

Pigengene: Computing and using eigengenes

Habil Zare

Modified: 26 April, 2016. Compiled: July 4, 2026

Contents

1	Introduction	1
2	How to run <i>Pigengene</i> ?	2
2.1	Installation	2
2.2	A quick overview	2
2.3	What is an eigengene?	2
2.4	A toy example	3
2.5	Running the <i>Pigengene</i> pipeline step by step	7
2.6	Citation	14
3	Session Information	14

1 Introduction

Gene expression profiling technologies such as microarray or RNA Seq provide valuable datasets, however, inferring biological information from these data remains cumbersome. *Pigengene* address two challenges:

1. **Curse of dimensionality:** The number of features in an expression profile is usually very high. For instance, there are about 20,000 genes in human. In contrast, the number of samples (patients) is often very limited in practice, and may not exceed a few hundreds. Yeung et al. have shown that standard data reduction methods such as principal component analysis (PCA) are not appropriate to directly apply on gene expression data [1]. Instead, *Pigengene* addresses this challenge by applying PCA on gene modules.

2. **Normalization:** Data produced using different technologies, or in different labs, are not easily comparable. *Pigengene* identifies *eigengenes*, informative biological signatures that are robust with respect to the profiling platform. For instance, it can identify the signatures (compute the eigengenes) on microarray data, and infer them on biologically-related RNA Seq data. The resulting signatures are directly comparable even if the set of samples (patients) are independent and disjoint in the two analyzed datasets.

2 How to run *Pigengene* ?

2.1 Installation

Pigengene is an *R* package that can be downloaded and installed from *Bioconductor* by the following commands in *R*:

```
if (!requireNamespace("BiocManager", quietly=TRUE)) install.packages("BiocManager")
BiocManager::install("Pigengene")
```

Alternatively, if the source code is already available, the package can be installed by the following command in Linux:

```
R CMD INSTALL Pigengene_x.y.z.tar.gz
```

where x.y.z. determines the version. The second approach requires all the dependencies be installed manually, therefore, the first approach is preferred.

2.2 A quick overview

Pigengene identifies gene modules (clusters), computes an eigengene for each module, and uses these biological signatures as features for classification. The main function is `one.step.pigengene` which requires a gene expression profile and the corresponding conditions (types). Individual functions are also provided to facilitate running the pipeline in a customized way. The inferred biological signatures (eigengenes) are useful for supervised or unsupervised analyses.

2.3 What is an eigengene?

In most functions of this package, eigengenes are computed or used as robust biological signatures. Briefly, each eigengene $\mathcal{E}$ is a weighted average of the expression of all genes in a given set of n genes (also known as a gene module or a cluster of genes).

$$\mathcal{E} = \alpha_1 g_1 + \alpha_2 g_2 + \cdots + \alpha_n g_n, \quad \mathbf{1}$$

where α_i represents the weight corresponding to gene g_i . The weights are adjusted in a way that the explained variance is maximized. This guarantees that the loss in the biological information is minimized.

2.4 A toy example

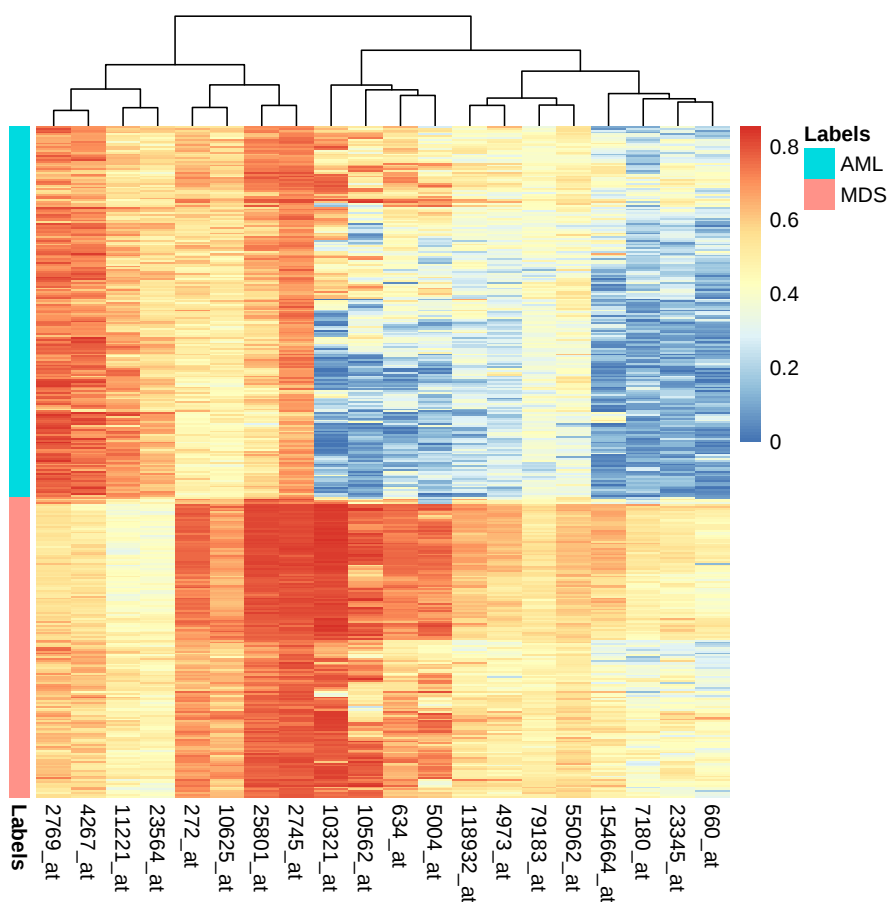
For a quick start, the application of *Pigengene* pipeline on some leukemia dataset is demonstrated below [2]. The first step is to load the package and data in R:

```
library(Pigengene)

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

Pigengene: Computing and using eigengenes

```
## Loading required package: BiocStyle
##
data(aml)
data(mds)
d1 <- rbind(aml,mds)
Labels <- c(rep("AML",nrow(aml)),rep("MDS",nrow(mds)))
names(Labels) <- rownames(d1)
Disease <- as.data.frame(Labels)
h1 <- pheatmap.type(d1[,1:20],annRow=Disease,show_rownames=FALSE)
```



Please note that the provided data in the package is sub-sampled for a quicker demonstration. For real applications, the expression of thousands of genes should be provided in order to the co-expression network analysis to be appropriate. It is common to first perform differential expression analysis, sort all the genes based on p-values, and use the top-third as the input [3]. Analyzing such input with *Pigengene* can take a few hours and may require 5-10 GB of memory. The following command runs Pigengene pipeline on the toy data:

Pigengene: Computing and using eigengenes

```
p1 <- one.step.pigengene(Data=d1,saveDir='pigengene', bnNum=0, verbose=1,
  seed=1, Labels=Labels, toCompact=FALSE, doHeat=FALSE)

## Pigengene started analizing 366 samples using 1000 genes...

## [1] "dataNum==1"

## Warning:  executing %dopar% sequentially:  no parallel backend registered

## Pigengene computation...

## Pigengene plots in:

## /tmp/RtmpJ1dzQG/Rbuild1e3730c5a89d/Pigengene/vignettes/pigengene/plots

## Making decision trees...

## costs:

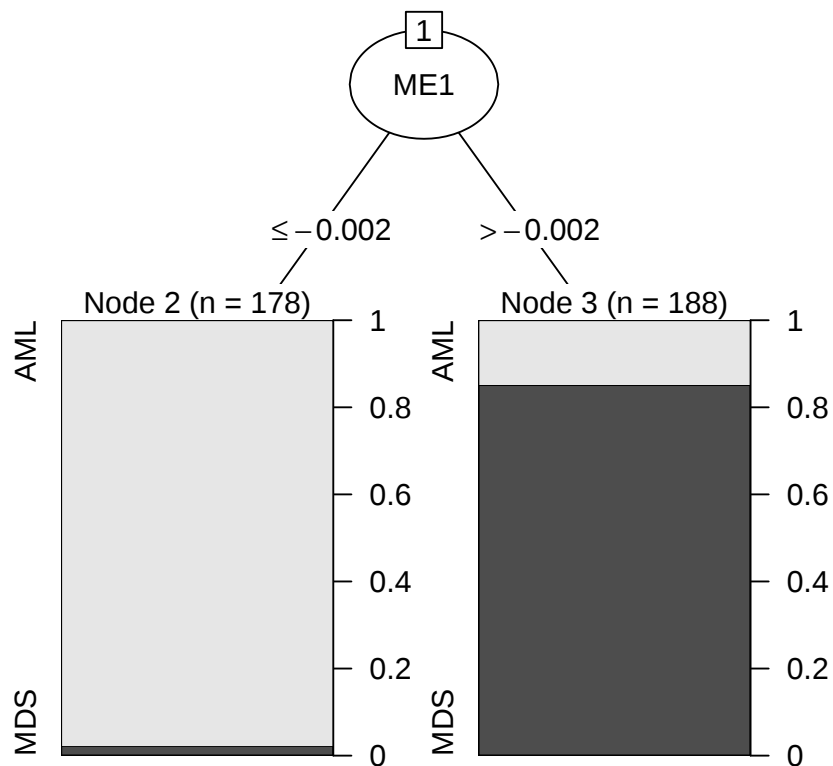
##      AML MDS
## AML    0   1
## MDS    1   0
```

Results and figures are saved in `pigengene` folder under the current directory. For more advanced applications, the user is encouraged to analyze the data step-by-step and customize the individual functions such as `compute.pigenegene` and `make.decision.tree`.

In addition to the provided decision trees, the user can also take alternative approaches to perform classification, clustering, survival analysis, etc. *using eigengenes as robust biological signatures (informative features)*. Eigengenes and other useful objects can be retrieved from the output. For instance, `c5treeRes` is a list containing the results of fitting decision trees to the eigengenes. As shown above, a couple of trees were fitted, one per a value for `minPerLeaf`. The following command plots the tree corresponding to 34, i.e., it was fitted requiring the minimum number of samples per every leaf to be at least 34.

```
plot(p1$c5treeRes$c5Trees[["34"]])
```

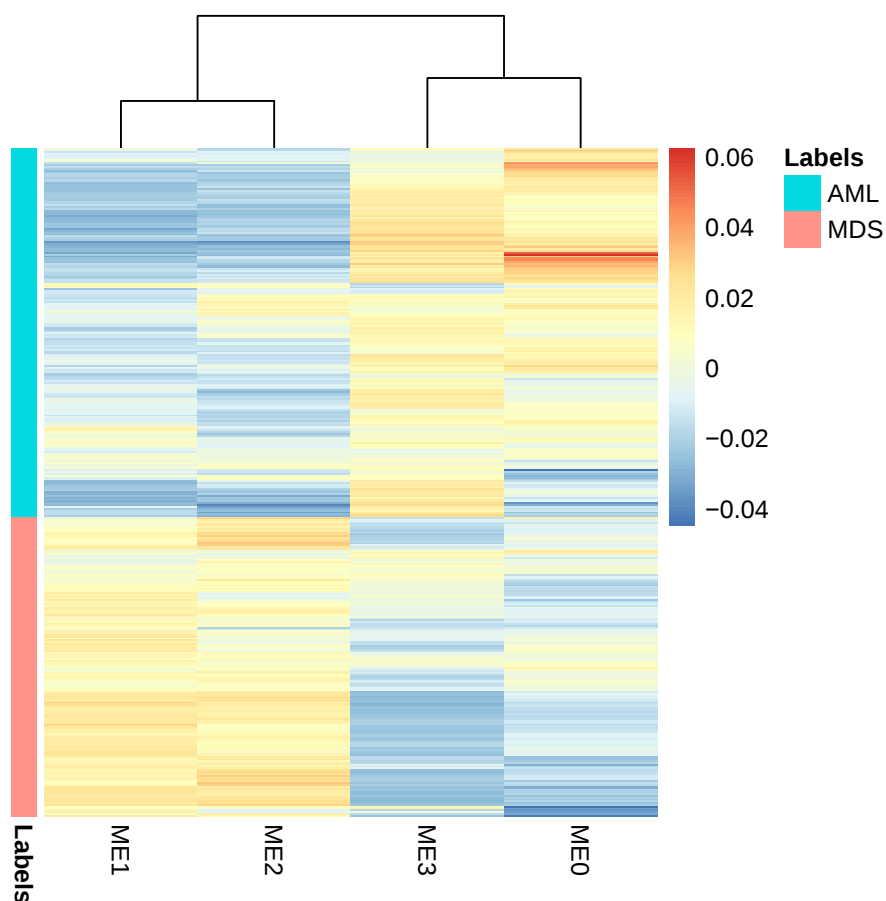
Pigengene: Computing and using eigengenes



The tree corresponding to other values are saved in `pigengene` folder. Of note, is the `pigenegene` object that contains the matrix of inferred eigengenes. Each row corresponds to a sample, and each column represents an eigengene.

```
dim(p1$pigengene$eigenenes)
## [1] 366 4

p1 <- pheatmap.type(p1$pigengene$eigenenes, annRow=Disease, show_rownames=FALSE)
```



2.5 Running the *Pigengene* pipeline step by step

If you are curious about the specific steps in the *Pigengene* pipeline, or you need to run some steps with different settings, you can follow the steps below. The results will be similar to the output of the `one.step.pigengene` function. The first step is quality control to make sure that the matrices have row and column names, and do not include too many NA values:

```
## QC:
checked <- check.pigengene.input(Data=d1, Labels=Labels)
DataI <- checked$Data
LabelsI <- checked$Labels
```

We oversample the data such that the number of cases in each condition is almost balanced.

```
wData <- balance(Data=DataI, Labels=LabelsI, verbose=1)$balanced
```

Pigengene: Computing and using eigengenes

```
## Balancing...  
## Oversampling to: 1818 of type AML , 1804 of type MDS ,
```

Weighted gene co-expression network analysis (WGCNA) [4] does not assume any cutoff (i.e., hard threshold) on the correlation values. Instead, it raises the correlation values to a power so that the correlation values that are close to zero diminish. This hyperparameter is called β , and can be estimated using as follows:

```
betaI <- calculate.beta(RsquaredCut=0.8, Data=wData, verbose=1)$power  
## Calculating beta...  
## beta: 9  
saveDir <- "steps" ## Results will be saved in this folder.  
dir.create(saveDir)
```

Once we have an estimate for β , WGCNA can be done in one step using the following function to identify gene modules (i.e., clusters):

```
## WGCNA  
wgRes <- wgcna.one.step(Data=wData, seed=1, power=betaI,  
  saveDir=saveDir, verbose=1)  
## Identifying the modules (WGCNA)...  
## power= 9  
## 4 modules were identified with the following sizes:  
## modules  
##    0    1    2    3  
##    1 440 333 226  
## save(net, file='steps/net.RData', compress=TRUE, ...)  
## save(wgOneStep, file='steps/wgOneStep.RData', compress=TRUE, ...)
```

The output of `wgcna.one.step` is a list of objects including `modules`, which is a numeric vector named with genes. Genes that map to the same numeric value are considered to be in the same module. We use this information to compute an eigengene for every module.

```
## Eigengenes:  
pigengene <- compute.pigengene(Data=DataI, Labels=LabelsI,  
  saveFile=file.path(saveDir, 'pigengene.RData'),
```

Pigengene: Computing and using eigengenes

```
modules=wgRes$modules, verbose=1)

## Pigengene computation...
## Pigengene plots in:
## steps/plots
class(pigengene)
## [1] "pigengene"
print(dim(pigengene$eigengenes)) ##This is the eigengenes matrix.
## [1] 366 4
print(pigengene$eigengenes[1:3,1:4])

##           ME1           ME2           ME3           ME0
## GSM376049 -0.006999271 -0.002890136 -0.006070692  0.004971246
## GSM376050 -0.003029049 -0.017728108  0.017265735 -0.001214384
## GSM376051 -0.004522860  0.008799329  0.003552811 -0.002341427
```

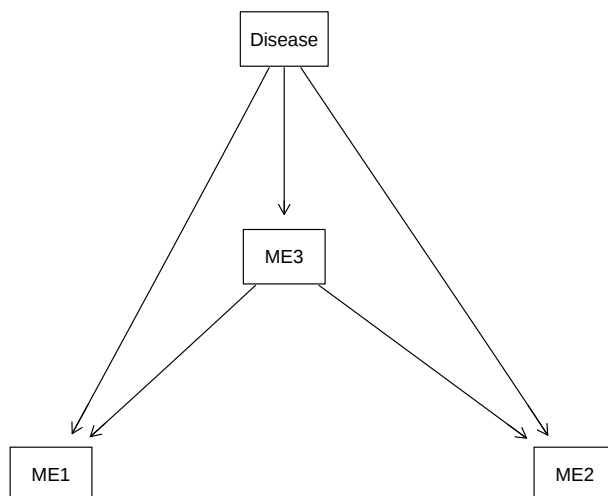
The number of columns in the eigengene matrix is equal to the number of modules, and the number of rows is equal to the number of samples.

Eigengenes can be used as robust biological signatures (i.e., features in machine learning terms) for classification, clustering, exploratory analysis, etc. For example, we can use them as random variables to fit a Bayesian network to data [5].

```
## Learning the BN structure:
library(bnlearn)
learnt <- learn.bn(pigengene=pigengene, bnPath=file.path(saveDir, "bn"),
                  bnNum=10, ## In real applications, at least 100-1000.
                  seed=1, verbose=1)

## learn.bn() with bnNum= 10 started at:
## 2026-07-04 00:33:41.171649

## Warning in discretize(as.data.frame(d), method = "interval", breaks
= length(unique(Disease))): at least one variable should be continuous
```



```
## learn.bn() took:  
## 7.49205 secs  
BN <- learnt$consensus1$BN
```

The `learn.bn` function has many arguments. See the corresponding documentation and publication [5] for technical details. Because usually thousands of individual networks are needed, it is wise to train many models in parallel on a cluster, which can be done with appropriate settings for the `learn.bn` input arguments. We can replot the consensus Bayesian network, which is already saved on disk, using the `draw.bn` function.

```
drawn <- draw.bn(BN)  
## By construction, the Disease node can have no parents.
```

The `learn.bn` function tries to find the best structure for the optimum Bayesian network. The conditional probability tables are yet to be inferred from the data.

```
## Fit the parameters of the Bayesian network:  
fit <- bn.fit(x=BN, data=learnt$consensus1$Data, method="bayes", iss=10)  
##where learnt$consensus1$Data is the discretized data matrix.  
  
## The conditional probability table for one of the children of the Disease node:  
selectedFeatures <- children("Disease", x=BN)  
print(fit[[selectedFeatures[1]]])  
  
##
```

```
## Parameters of node ME1 (multinomial distribution)
##
## Conditional probability table:
##
## , , Disease = AML
##
## ME3
## ME1 [-0.0294956, -0.00389314] (-0.00389314, 0.0151054]
## [-0.0361338, -0.018472] 0.219047619 0.139065817
## (-0.018472, 0.00862506] 0.476190476 0.826963907
## (0.00862506, 0.0308242] 0.304761905 0.033970276
## ME3
## ME1 (0.0151054, 0.0335374]
## [-0.0361338, -0.018472] 0.645833333
## (-0.018472, 0.00862506] 0.348039216
## (0.00862506, 0.0308242] 0.006127451
##
## , , Disease = MDS
##
## ME3
## ME1 [-0.0294956, -0.00389314] (-0.00389314, 0.0151054]
## [-0.0361338, -0.018472] 0.004642526 0.012437811
## (-0.018472, 0.00862506] 0.180129991 0.437810945
## (0.00862506, 0.0308242] 0.815227484 0.549751244
## ME3
## ME1 (0.0151054, 0.0335374]
## [-0.0361338, -0.018472] 0.119047619
## (-0.018472, 0.00862506] 0.547619048
## (0.00862506, 0.0308242] 0.333333333
```

The fitted Bayesian network can be used for predicting the labels (i.e., values of the Disease node).

```
l2 <- predict(object=fit, node="Disease", data=learnt$consensus1$Data, method="bayes-lw")
table(LabelsI, l2)

##      l2
## LabelsI AML MDS
##      AML 193  9
##      MDS  23 141
```

Pigengene: Computing and using eigengenes

Eigengenes can also be used in predictive models simpler than a Bayesian network. For example, a decision tree can be fitted using the `make.decision.tree` function [6].

```
## Decision trees:
treePath <- file.path(saveDir, 'C5Trees')
dir.create(path=treePath)
treeRes <- make.decision.tree(pigengene=pigengene, Data=DataI,
                             selectedFeatures=selectedFeatures, saveDir=treePath,
                             verbose=1, toCompact=FALSE)

## Making decision trees...

## costs:
##      AML MDS
## AML   0   1
## MDS   1   0

## Warning in normalizePath(dir): path[1]="/tmp/RtmpJ1dzQG/Rbuild1e3730c5a89d/Pigengene/v
No such file or directory

## Warning in normalizePath(dir): path[1]="/tmp/RtmpJ1dzQG/Rbuild1e3730c5a89d/Pigengene/v
No such file or directory

## Warning in normalizePath(dir): path[1]="/tmp/RtmpJ1dzQG/Rbuild1e3730c5a89d/Pigengene/v
No such file or directory
```

If `selectedFeatures="All"`, the `make.decision.tree` function automatically selects a “minimal” subset of eigengenes in order to prevent overfitting.

We can perform an over-representation pathway analysis based on the selected features in the decision tree using the `get.enriched.pw` function.

```
## Access tree results
usedFeatures <- get.used.features(c5Tree=treeRes$c5Trees[["17"]])
moduleMembers <- setNames(paste0("ME", wgRes$modules),
                          nm=sub("_[^_]+$", "", names(wgRes$modules)))
modMembersUsed <- moduleMembers[moduleMembers %in% usedFeatures]
moduList <- split(names(modMembersUsed), f=modMembersUsed)

library(org.Hs.eg.db)

## Loading required package: AnnotationDbi
## Loading required package: stats4
## 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)"'.

## Loading required package: IRanges

## Loading required package: S4Vectors

##
## Attaching package: 'S4Vectors'

## The following objects are masked from 'package:Rgraphviz':
##
##   from, to

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

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

##

pw1 <- get.enriched.pw(genes=moduList, idType="ENTREZID", pathwayDb="KEGG",
                      Org=NULL, OrgDb=org.Hs.eg.db, outPath=saveDir, verbose=1)

## Getting enriched pathways...

## Getting enriched pathways...

## Reading KEGG annotation online: "https://rest.kegg.jp/link/hsa/pathway"...
## Reading KEGG annotation online: "https://rest.kegg.jp/list/pathway/hsa"...
## kegg_category.rda is not found, download it online...

## No enriched pathways found using KEGG database

## Plots and an excel file are saved at: steps/ME1-enriched/enriched

## Getting enriched pathways...

## Loading required package: org.Mm.eg.db

##
```

```
## '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
## '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

## Plots and an excel file are saved at: steps/ME2_enriched/enriched
```

The output is saved for each selected module under the “moduleName_enrichment” folder. There is a subfolder that includes an excel file and plot(s). Each sheet in the excel file corresponds to a pathway database (KEGG in the above example). Each row is an overrepresented pathway.

2.6 Citation

The methodology and an interesting application of *Pigengene* on studying hematological malignancies is presented in the following reference [6].

```
citation("Pigengene")
```

To cite package 'Pigengene' in publications use:

Amir Froushani et al.(2016) Large-scale gene network analysis reveals the significance of extracellular matrix pathway and homeobox genes in acute myeloid leukemia: an introduction to the Pigengene package and its applications, Froushani et al., BMC Medical Genomics. URL: <https://bmcmmedgenomics.biomedcentral.com/articles/10.1186/s12920-017-0253-6>.

A BibTeX entry for LaTeX users is

@Article, author = Amir Froushani and et al., title = Large-scale gene network analysis reveals the significance of extracellular matrix pathway and homeobox genes in acute myeloid leukemia: an introduction to the Pigengene package and its applications, journal = BMC Medical Genomics, year = 2017, volume = 10, number = 1, pages = 16, month = 3,

3 Session Information

The output of `sessionInfo` on the system that compiled this document is as follows:

```
toLatex(sessionInfo())
```

- R version 4.6.1 (2026-06-24), x86_64-pc-linux-gnu
- Locale: LC_CTYPE=en_US.UTF-8, LC_NUMERIC=C, LC_TIME=en_US.UTF-8, LC_COLLATE=en_US.UTF-8, 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: Etc/UTC
- TZcode source: system (glibc)
- Running under: Ubuntu 26.04 LTS
- Matrix products: default
- BLAS: /usr/lib/x86_64-linux-gnu/openblas-pthread/libblas.so.3
- LAPACK:
/usr/lib/x86_64-linux-gnu/openblas-pthread/libopenblas-p0.3.32.so
; LAPACK version 3.12.0
- Base packages: base, datasets, graphics, grDevices, grid, methods, stats, stats4, utils
- Other packages: AnnotationDbi 1.75.0, Biobase 2.73.1, BiocGenerics 0.59.8, BiocStyle 2.41.0, bnlearn 5.1, generics 0.1.4, graph 1.91.0, IRanges 2.47.2, org.Hs.eg.db 3.23.1, org.Mm.eg.db 3.23.0, Pigengene 1.39.0, Rgraphviz 2.57.0, S4Vectors 0.51.5
- Loaded via a namespace (and not attached): aisdk 1.4.12, ape 5.8-1, aplot 0.3.0, backports 1.5.1, base64enc 0.1-6, BiocManager 1.30.27, Biostrings 2.81.3, bit 4.6.0, bit64 4.8.2, blob 1.3.0, buildtools 1.0.0, C50 0.2.0, cachem 1.1.0, callr 3.8.0, checkmate 2.3.4, cli 3.6.6, cluster 2.1.8.2, clusterProfiler 4.21.0, codetools 0.2-20, colorspace 2.1-2, compiler 4.6.1, crayon 1.5.3, Cubist 0.6.0, curl 7.1.0, data.table 1.18.4, DBI 1.3.0, digest 0.6.39, doParallel 1.0.17, DOSE 4.7.0, dplyr 1.2.1, dynamicTreeCut 1.63-1, enrichit 0.2.0, enrichplot 1.33.0, evaluate 1.0.5, farver 2.1.2, fastcluster 1.3.0, fastmap 1.2.0, fontBitstreamVera 0.1.1, fontLiberation 0.1.0, fontquiver 0.2.1, foreach 1.5.2, foreign 0.8-91, Formula 1.2-5, fs 2.1.0, gdata 3.0.1, gdtools 0.5.1, ggforce 0.5.0, ggfun 0.2.1, ggiraph 0.9.6, ggnewscale 0.5.2, ggplot2 4.0.3, ggplotify 0.1.3, ggraph 2.2.2, ggrepel 0.9.8, ggtangle 0.1.2, ggtree 4.3.0, glue 1.8.1, GO.db 3.23.1, GOSeqSim 2.39.2, graphite 1.59.0, graphlayouts 1.2.4, gridExtra 2.3.1, gridGraphics 0.5-1, gson 0.2.0,

gtable 0.3.6, gtools 3.9.5, highr 0.12, Hmisc 5.2-6, htmlTable 2.5.0, htmltools 0.5.9, htmlwidgets 1.6.4, httr 1.4.8, httr2 1.2.3, igraph 2.3.3, impute 1.87.0, inum 1.0-5, iterators 1.0.14, jsonlite 2.0.0, KEGGREST 1.53.4, knitr 1.51, labeling 0.4.3, lattice 0.22-9, lazyeval 0.2.3, libcoin 1.0-13, lifecycle 1.0.5, magrittr 2.0.5, maketools 1.3.2, MASS 7.3-65, Matrix 1.7-5, matrixStats 1.5.0, memoise 2.0.1, mvtnorm 1.4-1, nlme 3.1-169, nnet 7.3-20, openxlsx 4.2.8.1, otl 0.2.0, parallel 4.6.1, partykit 1.2-27, patchwork 1.3.2, pheatmap 1.0.13, pillar 1.11.1, pkgconfig 2.0.3, plyr 1.8.9, png 0.1-9, polyclip 1.10-7, preprocessCore 1.75.0, processx 3.9.0, ps 1.9.3, purrr 1.2.2, qvalue 2.45.0, R6 2.6.1, rappdirs 0.3.4, RColorBrewer 1.1-3, Rcpp 1.1.1-1.1, reactome.db 1.96.0, ReactomePA 1.57.0, reshape2 1.4.5, rlang 1.2.0, rmarkdown 2.31, rpart 4.1.27, RSQLite 3.53.3, rstudioapi 0.19.0, S7 0.2.2, scales 1.4.0, scatterpie 0.2.6, Seqinfo 1.3.0, splines 4.6.1, stringi 1.8.7, stringr 1.6.0, survival 3.8-6, sys 3.4.3, systemfonts 1.3.2, tibble 3.3.1, tidydr 0.0.6, tidygraph 1.3.1, tidyr 1.3.2, tidyselect 1.2.1, tidytree 0.4.8, tinytex 0.60, tools 4.6.1, treeio 1.37.0, tweenr 2.0.3, vctrs 0.7.3, viridis 0.6.5, viridisLite 0.4.3, WGCNA 1.74, withr 3.0.3, xfun 0.59, XVector 0.53.0, yaml 2.3.12, yulab.utils 0.2.4, zip 3.0.0

References

- [1] Ka Yee Yeung and Walter L. Ruzzo. Principal component analysis for clustering gene expression data. *Bioinformatics*, 17(9):763–774, 2001.
- [2] Ken I Mills, Alexander Kohlmann, P Mickey Williams, Lothar Wieczorek, Wei-min Liu, Rachel Li, Wen Wei, David T Bowen, Helmut Loeffler, Jesus M Hernandez, et al. Microarray-based classifiers and prognosis models identify subgroups with distinct clinical outcomes and high risk of aml transformation of myelodysplastic syndrome. *Blood*, 114(5):1063–1072, 2009.
- [3] Bin Zhang, Chris Gaiteri, Liviu-Gabriel Bodea, Zhi Wang, Joshua McElwee, Alexei A Podtelezhnikov, Chunsheng Zhang, Tao Xie, Linh Tran, Radu Dobrin, et al. Integrated systems approach identifies genetic nodes and networks in late-onset alzheimer’s disease. *Cell*, 153(3):707–720, 2013.
- [4] Peter Langfelder and Steve Horvath. Wgcna: an r package for weighted correlation network analysis. *BMC bioinformatics*, 9(1):559, 2008.

- [5] Rupesh Agrahari, Amir Foroushani, T Roderick Docking, Linda Chang, Gerben Duns, Monika Hudoba, Aly Karsan, and Habil Zare. Applications of bayesian network models in predicting types of hematological malignancies. *Scientific reports*, 8(1):6951, 2018.
- [6] A Foroushani, R Agrahari, R Docking, A Karsan, and H Zare. Large-scale gene network analysis reveals the significance of extracellular matrix pathway and homeobox genes in acute myeloid leukemia. *In preparation*.