

平均, 標準偏差, 最大値, 中央値, 最小値, 分位数

```
summary(データフレーム名)      または      apply(データフレーム名, 2, summary)
colMeans(データフレーム名)     または      apply(データフレーム名, 2, mean)
sapply(データフレーム名, sd)   または      apply(データフレーム名, 2, sd)
```

summary : 最小値, 四分位数, 最大値, 平均値
colMeans : 各列の平均値
mean : 平均値
sd : 標準偏差
max : 最大値
median : 中央値
min : 最小値
quantile(値) : 値で指定された累積比率のところのデータ
na.rm=TRUE : mean や sd などにおいて, 欠測値があるときはそれを除外する指定

```
> setwd("i:¥¥Rdocuments¥¥scripts¥¥")
> d1 <- read.table("統計図表データ.csv", header=T, sep=",")
> head(d1)
  番号 入学年度  学科 性別 モラトリアム 自己効力感 学習意欲 進路
1    1      20Y1 看護学  女         高          49      23 就職
2    2      20Y2 心理学  男         低          57      29 就職
3    3      20Y1  医学  女         高          42      23 進学
4    4      20Y1 看護学  女         高          41      23 就職
5    5      20Y2  医学  男         低          41      22 就職
6    6      20Y1 心理学  女         低          47      24 就職
```

```
> #分析に必要な変数だけを取り出す
> dtmp <- d1[,c("自己効力感", "学習意欲")]
>
```

```
> # 標本サイズ
> (ntmp <- nrow(dtmp))
[1] 270
>
```

```
> # 各変数の平均
> (mtmp <- colMeans(dtmp))
自己効力感   学習意欲
  49.60000    25.20741
>
```

```
> # 各変数の標準偏差
> (sd.d2 <- apply(dtmp, 2, sd))
自己効力感   学習意欲
  9.492866    5.099160
>
```

```
> # 各変数の最小値, 四分位数, 最大値
> (summarytmp <- summary(dtmp))
  自己効力感   学習意欲
Min.   :26.0   Min.    : 8.00
1st Qu.:42.0   1st Qu.:22.00
Median :49.0   Median :25.00
Mean   :49.6   Mean    :25.21
3rd Qu.:57.0   3rd Qu.:29.00
Max.   :77.0   Max.    :39.00
>
```

	A	B	C	D	E	F	G	H
1	番号	入学年度	学科	性別	モラトリアム	自己効力感	学習意欲	進路
2	1	20Y1	看護学	女	高	49	23	就職
3	2	20Y2	心理学	男	低	57	29	就職
4	3	20Y1	医学	女	高	42	23	進学
5	4	20Y1	看護学	女	高	41	23	就職
6	5	20Y2	医学	男	低	41	22	就職
7	6	20Y1	心理学	女	低	47	24	就職
8	7	20Y3	看護学	女	高	41	26	就職
9	8	20Y2	医学	男	低	43	29	就職
10	9	20Y2	看護学	女	低	54	21	就職
11	10	20Y3	医学	女	高	55	30	進学
12	11	20Y1	心理学	男	低	66	25	就職
13	12	20Y1	心理学	女	低	38	30	就職
14	13	20Y1	心理学	女	低	45	29	就職
15	14	20Y1	医学	男	高	53	23	就職
16	15	20Y1	心理学	男	高	47	25	就職

複数群あるときの要約統計量

集計したい変数が1つだけの場合

`tapply`(集計したい変数名, 群分け変数名, 関数, 関数のオプション群)

集計したい変数は1つしか指定できない。

集計したい変数が1つ以上の場合

byを使う方法

`by`(データフレーム名, 群分け変数名, 関数, 関数のオプション群)

aggregateを使う方法

`aggregate`(データフレーム名, `list`(群分け変数名), 関数, 関数のオプション群)

sapplyとtapplyを組み合わせて使う方法

`sapply`(データフレーム名, `tapply`, 群分け変数名, 関数, 関数のオプション群)

データフレームは, 集計したい変数だけを入れたものにするか, 集計したい変数を指定する。

```
> setwd("i:\\Rdocuments\\scripts\\")
> d1 <- read.table("統計図表データ.csv", header=T, sep=",")
> head(d1)
  番号 入学年度 学科 性別 モラトリアム 自己効力感 学習意欲 進路
1    1      20Y1 看護学 女      高          49        23 就職
2    2      20Y2 心理学 男      低          57        29 就職
3    3      20Y1 医学 女      高          42        23 進学
4    4      20Y1 看護学 女      高          41        23 就職
5    5      20Y2 医学 男      低          41        22 就職
6    6      20Y1 心理学 女      低          47        24 就職
>
>
```

	A	B	C	D	E	F	G	H
1	番号	入学年度	学科	性別	モラトリアム	自己効力感	学習意欲	進路
2		1 20Y1	看護学	女	高	49	23	就職
3		2 20Y2	心理学	男	低	57	29	就職
4		3 20Y1	医学	女	高	42	23	進学
5		4 20Y1	看護学	女	高	41	23	就職
6		5 20Y2	医学	男	低	41	22	就職
7		6 20Y1	心理学	女	低	47	24	就職
8		7 20Y3	看護学	女	高	41	26	就職
9		8 20Y2	医学	男	低	43	29	就職
10		9 20Y2	看護学	女	低	54	21	就職
11		10 20Y3	医学	女	高	55	30	進学
12		11 20Y1	心理学	男	低	66	25	就職
13		12 20Y1	心理学	女	低	38	30	就職
14		13 20Y1	心理学	女	低	45	29	就職
15		14 20Y1	医学	男	高	53	23	就職
16		15 20Y1	心理学	男	高	47	25	就職

> # 各群の標本サイズ

```
> table(d1$入学年度)
```

```
20Y1 20Y2 20Y3
  87   98   85
```

> # tapplyを使う方法

```
> tapply(d1$学習意欲, d1$入学年度, summary)
```

```
$`20Y1`
  Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
  8.00  21.50   23.00   24.08  28.50   37.00
```

```
$`20Y2`
  Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
 16.00  22.00   25.50   25.48  29.00   39.00
```

```
$`20Y3`
  Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
 15.00  23.00   26.00   26.05  29.00   38.00
```

```
> tapply(d1$学習意欲, d1$入学年度, sd)
```

```
 20Y1    20Y2    20Y3
5.528401 5.105545 4.445184
```

```

>
> # byを使う方法
> by(d1[,c("学習意欲")], d1$入学年度, summary)
d1$入学年度: 20Y1
  Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
   8.00   21.50   23.00   24.08   28.50   37.00
-----
d1$入学年度: 20Y2
  Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
  16.00   22.00   25.50   25.48   29.00   39.00
-----
d1$入学年度: 20Y3
  Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
  15.00   23.00   26.00   26.05   29.00   38.00

> by(d1[,c("学習意欲")], d1$入学年度, sd)
d1$入学年度: 20Y1
[1] 5.528401
-----
d1$入学年度: 20Y2
[1] 5.105545
-----
d1$入学年度: 20Y3
[1] 4.445184
>

> # aggregateを使う方法
> aggregate(d1[,c("学習意欲")], list(d1$入学年度), summary)
  Group.1 x. Min.  x. 1st Qu.  x. Median x. Mean x. 3rd Qu.  x. Max.
1    20Y1   8.00    21.50    23.00  24.08    28.50  37.00
2    20Y2  16.00    22.00    25.50  25.48    29.00  39.00
3    20Y3  15.00    23.00    26.00  26.05    29.00  38.00

> aggregate(d1[,c("学習意欲")], list(d1$入学年度), sd)
  Group.1 x
1    20Y1 5.528401
2    20Y2 5.105545
3    20Y3 4.445184
>

> # sapply と tapply を組み合わせて使う方法
> sapply(d1[,c("学習意欲", "自己効力感")], tapply, d1$入学年度, mean)
      学習意欲 自己効力感
20Y1 24.08046   47.40230
20Y2 25.47959   50.14286
20Y3 26.04706   51.22353

> sapply(d1[,c("学習意欲", "自己効力感")], tapply, d1$入学年度, sd)
      学習意欲 自己効力感
20Y1 5.528401   9.908253
20Y2 5.105545   9.649529
20Y3 4.445184   8.516617
>
>
>
> # 群別にデータフレームを分割する
> d1Y1 <- d1[d1$入学年度=="20Y1", ]
> d1Y2 <- d1[d1$入学年度=="20Y2", ]
> d1Y3 <- d1[d1$入学年度=="20Y3", ]
>
> # 各群における記述統計量
> nrow(d1Y1)
[1] 87

> summary(d1Y1[,c("学習意欲")])
  Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
   8.00   21.50   23.00   24.08   28.50   37.00

> sd(d1Y1[,c("学習意欲")])
[1] 5.528401

```

歪度・尖度

```
library(psych)
describe(データフレーム名[, c("変数名1", "変数名2", ...)])
```

群ごとの表示

byを使う方法

```
library(psych)
by(データフレーム名, 群分け変数名, describe)
```

describeByを使う方法

```
library(psych)
describeBy(データフレーム名, 群分け変数名)
```

あらかじめpsychパッケージをインストールしておく必要がある。

出力内容

item name 変数名
 item number 標本の大きさ
 number of valid cases 有効数
 mean 算術平均
 standard deviation 標準偏差 (分母=n-1)
 trimmed mean (with trim defaulting to .1) 調整平均(指定した割合の両端データを削除したときの平均)
 median (standard or interpolated) 中央値
 mad: median absolute deviation (from the median) 1.4826 × [中央値からの絶対偏差]の中央値
 minimum 最小値
 maximum 最大値
 skew 歪度
 kurtosis 尖度
 standard error 標準誤差

	A	B	C	D	E	F	G	H
1	番号	入学年度	学科	性別	モラトリアム	自己効力感	学習意欲	進路
2	1	20Y1	看護学	女	高	49	23	就職
3	2	20Y2	心理学	男	低	57	29	就職
4	3	20Y1	医学	女	高	42	23	進学
5	4	20Y1	看護学	女	高	41	23	就職
6	5	20Y2	医学	男	低	41	22	就職
7	6	20Y1	心理学	女	低	47	24	就職
8	7	20Y3	看護学	女	高	41	26	就職
9	8	20Y2	医学	男	低	43	29	就職
10	9	20Y2	看護学	女	低	54	21	就職
11	10	20Y3	医学	女	高	55	30	進学
12	11	20Y1	心理学	男	低	66	25	就職
13	12	20Y1	心理学	女	低	38	30	就職
14	13	20Y1	心理学	女	低	45	29	就職
15	14	20Y1	医学	男	高	53	23	就職
16	15	20Y1	心理学	男	高	47	25	就職

```
> setwd("i:¥¥¥documents¥¥¥scripts¥¥¥")
> d1 <- read.table("統計図表データ.csv", header=T, sep=",")
> head(d1)
  番号 入学年度 学科 性別 モラトリアム 自己効力感 学習意欲 進路
1    1      20Y1 看護学 女          高          49        23 就職
2    2      20Y2 心理学 男          低          57        29 就職
3    3      20Y1 医学 女          高          42        23 進学
4    4      20Y1 看護学 女          高          41        23 就職
5    5      20Y2 医学 男          低          41        22 就職
6    6      20Y1 心理学 女          低          47        24 就職
>
> #psych パッケージの読み込み
> library(psych)
>
```

> # 記述統計量の一括表示

```
> describe(d1[, c("学習意欲", "自己効力感")])
      vars   n  mean   sd median trimmed   mad min max range  skew
学習意欲    1 270 25.21 5.10    25   25.22  4.45   8  39    31 -0.01
自己効力感    2 270 49.60 9.49    49   49.51 10.38  26  77    51  0.10
      kurtosis   se
学習意欲    0.08 0.31
自己効力感   -0.31 0.58
```

```

>
> # 群別の記述統計量の一括表示
> # byを使う方法
> by(d1[, c("学習意欲", "自己効力感")], d1$入学年度, describe)
d1$入学年度: 20Y1
      vars  n mean   sd median trimmed   mad min max range skew
学習意欲    1 87 24.08 5.53    23   24.15  4.45   8 37   29 -0.12
自己効力感    2 87 47.40 9.91    46   47.04 10.38  26 73   47  0.32
      kurtosis   se
学習意欲    0.16 0.59
自己効力感   -0.58 1.06
-----
d1$入学年度: 20Y2
      vars  n mean   sd median trimmed   mad min max range skew
学習意欲    1 98 25.48 5.11    25.5   25.36  5.19  16 39   23 0.26
自己効力感    2 98 50.14 9.65    49.5   49.98 11.12  26 77   51 0.16
      kurtosis   se
学習意欲   -0.46 0.52
自己効力感    0.04 0.97
-----
d1$入学年度: 20Y3
      vars  n mean   sd median trimmed   mad min max range skew
学習意欲    1 85 26.05 4.45    26   26.04  4.45  15 38   23 0.10
自己効力感    2 85 51.22 8.52    52   51.35  8.90  29 72   43 -0.14
      kurtosis   se
学習意欲   -0.17 0.48
自己効力感   -0.34 0.92
>

> # describeByを使う方法
> describeBy(d1[, c("学習意欲", "自己効力感")], d1$入学年度)
group: 20Y1
      vars  n mean   sd median trimmed   mad min max range skew
学習意欲    1 87 24.08 5.53    23   24.15  4.45   8 37   29 -0.12
自己効力感    2 87 47.40 9.91    46   47.04 10.38  26 73   47  0.32
      kurtosis   se
学習意欲    0.16 0.59
自己効力感   -0.58 1.06
-----
group: 20Y2
      vars  n mean   sd median trimmed   mad min max range skew
学習意欲    1 98 25.48 5.11    25.5   25.36  5.19  16 39   23 0.26
自己効力感    2 98 50.14 9.65    49.5   49.98 11.12  26 77   51 0.16
      kurtosis   se
学習意欲   -0.46 0.52
自己効力感    0.04 0.97
-----
group: 20Y3
      vars  n mean   sd median trimmed   mad min max range skew
学習意欲    1 85 26.05 4.45    26   26.04  4.45  15 38   23 0.10
自己効力感    2 85 51.22 8.52    52   51.35  8.90  29 72   43 -0.14
      kurtosis   se
学習意欲   -0.17 0.48
自己効力感   -0.34 0.92
>

```

2 要因の人数, 平均, 標準偏差

自作関数を使う

desc2(data=データ名, x1="要因 1 変数名", x2="要因 2 変数名", y="従属変数名")

```
#-----
desc2 <- function(data, x1, x2, y) {
  dtmp <- data
  nc <- as.matrix(table(dtmp[, x1], dtmp[, x2]))
  nm1 <- as.matrix(margin.table(nc, 1))
  nm2 <- as.matrix(margin.table(nc, 2))
  nm <- as.matrix(margin.table(nc))
  ndtmp <- rbind(cbind(nc, nm1), c(nm2, nm))
  rownames(ndtmp) <- c(rownames(nc), "Total")
  colnames(ndtmp) <- c(colnames(nc), "Total")

  mc <- tapply(dtmp[, y], list(dtmp[, x1], dtmp[, x2]), mean)
  mm1 <- tapply(dtmp[, y], list(dtmp[, x1]), mean)
  mm2 <- tapply(dtmp[, y], list(dtmp[, x2]), mean)
  mm <- mean(dtmp[, y])
  mdtmp <- rbind(cbind(mc, mm1), c(mm2, mm))
  mdtmp <- round(mdtmp, 2)
  mcbar <- mc

  sc <- tapply(dtmp[, y], list(dtmp[, x1], dtmp[, x2]), sd)
  sm1 <- tapply(dtmp[, y], list(dtmp[, x1]), sd)
  sm2 <- tapply(dtmp[, y], list(dtmp[, x2]), sd)
  sm <- sd(dtmp[, y])
  sdtmp <- rbind(cbind(sc, sm1), c(sm2, sm))
  sdtmp <- round(sdtmp, 2)

  vr <- NULL
  ddesc2 <- matrix(c(""), (nrow(ndtmp)*4), ncol(ndtmp), byrow=T)
  for(i in 1:nrow(ndtmp)) {
    ic <- 4*(i-1)+1
    iN <- ic+1
    im <- iN+1
    is <- im+1
    ddesc2[iN,] <- ndtmp[i,]
    ddesc2[im,] <- mdtmp[i,]
    ddesc2[is,] <- sdtmp[i,]
    vr <- c(vr, c(rownames(ndtmp)[i], "N", "Mean", "SD"))
  }
  ddesc2 <- data.frame(vr, ddesc2)
  colnames(ddesc2) <- c("", colnames(ndtmp))
  print(ddesc2)
  ddesc2 <- ddesc2
}
#-----
```

対応のある要因がある場合は, stackデータを作成してから, この関数を用いる。
対応のある要因についても, 周辺度数は「人数×水準数」となることに注意。

結果の外部出力

write.table(ddesc2, "ファイル名.csv", row.names=F, sep=",")

```
> setwd("i:¥¥Rdocuments¥¥scripts")
> d1 <- read.table("2B平均値データ.csv", header=TRUE, sep=",")
> head(d1)
  id work  furture tekiou
1  1  few shushoku    10
2  2 many shingaku    12
3  3 many shushoku    12
4  4  few shushoku    12
```

```

5 5 few mitei 13
6 6 many shushoku 9
>

```

```

> # 2 要因をクロスさせたときの人数, 平均, 標準偏差
> desc2(data=d1, x1="work", x2="furture", y="tekiou")

```

```

      mitei shingaku shushoku Total
1  few
2  N      17       16       81    114
3  Mean 12.06    14.19    12.31 12.54
4  SD   1.82     2.48     2.06  2.18
5  many
6  N      25       24       63    112
7  Mean 10.4     14.38    12.83 12.62
8  SD   1.96     1.61     1.77  2.22
9  Total
10 N      42       40      144    226
11 Mean 11.07    14.3     12.53 12.58
12 SD   2.05     1.98     1.95  2.19
>
>

```

id	work	furture	tekiou
1	few	shushoku	10
2	many	shingaku	12
3	many	shushoku	12
4	few	shushoku	12
5	few	mitei	13
6	many	shushoku	9
7	few	shushoku	12
8	few	shushoku	11
9	few	shushoku	14
10	many	shushoku	12
11	many	shingaku	15
12	few	shushoku	14
13	many	shushoku	12
14	few	shushoku	14
15	many	mitei	11
16	few	shushoku	13
17	few	shingaku	17
18	many	mitei	10
19	few	shushoku	9
20	few	shushoku	19

```

> # 結果の外部ファイルへの出力
> write.table(ddesc2, "2BNMSD.csv", row.names=F, sep=",")
>
>

```

	mitei	shingaku	shushoku	Total
few				
N	17	16	81	114
Mean	12.06	14.19	12.31	12.54
SD	1.82	2.48	2.06	2.18
many				
N	25	24	63	112
Mean	10.4	14.38	12.83	12.62
SD	1.96	1.61	1.77	2.22
Total				
N	42	40	144	226
Mean	11.07	14.3	12.53	12.58
SD	2.05	1.98	1.95	2.19

共分散・相関係数

不偏共分散行列 (n-1で割る)

cov(データフレーム名)

相関係数行列

cor(データフレーム名, method="算出方法", use="欠測値の扱い方法")

method: "pearson" (default), "spearman", or "kendall"
 use: "everything" 当該2変数に欠測値がある場合, その箇所の値だけNAとなる.
 "complete.obs" 1つでも欠測値のある行を除外してすべての値を計算.
 "pairwise.complete.obs" 当該2変数に欠測値がある場合, その箇所だけ欠測値を除外して計算.

```
> setwd("i:¥¥Rdocuments¥¥scripts¥¥")
> d1 <- read.table("偏相関データ.csv", header=TRUE, sep=",")
> head(d1)
```

```
  grade kutsu math
1      3    21  49
2      2    20  36
3      3    20  47
4      3    20  54
5      6    26  73
6      1    17   2
```

```
>
> # 共分散
```

```
> cov(d1)
      grade      kutsu      math
grade 2.926421 4.573579 30.80100
kutsu 4.573579 9.073567 47.37457
math 30.801003 47.374571 400.18394
>
```

```
> # ピアソンの積率相関係数
```

```
> cor(d1)
      grade      kutsu      math
grade 1.0000000 0.887562 0.9000499
kutsu 0.8875620 1.000000 0.7861880
math 0.9000499 0.786188 1.0000000
>
```

```
> # スピアマンの順位相関係数
```

```
> cor(d1, method="spearman")
      grade      kutsu      math
grade 1.0000000 0.9004303 0.9094415
kutsu 0.9004303 1.0000000 0.8158021
math 0.9094415 0.8158021 1.0000000
>
```

```
> # ケンドールの順位相関係数
```

```
> cor(d1, method="kendall")
      grade      kutsu      math
grade 1.0000000 0.7881925 0.7766475
kutsu 0.7881925 1.0000000 0.6314266
math 0.7766475 0.6314266 1.0000000
>
>
```

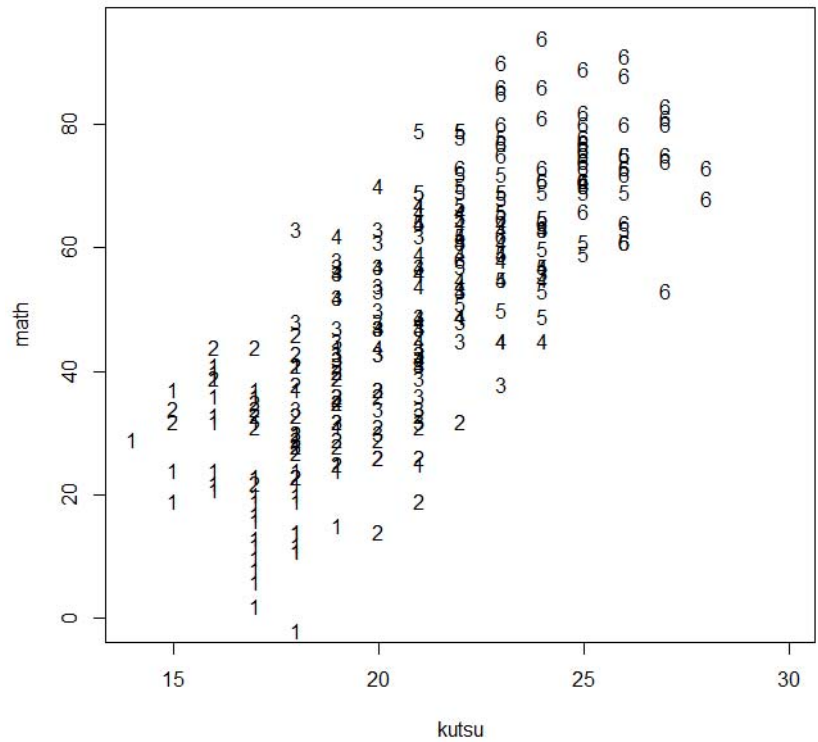
```
> 群別データ
```

```
> d11 <- d1[d1$grade==1,]
> d12 <- d1[d1$grade==2,]
> d13 <- d1[d1$grade==3,]
> d14 <- d1[d1$grade==4,]
> d15 <- d1[d1$grade==5,]
> d16 <- d1[d1$grade==6,]
>
```

	A	B	C
1	grade	kutsu	math
2	3	21	49
3	2	20	36
4	3	20	47
5	3	20	54
6	6	26	73
7	1	17	2
8	1	21	25
9	3	19	57
10	6	24	73
11	4	23	58
12	2	18	27
13	2	16	39
14	4	21	66
15	3	21	44
16	5	24	71
17	4	21	59
18	5	24	69
19	1	16	21
20	1	19	31
21	6	24	81

> # 群別散布図の重ね合わせ

```
> plot(d11$kutsu, d11$math, pch="1", axes=T, xlim=c(14, 30), ylim=c(0, 95),
+      xlab="kutsu", ylab="math")
> par(new=T)
> plot(d12$kutsu, d12$math, pch="2", axes=F, xlim=c(14, 30), ylim=c(0, 95), xlab="", ylab="")
> par(new=T)
> plot(d13$kutsu, d13$math, pch="3", axes=F, xlim=c(14, 30), ylim=c(0, 95), xlab="", ylab="")
> par(new=T)
> plot(d14$kutsu, d14$math, pch="4", axes=F, xlim=c(14, 30), ylim=c(0, 95), xlab="", ylab="")
> par(new=T)
> plot(d15$kutsu, d15$math, pch="5", axes=F, xlim=c(14, 30), ylim=c(0, 95), xlab="", ylab="")
> par(new=T)
> plot(d16$kutsu, d16$math, pch="6", axes=F, xlim=c(14, 30), ylim=c(0, 95), xlab="", ylab="")
>
>
>
>
>
```



複数群あるときの共分散・相関係数

不偏共分散行列 (n-1で割る)
by(データフレーム名, 群別変数名, cov)

相関係数行列
by(データフレーム名, 群別変数名, cor, use="欠測値の扱い方法")

```
method: "pearson" (default), "spearman", or "kendall"
use: "everything"      当該2変数に欠測値がある場合, その箇所の値だけNAとなる.
    "complete.obs"     1つでも欠測値のある行を除外してすべての値を計算.
    "pairwise.complete.obs" 当該2変数に欠測値がある場合, その箇所だけ欠測値を除外して計算.
```

データフレームは, 群別変数も含めた必要な変数だけにしておく.

```
> setwd("i:¥¥Rdocuments¥¥scripts¥¥")
> d1 <- read.table("偏相関データ.csv", header=TRUE, sep=",")
> head(d1)
  grade kutsu math
1      3    21  49
2      2    20  36
3      3    20  47
4      3    20  54
5      6    26  73
6      1    17   2
>
>
```

```
> # 各群の標本サイズ
> table(d1$grade)
```

```
 1  2  3  4  5  6
50 50 50 50 50 50
>
```

```
> # 各群の平均
> aggregate(d1[, c("kutsu", "math")], list(d1$grade), mean)
  Group.1 kutsu math
1      1 17.22 23.64
2      2 18.54 32.58
3      3 20.26 46.54
4      4 21.64 55.50
5      5 23.38 66.72
6      6 24.98 75.04
>
```

```
> # 各群のSD
> sapply(d1[, c("kutsu", "math")], tapply, d1$grade, sd)
      kutsu      math
1 1.266362 11.035675
2 1.501156  6.931413
3 1.208980  8.447219
4 1.396205  8.576784
5 1.496799  7.946351
6 1.477588  8.787654
>
```

```
> # 各群の相関係数行列
> by(d1[, c("kutsu", "math")], d1$grade, cor)
d1$grade: 1
      kutsu      math
kutsu 1.000000 -0.154852
math -0.154852 1.000000
```

	A	B	C
1	grade	kutsu	math
2	3	21	49
3	2	20	36
4	3	20	47
5	3	20	54
6	6	26	73
7	1	17	2
8	1	21	25
9	3	19	57
10	6	24	73
11	4	23	58
12	2	18	27
13	2	16	39
14	4	21	66
15	3	21	44
16	5	24	71
17	4	21	59
18	5	24	69
19	1	16	21
20	1	19	31
21	6	24	81

```

d1$grade: 2
      kutsu      math
kutsu 1.0000000 -0.2641161
math -0.2641161 1.0000000
-----

d1$grade: 3
      kutsu      math
kutsu 1.0000000 -0.06398709
math -0.06398709 1.0000000
-----

d1$grade: 4
      kutsu      math
kutsu 1.0000000 0.2164383
math 0.2164383 1.0000000
-----

d1$grade: 5
      kutsu      math
kutsu 1.0000000 -0.07494717
math -0.07494717 1.0000000
-----

d1$grade: 6
      kutsu      math
kutsu 1.0000000 -0.0738084
math -0.0738084 1.0000000
>
>
>
>

>
> # 群別データ
> d11 <- d1[d1$grade==1,]
> d12 <- d1[d1$grade==2,]
> d13 <- d1[d1$grade==3,]
> d14 <- d1[d1$grade==4,]
> d15 <- d1[d1$grade==5,]
> d16 <- d1[d1$grade==6,]
>

> # 第1群別の記述統計量
> dtmp <- d11[,c("kutsu", "math")]
> ntmp <- nrow(dtmp)
> mtmp <- colMeans(dtmp)
> stmp <- apply(dtmp, 2, sd)
> ctmp <- cor(dtmp)
> ktmp <- round(data.frame(ntmp, mtmp, stmp, ctmp), 2)
> colnames(ktmp) <- c("N", "Mean", "SD", colnames(ctmp))
> ktmp
      N Mean   SD kutsu  math
kutsu 50 17.22  1.27  1.00 -0.15
math  50 23.64 11.04 -0.15  1.00
>

> # 第2群別の記述統計量
> dtmp <- d12[,c("kutsu", "math")]
> ntmp <- nrow(dtmp)
> mtmp <- colMeans(dtmp)
> stmp <- apply(dtmp, 2, sd)
> ctmp <- cor(dtmp)
> ktmp <- round(data.frame(ntmp, mtmp, stmp, ctmp), 2)
> colnames(ktmp) <- c("N", "Mean", "SD", colnames(ctmp))
> ktmp
      N Mean   SD kutsu  math
kutsu 50 18.54  1.50  1.00 -0.26
math  50 32.58  6.93 -0.26  1.00
>

```

偏相関係数

影響を除く変数を指定した偏相関係数

```
library(psych)
partial.r(data, c(偏相関係数を求めたい変数群の列番号), c(影響を除きたい変数群の列番号))
```

あらかじめpsychパッケージをインストールしておく必要がある。
 data はデータフレームでもよいし、相関係数行列でもよい。
 偏相関係数を求めたい変数が2つ以上ある場合は、列番号をカンマで区切る。
 影響を除きたい変数が2つ以上ある場合は、列番号をカンマで区切る。

当該の2変数以外の全ての変数の影響を除いた偏相関係数行列

自作関数を使う方法

hensoukan(相関係数行列)

```
#-----
hensoukan <- function(x) {
  if (det(x) != 0) {
    inv.r <- solve(x)
    p <- nrow(x)
    par.r <- inv.r
    d.r <- diag(1/sqrt(diag(inv.r)), p, p)
    par.r <- d.r %*% -inv.r %*% d.r
    diag(par.r) <- rep(1, p)
    rownames(par.r) <- rownames(x)
    colnames(par.r) <- colnames(x)
  }
  return(par.r)
}
#-----
```

関数部分は、R起動中に1回行えば良い。

```
> setwd("i:¥¥Rdocuments¥¥scripts¥¥")
> d1 <- read.table("偏相関データ.csv", header=TRUE, sep=",")
> head(d1)
  grade kutsu math
1      3    21  49
2      2    20  36
3      3    20  47
4      3    20  54
5      6    26  73
6      1    17   2
>
>
> # 記述統計量
> dtmp <- d1
> ntmp <- nrow(dtmp)
> mtmp <- colMeans(dtmp)
> stmp <- apply(dtmp, 2, sd)
> ctmp <- cor(dtmp)
> ktmp <- round(data.frame(ntmp, mtmp, stmp, ctmp), 2)
> colnames(ktmp) <- c("N", "Mean", "SD", colnames(ctmp))
> ktmp
      N Mean    SD grade kutsu math
grade 300  3.5  1.71  1.00  0.89  0.90
kutsu 300 21.0  3.01  0.89  1.00  0.79
math  300 50.0 20.00  0.90  0.79  1.00>
```

	A	B	C
1	grade	kutsu	math
2	3	21	49
3	2	20	36
4	3	20	47
5	3	20	54
6	6	26	73
7	1	17	2
8	1	21	25
9	3	19	57
10	6	24	73
11	4	23	58
12	2	18	27
13	2	16	39
14	4	21	66
15	3	21	44
16	5	24	71
17	4	21	59
18	5	24	69
19	1	16	21
20	1	19	31
21	6	24	81

```

> # 偏相関係数
> library(psych)
> partial.r(d1, c(2,3), c(1))
partial correlations
      kutsu  math
kutsu  1.00 -0.06
math   -0.06  1.00

> partial.r(d1, c(1,3), c(2))
partial correlations
      grade math
grade  1.00  0.71
math   0.71  1.00

> partial.r(d1, c(1,2), c(3))
partial correlations
      grade kutsu
grade  1.00  0.67
kutsu  0.67  1.00
>
>
>

> # 当該の2変数以外の全ての変数の影響を除いた偏相関係数
>
> # 相関係数の逆行列を用いて算出する関数
> #-----
> #partial correlation matrix
> #
> hensoukan <- function(x, vn=c(1:ncol(x))) {
+   x <- x[vn,vn]
+   if (det(x) != 0) {
+     inv.r <- solve(x)
+     p <- nrow(x)
+     par.r <- inv.r
+     d.r <- diag(1/sqrt(diag(inv.r)), p, p)
+     par.r <- d.r %*% -inv.r %*% d.r
+     diag(par.r) <- rep(1, p)
+     rownames(par.r) <- rownames(x)
+     colnames(par.r) <- colnames(x)
+   }
+   return(par.r)
+ }
> #-----
>
>
> hensoukan(ctmp)
      grade      kutsu      math
grade 1.0000000 0.6682002 0.7104308
kutsu 0.6682002 1.0000000 -0.0630700
math  0.7104308 -0.0630700 1.0000000
>
>

```

四分相関係数・多分相関係数

```
library(psych)
library(polycor)
オブジェクト名1 <- polychoric(データフレーム名)
```

相関係数行列だけの取り出し
オブジェクト名1\$rho

あらかじめpsychとpolycorパッケージをインストールしておく必要がある。
データフレーム以外に、テーブルを代入したりすることもできる。

```
> setwd("i:¥¥Rdocuments¥¥scripts¥¥")
> d1 <- read.table("2値_データ.csv", header=TRUE, sep=",")
> head(d1)
  x1 x2 x3 x4 x5 x6 x7 x8
1  0  0  1  1  1  0  0  1
2  0  0  1  1  0  0  1  0
3  0  0  0  0  0  0  0  0
4  0  1  0  0  1  0  1  1
5  0  1  1  1  1  1  0  1
6  0  0  0  1  0  0  1  0
>
```

```
> # 標本サイズ, 平均, 標準偏差
> dtmp <- d1
> ntmp <- nrow(d1)
> mtmp <- apply(d1, 2, mean)
> stmp <- apply(d1, 2, sd)
> ktmp <- round(data.frame(ntmp, mtmp, stmp), 2)
> colnames(ktmp) <- c("N", "Mean", "SD")
> ktmp
      N Mean  SD
x1 346 0.48 0.5
x2 346 0.54 0.5
x3 346 0.51 0.5
x4 346 0.53 0.5
x5 346 0.49 0.5
x6 346 0.53 0.5
x7 346 0.53 0.5
x8 346 0.51 0.5>
```

```
> #積率相関係数
> round(cor(dtmp), 2)
      x1  x2  x3  x4  x5  x6  x7  x8
x1 1.00 0.18 0.25 0.16 0.03 0.05 0.05 0.01
x2 0.18 1.00 0.19 0.05 0.07 0.07 0.01 0.10
x3 0.25 0.19 1.00 0.14 0.09 0.04 0.03 0.07
x4 0.16 0.05 0.14 1.00 0.07 0.01 -0.02 0.00
x5 0.03 0.07 0.09 0.07 1.00 0.13 0.17 0.24
x6 0.05 0.07 0.04 0.01 0.13 1.00 0.10 0.03
x7 0.05 0.01 0.03 -0.02 0.17 0.10 1.00 0.12
x8 0.01 0.10 0.07 0.00 0.24 0.03 0.12 1.00
>
>
```

```
> #四分相関係数・多分相関係数
> library(psych)
> library(polycor)
```

	A	B	C	D	E	F	G	H
1	x1	x2	x3	x4	x5	x6	x7	x8
2	0	0	1	1	1	0	0	1
3	0	0	1	1	0	0	1	0
4	0	0	0	0	0	0	0	0
5	0	1	0	0	1	0	1	1
6	0	1	1	1	1	1	0	1
7	0	0	0	1	0	0	1	0
8	0	1	0	0	0	1	0	0
9	1	1	1	1	0	1	0	0
10	1	0	1	1	1	1	0	0
11	1	1	1	0	0	0	1	1
12	1	1	0	0	0	0	1	1
13	1	1	0	1	0	1	0	1
14	0	1	0	0	1	0	0	0
15	0	1	0	1	1	1	1	1
16	0	0	1	0	0	1	1	1
17	1	1	0	1	1	1	0	0
18	1	1	1	1	1	0	1	1
19	1	1	1	0	1	1	0	1
20	1	1	1	0	0	0	0	1
21	1	1	0	0	1	1	1	1

```

> (ptcor.d1 <- polychoric(d1))
Call: polychoric(x = d1)
Polychoric correlations
  x1    x2    x3    x4    x5    x6    x7    x8
x1  1.00
x2  0.28  1.00
x3  0.38  0.29  1.00
x4  0.25  0.07  0.22  1.00
x5  0.05  0.11  0.15  0.11  1.00
x6  0.08  0.10  0.07  0.01  0.21  1.00
x7  0.09  0.02  0.04 -0.03  0.27  0.16  1.00
x8  0.02  0.16  0.11  0.00  0.36  0.04  0.18  1.00

  with tau of
    0
x1  0.051
x2 -0.094
x3 -0.036
x4 -0.080
x5  0.014
x6 -0.073
x7 -0.065
x8 -0.014
>

```

```

> # 相関係数行列だけの取り出し
> pcor.d1 <- ptcor.d1$rho
> round(pcor.d1, 2)
  x1    x2    x3    x4    x5    x6    x7    x8
x1  1.00  0.28  0.38  0.25  0.05  0.08  0.09  0.02
x2  0.28  1.00  0.29  0.07  0.11  0.10  0.02  0.16
x3  0.38  0.29  1.00  0.22  0.15  0.07  0.04  0.11
x4  0.25  0.07  0.22  1.00  0.11  0.01 -0.03  0.00
x5  0.05  0.11  0.15  0.11  1.00  0.21  0.27  0.36
x6  0.08  0.10  0.07  0.01  0.21  1.00  0.16  0.04
x7  0.09  0.02  0.04 -0.03  0.27  0.16  1.00  0.18
x8  0.02  0.16  0.11  0.00  0.36  0.04  0.18  1.00
>
>

```

アルファ係数

```
変数リスト名 <- c("変数名1", "変数名2", ..., "変数名p")
library(psych)
alpha(データフレーム名[, 変数リスト名], check.keys=FALSE)
```

あらかじめpsychパッケージをインストールしておく必要がある。

check.keys=FALSE : 合計得点と負の相関を持つ項目の得点を逆転しないというオプション。デフォルトではTRUEとなっており、逆転項目処理を行っていないくても、逆転処理を行った場合のアルファ係数を算出してしまう。FALSEとしておいて、自分で逆転処理をするようにしておいたほうが、合計点の算出で間違いがない。

出力内容

```
total      a list containing
raw_alpha  alpha based upon the covariances
std.alpha  The standardized alpha based upon the correlations
G6(smc)    Guttman's Lambda 6 reliability
average_r  The average interitem correlation
mean       For data matrices, the mean of the scale formed by summing the items
sd         For data matrices, the standard deviation of the total score
alpha.drop A data frame with all of the above for the case of each item being removed one by one
item.stats A data frame including
r          The correlation of each item with the total score (not corrected for item overlap)
r.cor      Item whole correlation corrected for item overlap and scale reliability
r.drop     Item whole correlation for this item against the scale without this item
mean       For data matrices, the mean of each item
sd         For data matrices, the standard deviation of each item
response.freq For data matrices, the frequency of each item response (if less than 20)
```

coefficientパッケージの中にも α 係数を求めるalphaという関数があるが、項目分析をやってくれない。

```
> setwd("i:¥¥Rdocuments¥¥scripts¥¥")
> d1 <- read.table("統計分析力尺度データ.csv", header=TRUE, sep=",")
> head(d1)

  番号 x1 x2 x3 x4 x5 x6 x7 x8 x9 x10 y1 y2 y3 y4 y5 y6 y7 y8 y9 y10 xt yt
1    1  3  2  1  2  2  4  4  3  4  4  3  2  1  2  1  4  3  3  3  4 29 26
2    2  3  3  4  3  2  2  2  2  4  1  2  2  4  2  2  1  1  1  3  1 26 19
3    3  4  3  3  4  1  3  4  2  5  1  4  3  4  4  1  3  4  2  5  1 30 31
4    4  5  5  5  3  3  4  4  2  4  4  5  4  5  2  2  3  4  2  3  3 39 33
5    5  3  3  4  2  2  3  3  3  4  1  3  4  4  3  2  4  3  4  5  1 28 33
6    6  2  1  4  1  1  3  3  2  5  1  3  2  4  2  2  5  4  3  5  2 23 32
統計  数学  批判的思考力  国語  自己効力感
1   51   48             28   72             61
2   74   53             26   66             53
3   48   60             35   71             48
4   67   68             27   67             48
5   55   49             30   66             49
6   74   63             36   83             37
>

> #変数リスト名
> inames <- c("x1", "x2", "x3", "x4", "x5",
+            "x6", "x7", "x8", "x9", "x10")
>
```



```
> #α係数の推定
> library(psych)
> alpha(d1[, inames], check.keys=FALSE)
```

Reliability analysis

Call: alpha(x = d1[, inames], check.keys = FALSE)

```
raw_alpha std.alpha G6(smc) average_r S/N ase mean sd
0.84      0.84      0.84      0.35 5.3 0.02 3.1 0.61
```

```
lower alpha upper      95% confidence boundaries
0.8 0.84 0.88
```

Reliability if an item is dropped:

```
raw_alpha std.alpha G6(smc) average_r S/N alpha se
x1      0.82      0.82      0.82      0.34 4.7 0.022
x2      0.82      0.82      0.82      0.34 4.6 0.022
x3      0.83      0.83      0.83      0.34 4.7 0.022
x4      0.82      0.82      0.82      0.34 4.6 0.022
x5      0.82      0.82      0.82      0.34 4.6 0.022
x6      0.82      0.82      0.82      0.34 4.7 0.022
x7      0.83      0.83      0.83      0.35 4.8 0.022
x8      0.83      0.83      0.83      0.35 4.8 0.022
x9      0.83      0.83      0.83      0.35 4.9 0.022
x10     0.83      0.83      0.83      0.36 5.0 0.022
```

Item statistics

```
      n      r r.cor r.drop mean sd
x1 365 0.67 0.62 0.57 4.0 0.92
x2 365 0.67 0.63 0.57 3.1 0.96
x3 365 0.65 0.60 0.54 4.1 0.85
x4 365 0.68 0.64 0.57 3.0 1.07
x5 365 0.67 0.63 0.57 2.2 0.93
x6 365 0.65 0.61 0.54 3.0 1.01
x7 365 0.62 0.56 0.51 3.1 0.99
x8 365 0.63 0.58 0.53 2.2 0.89
x9 365 0.59 0.52 0.48 3.9 0.90
x10 365 0.57 0.51 0.46 2.1 0.95
```

Non missing response frequency for each item

```
      1      2      3      4      5 miss
x1 0.01 0.05 0.23 0.36 0.35 0
x2 0.04 0.22 0.42 0.25 0.07 0
x3 0.01 0.03 0.22 0.40 0.35 0
x4 0.09 0.23 0.36 0.24 0.08 0
x5 0.25 0.41 0.25 0.09 0.01 0
x6 0.06 0.24 0.38 0.25 0.07 0
x7 0.05 0.22 0.36 0.30 0.07 0
x8 0.25 0.42 0.25 0.07 0.00 0
x9 0.01 0.06 0.24 0.40 0.29 0
x10 0.30 0.37 0.25 0.08 0.01 0
>
>
```

	A	B	C	D	E	F	G	H	I	J	K
1	番号	x1	x2	x3	x4	x5	x6	x7	x8	x9	x10
2	1	3	2	1	2	2	4	4	3	4	4
3	2	3	3	4	3	2	2	2	2	4	1
4	3	4	3	3	4	1	3	4	2	5	1
5	4	5	5	5	3	3	4	4	2	4	4
6	5	3	3	4	2	2	3	3	3	4	1
7	6	2	1	4	1	1	3	3	2	5	1
8	7	5	3	5	5	4	5	5	3	5	3
9	8	4	3	4	3	2	3	2	2	4	1
10	9	4	2	4	2	2	2	2	2	4	2
11	10	4	5	5	4	4	3	3	3	4	3
12	11	4	2	4	4	2	4	4	2	4	3
13	12	5	3	4	5	2	3	3	2	3	2
14	13	3	3	5	3	3	2	3	2	3	2
15	14	4	4	5	3	3	3	4	2	3	1
16	15	4	2	4	3	1	4	2	1	3	3
17	16	4	3	4	1	1	2	1	1	3	1
18	17	3	2	4	3	2	2	2	2	3	2
19	18	3	3	5	2	2	3	2	1	4	1
20	19	3	3	5	2	2	4	4	2	4	1
21	20	3	2	5	1	2	3	3	1	4	2

オメガ係数

変数リスト名 <- c("変数名1", "変数名2", ..., "変数名p")

library(coefficientsalpha)

オブジェクト名 <- omega(データフレーム名[, 変数リスト名], se=TRUE)

summary.omega(オブジェクト名)

あらかじめcoefficientsalphaパッケージをインストールしておく必要がある。
se=TRUEをつけておかないと、summary.omegaで信頼区間を計算してくれない。
逆転項目は、あらかじめスコアを正しい方向に修正しておく必要がある。

psychパッケージにもomegaという関数があるが、因子分析を使うものである。

```
> setwd("i:\\Rdocuments\\scripts\\")
> d1 <- read.table("統計分析力尺度データ.csv", header=TRUE, sep=",")
```

```
> head(d1)
```

```
  番号 x1 x2 x3 x4 x5 x6 x7 x8 x9 x10 y1 y2 y3 y4 y5 y6 y7 y8 y9 y10 xt yt
1    1  3  2  1  2  2  4  4  3  4   4  3  2  1  2  1  4  3  3  3   4 29 26
2    2  3  3  4  3  2  2  2  2  4   1  2  2  4  2  2  1  1  1  3   1 26 19
3    3  4  3  3  4  1  3  4  2  5   1  4  3  4  4  1  3  4  2  5   1 30 31
4    4  5  5  5  3  3  4  4  2  4   4  5  4  5  2  2  3  4  2  3   3 39 33
5    5  3  3  4  2  2  3  3  3  4   1  3  4  4  3  2  4  3  4  5   1 28 33
6    6  2  1  4  1  1  3  3  2  5   1  3  2  4  2  2  5  4  3  5   2 23 32
```

```
統計 数学 批判的思考力 国語 自己効力感
1    51   48           28   72           61
2    74   53           26   66           53
3    48   60           35   71           48
4    67   68           27   67           48
5    55   49           30   66           49
6    74   63           36   83           37
```

```
> #変数リスト名
```

```
> inames <- c("x1", "x2", "x3", "x4", "x5",
+            "x6", "x7", "x8", "x9", "x10")
```

```
> # omega 係数の推定
```

```
> library(coefficientsalpha)
```

```
> o1 <- omega(d1[, inames], se=T)
```

```
Test of homogeneity
```

```
The robust F statistic is  3.288
```

```
with a p-value  0
```

```
**The F test rejected homogeneity**
```

```
The omega is 0.8409224 with the standard error 0.01255086.
```

```
About 6.85% of cases were downweighted.
```

```
> # 信頼区間の計算
```

```
> summary.omega(o1)
```

```
The estimated omega is
```

```
omega          0.841
se             0.013
p-value (omega>0) 0.000
Confidence interval 0.816      0.866
```

```
Test of homogeneity
```

```
The robust F statistic is  3.288
```

```
with a p-value  0
```

```
Note. The robust test rejected the assumption of homogeneity.
```

	A	B	C	D	E	F	G	H	I	J	K
1	番号	x1	x2	x3	x4	x5	x6	x7	x8	x9	x10
2	1	3	2	1	2	2	4	4	3	4	4
3	2	3	3	4	3	2	2	2	2	2	4
4	3	4	3	3	4	1	3	4	2	5	1
5	4	5	5	5	3	3	4	4	2	4	4
6	5	3	3	4	2	2	3	3	3	4	1
7	6	2	1	4	1	1	3	3	2	5	1
8	7	5	3	5	5	4	5	5	3	5	3
9	8	4	3	4	3	2	3	2	2	4	1
10	9	4	2	4	2	2	2	2	2	4	2
11	10	4	5	5	4	4	3	3	3	4	3
12	11	4	2	4	4	2	4	4	2	4	3
13	12	5	3	4	5	2	3	3	2	3	2
14	13	3	3	5	3	3	2	3	2	3	2
15	14	4	4	5	3	3	3	4	2	3	1
16	15	4	2	4	3	1	4	2	1	3	3
17	16	4	3	4	1	1	2	1	1	3	1
18	17	3	2	4	3	2	2	2	2	3	2
19	18	3	3	5	2	2	3	2	1	4	1
20	19	3	3	5	2	2	4	4	2	4	1
21	20	3	2	5	1	2	3	3	1	4	2

級内相関係数

```
library(psych)
ICC(データフレーム名)
```

あらかじめpsychパッケージをインストールしておく必要がある。

ICC1: 被験者内要因 (k回評価) の効果を誤差に含めるモデルにおける, 評価の一致度
 ICC2: 被験者内要因 (k回評価) の効果を変量効果とするモデルにおける, 評価の一致度
 ICC3: 被験者内要因 (k回評価) の効果を固定効果とするモデルにおける, 評価の一致度
 ICC1k: 被験者内要因 (k回評価) の効果を誤差に含めるモデルにおける, k回評価の平均値の信頼性
 ICC2k: 被験者内要因 (k回評価) の効果を変量効果とするモデルにおける, k回評価の平均値の信頼性
 ICC3k: 被験者内要因 (k回評価) の効果を固定効果とするモデルにおける, k回評価の平均値の信頼性

```
> setwd("i:\\Rdocuments\\scripts\\")
> d1 <- read.table("統計分析力尺度データ.csv", header=TRUE, sep=",")
> head(d1)
  番号 x1 x2 x3 x4 x5 x6 x7 x8 x9 x10 y1 y2 y3 y4 y5 y6 y7 y8 y9 y10 xt yt
1    1  1  3  2  1  2  2  4  4  3  4  4  3  2  1  2  1  4  3  3  3  4 29 26
2    2  2  3  3  4  3  2  2  2  2  4  1  2  2  4  2  2  1  1  1  3  1 26 19
3    3  3  4  3  3  4  1  3  4  2  5  1  4  3  4  4  1  3  4  2  5  1 30 31
4    4  4  5  5  5  3  3  4  4  2  4  4  5  4  5  2  2  3  4  2  3  3 39 33
5    5  5  3  3  4  2  2  3  3  3  4  1  3  4  4  3  2  4  3  4  5  1 28 33
6    6  6  2  1  4  1  1  3  3  2  5  1  3  2  4  2  2  5  4  3  5  2 23 32
  統計 数学 批判的思考力 国語 自己効力感
1    51   48           28   72         61
2    74   53           26   66         53
3    48   60           35   71         48
4    67   68           27   67         48
5    55   49           30   66         49
6    74   63           36   83         37
>
```

```
> #変数リスト名
> inames <- c("x1", "x2", "x3", "x4", "x5",
+            "x6", "x7", "x8", "x9", "x10")
>
>
```

```
> # 級内相関係数の推定
> library(psych)
> ICC(d1[, inames])
Call: ICC(x = d1[, inames])
```

```
Intraclass correlation coefficients
```

	type	ICC	F	df1	df2	p	lower bound	upper bound
Single_raters_absolute	ICC1	0.18	3.2	364	3285	0	0.15	0.22
Single_random_raters	ICC2	0.21	6.3	364	3276	0	0.13	0.29
Single_fixed_raters	ICC3	0.34	6.3	364	3276	0	0.31	0.39
Average_raters_absolute	ICC1k	0.69	3.2	364	3285	0	0.64	0.73
Average_random_raters	ICC2k	0.73	6.3	364	3276	0	0.61	0.81
Average_fixed_raters	ICC3k	<u>0.84</u>	6.3	364	3276	0	0.81	0.86

Number of subjects = 365 Number of Judges = 10>

```
> # ICC3Kはα係数に等しい
> library(psych)
> alpha(d1[, inames], check.keys=FALSE)
```

```
Reliability analysis
Call: alpha(x = d1[, inames], check.keys = FALSE)
```

```
raw_alpha std.alpha G6(smc) average_r S/N ase mean sd
  0.84      0.84      0.84      0.35 5.3 0.02  3.1 0.61

lower alpha upper      95% confidence boundaries
0.8 0.84 0.88
```