

Identifying differential distributions in single-cell RNA-seq experiments with scDD

*Keegan Korthauer**

*keegan@jimmy.harvard.edu

October 7, 2025

Abstract

The *scDD* package models single-cell gene expression data (from single-cell RNA-seq) using flexible nonparametric Bayesian mixture models in order to explicitly handle heterogeneity within cell populations. In bulk RNA-seq data, where each measurement is an average over thousands of cells, distributions of expression over samples are most often unimodal. In single-cell RNA-seq data, however, even when cells represent genetically homogeneous populations, multimodal distributions of gene expression values over samples are common [1]. This type of heterogeneity is often treated as a nuisance factor in studies of differential expression in single-cell RNA-seq experiments. Here, we explicitly accommodate it in order to improve power to detect differences in expression distributions that are more complicated than a mean shift.

Package

scDD 1.33.0

Report issues on <https://github.com/kdkorthauer/scDD/issues>

Contents

1	Background	3
2	Identify and Classify DD genes	4
3	Alternate test for Differential Distributions	6
4	Simulation	7
5	Formatting and Preprocessing	10
	5.1 Constructing a SingleCellExperiment object	11
	5.2 Filtering and Normalization	12
6	Plotting	13
7	Session Info	15

1 Background

Our aim is two-fold: (1) to detect which genes have different expression distributions across two biological conditions and (2) to classify those differences into informative patterns. Note that in (1) we explicitly say differences in 'distributions' rather than differences in 'average', which would correspond to traditional DE (differential expression) analysis in bulk RNA-seq. By examining the entire distribution, we are able to detect more subtle differences as well as describe complex patterns, such as the existence of subgroups of cells within and across condition that express a given gene at a different level.

We start by assuming that the log-transformed nonzero expression values arise out of a Dirichlet Process Mixture of normals model. This allows us to characterize expression distributions in terms of the number of modes (or clusters). To detect differences in these distributions across conditions, an approximate Bayes Factor score is used which compares the conditional likelihood under the hypothesis of Equivalent distributions (ED) where one clustering process governs both conditions jointly, with the hypothesis of Differential distributions (DD) where each condition is generated from its own clustering process. In the full framework, significance of the scores for each gene are evaluated via empirical p-values after permutation. Optionally, a fast implementation obtains the p-values from the non-parametric Kolmogorov-Smirnov test. Zero values are considered by also implementing a χ^2 test of whether the proportion of zero values differs by condition (adjusted for overall sample detection rate). More details are provided in [1].

After the detection step is carried out, the significantly DD genes are classified into four informative patterns based on the number of clusters detected and whether they overlap. These patterns, depicted in Figure 1, include (a) DE (differential expression of unimodal genes), (b) DP (differential proportion for multimodal genes), (c) DM (differential modality), and (d) DB (both differential modality and different component means). Genes where a differential proportion of zeroes were identified are classified as DZ (differential zero). Genes that are identified as significantly differentially distributed but do not fall into one of the above categories are abbreviated NC (for no call). This includes genes with the same number of components with similar component means, but differential variance. For reasons detailed in [1], we do not aim to interpret this type of pattern.

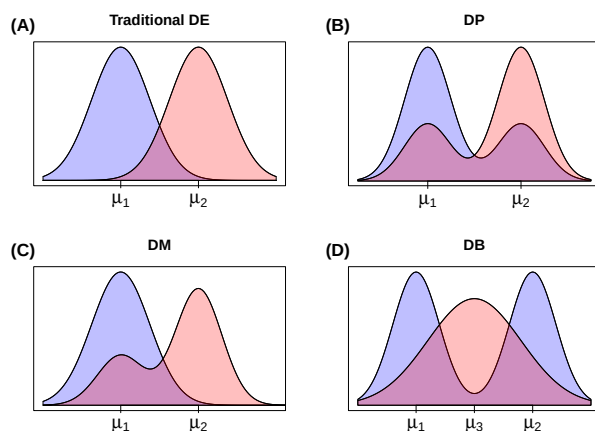


Figure 1: Illustration of informative DD patterns

The rest of this vignette outlines the main functionality of the *scDD* package. This includes:

- Identifying genes that are expressed differently between two biological conditions and classifying them into informative patterns.
- Simulating single-cell RNA-seq data with differential expression that exhibits multi-modal patterns.
- Preprocessing and formatting of single-cell RNA-seq data to facilitate analysis
- Visualizing the expression patterns using a violin plotting scheme

2 Identify and Classify DD genes

In this section, we demonstrate how to use the main function `scDD` to find genes with differential distributions and classify them into the patterns of interest described in the previous section.

First, we need to load the `scDD` package. For each of the following sections in this vignette, we assume this step has been carried out.

```
library(scDD)
```

Next, we load the toy simulated example *SingleCellExperiment* object that we will use for identifying and classifying DD genes.

```
data(scDatExSim)
```

Verify that this object is a member of the *SingleCellExperiment* class and that it contains 200 samples and 30 genes. The `colData` slot (which contains a dataframe of metadata for the cells) should have a column that contains the biological condition or grouping of interest. In this example data, that variable is the 'condition' variable. Note that the input gene set needs to be in *SingleCellExperiment* format, and should contain normalized counts. In practice, it is also advisable to filter the input gene set to remove genes that have an extremely high proportion of zeroes (see Section 6). More specifically, the test for differential distributions of the expressed measurements will not be carried out on genes where only one or fewer cells had a nonzero measurement (these genes will still be tested for differential proportion of zeroes (DZ) if the `testZeroes` parameter is set to `TRUE`, however).

```
class(scDatExSim)
## [1] "SingleCellExperiment"
## attr(,"package")
## [1] "SingleCellExperiment"
dim(scDatExSim)
## [1] 30 200
```

Next, specify the hyperparameter arguments that we'll pass to the `scDD` function. These values reflect heavy-tailed distributions over the parameters and are robust to many different settings in simulation (see [1] for more details).

```
prior_param=list(alpha=0.01, mu0=0, s0=0.01, a0=0.01, b0=0.01)
```

Finally, call the `scDD` function to test for differential distributions, classify DD genes, and return the results. If the biological condition or grouping variable in the `colData` slot is named something other than 'condition', you'll need to specify the name of the variable as

an argument to the `scDD` function (set the `condition` argument equal to the name of the relevant column). We won't perform the test for a difference in the proportion of zeroes since none exists in this simulated toy example data, but this option can be invoked by changing the `testZeroes` option to `TRUE`. Note that the default option is to use a fast test of differential distributions that involves the Kolmogorov-Smirnov test instead of the full permutation testing framework. This provides a fast implementation of the method at the cost of potentially slightly decreased power compared to the full scDD framework described in the manuscript (see Section 4 for more details).

Note that if you are only interested in obtaining the results of the test for significance, and not in the classification of genes to the patterns mentioned above, you can achieve further computational speedup by setting `categorize` to `FALSE`.

```
scDatExSim <- scDD(scDatExSim, prior_param=prior_param, testZeroes=FALSE)
## Setting up parallel back-end using 4 cores
## Clustering observed expression data for each gene
## Notice: Number of permutations is set to zero; using
##           Kolmogorov-Smirnov to test for differences in distributions
##           instead of the Bayes Factor permutation test
## Classifying significant genes into patterns
```

Four results objects are added to the `scDatExSim` `SingleCellExperiment` object in the `metadata` slot. For convenience, the results objects can be extracted with the `results` function.

The main results object is the "Genes" object which is a `data.frame` containing the following columns:

- `gene`: gene name (matches rownames of `SCdat`)
- `DDcategory`: name of the DD pattern (DE, DP, DM, DB, DZ), or NC (no call), or NS (not significant).
- `Clusters.combined`: the number of clusters identified when pooling condition 1 and 2 together
- `Clusters.c1`: the number of clusters identified in condition 1 alone
- `Clusters.c2`: the number of clusters identified in condition 2 alone
- `nonzero.pvalue`: p-value for KS test of differential distributions of expressed cells
- `nonzero.pvalue.adj`: Benjamini-Hochberg adjusted p-value for KS test of differential distributions
- `zero.pvalue`: p-value for test of difference in dropout rate (only if `testZeroes==TRUE`)
- `zero.pvalue.adj`: Benjamini-Hochberg adjusted p-value for test of difference in dropout rate (only if `testZeroes==TRUE`)
- `combined.pvalue`: Fisher's combined p-value for a difference in nonzero or zero values (only if `testZeroes==TRUE`)
- `combined.pvalue.adj`: Benjamini-Hochberg adjusted Fisher's combined p-value for a difference in nonzero or zero values (only if `testZeroes==TRUE`)

This can be extracted using the following call to `results`:

```
RES <- results(scDatExSim)
head(RES)

##      gene DDcategory Clusters.combined Clusters.c1 Clusters.c2 nonzero.pvalue
## DE1  DE1          DB                1            2            1  4.215186e-07
## DE2  DE2          DE                1            1            1  1.663921e-08
## DE3  DE3          DE                1            1            1  6.390733e-20
## DE4  DE4          DE                1            1            1  1.656735e-20
## DE5  DE5          DE                1            1            1  1.274147e-07
## DP6  DP6          NC                2            2            2  8.445703e-06
##      nonzero.pvalue.adj
## DE1      2.529112e-06
## DE2      1.663921e-07
## DE3      9.586100e-19
## DE4      4.970206e-19
## DE5      9.556106e-07
## DP6      3.167138e-05
```

The remaining three results objects are matrices (first for condition 1 and 2 combined, then condition 1 alone, then condition 2 alone) that contain the cluster memberships (partition estimates) for each sample (for clusters 1,2,3,...) in columns and genes in rows. Zeroes, which are not involved in the clustering, are labeled as zero. These can be extracted by specifying an alternative `type` when calling the `results` function. For example, we can extract the partition estimates for condition 1 with the following:

```
PARTITION.C1 <- results(scDatExSim, type="Zhat.c1")
PARTITION.C1[1:5,1:5]

##      Sample1 Sample2 Sample3 Sample4 Sample5
## DE1        0        0        1        2        2
## DE2        1        1        1        1        0
## DE3        1        1        1        1        1
## DE4        1        1        1        1        1
## DE5        1        1        0        1        0
```

3 Alternate test for Differential Distributions

The first step in the scDD framework that identifies Differential Distributions was designed to have optimal power to detect differences in expression distributions, but the utilization of a permutation test on the Bayes Factor can be computationally demanding. While this is not an issue when machines with multiple cores are available since the code takes advantage of parallel processing, we also provide the option to use an alternate test to detect distributional differences that avoids the use of a permutation test. This option (default) uses the Kolmogorov-Smirnov test, which examines the null hypothesis that two samples are generated from the same continuous distribution. While the use of this test yielded slightly lower power in simulations than the full permutation testing framework at lower sample sizes (50-75 cells in each condition) and primarily affected the DB pattern genes, it does not require permutations and thus is orders of magnitude faster. The overall power to detect DD genes in simulation was still comparable or favorable to existing methods for differential expression analysis of scRNA-seq experiments.

The remaining steps of the scDD framework remain unchanged if the alternate test is used. That is, the Dirichlet process mixture model is still fit to the observed expression measurements so that the significant DD genes can be categorized into patterns that represent the major distributional changes, and results can still be visualized with violin plots using the `sideViolin` function described in the Plotting section.

The option to use the full permutation testing procedure instead of the Kolmogorov-Smirnov test is invoked by setting the number of permutations to something other than zero (the `permutations` argument in `scDD`) when calling the main `scDD` function as follows:

```
scDatExSim <- scDD(scDatExSim, prior_param=prior_param,
                  testZeroes=FALSE, permutations=100)

## Setting up parallel back-end using 4 cores
## Clustering observed expression data for each gene
## Performing permutations to evaluate independence of clustering
##               and condition for each gene
## Parallelizing by Genes
## Classifying significant genes into patterns
```

The line above will run 100 permutations of every gene. In practice, it is recommended that at least 1000 permutations are carried out if using the full permutation testing option. Note that this option will take significantly longer than the default option to use the alternate KS test, and computation time will increase with more genes and/or more permutations, but multiple cores will automatically be utilized (if available) via the *BiocParallel* package. By default, an OS appropriate back-end using the number of cores on the machine minus 2 is chosen automatically. Alternatively, you can specify the number of cores to use by passing in a `param` argument in the `scDD` function call (where the `param` argument is an object of class `MulticoreParam` for Linux-like OS or `SnowParam` for Windows). For example, to use 12 cores on a Linux-like OS, specify `param=MulticoreParam(workers=12)`.

The results returned by `scDD` remain exactly as described in the previous section, with the exception that the `nonzero.pvalue` and `nonzero.pvalue.adj` columns of the `Genes` data frame now contain the p-values and Benjamini-Hochberg adjusted p-values of the permutation test of the Bayes Factor for independence of condition membership with clustering.

4 Simulation

Here we show how to generate a simulated single-cell RNA-seq dataset which contains multi-modal genes. The `simulateSet` function simulates data from a two-condition experiment with a specified number of genes that fall into each of the patterns of interest. For DD genes, these include DE (differential expression of unimodal genes), DP (differential proportion for multimodal genes), DM (differential modality), and DB (both differential modality and mean expression levels), and for ED genes these include EE (equivalent expression for unimodal genes) and EP (equivalent proportion for multimodal genes). The simulation parameters are based on observed data from two conditions, so the function requires an *SingleCellExperiment* formatted dataset as input. The output of the function is also a *SingleCellExperiment* object with information about the true category of each gene and its simulated fold change stored in the `rowData` slot.

First, we load the toy example *SingleCellExperiment* to simulate from

```
data(scDatEx)
```

We'll verify that this object is a member of the *SingleCellExperiment* class and that it contains 142 samples and 500 genes

```
class(scDatEx)
## [1] "SingleCellExperiment"
## attr(,"package")
## [1] "SingleCellExperiment"

dim(scDatEx)
## [1] 500 142
```

Next we need to set the arguments that will be passed to the `simulateSet` function. In this example we will simulate 30 genes total, with 5 genes of each type and 100 samples in each of two conditions. We also set a random seed for reproducibility.

```
nDE <- 5
nDP <- 5
nDM <- 5
nDB <- 5
nEE <- 5
nEP <- 5
numSamples <- 100
seed <- 816
```

Finally, we'll create the simulated set with specified numbers of DE, DP, DM, DM, EE, and EP genes and specified number of samples, where DE gene fold changes represent 2 standard deviations of the observed fold change distribution, and multimodal genes have cluster mean distance of 4 standard deviations.

```
SD <- simulateSet(scDatEx, numSamples=numSamples,
                  nDE=nDE, nDP=nDP, nDM=nDM, nDB=nDB,
                  nEE=nEE, nEP=nEP, sd.range=c(2,2), modeFC=4, plots=FALSE,
                  random.seed=seed)

## Setting up parallel back-end using 4 cores
## Identifying a set of genes to simulate from...
## Simulating DE fold changes...
## Simulating individual genes...
## Done! Simulated 5 DE, 5 DP, 5 DM, 5 DB, 5 EE, and 5 EP genes
# load the SingleCellExperiment package to use rowData method
library(SingleCellExperiment)

## Loading required package: SummarizedExperiment
## Loading required package: MatrixGenerics
## Loading required package: matrixStats
##
## Attaching package: 'MatrixGenerics'
```



```

## The following objects are masked from 'package:matrixStats':
##
##   colAlls, colAnyNAs, colAnys, colAvgPerRowSet, colCollapse,
##   colCounts, colCummaxs, colCummins, colCumprods, colCumsums,
##   colDiffs, colIQRDiffs, colIQRs, colLogSumExps, colMadDiffs,
##   colMads, colMaxs, colMeans2, colMedians, colMins, colOrderStats,
##   colProds, colQuantiles, colRanges, colRanks, colSdDiffs, colSds,
##   colSums2, colTabulates, colVarDiffs, colVars, colWeightedMads,
##   colWeightedMeans, colWeightedMedians, colWeightedSds,
##   colWeightedVars, rowAlls, rowAnyNAs, rowAnys, rowAvgPerColSet,
##   rowCollapse, rowCounts, rowCummaxs, rowCummins, rowCumprods,
##   rowCumsums, rowDiffs, rowIQRDiffs, rowIQRs, rowLogSumExps,
##   rowMadDiffs, rowMads, rowMaxs, rowMeans2, rowMedians, rowMins,
##   rowOrderStats, rowProds, rowQuantiles, rowRanges, rowRanks,
##   rowSdDiffs, rowSds, rowSums2, rowTabulates, rowVarDiffs,
##   rowVars, rowWeightedMads, rowWeightedMeans, rowWeightedMedians,
##   rowWeightedSds, rowWeightedVars

## Loading required package: GenomicRanges

## Loading required package: stats4

## Loading required package: BiocGenerics

## Loading required package: generics

##
## Attaching package: 'generics'

## The following objects are masked from 'package:base':
##
##   as.difftime, as.factor, as.ordered, intersect, is.element,
##   setdiff, setequal, union

##
## Attaching package: 'BiocGenerics'

## The following objects are masked from 'package:stats':
##
##   IQR, mad, sd, var, xtabs

## The following objects are masked from 'package:base':
##
##   Filter, Find, Map, Position, Reduce, anyDuplicated, aperm,
##   append, as.data.frame, basename, cbind, colnames, dirname,
##   do.call, duplicated, eval, evalq, get, grep, grepl, is.unsorted,
##   lapply, mapply, match, mget, order, paste, pmax, pmax.int, pmin,
##   pmin.int, rank, rbind, rownames, sapply, saveRDS, table, tapply,
##   unique, unsplit, which.max, which.min

## Loading required package: S4Vectors

##
## Attaching package: 'S4Vectors'

```

```
## The following object is masked from 'package:utils':
##
##   findMatches

## The following objects are masked from 'package:base':
##
##   I, expand.grid, unname

## Loading required package: IRanges
## Loading required package: Seqinfo
## Loading required package: Biobase

## Welcome to Bioconductor
##
##   Vignettes contain introductory material; view with
##   'browseVignettes()'. To cite Bioconductor, see
##   'citation("Biobase)", and for packages 'citation("pkgname)".

##
## Attaching package: 'Biobase'

## The following object is masked from 'package:MatrixGenerics':
##
##   rowMedians

## The following objects are masked from 'package:matrixStats':
##
##   anyMissing, rowMedians

head(rowData(SD))

## DataFrame with 6 rows and 2 columns
##      Category      FC
##   <character> <numeric>
## DE1         DE         1
## DE2         DE         2
## DE3         DE        -1
## DE4         DE        -1
## DE5         DE         1
## DP6         DP         4
```

The `normcounts` assay `SD` object (of class *SingleCellExperiment*) contains simulated expression values. The `rowData` slot stores the fold change/modal distance values and gene expression categories, which are useful in assessing performance of a differential expression method.

5 Formatting and Preprocessing

Before beginning an analysis using *scDD*, you will need to carry out a few preprocessing steps. This includes normalization, filtering of genes that are mostly zero, and getting the data into format that is expected by the `scDD` function. The following subsections will detail these steps.

5.1 Constructing a SingleCellExperiment object

In this subsection, we provide a quick example of how to construct an object of the *SingleCellExperiment* class. For more detailed instructions, refer to the *SingleCellExperiment* package documentation.

Here we will convert the toy example data object `scDatExList` into a *SingleCellExperiment* object. This object is a list of two matrices (one for each condition) of normalized counts. Each matrix has genes in rows and cells in columns, and is named "C1" or "C2" (for condition 1 or 2).

First, load the *SingleCellExperiment* package:

```
library(SingleCellExperiment)
```

Next, load the toy example data list:

```
data(scDatExList)
```

Next, create a vector of condition membership labels (these should be 1 or 2). In our example data list, we have 78 cells in condition 1, and 64 cells in condition 2.

```
condition <- c(rep(1, ncol(scDatExList$C1)), rep(2, ncol(scDatExList$C2)))
```

The rows and columns of the expression matrix should have unique names. This is already the case for our toy example dataset. The names of the columns should also correspond to the names of the condition membership labels in `condition`.

```
# Example of row and column names
head(rownames(scDatExList$C1))
## [1] "MKL2" "CD109" "ABTB1" "MAST2" "KAT5" "WWC2"

head(colnames(scDatExList$C2))
## [1] "C2.001" "C2.002" "C2.003" "C2.004" "C2.005" "C2.006"

names(condition) <- c(colnames(scDatExList$C1), colnames(scDatExList$C2))
```

Once our labeling is intact, we can call the `SingleCellExperiment` function and specify the two relevant pieces of information. The `normcounts` assays slot should contain one matrix, so we use `rbind` here to combine both conditions. Optionally, additional experiment information can be stored in additional slots; see the *SingleCellExperiment* package for more details.

```
sce <- SingleCellExperiment(assays=list(normcounts=cbind(scDatExList$C1,
                                                         scDatExList$C2)),
                           colData=data.frame(condition))

show(sce)

## class: SingleCellExperiment
## dim: 100 142
## metadata(0):
## assays(1): normcounts
## rownames(100): MKL2 CD109 ... ASB10 HBP1
## rowData names(0):
## colnames(142): C1.073 C1.074 ... C2.068 C2.070
## colData names(1): condition
```

```
## reducedDimNames(0):
## mainExpName: NULL
## altExpNames(0):
```

5.2 Filtering and Normalization

In this subsection, we demonstrate the utility of the `preprocess` function, which can be helpful if working with raw counts, or data which contains genes that are predominantly zero (common in single-cell RNA-seq experiments). This function takes as input a list of data matrices, one for each condition.

First, load the toy example data and verify it is a *SingleCellExperiment* object:

```
data(scDatEx)
show(scDatEx)

## class: SingleCellExperiment
## dim: 500 142
## metadata(0):
## assays(2): normcounts counts
## rownames(500): RHOF FAM161B ... BHLHB9 GLIPR2
## rowData names(0):
## colnames(142): C1.073 C1.074 ... C2.068 C2.070
## colData names(1): condition
## reducedDimNames(0):
## mainExpName: NULL
## altExpNames(0):
```

Now, apply the `preprocess` function with the `zero.thresh` argument set to 0.9 so that genes are filtered out if they are 90 (or more) percent zero.

```
scDatEx <- preprocess(scDatEx, zero.thresh=0.9)
show(scDatEx)

## class: SingleCellExperiment
## dim: 500 142
## metadata(0):
## assays(2): normcounts counts
## rownames(500): RHOF FAM161B ... BHLHB9 GLIPR2
## rowData names(0):
## colnames(142): C1.073 C1.074 ... C2.068 C2.070
## colData names(1): condition
## reducedDimNames(0):
## mainExpName: NULL
## altExpNames(0):
```

We can see that no genes were removed, since all have fewer than 90 percent of zeroes to begin with.

Now, apply the `preprocess` function again, but this time use a more stringent threshold on the proportion of zeroes and apply normalization using size factors calculated using the `scrn`. In this example, we set the `zero.thresh` argument to 0.50 so that genes with more than 50 percent zeroes are filtered out and we set the `scrn_norm` argument to `TRUE` to return `scrn` normalized counts.

```
scDatEx.scrn <- preprocess(scDatEx, zero.thresh=0.5, scrn_norm=TRUE)

## Warning in preprocess(scDatEx, zero.thresh = 0.5, scrn_norm = TRUE): median
## or scrn norm is specified and the 'normcounts' assay already exists; replacing
## 'normcounts' in output with the specified normalization method. Original contents
## of 'normcounts' are now in 'normcounts-orig'.

## Performing scrn Normalization

show(scDatEx.scrn)

## class: SingleCellExperiment
## dim: 462 142
## metadata(0):
## assays(3): normcounts counts normcounts-orig
## rownames(462): FAM161B EIF2AK4 ... BHLHB9 GLIPR2
## rowData names(0):
## colnames(142): C1.073 C1.074 ... C2.068 C2.070
## colData names(2): condition sizeFactor
## reducedDimNames(0):
## mainExpName: NULL
## altExpNames(0):
```

We can see that 38 genes were removed due to having more than 50 percent zeroes. The warning message here is letting us know that our `SingleCellExperiment` object already contained a `normcounts` assay, and that by specifying `scrn_norm=TRUE`, we are replacing that with the `scrn` normalized counts and moving the original values to a new slot called `normcounts-orig`.

Also included is the option to use median normalization, invoked by setting `median_norm` to `TRUE`.

6 Plotting

Next we demonstrate the plotting routine that is implemented in the `sideViolin` function. This function produces side-by-side violin plots (where the curves represent a smoothed kernel density estimate) of the log-transformed data. A count of 1 is added before log-transformation so that zeroes can be displayed, but they are not included in the density estimation. Each condition is represented by one violin plot. Individual data points are plotted (with jitter) on top.

We illustrate this function by displaying the six types of simulated genes using the toy example simulated dataset. First, load the toy simulated dataset:

```
data(scDatExSim)
```

Next, load the `SingleCellExperiment` package to facilitate subset operations on `SingleCellExperiment` class objects:

```
library(SingleCellExperiment)
```

The following lines will produce the figures in Figure 2.

Plot side by side violin plots for Gene 1 (DE):

```
de <- sideViolin(normcounts(scDatExSim)[1,], scDatExSim$condition,
  title.gene=rownames(scDatExSim)[1])
```

Plot side by side violin plots for Gene 6 (DP):

```
dp <- sideViolin(normcounts(scDatExSim)[6,], scDatExSim$condition,
  title.gene=rownames(scDatExSim)[6])
```

Plot side by side violin plots for Gene 11 (DM):

```
dm <- sideViolin(normcounts(scDatExSim)[11,], scDatExSim$condition,
  title.gene=rownames(scDatExSim)[11])
```

Plot side by side violin plots for Gene 16 (DB):

```
db <- sideViolin(normcounts(scDatExSim)[16,], scDatExSim$condition,
  title.gene=rownames(scDatExSim)[16])
```

Plot side by side violin plots for Gene 21 (EP):

```
ep <- sideViolin(normcounts(scDatExSim)[21,], scDatExSim$condition,
  title.gene=rownames(scDatExSim)[21])
```

Plot side by side violin plots for Gene 26 (EE):

```
ee <- sideViolin(normcounts(scDatExSim)[26,], scDatExSim$condition,
  title.gene=rownames(scDatExSim)[26])
```

The plot objects returned by `sideViolin` are standard *ggplot2* objects, and thus can be manipulated into multipanel figures with the help of the *gridExtra* or *cowplot* packages. Here we use `grid.arrange` from the *gridExtra* package to visualize all the plots generated above. The end result is shown in Figure 2.

```
library(gridExtra)
grid.arrange(de, dp, dm, db, ep, ee, ncol=2)
```

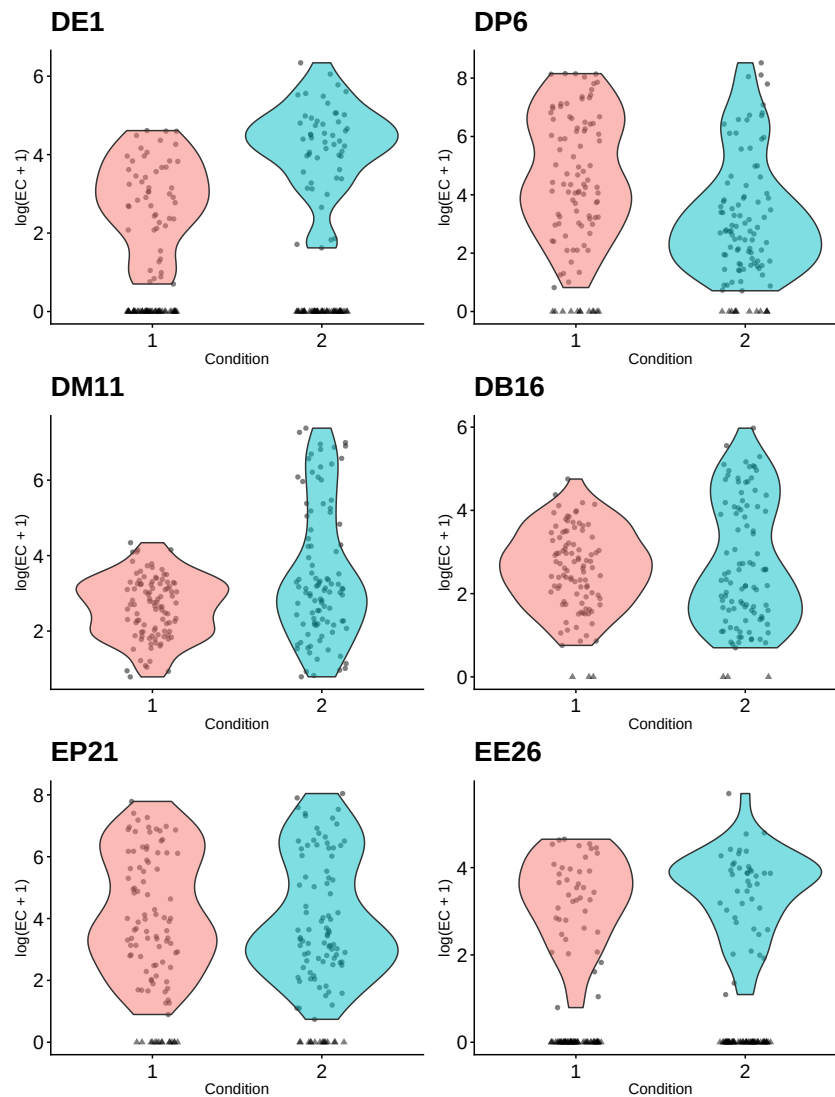


Figure 2: Example Simulated DD genes

7 Session Info

Here is the output of `sessionInfo` on the system where this document was compiled:

```
sessionInfo()

## R version 4.5.1 Patched (2025-08-23 r88802)
## Platform: x86_64-pc-linux-gnu
## 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  LAPACK version 3.12.0
##
```

```

## locale:
## [1] LC_CTYPE=en_US.UTF-8      LC_NUMERIC=C
## [3] LC_TIME=en_GB             LC_COLLATE=C
## [5] LC_MONETARY=en_US.UTF-8   LC_MESSAGES=en_US.UTF-8
## [7] LC_PAPER=en_US.UTF-8      LC_NAME=C
## [9] LC_ADDRESS=C              LC_TELEPHONE=C
## [11] LC_MEASUREMENT=en_US.UTF-8 LC_IDENTIFICATION=C
##
## time zone: America/New_York
## tzcode source: system (glibc)
##
## attached base packages:
## [1] stats4      stats      graphics  grDevices  utils      datasets  methods
## [8] base
##
## other attached packages:
## [1] gridExtra_2.3             SingleCellExperiment_1.31.1
## [3] SummarizedExperiment_1.39.2 Biobase_2.69.1
## [5] GenomicRanges_1.61.5      Seqinfo_0.99.2
## [7] IRanges_2.43.5            S4Vectors_0.47.4
## [9] BiocGenerics_0.55.1        generics_0.1.4
## [11] MatrixGenerics_1.21.0     matrixStats_1.5.0
## [13] scDD_1.33.0
##
## loaded via a namespace (and not attached):
## [1] tidyselect_1.2.1          viridisLite_0.4.2         dplyr_1.1.4
## [4] farver_2.1.2              fields_17.1                S7_0.2.0
## [7] bitops_1.0-9              fastmap_1.2.0             bluster_1.19.0
## [10] rsvd_1.0.5                dotCall64_1.2             digest_0.6.37
## [13] lifecycle_1.0.4          cluster_2.1.8.1           statmod_1.5.0
## [16] magrittr_2.0.4            compiler_4.5.1            rlang_1.1.6
## [19] tools_4.5.1              igraph_2.1.4              yaml_2.3.10
## [22] EBSeq_2.7.0               knitr_1.50                labeling_0.4.3
## [25] dqrng_0.4.1              S4Arrays_1.9.1            mclust_6.1.1
## [28] DelayedArray_0.35.3       RColorBrewer_1.1-3        abind_1.4-8
## [31] BiocParallel_1.43.4       KernSmooth_2.23-26        withr_3.0.2
## [34] grid_4.5.1               caTools_1.18.3            beachmat_2.25.5
## [37] edgeR_4.7.5              ggplot2_4.0.0             scales_1.4.0
## [40] gtools_3.9.5             MASS_7.3-65               blockmodeling_1.1.8
## [43] dichromat_2.0-0.1        tinytex_0.57              cli_3.6.5
## [46] rmarkdown_2.30           crayon_1.5.3              reformulas_0.4.1
## [49] metapod_1.17.0           scuttle_1.19.0            minqa_1.2.8
## [52] outliers_0.15            splines_4.5.1             maps_3.4.3
## [55] parallel_4.5.1           BiocManager_1.30.26       XVector_0.49.1
## [58] vctrs_0.6.5              boot_1.3-32               Matrix_1.7-4
## [61] BiocSingular_1.25.0       BiocNeighbors_2.3.1       irlba_2.3.5.1
## [64] locfit_1.5-9.12          testthat_3.2.3            limma_3.65.5
## [67] glue_1.8.0               spam_2.11-1               nloptr_2.2.1
## [70] codetools_0.2-20         gtable_0.3.6              ScaledMatrix_1.17.0
## [73] lme4_1.1-37              tibble_3.3.0              pillar_1.11.1
## [76] htmltools_0.5.8.1        gplots_3.2.0              brio_1.1.5

```



```
## [79] R6_2.6.1           RcppEigen_0.3.4.0.2 Rdpack_2.6.4
## [82] evaluate_1.0.5      lattice_0.22-7      highr_0.11
## [85] rbibutils_2.3       BiocStyle_2.37.1    arm_1.14-4
## [88] scrn_1.37.0         Rcpp_1.1.0          coda_0.19-4.1
## [91] SparseArray_1.9.1   nlme_3.1-168        xfun_0.53
## [94] pkgconfig_2.0.3
```

References

- [1] Keegan D Korthauer, Li-Fang Chu, Michael A. Newton, Yuan Li, James Thomson, Ron Stewart, and Christina Kendziorski. A statistical approach for identifying differential distributions in single-cell RNA-seq experiments. *Genome Biology*, 17:222, 10 2016.