

~\Downloads\R基础和基本统计.r

```

1  #修改C盘临时文件夹下面中文的名称问题
2  #https://zhidao.baidu.com/question/748080860947976012.html
3  #安装包 options(warning=-1)
4  Sys.setenv(LANGUAGE= "en")
5  #https://rdr.io/snippets/ 网络R
6  #初学者很好的学习材料
7  #https://tianyishi2001.github.io/r-and-tidyverse-book/what-is-R.html#youshi
8  #https://www.rdocumentation.org/
9  #搜R语言和来自CRAN, GitHub和Bioconductor上的近18000个包的所有的函数的说明和使用例。
10 install.packages("vegan")
11 #install.packages(c("vegan","hier.part"))
12 library(vegan)
13 #require(vegan);citation("vegan")
14 cca
15 ?cca
16 example(cca)
17 ?hier.part
18 #如果两个??,加函数还是搜不到,用google,bing,百度
19 #在当前R里面获取任何相关字节
20 getAnywhere(t.test)
21 #R官网的初学者材料
22 help.start()
23 help(package='vegan') #显示已经安装的包的描述和函数说明
24 RSiteSearch("pca") #在官方网站上联网搜索
25 Sys.getenv('R_HOME') #报告R主程序安装目录
26 (.packages())#查看工作空间已加载的包
27
28 #向量vector
29 a <- c(1, 2, 5, 3, 6, -2, 4)      # 数值型向量
30 b <- c("one", "two", "three")    # 字符型向量
31 c <- c(TRUE, NA, FALSE)         # 逻辑型向量
32
33 #因子factor
34 province <- c("四川", "湖南", "江苏", "四川", "四川", "四川", "湖南", "江苏", "湖南", "江苏")
35 pf <- factor(province)          # 创建 province 的因子 pf
36 pf
37 #向量选择
38 a <- c(1, 2, 5, 3, 6, -2, 4)
39 a[3]
40 a[c(1, 3, 5)]
41 a[2:6]
42 a[a>3&a<5]
43 a[-1]
44 a[-1:-3]
45 a#产生向量
46 rep(2:5, 2) # 等价于 rep(2:5, times = 2)
47 rep(2:5, each=2)
48 rep(1:3, times = 4, each = 2) #这个在自己编写差异表达分析的分组信息时可以用到
49 seq(1,10,0.1)
50 #想一下,如何保存这些系列呢,赋值!
51 #0.1-1, 间隔0.1, 重复1: 10
52
53 #矩阵matrix必须是同一属性数据,如果不是,会自动调整为同一属性, vector类同
54 matrix(1:4, nrow = 2, ncol = 2) #默认所有元素是按列向
55 (mat1<-matrix(1:12,6, byrow=T))
56 mat1

```

```

57 mat2<-matrix(1:12,3,4, byrow=F)
58 mat2
59 x <- matrix(1:6, 2, 3)
60 x[2,2]
61 x[2,]
62 x[,2]
63 x[1:2,2:3]
64 ## %% does matrix multiplication
65 m <- t(x)      # Assign m the transposed m
66 x %% t(x)     # %% does matrix multiplication
67 m * m         # * 哈达马积(hadamard product)
68 #数组
69 a <- array(data=1:18,dim=c(3,3,2)) # 3d with dimensions 3x3x2
70 a <- array(1:18,c(3,3,2))         # the same array
71 a
72 #选取数字为17怎么选
73
74 #数据框data frame
75 patientID <- c(1, 2, 3, 4)
76 age <- c(25, 34, 28, 52)
77 diabetes <- c("Type1", "Type2", "Type1", "Type1")
78 status <- c("Poor", "Improved", "Excellent", "Poor")
79 patientdata <- data.frame(patientID, age, diabetes, status)
80 patientdata
81 patientdata[1:2,] # 第 1、2 行的所有元素
82 patientdata[,1:2]
83 patientdata$age   #"$"符合用于选取一个指定的变量
84 patientdata[, "age"] #试一下去逗号
85
86 patientdata[5,]=c(4,52,"Type1","Poor")
87 unique(patientdata) #去掉重复的行,常用于数据清洗
88 #但是不要乱加
89 a=unique(patientdata)
90 a$age #现在的a$age这个向量里面都是字符,而不是数值了,因为前面patientdata[5,]这一步操作相当于加
    的是一个字符型的向量。因为向量里面的元素类型应该是一致的。
91
92 # 列表list
93 #list是储物箱, list里面可以装list, 可以无限延展。
94 g <- "My First List"
95 h <- c(25, 26, 18, 39)
96 j <- matrix(1:10, nrow = 5)
97 k <- c("one", "two", "three")
98 mylist <- list(title = g, ages = h, j, k, patientdata)
99 mylist[[3]] #提取list的第3个元素
100 mylist[3]   #也是提取的list的第3个元素, 但提取出来后, 仍是作为列表的第1个元素。少了一个套, 相当
    于没有提出真正的元素。
101 #观察一下有什么不同吗?
102 #用class来判断一个对象的类型。 class()
103
104
105 #工作目录
106 getwd()
107 # 导入 text 和csv 文件, 当前工作目录下
108 # mydata2 <- read.csv("oats.csv")
109 mydata2 <- read.csv("../learning/lai_jiangshan/data/oats.csv")
110 mydata2
111
112 # mydata3 <- read.table("oats.txt", header=TRUE)
113 mydata3 <- read.table("../learning/lai_jiangshan/data/oats.txt", header = TRUE)
114 mydata3

```

```

115
116 mydata4 <- read.csv(file.choose(), header=T)
117 mydata4
118 class(mydata3)
119
120 #mydata4 <- read.csv(file.choose(), header=T)
121 #row.names=1
122 #读excel表格，两个方式任选一种
123 #install.packages("openxlsx")
124 library(openxlsx)
125 # mydata <- read.xlsx("oats.xlsx",sheet = 1)
126 mydata <- read.xlsx("./learning/lai_jiangshan/data/oats.xlsx", sheet = 1)
127 mydata
128 class(mydata)
129 #如果是openxlsx不行，尝试一下readxl函数
130 #install.packages("readxl")
131 #library(readxl)
132 #mydata <- read_excel("oats.xlsx",sheet = 1)
133 #mydata
134 #mydata=as.data.frame(mydata)
135 #class(mydata)
136 #http://www.milanor.net/blog/read-excel-files-from-r/ 读excel数据网站
137
138 #加盘符"", 其他类型数据输入看ppt
139 #install.packages("foreign")
140 library(foreign)
141 x <-read.spss("sps.sav", to.data.frame=T) # Read file sps.sav
142 x
143 #尝试一下去掉to.data.frame=T，读入是列表
144 #https://cran.r-project.org/doc/manuals/r-release/R-data.html
145 #http://rsoftware.h.baike.com/article-1867733.html"
146 #练习读入数据你们自己的数据，发现问题
147
148 getwd()
149 #2.3 数据的存储
150 write.table(mydata2, "D:/fg.txt", row.names = F)
151 write.csv(mydata2,"D:/fg.csv", row.names = F)
152
153 library(openxlsx)#优势，可以保留列名重复
154 write.xlsx(mydata2,"D:/fg.xlsx")
155 ?write.xlsx
156 save(mydata2, file = "D:/fg.Rdata") #练习一下多个文件,这个"file ="不能省略
157 save(mydata2, x, mydata, file = "D:/fg.Rdata") #save 可以存储多个文件,非常有用,
158 #而且这种格式".Rdata"直接打开，不用read.table, read.csv之类的函数来读取；另一个好处是可以保存
  其他数据类型用于下一次直接打开，vevtor, matrix, array都可以
159 #而read.csv(), read.txt()只能读入数据框data frame。
160 #所以，如果不是为了给其他人数据，以后都用save()来存.Rdata数据，这个非常有用啊。
161 #载入.Rdata数据，用"load()"函数即可。
162
163 #fix()这个函数可以直接把某个数据调出进行修改,不过服务器上似乎用不了。
164 fix(mydata2)
165
166 ls() #显示当前工作目录下所有对象，包括：数据、函数等
167
168
169 #3 数据的排序 (order)
170 head(mydata2)
171 # plyr 包的 arrange( )函数
172 library(plyr)
173 arrange(mydata2, yield)

```

```
174 arrange(mydata2, Variety, -yield)
175 arrange(mydata2, desc(Variety), -yield)
176 #https://stackoverflow.com/questions/1296646/how-to-sort-a-dataframe-by-multiple-columns
177
178 #数据集的合并
179 #案例分析
180 # chinaIUCN<- read.csv("chinaIUCN.csv")
181 chinaIUCN <- read.csv("./learning/lai_jiangshan/data/chinaIUCN.csv")
182 head(chinaIUCN)
183
184 # focussp <- read.csv("shennongjia.csv") #如果文件中有中文，那么用记事本保存为"UTF-8"编码，再
    在read.csv()函数中加是"encoding = "UTF-8""
185 focussp <- read.csv("./learning/lai_jiangshan/data/shennongjia-UTF-8.csv", encoding = "UTF-
    8")
186 head(focussp)
187 bb <- merge(chinaIUCN, focussp, by.x = "species", by.y ="name", all.y=T)
188 head(bb)
189 class(bb)
190 #练习，结果存为CSV和R的格式
191
192 #多个数据列是否可以
193 #http://rstudio-pubs-static.s3.amazonaws.com/13602_96265a9b3bac4cb1b214340770aa18a1.html
194
195 #数据结构查看
196 data() #查看R里面的自带数据集dim(iris)      # 数据集的维度，有多少行多少列？
197 names(iris) # 数据有哪些列？
198 str(iris)   # 数据的结构如何？
199 iris[1:5,]  # 查看数据的前 5 行
200 head(iris)  # 查看数据的前 6 行
201 tail(iris)  # 查看数据的最后 6 行
202 class(iris) # 查看类别
203
204 #巧用批处理，对所有行/列作同一种运算。
205 sapply(iris[, -5], mean, na.rm=T)
206 apply(iris[, -5], 2, mean, na.rm=T) #margin=2对列
207 apply(iris[, -5], 1, mean, na.rm=T) #margin=1对行
208
209 #tapply第一个参数为统计的对象，第二个参数为分组的因子，第三参数为统计函数
210 tapply(iris$Sepal.Length, iris$Species, var) #针对某一列的作一 种统计分析，可以按一个或多个因素
    分组，分组用list(),这一个或多个因素都会强行分类，
211 #如果某个因素组合下没有统计结果，就用“NA”表示。最终生成一个长行表。
212 aggregate(iris[, -5], list(iris$Species), sd) #可以对多列进行一种统计分析,分组用list(), 只保留有
    统计结果的组合，但最终data.frame没有列名，最终生成长列表
213
214 #不同的列做不同的运算，即多列数据进行多种统计运算，这个最好，包罗前面的，同时生成的结果
    data.frame也有列名，代码量也小一些。
215 library(plyr)
216 ddply(iris, .
    (Species), summarise, sum.a=sum(Sepal.Length), mean.c=mean(Sepal.Width), pl.sd=sd(Petal.Length))
217
218 #双分组因子
219 library(vegan)
220 data(dune.env)
221 head(dune.env)
222 ddply(dune.env, .(Management, Use), summarise, sum.a1=sum(A1))
223
224 #处理多变量的分组
225 #我是xxx,最近在整理大样地的数据。刚好读了您的博文《将样方多度数据转为“样方-物种“矩阵》，有个问
    题想请教一下老师。
226 #我现在的数据有两列，一列是样方号（site），一列是物种名(species)， 怎么转换成三列数据样方号
    site, 物种名species, 物种多度abund。
227 #我的部分数据见附件。
```

```

228 #dfx <- read.csv("样地原始数据.csv")
229 dfx <- read.csv("../learning/lai_jiangshan/data/样地原始数据.csv")
230 head(dfx)
231 #方法1, tapply(), 单列单运算法, 可以分组
232 zqdata <- tapply(dfx$abund, list(dfx$site, dfx$species), sum)
233 head(zqdata)
234 #把NA变为0
235 zqdata[is.na(zqdata)] = 0
236 #ifelse()和“vegan”包中的decostand()两个函数可以转换原始数据为0、1数据。decostand()主要是用于数
    据转换的函数。
237 #其他可能用到的转换数据: ifelse(zqdata > 0, 1, zqdata)
238
239 #方法2, aggregate(), 多列单运算法, 可以分组
240 zqdata1 <- aggregate(dfx[, 3], list(dfx$site, dfx$species), sum)
241 head(zqdata1)
242
243 #方法3, ddply(), 多列多运算法, 可以分组
244 zqdata2 <- ddply(dfx, .(site, species), summarise, sum.abund = sum(abund))
245 head(zqdata2)
246
247 #思考一下如何把0的地方变为NA?
248 zqdata[zqdata == 0] <- NA
249
250 #####
251 install.packages(c("adephylo", "adespatial", "FD", "ggpubr"))
252 #####
253
254 save(zqdata, file = "D:/ourresult.Rdata")
255 #练习存储你自己的数据和计算平均值
256 #options(warn=-1)
257 #描述统计
258 #https://www.statmethods.net/stats/descriptives.html
259 #install.packages("pastecs")
260 library(pastecs)
261 #Sepal.Width花萼宽度
262 x <- iris$Sepal.Width
263 stat.desc(x)
264 #x可以是矩阵
265 #标准误等于标准差/样本量的平方根/置信区间值为qt(0.975, 150-1)*SE
266 0.43586628/sqrt(150)
267 qt(0.975, 150-1)*0.03558833 #计算置信区间qt(0.975, n-1)*SE
268 #多个物种的情况
269 tapply(iris[, 2], iris[, 5], stat.desc)
270 #练习, 样本量为100, 平均值为5, 标准差为0.6, 求平均值95%的置信区间
271 #尝试一下, 是否可以给矩阵matrix或数据框data.frame
272 #Hit <Return> to see next plot: par(ask=F) #
273 #histogram
274 x1 <- iris$Sepal.Width
275 h<-hist(x1, breaks=10, col=2, xlab="", main="Histogram with Normal Curve")
276 hist(x1, main="Histogram of observed data", prob=T) ##probability densities:
    h$counts/150/diff(h$breaks)
277 #Kernel Density Estimation
278 lines(density(x1), col=2, lwd=3, lty=2)
279
280 #par(ask=F)
281 plot(density(x1), col=1, lty=1, lwd=2)
282 #dorm后面这个x是个表达式
283 curve(dnorm(x, mean=3.057333, sd=0.4358663), add=T, col=4, lwd=3, lty=2)
284 #for normality test
285 qqnorm(x1)

```

```

286 qqline(x1)
287
288 #线上发散，线下收敛
289 #W = 是观测与实际重叠的面积
290 #http://www.real-statistics.com/tests-normality-and-symmetry/statistical-tests-normality-symmetry/shapiro-wilk-test/
291 #样本量必须在3-5000之间
292 ?shapiro.test
293 shapiro.test(x1)
294 shapiro.test(log(x1))
295 #17 popular distribution
296 ?distribution
297 dnorm(0,mean=0,sd=1)
298 pnorm(1.96)
299 qnorm(0.975)
300 rnorm(10)
301 #dname( ) density or probability function
302 #pname( ) cumulative density function
303 #qname( ) quantile function
304 #rname( ) random deviates
305 #Probability Plots
306 #https://www.statmethods.net/advgraphs/probability.html
307 #分布拟合
308 #https://stats.stackexchange.com/questions/132652/how-to-determine-which-distribution-fits-my-data-best
309 #options(warn=-1)
310
311 x=iris$Petal.Width
312 #看一个定量向量是否属于某一个分布，先拟合求参数，再把参数代到ks.test里面
313 library(MASS)
314 #把x作为gamma分布进行拟合，得到参数
315 fitdistr(x,'gamma')
316 #用上一步得到拟合参数，作为真正的gamma分布，与实际的x向量值进行比较，确定x是否是gamma分布
317 ks.test(x,'pgamma',1.5578249,1.2989101)
318 fitdistr(x,'poisson')
319 ks.test(x,'ppois',1.19933333)
320
321 #练习，检验你们的数据是否属于正态分布
322
323 #####单样本T检验#####
324 #杨树的树高理论值为8m
325 height<-c(8.0, 7.9,7.9,8.1, 8.2, 8.3, 8.3, 8.5, 8.6, 8.8, 8.3, 8.5,8.7,8.7,
326           8.7, 8.8,7.2, 8.0, 7.5, 7.6, 9.0, 8.7, 8.5, 8.8, 8.7, 8.6, 8.6,
327           8.5, 8.2, 8.2, 8.0, 8.3)
328 shapiro.test(height)
329 #t=(x平均值-mu)/se
330 t.test(height, mu = 8, alternative = "two.sided")
331 #结果为: 2*(1-pt(4.5186, 31), 2*(1-pt(mu, df))
332 #t.test(height, mu = 8, alternative = "greater")
333 #t.test(height, mu = 8, alternative = "less")
334 sink("test.txt")
335 t.test(height, mu = 8, alternative = "two.sided")
336 sink()
337 #sink()函数一定是要闭合，跟括号一样，并且不会在屏幕上显示
338 #练习读入你们自己的数据，然后进行单样本T检验
339 #change rank to test威尔科克森秩检验（非参数）
340 # c("two.sided", "less", "greater")
341 wilcox.test(height, mu = 8)
342
343 ##### chisq test #####

```

```
344 freq = c(22,21,22,27,22,36)
345 prob = rep(1,6)/6
346 chisq.test(freq,p=prob)
347 #卡方检验都检验一切计算实测与预测值之间是否显著差异检验
348 #sum((freq-期望值)^2/期望值)
349 #1-pchisq(6.72,5)
350
351 #set union, intersection, (asymmetric!) difference, equality and membership on two vectors.
352 x <- 1:10
353 y <- 5:14
354 #并集
355 union(x, y)
356 #交集
357 intersect(x, y)
358 setdiff(x, y)# which in x, not in y
359 setdiff(y, x)
360 #ks.test()
361 ##### 两个样本平均值T检验#####
362 #whether is mean of A equal to mean of B
363 A <- c(79.98, 80.04, 80.02, 80.04, 80.03, 80.03, 80.04, 79.97,
364 80.05, 80.03, 80.02, 80.00, 80.02)
365
366 B <-c(80.02,79.94, 79.98, 79.97, 79.97, 80.03, 79.95, 79.97)
367 #ks.test()
368 #####
369 #t.test的正态分布+方差齐性检验
370 shapiro.test(A)
371 shapiro.test(B)
372 var.test(A,B)
373 dev.off()#删除当前的做图窗口
374 boxplot(A,B)
375 #零假设为差值=0, 等价于95%置信区间是否互相包括对方的均值
376 t.test(A,B,var.equal=T)
377 #零假设设定为A-B=0.02
378 t.test(A,B,mu=0.02,var.equal=T)
379 #var.equal =T, 加这个参数, 统计结果更容易显著
380
381 #if no norm 威尔科克森秩检验
382 wilcox.test(A,B)
383
384 #如果不正态, 用置换检验的比较即可
385 library(RVAideMemoire)
386 perm.t.test(A,B,nperm=999)
387 #成对T检验
388 #Testdata=read.csv("ravens.csv", header=T)
389 Testdata=read.csv("../learning/lai_jiangshan/data/ravens.csv", header=T)
390 Testdata
391 t.test(Testdata$after, Testdata$before, paired=T)
392 t.test(Testdata$after, Testdata$before)
393
394 #查看变量之间的相关性
395 #Practice Correlations
396 ?cor
397 cor(iris$Sepal.Length, iris$Sepal.Width)
398 cor.test(iris$Sepal.Length, iris$Sepal.Width)
399
400 plot(iris$Sepal.Length, iris$Sepal.Width)
401 cor.test(iris$Sepal.Length, iris$Sepal.Width,method= "spearman")
402 cor(iris[,1:4]) #default pearson
403
```

```

404 # 用相关系数算显著性
405 #对数据不要求正态分布和方差齐性
406 r=-0.1175698;n=150
407 #p值的算法
408 (1-pt(abs(r)*sqrt((n-2)/(1-r^2)),n-2))*2
409
410 #多变量的相关
411 library(psych)
412 corr.test(iris[,1:4])
413 cord <- corr.test(iris[,1:4],adjust="none")
414 cord$r #提取相关系数
415 cord$p #提取p值
416 (domG <- read.csv("./learning/lai_jiangshan/data/domG.csv",header=T))
417 (enzydomG <- read.csv("./learning/lai_jiangshan/data/enzydomG.csv"))
418 #计算酶与微生物的相关
419 library(psych)
420 pvalue <- corr.test(domG,enzydomG,adjust="none")
421 pvalue
422 round(pvalue$p,4)
423 library(corrplot)
424 dev.new()
425 corrplot(pvalue$r, method = "circle")
426 #?ccf #研究时间系列相关cross-correlation (互相关)
427 # pairs() is a function to plot a matrix of bivariate scatter plots.
428 # panelutils.R is a set of functions that add useful features to pairs():
429 # upper.panel=panel.cor: to print correlation coefficients in the upper panel,
430 # with significance levels;
431 # diag.panel=panel.hist: to plot histograms of the variables in the diagonal.
432 # Specify method for the choice of the correlation coefficient:
433 # by default, method = "pearson", other choices are "spearman" and "kendall".
434 # To get a plot in grey tones instead of colors, use no.col=TRUE and
435 # lower.panel=panel.smoothb.
436
437 source("./learning/lai_jiangshan/data/panelutils.R")
438 #去掉趋势线, 请去掉 lower.panel参数
439 pairs(iris[,1:4], lower.panel=panel.smooth, upper.panel=panel.cor,
440       diag.panel=panel.hist, main="Pearson Correlation Matrix",method="pearson")
441 #https://cran.r-project.org/web/packages/corrplot/vignettes/corrplot-intro.html
442 #http://uc-r.github.io/correlations
443 #install.packages("GGally")
444 library(GGally)
445 ggpairs(iris[,1:4])
446
447
448 #Partial Correlation 偏相关
449 # data
450 library(ppcor)
451 data=read.csv("./learning/lai_jiangshan/data/pcordata.csv")
452 data
453 # partial correlation between "x" and "y" given "v1"
454 #做剔除v1,v2影响后的x和y的相关
455 pcor.test(data$x,data$y,data$v1)
456
457 #attach()
458 a=resid(lm(x~v1,data))#回归获得残差
459 b=resid(lm(y~v1,data))#回归获得残差
460 cor.test(a,b)
461 cor.test(data$x,data$y)
462
463

```

```
464 #
465 ###图形输出
466 #ggplot2包下面的函数，只能保存ggplot2刚画出来的图
467 ggsave("./project/Temp/lai.pdf")
468
469 #三步：创建画布，画图，关闭画布
470 pdf(file="output.pdf") #输出文件名为output 的pdf文件
471 par(mfrow=c(2,2))
472 plot(1:10)
473 plot(1:10)
474 plot(1:10)
475 plot(1:10)
476 dev.off() #关闭pdf 绘图设备 输出的图形请在工作路径下查找
477 #拼接pdf的网站
478 #https://smallpdf.com/cn/merge-pdf
479 #https://www.ilovepdf.com/zh-cn/merge_pdf
480
481 #三步：创建画布，画图，关闭画布
482 jpeg(filename="output.jpg") #输出文件名为output 的jpg文件
483 plot(1:10) #绘图
484 dev.off() #关闭jpeg绘图设备 输出的图形请在工作路径下查找
485 dev.new() #如果Rstudio左边框里看不到图，就输这个命令
486 #输出高分辨率的图，用于发表
487 tiff(filename = "test.tif", width = 1200, height = 1200,units = "px",
488 pointsize = 3,bg = "white", res = 600)
489 ?tiff
```