

Reverse engineering transcriptional regulatory networks from gene expression microarray data using *qpgraph*

Robert Castelo¹ and Alberto Roverato²

October 7, 2025

1. Universitat Pompeu Fabra, Barcelona, Spain.

2. Università di Bologna, Bologna, Italy.

1 Introduction

This vignette describes how to use the package *qpgraph* in order to reverse engineer a transcriptional regulatory network from a particular gene expression microarray data set of *Escherichia coli* (*E. coli*). Concretely, the data corresponds to $n = 43$ experiments of various mutants under oxygen deprivation (Covert et al., 2004). The mutants were designed to monitor the response from *E. coli* during an oxygen shift in order to target the *a priori* most relevant part of the transcriptional network by using six strains with knockouts of the following key transcriptional regulators in the oxygen response: $\Delta arcA$, $\Delta appY$, Δfnr , $\Delta oxyR$, $\Delta soxS$ and the double knockout $\Delta arcA \Delta fnr$. To get started, load the following packages:

```
> library(Biobase)
> library(annotate)
> library(genefilter)
> library(org.EcK12.eg.db)
> library(graph)
> library(qpgraph)
```

Within the *qpgraph* package there is a data file called *EcoliOxygen* in which we will find the following objects stored:

```
> data(EcoliOxygen)
> ls()

[1] "annot"           "chr"
[3] "cross"           "eqtlnet.q0"
[5] "eqtlnet.q0.fdr"  "eqtlnet.q0.fdr.nrr"
[7] "eqtlnet.q0.fdr.nrr.sel" "eqtls"
[9] "filtered.regulon6.1" "gds680.eset"
[11] "genome"          "i"
[13] "j"               "map"
[15] "pMap"            "param"
[17] "sim.eqtl"        "subset.filtered.regulon6.1"
[19] "subset.gds680.eset"
```

where `filtered.regulon6.1` contains a subset of the *E. coli* transcriptional network from RegulonDB 6.1 ([Gama-Castro et al., 2008](#)) obtained through the filtering steps described in ([Castelo and Roverato, 2009](#)) and `gds680.eset` is an `ExpressionSet` object with the $n = 43$ microarray experiments of [Covert et al. \(2004\)](#) described before. These experiments provide expression profiles for $p = 4205$ genes derived from the original data set downloaded from the Gene Expression Omnibus ([Barrett et al., 2007](#)) with accession GDS680 by applying the filtering described in ([Castelo and Roverato, 2009](#)). You can see a summary of the data contained in this object by simply typing its name on the R-shell:

```
> gds680.eset

ExpressionSet (storageMode: lockedEnvironment)
assayData: 4205 features, 43 samples
  element names: exprs
protocolData: none
phenoData
  rowNames: GSM18235 GSM18236 ... GSM18289 (43 total)
  varLabels: Strain GrowthProtocol GenotypeVariation Description
  varMetadata: labelDescription
featureData: none
experimentData: use 'experimentData(object)'
  pubMedIds: 15129285
Annotation: org.EcK12.eg.db
```

where the usual probeset identifiers in the `featureNames` slot have been already replaced by the corresponding Entrez IDs according to the filtering steps taken in ([Castelo and Roverato, 2009](#)).

2 Preprocessing steps

In order to keep time and space requirements of the calculations at a manageable level for a vignette, we will use a subset of these data. Concretely, we will consider first those genes forming part in RegulonDB of the regulatory modules of the five knocked-out transcription factors and select the 100 genes with largest variability measured by the interquartile range (IQR). In the *qpgraph* package the filtered RegulonDB data is stored in the form of a data frame where each row corresponds to a transcriptional regulatory relationship, the first two columns contain Blattner IDs of the transcription factor (TF) and target (TG) genes, respectively, and the following two correspond to the same genes but specified by Entrez IDs. The fifth column contains the direction of the regulation according to RegulonDB and this is how the first rows look like:

```
> head(filtered.regulon6.1)

  BLID_TF BLID_TG EgID_TF EgID_TG Direction
1  b0464   b0463  945516  945112         -
2  b0464   b0462  945516  945108         -
5  b2213   b4187  946710  948710         +
6  b2213   b2068  946710  947371         +
7  b2213   b2212  946710  946708        +-
8  b4116   b4117  948627  948638         +
```

We select the rows of `filtered.regulon6.1` that correspond to the subnetwork of the 5 knocked-out TFs as follows. First, obtain the Entrez IDs of these genes from their symbols:

```
> knockoutsyms <- c("arcA", "appY", "oxyR", "soxS", "fnr")
> rmap <- revmap(getAnnMap("SYMBOL", "org.Eck12.eg.db"))
> knockoutEgIDs <- unlist(mget(knockoutsyms, rmap))
> knockoutEgIDs

      arcA      appY      oxyR      soxS      fnr
"948874" "948797" "948462" "948567" "945908"
```

Next, get all transcriptional regulatory relationships from these TFs and obtain the subset of non-redundant genes involved in this subnetwork:

```
> mt <- match(filtered.regulon6.1[, "EgID_TF"], knockoutEgIDs)
> cat("These 5 TFs are involved in", sum(!is.na(mt)), "TF-TG interactions\n")

These 5 TFs are involved in 462 TF-TG interactions

> genes02net <- as.character(unique(as.vector(
+       as.matrix(filtered.regulon6.1[!is.na(mt), c("EgID_TF", "EgID_TG")]))))
> cat("There are", length(genes02net), "different genes in this subnetwork\n")

There are 378 different genes in this subnetwork
```

and, finally, select the 100 most variable genes by using the IQR:

```
> IQRs <- apply(exprs(gds680.eset[genes02net,]), 1, IQR)
> largestIQRgenes02net <- names(sort(IQRs, decreasing=TRUE)[1:100])
```

Using these genes we create a new ExpressionSet object, which we shall call `subset.gds680.eset` by subsetting directly from `gds680.eset`:

```
> dim(gds680.eset)

Features Samples
   4205      43

> subset.gds680.eset <- gds680.eset[largestIQRgenes02net,]
> dim(subset.gds680.eset)

Features Samples
   100      43

> subset.gds680.eset

ExpressionSet (storageMode: lockedEnvironment)
assayData: 100 features, 43 samples
  element names: exprs
protocolData: none
phenoData
  rowNames: GSM18235 GSM18236 ... GSM18289 (43 total)
  varLabels: Strain GrowthProtocol GenotypeVariation Description
  varMetadata: labelDescription
featureData: none
experimentData: use 'experimentData(object)'
  pubMedIds: 15129285
Annotation: org.Eck12.eg.db
```

In order to compare later our results against the transcriptional network from RegulonDB we will extract the subnetwork that involves exclusively these selected 100 genes as follows. First extract the corresponding rows:

```
> mtTF <- match(filtered.regulon6.1[, "EgID_TF"], largestIQRgenes02net)
> mtTG <- match(filtered.regulon6.1[, "EgID_TG"], largestIQRgenes02net)
> cat(sprintf("The 100 genes are involved in %d RegulonDB interactions\n",
+   sum(!is.na(mtTF) & !is.na(mtTG))))
The 100 genes are involved in 128 RegulonDB interactions
> subset.filtered.regulon6.1 <- filtered.regulon6.1[!is.na(mtTF) & !is.na(mtTG),]
```

Next, we need to build an incidence matrix of this subset of interactions, which we shall call `subset.filtered.regulon6.1.I`, in order to ease posterior comparisons with reverse-engineered networks and for this purpose we should first map the Entrez IDs to the indexed position they have within the `ExpressionSet` object and then build the incidence matrix:

```
> TFi <- match(subset.filtered.regulon6.1[, "EgID_TF"],
+   featureNames(subset.gds680.eset))
> TGi <- match(subset.filtered.regulon6.1[, "EgID_TG"],
+   featureNames(subset.gds680.eset))
> subset.filtered.regulon6.1 <- cbind(subset.filtered.regulon6.1,
+   idx_TF=TFi, idx_TG=TGi)
> p <- dim(subset.gds680.eset)["Features"]
> subset.filtered.regulon6.1.I <- matrix(FALSE, nrow=p, ncol=p)
> rownames(subset.filtered.regulon6.1.I) <- featureNames(subset.gds680.eset)
> colnames(subset.filtered.regulon6.1.I) <- featureNames(subset.gds680.eset)
> idxTFTG <- as.matrix(subset.filtered.regulon6.1[, c("idx_TF", "idx_TG")])
> subset.filtered.regulon6.1.I[idxTFTG] <-
+   subset.filtered.regulon6.1.I[cbind(idxTFTG[,2], idxTFTG[,1])] <- TRUE
```

3 Reverse engineer a transcriptional regulatory network

We are set to reverse engineer a transcriptional regulatory network from the subset of the oxygen deprivation microarray data formed by the selected 100 genes and we will use three methods: 1. the estimation of Pearson correlation coefficients (PCCs); 2. the estimation of average non-rejection rates (avgNRRs); and, as a baseline comparison, 3. the assignment of random correlations drawn from a uniform distribution between -1 and +1 to every pair of genes. We can estimate PCCs for all gene pairs with the function `qpPCC` from the *qpgraph* package as follows:

```
> pcc.estimates <- qpPCC(subset.gds680.eset)
```

which returns a list with two members, one called `R` with the PCCs and another called `P` with the corresponding two-sided P-values for the null hypothesis of zero correlation. Let's take a look to the distribution of absolute PCCs between all possible TF-TG pairs in this subset of 100 genes:

```

> largestIQRgenes02net_i <- match(largestIQRgenes02net,
+                               featureNames(subset.gds680.eset))
> largestIQRgenes02netTFs <- largestIQRgenes02net[!is.na(
+                               match(largestIQRgenes02net, filtered.regulon6.1[, "EgID_TF"])]
> largestIQRgenes02netTFs_i <- match(largestIQRgenes02netTFs,
+                                   featureNames(subset.gds680.eset))
> TFsbyTGs <- as.matrix(expand.grid(largestIQRgenes02netTFs_i,
+                                   setdiff(largestIQRgenes02net_i, largestIQRgenes02netTFs_i)))
> TFsbyTGs <- rbind(TFsbyTGs, t(combn(largestIQRgenes02netTFs_i, 2)))
> summary(abs(pcc.estimates$R[TFsbyTGs]))

      Min.   1st Qu.   Median     Mean   3rd Qu.    Max.
0.0001023 0.3026803 0.5145512 0.5036736 0.7181577 0.9862190

```

Note that they are distributed almost uniformly at random throughout the entire range [0,1] while if we look at the distribution of the PCC estimates for the entire RegulonDB data, i.e., for all possible TF-TG pairs among the initial $p = 4205$ genes:

```

> regulonDBgenes <- as.character(unique(c(filtered.regulon6.1[, "EgID_TF"],
+                                       filtered.regulon6.1[, "EgID_TG"])))
> cat(sprintf("The RegulonDB transcriptional network involves %d genes",
+             length(regulonDBgenes)))

The RegulonDB transcriptional network involves 1428 genes

> pcc.allRegulonDB.estimates <- qpPCC(gds680.eset[regulonDBgenes,])
> allTFs_i <- match(unique(filtered.regulon6.1[, "EgID_TF"]), regulonDBgenes)
> allTFsbyTGs <- as.matrix(expand.grid(allTFs_i,
+                                       setdiff(1:length(regulonDBgenes), allTFs_i)))
> allTFsbyTGs <- rbind(allTFsbyTGs, t(combn(allTFs_i, 2)))
> summary(abs(pcc.allRegulonDB.estimates$R[allTFsbyTGs]))

      Min.   1st Qu.   Median     Mean   3rd Qu.    Max.
2.390e-06 1.038e-01 2.204e-01 2.555e-01 3.739e-01 9.862e-01

```

we see that, opposite to what happens in the subset of 100 genes, most of the absolute PCC values for all (i.e., present and absent from RegulonDB) TF-TG pairs are small. The high level of correlation among most of the 100 genes is probably due to the coordinated transcriptional program to which all these genes belong to, since they form part of some of the key regulatory modules in the response to oxygen deprivation. Recall that five TFs in these regulatory modules were knocked-out in the assayed experimental conditions and we selected the most variable 100 genes. Concretely, among the five TFs the following ones were finally included in these 100 most variable genes:

```

> mt <- match(knockoutEgIDs, largestIQRgenes02net)
> unlist(mget(largestIQRgenes02net[mt[!is.na(mt)]], org.Eck12.egSYMBOL))

948874 948797 948567 945908
"arcA" "appY" "soxS" "fnr"

```

If we look now to the distribution of absolute PCC values for only those TF-TG pairs that are present in the subset of RegulonDB involved in the 100 genes:

```
> maskRegulonTFTG <- subset.filtered.regulon6.1.I & upper.tri(subset.filtered.regulon6.1.I)
> summary(abs(pcc.estimates$R[maskRegulonTFTG]))
```

```
      Min.   1st Qu.   Median     Mean   3rd Qu.     Max.
0.0001023 0.1561148 0.2896006 0.3303612 0.4573992 0.9862190
```

they show much lower values ($50\% < 0.3$) and thus we can expect that a substantial number of TF-TG pairs absent from RegulonDB but with strong PCC values will sneak in as false positives in our assessment below of the estimation of PCCs as a reverse engineering method. If we look at the distribution of the PCC values from the RegulonDB interactions separately by each of the regulatory modules within these 100 genes (i.e., by each of the TFs) we can see that *fnr* is one of the responsables for having low PCCs in a large fraction of this subset of RegulonDB. We have used the R code below to produce Figure 1 where this is shown.

```
> par(mar=c(5,4,5,2))
> pccsbyTF <- list()
> for (TFi in subset.filtered.regulon6.1[, "idx_TF"])
+   pccsbyTF[[featureNames(subset.gds680.eset)[TFi]]] <-
+     abs(pcc.estimates$R[TFi, subset.filtered.regulon6.1.I[TFi,]])
> bp <- boxplot(pccsbyTF, names=sprintf("%s", mget(names(pccsbyTF), org.EcK12.egSYMBOL)),
+               ylab="Pearson correlation coefficient (PCC)",
+               main=paste("Distribution of PCCs in each RegulonDB",
+                           "regulatory module within the 100 genes data set", sep="\n"))
> nint <- sprintf("(%d)", sapply(names(pccsbyTF), function(x)
+                               sum(subset.filtered.regulon6.1.I[x,])))
> mtext(nint, at=seq(bp$n), line=+2, side=1)
> mtext("Transcription factor (# RegulonDB interactions)", side=1, line=+4)
```

As observed by [Covert et al. \(2004\)](#) when *fnr* becomes active under anaerobic conditions its mRNA level is significantly reduced and we hypothesize that this fact probably leads to weak correlations of the expression level with its target genes.

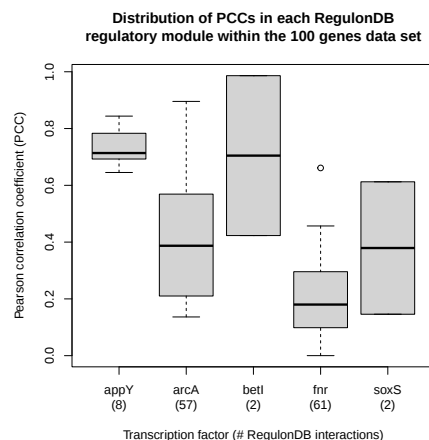


Figure 1: Distribution of Pearson correlation coefficients (PCCs) calculated from the [Covert et al. \(2004\)](#) oxygen deprivation data between genes forming RegulonDB interactions. Distributed values are shown separately by each regulatory module defined as a transcription factor (TF) and its set of target genes.

Now we will show how can we use qp-graphs to tackle such a challenging situation. We should start by estimating avgNRRs with the function `qpAvgNrr()` but before we do that, and for the sake of reproducibility of the results of this vignette, we should take into account that because the non-rejection rate is estimated by a random sampling procedure (see [Castelo and Roverato, 2006](#)), its value may vary slightly from run to run and thus edges with very similar avgNRR values may alternate their positions when ranking them and thus show up differently in different qp-graphs obtained from different runs if, within the ranking, they lie at the boundary of the precision threshold we may be using later. For this reason, and in order to let the reader reproduce exactly the results contained in this vignette, we will specify a particular seed to the random number generator as follows:

```
> set.seed(123)
```

Moreover, in this exercise, we are only interested in TF-TG relationships and thus we will speed-up the calculations by restricting the formation of gene pairs with the parameters `pairup.i` and `pairup.j` in the following way:

```
> avgnrr.estimate <- qpAvgNrr(subset.gds680.eset,  
+                             pairup.i=largestIQRgenes02netTFs,  
+                             pairup.j=largestIQRgenes02net, verbose=FALSE)
```

The function `qpAvgNrr()` uses by default four equidistant q-values along the available range and returns a matrix with the estimates for all gene pairs except when, as in this case, we restrict the genes allowed to pair with each other. In order to assess the accuracy of the PCC and qp-graph methods we will use the transcriptional regulatory relationships in the subset of RegulonDB that we selected before and calculate precision-recall curves ([Fawcett, 2006](#)) using the `qpPrecisionRecall` function from the *qpgraph* package.

We have to be careful with the fact that while we calculated avgNRRs only for TF-TG pairs, the matrix `pcc.estimate$R` contains PCC values for all pairs of genes and thus in order to obtain comparable precision-recall curves we will have to inform `qpPrecisionRecall` of the pairs that should be considered when giving it the matrix of PCC values. This is not necessary with avgNRRs as the matrix has NA values on the cells corresponding to pairs where no calculation was performed (on the pairs of non-transcription factor genes).

```
> pcc.prerec <- qpPrecisionRecall(abs(pcc.estimate$R), subset.filtered.regulon6.1.I,  
+                               decreasing=TRUE, pairup.i=largestIQRgenes02netTFs,  
+                               pairup.j=largestIQRgenes02net,  
+                               recallSteps=c(seq(0,0.1,0.01),seq(0.2,1,0.1)))
```

Note also that, opposite to PCCs, in avgNRR estimates the value indicating the smallest strength of the interaction is 1 instead of 0 and therefore we should set `decreasing=FALSE`:

```
> avgnrr.prerec <- qpPrecisionRecall(avgnrr.estimate, subset.filtered.regulon6.1.I,  
+                                  decreasing=FALSE,  
+                                  recallSteps=c(seq(0,0.1,0.01),seq(0.2,1,0.1)))
```

Finally, in order to have the assignment of random correlations as a baseline comparison we should do the following:

```
> set.seed(123)  
> rndcor <- qpUnifRndAssociation(100, featureNames(subset.gds680.eset))  
> random.prerec <- qpPrecisionRecall(abs(rndcor), subset.filtered.regulon6.1.I,  
+                                   decreasing=TRUE, pairup.i=largestIQRgenes02netTFs,
```

```
+ pairup.j=largestIQRgenes02net,
+ recallSteps=c(seq(0,0.1,0.01),seq(0.2,1,0.1)))
```

where again we have specified a seed for the random number generator in order to enforce reproducing the same random correlations each time we run this vignette.

A way to quantitatively compare these three precision-recall curves is to calculate the area under these curves where the larger it is, the more accurate the method is:

```
> f <- approxfun(pcc.prerec[,c("Recall","Precision")])
> area <- integrate(f,0,1)$value
> f <- approxfun(avgnrr.prerec[,c("Recall","Precision")])
> area <- cbind(area, integrate(f,0,1)$value)
> f <- approxfun(random.prerec[,c("Recall","Precision")])
> area <- cbind(area, integrate(f,0,1)$value)
> colnames(area) <- c("PCC", "avgNRR", "Random")
> rownames(area) <- "AreaPrecisionRecall"
> printCoefmat(area)
```

```
                PCC  avgNRR Random
AreaPrecisionRecall 0.13747 0.26436 0.189
```

From these values we may conclude that, for these data ($n = 43$ microarray experiments on $p = 100$ genes among which 7 are TFs, and with 128 transcriptional regulatory relationships from RegulonDB for comparison), the random method outperforms the usage of PCCs but it performs worse than the qp-graph method with avgNRRs which, therefore, constitutes the best solution among these three approaches. While it may sound a bit counter-intuitive that the assignment of a random correlation provides better results than using PCCs, the reason for this lies in the fact that with these data we have $7 \times 93 + \binom{7}{2} = 672$ possible TF-TG interactions out of which 128 from RegulonDB form our gold-standard. This yields a bottomline precision of $(128/672) \times 100 \approx 19\%$ which is quickly attained by drawing random correlations. However, we saw before that absolute PCCs of the RegulonDB interactions forming our gold-standard are most of them distributed under 0.5 and this yields, for this particular data set, a performance that is worse than random at regions of high-precision. We may see this situation depicted in Figure 2 whose left panel has been produced with the following R code:

```
> par(mai=c(.5,.5,1,.5),mar=c(5,4,7,2)+0.1)
> plot(avgnrr.prerec[,c(1,2)], type="b", lty=1, pch=19, cex=0.65, lwd=4, col="red",
+      xlim=c(0,0.1), ylim=c(0,1), axes=FALSE,
+      xlab="Recall (% RegulonDB interactions)", ylab="Precision (%)")
> axis(1, at=seq(0,1,0.01), labels=seq(0,100,1))
> axis(2, at=seq(0,1,0.10), labels=seq(0,100,10))
> axis(3, at=avgnrr.prerec[, "Recall"],
+      labels=round(avgnrr.prerec[, "Recall"]*dim(subset.filtered.regulon6.1)[1],
+                  digits=0))
> title(main="Precision-recall comparison", line=+5)
> lines(pcc.prerec[,c(1,2)], type="b", lty=1, pch=22, cex=0.65, lwd=4, col="blue")
> lines(random.prerec[,c(1,2)], type="l", lty=2, lwd=4, col="black")
> mtext("Recall (# RegulonDB interactions)", 3, line=+3)
> legend(0.06, 1.0, c("qp-graph", "PCC", "Random"), col=c("red", "blue", "black"),
+       lty=c(1,1,2), pch=c(19,22,-1), lwd=3, bg="white", pt.cex=0.85)
```

```

> par(mai=c(.5,.5,1,.5),mar=c(5,4,7,2)+0.1)
> plot(avgnrr.prerec[,c(1,2)], type="b", lty=1, pch=19, cex=0.65, lwd=4, col="red",
+      xlim=c(0,1), ylim=c(0,1), axes=FALSE,
+      xlab="Recall (% RegulonDB interactions)", ylab="Precision (%)")
> axis(1, at=seq(0,1,0.10), labels=seq(0,100,10))
> axis(2, at=seq(0,1,0.10), labels=seq(0,100,10))
> axis(3, at=avgnrr.prerec[, "Recall"],
+      labels=round(avgnrr.prerec[, "Recall"]*dim(subset.filtered.regulon6.1)[1],
+                  digits=0))
> title(main="Precision-recall comparison", line=+5)
> lines(pcc.prerec[,c(1,2)], type="b", lty=1, pch=22, cex=0.65, lwd=4, col="blue")
> lines(random.prerec[,c(1,2)], type="l", lty=2, lwd=4, col="black")
> mtext("Recall (# RegulonDB interactions)", 3, line=+3)
> legend(0.6, 1.0, c("qp-graph", "PCC", "Random"), col=c("red", "blue", "black"),
+      lty=c(1,1,2), pch=c(19,22,-1), lwd=3, bg="white", pt.cex=0.85)

```

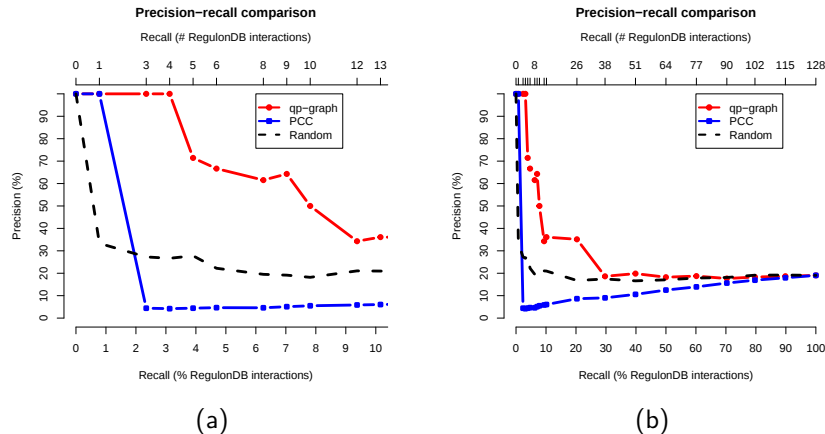


Figure 2: Comparison of precision-recall curves for various reverse-engineering methods with panel (a) showing a high-precision recall region of [0,0.1] and panel(b) showing the entire recall range.

The final step in this analysis is to get a transcriptional regulatory network from a qp-graph using avgNNRs and, if possible, obtain estimates of partial correlation coefficients (PAC) for the interactions. A qp-graph can be obtained by thresholding on the avgNNRs using the function `qpGraph`. When, as in our case now, we have a gold-standard network like RegulonDB, a sensible strategy to decide on a particular threshold is to derive it from a nominal precision level with respect to the gold-standard network. We can do this with the function `qpPRscoreThreshold` which reads the output of `qpPrecisionRecall` and takes a desired precision or recall level. We will use it with a nominal precision level of 50%:

```

> thr <- qpPRscoreThreshold(avgnrr.prerec, level=0.5, recall.level=FALSE, max.score=0)
> thr

ScoreThreshold
0.5065827

```

In order to manipulate the final reverse engineered transcriptional regulatory network from this 50%-precision qp-graph we will obtain a `graphNEL` object through the `qpGraph()` function:

Reverse engineering networks using *qpgraph*

```
> g <- qpGraph(avgnrr.estimated, threshold=thr, return.type="graphNEL")
> g

A graphNEL graph with undirected edges
Number of Nodes = 100
Number of Edges = 20
```

We are going to estimate now the corresponding PACs for the interactions. First, we should see if this is at all possible by calculating the size of the largest clique in this undirected graph with the `qpCliqueNumber` function from the *qpgraph* package:

```
> qpCliqueNumber(g, verbose=FALSE)

[1] 2
```

The maximum clique size (aka clique number) is smaller than the number of observations in the data ($n = 43$) and therefore we can go on with the PAC estimation (see [Lauritzen, 1996](#), for further details on this):

```
> pac.estimated <- qpPAC(subset.gds680.eset, g, verbose=FALSE)
```

Before making a graphical representation of the transcriptional regulatory network we have in `g` we would like to make a text-based summary of the interactions, more amenable for an occasional automatic processing of them outside R, including their presence or absence of RegulonDB and corresponding avgNRRs, PACs and PCCs. We start by building a matrix of the directed edges,

```
> edL <- edges(g)[names(edges(g))[unlist(lapply(edges(g), length)) > 0]]
> edM <- matrix(unlist(sapply(names(edL),
+                           function(x) t(cbind(x, edL[[x]])), USE.NAMES=FALSE)),
+              ncol=2, byrow=TRUE)
```

and continue by gathering all the necessary information on these edges,

```
> edSymbols <- cbind(unlist(mget(edM[,1], org.EcK12.egSYMBOL)),
+                   unlist(mget(edM[,2], org.EcK12.egSYMBOL)))
> idxTF <- match(edM[,1], featureNames(subset.gds680.eset))
> idxTG <- match(edM[,2], featureNames(subset.gds680.eset))
> nrrs <- avgnrr.estimated[cbind(idxTF, idxTG)]
> pacs.rho <- pac.estimated$R[cbind(idxTF, idxTG)]
> pacs.pva <- pac.estimated$P[cbind(idxTF, idxTG)]
> pccs.rho <- pcc.estimated$R[cbind(idxTF, idxTG)]
> pccs.pva <- pcc.estimated$P[cbind(idxTF, idxTG)]
> idxRegDB <- apply(edM, 1, function(x) {
+   regdbmask <-
+     apply(
+       cbind(match(subset.filtered.regulon6.1[, "EgID_TF"], x[1]),
+             match(subset.filtered.regulon6.1[, "EgID_TG"], x[2])),
+       1, function(y) sum(!is.na(y))) == 2 ;
+   if (sum(regdbmask) > 0)
+     (1:dim(subset.filtered.regulon6.1)[1])[regdbmask]
+   else
+     NA
+ })
```

Reverse engineering networks using *qpgraph*

```
+
+           })
> isinRegDB <- matrix(c("present","absent"),
+                     nrow=2, ncol=length(idxRegDB))[t(cbind(!is.na(idxRegDB),
+                     is.na(idxRegDB)))]
```

to end up creating a data frame that includes all the information,

```
> txregnet <- data.frame(RegulonDB=isinRegDB,
+                         RegDBdir=subset.filtered.regulon6.1[idxRegDB,"Direction"],
+                         AvgNRR=round(nrns,digits=2),
+                         PCC.rho=round(pccs.rho,digits=2),
+                         PCC.pva=format(pccs.pva,scientific=TRUE,digits=3),
+                         PAC.rho=round(pacs.rho,digits=2),
+                         PAC.pva=format(pacs.pva,scientific=TRUE,digits=3))
> rownames(txregnet) <- paste(edSymbols[,1],edSymbols[,2],sep=" -> ")
```

and which allows us to display the transcriptional regulatory network as a list of edges ordering them, for instance, by the avgNRR from the stronger (0.0) to the weaker (1.0) support for the presence of that interaction in the network:

```
> txregnet[sort(txregnet[["AvgNRR"]],index.return=TRUE)$ix,]
      RegulonDB RegDBdir AvgNRR PCC.rho PCC.pva PAC.rho PAC.pva
betB -> betI   absent   <NA>   0.09   0.99 0.00e+00   0.87 4.13e-06
betI -> betB   present    -   0.09   0.99 0.00e+00   0.87 4.13e-06
grcA -> fnr    absent   <NA>   0.11   0.66 1.39e-06   0.61 2.98e-04
fnr -> grcA    present   +-   0.11   0.66 1.39e-06   0.61 2.98e-04
appC -> appY   absent   <NA>   0.12   0.84 1.21e-12   0.50 2.35e-05
appY -> appC   present    +   0.12   0.84 1.21e-12   0.50 2.35e-05
arcA -> fadB   present    -   0.15  -0.90 5.23e-16  -0.74 1.22e-05
fadB -> arcA   absent   <NA>   0.15  -0.90 5.23e-16  -0.74 1.22e-05
fnr -> flu     absent   <NA>   0.19   0.32 3.90e-02   0.18 4.82e-02
flu -> fnr     absent   <NA>   0.19   0.32 3.90e-02   0.18 4.82e-02
flu -> appY    absent   <NA>   0.24   0.66 1.48e-06   0.27 3.04e-04
appY -> flu    absent   <NA>   0.24   0.66 1.48e-06   0.27 3.04e-04
appB -> appY   absent   <NA>   0.32   0.81 2.98e-11   0.45 3.45e-05
appY -> appB   present    +   0.32   0.81 2.98e-11   0.45 3.45e-05
arcA -> lpd    present    -   0.36  -0.64 4.58e-06  -0.17 4.33e-04
caiF -> narG   absent   <NA>   0.36   0.81 6.56e-11   0.57 3.84e-05
narG -> caiF   absent   <NA>   0.36   0.81 6.56e-11   0.57 3.84e-05
lpd -> arcA    absent   <NA>   0.36  -0.64 4.58e-06  -0.17 4.33e-04
caiF -> dmsC   absent   <NA>   0.42   0.86 1.63e-13   0.70 1.92e-05
dmsC -> caiF   absent   <NA>   0.42   0.86 1.63e-13   0.70 1.92e-05
hyaF -> appY   absent   <NA>   0.44   0.75 6.18e-09   0.36 8.12e-05
appY -> hyaF   present    +   0.44   0.75 6.18e-09   0.36 8.12e-05
hyaD -> appY   absent   <NA>   0.47   0.71 1.00e-07   0.32 1.48e-04
sucC -> betI   absent   <NA>   0.47   0.92 0.00e+00   0.35 8.62e-06
appY -> hyaD   present    +   0.47   0.71 1.00e-07   0.32 1.48e-04
betI -> sucC   absent   <NA>   0.47   0.92 0.00e+00   0.35 8.62e-06
arcA -> glcB   present    -   0.48  -0.78 4.83e-10  -0.38 5.18e-05
lldR -> glcB   absent   <NA>   0.48   0.75 8.14e-09   0.57 8.57e-05
lpd -> betI    absent   <NA>   0.48   0.89 3.11e-15   0.26 1.38e-05
```

betI -> lpd	absent	<NA>	0.48	0.89	3.11e-15	0.26	1.38e-05
glcB -> arcA	absent	<NA>	0.48	-0.78	4.83e-10	-0.38	5.18e-05
glcB -> lldR	absent	<NA>	0.48	0.75	8.14e-09	0.57	8.57e-05
gadA -> appY	absent	<NA>	0.50	0.43	4.47e-03	0.15	1.03e-02
fnr -> sufC	absent	<NA>	0.50	-0.43	4.45e-03	-0.32	1.03e-02
sufC -> fnr	absent	<NA>	0.50	-0.43	4.45e-03	-0.32	1.03e-02
appY -> gadA	absent	<NA>	0.50	0.43	4.47e-03	0.15	1.03e-02
fdhF -> betI	absent	<NA>	0.51	-0.73	2.57e-08	-0.16	1.09e-04
hyaE -> appY	absent	<NA>	0.51	0.71	8.68e-08	0.32	1.43e-04
appY -> hyaE	present	+	0.51	0.71	8.68e-08	0.32	1.43e-04
betI -> fdhF	absent	<NA>	0.51	-0.73	2.57e-08	-0.16	1.09e-04

We can plot the network with the function `qpPlotNetwork` as follows and obtain the result shown in Figure 3.

```
> qpPlotNetwork(g, pairup.i=largestIQRgenes02netTFs, pairup.j=largestIQRgenes02net,
+               annotation="org.Eck12.eg.db")
```

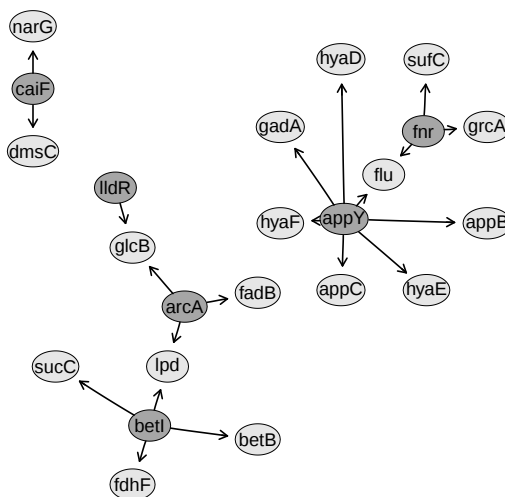


Figure 3: Reverse-engineered transcriptional network using a qp-graph at a nominal 50% precision.

4 Session Information

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

- R version 4.5.1 Patched (2025-08-23 r88802), x86_64-pc-linux-gnu
- Locale: LC_CTYPE=en_US.UTF-8, LC_NUMERIC=C, LC_TIME=en_GB, LC_COLLATE=C, LC_MONETARY=en_US.UTF-8, LC_MESSAGES=en_US.UTF-8, LC_PAPER=en_US.UTF-8, LC_NAME=C, LC_ADDRESS=C, LC_TELEPHONE=C, LC_MEASUREMENT=en_US.UTF-8, LC_IDENTIFICATION=C
- Time zone: America/New_York

- TZcode source: system (glibc)
- Running under: Ubuntu 24.04.3 LTS
- Matrix products: default
- BLAS: /home/biocbuild/bbs-3.22-bioc/R/lib/libRblas.so
- LAPACK: /usr/lib/x86_64-linux-gnu/lapack/liblapack.so.3.12.0
- Base packages: base, datasets, grDevices, graphics, grid, methods, stats, stats4, utils
- Other packages: AnnotationDbi 1.71.1, Biobase 2.69.1, BiocGenerics 0.55.1, IRanges 2.43.5, Rgraphviz 2.53.0, S4Vectors 0.47.4, Seqinfo 0.99.2, XML 3.99-0.19, annotate 1.87.0, genefilter 1.91.0, generics 0.1.4, graph 1.87.0, org.EcK12.eg.db 3.21.0, qpgraph 2.43.1, qtl 1.70
- Loaded via a namespace (and not attached): BiocIO 1.19.0, BiocManager 1.30.26, BiocParallel 1.43.4, BiocStyle 2.37.1, Biostrings 2.77.2, DBI 1.2.3, DelayedArray 0.35.3, GenomicAlignments 1.45.4, GenomicFeatures 1.61.6, GenomicRanges 1.61.5, KEGGREST 1.49.1, Matrix 1.7-4, MatrixGenerics 1.21.0, R6 2.6.1, RCurl 1.98-1.17, RSQLite 2.4.3, Rsamtools 2.25.3, S4Arrays 1.9.1, SparseArray 1.9.1, SummarizedExperiment 1.39.2, XVector 0.49.1, abind 1.4-8, bit 4.6.0, bit64 4.6.0-1, bitops 1.0-9, blob 1.2.4, cachem 1.1.0, cli 3.6.5, codetools 0.2-20, compiler 4.5.1, crayon 1.5.3, curl 7.0.0, digest 0.6.37, evaluate 1.0.5, fastmap 1.2.0, htmltools 0.5.8.1, httr 1.4.7, knitr 1.50, lattice 0.22-7, matrixStats 1.5.0, memoise 2.0.1, mvtnorm 1.3-3, parallel 4.5.1, pkgconfig 2.0.3, png 0.1-8, restfulr 0.0.16, rjson 0.2.23, rlang 1.1.6, rmarkdown 2.30, rtracklayer 1.69.1, splines 4.5.1, survival 3.8-3, tools 4.5.1, vctrs 0.6.5, xfun 0.53, xtable 1.8-4, yaml 2.3.10

References

- Barrett, T., Troup, D. B., Wilhite, S. E., Ledoux, P., Rudnev, D., Evangelista, C., Kim, I. F., Soboleva, A., Tomashevsky, M., and Edgar, R. (2007). NCBI GEO: mining tens of millions of expression profiles—database and tools update. *Nucleic Acids Res*, 35(Database issue):D760–5.
- Castelo, R. and Roverato, A. (2006). A robust procedure for gaussian graphical model search from microarray data with p larger than n . *J Mach Learn Res*, 7:2621–2650.
- Castelo, R. and Roverato, A. (2009). Reverse engineering molecular regulatory networks from microarray data with qp-graphs. *J Comput Biol*, 16(2):213–27.
- Covert, M. W., Knight, E. M., Reed, J. L., Herrgard, M. J., and Palsson, B. O. (2004). Integrating high-throughput and computational data elucidates bacterial networks. *Nature*, 429(6987):92–96.
- Fawcett, T. (2006). An introduction to ROC analysis. *Pattern Recogn Lett*, 27:861–874.
- Gama-Castro, S., Jimenez-Jacinto, V., Peralta-Gil, M., Santos-Zavaleta, A., Penaloza-Spinola, M. I., Contreras-Moreira, B., Segura-Salazar, J., Muniz-Rascado, L., Martinez-Flores, I., Salgado, H., Bonavides-Martinez, C., Abreu-Goodger, C., Rodriguez-Penagos, C., Miranda-Rios, J., Morett, E., Merino, E., Huerta, A. M., Trevino-Quintanilla, L., and

Collado-Vides, J. (2008). RegulonDB (version 6.0): gene regulation model of Escherichia coli K-12 beyond transcription, active (experimental) annotated promoters and textpresso navigation. *Nucleic Acids Res*, 36(Database issue):D120–4.

Lauritzen, S. (1996). *Graphical models*. Oxford University Press.