

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

Mattia Chiesa¹, Giada Maioli², and Luca Piacentini¹

¹Immunology and Functional Genomics Unit, Centro Cardiologico Monzino, IRCCS, Milan, Italy;

²Università degli Studi di Pavia, Pavia, Italy

April 15, 2025

Abstract

Feature selection aims to identify and, remove redundant, irrelevant and noisy variables from high-dimensional datasets. Selecting informative features affects the subsequent classification and regression analyses by improving their overall performances. Several methods have been proposed to perform feature selection: most of them relies on univariate statistics, correlation, entropy measurements or the usage of backward/forward regressions.

Herein, we propose an efficient, robust and fast method that adopts stochastic optimization approaches for high-dimensional. Genetic algorithms, a type of evolutionary algorithms, are often used to find solutions for optimization and search problems and promise to be effective on complex data. They operate on a population of potential solutions and apply the “principle of survival of the fittest” to produce the better approximation of the optimal solution.

GARS is an innovative implementation of a genetic algorithm that selects robust features in high-dimensional and challenging datasets.

Package

GARS 1.28.0

Contents

- 1 Introduction 3
 - 1.1 Citation and code. 4
- 2 Using GARS: a classification analysis 5
 - 2.1 The testing RNA-Seq dataset 5
 - 2.2 Launch GARS 5
 - 2.3 Test the robustness of the feature set 10
 - 2.4 Find the best features set 13
- 3 Build your custom GA 15
- 4 Session Info 15

1 Introduction

A crucial step in a Data Mining analysis is Feature Selection, which is the process of identifying the most informative predictors to build accurate classification models. Indeed, many features could inadvertently introduce bias in a prediction model, lead to overfitting and increase the complexity in further downstream analysis. In last decades, several methods have been proposed to perform feature selection; they are usually grouped in three main classes [1, 2]:

- **Filter Methods** - A measure based on statistics or entropy is used to rank and then select variables. The most popular filter methods are the *Information Gain*, the *ReliefF*, the *Gini Index* and the χ^2 test;
- **Wrapper Methods** - The set of important features is identified using a classifier to evaluate the accuracy of several combinations of variables. *Backward/Forward/Stepwise-elimination* strategies belong to this category;
- **Embedded Methods** - As for wrapper methods, a classifier is used to identify the best set of variables; however, in this case the procedure to identify the features is totally joined with the construction of the classifier. The “tree-based” classifiers (e.g. *Decision trees* and the *Random Forest*) are the widely used embedded methods.

Genetic Algorithms (GAs) are heuristic adaptive search algorithms that simulate the Darwinian law of “the survival of the fittest” among individuals, over consecutive generations, for solving hard problems, such as pattern recognition and feature selection. A GA is traditionally composed of three main consecutive stages: first, a random set of candidate solutions, *i.e.* chromosomes, is generated. Then, each chromosome is evaluated by a custom score, *i.e.* *fitness function* that reflects how good a solution is. Finally, the evolutionary operators are sequentially applied to the entire population: *Selection*, *Crossover* and *Mutation*. To find the optimal solution, this process has to be repeated several times: the starting chromosome population of a certain generation corresponds to the resulting chromosome population of the previous generation.

The idea to use a GA to perform Feature Selection is not novel; however, all the developed GA-based methods need a classifier to evaluate the goodness of a set of features (namely, they belong to the “Wrapper” or the “Embedded” category). The classifiers mainly used to assess the selected features are the Support Vector Machines (GA-SVM) [3], the k-Nearest Neighbours (GA-KNN) [4], the Random Forest (rfGA, see the [caret](#) package for details), the LDA (caretGA, see the [caret](#) package for details) [5] and maximum-likelihood based methods (GA-MLHD) [6].

One of the most relevant contexts where the feature selection is becoming more and more essential, is the -OMICS field: in fact, datasets coming from genomics, transcriptomics, proteomics and metabolomics experiments are typically composed of a large number of features compared to the sample size; this poses a big challenge for a data mining analysis.

In this context, we developed an innovative implementation of a **Genetic Algorithm** that selects **Robust Subsets** of features (**GARS**) in high-dimensional and challenging datasets. GARS has several benefits:

1. it does not need any classifier to evaluate the goodness of the selected feature. The fitness is calculated by the averaged Silhouette Index [7], after computing a Multi-Dimensional Scaling of the data. This allows being less prone to overfitting and local optima;
2. it is relatively fast, when the number of features is relatively high (tens of thousands);
3. it can be used in multi-class classification problems;

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

4. even though it has been thought for solving -OMICs tasks, it can be easily used in several other contexts (e.g. low-dimensional data);
5. it can be easily integrated with other R and Bioconductor packages for Data Mining (e.g. [caret](#) and [DaMiRseq](#)).

1.1 Citation and code

Users can find useful details about the GARS algorithm, [downloading the publication](#) ¹. This paper should be used to cite GARS, as well.

In addition, we provided in GitHub the code of all the analyses performed for the publication: https://github.com/BioinfoMonzino/GARS_paper_Code. Users are invited to download and customize the proposed workflows in which GARS is used to perform feature selection in three different machine learning analyses.

¹See: Chiesa et al. *GARS: Genetic Algorithm for the identification of a Robust Subset of features in high-dimensional datasets* [8].

2 Using GARS: a classification analysis

2.1 The testing RNA-Seq dataset

The dataset used in this vignette comes from a miRNA-Seq experiment performed on cervical tissues [9]; the dataset is composed of 714 miRNAs and 58 samples: 29 Tumor (T) and 29 Non-Tumor (N) cervical samples, respectively. In order to obtain a normalized gene expression matrix, we used the `DaMiR.normalization` function of the *DaMiRseq* package with default parameters.

```
library(MLSeq)
library(DaMiRseq)
library(GARS)

# load dataset
filepath <- system.file("extdata/cervical.txt", package = "MLSeq")
cervical <- read.table(filepath, header=TRUE)

# replace "wild-card" characters with other characters
rownames(cervical) <- gsub("*", "x", rownames(cervical), fixed = TRUE)
rownames(cervical) <- gsub("-", "_", rownames(cervical), fixed = TRUE)

# create the "class" vector
class_vector <- data.frame(gsub('[0-9]+', '', colnames(cervical)))
colnames(class_vector) <- "class"
rownames(class_vector) <- colnames(cervical)

# create a Summarized Experiment object
SE_obj <- DaMiR.makeSE(cervical, class_vector)

## Your dataset has:
## 714 Features;
## 58 Samples, divided in:
## 1 variables: class ;
## 'class' included.

# filter and normalize the dataset
datanorm <- DaMiR.normalization(SE_obj)

## 545 Features have been filtered out by expression. 169 Features remained.
## 0 'Hypervariant' Features have been filtered out. 169 Features remained.
## Performing Normalization by 'vst' with dispersion parameter: parametric
```

2.2 Launch GARS

After filtering and normalizing data we got a dataset with 161 expressed miRNAs and 58 samples. The best way to use *GARS* for selecting a robust set of features from an high-dimensional dataset is to exploit the wrapper function `GARS_GA`. A dataset must be provided, as well as a vector containing the class information. *GARS* gives the opportunity to provide input data in the form of `SummarizedExperiment`, `matrix` or `data.frame` objects. On one

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

hand, `SummarizedExperiment` is preferred in the case of a RNA-Seq experiment; on the other hand, a `data.frame` allows the user to integrate expression data with other numerical and/or categorical features.

In addition, several other parameters have to be set:

- `chr.num` - The number of chromosomes in each population. If the number of chromosomes is too small, the GA will explore a small part of the “solution space” (each chromosome is a candidate solution); conversely, if the number is too high, the GA will produce results very slowly. Default: 1000
- `chr.len` - The length of each chromosome. **This argument is the most important in GARS: it corresponds to the length of the desired feature set.** Usually, in data mining analysis the number of features, needed to build a classification model, has to be much smaller than the number of observations (i.e. samples).
- `generat` - The maximum number of generations. This number is usually high (hundreds to thousands): the higher the number of generations, the higher the probability to reach the best solution. Default: 500
- `co.rate` - The probability to perform the crossover for each random couple of chromosomes. This parameter allows the evolution rate to be controlled and tuned. Default: 0.8
- `mut.rate` - The probability to mutate each chromosome base. This parameter allows the evolution rate to be controlled and tuned. Default: 0.01
- `n.elit` - The number of best chromosomes that must be “preserved from the evolution”. This number is usually small compared to the number of chromosomes, in order to enhance the evolution. Default: 10
- `type.sel` - The algorithm that performs the Selection step. “*Roulette Wheel*” and “*Tournament*” selections are implemented. Default: *Roulette Wheel*.
- `type.co` - The algorithm that performs the Crossover step. “*One point*” and “*Two points*” crossover are implemented. Default: *One point*.
- `type.one.p.co` - In the case of “*One point*” crossover, this argument allows setting the quartile where the crossover has to be applied. The user can choose among the first, the second and the third quartile. Default: First quartile.
- `n.gen.conv` - The maximum number of consecutive generations with the same maximum fitness score. When the maximum fitness scores are the same for several generation, this means that the GA found the optimal solution (i.e. reached the convergence). This argument is useful to stop GARS when the convergence is reached. Default: 80.
- `plots` - Whether generating plots or not. Default: yes.
- `verbose` - Whether printing information in the console or not. Default: yes.
- `n.Feat_plot` - If `plots = yes`, the number of features to be plotted by `GARS_PlotFeaturesUsage`

To speed up the execution time of the function, here we set `generat = 20`, `chr.num = 100` and `chr.len = 8`; however, for a typical -omic experiment (thousands of features), this is probably not sufficient to find the best feature set. We strongly recommend to set accurately each parameter, trying different combinations of them (especially `chr.len`, See Section 2.4).

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

```
set.seed(123)
res_GA <- GARS_GA(data=datanorm,
                  classes = colData(datanorm),
                  chr.num = 100,
                  chr.len = 8,
                  generat = 20,
                  co.rate = 0.8,
                  mut.rate = 0.1,
                  n.elit = 10,
                  type.sel = "RW",
                  type.co = "one.p",
                  type.one.p.co = "II.quart",
                  n.gen.conv = 150,
                  plots="no",
                  verbose="yes")

## GARS has been set with these parameters:
## Number of starting features: 169
## Number of samples: 58
## Number of classes: 2
## Number of chromosomes: 100
## Length of chromosomes (i.e. number of desired features): 8
## Number of maximum generations: 20
## Crossing-Over rate: 0.8
## Mutation rate: 0.1
## Number of chromosomes kept by elitism: 10
## Type of Selection method: RW
## Type of CrossingOver method: one.p
## Position of the one-Point Crossover: II.quart
## Number of max generations allowed with the same Fitness: 150
## Produce graphs: no
##
## #####
## ##### GARS is running #####
## #####
##
## Reached 10 iterations...
## Reached 20 iterations...
## GARS found a solution after 20 iterations.
## With these parameters, the best solution:
##
## 1. is reached after 20 iterations;
## 2. is reached looking at the 100 % of the 169 features;
## 3. got a maximum fitness score = 0.57
## 4. is composed of the following features:
## miR_25 miR_181a miR_145x miR_7 miR_21x miR_203 miR_148a miR_143
```

The results of `GARS_GA` are stored in a `GarsSelectedFeatures` object, herein `res_GA`, where the informations could be extracted by 4 Assessor methods:

- `MatrixFeatures()` - Extracts the `matrix` containing the expression values for the selected features;

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

- `LastPop()` - Extracts the `matrix` containing the chromosome population of the last generation. The first column of this matrix represent the best solution, found by the GA;
- `AllPop()` - Extracts the `list` containing all the populations produced over the generations;
- `FitScore()` - Extracts the `vector` containing the maximum fitness scores, computed in each generation.

The information stored in the `GarsSelectedFeatures` object could be used for downstream analysis (See Section 2.3) or for generating plots (before, we set `plots = no`). In GARS the functions `GARS_PlotFitnessEvolution()` allows the user to plot the fitness evolution over the generations, while the function `GARS_PlotFeaturesUsage()` allows representing the frequency of each feature in a bubble chart:

```
# Plot Fitness Evolution
fitness_scores <- FitScore(res_GA)
GARS_PlotFitnessEvolution(fitness_scores)
#Plot the frequency of each features over the generations
Allfeat_names <- rownames(datanorm)
Allpopulations <- AllPop(res_GA)
GARS_PlotFeaturesUsage(Allpopulations,
                        Allfeat_names,
                        nFeat = 10)
```

As mentioned before, in this example the number of chromosomes (100) and the number of generations (20) were intentionally small. Nevertheless, the population evolved over the generations: indeed, as shown in Figure 1, the maximum fitness score is equal to 0.41 in the first generation and reaches the value of 0.55 in the last generation (an increasing of 30%). Moreover, the Figure 2 shows the most recurring (i.e. “conserved”) miRNAs over the iterations.

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

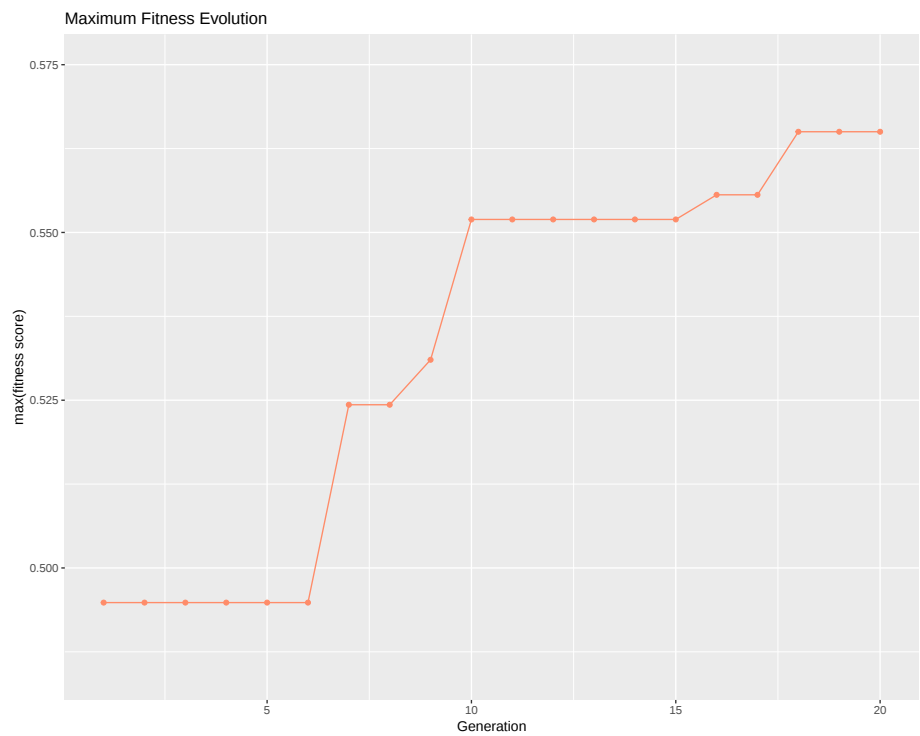


Figure 1: Fitness Evolution plot. The plot shows the evolution of the maximum fitness across the generations

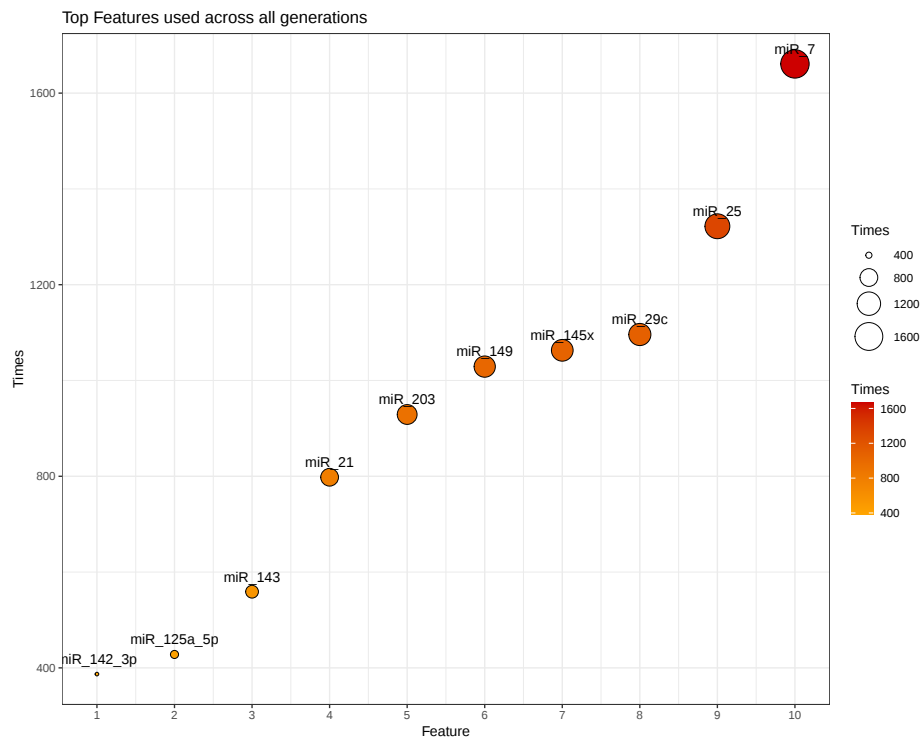


Figure 2: Recurring Features. Each circle in the plot represents a feature. The color and size of each circle are, respectively, darker and bigger when a feature is more recurring.

2.3 Test the robustness of the feature set

Besides the maximum fitness score of the last population, we can assess the quality of the results, through a classification analysis. To perform this task, we used the functions implemented in the *DaMiRseq* package, which offers several easy-to-use and efficient functions for Data Mining; however, the user may perform the analysis, exploiting other packages for data mining, such as the *caret* package.

First, we extracted data from the `MatrixFeatures(res_GA)` object where the expression values for the selected features are stored. Then, we transformed this matrix and the `classes_GARS` in a `data.frame` object that we used as input for the `DaMiR.EnsembleLearning` function. We set `iter = 5` for practical reasons.

```
# expression data of selected features
data_reduced_GARS <- MatrixFeatures(res_GA)

# Classification
data_reduced_DaMiR <- as.data.frame(data_reduced_GARS)
classes_DaMiR <- as.data.frame(colData(datanorm))
colnames(classes_DaMiR) <- "class"
rownames(classes_DaMiR) <- rownames(data_reduced_DaMiR)

DaMiR.MDSplot(data_reduced_DaMiR, classes_DaMiR)
DaMiR.Clustplot(data_reduced_DaMiR, classes_DaMiR)
set.seed(12345)
Classification.res <- DaMiR.EnsembleLearning(data_reduced_DaMiR,
                                             as.factor(classes_DaMiR$class),
                                             iter=5)

## You select: RF LR kNN LDA NB SVM weak classifiers for creating
##           the Ensemble meta-learner.
## Ensemble classification is running. 5 iterations were chosen:
## Accuracy [%]:
## Ensemble RF SVM NB LDA LR kNN
## Mean: 0.91 0.93 0.9 0.9 0.89 0.92 0.94
## St.Dev. 0.05 0.05 0.07 0.07 0.07 0.03 0.04
## MCC score:
## Ensemble RF SVM NB LDA LR kNN
## Mean: 0.83 0.88 0.81 0.8 0.79 0.85 0.89
## St.Dev. 0.1 0.08 0.14 0.14 0.13 0.05 0.08
## Specificity:
## Ensemble RF SVM NB LDA LR kNN
## Mean: 0.92 0.96 0.9 0.89 0.91 0.94 0.94
## St.Dev. 0.05 0.05 0.08 0.02 0.05 0.05 0.06
## Sensitivity:
## Ensemble RF SVM NB LDA LR kNN
## Mean: 0.92 0.93 0.92 0.92 0.89 0.93 0.96
## St.Dev. 0.1 0.1 0.11 0.12 0.12 0.1 0.06
## PPV:
## Ensemble RF SVM NB LDA LR kNN
## Mean: 0.91 0.96 0.89 0.89 0.91 0.93 0.93
## St.Dev. 0.05 0.06 0.08 0 0.05 0.06 0.06
## NPV:
```

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

```
## Ensemble RF SVM NB LDA LR kNN  
## Mean: 0.91 0.91 0.91 0.91 0.87 0.91 0.96  
## St.Dev. 0.12 0.12 0.12 0.14 0.14 0.12 0.06
```

The features selected by *GARS* allowed us to clearly discriminate the N and T classes (See Figures 3 and 4). Moreover, we obtained high classification accuracy for all the classifiers built using the features set (See Figure 5).

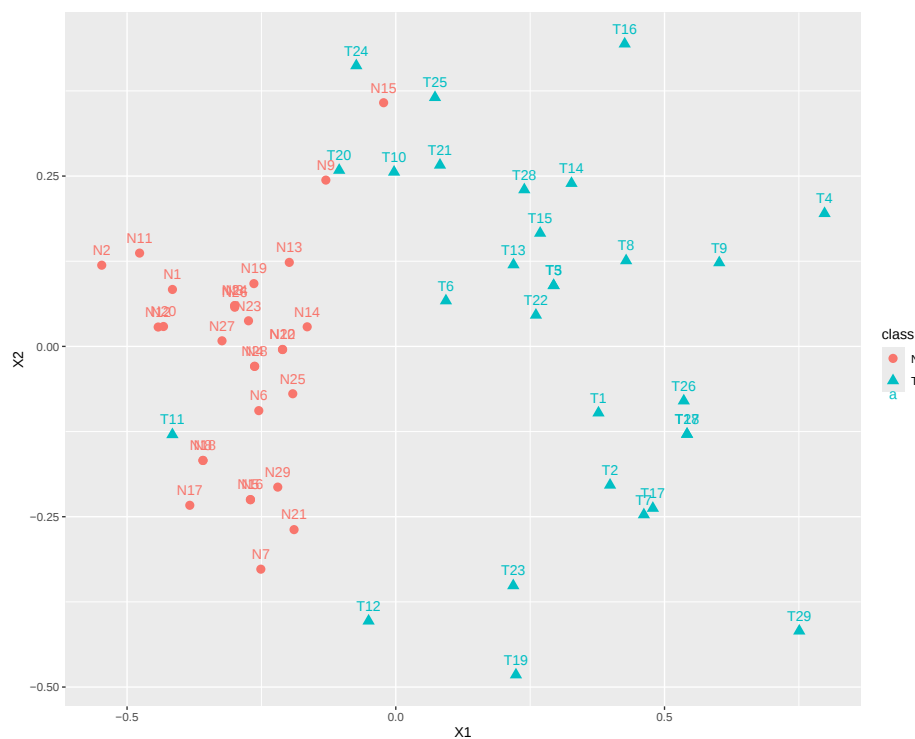


Figure 3: MultiDimensional Scaling plot. The MDS is drawn using the 8 features selected by *GARS*. The averaged Silhouette Index (i.e. the maximum fitness function of the last population) is equal to 0.55.

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

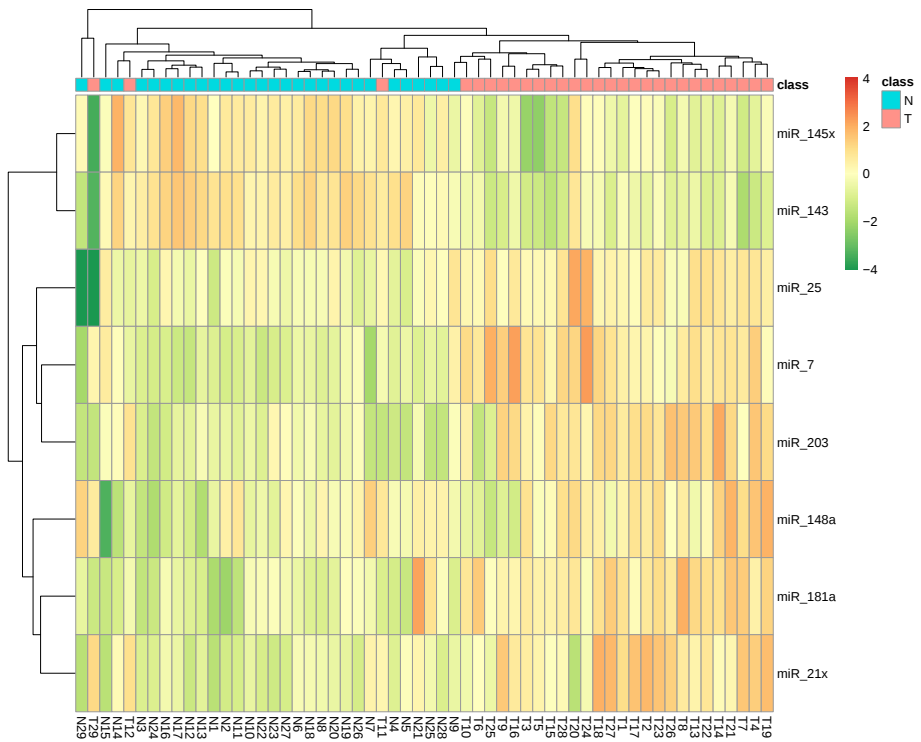


Figure 4: Clustergram. The clustergram is drawn using the 8 features selected by GARS.

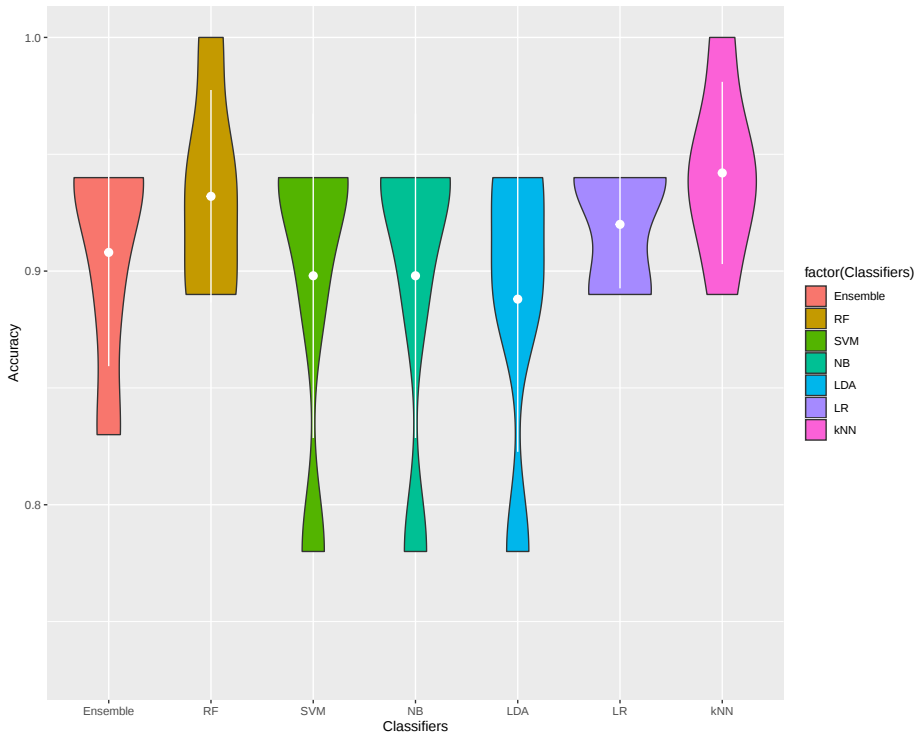


Figure 5: Violin plot. The violin plot highlights the classification accuracy of each classifier. Using the features set, selected by GARS, the averaged classification accuracy is always high, despite the small number of iterations.

2.4 Find the best features set

In the previous Section, we run *GARS* setting `chr.len = 8`. In this way, we forced the algorithm to find the best solution consisting of 8 features. However, this is probably not the best solution ever, but rather the optimal solution with 8 features. In order to find the best solution, we need to try several values of `chr.len`.

A practical solution is to insert the *GARS_GA* function inside a for loop. In the next example, we run *GARS* with `chr.len` equal to 7, 8 and 9. Finally, we search for the best solution.

```
populs <- list()
k=1
for (ik in c(7,8,9)){
  set.seed(1)
  cat(ik, "features","\n")
  populs[[k]] <- GARS_GA(data=datanorm,
                        classes = colData(datanorm),
                        chr.num = 100,
                        chr.len = ik,
                        generat = 20,
                        co.rate = 0.8,
                        mut.rate = 0.1,
                        n.elit = 10,
                        type.sel = "RW",
                        type.co = "one.p",
                        type.one.p.co = "II.quart",
                        n.gen.conv = 150,
                        plots = "no",
                        verbose="no")

  k <- k +1
}

## 7 features
## Reached 10 iterations...
## Reached 20 iterations...
## GARS found a solution after 20 iterations.
## 8 features
## Reached 10 iterations...
## Reached 20 iterations...
## GARS found a solution after 20 iterations.
## 9 features
## Reached 10 iterations...
## Reached 20 iterations...
## GARS found a solution after 20 iterations.

# find the maximum fitness for each case
max_fit <- 0

for (i in seq_len(length(populs))){
  max_fit[i] <- max(FitScore(populs[[i]]))
}
max_fit
```

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

```
## [1] 0.5471264 0.5596298 0.4976552  
best_popul <- populs[[which(max_fit == max(max_fit))]]  
  
# number of features (best solution)  
dim(MatrixFeatures(best_popul))[2]  
## [1] 8
```

Now, we can compare the results obtained from several `chr.len` values and select the best solution, looking at the maximum fitness scores and, eventually, applying the “law of parsimony” principle (*Occam’s razor*).

3 Build your custom GA

As mentioned in Section 2, the best way to use *GARS* is to run the `GARS_GA` function. However, the *GARS* package allows the user to build a custom GA (e.g. avoiding the Crossover step), joining the functions embedded in `GARS_GA`:

- `GARS_create_rnd_population()` - allows creating a random chromosome population;
- `GARS_FitFun()` - allows computing the fitness function, given a chromosome population;
- `GARS_Elitism()` - allows splitting a chromosome population, ordered by fitness scores;
- `GARS_Selection()` - allows selecting the best chromosomes, given a chromosome population;
- `GARS_Crossover()` - allows performing the Crossover step;
- `GARS_Mutation()` - allows performing the Mutation step;
- `GARS_PlotFitnessEvolution()` - allows plotting the evolution of the maximum fitness over the generations;
- `GARS_PlotFeaturesUsage()` - allows plotting how many times a feature is present over the generations;

4 Session Info

- R version 4.5.0 RC (2025-04-04 r88126), 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=en_US.UTF-8, LC_ADDRESS=en_US.UTF-8, LC_TELEPHONE=en_US.UTF-8, LC_MEASUREMENT=en_US.UTF-8, LC_IDENTIFICATION=en_US.UTF-8
- Time zone: America/New_York
- TZcode source: system (glibc)
- Running under: Ubuntu 24.04.2 LTS
- Matrix products: default
- BLAS: /home/biocbuild/bbs-3.21-bioc/R/lib/libRblas.so
- LAPACK: /usr/lib/x86_64-linux-gnu/lapack/liblapack.so.3.12.0
- Base packages: base, datasets, grDevices, graphics, methods, stats, stats4, utils
- Other packages: Biobase 2.68.0, BiocGenerics 0.54.0, DaMiRseq 2.20.0, GARS 1.28.0, GenomeInfoDb 1.44.0, GenomicRanges 1.60.0, IRanges 2.42.0, MLSeq 2.26.0, MatrixGenerics 1.20.0, S4Vectors 0.46.0, SummarizedExperiment 1.38.0, caret 7.0-1, cluster 2.1.8.1, generics 0.1.3, ggplot2 3.5.2, knitr 1.50, lattice 0.22-7, matrixStats 1.5.0
- Loaded via a namespace (and not attached): AnnotationDbi 1.70.0, BiocFileCache 2.16.0, BiocIO 1.18.0, BiocManager 1.30.25, BiocParallel 1.42.0, BiocStyle 2.36.0, Biostrings 2.76.0, DBI 1.2.3, DESeq2 1.48.0, DT 0.33, DelayedArray 0.34.0, EDASeq 2.42.0, FSelector 0.34, FactoMineR 2.11, Formula 1.2-5, GenomeInfoDbData 1.2.14, GenomicAlignments 1.44.0,

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

GenomicFeatures 1.60.0, Hmisc 5.2-3, KEGGREST 1.48.0, MASS 7.3-65, Matrix 1.7-3, ModelMetrics 1.2.2.2, R.methodsS3 1.8.2, R.oo 1.27.0, R.utils 2.13.0, R6 2.6.1, RColorBrewer 1.1-3, RCurl 1.98-1.17, RSNNS 0.4-17, RSQLite 2.3.9, RWeka 0.4-46, RWekajars 3.9.3-2, Rcpp 1.0.14, Rdpack 2.6.4, Rsamtools 2.24.0, S4Arrays 1.8.0, ShortRead 1.66.0, SparseArray 1.8.0, TH.data 1.1-3, UCSC.utils 1.4.0, XML 3.99-0.18, XVector 0.48.0, abind 1.4-8, annotate 1.86.0, arm 1.14-4, aroma.light 3.38.0, backports 1.5.0, base64enc 0.1-3, bdsmatrix 1.3-7, biomaRt 2.64.0, bit 4.6.0, bit64 4.6.0-1, bitops 1.0-9, blob 1.2.4, boot 1.3-31, cachem 1.1.0, checkmate 2.3.2, class 7.3-23, cli 3.6.4, coda 0.19-4.1, codetools 0.2-20, colorspace 2.1-1, compiler 4.5.0, corrplot 0.95, crayon 1.5.3, curl 6.2.2, data.table 1.17.0, dbplyr 2.5.0, deldir 2.0-4, digest 0.6.37, dplyr 1.1.4, e1071 1.7-16, edgeR 4.6.0, emmeans 1.11.0, entropy 1.3.2, estimability 1.5.1, evaluate 1.0.3, farver 2.1.2, fastmap 1.2.0, filelock 1.0.3, flashClust 1.01-2, foreach 1.5.2, foreign 0.8-90, future 1.40.0, future.apply 1.11.3, genalg 0.2.1, genefilter 1.90.0, ggrepel 0.9.6, globals 0.16.3, glue 1.8.0, gower 1.0.2, grid 4.5.0, gridExtra 2.3, gtable 0.3.6, hardhat 1.4.1, highr 0.11, hms 1.1.3, htmlTable 2.4.3, htmltools 0.5.8.1, htmlwidgets 1.6.4, httr 1.4.7, httr2 1.1.2, hwriter 1.3.2.1, igraph 2.1.4, ineq 0.2-13, interp 1.1-6, ipred 0.9-15, iterators 1.0.14, jpeg 0.1-11, jsonlite 2.0.0, kkn 1.3.1, labeling 0.4.3, latticeExtra 0.6-30, lava 1.8.1, leaps 3.2, lifecycle 1.0.4, limma 3.64.0, listenv 0.9.1, lme4 1.1-37, locfit 1.5-9.12, lubridate 1.9.4, magrittr 2.0.3, memoise 2.0.1, mgcv 1.9-3, minqa 1.2.8, multcomp 1.4-28, multcompView 0.1-10, munsell 0.5.1, mvtnorm 1.3-3, nlme 3.1-168, nloptr 2.2.1, nnet 7.3-20, pROC 1.18.5, parallel 4.5.0, parallelly 1.43.0, pheatmap 1.0.12, pillar 1.10.2, pkgconfig 2.0.3, pls 2.8-5, plsVarSel 0.9.13, plyr 1.8.9, png 0.1-8, praznik 11.0.0, prettyunits 1.2.0, prodlim 2024.06.25, progress 1.2.3, proxy 0.4-27, purrr 1.0.4, pwalgn 1.4.0, rJava 1.0-11, randomForest 4.7-1.2, rappdirs 0.3.3, rbitutils 2.3, recipes 1.2.1, reformulas 0.4.0, reshape2 1.4.4, restfulr 0.0.15, rjson 0.2.23, rlang 1.1.6, rmarkdown 2.29, rpart 4.1.24, rstudioapi 0.17.1, rtracklayer 1.68.0, sSeq 1.46.0, sandwich 3.1-1, scales 1.3.0, scatterplot3d 0.3-44, splines 4.5.0, statmod 1.5.0, stringi 1.8.7, stringr 1.5.1, survival 3.8-3, sva 3.56.0, tibble 3.2.1, tidyselect 1.2.1, timeDate 4041.110, timechange 0.3.0, tinytex 0.57, tools 4.5.0, vctrs 0.6.5, withr 3.0.2, xfun 0.52, xml2 1.3.8, xtable 1.8-4, yaml 2.3.10, zoo 1.8-14

References

- [1] Yvan Saeys, Inaki Inza, and Pedro Larranaga. A review of feature selection techniques in bioinformatics. *bioinformatics*, 23(19):2507–2517, 2007.
- [2] Zena M Hira and Duncan F Gillies. A review of feature selection and feature extraction methods applied on microarray data. *Advances in bioinformatics*, 2015, 2015.
- [3] Mohd Saberi Mohamad, Safaai Deris, and Rosli Md Illias. A hybrid of genetic algorithm and support vector machine for features selection and classification of gene expression microarray. *International Journal of Computational Intelligence and Applications*, 5(01):91–107, 2005.
- [4] Leping Li, Clarice R Weinberg, Thomas A Darden, and Lee G Pedersen. Gene selection for sample classification based on gene expression data: study of sensitivity to choice of parameters of the ga/knn method. *Bioinformatics*, 17(12):1131–1142, 2001.
- [5] Max Kuhn et al. The caret package. *Journal of Statistical Software*, 28(5):1–26, 2008.

GARS: a Genetic Algorithm for the identification of Robust Subsets of variables in high-dimensional and challenging datasets

- [6] CH Ooi and Patrick Tan. Genetic algorithms applied to multi-class prediction for the analysis of gene expression data. *Bioinformatics*, 19(1):37–44, 2003.
- [7] Peter J Rousseeuw. Silhouettes: a graphical aid to the interpretation and validation of cluster analysis. *Journal of computational and applied mathematics*, 20:53–65, 1987.
- [8] Mattia Chiesa, Giada Maioli, Gualtiero I Colombo, and Luca Piacentini. Gars: Genetic algorithm for the identification of a robust subset of features in high-dimensional datasets. *BMC bioinformatics*, 21(1):54, 2020.
- [9] Daniela Witten, Robert Tibshirani, Sam G Gu, Andrew Fire, and Weng-Onn Lui. Ultra-high throughput sequencing-based small rna discovery and discrete statistical biomarker analysis in a collection of cervical tumours and matched controls. *BMC biology*, 8(1):58, 2010.