

R/GSEPD Tutorial

Gene Set Enrichment and Projection Displays

Karl Stamm – karl.stamm@gmail.com –

October 7, 2025

1 Introduction

GSEPD is a package for streamlining RNA-Seq data analysis, targeting complex samples with low replicate count such as human tissues, where all factors (metabolic, genetic, etc) cannot be controlled statistically [Stamm et al., 2019]. As a prerequisite, you need only your multiple samples’ count data as a matrix whose columns are samples and rows are RefSeq NM and NR transcript identifiers. A second matrix associates sample identifiers with treatment/condition. Given both datasets, GSEPD will automate differential expression via DESeq2 [Love et al., 2014], functional analysis via GOSec [Young et al., 2010], generate heatmaps of gene expression for significantly differentially expressed genes, and subsets of genes defined by the significantly enriched GO Terms.

After gene sets are detected from a differential expression analysis, the results are merged into a novel ‘projection display’ wherein each sample is scored according to each condition’s multidimensional average expression. When the treatment samples are found to have a perturbed expression profile for a particular GO Term (geneset), all samples are scored on an axis ranging from control to treatment condition, and outliers or anomalous samples are readily apparent. Clustering quality of samples in a given geneset-space is quantified by the cluster’s “Validity score” [Rosenberg and Hirschberg, 2007] and an empirical permutation derived p-value. GO Terms with more genes than samples in your comparison will randomly appear enriched, so the Segregation P-value is used to determine if a GO Term is significantly segregating your samples.

2 Usage

You’ll need a prepared matrix of read-counts per transcript (Table 1.) You can use HTSeq or RSEM or coverageBed, or any other generation method, so long as it ends with a table of counts by transcript ID. This software comes pre-packaged with a dataset based on the IlluminaBodyMap project, counted with coverageBed. The second prerequisite is a metadata table associating sample identifiers with their test-condition and a nickname to annotate figures with (see Table 2 for the included sample). Alternatively, the manual/Vignette for DESeq2 describes how to generate a dataset object from HTSeq read counts, you can also

initialize GSEPD with the DESeqDataSet object instead of the counts matrix. By default this table will be normalized by DESeq during processing, but if you have a pre-normalized table, you can prevent the double normalization by setting `renormalize=FALSE` in the `INIT()` routine.

2.1 Naming Conventions

In R, column names are not allowed to have spaces or certain special characters like `+` or `.` As your sample names are the columns of the count table, this implies your sample names may not have special characters or spaces. When you load a table with invalid column headers, R may silently convert invalid characters into periods, thus “Sample 1” becomes “Sample.1”. If your metadata table (where samples must be annotated) matches the original count table, it won’t match this converted name, and you’ll either get an error about samples not being found, or worse, an apparently successful run with invalid data. In later stages of the GSEPD process your test conditions also become column headers, and spaces will cause an error before completion. It’s better to compare groups ‘test’ vs ‘control’ than ‘lung tissue’ vs ‘healthy(-ish) person’.

```
library(rgsepd)

## Loading required package: DESeq2
## Loading required package: S4Vectors
## 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
```

```

##
## 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: GenomicRanges
## Loading required package: Seqinfo
## 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, colAugsPerRowSet, 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, rowAugsPerColSet,
##   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: 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
## Loading required package: goseq
## Loading required package: BiasedUrn
## Loading required package: geneLenDataBase
##
## Attaching package: 'geneLenDataBase'
## The following object is masked from 'package:S4Vectors':
##
##      unfactor
##
##
## Loading R/GSEPD 1.41.0
## Building human gene name caches
## 'select()' returned 1:1 mapping between keys and columns
## 'select()' returned 1:1 mapping between keys and columns

set.seed(1000) #fixed randomness
data("IlluminaBodymap" , package="rgsepd")
data("IlluminaBodymapMeta" , package="rgsepd")
```

	adipose.1	adipose.2	adipose.3	adipose.4	blood.1	heart.1	skeletal_muscle.1
NM.000014	91717	91406	91667	91788	290	67285	15609
NM.000015	7	7	7	0	0	3	0
NM.000016	1395	1410	1505	1368	2592	15851	2255
NM.000017	1400	1439	1393	1334	460	1992	2759
NM.000018	9505	9432	9217	9780	4432	33252	15478
NM.000019	5609	5580	5707	5622	1530	16616	11988
NM.000020	5961	5625	5774	5862	98	1658	497
NM.000021	3935	4092	3931	3964	12021	1570	2009
NM.000022	335	267	292	286	723	144	75
NM.000023	103	108	97	85	16	1463	13083

Table 1: First few rows of the included IlluminaBodymap dataset. See ?IlluminaBodymap for more details.

Next we'll sub-select 5,000 genes from this set for speed. Initialize the GSEPD object with the GSEPD_INIT function. Finally, to indicate which conditions will be tested as this dataset includes samples from 'condition' A, B, and C, we use GSEPD_ChangeConditions.

```
isoform_ids <- Name_to_RefSeq(c("GAPDH","HIF1A","EGFR","MYH7","CD33","BRCA2"))
rows_of_interest <- unique( c( isoform_ids ,
                               sample(rownames(IlluminaBodymap),
                                       size=5000,replace=FALSE)))
G <- GSEPD_INIT(Output_Folder="OUT",
                 finalCounts=round(IlluminaBodymap[rows_of_interest , ]),
```

	Sample	Condition	SHORTNAME
1	adipose.1	A	AD1
2	adipose.2	A	AD2
3	adipose.3	A	AD3
4	adipose.4	A	AD4
5	blood.1	C	BL1
6	blood.2	C	BL2

Table 2: First few rows of the included `IlluminaBodymapMeta` dataset. See `?IlluminaBodymapMeta` for more details. These are easy to build with a spreadsheet, saved to csv and R's builtin `?read.csv`

```
sampleMeta=IlluminaBodymapMeta,
COLORS=c(blue="#4DA3FF",black="#000000",gold="#FFFF4D"))

## Keeping rows with counts (4595 of 5006)

G <- GSEPD_ChangeConditions( G, c("A","B"))
```

This should only take a moment, and create the folder `OUT` which will hold your generated results. If you're familiar with R objects, you can explore the `G` object here and change default parameters, all set by `GSEPD_INIT`. We'll change some parameters now to demonstrate:

```
G$MAX_Genes_for_Heatmap <- 30
G$MAX_GOs_for_Heatmap <- 25
G$MaxGenesInSet <- 12
G$LIMIT$LFC <- log( 2.50 , 2 )
G$LIMIT$HARD <- FALSE
```

Here we changed five default settings on the `G` master object. The parameters `MAX_Genes_for_Heatmap` and `MAX_GOs_for_Heatmap` cap how many rows you'll see on your final differential expression heatmap (Figure 6) and the projection HMA file (Figure 4), choosing the most significant rows so your figures are shorter. If you'd like a figure containing everything, make these values large.

The parameter `MaxGenesInSet` controls the size of evaluated GO-Terms. Default is 30, here we reduce it to 12 for speed. Calculating projection significance for large sets can be slow. Also see `MinGenesInSet` for culling niche gene sets. The goal here is to limit our results to gene sets which are not too broad and not too narrow.

The parameter `LIMIT$LFC` is the $\log_2$ minimum foldchange required for significance, here we've set it to require 250% expression up or down (default is 200%, at $LFC=1$). Finally `LIMIT$HARD` if `TRUE` (the default), means figures and plots will respect the specified p-value limits. Sometimes your comparison won't have any significant genes or GO-terms and later stages of the pipeline will error or quit. To force generation of all stages and plots of less-than-strictly significant sets, we have set the `LIMIT$HARD=FALSE`. You'll see messages during processing if very few genes would be strictly significant, as the system adjusts the threshold automatically. By default, limits are hard at $p = 0.05$.

Now we're ready to run the pipeline:

```
G <- GSEPD_Process( G )

## converting counts to integer mode
## Generating OUT/DESEQ.counts.Ax4.Bx8.csv
## converting counts to integer mode
## estimating size factors
## estimating dispersions
## gene-wise dispersion estimates
## mean-dispersion relationship
## final dispersion estimates
## fitting model and testing
## using 'normal' for LFC shrinkage, the Normal prior from Love et al (2014).
##
## Note that type='apeglm' and type='ashr' have shown to have less bias than
type='normal'.
## See ?lfcShrink for more details on shrinkage type, and the DESeq2 vignette.
## Reference: https://doi.org/10.1093/bioinformatics/bty895
## Generating OUT/DESEQ.Volcano.Ax4.Bx8.png
## Generating OUT/DESEQ.RES.Ax4.Bx8.csv
## dropping rows with raw PVAL>0.990 OR baseMean < 1 OR on excludes list
## deseq txid rows remaining: 2844
## Converting identifiers with local DB
## Generating OUT/DESEQ.RES.Ax4.Bx8.Annote.csv
## Too many genes found differentially expressed, changing the threshold to
raw p=0.000000 so we can use 30 genes in the heatmap.
## Generating OUT/HM.Ax4.Bx8.30.pdf
## Loading hg19 length data...
## Fetching GO annotations...
## Loading required package: AnnotationDbi
## Calculating the p-values...
## 'select()' returned 1:1 mapping between keys and columns
## Loading hg19 length data...
## Fetching GO annotations...
## Calculating the p-values...
## 'select()' returned 1:1 mapping between keys and columns
## 'select()' returned 1:1 mapping between keys and columns
## 'select()' returned 1:1 mapping between keys and columns
## 'select()' returned 1:1 mapping between keys and columns
## Generating OUT/GOSEQ.RES.Ax4.Bx8.GO.csv
```

```
## Written GO categories, now reverse mapping
## Generating OUT/GSEPD.RES.Ax4.Bx8.GO2.csv
## Generating OUT/GSEPD.RES.Ax4.Bx8.MERGE.csv
## Generating OUT/GSEPD.PCA_AG.Ax4.Bx8.pdf
## Generating OUT/GSEPD.PCA_DEG.Ax4.Bx8.pdf
## Generating OUT/SCA.GSEPD.Ax4.Bx8.pdf
## Calculating Projections and Segregation Significance
## Generating OUT/GSEPD.HMA.Ax4.Bx8.pdf
## Generating OUT/GSEPD.HMA.Ax4.Bx8.csv
## All Done!
```

This step can take a few minutes on a full genome-wide dataset. If you change something and re-run GSEPD will reuse any files it finds with the same filename, so you don't have to wait for each step again, unless the filenames change. GSEPD's automation steps will convert gene identifiers, and GOSep can take a few minutes as it runs on three differential expression sets (the upregulated, the downregulated, and the combined). All of these results are saved as CSV files under the OUT folder.

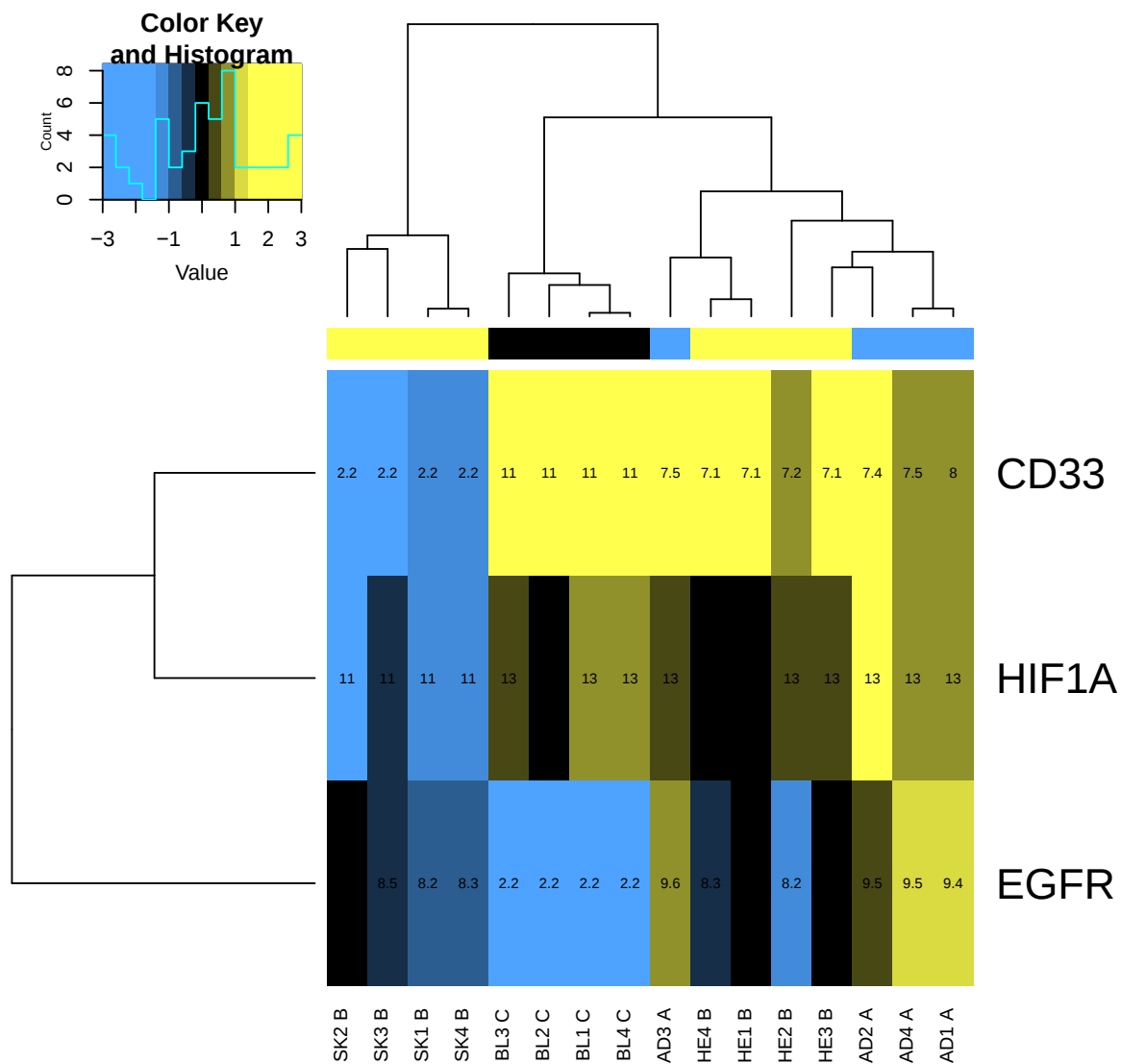
We save the G object from a GSEPD_Process routine, to retain information such as the normalized counts (wherein DESeq adjusted for library sequencing depths). This object will be passed to further visualizations, such as heatmaps and PCA.

```
print(isoform_ids)

## [1] "NR_152150"      "NM_181054"      "NM_201284"      "NM_001407004"  "NM_001772"
## [6] "NR_176251"

GSEPD_Heatmap(G,isoform_ids)

## Warning: 3 gene ids are not found in count data
```



```
GSEPD_PCA_Plot(G)
```

```
## Generating OUT/GSEPD.PCA_AG.Ax4.Bx8.pdf, overwriting previous existing version.
```

```
## Generating OUT/GSEPD.PCA_DEG.Ax4.Bx8.pdf, overwriting previous existing version.
```

```
## pdf
```

```
## 2
```

3 Results

Several files are generated from each run. When `GSEPD_Processis` invoked the pre-specified conditions are compared, or when `GSEPD_ProcessAllis` invoked, each condition is tested against all others. For each comparison, files with the comparison listed in their filename are generated. For a condition named “A” with N samples, versus “B” with M samples, your normalized counts file will be written to `OUT\DESEQ.counts.AxN.BxM.csv` for example. If one of your conditions has the letter `x` in it, please change the delimiter with something like `G$C2T_Delimiter <- 'z'` or other unused character.

3.1 Heatmap Organizational Clustering

Each generated heatmap figure organizes the rows and columns such that similar profiles are adjacent. For this we use the default methods of `gplots::heatmap.2` which calls `hclust` on the supplied data. For gene heatmaps where you see numbers in each cell, representing the gene’s expression value calculated by `DESeq2::varianceStabilizingTransformation`. For users who have pre-normalized their datasets, systemic re-normalization can be disabled at initialization with the `renormalize` parameter. The magnitude of a gene’s expression might dominate the heirarchical clustering, so scaling is warranted. `GSEPD` will scale gene expression values within each row(gene) as a normal Gaussian by subtracting the row’s mean and dividing out the standard deviation. Therefore it is dependent on the samples used in the figure, and heirarchical clustering with complete linkage is not guaranteed to be stable. To ensure some values of each specified color, the normalized color data is capped to a specifiable minimum and maximum (default 3) before `heatmap.2`’s clustering is performed.

```
Annotated_Filtered <- read.csv("OUT/DESEQ.RES.Ax4.Bx8.Annote_Filter.csv",
                              header=TRUE, as.is=TRUE)
```

REFSEQ	baseMean	Ax4	Bx8	X.X.Y.	LOG2.X.Y.	lfcSE	PVAL	PADJ	HGNC	ENTREZ
NM.000016	6263.66	10.40	12.87	0.13	-2.98	0.56	0.00	0.01	ACADM	34
NM.000028	5294.45	10.44	13.13	0.15	-2.70	0.05	0.00	0.00	AGL	178
NM.000029	2769.14	10.22	12.22	0.24	-2.06	0.22	0.00	0.00	AGT	183
NM.000034	87627.22	14.42	16.91	0.15	-2.78	0.43	0.00	0.00	ALDOA	226
NM.000045	327.85	10.26	3.08	295.78	8.21	0.29	0.00	0.00	ARG1	383
NM.000050	2269.53	12.66	10.34	4.99	2.32	0.08	0.00	0.00	ASS1	445
NM.000062	20542.97	15.71	13.75	3.89	1.96	0.03	0.00	0.00	SERPING1	710
NM.000092	332.62	5.25	8.97	0.05	-4.30	0.47	0.00	0.00	COL4A4	1286
NM.000164	41.22	7.15	4.41	8.06	3.01	0.54	0.00	0.01	GIPR	2696
NM.000206	4859.52	9.22	7.33	3.82	1.93	0.12	0.00	0.00	IL2RG	3561

Table 3: First few rows of `OUT/DESEQ.RES.Ax4.Bx8.Annote_Filter.csv` which contains the DESeq results, cropped for significant results, and annotated with gene names (the HGNC Symbol).

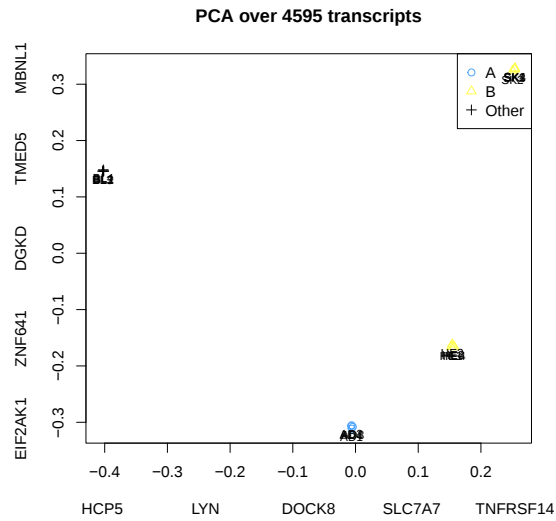


Figure 1: OUT\GSEPD.PCA_AG.Ax4.Bx8.pdf is the Principle Components Analysis of All Genes, from the comparison Ax4 vs Bx8 under run OUT. This function annotates the top four major genes underlying each principle component dimension along each axis (by maximum absolute weight). Samples from the tested condition are marked in the comparison colors, here blue and gold. All genes and all samples are included. A true outlier sample can direct the principle components away from the differentiating genes, and all tested samples can show as a single cluster.

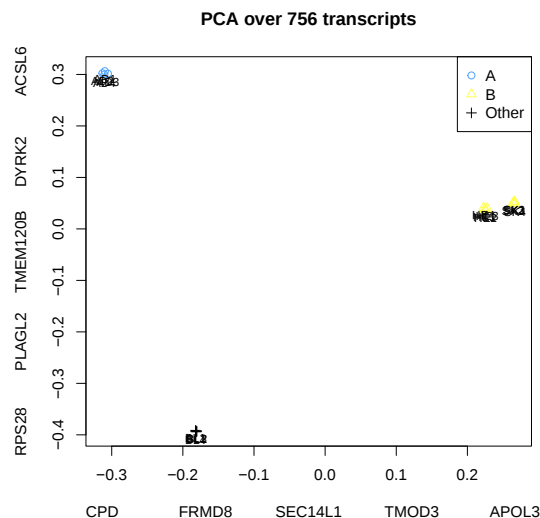


Figure 2: OUT\GSEPD.PCA_DEG.Ax4.Bx8.pdf is the Principle Components Analysis of only Differentially Expressed Genes, from the comparison Ax4 vs Bx8 under run OUT. This function annotates the top four major genes underlying each principle component dimension along each axis (by maximum absolute weight). Samples from the tested condition are marked in the comparison colors, here blue and gold, with the non-comparison samples marked as black ‘Other’. **Because we’re only reviewing genes found to be differentially expressed, this plot is guaranteed to show separation of your samples, sometimes spuriously.**

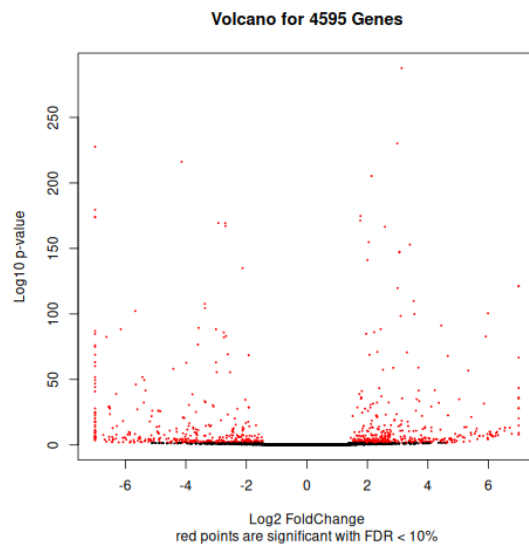


Figure 3: OUT\DESEQ.Volcano.Ax4.Bx8.png is the ‘volcano’ plot from the comparison Ax4 vs Bx8. Here, the horizontal axis is the relative fold-change between conditions, and the vertical is significance. A left-right balanced figure indicates similar numbers of genes found up and down-regulated. The bundled IlluminaBodymap dataset does not have biological replicates, causing the Volcano plot to show too many significant genes.

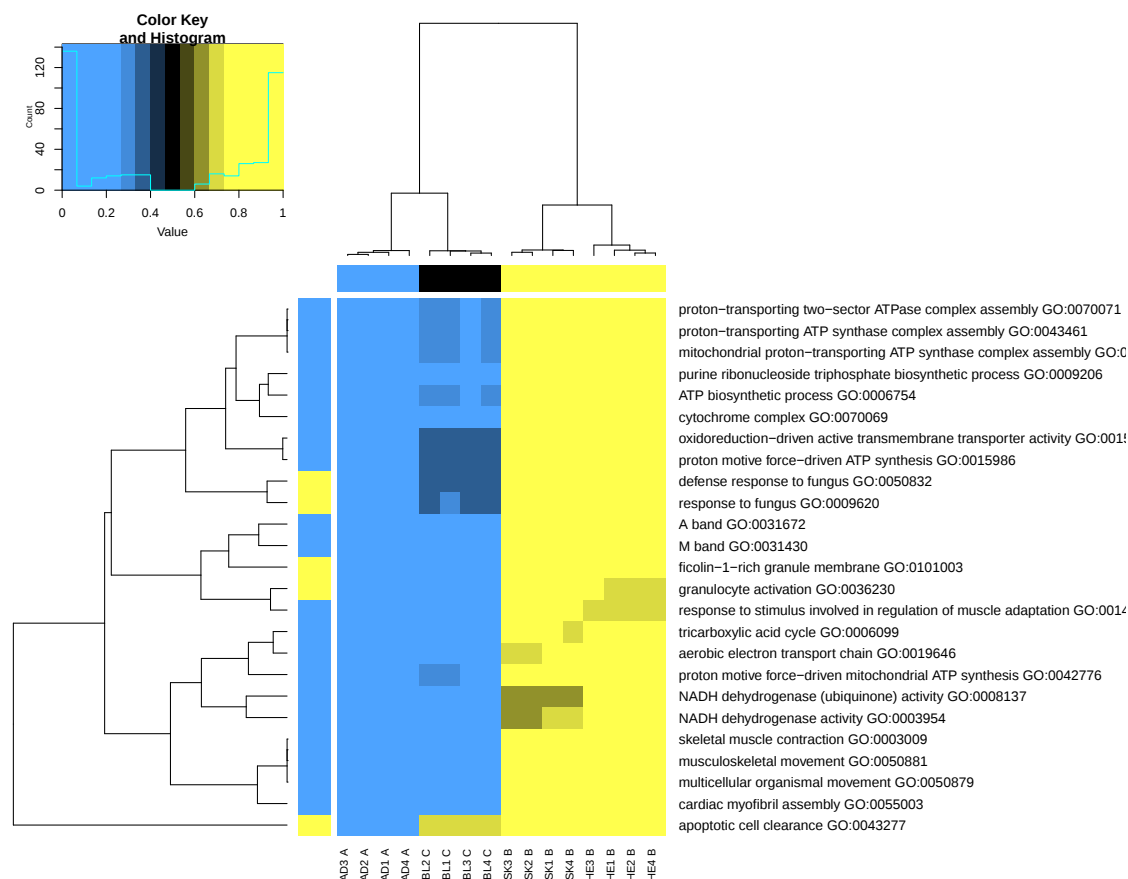


Figure 4: OUT\GSEPD.HMA.Ax4.Bx8.pdf is the projection display summary heatmap from the comparison Ax4 vs Bx8. Here, as in a normal heatmap, your samples are columns, and rows are GO Terms with significant segregation ability. All rows and columns are arranged to cluster. The color in each cell represents the sample's Alpha score for that GO Term, with blue indicating similarity to class 'A', and the gold indicating similarity to class 'B'. The unlabeled top row indicates the comparison categories, also seen on the sample labels along the bottom (A, B, or C). Any white dots indicate the sample was distant from the axis between conditions, so the color should be interpreted with caution or investigated further. The most interesting results from GSEPD are here, when unclassified samples (Blood/C) are scored as similar to either of the tested conditions on a geneset-by-geneset basis. Both the Alpha table and Beta table are summarized in the HMA figure.

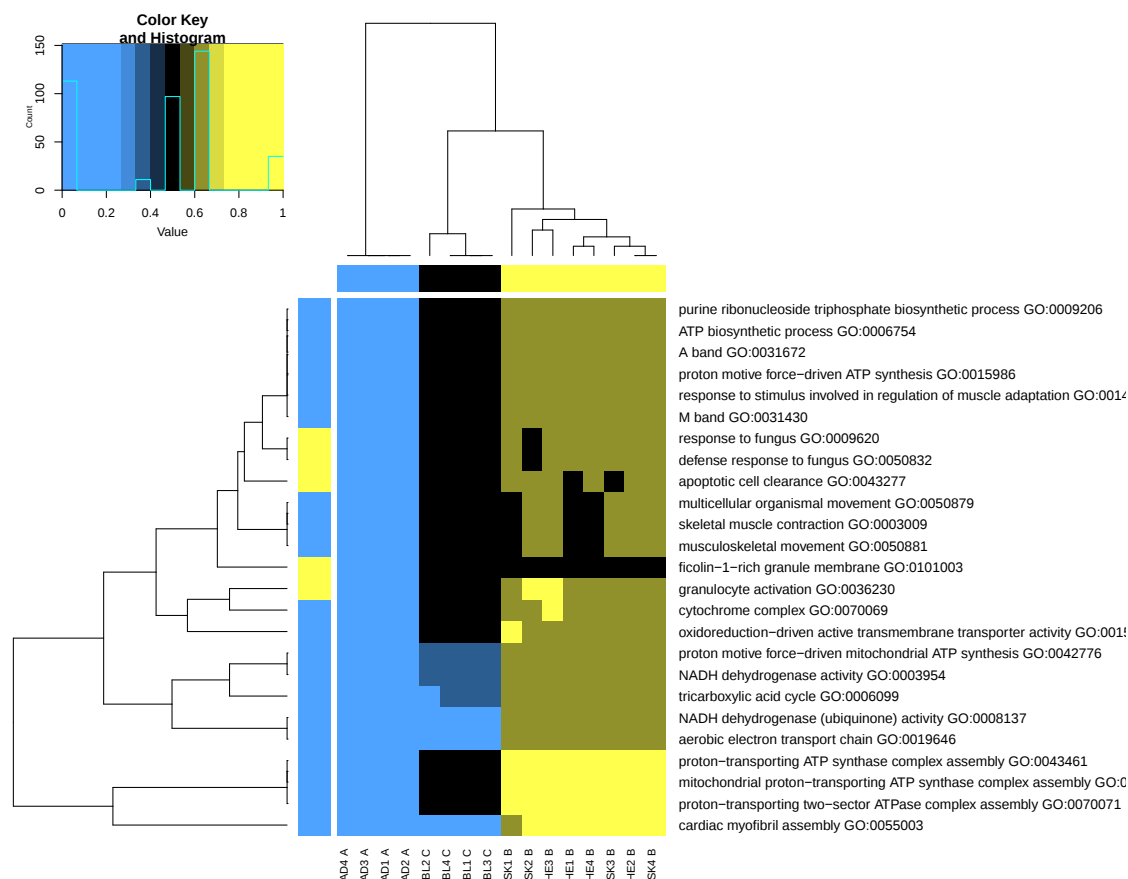


Figure 5: OUT\GSEPD.HMG.Ax4.Bx8.pdf is a simplified projection display summary heatmap from the comparison Ax4 vs Bx8. Here, as in a normal heatmap, your samples are columns, and rows are GO Terms with significant segregation ability. All rows and columns are arranged to cluster. The color in each cell represents the sample's $\Gamma(1/2)$ score for that GO Term, with blue indicating similarity to class 'A', and the gold indicating similarity to class 'B'. The unlabeled top row indicates the comparison categories, also seen on the sample labels along the bottom (A, B, or C). Black areas indicate the sample was distant from both conditions. The most interesting results from GSEPD are here, when unclassified samples (Blood/C) are scored as similar to either of the tested conditions on a geneset-by-geneset basis. Both the Γ_1 table and Γ_2 table are summarized in this HMG figure.

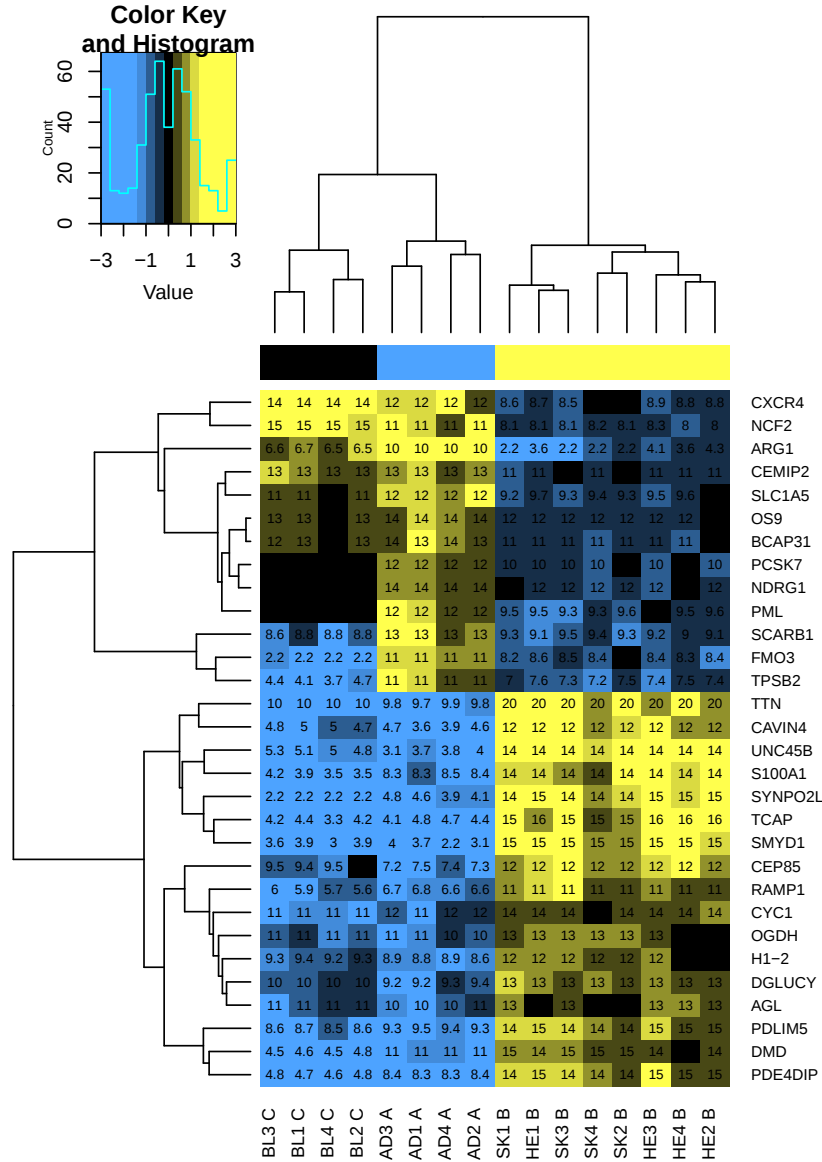


Figure 6: OUT\HM.Ax4.Bx8.30.pdf is a heatmap of your comparisons with full expression details for those genes found significantly expressed. The number of genes is given in the filename. Each row is scaled such that the lowest values are blue and the highest are gold (default colors are changable in GSEPD-INIT). The numbers within each cell are the expression values as variance-stabilized normalized counts, similar to a log transform, provided by DESeq2. Across the top, between the sample clustering dendrogram and the heatmap itself is a colorbar annotating which samples belong to which comparison group, for this differential expression blue vs gold. Black corresponds to extraneous samples, contributing context. Other variations of this plot with less information are generated as HMS files (compared samples only) and HM- files (minimal figure without details).

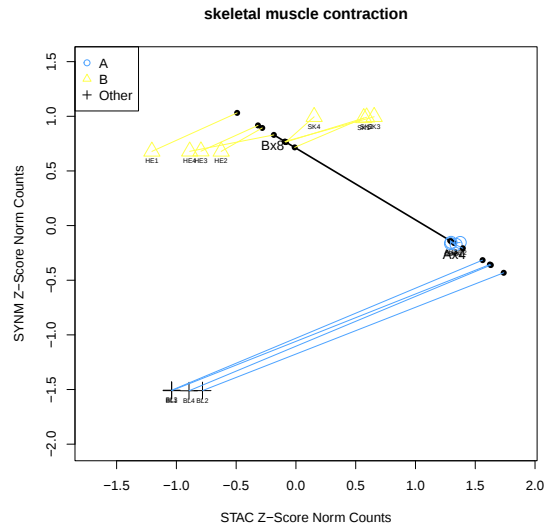


Figure 7: OUT\SCGO\GSEPD.Ax4.Bx8.GO0003009.pdf First page of the projections file for one GO Term. All significant sets have this figure generated displaying the central axis between two groups as a black line (with ends annotated Ax4 and Bx8), and each sample's closest point along the line. These files have half as many pages as genes in the set, so we can display pairs as two-dimensional scatterplots. It's a workaround to the problem of displaying an N-dimensional scatterplot. For sets with fewer than 11 genes, full pairings are viewable with the Pairs file, Figure 8

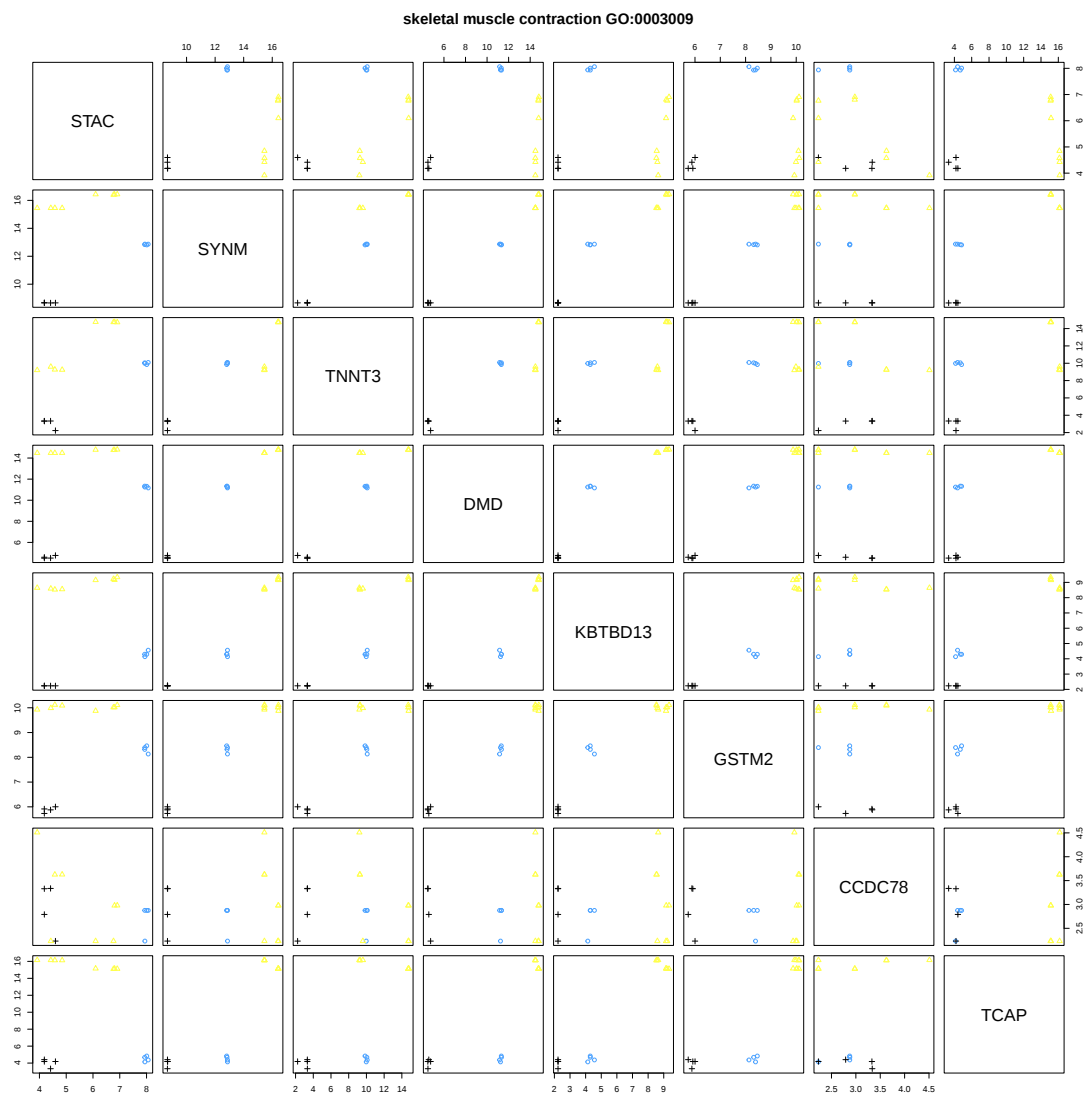


Figure 8: OUT\SCGO\Pairs.Ax4.Bx8.GO0003009.pdf The Pairs file is generated for significant GO terms with between two and ten genes, making it easy to review correlations or subgroups of gene expression among low-level gene sets.

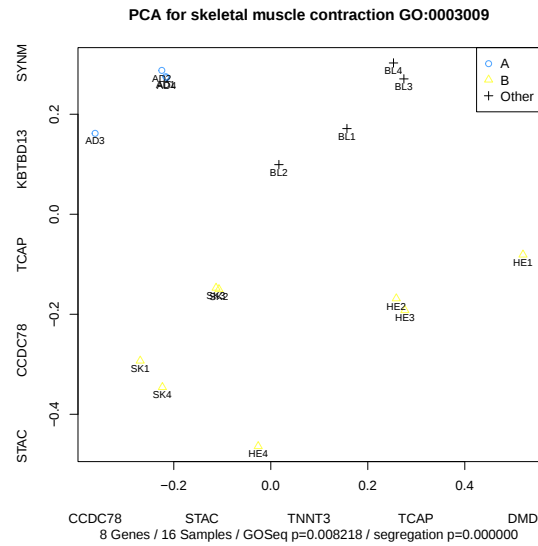


Figure 9: OUT\SCGO\Scatter.Ax4.Bx8.GO0003009.pdf The Scatter file is similar to the PCA files, but the gene set is restricted to those within a given, significant, GO Term. As in the PCA file, genes driving the principle components are annotated along the axes. At the bottom of the figure, statistics are given pertaining to this group. All such data is available in tables, but this Scatter plot helps the user see which samples behave like which groups. Depending on the number of genes relative to the number of samples, different PCA formulas may be used. For a set with only two genes, this file can directly display the expression values as normalized counts, the same as found on the expression heatmap (Fig. 6.)

X	category	over_represented_padj.x	Term.x	GOSEQ_DEG_Type	HGNC	LOG2.X.Y.	REFSEQ	Ax4	Bx8	PVAL	PADJ
1817	GO:0003009	0.01	skeletal muscle contraction	Down	STAC	2.14	NM_003149	7.99	5.54	0.08	0.33
1818	GO:0003009	0.01	skeletal muscle contraction	Down	SYNM	-3.17	NM_145728	12.80	15.90	0.00	0.00
1819	GO:0003009	0.01	skeletal muscle contraction	Down	TNNT3	-3.73	NM_006757	9.99	12.00	0.01	0.04
1820	GO:0003009	0.01	skeletal muscle contraction	Down	DMD	-3.37	NM_004006	11.30	14.60	0.00	0.00
1821	GO:0003009	0.01	skeletal muscle contraction	Down	KBTBD13	-5.33	NM_001101362	4.32	8.90	0.00	0.00
1822	GO:0003009	0.01	skeletal muscle contraction	Down	GSTM2	-1.72	NM_000848	8.33	10.00	0.00	0.00
1823	GO:0003009	0.01	skeletal muscle contraction	Down	CCDC78	-2.00	NM_001031737	2.71	3.05	0.27	0.87
1824	GO:0003009	0.01	skeletal muscle contraction	Down	TCAP	-11.80	NM_003673	4.51	15.60	0.00	0.00
2696	GO:0003954	0.00	NADH dehydrogenase activity	Down	NDUFB3	-2.78	NM_002491	9.41	11.60	0.00	0.03
2697	GO:0003954	0.00	NADH dehydrogenase activity	Down	NDUFA4	-2.90	NM_002489	10.60	13.20	0.00	0.00
2698	GO:0003954	0.00	NADH dehydrogenase activity	Down	NDUFB6	-2.13	NM_002493	9.40	11.40	0.01	0.04
2699	GO:0003954	0.00	NADH dehydrogenase activity	Down	NDUFC1	-2.49	NM_002494	9.03	11.40	0.00	0.00
2700	GO:0003954	0.00	NADH dehydrogenase activity	Down	ENOX1	-2.12	NM_001242863	6.75	8.78	0.00	0.00
5450	GO:0006099	0.01	tricarboxylic acid cycle	Down	MDH1B	-1.34	NM_001039845	2.86	3.22	0.47	1.00
5451	GO:0006099	0.01	tricarboxylic acid cycle	Down	CS	-1.66	NM_004077	12.80	14.50	0.00	0.00

Table 4: First few rows of OUT/GSEPD.RES.Ax4.Bx8.MERGE.csv showing enriched GO Terms, and each terms’ underlying gene expression averages per group. This data is central to the rgsepd package, defining the group centroids per GO-Term. It consists of the cross-product of the GO enrichment statistics and the DESeq differential expression and summarization.

	adipose.1	adipose.2	adipose.3	adipose.4	blood.1	heart.1	skeletal_muscle.1
GO:0003009	0.02	0.00	-0.05	0.02	-0.20	1.20	0.88
GO:0003954	-0.01	-0.02	0.01	0.02	0.04	1.34	0.67
GO:0006099	-0.01	0.00	0.01	-0.00	0.19	1.27	0.74
GO:0006754	-0.01	-0.00	-0.00	0.02	0.27	1.14	0.87
GO:0008137	-0.02	-0.02	0.00	0.03	0.15	1.37	0.64
GO:0009206	-0.01	-0.00	-0.00	0.02	0.24	1.09	0.92
GO:0009620	-0.00	0.00	0.00	-0.00	0.30	0.79	1.10
GO:0014874	0.02	0.01	-0.02	-0.01	-0.42	0.71	1.28
GO:0015453	-0.01	-0.01	0.01	0.02	0.35	1.18	0.82
GO:0015986	-0.01	-0.01	-0.00	0.02	0.35	1.19	0.82

Table 5: First ten rows of OUT/GSEPD.Alpha.Ax4.Bx8.csv showing the group projection scores for each sample, these directly correspond to the colors in the HMA file. Where the HMA displays only significant sets, the Alpha table continues for all tested GO Terms. Both the Alpha table and Beta table are summarized in Figure 4.

References

Michael I Love, Wolfgang Huber, and Simon Anders. Moderated estimation of fold change and dispersion for rna-seq data with deseq2. *bioRxiv*, 2014. doi: 10.1101/002832. URL <http://dx.doi.org/10.1101/002832>.

Andrew Rosenberg and Julia Hirschberg. *V-Measure: A conditional entropy-based external cluster evaluation measure*, 2007. URL http://www1.cs.columbia.edu/~amaxwell/pubs/v_measure-emnlp07.pdf. Department of Computer Science Columbia University New York, NY 10027.

Karl Stamm, Aoy Tomita-Mitchell, and Serdar Bozdag. Gsepd: a bioconductor package for rna-seq gene set enrichment and projection display. *BMC Bioinformatics*, 20(1):115, 2019. ISSN 1471-2105. doi: 10.1186/s12859-019-2697-5. URL <https://bmcbioinformatics.biomedcentral.com/articles/10.1186/s12859-019-2697-5>.

Matthew Young, Matthew Wakefield, Gordon Smyth, and Alicia Oshlack. Gene

	adipose.1	adipose.2	adipose.3	adipose.4	blood.1	heart.1	skeletal_muscle.1
GO:0003009	0.03	0.04	0.09	0.03	0.48	0.31	0.20
GO:0003954	0.02	0.04	0.02	0.06	0.32	0.12	0.11
GO:0006099	0.09	0.09	0.08	0.09	0.21	0.18	0.17
GO:0006754	0.02	0.02	0.01	0.03	0.22	0.17	0.16
GO:0008137	0.01	0.05	0.01	0.05	0.13	0.12	0.10
GO:0009206	0.02	0.02	0.01	0.02	0.21	0.19	0.19
GO:0009620	0.01	0.01	0.01	0.00	0.40	0.17	0.15
GO:0014874	0.01	0.00	0.01	0.02	0.46	0.13	0.13
GO:0015453	0.01	0.03	0.00	0.03	0.35	0.16	0.14
GO:0015986	0.03	0.03	0.01	0.03	0.23	0.17	0.17

Table 6: First ten rows of OUT/GSEPD.Beta.Ax4.Bx8.csv showing the linear divergence (distance to axis) for each sample, high values here would be annotated with white dots on the HMA file to indicate that a sample is not falling on the axis. Non-tested samples are expected to frequently have high values here, the C group was not part of the A vs B comparison. Where the HMA displays only significant sets, the Beta table continues for all tested GO Terms. Both the Alpha table and Beta table are summarized in Figure 4.

ontology analysis for rna-seq: accounting for selection bias. *Genome Biology*, 11(2):R14, 2010. ISSN 1465-6906. doi: 10.1186/gb-2010-11-2-r14. URL <http://genomebiology.com/2010/11/2/R14>.

Session Info

```
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)
```

	adipose.1	adipose.2	adipose.3	adipose.4	blood.1	heart.1	skeletal_muscle.1
GO:0003009	0.08	0.09	0.23	0.08	1.15	1.41	1.00
GO:0003954	0.02	0.06	0.02	0.08	0.44	1.35	0.69
GO:0006099	0.14	0.14	0.14	0.14	0.40	1.30	0.79
GO:0006754	0.04	0.04	0.01	0.05	0.47	1.17	0.91
GO:0008137	0.02	0.06	0.01	0.07	0.21	1.38	0.65
GO:0009206	0.04	0.04	0.01	0.05	0.49	1.15	1.00
GO:0009620	0.01	0.03	0.02	0.01	0.88	0.87	1.14
GO:0014874	0.03	0.01	0.03	0.04	1.00	0.75	1.31
GO:0015453	0.02	0.04	0.01	0.05	0.70	1.20	0.86
GO:0015986	0.04	0.04	0.01	0.05	0.49	1.22	0.85

Table 7: First ten rows of OUT/GSEPD.HMG1.Ax4.Bx8.csv showing the z-scaled distance to the Group1 centroid for each sample, these directly correspond to the colors in the HMG file. Where the HMG displays only significant sets, the Gamma table continues for all tested GO Terms. Both the Gamma1 and Gamma2 tables are summarized in Figure 5. Distance is normalized to dimensionality by scaling between the centroids. Thus a score of 0 means the sample resides on the centroid, and a score of 1 means it resides on the opposite class centroid, or equidistant.

```
##
## attached base packages:
## [1] stats4      stats      graphics  grDevices  utils      datasets  methods
## [8] base
##
## other attached packages:
## [1] org.Hs.eg.db_3.21.0      AnnotationDbi_1.71.1
## [3] rgsepd_1.41.0           goseq_1.61.1
## [5] geneLenDataBase_1.45.0   BiasedUrn_2.0.12
## [7] DESeq2_1.49.4           SummarizedExperiment_1.39.2
## [9] Biobase_2.69.1          MatrixGenerics_1.21.0
## [11] matrixStats_1.5.0       GenomicRanges_1.61.5
## [13] Seqinfo_0.99.2          IRanges_2.43.5
## [15] S4Vectors_0.47.4        BiocGenerics_0.55.1
## [17] generics_0.1.4          xtable_1.8-4
##
## loaded via a namespace (and not attached):
## [1] tidyselect_1.2.1        dplyr_1.1.4            farver_2.1.2
## [4] blob_1.2.4              filelock_1.0.3         Biostrings_2.77.2
## [7] S7_0.2.0                bitops_1.0-9           fastmap_1.2.0
## [10] RCurl_1.98-1.17         BiocFileCache_2.99.6   GenomicAlignments_1.45.4
## [13] XML_3.99-0.19           lifecycle_1.0.4        KEGGREST_1.49.1
## [16] RSQLite_2.4.3           magrittr_2.0.4         compiler_4.5.1
## [19] rlang_1.1.6             progress_1.2.3         tools_4.5.1
## [22] yaml_2.3.10             rtracklayer_1.69.1     knitr_1.50
## [25] prettyunits_1.2.0       S4Arrays_1.9.1         bit_4.6.0
```

	adipose.1	adipose.2	adipose.3	adipose.4	blood.1	heart.1	skeletal_muscle.1
GO:0003009	0.98	1.00	1.07	0.98	1.65	0.77	0.48
GO:0003954	1.01	1.02	0.99	0.98	1.06	0.38	0.36
GO:0006099	1.02	1.01	1.00	1.01	0.88	0.39	0.37
GO:0006754	1.01	1.01	1.01	0.98	0.82	0.32	0.31
GO:0008137	1.02	1.02	1.00	0.97	0.86	0.40	0.38
GO:0009206	1.01	1.01	1.00	0.98	0.87	0.38	0.38
GO:0009620	1.00	1.00	1.00	1.00	1.09	0.42	0.33
GO:0014874	0.98	0.99	1.02	1.01	1.69	0.39	0.39
GO:0015453	1.01	1.01	0.99	0.98	0.89	0.32	0.30
GO:0015986	1.01	1.01	1.00	0.98	0.73	0.32	0.31

Table 8: First ten rows of OUT/GSEPD.HMG2.Ax4.Bx8.csv showing the z-scaled distance to the Group2 centroid for each sample, these directly correspond to the colors in the HMG file. Where the HMG displays only significant sets, the Gamma table continues for all tested GO Terms. Both the Gamma1 and Gamma2 tables are summarized in Figure 5. Distance is normalized to dimensionality by scaling between the centroids. Thus a score of 0 means the sample resides on the centroid, and a score of 1 means it resides on the opposite class centroid, or equidistant.

## [28]	curl_7.0.0	DelayedArray_0.35.3	RColorBrewer_1.1-3
## [31]	KernSmooth_2.23-26	abind_1.4-8	BiocParallel_1.43.4
## [34]	txdbmaker_1.5.6	grid_4.5.1	caTools_1.18.3
## [37]	G0.db_3.21.0	ggplot2_4.0.0	gtools_3.9.5
## [40]	scales_1.4.0	dichromat_2.0-0.1	biomaRt_2.65.16
## [43]	cli_3.6.5	crayon_1.5.3	httr_1.4.7
## [46]	rjson_0.2.23	DBI_1.2.3	cachem_1.1.0
## [49]	stringr_1.5.2	splines_4.5.1	parallel_4.5.1
## [52]	XVector_0.49.1	restfulr_0.0.16	vctrs_0.6.5
## [55]	Matrix_1.7-4	jsonlite_2.0.0	hms_1.1.3
## [58]	bit64_4.6.0-1	GenomicFeatures_1.61.6	locfit_1.5-9.12
## [61]	glue_1.8.0	codetools_0.2-20	stringi_1.8.7
## [64]	gtable_0.3.6	GenomeInfoDb_1.45.12	BiocIO_1.19.0
## [67]	UCSC.utils_1.5.0	tibble_3.3.0	pillar_1.11.1
## [70]	gplots_3.2.0	rappdirs_0.3.3	R6_2.6.1
## [73]	dbplyr_2.5.1	httr2_1.2.1	evaluate_1.0.5
## [76]	lattice_0.22-7	highr_0.11	png_0.1-8
## [79]	Rsamtools_2.25.3	memoise_2.0.1	Rcpp_1.1.0
## [82]	nlme_3.1-168	SparseArray_1.9.1	mgcv_1.9-3
## [85]	xfun_0.53	pkgconfig_2.0.3	