all off-target reads is only including variants in the flanking regions of targets (between 50-100bp). This will usually significantly increase the number of heterozygous SNPs (see Section 12.2). These SNPs are automatically added if the variant caller was run without interval file or with interval padding.

10 Advanced usage

10.1 Custom normalization and segmentation

Copy number normalization and segmentation are crucial for obtaining good purity and ploidy estimates. If you have a well-tested pipeline that produces clean results for your data, you might want to use *PureCN* as add-on to your pipeline. By default, we will use *DNAcopy* [12] to segment normalized target-level coverage log2-ratios. It is straightforward to replace the default with other methods and the `segmentationCBS` function can serve as an example.

The next section describes how to replace the default segmentation. For the probably more uncommon case that only the coverage normalization is performed by third-party tools, see Section 10.1.2.

10.1.1 Custom segmentation

It is possible to provide already segmented data, which is especially recommended when third-party segmentation tools are not written in R. Otherwise it is usually however better to customize the default segmentation function, since the algorithm then has access to the raw log2-ratio distribution⁵. The expected file format for already segmented copy number data parsed by `readSegmentationFile` is⁶:

ID	chrom	loc.start	loc.end	num.mark	seg.mean
Sample1	1	61723	5773942	2681	0.125406444072723
Sample1	1	5774674	5785170	10	-0.756511807441712

Since its likelihood model is exon-based, *PureCN* currently still requires an interval file to generate simulated target-level log2-ratios from a segmentation file. For simplicity, this interval file is expected either via the `tumor.coverage.file` or via the `interval.file` argument (see Figure 7). Note that *PureCN* will re-segment the simulated log2-ratios using the default `segmentationCBS` function, in particular to identify regions of copy-number neutral LOH and to cluster segments with similar allelic imbalance and log2-ratio. The provided interval file should therefore cover all significant copy number alterations⁷. Please check that the log2-ratios are similar to the ones obtained by the default *PureCN* segmentation and normalization.

⁵If the third-party tool provides target-level log2-ratios, then these can be provided via the `log.ratio` argument in addition to `seg.file` though. See also Section 10.1.2.

⁶This segmentation file can contain multiple samples, in which case the provided `sampleid` must match a sample in the column `ID`

⁷If this behaviour is not wanted, because maybe the custom function already identifies CNNLOH reliably, `segmentationCBS` can be replaced with a minimal version.

Copy number calling and SNV classification using targeted short read sequencing

```
retSegmented <- runAbsoluteCN(seg.file = seg.file,
  interval.file = interval.file, vcf.file = vcf.file,
  max.candidate.solutions = 1, genome = "hg19",
  test.purity = seq(0.3,0.7,by = 0.05), verbose = FALSE,
  plot.cnv = FALSE)

## WARN [2025-10-07 19:32:27] Allosome coverage missing, cannot determine sex.
## WARN [2025-10-07 19:32:27] Allosome coverage missing, cannot determine sex.
```

The `max.candidate.solutions` and `test.purity` arguments are set to non-default values to reduce the runtime of this vignette.

```
plotAbs(retSegmented, 1, type = "BAF")
```

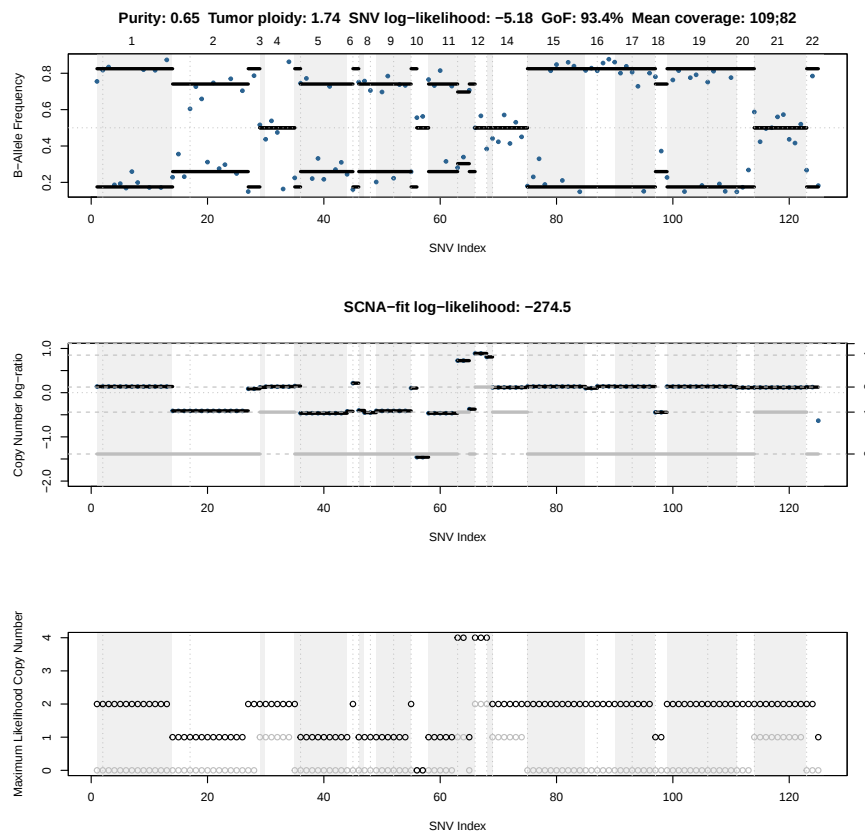


Figure 7: B-allele frequency plot for segmented data. This plot shows the maximum likelihood solution for an example where segmented data are provided instead of coverage data. Note that the middle panel shows no variance in log₂-ratios, since only segment-level log₂-ratios are available.

10.1.2 Custom normalization

If third-party tools such as *GATK4* are used to calculate target-level copy number log₂-ratios, and *PureCN* should be used for segmentation and purity/ploidy inference only, it is possible to provide these log₂-ratios:

Copy number calling and SNV classification using targeted short read sequencing

```
# We still use the log2-ratio exactly as normalized by PureCN for this
# example
log.ratio <- calculateLogRatio(readCoverageFile(normal.coverage.file),
                              readCoverageFile(tumor.coverage.file))

retLogRatio <- runAbsoluteCN(log.ratio = log.ratio,
                             interval.file = interval.file, vcf.file = vcf.file,
                             max.candidate.solutions = 1, genome = "hg19",
                             test.purity = seq(0.3, 0.7, by = 0.05), verbose = FALSE,
                             normalDB = normalDB, plot.cnv = FALSE)

## WARN [2025-10-07 19:32:47] Allosome coverage missing, cannot determine sex.
## WARN [2025-10-07 19:32:47] Allosome coverage missing, cannot determine sex.
```

Again, the `max.candidate.solutions` and `test.purity` arguments are set to non-default values to reduce the runtime of this vignette. Note that this example uses a pool of normals to filter low quality targets. Interval coordinates are again expected in either a `interval.file` or a `tumor.coverage.file`.

The `readLogRatioFile` function provides support for parsing *GATK3*-style and *GATK4 DenoiseReadCounts* formats:

```
# GATK3-style
Target log_ratio
1:3682750-3683489 0.109473
1:3690585-3691264 0.126205

# GATK4
@HD VN:1.6
@SQ SN:1 LN:248956422
@SQ SN:2 LN:242193529
...
CONTIG START END LOG2_COPY_RATIO
1 3682750 3683489 0.109473
1 3690585 3691264 0.126205
```

It is highly recommended to compare the log2-ratios obtained by *PureCN* and the third-party tool, since some pipelines automatically adjust log2-ratios for a default purity value.

10.2 COSMIC annotation

If a matched normal is not available, it is also helpful to provide `runAbsoluteCN` the COSMIC database [13] via `cosmic.vcf.file` (or via a `Cosmic.CNT` INFO field in the VCF). While this has limited effect on purity and ploidy estimation due the sparsity of hotspot mutations, it often helps in the manual curation to compare how well high confidence germline (based on dbSNP membership or POPAF) vs. somatic (COSMIC) variants fit a particular purity/ploidy combination.

For variant classification (Section 6.2.1), providing COSMIC annotation also avoids that hotspot mutations in germline databases get assigned a very low prior probability of being somatic.

10.3 ExAC and gnomAD annotation

PureCN is not automatically annotating input VCFs with data from common germline databases such as ExAC. See Section 2.1 for ways to tell *PureCN* where to find either a summary binary flag (i.e. likely germline yes/no) or population allele frequencies.

10.4 Mutation burden

The `predictSomatic` function described in Section 6.2.1 can be used to efficiently remove private germline mutations. This in turn allows the calculation of mutation burden for unmatched tumor samples. A wrapper function for this specific task is included as `callMutationBurden`:

```
callableBed <- import(system.file("extdata", "example_callable.bed.gz",
  package = "PureCN"))

callMutationBurden(ret, callable = callableBed)

## somatic.ontarget somatic.all private.germline.ontarget
## 1 0 0 0
## private.germline.all callable.bases.ontarget callable.bases.flanking
## 1 0 1493367 2194877
## callable.bases.all somatic.rate.ontarget somatic.rate.ontarget.95.lower
## 1 3063762 0 0
## somatic.rate.ontarget.95.upper private.germline.rate.ontarget
## 1 2.470173 0
## private.germline.rate.ontarget.95.lower
## 1 0
## private.germline.rate.ontarget.95.upper
## 1 2.470173
```

The `callableBed` file should be a file parsable by *rtracklayer*. This file can specify genomic regions that are callable, for example as obtained by *GATK3 CallableLoci*. This is optional, but if provided can be used to accurately calculate mutation rates per megabase. Variants outside the callable regions are not counted. Private germline rates should be fairly constant across samples; outliers here should be manually inspected.

The output columns are explained in Table 6.

10.5 Chromosomal Instability

Chromosomal Instability (CIN) is usually defined as the fraction of the genome that is altered. The `callCIN` function can be used to estimate this fraction.

Parameters define regions that are altered. First, `allele.specific` defines whether only total vs. both minor and major copy number define a state. Copy number neutral LOH would count as altered only when this parameter is set to `TRUE`. Second, `reference.state` defines the unaltered copy number state. This can be either `normal` for 2/1, or `dominant` for the most common state. While technically potentially wrong, the latter is robust to errors in ploidy and is thus recommended without careful manual curation. Similarly, setting `allele.specific` to `FALSE` makes this metric more robust to purity and ploidy errors, but usually to a much lesser extent.

Copy number calling and SNV classification using targeted short read sequencing

Table 6: `callMutationBurden` output columns.

Column name	Description
<code>somatic.ontarget</code>	Number of mutations in target regions
<code>somatic.all</code>	Total number of mutations. Depending on input VCF and <code>runAbsoluteCN</code> arguments, this might include calls in flanking regions and off-targets reads.
<code>private.germline.ontarget</code>	Number of private germline SNPs in targets. For matched tumor/normal samples, this is the number of variants annotated as somatic, but classified as germline.
<code>private.germline.all</code>	Total number of private germline SNPs
<code>callable.bases.ontarget</code>	Number of callable on-target bases
<code>callable.bases.flanking</code>	Number of callable on-target and flanking bases
<code>callable.bases.all</code>	Total number of callable bases. With default parameters includes off-target regions that were ignored by <code>runAbsoluteCN</code> .
<code>somatic.rate.ontarget</code>	Somatic mutations per megabase in target regions
<code>somatic.rate.ontarget.95.lower</code>	Lower 95% of confidence interval
<code>somatic.rate.ontarget.95.upper</code>	Upper 95% of confidence interval
<code>private.germline.rate.ontarget</code>	Private germline mutations per megabase in target regions
<code>private.germline.rate.ontarget.95.lower</code>	Lower 95% of confidence interval
<code>private.germline.rate.ontarget.95.upper</code>	Upper 95% of confidence interval

It is recommended to test for a relationship of tumor purity and CIN and if necessary exclude low purity samples with uncertain CIN.

```
# All 4 possible ways to calculate fraction of genome altered
cin <- data.frame(
  cin = callCIN(ret, allele.specific = FALSE, reference.state = "normal"),
  cin.allele.specific = callCIN(ret, reference.state = "normal"),
  cin.ploidy.robust = callCIN(ret, allele.specific = FALSE),
  cin.allele.specific.ploidy.robust = callCIN(ret)
)
```

10.6 Detect cross-sample contamination

It is important to correctly handle heterozygous common SNPs that do not have an expected allelic fraction of 0.5 in normal samples. These can be SNPs in poor quality regions (as already described, see Section 4.2.1), but also SNPs from cross-sample contaminated DNA. Without matched normals, detection of those problematic SNPs is not trivial.

Copy number calling and SNV classification using targeted short read sequencing

For cross-sample contamination, *PureCN* by default always tests for a 1% contamination and assigns common SNPs to a contamination state when allelic fractions are either close to 0 or close to 1 and when this cannot be explained by CNAs. The main purpose of these states is to provide a bin for common SNPs that for artifactual reasons do not fit any other state.

This tool applies a simple heuristic to flag samples for cross-contamination: Given the coverage and putative contamination rate based on allelic fractions of potentially contaminated SNPs, how many SNPs do we expect to detect based on our power to detect variants at that contamination rate? If the expected number is much higher than observed, then significant contamination is unlikely; observed SNPs close to 0 or 1 are more likely artifacts or the contamination rate is much lower than the minimum tested. Otherwise *PureCN* will perform a post-optimization in which contamination rate is optimized in additional variant fitting steps.

Cross-sample contamination can also result in increased observed heterozygosity on chrX for males, which in turn often results in a *PureCN* warning that sex inferred from coverage and VCF are in conflict.

By default, cross-contamination is tested in the range from 1 to 7.5%. Catastrophic failures with higher contamination might not get flagged.

10.7 Power to detect somatic mutations

As final quality control step, we can test if coverage and tumor purity are sufficient to detect mono-clonal or even sub-clonal somatic mutations. We strictly follow the power calculation by Carter et al. [2].

The following Figure 8 shows the power to detect mono-clonal somatic mutations as a function of tumor purity and sequencing coverage (reproduced from [2]):

```
purity <- c(0.1,0.15,0.2,0.25,0.4,0.6,1)
coverage <- seq(5,35,1)
power <- lapply(purity, function(p) sapply(coverage, function(cv)
  calculatePowerDetectSomatic(coverage = cv, purity = p, ploidy = 2,
    verbose = FALSE)$power))

# Figure S7b in Carter et al.
plot(coverage, power[[1]], col = 1, xlab = "Sequence coverage",
  ylab = "Detection power", ylim = c(0, 1), type = "l")

for (i in 2:length(power)) lines(coverage, power[[i]], col = i)
abline(h = 0.8, lty = 2, col = "grey")
legend("bottomright", legend = paste("Purity", purity),
  fill = seq_along(purity))
```

Figure 9 then shows the same plot for sub-clonal mutations present in 20% of all tumor cells:

```
coverage <- seq(5,350,1)
power <- lapply(purity, function(p) sapply(coverage, function(cv)
  calculatePowerDetectSomatic(coverage = cv, purity = p, ploidy = 2,
    cell.fraction = 0.2, verbose = FALSE)$power))
plot(coverage, power[[1]], col = 1, xlab = "Sequence coverage",
  ylab="Detection power", ylim = c(0, 1), type = "l")

for (i in 2:length(power)) lines(coverage, power[[i]], col = i)
```

Copy number calling and SNV classification using targeted short read sequencing

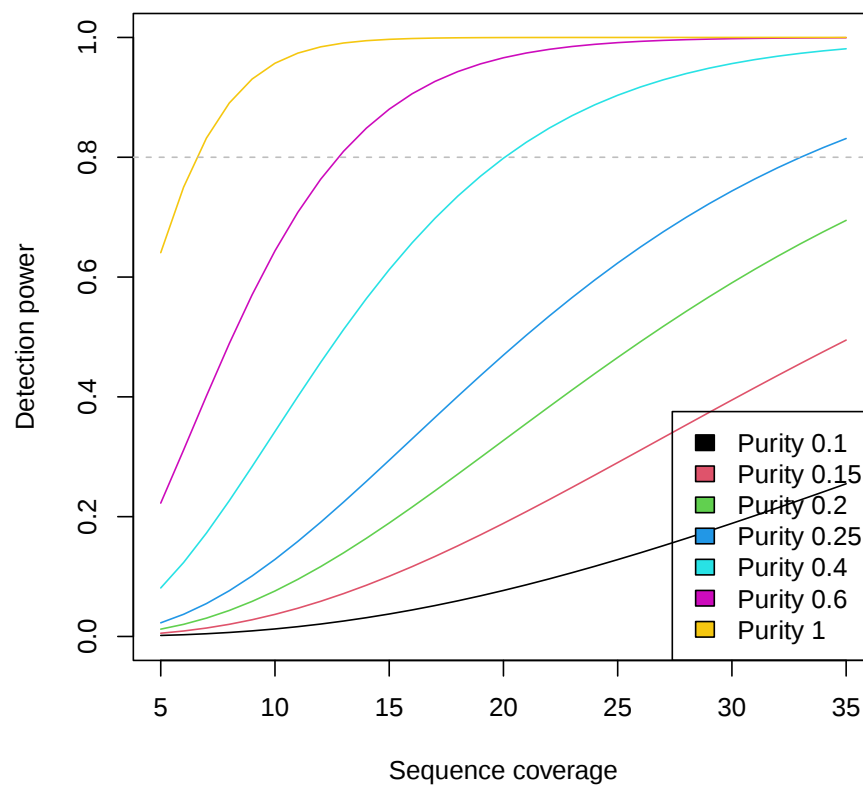


Figure 8: Power to detect mono-clonal somatic mutations. Reproduced from [2].

```
abline(h = 0.8, lty = 2, col = "grey")
legend("bottomright", legend = paste("Purity", purity),
      fill = seq_along(purity))
```

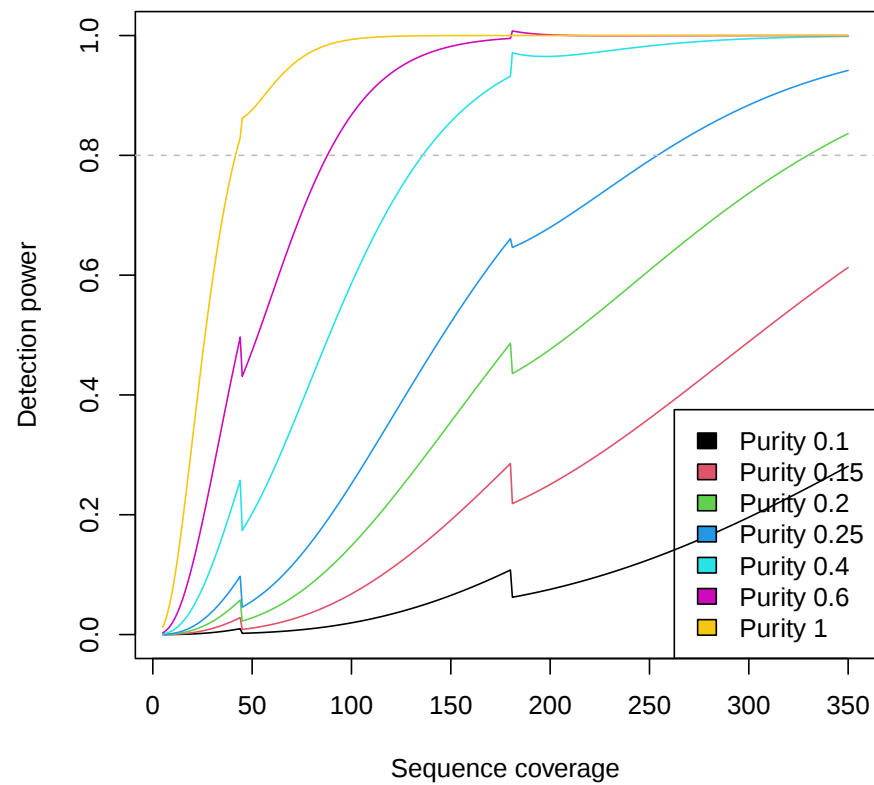


Figure 9: Power to detect sub-clonal somatic mutations present in 20% of all tumor cells. Reproduced from [2].

11 Limitations

PureCN currently assumes a completely diploid normal genome. For human samples, it tries to detect sex by calculating the coverage ratio of chromosomes X and Y and will then remove sex chromosomes in male samples⁸. For non-human samples, the user needs to manually remove all non-diploid chromosomes from the coverage data and specify `sex="diploid"` in the *PureCN* call.

While *PureCN* supports and models sub-clonal somatic copy number alterations, it currently assumes that the majority of alterations are mono-clonal. For most clinical samples, this is reasonable, but very heterogeneous samples are likely not possible to call without manual curation. Poly-genomic tumors are often called as high ploidy or low purity. The former usually happens when sub-clonal losses are called as 2 copies and mono-clonal losses correctly as 1 copy. The latter when sub-clonal losses are called mono-clonal, which only happens when there are far more sub-clonal than mono-clonal losses. Please note however that unless purities are very high, algorithms that model poly-genomic tumors do not necessarily have a higher call rate, since they tend to overfit noisy samples or similarly confuse true high-ploidy with poly-genomic tumors. Due to the lack of signal, manual curation is also recommended in low purity samples or very quiet genomes.

12 Support

If you encounter bugs or have problems running *PureCN*, please post them at

- https://support.bioconductor.org/p/new/post/?tag_val=PureCN or
- <https://github.com/lima1/PureCN/issues>.

If *PureCN* throws user errors, then there is likely a problem with the input files. If the error message is not self-explanatory, feel free to seek help at the support site.

In your report, please add the outputs of the `runAbsoluteCN` call (with `verbose=TRUE`, or the `*.log` file in *PureCN.R*) and `sessionInfo()`. Please also check that your problem is not already covered in the following sections.

For general feedback such as suggestions for improvements, feature requests, complaints, etc. please do not hesitate to send us an email.

12.1 Checklist

- Used the correct interval files provided by the manufacturer of the capture kit and the genome version of the interval file matches the reference. Ideally used the baits file, not the targets file (in Agilent data, the baits files are called "covered" and the targets are "regions").
- For hybrid capture data, included off-target reads in the coverage calculation
- BAM files were generated following established best practices and tools finished successfully.
- Checked standard QC metrics such as AT/GC dropout and duplication rates.

⁸Loss of Y chromosome (LOY) can result in wrong female calls, especially in high purity samples or if LOY is in both tumor and contaminating normal cells. The software throws a warning in this case when germline SNPs are provided.

Copy number calling and SNV classification using targeted short read sequencing

- Tumor and normal data were obtained using the same capture kit and pipeline.
- Coverage data of tumor and normal were GC-normalized.
- The VCF file contains germline variants (i.e. not only somatic calls).
- Maximized the number of high coverage heterozygous SNPs, for example by running *MuTect* with a 50-75bp interval padding (Section 9). The `runAbsoluteCN` output lists the percentage of targets with variants and this should be around 10-15%. Ultra-deep sequencing data can provide good SNP allelic fractions in the 100-200bp flanking regions.
- If a pool of normal samples is available, followed the steps in Section 4.2.
- Read the output of `runAbsoluteCN` with `verbose = TRUE`, fixed all warnings.
- If third-party segmentation tools are used, checked that normalized log2-ratios are not biased, i.e. very similar compared to *PureCN* log2-ratios (some pipelines already adjust for a default normal contamination).
- Followed the best practices described in the Quick vignette

12.2 FAQ

If the ploidy is frequently too high, please check:

- Does the log2-ratio histogram (Figure 3) look noisy? If yes, then
 - Is the coverage sufficient? Tumor coverages below 80X can be difficult, especially in low purity samples. Normal coverages below 50X might result in high variance of log2-ratios. See Section 4.1 for finding a good normal sample for log2-ratio calculation.
 - Is the coverage data of both tumor and normal GC-normalized? If not, see `correctCoverageBias`. It can be safe to skip the normalization, with large panels of normals potentially even slightly beneficial, but make sure that GC-normalization was either skipped or performed in both tumor sample and normal coverage database.
 - Is the quality of both tumor and normal sufficient? A high AT or GC-dropout might result in high variance of log2-ratios. Challenging FFPE samples also might need parameter tuning of the segmentation function. See `segmentationCBS`. A high expected tumor purity allows more aggressive segmentation parameters, such as `prune.hclust.h=0.2` or higher.
 - Was the correct target interval file used (genome version and capture kit, see Section 2.4)? If unsure, ask the help desk of your sequencing center.
 - Were the normal samples run with the same assay and pipeline?
 - Did you provide `runAbsoluteCN` all the recommended files as described in Section 5?
 - For whole-genome data, you will get better results using a specialized third-party segmentation method as described in section 10.1, since our default is optimized for targeted sequencing. In general, you should probably start with tools optimized for WGS data, such as Battenberg [14], ABSOLUTE [2], ACEseq [7], or TitanCNA [8]. We are planning to incorporate proper support for WGS once high coverage diagnostic WGS becomes more common.

Copy number calling and SNV classification using targeted short read sequencing

- Otherwise, if log2-ratio peaks are clean as in Figure 3:
 - Was MuTect run without a matched normal? If yes, then make sure to provide either a pool of normal VCF or a SNP blacklist (if no pool of normal samples is available) as described in Sections 4.2 and 4.3.
 - A high fraction of sub-clonal copy-number alterations might also result in a low ranking of correct low ploidy solutions (see Section 11).

If the ploidy is frequently too low:

- *PureCN* with default parameters is conservative in calling genome duplications.
- This should only affect low purity samples (< 35%), since in higher purity samples the duplication signal is usually strong enough to reliably detect it.
- In whole-exome data, it is usually safe to decrease the `max.homozygous.loss` default, since such large losses are rare.

Will PureCN work with my data?

- *PureCN* was designed for medium-sized (>2-3Mb) targeted panels. The more data, the better, best results are typically achieved in whole-exome data.
- The number of heterozygous SNPs is also important (>1000 per sample). Copy number baits enriched in SNPs are therefore very helpful (see Section 9).
- Some users got acceptable results with small (<1Mb) panels. Try to find a perfect off-target bin width (`average.off.target.width` in `preprocessIntervals`) and maximize the number of heterozygous SNPs by including as much padding as possible. Keep in mind that without tiling baits, you will only have a poor resolution to detect LOH.
- Coverages below 80X are difficult unless purities are high and coverages are even.
- *PureCN* also needs process-matched normal samples, again, the more the better.
- Samples with tumor purities below 15-20% usually cannot be analyzed with this algorithm and *PureCN* might return very wrong purity estimates. In high coverage samples with low duplication rates, this limit can be close to 10%.
- Whole-genome data is not officially supported and specialized tools will likely provide better results. Third-party segmentation tools designed for this data type would be again required.
- Amplicon sequencing data is also not officially supported. If the assay contains tiling probes (at least with 1Mb spacing) and uses a barcode protocol that reduces PCR bias of measured allelic fractions, then this method might produce acceptable results. Setting the `model` argument of `runAbsoluteCN` to `betabin` is recommended. Specialized segmentation tools might be again better than our default.
- This tool was originally designed for single sample pipelines. It becomes more common to have multiple biopsies per patient and sharing information across biopsies might be beneficial. Other tools worth checking out are *lumosVar* 2.0 [15] and *SuperFreq* [16].

If you have trouble generating input data PureCN accepts, please check:

Copy number calling and SNV classification using targeted short read sequencing

- For problems related to generating valid coverage data, either consult the *GATK* manual for the *DepthOfCoverage* or *CollectFragmentCounts* tools or Section 2.3 for the equivalent function in *PureCN*. If you use *DepthOfCoverage* and off-target intervals as generated by *IntervalFile.R* (See Quick vignette), make sure to run it with parameters `-omitDepthOutputAtEachBase` and `-interval_merging OVERLAPPING_ONLY`.
- Currently only VCF files generated by *MuTect 1* are officially supported and well tested. A minimal example *MuTect* call would be:

```
$JAVA7 -Xmx6g -jar $MUTECT \  
--analysis_type MuTect -R $REFERENCE \  
--dbSNP $DBSNP_VCF \  
--cosmic $COSMIC_VCF \  
-I:normal $BAM_NORMAL \  
-I:tumor $BAM_TUMOR \  
-o $OUT/${ID}_bwa_mutect_stats.txt \  
-vcf $OUT/${ID}_bwa_mutect.vcf
```

The normal needs to be matched; without matched normal, only provide the tumor BAM file (do NOT provide a process-matched normal here). The default output file is the stats or call-stats file; this can be provided in addition to the required VCF file via `args.filterVcf` in *runAbsoluteCN*. If provided, it may help *PureCN* filter artifacts. This requires *MuTect* in version 1.1.7. This version is currently available at <https://software.broadinstitute.org/gatk/download/mutect> and requires Java 1.7 (it does not work with Java 1.8).

Note that this *MuTect* VCF will contain variants in off-target reads. By default, *PureCN* will remove variants outside the provided targets and their 50bp flanking regions. We highly recommend finding good values for each assay. A good cutoff will maximize the number of heterozygous SNPs and keep only an acceptable number of lower quality calls. This cutoff is set via `interval.padding` in `args.filterVcf`. See Section 9.

- Support for *GATK4* and *Mutect2* is in beta. Version 4.1.7.0 and higher as well as specifying the `-genotype-germline-sites` flag in *Mutect2* is required. Please follow the Broad Best Practices as closely as possible as we only test those workflows.
- For VCFs generated by other callers, the required dbSNP annotation can be added for example with *bcftools*:

```
bcftools annotate --annotation $DBSNP_VCF \  
--columns ID --output $OUT/${ID}_bwa_dbsnp.vcf.gz --output-type z \  
$OUT/${ID}_bwa.vcf.gz
```

- To generate a mappability file with the *GEM* library:
 - Calculate mappability, set kmer size to length of mapped reads.

```
REFERENCE="hg38.fa"  
PREFIX=`basename $REFERENCE .fa`  
THREADS=4  
KMER=100  
gem-indexer -T ${THREADS} -c dna -i ${REFERENCE} -o ${PREFIX}_index  
gem-mappability -T ${THREADS} -I ${PREFIX}_index.gem -l ${KMER} \  
-o ${PREFIX}_${KMER} -m 2 -e 2  
gem-2-wig -I ${PREFIX}_index.gem -i ${PREFIX}_${KMER}.mappability \  
-
```

Copy number calling and SNV classification using targeted short read sequencing

```
-o ${PREF}_${KMER}
```

- Convert to bigWig format, for example using the UCSC *wigToBigWig* tool:

```
cut -f1,2 ${REFERENCE}.fai > ${PREF}.sizes
wigToBigWig ${PREF}_${KMER}.wig ${PREF}.sizes ${PREF}_${KMER}.bw
```

- To generate the normal panel VCF (`normal.panel.vcf.file`) with *GATK3*:

- Run *MuTect* on the normal with `-I:tumor $BAM_NORMAL` and optionally with the `-artifact_detection_mode` flag.

- Remove the empty `none` sample from the VCF:

```
$JAVA -Xmx6g -jar $GATK3 \
--analysis_type SelectVariants -R $REFERENCE \
--exclude_sample_expressions none \
-V $OUT/${ID}_bwa_mutect_artifact_detection_mode.vcf \
-o $OUT/${ID}_bwa_mutect_artifact_detection_mode_no_none.vcf
```

- Merge the VCFs:

```
CMD="java -Xmx12g -jar $GATK3 -T CombineVariants --minimumN 3 -R $REFERENCE"
for VCF in $OUT/*bwa_mutect_artifact_detection_mode_no_none.vcf;
do
    CMD="$CMD --variant $VCF"
done
CMD="$CMD -o $OUT/normals_merged_min3.vcf"
echo $CMD > $OUT/merge_normals_min3.sh
bgzip $OUT/normals_merged_min3.vcf
tabix $OUT/normals_merged_min3.vcf.gz
```

Questions related to the output.

- *Where is the major and minor segment copy number?* The minor copy number is `ML.M.SEGMENT` in `predictSomatic` (Table 1) and `M` in `callAlterations` (Table 2) and `callLOH` (Table 4). The major copy number is simply the total copy number (column `C` or `ML.C` in `predictSomatic`) minus the minor copy number.
- *Why are there two different amplification output files?* See Section 6.2.2. Gene-level p-values are not very useful for calling high-level amplifications in tissue samples with significant amount of tumor cells because even regions of single gains will have significant read pile-ups. These p-values can be however useful in low purity samples commonly observed in cfDNA samples. There is no corresponding file for homozygous deletions since detecting those is usually impossible in very low purity samples.

Questions related to pool of normals. As described in detail in Sections 4.1 and 4.2, a pool of normal samples is used in *PureCN* for coverage normalization (to adjust for target-specific capture biases) and for artifact filtering. A few recommendations:

- Matched normals are obviously perfect for identifying most alignment artifacts and mapping biases. While not critical, we still recommend generating a `normal.panel.vcf.file` for *MuTect* and `calculateMappingBiasVcf` using the available normals.

Copy number calling and SNV classification using targeted short read sequencing

- Without matched normals, we need process-matched normals for coverage normalization. We recommend more than 5 from 2-3 different library preparation and sequencing batches. The more, the better.
- If you are able to add new normal samples over time capturing potential subtle changes in the lab, do so!
- These normals used for coverage normalization should be ideally sequenced to at least half the coverage as the tumor samples. This is completely different from matched normals for germline calling where 30-40X provide sufficient power to detect heterozygosity.
- For artifact removal, the more normals available, the more rare artifacts are removed. We recommend at least 10, 50 or more are ideal. The more artifacts are removed, the less likely *PureCN* output requires manual curation (Section 4.2).
- For position-specific mapping bias correction, the more normals are available, the more rare SNPs will have reliable mapping bias estimates. This requires again at least about 10 normals to be useful, 100 or more are ideal.
- With smaller pool of normals, we additionally recommend filtering SNPs from low quality regions (Section 4.3).
- It is worth trying the beta-binomial function instead of the default in the `model` argument of `runAbsoluteCN`. This will incorporate uncertainty of observed variant allelic fractions in the variant fitting step. The more normals, the more SNPs will have reliable estimates of allelic fraction overdispersion.
- *Do I really need a pool of normals? I only have tumor samples.* Unfortunately, yes. If you used a commercial capture kit, you might be able to obtain control samples from the vendor or the public domain. This is not optimal, but usually better than nothing.
- It is safe to include multiple normals from the same individual. Fewer common germline CNVs in `createNormalDatabase` and fewer SNPs in `calculateMappingBiasVcf` will be detected. But especially when the controls were sequenced in multiple batches, these replicates will still provide useful information for coverage normalization.

Questions related to manual curation. *PureCN*, like most other related tools, essentially finds the most simple explanation of the data. There are three major problems with this approach:

- First, hybrid capture data can be noisy and the algorithm must distinguish signal from noise; if the algorithm mistakes noise for signal, then this often results in wrong high ploidy calls (see Sections 4.2 and 4.3). If all steps in this vignette were followed, then *PureCN* should ignore common artifacts. Noisy samples thus often have outlier ploidy values and are often automatically flagged by *PureCN*. The correct solution is in most of these cases ranked second or third.
- The second problem is that signal can be sparse, i.e. when the tumor purity is very low or when there are only few somatic events. Manual curation is often easy in the latter case. For example when small losses are called as homozygous, but corresponding germline allele-frequencies are unbalanced (a complete loss would result in balanced germline allele frequencies, since only normal DNA is left). Future versions might improve calling in these cases by underweighting uninformative genomic regions.

Copy number calling and SNV classification using targeted short read sequencing

- The third problem is that tumor evolution is fast and complex and very difficult to incorporate into general likelihood models. Sometimes multiple solutions explain the data equally well, but one solution is then often clearly more consistent with known biology, for example LOH in tumor suppressor genes such as *TP53*. A basic understanding of both the algorithm and the tumor biology of the particular cancer type are thus important for curation. Fortunately, in most cancer types, such ambiguity is rather rare. See also Section 11.

Questions related to matched normals.

Coverage normalization: Even when matched normals are available, we recommend building a normal database for coverage normalization. This usually produces cleaner coverage profiles than the matched normal ([9]).

VCFs: When matched normals are available, simply provide this information to the variant caller and make sure that germline SNPs are not filtered out. *PureCN* should automatically find the matched information.

Questions related to off-target reads. Hybrid capture assays usually provide a significant amount of reads mapping outside the captured regions. This means these assays provide shallow whole genome sequencing data for free. *PureCN* makes it straightforward to use that information, but it might not always make sense to do so:

- Does the assay contain a fairly large number of big gaps (>1Mb) between neighboring baits? For whole exome data or targeted panels with copy number tiling probes, the answer can be no.
- Off-target reads are usually less affected by library-specific biases (Section 3) than on-target reads. In noisy data, off-target reads can thus provide relatively clean additional information.
- Very efficient assays with low off-target mapping might not provide a large number of reads in useful window sizes (usually below 0.2Mb), especially when total sequencing coverages are rather low. This can result in off-target data that are significantly noisier than on-target data.
- It is recommended to check the *runAbsoluteCN* output for the standard deviations of log2 copy number ratios for both on- and off-target reads. If the off-target noise is consistently larger than the on-target noise, you can try increasing the off-target bin width - and vice versa make it smaller when the on-target noise is larger.

Questions related to GC normalization.

- Explicit GC-normalization is useful whenever a significant difference in tumor and normal controls is suspected. One example would be FFPE tumors normalized with blood controls. Otherwise the normalization using the panel of normals should automatically remove any GC bias.
- Smaller panels of normals might also benefit from explicit GC-normalization.
- Otherwise it is safe skip the GC-normalization. In the command line tools, set `-skip-gc-norm` in *Coverage.R*.

If all or most of the samples are flagged as:

Copy number calling and SNV classification using targeted short read sequencing

Noisy segmentation: The default of 300 for `max.segments` is calibrated for high quality and high coverage whole-exome data. For whole-genome data or lower coverage data, this value needs to be re-calibrated. In case the copy number data looks indeed noisy, please see the first FAQ. Please be aware that *PureCN* will apply more aggressive segmentation parameters when the number of segments exceeds this cutoff. If the high segment count is real, this might confound downstream analyses.

High AT/GC dropout: If the data are GC-normalized, then there might be issues with either the target intervals or the provided GC content. Please double check that all files are correct and that all the coverage files are GC-normalized (Section 3).

Sex mismatch of coverage and VCF: If the panel contains baits for chromosome Y, then the interval file was probably generated without mappability file (Section 2.2). Similarly when third-party tools were used for coverage normalization and segmentation, this usually means probes on chromosome Y were not filtered for mappability. Cross-sample contamination (Section 10.6) can also cause sex mismatches.

References

- [1] Markus Riester, Angad P. Singh, A. Rose Brannon, Kun Yu, Catarina D. Campbell, Derek Y. Chiang, and Michael P. Morrissey. *PureCN: copy number calling and SNV classification using targeted short read sequencing. Source code for biology and medicine*, 11:13, Dec 2016.
- [2] Scott L. Carter, Kristian Cibulskis, Elena Helman, Aaron McKenna, Hui Shen, Travis Zack, Peter W. Laird, Robert C. Onofrio, Wendy Winckler, Barbara A. Weir, Rameen Beroukhi, David Pellman, Douglas A. Levine, Eric S. Lander, Matthew Meyerson, and Gad Getz. Absolute quantification of somatic DNA alterations in human cancer. *Nature biotechnology*, 30(5):413–421, May 2012.
- [3] Aaron McKenna, Matthew Hanna, Eric Banks, Andrey Sivachenko, Kristian Cibulskis, Andrew Kernysky, Kiran Garimella, David Altshuler, Stacey Gabriel, Mark Daly, and Mark A. DePristo. The Genome Analysis Toolkit: a MapReduce framework for analyzing next-generation DNA sequencing data. *Genome research*, 20(9):1297–1303, Sep 2010.
- [4] Kristian Cibulskis, Michael S. Lawrence, Scott L. Carter, Andrey Sivachenko, David Jaffe, Carrie Sougnez, Stacey Gabriel, Matthew Meyerson, Eric S. Lander, and Gad Getz. Sensitive detection of somatic point mutations in impure and heterogeneous cancer samples. *Nature biotechnology*, 31(3):213–219, Mar 2013.
- [5] Eric Talevich, A. Hunter Shain, Thomas Botton, and Boris C. Bastian. CNVkit: Genome-wide copy number detection and visualization from targeted DNA sequencing. *PLoS computational biology*, 12(4):e1004873, Apr 2016.
- [6] Amnon Koren, Robert E. Handsaker, Nolan Kamitaki, Rosa Karlić, Sulagna Ghosh, Paz Polak, Kevin Eggan, and Steven A. McCarroll. Genetic variation in human dna replication timing. *Cell*, 159(5):1015–1026, Nov 2014.
- [7] Kortine Kleinheinz, Isabell Bludau, Daniel Huebschmann, Michael Heinold, Philip Kensche, Zuguang Gu, Cristina Lopez, Michael Hummel, Wolfram Klapper, Peter Moeller, Inga Vater, Rabea Wagener, Benedikt Brors, Reiner Siebert, Roland Eils, and

- Matthias Schlesner. ACEseq - allele specific copy number estimation from whole genome sequencing. *bioRxiv*, 2017. URL: <https://www.biorxiv.org/content/early/2017/10/29/210807>, doi:10.1101/210807.
- [8] Gavin Ha, Andrew Roth, Jaswinder Khattra, Julie Ho, Damian Yap, Leah M. Prentice, Nataliya Melnyk, Andrew McPherson, Ali Bashashati, Emma Laks, Justina Biele, Jiarui Ding, Alan Le, Jamie Rosner, Karey Shumansky, Marco A. Marra, C. Blake Gilks, David G. Huntsman, Jessica N. McAlpine, Samuel Aparicio, and Sohrab P. Shah. TITAN: inference of copy number architectures in clonal cell populations from tumor whole-genome sequence data. *Genome research*, 24(11):1881–1893, Nov 2014.
- [9] Barbara Tabak, Gordon Saksena, Coyin Oh, Galen F. Gao, Barbara Hill Meyers, Michael Reich, Steven E. Schumacher, Lindsay C. Westlake, Ashton C. Berger, Scott L.b Carter, Andrew D. Cherniack, Matthew Meyerson, Rameen Beroukhim, and Gad Getz. The tangent copy-number inference pipeline for cancer genome analyses. *bioRxiv*, 2019. URL: <https://www.biorxiv.org/content/early/2019/03/05/566505>, doi:10.1101/566505.
- [10] Nicholas McGranahan, Francesco Favero, Elza C. de Bruin, Nicolai Juul Birkbak, Zoltan Szallasi, and Charles Swanton. Clonal status of actionable driver events and the timing of mutational processes in cancer evolution. *Science translational medicine*, 7(283):283ra54, Apr 2015.
- [11] P. Andrew Futreal, Lachlan Coin, Mhairi Marshall, Thomas Down, Timothy Hubbard, Richard Wooster, Nazneen Rahman, and Michael R. Stratton. A census of human cancer genes. *Nature reviews. Cancer*, 4(3):177–183, Mar 2004.
- [12] E. S. Venkatraman and Adam B. Olshen. A faster circular binary segmentation algorithm for the analysis of array CGH data. *Bioinformatics (Oxford, England)*, 23(6):657–663, Mar 2007.
- [13] Simon A. Forbes, David Beare, Prasad Gunasekaran, Kenric Leung, Nidhi Bindal, Harry Boutselakis, Minjie Ding, Sally Bamford, Charlotte Cole, Sari Ward, Chai Yin Kok, Mingming Jia, Tisham De, Jon W. Teague, Michael R. Stratton, Ultan McDermott, and Peter J. Campbell. COSMIC: exploring the world’s knowledge of somatic mutations in human cancer. *Nucleic acids research*, 43(Database issue):D805–D811, Jan 2015.
- [14] Serena Nik-Zainal, Peter Van Loo, David C. Wedge, Ludmil B. Alexandrov, Christopher D. Greenman, King Wai Lau, Keiran Raine, David Jones, John Marshall, Manasa Ramakrishna, Adam Shlien, Susanna L. Cooke, Jonathan Hinton, Andrew Menzies, Lucy A. Stebbings, Catherine Leroy, Mingming Jia, Richard Rance, Laura J. Mudie, Stephen J. Gamble, Philip J. Stephens, Stuart McLaren, Patrick S. Tarpey, Elli Papaemmanuil, Helen R. Davies, Ignacio Varela, David J. McBride, Graham R. Bignell, Kenric Leung, Adam P. Butler, Jon W. Teague, Sancha Martin, Goran Jönsson, Odette Mariani, Sandrine Boyault, Penelope Miron, Aquila Fatima, Anita Langerød, Samuel A. J. R. Aparicio, Andrew Tutt, Anieta M. Sieuwerts, Åke Borg, Gilles Thomas, Anne Vincent Salomon, Andrea L. Richardson, Anne-Lise Børresen-Dale, P. Andrew Futreal, Michael R. Stratton, Peter J. Campbell, and Breast Cancer Working Group of the International Cancer Genome Consortium. The life history of 21 breast cancers. *Cell*, 149(5):994–1007, May 2012. doi:10.1016/j.cell.2012.04.023.
- [15] R. F. Halperin, W. S. Liang, S. Kulkarni, E. E. Tassone, J. Adkins, D. Enriquez, N. L. Tran, N. C. Hank, J. Newell, C. Kodira, R. Korn, M. E. Berens, S. Kim, and S. A. Byron. Leveraging Spatial Variation in Tumor Purity for Improved Somatic Variant Calling of Archival Tumor Only Samples. *Front Oncol*, 9:119, 2019.

- [16] C. Flensburg, T. Sargeant, A. Oshlack, and I. J. Majewski. SuperFreq: Integrated mutation detection and clonal tracking in cancer. *PLoS Comput. Biol.*, 16(2):e1007603, 02 2020.

Session Info

- R version 4.5.1 Patched (2025-08-23 r88802), x86_64-pc-linux-gnu
- Locale: LC_CTYPE=en_US.UTF-8, LC_NUMERIC=C, LC_TIME=en_GB, LC_COLLATE=C, LC_MONETARY=en_US.UTF-8, LC_MESSAGES=en_US.UTF-8, LC_PAPER=en_US.UTF-8, LC_NAME=C, LC_ADDRESS=C, LC_TELEPHONE=C, LC_MEASUREMENT=en_US.UTF-8, LC_IDENTIFICATION=C
- Time zone: America/New_York
- TZcode source: system (glibc)
- Running under: Ubuntu 24.04.3 LTS
- Matrix products: default
- BLAS: /home/biocbuild/bbs-3.22-bioc/R/lib/libRblas.so
- LAPACK: /usr/lib/x86_64-linux-gnu/lapack/liblapack.so.3.12.0
- Base packages: base, datasets, grDevices, graphics, methods, stats, stats4, utils
- Other packages: Biobase 2.69.1, BiocGenerics 0.55.1, BiocStyle 2.37.1, Biostrings 2.77.2, DNACopy 1.83.0, GenomicRanges 1.61.5, IRanges 2.43.5, MatrixGenerics 1.21.0, PureCN 2.15.4, Rsamtools 2.25.3, S4Vectors 0.47.4, Seqinfo 0.99.2, SummarizedExperiment 1.39.2, VariantAnnotation 1.55.1, XVector 0.49.1, generics 0.1.4, matrixStats 1.5.0
- Loaded via a namespace (and not attached): AnnotationDbi 1.71.1, BSgenome 1.77.2, BiocIO 1.19.0, BiocManager 1.30.26, BiocParallel 1.43.4, DBI 1.2.3, DelayedArray 0.35.3, GenomeInfoDb 1.45.12, GenomicAlignments 1.45.4, GenomicFeatures 1.61.6, KEGGREST 1.49.1, Matrix 1.7-4, R.methodsS3 1.8.2, R.oo 1.27.1, R.utils 2.13.0, R6 2.6.1, RColorBrewer 1.1-3, RCurl 1.98-1.17, RSQLite 2.4.3, Rhdf5lib 1.31.0, S4Arrays 1.9.1, S7 0.2.0, SparseArray 1.9.1, UCSC.utils 1.5.0, VGAM 1.1-13, XML 3.99-0.19, abind 1.4-8, bit 4.6.0, bit64 4.6.0-1, bitops 1.0-9, blob 1.2.4, bookdown 0.45, bslib 0.9.0, cachem 1.1.0, cli 3.6.5, codetools 0.2-20, compiler 4.5.1, crayon 1.5.3, curl 7.0.0, data.table 1.17.8, dichromat 2.0-0.1, digest 0.6.37, dplyr 1.1.4, evaluate 1.0.5, farver 2.1.2, fastmap 1.2.0, formatR 1.14, futile.logger 1.4.3, futile.options 1.0.1, ggplot2 4.0.0, glue 1.8.0, grid 4.5.1, gridExtra 2.3, gtable 0.3.6, highr 0.11, htmltools 0.5.8.1, httr 1.4.7, jquerylib 0.1.4, jsonlite 2.0.0, knitr 1.50, labeling 0.4.3, lambda.r 1.2.4, lattice 0.22-7, lifecycle 1.0.4, magrittr 2.0.4, mclust 6.1.1, memoise 2.0.1, parallel 4.5.1, pillar 1.11.1, pkgconfig 2.0.3, png 0.1-8, restfulr 0.0.16, rhdf5 2.53.5, rhdf5filters 1.21.0, rjson 0.2.23, rlang 1.1.6, rmarkdown 2.30, rtracklayer 1.69.1, sass 0.4.10, scales 1.4.0, splines 4.5.1, tibble 3.3.0, tidyselect 1.2.1, tinytex 0.57, tools 4.5.1, vctrs 0.6.5, withr 3.0.2, xfun 0.53, yaml 2.3.10