

Using clusterProfiler to characterize multiomics data

In the format provided by the
authors and unedited

Supplementary Note 1:

Parameters for differential analysis in case study I

DA function:

```
DA <- function(expr,
               meta,
               abundance = 'Abundance',
               group_colname = 'Diagnosis',
               force = TRUE,
               relative = FALSE,
               subset_group,
               diff_group,
               filter.p = 'pvalue', ...) {
  sign_group_colname <- paste0('Sign_', group_colname)
  mpse <- MPSE(expr)
  mpse <- mpse |> left_join(meta, by = "Sample")
  mpse |>
    dplyr::filter(!as.symbol(group_colname) %in% subset_group) |>
    mp_diff_analysis(
      .abundance = !!as.symbol(abundance),
      .group = !!as.symbol(group_colname),
      force = force,
      relative = relative,
      filter.p = filter.p,
```

```

26     ...
27   ) |>
28   mp_extract_feature() |>
29   dplyr::filter(!as.symbol(sign_group_colname) == diff_group) |>
30   dplyr::pull(OTU)
31 }

```

32

33 We defined the function DA, which is essentially a second-level encapsulated
34 function, to enhance code conciseness. The main approach involves utilizing the feature
35 abundance table and sample information to construct an MPSE object. Following this,
36 the mp_diff_analysis function is applied to conduct a differential analysis, from which
37 the differential features are subsequently extracted. The parameters of this function are
38 as follows:

39 **expr:** Stores the abundance matrix of microbial genes or metabolites, where rows
40 represent features and columns represent samples.

41 **meta:** Sample grouping information, comprising at least two columns. The first column
42 is the sample identifier, corresponding to the column names in the expr parameter. The
43 second column represents sample grouping.

44 **group_colname:** The column name in the meta parameter where sample grouping
45 information is located.

46 **force:** Whether to use the original abundance matrix to perform differential analysis
47 without rarefaction.

48 **relative:** Whether to calculate the relative abundance to perform differential analysis.

49 **subset_group:** A vector of length 2, representing the two group names for inter-group
50 differential analysis – the treatment group and the control group.

51 **diff_group:** The name of the treatment group within subset_group.

52 filter.p: A string specifying the method for filtering p-values, with the default being fdr.

53 This implies that features with $\text{fdr} \leq .\text{first.test.alpha}$ will be retained. If set to pvalue,

54 features with p-value $\leq .\text{first.test.alpha}$ will be kept.

55 ...: Additional parameters passed to the MicrobiotaProcess::mp_diff_analysis function.

56

57 Given that we use a second-level encapsulated function, the primary goal is

58 differential analysis and filtering of significantly different features between different

59 groups. For a comprehensive understanding of the function used in this instance, please

60 refer to this link:

61 [https://www.bioconductor.org/packages/release/bioc/vignettes/MicrobiotaProcess/inst/](https://www.bioconductor.org/packages/release/bioc/vignettes/MicrobiotaProcess/inst/doc/MicrobiotaProcess.html#biomarker-discovery)

62 [doc/MicrobiotaProcess.html#biomarker-discovery](https://www.bioconductor.org/packages/release/bioc/vignettes/MicrobiotaProcess/inst/doc/MicrobiotaProcess.html#biomarker-discovery)

63

64 **Supplementary Note 2**

65 This is the source code of the procedures with syntax highlighted.

66 **Procedure 1: Functional Enrichment Analyses in Metabolomic and Metagenomic**

67 **Data**

68 Download data:

```
69 url <- "https://yulab-smu.top/clusterProfiler_protocol/examples"
70
71 IBD_files <- c("mg.meta.csv", "mg.expr.csv",
72               "metabolism_meta.csv", "metabolism_expr.csv")
73
74 for (f in IBD_files) {
```

```
75   download.file(file.path(url, "IBD_2_subtypes", f), destfile =  
76   f)  
77 }
```

78 **1|** Load the R packages into the R environment:

```
79 library(MicrobiotaProcess)  
80 library(clusterProfiler)  
81 library(ggplot2)  
82 library(enrichplot)
```

83 **2|** Import metagenomic data:

```
84 meta_mg <- read.csv("mg.meta.csv")  
85 metagenome <- read.csv("mg.expr.csv",  
86   row.names = 1,  
87   check.name = FALSE)
```

88 **3|** Define a function to perform differential analysis:

```
89 DA <- function(expr,  
90   meta,  
91   abundance = 'Abundance',  
92   group_colname = 'Diagnosis',  
93   force = TRUE,  
94   relative = FALSE,  
95   subset_group,  
96   diff_group,  
97   filter.p = 'pvalue', ...) {
```

```

98   sign_group_colname <- paste0('Sign_', group_colname)
99   mpse <- MPSE(expr)
100  mpse <- mpse |> left_join(meta, by = "Sample")
101  mpse |>
102  dplyr::filter(!!as.symbol(group_colname) %in% subset_group) |>
103  mp_diff_analysis(
104    .abundance = !!as.symbol(abundance),
105    .group = !!as.symbol(group_colname),
106    force = force,
107    relative = relative,
108    filter.p = filter.p,
109    ...
110  ) |>
111  mp_extract_feature() |>
112  dplyr::filter(!!as.symbol(sign_group_colname) == diff_group)
113  |>
114  dplyr::pull(OTU) |> suppressMessages()
115  }

```

116 **4|** Identify differentially abundant microbial genes:

```

117  groups <- c(CD = 'CD', UC = 'UC')
118  de_gene <- lapply(groups, function(x) {
119    DA(expr = metagenome,
120      meta = meta_mg,
121      subset_group = c(x, 'Control'),

```

```
122   diff_group = x)
```

```
123   })
```

124 **5|** Perform enrichment analysis of differential genes:

```
125 gene_enrich_result <- compareCluster(geneClusters = de_gene,
```

```
126   fun = "enrichKEGG",
```

```
127   organism = "ko")
```

128 **6|** Visualize the functional enrichment result of differential genes:

```
129 dotplot(gene_enrich_result, facet = 'intersect', showCategory =  
130   10,
```

```
131   split = "intersect", label_format = 60) +
```

```
132   ggtitle("Functional enrichment of intestinal genes") +
```

```
133   theme(plot.title = element_text(hjust = 1))
```

134 **7|** Import metabolomic data

```
135 meta_mb <- read.csv("metabolism_meta.csv")
```

```
136 metabolism <- read.csv("metabolism_expr.csv",
```

```
137   row.names = 1,
```

```
138   check.name = FALSE)
```

139 **8|** Identify differentially abundance metabolites

```
140 groups <- c(CD = 'CD', UC = 'UC')
```

```
141 de_cpd <- lapply(groups, function(x) {
```

```
142   DA(expr = metabolism,
```

```
143   meta = meta_mb,
```

```
144   subset_group = c(x, 'Control'),
```

```
145   diff_group = x)
```

```
146 })
```

147 **9|** Perform enrichment analysis of differential metabolites

```
148 cpd_enrich_result <- compareCluster(geneClusters = de_cpd,
```

```
149   fun = "enrichKEGG",
```

```
150   organism = "cpd")
```

151 **10|** Visualize the functional enrichment results of differential metabolites (Figure 2B):

```
152 dotplot(cpd_enrich_result, facet = 'intersect', showCategory =
153 10,
```

```
154   split = "intersect", label_format = 60) +
```

```
155   ggtitle("Functional enrichment of chemical compounds") +
```

```
156   theme(plot.title = element_text(hjust = 1))
```

```
157
```

158 **Procedure 2: Transcription Factor Analysis Pertaining to Cold Tolerance in**

159 *Phyllostachys edulis* (PE)

160 Download data:

```
161 count_files <- c("counts.txt", "group_info.txt")
```

```
162 url <- "https://yulab-smu.top/clusterProfiler_protocol/examples
163 /Phyllostachys_heterocycla"
```

```
164
```

```
165 for (f in count_files) {
```

```
166   download.file(file.path(url, f), destfile = f)
```

```
167 }
```

```

168 annot_files <- c("regulation_from_motif_CE_Phe.txt",
169   "Phe_TF_list.txt", "Phe_GO_annotation.txt"
170 )
171 for (f in annot_files) {
172   download.file(file.path(url, "annot_data", f), destfile = f)
173 }

```

174 **1|** Load the R packages into the R environment:

```

175 library(DESeq2)
176 library(clusterProfiler)
177 library(ggplot2)
178 library(enrichplot)
179 library(aplot)
180 library(ggfun)

```

181 **2|** Read the expression matrix and sample grouping information:

```

182 counts <- read.delim("counts.txt")
183
184 group_info <- read.delim("group_info.txt")
185
186 group_info$group <- factor(group_info$group,
187   levels = c("0h", "2h", "24h", "168h")
188 )

```

189 **3|** Create a DESeqDataSet object:

```

190 count_dds <- DESeqDataSetFromMatrix(countData = counts,
191   colData = group_info,
192   design = ~ group)

```

193 4| Use DESeq to perform default differential expression analysis, including logarithmic
194 fold changes that incorporate data-driven prior distributions:

```
195 count_dds <- DESeq(count_dds)
```

196 5| Define groups (different time points) to be compared with the baseline group (i.e.,
197 0h):

```
198 time_points <- setNames(object = c("168h", "24h", "2h"),  
199 nm = c("168h_vs_0h", "24h_vs_0h", "2h_vs_0h"))
```

200 6| Extract log2FC values:

```
201 all_result <- lapply(time_points, function(time_point) {  
202   result <- results(count_dds, tidy = TRUE,  
203     contrast = c("group", time_point, "0h"))  
204   setNames(object = result$log2FoldChange, nm = result$row) |>  
205   sort(decreasing = TRUE)  
206 })
```

207 7| Prepare transcription factor annotation data:

```
208 tf_db <- read.delim(  
209   "regulation_from_motif_CE_Phe.txt", header = FALSE,  
210   colClasses = c("character", "NULL", "character", rep("NULL",  
211   4))  
212 ) |> setNames(c("TF", "targetGene"))
```

213 8| Use compareCluster to perform batch GSEA for identifying perturbed transcription
214 factors from each time point:

```
215 perturbed_TF_result <- compareCluster(all_result, fun = "GSEA",
```

```

216 pvalueCutoff = .05,
217 TERM2GENE = tf_db,
218 seed = 1234)

```

219 **9|** Visualize the enrichment result using dotplot:

```

220 perturbed_TF_plot <- dotplot(perturbed_TF_result,
221   showCategory = 25) +
222   aes(shape = I(22)) +
223   coord_flip() +
224   theme_minimal() +
225   theme_noxaxis() +
226   xlab(NULL) +
227   set_enrichplot_color(
228     c("#6C8FAD", "#84ADA7", "#C7B398"),
229     .fun = ggplot2::scale_fill_gradientn
230   )

```

231 **10 |** Construct a named list of TF target genes and extract a subset of the most
 232 significant TFs that are related to cold responses:

```

233 tf_id <- unique(get_plot_data(perturbed_TF_plot, "ID")[, 1])
234 tf_genes <- split(tf_db$targetGene, tf_db$TF)[tf_id]

```

235 **11|** Import the GO annotations of PE genes:

```

236 go_db <- read.delim(file = "Phe_GO_annotation.txt")

```

237 **12** | Characterize TF functions based on GO enrichment analysis of TF target genes:

```
238 TF_GO_result <- compareCluster(tf_genes, fun = "enricher",  
239   TERM2GENE = go_db[, c(2, 1)],  
240   TERM2NAME = go_db[, c(2, 3)])
```

241 **13** | Visualize TF function enrichment result:

```
242 TF_GO_plot <- dotplot(TF_GO_result, by = "count",  
243   showCategory = 3,  
244   label_format = 40) +  
245   theme_minimal() +  
246   theme(axis.text.x = element_text(vjust = 1, hjust = 1,  
247     angle = 30, size = 10)) +  
248   xlab(NULL) + ggtitle(NULL)
```

249 **14** | Import TF family annotation information:

```
250 tf_family <- read.delim("Phe_TF_list.txt", row.names = 1)
```

251 **15** | Prepare TF family data for visualization:

```
252 family_data <- subset(tf_family, Gene_ID %in% tf_id)  
253 family_data <- family_data[order(family_data$Family), ]  
254 family_data$Gene_ID <- factor(family_data$Gene_ID,  
255   levels = family_data$Gene_ID)
```

256 **16** | Visualize the TF family information:

```
257 tf_family_plot <- ggplot(data = family_data,  
258   aes(x = Gene_ID, y = 1, fill = Family)) +
```

```

259 geom_tile() +
260 scale_fill_discrete(type=c('#B3D3AA', '#4B6B5C', '#E88A71',
261 '#DEAB76', '#CD574D', '#85C1BF', '#BF38AE',
262 '#176D84', '#7D83B7', '#4040C6', '#994B41')) +
263 ggfun::theme_nothing()

```

264 **17|** Create a composite plot that integrates all the pieces of information (Figure 3):

```

265 insert_top(tf_family_plot, perturbed_TF_plot, height = 5) |>
266 insert_bottom(TF_GO_plot, height = 50)

```

267

268 **Procedure 3: Single-cell transcriptomic cell type annotation.**

269 Download data:

```

270 url <- "https://yulab-smu.top/clusterProfiler_protocol/examples
271 /single_cell/"
272 dir.create("hg19")
273 pbmc_files <- c("barcodes.tsv", "genes.tsv", "matrix.mtx")
274 for (f in pbmc_files) {
275   download.file(
276     file.path(url, "filtered_gene_bc_matrices/hg19/", f),
277     method = "auto",
278     destfile = file.path("hg19", f)
279   )
280 }
281 download.file(

```

```

282   file.path(url, "cell_marker_db/c8.all.v2023.1.Hs.symbols.gmt
283   "),
284   method = "auto",
285   destfile = "c8.all.v2023.1.Hs.symbols.gmt"
286 )

```

287 **1|** Load the R packages into the R environment:

```

288 library(Seurat)
289 library(CelliD)
290 library(clusterProfiler)
291 library(ggplot2)
292 library(ggrepel)
293 library(ggsc)

```

294 **2|** Import example PBMC data:

```

295 pbmc_counts <- Read10X(data.dir = "hg19")
296 pbmc <- CreateSeuratObject(counts = pbmc_counts, project = "pbm
297 c3k",
298   min.cells = 3)

```

299 **3|** Quality control by removing unwanted cells:

```

300 pbmc[["percent.mt"]] <- PercentageFeatureSet(pbmc, pattern = "^
301 MT-")
302 pbmc <- subset(pbmc,
303   subset = nFeature_RNA > 200 & nFeature_RNA < 2500 & percent.mt
304   < 5
305 )

```

306 **4|** Normalize the data:

```
307 pbmc <- NormalizeData(pbmc, normalization.method = "LogNormaliz  
308 e")  
309 pbmc <- ScaleData(pbmc)
```

310 **5|** Identify highly variable genes:

```
311 pbmc <- FindVariableFeatures(pbmc, selection.method = "vst",  
312 nfeatures = 2000)
```

313 **6|** Perform dimensional reduction by typing:

```
314 pbmc <- RunPCA(pbmc, features = VariableFeatures(object = pbm  
315 c))  
316 pbmc <- RunUMAP(pbmc, dims = 1:10)
```

317 **7|** Construct a K-nearest neighbor (KNN) graph:

```
318 pbmc <- FindNeighbors(pbmc, dims = 1:10)
```

319 **8|** Cluster cells into different groups:

```
320 pbmc <- FindClusters(pbmc, resolution = 0.5)
```

321 **9|** Find cluster biomarkers

```
322 pbmc <- RunMCA(pbmc)  
323 cluster_markers <- GetGroupGeneSet(X = pbmc, n.features = 20)
```

324 **10|** Import cell marker gene set:

```
325 cell_marker_db <- read.gmt("c8.all.v2023.1.Hs.symbols.gmt")
```

326 **11|** Perform cell type enrichment analysis by typing:

```
327 cell_type_enrich_result <- compareCluster(cluster_markers,
```

```
328 fun = "enricher", TERM2GENE = cell_marker_db
329 )
```

330 **12|** Predict cell type:

```
331 predict_cell_type <- function(enrich_result) {
332   enrich_result <- as.data.frame(enrich_result)
333   result <- split(
334     enrich_result, enrich_result$Cluster
335   ) |>
336   vapply(function(x) {
337     x$ID[which.min(x$p.adjust)]
338   }, FUN.VALUE = character(1))
339   cell_type <- gsub("_", " ", result) |>
340   yulab.utils::str_wrap(18)
341   names(cell_type) <- names(result)
342   return(cell_type)
343 }
344 cell_type_predict <- predict_cell_type(cell_type_enrich_result)
```

345 **13|** Assign cell type identity to clusters

```
346 pbmc <- RenameIdents(pbmc, cell_type_predict)
```

347 **14|** Visualize predicted cell type identity (Figure 4B):

```
348 cols <- c('#B3D3AA', '#E88A71', '#DEAB76', '#CD574D',
349   '#BF38AE', '#176D84', '#7D83B7', '#4040C6', '#994B41')
350 sc_dim(pbmc) +
```

```
351 sc_dim_geom_label(geom = ggrepel::geom_text_repel,  
352 color = "black", bg.color = "white") +  
353 scale_color_discrete(type=cols) +  
354 theme(legend.position = "none")
```

355