

多重指標モデル

```
library(sem)
opt <- options(fit.indices = c("GFI", "AGFI", "RMSEA", "NFI", "NNFI", "CFI", "RNI", "IFI",
                               "SRMR", "AIC", "AICc", "BIC", "CAIC"))
```

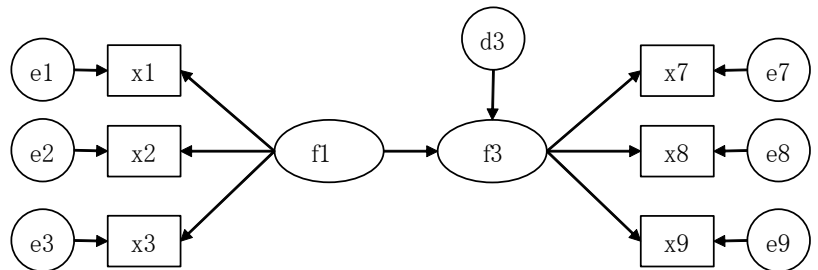
```
モデル名 <- specifyEquations()
構造方程式
```

```
sem(モデル名, 共分散(相関)行列, N=標本サイズ)
```

あらかじめsemパッケージをインストールしておく必要がある。
 適合度指標は、デフォルトではAICくらいしか出力されない。必要なものを出力するように指定する。
 パス係数の頭文字は b にしておくのが無難。他の文字を用いると、それだけで不具合になる場合がある。
 空白行があるところで、モデルの設定が修了したとみなされる。
 (逆に、行をあけないと、モデル設定が終了したことになる)

モデル部分のスキプットの例

```
seq.1 <- specifyEquations()
x1 = 1*f1
x2 = b2*f1
x3 = b3*f1
x7 = 1*f3
x8 = b8*f3
x9 = b9*f3
f3 = a31*f1
V(x1) = ev1
V(x2) = ev2
V(x3) = ev3
V(x7) = ev7
V(x8) = ev8
V(x9) = ev9
V(f3) = evf3
V(f1) = vf1
```



```
> setwd("i:¥¥documents¥¥scripts¥¥")
> d1 <- read.table("拡張SEMデータ.csv", header=TRUE, sep=",")
> head(d1)
  x1 x2 x3 x4 x5 x6 x7 x8 x9 x10 x11 x12
1  4  3  2  3  4  4  2  4  2   4   4   3
2  2  3  2  5  3  4  5  3  3   4   4   4
3  1  4  2  2  3  3  4  3  2   3   4   4
4  1  5  2  2  3  2  3  2  2   5   4   3
5  2  3  1  4  2  1  3  3  2   1   1   1
6  2  4  3  3  5  4  4  3  4   2   4   1
>
>
> # 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	x1	x2	x3	x4	x5	x6	x7	x8	x9	x10	x11	x12
x1	276	3.01	0.98	1.00	0.28	0.19	0.27	-0.01	0.03	0.11	0.07	0.08	0.06	0.07	0.07
x2	276	3.01	1.03	0.28	1.00	0.22	-0.01	0.02	0.04	0.28	0.06	0.09	0.08	0.09	0.10
x3	276	3.00	0.99	0.19	0.22	1.00	-0.01	0.03	0.01	0.12	0.08	0.09	0.21	0.01	0.07
x4	276	2.99	1.02	0.27	-0.01	-0.01	1.00	0.20	0.23	0.03	0.04	0.02	0.09	0.08	0.08

```

x5 276 3.00 1.00 -0.01 0.02 0.03 0.20 1.00 0.21 0.05 0.23 0.02 0.08 0.08 0.10
x6 276 3.03 1.04 0.03 0.04 0.01 0.23 0.21 1.00 0.09 0.04 0.07 0.09 0.29 0.11
x7 276 3.03 0.99 0.11 0.28 0.12 0.03 0.05 0.09 1.00 0.21 0.23 0.06 0.05 0.07
x8 276 3.01 1.01 0.07 0.06 0.08 0.04 0.23 0.04 0.21 1.00 0.24 0.09 0.03 0.11
x9 276 3.00 1.01 0.08 0.09 0.09 0.02 0.02 0.07 0.23 0.24 1.00 0.06 0.04 0.22
x10 276 2.99 1.03 0.06 0.08 0.21 0.09 0.08 0.09 0.06 0.09 0.06 1.00 0.17 0.22
x11 276 2.99 1.00 0.07 0.09 0.01 0.08 0.08 0.29 0.05 0.03 0.04 0.17 1.00 0.24
x12 276 2.99 1.01 0.07 0.10 0.07 0.08 0.10 0.11 0.07 0.11 0.22 0.22 0.24 1.00
>
>
> # 共分散行列・相関係数行列
> cov.d1 <- cov(d1)
> cor.d1 <- cor(d1)
>
>
> # 多重指標モデル
> library(sem)
> opt <- options(fit.indices = c("GFI", "AGFI", "RMSEA", "NFI", "NNFI", "CFI", "RNI", "IFI", "SRM
R", "AIC", "AICc", "BIC", "CAIC"))

```

```
seq.1 <- specifyEquations()
```

```

1: x1 = 1*f1
2: x2 = b2*f1
3: x3 = b3*f1
4: x7 = 1*f3
5: x8 = b8*f3
6: x9 = b9*f3
7: f3 = a31*f1
8: V(x1) = ev1
9: V(x2) = ev2
10: V(x3) = ev3
11: V(x7) = ev7
12: V(x8) = ev8
13: V(x9) = ev9
14: V(f3) = evf3
15: V(f1) = vf1
16:

```

Read 15 items

```

> sem.seq.1 <- sem(seq.1, cov.d1, N=nrow(d1))
> summary(sem.seq.1)

```

	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

```

Model Chisquare = 9.187096 Df = 8 Pr(>Chisq) = 0.3267596
Goodness-of-fit index = 0.9885242
Adjusted goodness-of-fit index = 0.969876
RMSEA index = 0.02322905 90% CI: (NA, 0.07683467)
Bentler-Bonett NFI = 0.9125166
Tucker-Lewis NNFI = 0.975273
Bentler CFI = 0.9868123
Bentler RNI = 0.9868123
Bollen IFI = 0.9877638
SRMR = 0.03268035
AIC = 35.1871
AICc = 10.57641
BIC = -35.77611
CAIC = -43.77611

```

```

Normalized Residuals
Min. 1st Qu. Median Mean 3rd Qu. Max.
-1.123000 -0.332200 -0.005191 -0.031210 0.056290 1.362000

```

```

R-square for Endogenous Variables
x1 x2 x3 f3 x7 x8 x9
0.1912 0.3961 0.1360 0.2738 0.3787 0.1409 0.1694

```

```

Parameter Estimates
Estimate Std Error z value Pr(>|z|)
b2 1.5186975 0.43869639 3.461842 5.364926e-04 x2 <--- f1
b3 0.8561944 0.25021723 3.421804 6.220703e-04 x3 <--- f1
b8 0.6234700 0.19324808 3.226267 1.254161e-03 x8 <--- f3
b9 0.6838182 0.20630464 3.314604 9.177304e-04 x9 <--- f3
a31 0.7444214 0.23023345 3.233333 1.223549e-03 f3 <--- f1

```

```

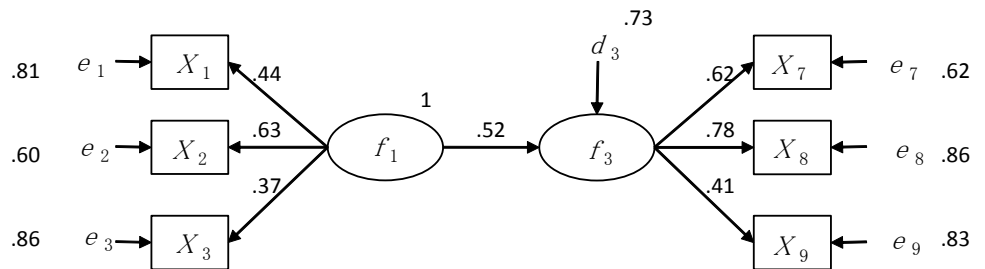
ev1  0.7762467 0.08607366  9.018401 1.908553e-19 x1 <--> x1
ev2  0.6455500 0.13005700  4.963593 6.920086e-07 x2 <--> x2
ev3  0.8545419 0.08520639 10.029082 1.135676e-23 x3 <--> x3
ev7  0.6094471 0.11954549  5.098035 3.431974e-07 x7 <--> x7
ev8  0.8808253 0.08941441  9.851044 6.783564e-23 x8 <--> x8
ev9  0.8517256 0.09136838  9.321886 1.142907e-20 x9 <--> x9
evf3 0.2698155 0.11038347  2.444347 1.451145e-02 f3 <--> f3
vfl  0.1835426 0.07310119  2.510801 1.204575e-02 f1 <--> f1

```

```

Iterations = 31
> stdCoef(sem.seq.1)
      Std. Estimate
1      0.4373010 x1 <--- f1
2      b2      0.6293252 x2 <--- f1
3      b3      0.3688269 x3 <--- f1
4      0.6154132 x7 <--- f3
5      b8      0.3753166 x8 <--- f3
6      b9      0.4116027 x9 <--- f3
7      a31     0.5232285 f3 <--- f1
8      ev1     0.8087679 x1 <--> x1
9      ev2     0.6039498 x2 <--> x2
10     ev3     0.8639667 x3 <--> x3
11     ev7     0.6212666 x7 <--> x7
12     ev8     0.8591374 x8 <--> x8
13     ev9     0.8305832 x9 <--> x9
14     evf3    0.7262319 f3 <--> f3
15     vfl     1.0000000 f1 <--> f1
>
>
>
>

```



$\chi^2=9.19$, $df=8$, $p=0.327$,
 AGFI=0.97, RMSEA=0.02, CFI=0.99

潜在変数の構造方程式モデリング — semパッケージを使う方法

```
library(sem)
opt <- options(fit.indices = c("GFI", "AGFI", "RMSEA", "NFI", "NNFI", "CFI", "RNI", "IFI",
                               "SRMR", "AIC", "AICc", "BIC", "CAIC"))
```

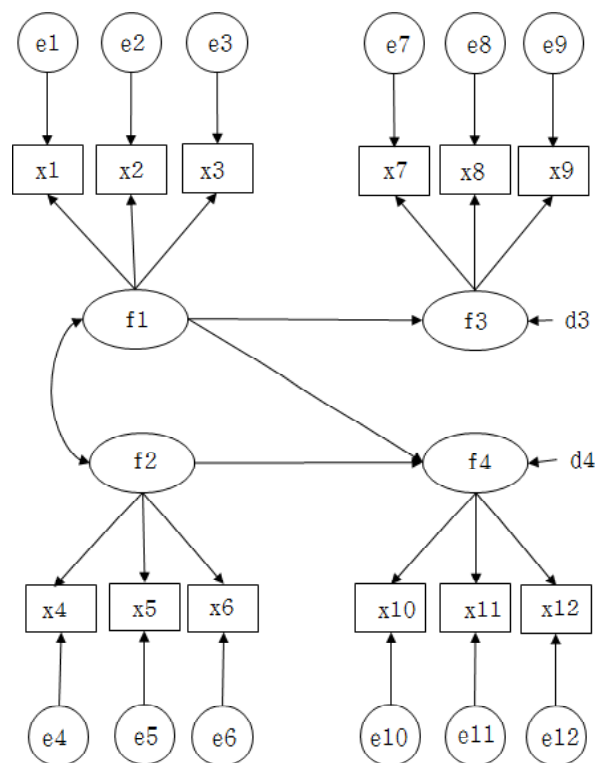
```
モデル名 <- specifyEquations()
構造方程式
```

```
sem(モデル名, 共分散(相関)行列, N=標本サイズ)
```

あらかじめsemパッケージをインストールしておく必要がある。
 適合度指標は、デフォルトではAICくらいしか出力されない。必要なものを出力するように指定する。
 パス係数の頭文字は b にしておくのが無難。他の文字を用いると、それだけで不具合になる場合がある。
 空白行があるところで、モデルの設定が修了したとみなされる。
 (逆に、行をあげないと、モデル設定が終了したことにならない)

モデル部分のスキプトの例

```
seq.1 <- specifyEquations()
x1 = 1*f1                                # 測定方程式
x2 = b2*f1
x3 = b3*f1
x4 = 1*f2
x5 = b5*f2
x6 = b6*f2
x7 = 1*f3
x8 = b8*f3
x9 = b9*f3
x10 = 1*f4
x11 = b11*f4
x12 = b12*f4
f3 = a31*f1                                # 潜在変数間の予測式
f4 = a41*f1 + a42*f2
V(x1) = ev1                                # 内生変数の誤差分散
V(x2) = ev2
V(x3) = ev3
V(x4) = ev4
V(x5) = ev5
V(x6) = ev6
V(x7) = ev7
V(x8) = ev8
V(x9) = ev9
V(x10) = ev10
V(x11) = ev11
V(x12) = ev12
V(f3) = evf3
V(f4) = evf4
V(f1) = vf1                                # 外生変数の分散
V(f2) = vf2
C(f1, f2) = cf12                           # 外生変数の共分散
```



```
> setwd("i:¥¥Rdocuments¥¥scripts¥¥")
> d1 <- read.table("拡張SEMデータ.csv", header=TRUE, sep=",")
> head(d1)
  x1 x2 x3 x4 x5 x6 x7 x8 x9 x10 x11 x12
1  4  3  2  3  4  4  2  4  2  4  4  3
2  2  3  2  5  3  4  5  3  3  4  4  4
3  1  4  2  2  3  3  4  3  2  3  4  4
4  1  5  2  2  3  2  3  2  2  5  4  3
5  2  3  1  4  2  1  3  3  2  1  1  1
6  2  4  3  3  5  4  4  3  4  2  4  1
>
>
```

```

> # 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  x1  x2  x3  x4  x5  x6  x7  x8  x9  x10  x11  x12
x1 276 3.01 0.98 1.00 0.28 0.19 0.27 -0.01 0.03 0.11 0.07 0.08 0.06 0.07 0.07
x2 276 3.01 1.03 0.28 1.00 0.22 -0.01 0.02 0.04 0.28 0.06 0.09 0.08 0.09 0.10
x3 276 3.00 0.99 0.19 0.22 1.00 -0.01 0.03 0.01 0.12 0.08 0.09 0.21 0.01 0.07
x4 276 2.99 1.02 0.27 -0.01 -0.01 1.00 0.20 0.23 0.03 0.04 0.02 0.09 0.08 0.08
x5 276 3.00 1.00 -0.01 0.02 0.03 0.20 1.00 0.21 0.05 0.23 0.02 0.08 0.08 0.10
x6 276 3.03 1.04 0.03 0.04 0.01 0.23 0.21 1.00 0.09 0.04 0.07 0.09 0.29 0.11
x7 276 3.03 0.99 0.11 0.28 0.12 0.03 0.05 0.09 1.00 0.21 0.23 0.06 0.05 0.07
x8 276 3.01 1.01 0.07 0.06 0.08 0.04 0.23 0.04 0.21 1.00 0.24 0.09 0.03 0.11
x9 276 3.00 1.01 0.08 0.09 0.09 0.02 0.02 0.07 0.23 0.24 1.00 0.06 0.04 0.22
x10 276 2.99 1.03 0.06 0.08 0.21 0.09 0.08 0.09 0.06 0.09 0.06 1.00 0.17 0.22
x11 276 2.99 1.00 0.07 0.09 0.01 0.08 0.08 0.29 0.05 0.03 0.04 0.17 1.00 0.24
x12 276 2.99 1.01 0.07 0.10 0.07 0.08 0.10 0.11 0.07 0.11 0.22 0.22 0.24 1.00
>
>
> # 共分散行列・相関係数行列
> cov.d1 <- cov(d1)
> cor.d1 <- cor(d1)
>
>
> # 因子間構造モデル
> library(sem)
> opt <- options(fit.indices = c("GFI", "AGFI", "RMSEA", "NFI", "NNFI", "CFI", "RNI", "IFI", "SRMR", "AIC", "AICc", "BIC", "CAIC"))
>
> seq.1 <- specifyEquations()
1: x1 = 1*f1
2: x2 = b2*f1
3: x3 = b3*f1
4: x4 = 1*f2
5: x5 = b5*f2
6: x6 = b6*f2
7: x7 = 1*f3
8: x8 = b8*f3
9: x9 = b9*f3
10: x10 = 1*f4
11: x11 = b11*f4
12: x12 = b12*f4
13: f3 = a31*f1
14: f4 = a41*f1 + a42*f2
15: V(x1) = ev1
16: V(x2) = ev2
17: V(x3) = ev3
18: V(x4) = ev4
19: V(x5) = ev5
20: V(x6) = ev6
21: V(x7) = ev7
22: V(x8) = ev8
23: V(x9) = ev9
24: V(x10) = ev10
25: V(x11) = ev11
26: V(x12) = ev12
27: V(f3) = evf3
28: V(f4) = evf4
29: V(f1) = vf1
30: V(f2) = vf2
31: C(f1, f2) = cf12
32:
Read 31 items

> sem.seq.1 <- sem(seq.1, cor.d1, N=nrow(d1))
> summary(sem.seq.1)

```

	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

Model Chisquare = 83.02649 Df = 50 Pr(>Chisq) = 0.002305156
 Goodness-of-fit index = 0.953118
 Adjusted goodness-of-fit index = 0.9268641
 RMSEA index = 0.04900945 90% CI: (0.02933125, 0.06721064)
 Bentler-Bonett NFI = 0.6987328
 Tucker-Lewis NNFI = 0.7919997
 Bentler CFI = 0.842424
 Bentler RNI = 0.842424
 Bollen IFI = 0.8536001
 SRMR = 0.05428159
 AIC = 139.0265
 AICc = 89.60139
 BIC = -197.9936
 CAIC = -247.9936

Normalized Residuals

Min.	1st Qu.	Median	Mean	3rd Qu.	Max.
-1.1980	-0.4154	-0.1110	0.1063	0.2682	3.7710

R-square for Endogenous Variables

x1	x2	x3	x4	x5	x6	f3	x7	x8	x9	f4	x10	x11
0.2005	0.3366	0.1483	0.1676	0.1510	0.3358	0.3151	0.3396	0.1550	0.1890	0.3986	0.1587	0.2398
x12												
0.2356												

Parameter Estimates

	Estimate	Std Error	z value	Pr(> z)	
b2	1.29569918	0.34075086	3.802482	1.432534e-04	x2 <--- f1
b3	0.86015848	0.24459171	3.516712	4.369283e-04	x3 <--- f1
b5	0.94917157	0.29318375	3.237463	1.205975e-03	x5 <--- f2
b6	1.41554193	0.43270252	3.271397	1.070174e-03	x6 <--- f2
b8	0.67557465	0.19781491	3.415186	6.373854e-04	x8 <--- f3
b9	0.74592556	0.21220194	3.515168	4.394752e-04	x9 <--- f3
b11	1.22951694	0.35459376	3.467396	5.255274e-04	x11 <--- f4
b12	1.21854597	0.35180875	3.463660	5.328791e-04	x12 <--- f4
a31	0.73053457	0.22194276	3.291545	9.963869e-04	f3 <--- f1
a41	0.26687006	0.13446825	1.984633	4.718537e-02	f4 <--- f1
a42	0.48420288	0.19075744	2.538317	1.113870e-02	f4 <--- f2
ev1	0.79951327	0.08802143	9.083166	1.054612e-19	x1 <--> x1
ev2	0.66341533	0.10181173	6.516099	7.215927e-11	x2 <--> x2
ev3	0.85166528	0.08572231	9.935165	2.926898e-23	x3 <--> x3
ev4	0.83243627	0.08985916	9.263788	1.973045e-20	x4 <--> x4
ev5	0.84903732	0.08863006	9.579564	9.745538e-22	x5 <--> x5
ev6	0.66424202	0.11484738	5.783693	7.307822e-09	x6 <--> x6
ev7	0.66040641	0.10968439	6.020970	1.733749e-09	x7 <--> x7
ev8	0.84500901	0.08741897	9.666197	4.196854e-22	x8 <--> x8
ev9	0.81104876	0.08962717	9.049139	1.440999e-19	x9 <--> x9
ev10	0.84134121	0.08760173	9.604162	7.678009e-22	x10 <--> x10
ev11	0.76015347	0.09371818	8.111057	5.018146e-16	x11 <--> x11
ev12	0.76441388	0.09327267	8.195475	2.496052e-16	x12 <--> x12
evf3	0.23259733	0.09569770	2.430542	1.507624e-02	f3 <--> f3
evf4	0.09541110	0.04785396	1.993797	4.617419e-02	f4 <--> f4
vf1	0.20048687	0.07564876	2.650234	8.043615e-03	f1 <--> f1
vf2	0.16756383	0.07258184	2.308619	2.096472e-02	f2 <--> f2
cf12	0.03747028	0.02461593	1.522197	1.279598e-01	f2 <--> f1

Iterations = 47

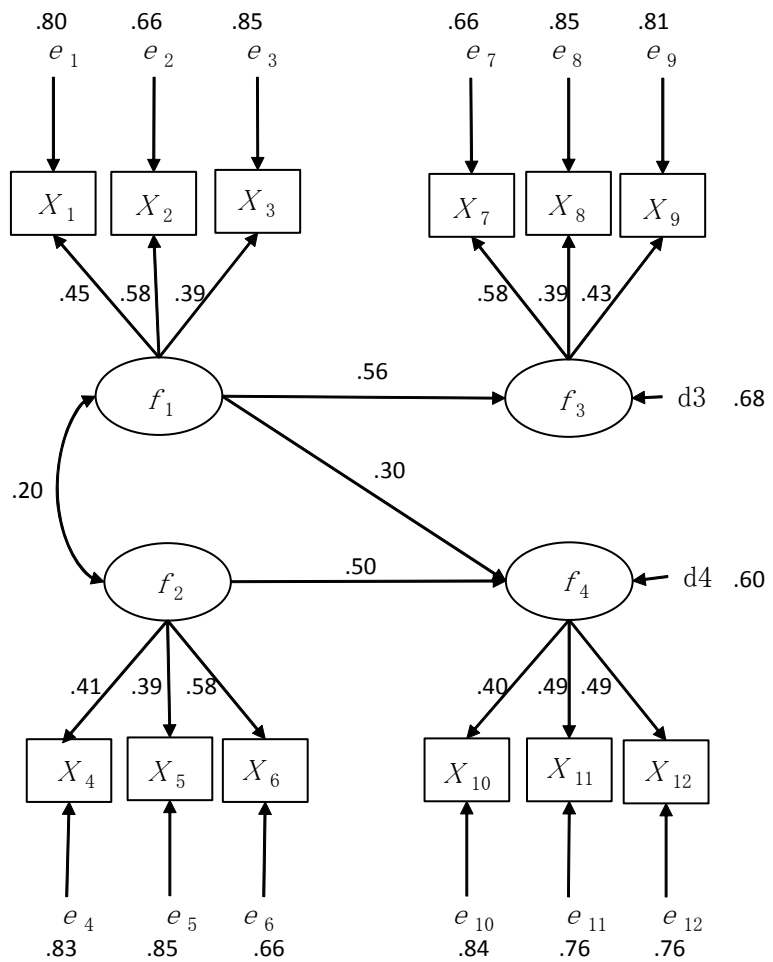
> stdCoef(sem. seq. 1)

	Std. Estimate	
1	0.4477576	x1 <--- f1
2	0.5801592	x2 <--- f1
3	0.3851425	x3 <--- f1
4	0.4093456	x4 <--- f2
5	0.3885392	x5 <--- f2
6	0.5794460	x6 <--- f2
7	0.5827464	x7 <--- f3
8	0.3936887	x8 <--- f3
9	0.4346854	x9 <--- f3
10	0.3983204	x10 <--- f4
11	0.4897417	x11 <--- f4
12	0.4853719	x12 <--- f4

```

13 a31      0.5613118    f3 <--- f1
14 a41      0.2999924    f4 <--- f1
15 a42      0.4976051    f4 <--- f2
16 ev1      0.7995132    x1 <--> x1
17 ev2      0.6634154    x2 <--> x2
18 ev3      0.8516653    x3 <--> x3
19 ev4      0.8324362    x4 <--> x4
20 ev5      0.8490373    x5 <--> x5
21 ev6      0.6642423    x6 <--> x6
22 ev7      0.6604066    x7 <--> x7
23 ev8      0.8450092    x8 <--> x8
24 ev9      0.8110486    x9 <--> x9
25 ev10     0.8413409    x10 <--> x10
26 ev11     0.7601531    x11 <--> x11
27 ev12     0.7644141    x12 <--> x12
28 evf3     0.6849291    f3 <--> f3
29 evf4     0.6013587    f4 <--> f4
30 vf1      1.0000000    f1 <--> f1
31 vf2      1.0000000    f2 <--> f2
32 cf12     0.2044343    f2 <--> f1
>
>
>

```



$\chi^2=83.03, df=50, p=0.002,$
AGFI=0.93, RMSEA=0.05, CFI=0.84

潜在変数の構造方程式モデリング — lavaanパッケージを使う方法

パッケージの読み込み

```
library(lavaan)
```

あらかじめlavaanパッケージをインストールしておく必要がある。

モデルの設定

```
モデル名 <- '
# 測定変数の指定 (= ~ を使う)
潜在変数名1 ~ 観測変数1 + 観測変数2 + ...
潜在変数名2 ~ 観測変数1 + 観測変数2 + ...

# 回帰式 (~ を使う)
従属変数名1 ~ 説明変数1 + 説明変数2 + ...
従属変数名2 ~ 説明変数1 + 説明変数2 + ...

# 分散, 共分散 (~ ~ を使う)
変数名i ~ 変数名i      # 分散
, 変数名j ~ 変数名k      # 共分散
,
```

モデル部分はシングルカンマ (') でくくる。

すべての変数の分散を推定する場合は分散の指定は省略できる。

パラメタ値の推定

```
lavaanオブジェクト名 <- sem(モデル名, data=データ名)      または
```

```
lavaanオブジェクト名 <- lavaan(モデル名, data=データ名, model.type="sem",
                                auto.var=TRUE, auto.fix.first=TRUE, auto.cov.lv.x=FALSE, auto.cov.y=FALSE)
```

auto.var=TRUE : すべての変数の分散または残差分散を推定する

auto.fix.first=TRUE : 非標準化解において測定変数のうち最初の1つのパス係数を1に固定する

auto.cov.lv.x=FALSE : 潜在変数間の共分散を自動的に推定しない

auto.cov.y=FALSE : 従属変数の残差間の共分散を自動的に推定しない

```
summary(lavaanオブジェクト名, fit.measures=TRUE, standardized=TRUE)
```

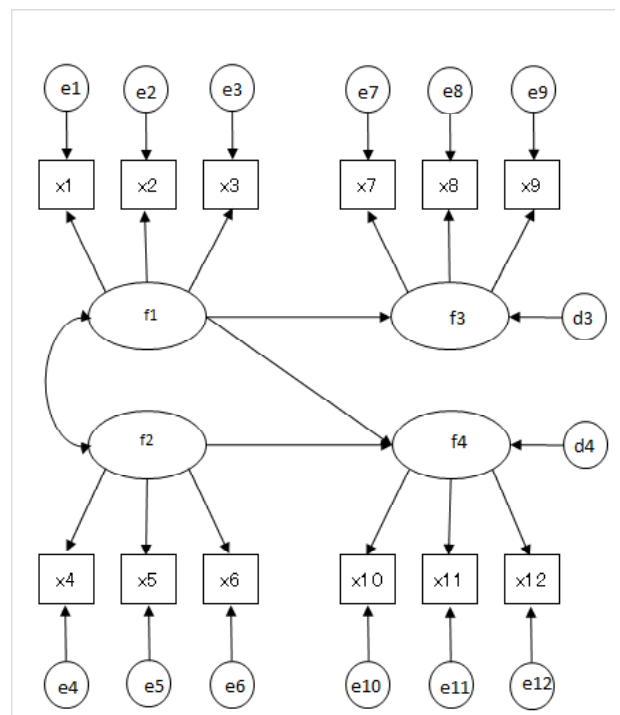
モデル部分のスキプトの例

```
model.1 <- '
# latent variable definitions
f1 ~ x1 + x2 + x3
f2 ~ x4 + x5 + x6
f3 ~ x7 + x8 + x9
f4 ~ x10 + x11 + x12

# regression
f3 ~ f1
f4 ~ f1 + f2

# variances and covariances
f1 ~ f2
,
```

すべての分散を推定するので、分散の式は省略している



```
> setwd("d:¥¥")
> d1 <- read.table("共分散構造分析_データ.csv", header=TRUE, sep=",")
> head(d1)
  x1 x2 x3 x4 x5 x6 x7 x8 x9 x10 x11 x12
1  4  3  2  3  4  4  2  4  2   4   4   3
2  2  3  2  5  3  4  5  3  3   4   4   4
3  1  4  2  2  3  3  4  3  2   3   4   4
4  1  5  2  2  3  2  3  2  2   5   4   3
5  2  3  1  4  2  1  3  3  2   1   1   1
6  2  4  3  3  5  4  4  3  4   2   4   1
>
```

```
> #データフレームの行数（標本の大きさ）、平均、標準偏差、共分散、相関
> n.d1 <- nrow(d1)
> mean.d1 <- apply(d1, 2, mean)
> sd.d1 <- apply(d1, 2, sd)
> cor.d1 <- cor(d1)
> round(data.frame(n.d1, mean.d1, sd.d1, cor.d1), 2)
```

```
   n.d1 mean.d1 sd.d1   x1   x2   x3   x4   x5   x6   x7   x8   x9
x1    276    3.01  0.98  1.00  0.28  0.19  0.27 -0.01  0.03  0.11  0.07  0.08
x2    276    3.01  1.03  0.28  1.00  0.22 -0.01  0.02  0.04  0.28  0.06  0.09
x3    276    3.00  0.99  0.19  0.22  1.00 -0.01  0.03  0.01  0.12  0.08  0.09
x4    276    2.99  1.02  0.27 -0.01 -0.01  1.00  0.20  0.23  0.03  0.04  0.02
x5    276    3.00  1.00 -0.01  0.02  0.03  0.20  1.00  0.21  0.05  0.23  0.02
x6    276    3.03  1.04  0.03  0.04  0.01  0.23  0.21  1.00  0.09  0.04  0.07
x7    276    3.03  0.99  0.11  0.28  0.12  0.03  0.05  0.09  1.00  0.21  0.23
x8    276    3.01  1.01  0.07  0.06  0.08  0.04  0.23  0.04  0.21  1.00  0.24
x9    276    3.00  1.01  0.08  0.09  0.09  0.02  0.02  0.07  0.23  0.24  1.00
x10   276    2.99  1.03  0.06  0.08  0.21  0.09  0.08  0.09  0.06  0.09  0.06
x11   276    2.99  1.00  0.07  0.09  0.01  0.08  0.08  0.29  0.05  0.03  0.04
x12   276    2.99  1.01  0.07  0.10  0.07  0.08  0.10  0.11  0.07  0.11  0.22
  x10 x11 x12
x1  0.06 0.07 0.07
x2  0.08 0.09 0.10
x3  0.21 0.01 0.07
x4  0.09 0.08 0.08
x5  0.08 0.08 0.10
x6  0.09 0.29 0.11
x7  0.06 0.05 0.07
x8  0.09 0.03 0.11
x9  0.06 0.04 0.22
x10 1.00 0.17 0.22
x11 0.17 1.00 0.24
x12 0.22 0.24 1.00
>
>
```

```
> #共分散構造分析の実行
> #lavaanパッケージの読み込み
> library(lavaan)
>
> # モデルの設定
> model.1 <- '
+   # latent variable definitions
+   f1 =~ x1 + x2 + x3
+   f2 =~ x4 + x5 + x6
+   f3 =~ x7 + x8 + x9
+   f4 =~ x10 + x11 + x12
+
+   # regression
+   f3 =~ f1
+   f4 =~ f1 + f2
+
+   # variances and covariances
+   , f1 ~~ f2
+
>
```

	A	B	C	D	E	F	G	H	I	J	K	L
1	x1	x2	x3	x4	x5	x6	x7	x8	x9	x10	x11	x12
2	4	3	2	3	4	4	2	4	2	4	4	3
3	2	3	2	5	3	4	5	3	3	4	4	4
4	1	4	2	2	3	3	4	3	2	3	4	4
5	1	5	2	2	3	2	3	2	2	5	4	3
6	2	3	1	4	2	1	3	3	2	1	1	1
7	2	4	3	3	5	4	4	3	4	2	4	1
8	4	4	4	4	4	3	2	1	2	2	2	1
9	3	4	3	3	3	4	3	4	3	4	4	3
10	3	4	3	2	3	5	4	4	3	2	5	4
11	3	5	2	2	4	4	4	4	5	4	1	3
12	4	2	3	4	4	5	3	4	2	5	3	2
13	4	4	2	3	4	2	2	4	4	3	3	4
14	3	4	3	2	3	3	4	4	4	4	2	5
15	4	3	5	4	2	3	3	2	2	4	3	3
16	1	4	3	2	3	2	4	1	3	2	3	2
17	4	4	3	2	3	4	4	4	3	3	3	3
18	3	2	2	4	1	4	4	4	4	3	2	2
19	4	4	4	4	4	3	3	4	2	2	1	3
20	3	3	3	3	3	4	2	3	4	4	3	3
21	4	2	5	2	2	2	4	2	4	4	3	5

```
> # lavaan関数を使って計算
> lav.model.1 <- lavaan(model.1, data=d1, model.type="sem",
+                         fixed.x=F, meanstructure=F,
+                         auto.var=TRUE, auto.fix.first=TRUE)
> summary(lav.model.1, fit.measures=TRUE, standardized=TRUE)
lavaan (0.5-10) converged normally after 60 iterations
```

Number of observations	276
Estimator	ML
Minimum Function Chi-square	83.328
Degrees of freedom	50
P-value	0.002

Chi-square test baseline model:

Minimum Function Chi-square	276.593
Degrees of freedom	66
P-value	0.000

Full model versus baseline model:

Comparative Fit Index (CFI)	0.842
Tucker-Lewis Index (TLI)	0.791

Loglikelihood and Information Criteria:

Loglikelihood user model (H0)	-4630.239
Loglikelihood unrestricted model (H1)	-4588.575
Number of free parameters	28
Akaike (AIC)	9316.478
Bayesian (BIC)	9417.849
Sample-size adjusted Bayesian (BIC)	9329.066

Root Mean Square Error of Approximation:

RMSEA	0.049
90 Percent Confidence Interval	0.030 0.067
P-value RMSEA <= 0.05	0.508

Standardized Root Mean Square Residual:

SRMR	0.054
------	-------

Parameter estimates:

	Information Standard Errors	Expected Standard					
	Estimate	Std. err	Z-value	P(> z)	Std. lv	Std. all	
Latent variables:							
f1 =							
x1	1.000				0.438	0.448	
x2	1.367	0.359	3.809	0.000	0.599	0.580	
x3	0.873	0.248	3.523	0.000	0.382	0.385	
f2 =							
x4	1.000				0.418	0.409	
x5	0.931	0.287	3.243	0.001	0.389	0.389	
x6	1.442	0.440	3.277	0.001	0.603	0.579	
f3 =							
x7	1.000				0.576	0.583	
x8	0.691	0.202	3.421	0.001	0.398	0.394	
x9	0.763	0.217	3.522	0.000	0.439	0.435	
f4 =							
x10	1.000				0.408	0.398	
x11	1.193	0.343	3.474	0.001	0.487	0.490	
x12	1.195	0.344	3.470	0.001	0.488	0.485	

Regressions:

f3 =						
f1	0.739	0.224	3.298	0.001	0.561	0.561

f4 ~						
f1	0.280	0.141	1.988	0.047	0.300	0.300
f2	0.486	0.191	2.543	0.011	0.498	0.498
Covariances:						
f1						
f2	0.037	0.025	1.525	0.127	0.204	0.204
Variances:						
x1	0.765	0.084			0.765	0.800
x2	0.707	0.108			0.707	0.663
x3	0.839	0.084			0.839	0.852
x4	0.868	0.094			0.868	0.832
x5	0.852	0.089			0.852	0.849
x6	0.719	0.124			0.719	0.664
x7	0.645	0.107			0.645	0.660
x8	0.863	0.089			0.863	0.845
x9	0.829	0.091			0.829	0.811
x10	0.884	0.092			0.884	0.841
x11	0.752	0.093			0.752	0.760
x12	0.773	0.094			0.773	0.764
f1	0.192	0.072			1.000	1.000
f2	0.175	0.076			1.000	1.000
f3	0.227	0.093			0.685	0.685
f4	0.100	0.050			0.601	0.601

>

成長曲線モデル — lavaanパッケージを使う方法

パッケージの読み込み
library(lavaan)

あらかじめlavaanパッケージをインストールしておく必要がある。

モデルの設定
モデル名 <- '

```
# 測定変数の指定 (= ~ を使う)
切片潜在変数名 =~ 1*観測変数1 + 1*観測変数2 + 1*観測変数3 + ...
傾き潜在変数名 =~ 0*観測変数1 + 1*観測変数2 + 2*観測変数3 + ...

# 回帰式 (~ を使う)
切片潜在変数名 ~ 説明変数1 + 説明変数2 + ...
傾き潜在変数名 ~ 説明変数1 + 説明変数2 + ...
観測変数名1 ~ 説明変数1 + 説明変数2 + ...

# 分散, 共分散 (~~ を使う)
変数名i ~~ 変数名i      # 分散
変数名j ~~ 変数名k      # 共分散
```

モデル部分はシングルカンマ (') でくくる。
すべての変数の分散を推定する場合は分散の指定は省略できる。

パラメタ値の推定

lavaanオブジェクト名 <- growth(モデル名, data=データ名) または

lavaanオブジェクト名 <- lavaan(モデル名, data=データ名, model.type="growth",
fixed.x=F, int.lv.free=TRUE, auto.var=TRUE)

fixed.x=FALSE : 外生変数となる観測変数の分散, 共分散, 平均を標本平均で固定する
int.lv.free = TRUE : 潜在変数の切片を推定する
auto.var=TRUE : すべての変数の分散または残差分散を推定する

summary(lavaanオブジェクト名, fit.measures=TRUE)

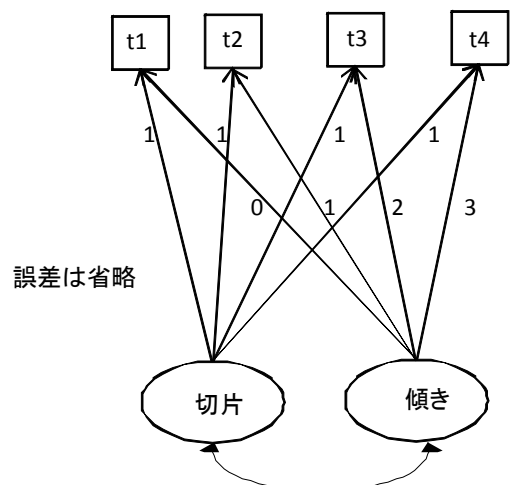
結果の解釈には非標準化解を用いるので standardized=TRUE はつけていない。
切片変数と傾き変数の相関を求めるために標準化解を求めることはある。

モデル部分のスキプトの例

```
model.1 <- '
# latent variable definitions
f.inter =~ 1*t1 + 1*t2 + 1*t3 + 1*t4
f.slope =~ 0*t1 + 1*t2 + 2*t3 + 3*t4

# regression

# variances and covariances
f.inter ~~ f.slope
,
# すべての分散を推定するので, 分散の式は省略している
```



```
> setwd("d:¥¥")
> d1 <- read.table("成長曲線モデル_データ.csv", header=TRUE, sep=",")
> head(d1)
  id t1 t2 t3 t4
1  1  3  6  8 12
2  2  8 12 17 22
3  3  8 11 13 16
4  4  7 13 20 25
5  5  7 13 20 25
6  6  7 14 20 26
>

> # id 列の削除
> d1 <- d1[,c(-1)]
>
```

```
> #データフレームの行数（標本の大きさ）、平均、標準偏差、共分散、相関
> n.d1 <- nrow(d1)
> mean.d1 <- apply(d1, 2, mean)
> sd.d1 <- apply(d1, 2, sd)
> cor.d1 <- cor(d1)
> round(data.frame(n.d1, mean.d1, sd.d1, cor.d1), 2)
  n.d1 mean.d1 sd.d1  t1  t2  t3  t4
t1    60     5.62  1.53 1.00 0.74 0.57 0.50
t2    60    10.10  2.20 0.74 1.00 0.93 0.91
t3    60    14.73  3.44 0.57 0.93 1.00 0.98
t4    60    19.15  4.68 0.50 0.91 0.98 1.00
>
```

	A	B	C	D	E
1	id	t1	t2	t3	t4
2	1	3	6	8	12
3	2	8	12	17	22
4	3	8	11	13	16
5	4	7	13	20	25
6	5	7	13	20	25
7	6	7	14	20	26
8	7	7	13	18	22
9	8	8	12	16	20
10	9	6	13	20	28
11	10	7	13	21	27
12	11	7	12	19	25
13	12	7	13	18	25
14	13	6	11	17	22
15	14	6	11	16	22
16	15	6	10	14	19
17	16	6	11	14	19
18	17	6	11	14	19
19	18	6	9	13	15
20	19	7	10	13	16
21	20	6	8	11	12

```
> #共分散構造分析の実行
> #lavaanパッケージの読み込み
> library(lavaan)
>
> # モデルの設定
> model.1 <- '
+ # latent variable definitions
+ f.inter =~ 1*t1 + 1*t2 + 1*t3 + 1*t4
+ f.slope =~ 0*t1 + 1*t2 + 2*t3 + 3*t4
+
+ # regression
+
+ # variances and covariances
+ ,f.inter =~ f.slope
+
>
```

```
> # growth関数を使って計算
> gro.model.1 <- growth(model.1, data=d1)
>
> summary(gro.model.1, fit.measures=TRUE)
lavaan (0.4-11) converged normally after 67 iterations
```

```
Number of observations                    60

Estimator                                ML
Minimum Function Chi-square              2.609
Degrees of freedom                       5
P-value                                 0.760
```

Chi-square test baseline model:

```
Minimum Function Chi-square              390.102
Degrees of freedom                       6
P-value                                 0.000
```

Full model versus baseline model:

```
Comparative Fit Index (CFI)              1.000
Tucker-Lewis Index (TLI)                 1.007
```

Loglikelihood and Information Criteria:

Loglikelihood user model (H0)	-384.343
Loglikelihood unrestricted model (H1)	-383.039
Number of free parameters	9
Akaike (AIC)	786.687
Bayesian (BIC)	805.536
Sample-size adjusted Bayesian (BIC)	777.228

Root Mean Square Error of Approximation:

RMSEA	0.000
90 Percent Confidence Interval	0.000 0.124
P-value RMSEA $\leq$ 0.05	0.813

Standardized Root Mean Square Residual:

SRMR	0.017
------	-------

Parameter estimates:

Information				Expected
Standard Errors				Standard
	Estimate	Std. err	Z-value	P(> z)
Latent variables:				
f. inter = ~				
t1	1.000			
t2	1.000			
t3	1.000			
t4	1.000			
f. slope = ~				
t1	0.000			
t2	1.000			
t3	2.000			
t4	3.000			

Covariances:

f. inter ~				
f. slope	0.482	0.261	1.843	0.065

Intercepts:

t1	0.000			
t2	0.000			
t3	0.000			
t4	0.000			
f. inter	5.610	0.185	30.268	0.000
f. slope	4.515	0.177	25.470	0.000

Variances:

t1	0.355	0.141
t2	0.164	0.060
t3	0.360	0.096
t4	0.044	0.187
f. inter	1.886	0.378
f. slope	1.858	0.346

```
> lav.model.1 <- lavaan(model.1, data=d1, model.type="growth",
+                         fixed.x=F, int.lv.free=TRUE, auto.var=TRUE)
> summary(lav.model.1, fit.measures=TRUE)
lavaan (0.5-10) converged normally after 67 iterations
```

Number of observations	60
Estimator	ML
Minimum Function Chi-square	2.609
Degrees of freedom	5
P-value	0.760

Chi-square test baseline model:

Minimum Function Chi-square	390.102
Degrees of freedom	6
P-value	0.000

Full model versus baseline model:

Comparative Fit Index (CFI)	1.000
Tucker-Lewis Index (TLI)	1.007

Loglikelihood and Information Criteria:

Loglikelihood user model (H0)	-384.343
Loglikelihood unrestricted model (H1)	-383.039
Number of free parameters	9
Akaike (AIC)	786.687
Bayesian (BIC)	805.536
Sample-size adjusted Bayesian (BIC)	777.228

Root Mean Square Error of Approximation:

RMSEA	0.000
90 Percent Confidence Interval	0.000 0.124
P-value RMSEA $\leq$ 0.05	0.813

Standardized Root Mean Square Residual:

SRMR	0.017
------	-------

Parameter estimates:

	Information			Expected
	Standard Errors			Standard
	Estimate	Std. err	Z-value	P(> z)
Latent variables:				
f. inter = ~				
t1	1.000			
t2	1.000			
t3	1.000			
t4	1.000			
f. slope = ~				
t1	0.000			
t2	1.000			
t3	2.000			
t4	3.000			
Covariances:				
f. inter ~				
f. slope	0.482	0.261	1.843	0.065
Intercepts:				
t1	0.000			
t2	0.000			
t3	0.000			
t4	0.000			
f. inter	5.610	0.185	30.268	0.000
f. slope	4.515	0.177	25.470	0.000
Variances:				
t1	0.355	0.141		
t2	0.164	0.060		
t3	0.360	0.096		
t4	0.044	0.187		
f. inter	1.886	0.378		
f. slope	1.858	0.346		

多母集団分析 — semパッケージを用いる方法

```
library(sem)
opt <- options(fit.indices = c("GFI", "AGFI", "RMSEA", "NFI", "NNFI", "CFI", "RNI", "IFI",
                               "SRMR", "AIC", "AICc", "BIC", "CAIC"))
```

```
モデル名1 <- specifyEquations()
  構造方程式
```

```
モデル名2 <- specifyEquations()
  構造方程式
```

```
...
```

```
mgオブジェクト名 <- multigroupModel(モデル名1, モデル名2, ..., groups=c("群名1", "群名2"...))
群分け変数名 <- factor(群分け変数名)
```

```
semオブジェクト名 <- sem(mgオブジェクト名, データ行列, N=標本サイズ)
summary(semオブジェクト名)
```

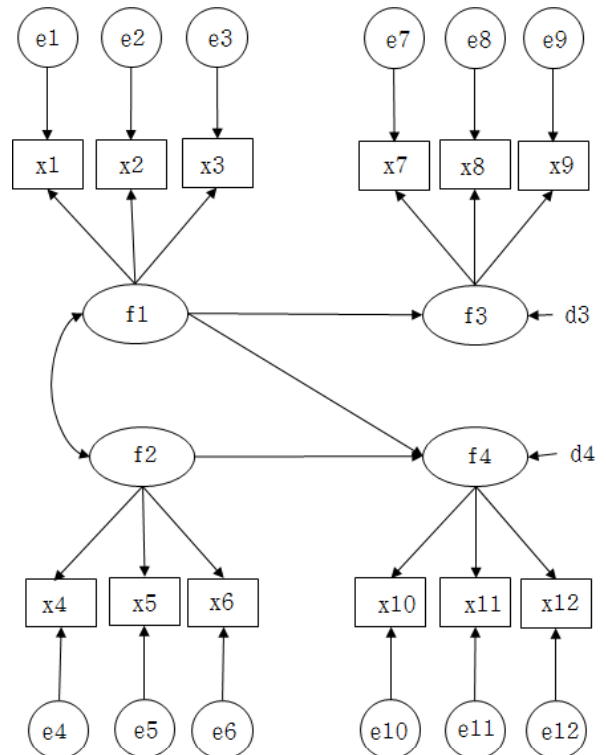
```
stdCoef(semオブジェクト名)
```

各群のモデルで、同じパラメタ名にすると等値制約になる
 群分け変数名はfactor関数を使ってfactor化しておかなければならない
 群分け変数名に「group」を用いると不具合が生じることがあるので他の名前を使う

あらかじめsemパッケージをインストールしておく必要がある。
 適合度指標は、デフォルトではAICくらいしか出力されない。必要なものを出力するように指定する。
 パス係数の頭文字は b にしておくのが無難。他の文字を用いると、それだけで不具合になる場合がある。
 空白行があるところで、モデルの設定が修了したとみなされる。
 (逆に、行をあげないと、モデル設定が終了したことになる)

モデル部分のスキプトの例

```
seq.1 <- specifyEquations()
x1 = 1*f1
x2 = b2.1*f1
x3 = b3.1*f1
x4 = 1*f2
x5 = b5.1*f2
x6 = b6.1*f2
x7 = 1*f3
x8 = b8.1*f3
x9 = b9.1*f3
x10 = 1*f4
x11 = b11.1*f4
x12 = b12.1*f4
f3 = a31.1*f1
f4 = a41.1*f1 + a42.1*f2
V(x1) = ev1.1
V(x2) = ev2.1
V(x3) = ev3.1
V(x4) = ev4.1
V(x5) = ev5.1
V(x6) = ev6.1
V(x7) = ev7.1
V(x8) = ev8.1
V(x9) = ev9.1
V(x10) = ev10.1
V(x11) = ev11.1
V(x12) = ev12.1
V(f3) = evf3.1
V(f4) = evf4.1
V(f1) = vf1.1
V(f2) = vf2.1
C(f1, f2) = cf12.1
```



```

seq.2 <- specifyEquations()
x1 = 1*f1
x2 = b2.2*f1
x3 = b3.2*f1
x4 = 1*f2
x5 = b5.2*f2
x6 = b6.2*f2
x7 = 1*f3
x8 = b8.2*f3
x9 = b9.2*f3
x10 = 1*f4
x11 = b11.2*f4
x12 = b12.2*f4
f3 = a31.2*f1
f4 = a41.2*f1 + a42.2*f2
V(x1) = ev1.2
V(x2) = ev2.2
V(x3) = ev3.2
V(x4) = ev4.2
V(x5) = ev5.2
V(x6) = ev6.2
V(x7) = ev7.2
V(x8) = ev8.2
V(x9) = ev9.2
V(x10) = ev10.2
V(x11) = ev11.2
V(x12) = ev12.2
V(f3) = evf3.2
V(f4) = evf4.2
V(f1) = vf1.2
V(f2) = vf2.2
C(f1, f2) = cf12.2

```

```

> setwd("d:¥¥")
> d1 <- read.table("多母集団分析_データ.csv", header=TRUE, sep=",")
> head(d1)
  id x1 x2 x3 x4 x5 x6 x7 x8 x9 x10 x11 x12 era
1  1  4  3  2  3  4  4  2  4  2  4  4  3 2010
2  2  2  3  2  5  3  4  5  3  3  4  4  4 2010
3  3  1  4  2  2  3  3  4  3  2  3  4  4 2000
4  4  1  5  2  2  3  2  3  2  2  5  4  3 2000
5  5  2  3  1  4  2  1  3  3  2  1  1  1 2000
6  6  2  4  3  3  5  4  4  3  4  2  4  1 2010

```

```

> # id 列の削除
> d1 <- d1[,c(-1)]
>

```

```

> #データフレームの行数（標本サイズ）、平均、標準偏差、共分散、相関
> n.d1 <- nrow(d1)
> mean.d1 <- apply(d1, 2, mean)
> sd.d1 <- apply(d1, 2, sd)
> cov.d1 <- cov(d1)
> cor.d1 <- cor(d1)
> round(data.frame(n.d1, mean.d1, sd.d1, cov.d1, cor.d1), 2)

```

	n.d1	mean.d1	sd.d1	x1	x2	x3	x4	x5	x6	x7	x8	x9	x10	x11	x12
x1	276	3.01	0.98	0.96	0.28	0.19	0.27	-0.01	0.03	0.11	0.07	0.08	0.06	0.07	0.07
x2	276	3.01	1.03	0.28	1.07	0.22	-0.01	0.03	0.04	0.28	0.06	0.09	0.08	0.09	0.11
x3	276	3.00	0.99	0.19	0.22	0.99	-0.01	0.03	0.01	0.12	0.08	0.09	0.21	0.01	0.07
x4	276	2.99	1.02	0.27	-0.01	-0.01	1.05	0.20	0.24	0.03	0.04	0.02	0.09	0.08	0.09
x5	276	3.00	1.00	-0.01	0.03	0.03	0.20	1.01	0.22	0.05	0.23	0.02	0.08	0.08	0.11
x6	276	3.03	1.04	0.03	0.04	0.01	0.24	0.22	1.09	0.09	0.05	0.07	0.09	0.30	0.12
x7	276	3.03	0.99	0.11	0.28	0.12	0.03	0.05	0.09	0.98	0.21	0.23	0.07	0.05	0.07
x8	276	3.01	1.01	0.07	0.06	0.08	0.04	0.23	0.05	0.21	1.03	0.24	0.09	0.03	0.11
x9	276	3.00	1.01	0.08	0.09	0.09	0.02	0.02	0.07	0.23	0.24	1.03	0.06	0.04	0.22
x10	276	2.99	1.03	0.06	0.08	0.21	0.09	0.08	0.09	0.07	0.09	0.06	1.05	0.17	0.23

```
x11 276 2.99 1.00 0.07 0.09 0.01 0.08 0.08 0.30 0.05 0.03 0.04 0.17 0.99 0.24
x12 276 2.99 1.01 0.07 0.11 0.07 0.09 0.11 0.12 0.07 0.11 0.22 0.23 0.24 1.01
era 276 2004.93 5.01 -0.04 0.11 0.22 -0.26 0.45 0.13 0.26 0.36 0.11 -0.29 -0.09 0.13
```

```
era x1.1 x2.1 x3.1 x4.1 x5.1 x6.1 x7.1 x8.1 x9.1 x10.1 x11.1 x12.1 era.1
x1 -0.04 1.00 0.28 0.19 0.27 -0.01 0.03 0.11 0.07 0.08 0.06 0.07 0.07 -0.01
x2 0.11 0.28 1.00 0.22 -0.01 0.02 0.04 0.28 0.06 0.09 0.08 0.09 0.10 0.02
x3 0.22 0.19 0.22 1.00 -0.01 0.03 0.01 0.12 0.08 0.09 0.21 0.01 0.07 0.04
x4 -0.26 0.27 -0.01 -0.01 1.00 0.20 0.23 0.03 0.04 0.02 0.09 0.08 0.08 -0.05
x5 0.45 -0.01 0.02 0.03 0.20 1.00 0.21 0.05 0.23 0.02 0.08 0.08 0.10 0.09
x6 0.13 0.03 0.04 0.01 0.23 0.21 1.00 0.09 0.04 0.07 0.09 0.29 0.11 0.02
x7 0.26 0.11 0.28 0.12 0.03 0.05 0.09 1.00 0.21 0.23 0.06 0.05 0.07 0.05
x8 0.36 0.07 0.06 0.08 0.04 0.23 0.04 0.21 1.00 0.24 0.09 0.03 0.11 0.07
x9 0.11 0.08 0.09 0.09 0.02 0.02 0.07 0.23 0.24 1.00 0.06 0.04 0.22 0.02
x10 -0.29 0.06 0.08 0.21 0.09 0.08 0.09 0.06 0.09 0.06 1.00 0.17 0.22 -0.06
x11 -0.09 0.07 0.09 0.01 0.08 0.08 0.29 0.05 0.03 0.04 0.17 1.00 0.24 -0.02
x12 0.13 0.07 0.10 0.07 0.08 0.10 0.11 0.07 0.11 0.22 0.22 0.24 1.00 0.03
era 25.09 -0.01 0.02 0.04 -0.05 0.09 0.02 0.05 0.07 0.02 -0.06 -0.02 0.03 1.00
```

```
>
>
```

```
> #多母集団パス解析
> #semパッケージの読み込み
> library(sem)
> opt <- options(fit.indices = c("GFI", "AGFI", "RMSEA", "NFI", "NNFI", "CFI", "RNI", "IFI",
+ "SRMR", "AIC", "AICc", "BIC", "CAIC"))
```

```
> seq.1 <- specifyEquations()
```

```
1: x1 = 1*f1
2: x2 = b2.1*f1
3: x3 = b3.1*f1
4: x4 = 1*f2
5: x5 = b5.1*f2
6: x6 = b6.1*f2
7: x7 = 1*f3
8: x8 = b8.1*f3
9: x9 = b9.1*f3
10: x10 = 1*f4
11: x11 = b11.1*f4
12: x12 = b12.1*f4
13: f3 = a31.1*f1
14: f4 = a41.1*f1 + a42.1*f2
15: V(x1) = ev1.1
16: V(x2) = ev2.1
17: V(x3) = ev3.1
18: V(x4) = ev4.1
19: V(x5) = ev5.1
20: V(x6) = ev6.1
21: V(x7) = ev7.1
22: V(x8) = ev8.1
23: V(x9) = ev9.1
24: V(x10) = ev10.1
25: V(x11) = ev11.1
26: V(x12) = ev12.1
27: V(f3) = evf3.1
28: V(f4) = evf4.1
29: V(f1) = vf1.1
30: V(f2) = vf2.1
31: C(f1, f2) = cf12.1
32:
Read 31 items
```

	A	B	C	D	E	F	G	H	I	J	K	L	M	N
1	id	x1	x2	x3	x4	x5	x6	x7	x8	x9	x10	x11	x12	era
2	1	4	3	2	3	4	4	2	4	2	4	4	3	2010
3	2	2	3	2	5	3	4	5	3	3	4	4	4	2010
4	3	1	4	2	2	3	3	4	3	2	3	4	4	2000
5	4	1	5	2	2	3	2	3	2	2	5	4	3	2000
6	5	2	3	1	4	2	1	3	3	2	1	1	1	2000
7	6	2	4	3	3	5	4	4	3	4	2	4	1	2010
8	7	4	4	4	4	4	3	2	1	2	2	2	1	2000
9	8	3	4	3	3	3	4	3	4	3	4	4	3	2000
10	9	3	4	3	2	3	5	4	4	3	2	5	4	2010
11	10	3	5	2	2	4	4	4	4	5	4	1	3	2010
12	11	4	2	3	4	4	5	3	4	2	5	3	2	2000
13	12	4	4	2	3	4	2	2	4	4	3	3	4	2010
14	13	3	4	3	2	3	3	4	4	4	4	2	5	2010
15	14	4	3	5	4	2	3	3	2	2	4	3	3	2010
16	15	1	4	3	2	3	2	4	1	3	2	3	2	2010
17	16	4	4	3	2	3	4	4	4	3	3	3	3	2000
18	17	3	2	2	4	1	4	4	4	4	3	2	2	2000
19	18	4	4	4	4	4	3	3	4	2	2	1	3	2010
20	19	3	3	3	3	3	4	2	3	4	4	3	3	2010
21	20	4	2	5	2	2	2	4	2	4	4	3	5	2010

```
> seq.2 <- specifyEquations()
```

```
1: x1 = 1*f1
2: x2 = b2.2*f1
3: x3 = b3.2*f1
4: x4 = 1*f2
5: x5 = b5.2*f2
6: x6 = b6.2*f2
7: x7 = 1*f3
8: x8 = b8.2*f3
```

```

9: x9 = b9.2*f3
10: x10 = l*f4
11: x11 = b11.2*f4
12: x12 = b12.2*f4
13: f3 = a31.2*f1
14: f4 = a41.2*f1 + a42.2*f2
15: V(x1) = ev1.2
16: V(x2) = ev2.2
17: V(x3) = ev3.2
18: V(x4) = ev4.2
19: V(x5) = ev5.2
20: V(x6) = ev6.2
21: V(x7) = ev7.2
22: V(x8) = ev8.2
23: V(x9) = ev9.2
24: V(x10) = ev10.2
25: V(x11) = ev11.2
26: V(x12) = ev12.2
27: V(f3) = evf3.2
28: V(f4) = evf4.2
29: V(f1) = vf1.2
30: V(f2) = vf2.2
31: C(f1, f2) = cf12.2
32:
Read 31 items
>

```

```

> mg.seq <- multigroupModel(seq.1, seq.2, groups=c("2000", "2010"))
> d1$era <- factor(d1$era)
> sem.mg.seq <- sem(mg.seq, data=d1, group="era")
> summary(sem.mg.seq)

```

```

Model Chisquare = 139.27 Df = 100 Pr(>Chisq) = 0.0057848
Chisquare (null model) = 351.43 Df = 132
Goodness-of-fit index = 0.92389
Adjusted goodness-of-fit index = 0.8904
RMSEA index = 0.053541 90% CI: (0.029782, 0.07371)
Bentler-Bonnett NFI = 0.6037
Tucker-Lewis NNFI = 0.76375
Bentler CFI = 0.82102
SRMR = 0.07014
AIC = 251.27
AICc = 168.42
BIC = 454.02

```

```
Iterations: initial fits, 55 57 final fit, 0
```

```
era: 2000
```

```

Model Chisquare = 55.626 Df = 50 Pr(>Chisq) = 0.27128
Chisquare (null model) = 201.42 Df = 66
Goodness-of-fit index = 0.93966
Adjusted goodness-of-fit index = 0.90587
RMSEA index = 0.028452 90% CI: (NA, 0.06355)
Bentler-Bonnett NFI = 0.72383
Tucker-Lewis NNFI = 0.94516
Bentler CFI = 0.95845
SRMR = 0.059422
AIC = 111.63
AICc = 70.257
BIC = 193.99
CAIC = -241.46

```

```
Normalized Residuals
```

```

Min. 1st Qu. Median Mean 3rd Qu. Max.
-1.58000 -0.34000 0.00361 0.04690 0.52400 2.49000

```

```
R-square for Endogenous Variables
```

```

x1 x2 x3 x4 x5 x6 f3 x7 x8 x9 f4 x10 x11
0.2201 0.3168 0.2498 0.1351 0.2667 0.4718 0.3518 0.3523 0.2093 0.2718 0.4097 0.1134 0.3386
x12
0.3100

```

Parameter Estimates

	Estimate	Std Error	z value	Pr(> z)		
b2.1	1.387211	0.431740	3.2131	1.3132e-03	x2 <---	f1
b3.1	1.108615	0.354829	3.1244	1.7819e-03	x3 <---	f1
b5.1	1.419028	0.498984	2.8438	4.4574e-03	x5 <---	f2
b6.1	1.901602	0.689501	2.7579	5.8167e-03	x6 <---	f2
b8.1	0.840790	0.270683	3.1062	1.8952e-03	x8 <---	f3
b9.1	0.976374	0.302761	3.2249	1.2602e-03	x9 <---	f3
b11.1	1.641288	0.628409	2.6118	9.0063e-03	x11 <---	f4
b12.1	1.536786	0.589578	2.6066	9.1450e-03	x12 <---	f4
a31.1	0.715202	0.264344	2.7056	6.8187e-03	f3 <---	f1
a41.1	0.275824	0.165362	1.6680	9.5315e-02	f4 <---	f1
a42.1	0.433487	0.236005	1.8368	6.6244e-02	f4 <---	f2
ev1.1	0.734696	0.110091	6.6735	2.4974e-11	x1 <-->	x1
ev2.1	0.860047	0.153842	5.5905	2.2647e-08	x2 <-->	x2
ev3.1	0.764970	0.120229	6.3626	1.9838e-10	x3 <-->	x3
ev4.1	0.901075	0.120738	7.4631	8.4519e-14	x4 <-->	x4
ev5.1	0.779726	0.129340	6.0285	1.6547e-09	x5 <-->	x5
ev6.1	0.569972	0.167696	3.3988	6.7674e-04	x6 <-->	x6
ev7.1	0.554234	0.111914	4.9523	7.3333e-07	x7 <-->	x7
ev8.1	0.805212	0.120075	6.7059	2.0014e-11	x8 <-->	x8
ev9.1	0.769747	0.128642	5.9837	2.1819e-09	x9 <-->	x9
ev10.1	1.024657	0.135168	7.5806	3.4385e-14	x10 <-->	x10
ev11.1	0.689958	0.134091	5.1454	2.6690e-07	x11 <-->	x11
ev12.1	0.689099	0.124773	5.5228	3.3363e-08	x12 <-->	x12
evf3.1	0.195389	0.094658	2.0642	3.9003e-02	f3 <-->	f3
evf4.1	0.077390	0.054608	1.4172	1.5643e-01	f4 <-->	f4
vf1.1	0.207288	0.096718	2.1432	3.2095e-02	f1 <-->	f1
vf2.1	0.140805	0.082652	1.7036	8.8458e-02	f2 <-->	f2
cf12.1	0.048003	0.030753	1.5609	1.1854e-01	f2 <-->	f1

era: 2010

Model Chisquare = 83.647 Df = 50 Pr(>Chisq) = 0.0020042
 Chisquare (null model) = 150.02 Df = 66
 Goodness-of-fit index = 0.90766
 Adjusted goodness-of-fit index = 0.85595
 RMSEA index = 0.070603 90% CI: (0.042712, 0.096505)
 Bentler-Bonnett NFI = 0.44242
 Tucker-Lewis NNFI = 0.47137
 Bentler CFI = 0.59952
 SRMR = 0.081173
 AIC = 139.65
 AICc = 98.825
 BIC = 221.2
 CAIC = -211.99

Normalized Residuals

Min.	1st Qu.	Median	Mean	3rd Qu.	Max.
-2.4200	-0.6220	-0.0153	0.0057	0.3470	2.9700

R-square for Endogenous Variables

x1	x2	x3	x4	x5	x6	f3	x7	x8	x9	f4	x10	x11
0.1830	0.2500	0.1217	0.8147	0.0245	0.0576	0.3041	0.3805	0.0861	0.0885	0.2979	0.5350	0.0562
x12												
0.1042												

Parameter Estimates

	Estimate	Std Error	z value	Pr(> z)		
b2.2	1.104586	0.444171	2.48685	1.2888e-02	x2 <---	f1
b3.2	0.805446	0.366318	2.19876	2.7895e-02	x3 <---	f1
b5.2	0.163661	0.210236	0.77847	4.3629e-01	x5 <---	f2
b6.2	0.271632	0.328033	0.82806	4.0763e-01	x6 <---	f2
b8.2	0.457730	0.275959	1.65869	9.7179e-02	x8 <---	f3
b9.2	0.457250	0.274247	1.66729	9.5456e-02	x9 <---	f3
b11.2	0.323457	0.201419	1.60589	1.0830e-01	x11 <---	f4
b12.2	0.460724	0.251944	1.82867	6.7448e-02	x12 <---	f4
a31.2	0.844737	0.388132	2.17642	2.9524e-02	f3 <---	f1
a41.2	0.837488	0.388973	2.15307	3.1313e-02	f4 <---	f1
a42.2	0.105691	0.163408	0.64679	5.1777e-01	f4 <---	f2
ev1.2	0.804811	0.125131	6.43173	1.2616e-10	x1 <-->	x1

```

ev2.2  0.660110 0.118741 5.55924 2.7095e-08 x2 <--> x2
ev3.2  0.844021 0.118422 7.12723 1.0241e-12 x3 <--> x3
ev4.2  0.195468 0.999339 0.19560 8.4493e-01 x4 <--> x4
ev5.2  0.917287 0.115124 7.96778 1.6155e-15 x5 <--> x5
ev6.2  1.037627 0.146550 7.08038 1.4376e-12 x6 <--> x6
ev7.2  0.688786 0.251684 2.73671 6.2056e-03 x7 <--> x7
ev8.2  0.940539 0.129059 7.28768 3.1533e-13 x8 <--> x8
ev9.2  0.911040 0.125573 7.25508 4.0144e-13 x9 <--> x9
ev10.2 0.442294 0.263279 1.67994 9.2968e-02 x10 <--> x10
ev11.2 0.894038 0.114490 7.80888 5.7697e-15 x11 <--> x11
ev12.2 0.928807 0.127870 7.26367 3.7673e-13 x12 <--> x12
evf3.2 0.294469 0.235574 1.25001 2.1130e-01 f3 <--> f3
evf4.2 0.357242 0.261303 1.36716 1.7158e-01 f4 <--> f4
vf1.2  0.180320 0.104133 1.73163 8.3339e-02 f1 <--> f1
vf2.2  0.859379 1.006991 0.85341 3.9343e-01 f2 <--> f2
cf12.2 0.087748 0.060126 1.45939 1.4446e-01 f2 <--> f1

```

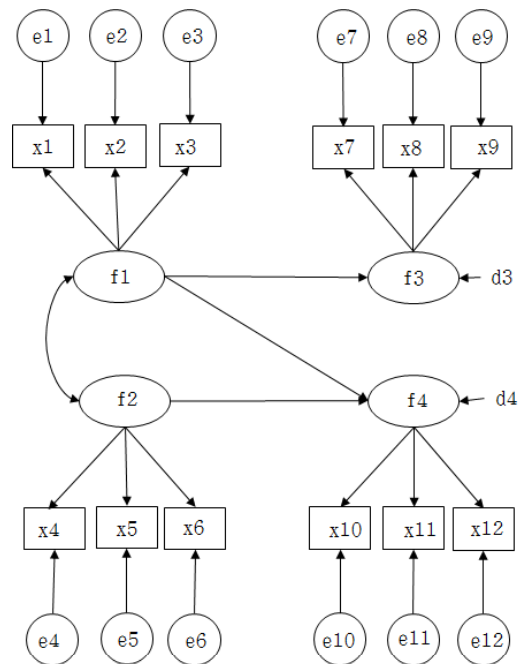
> stdCoef(sem.mg.seq)

Group: 2000

	Std.	Estimate	
1		0.4690996	x1 <--- f1
2	b2.1	0.5628938	x2 <--- f1
3	b3.1	0.4998324	x3 <--- f1
4		0.3676214	x4 <--- f2
5	b5.1	0.5163943	x5 <--- f2
6	b6.1	0.6868962	x6 <--- f2
7		0.5935218	x7 <--- f3
8	b8.1	0.4574426	x8 <--- f3
9	b9.1	0.5213693	x9 <--- f3
10		0.3367948	x10 <--- f4
11	b11.1	0.5818584	x11 <--- f4
12	b12.1	0.5567894	x12 <--- f4
13	a31.1	0.5931025	f3 <--- f1
14	a41.1	0.3468331	f4 <--- f1
15	a42.1	0.4492499	f4 <--- f2
16	ev1.1	0.7799455	x1 <--> x1
17	ev2.1	0.6831505	x2 <--> x2
18	ev3.1	0.7501675	x3 <--> x3
19	ev4.1	0.8648545	x4 <--> x4
20	ev5.1	0.7333369	x5 <--> x5
21	ev6.1	0.5281736	x6 <--> x6
22	ev7.1	0.6477319	x7 <--> x7
23	ev8.1	0.7907463	x8 <--> x8
24	ev9.1	0.7281741	x9 <--> x9
25	ev10.1	0.8865693	x10 <--> x10
26	ev11.1	0.6614408	x11 <--> x11
27	ev12.1	0.6899855	x12 <--> x12
28	evf3.1	0.6482294	f3 <--> f3
29	evf4.1	0.5903208	f4 <--> f4
30	vf1.1	1.0000000	f1 <--> f1
31	vf2.1	1.0000000	f2 <--> f2
32	cf12.1	0.2809765	f2 <--> f1

Group: 2010

	Std.	Estimate	
1		0.4278336	x1 <--- f1
2	b2.2	0.4999775	x2 <--- f1
3	b3.2	0.3488959	x3 <--- f1
4		0.9026048	x4 <--- f2
5	b5.2	0.1564604	x5 <--- f2
6	b6.2	0.2399789	x6 <--- f2
7		0.6168856	x7 <--- f3
8	b8.2	0.2934971	x8 <--- f3
9	b9.2	0.2975113	x9 <--- f3
10		0.7314283	x10 <--- f4
11	b11.2	0.2370683	x11 <--- f4
12	b12.2	0.3227633	x12 <--- f4
13	a31.2	0.5514428	f3 <--- f1
14	a41.2	0.4985463	f4 <--- f1



```

15 a42.2      0.1373522    f4 <--- f2
16 ev1.2      0.8169584     x1 <--> x1
17 ev2.2      0.7500225     x2 <--> x2
18 ev3.2      0.8782716     x3 <--> x3
19 ev4.2      0.1853045     x4 <--> x4
20 ev5.2      0.9755201     x5 <--> x5
21 ev6.2      0.9424101     x6 <--> x6
22 ev7.2      0.6194522     x7 <--> x7
23 ev8.2      0.9138594     x8 <--> x8
24 ev9.2      0.9114870     x9 <--> x9
25 ev10.2     0.4650126    x10 <--> x10
26 ev11.2     0.9437986    x11 <--> x11
27 ev12.2     0.8958239    x12 <--> x12
28 evf3.2     0.6959108     f3 <--> f3
29 evf4.2     0.7020583     f4 <--> f4
30 vf1.2      1.0000000     f1 <--> f1
31 vf2.2      1.0000000     f2 <--> f2
32 cf12.2     0.2229062     f2 <--> f1
>

```

多母集団分析 — lavaanパッケージを使う方法

パッケージの読み込み

```
library(lavaan)
```

あらかじめlavaanパッケージをインストールしておく必要がある。

lavaanパッケージをインストールするために、mnormtパッケージをインストールしておく。

モデルの設定

```
モデル名 <- ' '
```

```
# 測定変数の指定 (= ~ を使う)
```

```
潜在変数名1 ~ 観測変数1 + 観測変数2 + ...
```

```
潜在変数名2 ~ 観測変数1 + 観測変数2 + ...
```

```
# 回帰式 (~ を使う)
```

```
従属変数名1 ~ 説明変数1 + 説明変数2 + ...
```

```
従属変数名2 ~ 説明変数1 + 説明変数2 + ...
```

```
# 分散, 共分散 (~~ を使う)
```

```
変数名i ~~ 変数名i # 分散
```

```
, 変数名j ~~ 変数名k # 共分散
```

モデル部分はシングルカンマ (') でくくる。

すべての変数の分散を推定する場合は分散の指定は省略できる。

パラメタ値の推定

```
lavaanオブジェクト名 <- lavaan(モデル名, data=データ名, model.type="sem", group="群分け変数名",
fixed.x=FALSE, meanstructure=FALSE, auto.var=TRUE, auto.fix.first=TRUE)
```

group="群分け変数名" : 母集団を分ける変数の名前を指定する

fixed.x=FALSE : 外生変数となる観測変数の分散, 共分散, 平均を標本平均で固定する

meanstructure=False : 平均構造の検討をしない

auto.var=TRUE : すべての変数の分散または残差分散を推定する

auto.fix.first=TRUE : 非標準化解において測定変数のうち最初の1つのパス係数を1に固定する

```
summary(lavaanオブジェクト名, fit.measures=TRUE, standardized=TRUE)
```

モデル部分のスキプトの例

```
model.1 <- ' '
```

```
# latent variable definitions
```

```
f1 =~ x1 + x2 + x3
```

```
f2 =~ x4 + x5 + x6
```

```
f3 =~ x7 + x8 + x9
```

```
f4 =~ x10 + x11 + x12
```

```
# regression
```

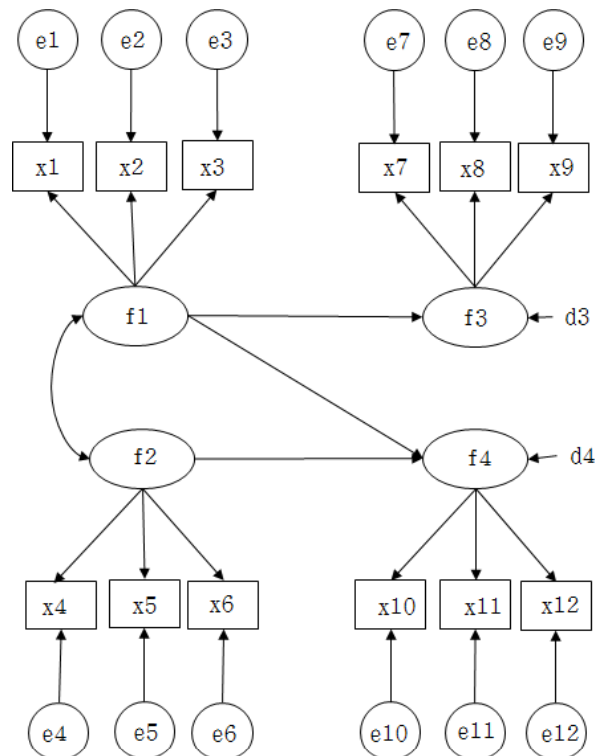
```
f3 =~ f1
```

```
f4 =~ f1 + f2
```

```
# variances and covariances
```

```
f1 ~~ f2
```

すべての分散を推定するので, 分散の式は省略している



```

> setwd("d:¥¥¥")
> d1 <- read.table("多母集団分析_データ.csv", header=TRUE, sep=",")
>
> # id 列の削除
> d1 <- d1[,c(-1)]
>
> #データフレームの行数（標本の大きさ）、平均、標準偏差、共分散、相関
> n.d1 <- nrow(d1)
> mean.d1 <- apply(d1, 2, mean)
> sd.d1 <- apply(d1, 2, sd)
> cov.d1 <- cov(d1)
> cor.d1 <- cor(d1)
> round(data.frame(n.d1, mean.d1, sd.d1, cov.d1, cor.d1), 2)
>

```

```

> # 群ごとのデータフレーム作成と記述統計量

```

```

> ds1 <- d1[d1$era==2000,]
> n.ds1 <- nrow(ds1)
> mean.ds1 <- apply(ds1, 2, mean)
> sd.ds1 <- apply(ds1, 2, sd)
> cov.ds1 <- cov(ds1)
> cor.ds1 <- cor(ds1)
>
> ds2 <- d1[d1$era==2010,]
> n.ds2 <- nrow(ds2)
> mean.ds2 <- apply(ds2, 2, mean)
> sd.ds2 <- apply(ds2, 2, sd)
> cov.ds2 <- cov(ds2)
> cor.ds2 <- cor(ds2)

```

	A	B	C	D	E	F	G	H	I	J	K	L	M	N
1	id	x1	x2	x3	x4	x5	x6	x7	x8	x9	x10	x11	x12	era
2	1	4	3	2	3	4	4	2	4	2	4	4	3	2010
3	2	2	3	2	5	3	4	5	3	3	4	4	4	2010
4	3	1	4	2	2	3	3	4	3	2	3	4	4	2000
5	4	1	5	2	2	3	2	3	2	2	5	4	3	2000
6	5	2	3	1	4	2	1	3	3	2	1	1	1	2000
7	6	2	4	3	3	5	4	4	3	4	2	4	1	2010
8	7	4	4	4	4	4	3	2	1	2	2	2	1	2000
9	8	3	4	3	3	3	4	3	4	3	4	4	3	2000
10	9	3	4	3	2	3	5	4	4	3	2	5	4	2010
11	10	3	5	2	2	4	4	4	4	5	4	1	3	2010
12	11	4	2	3	4	4	5	3	4	2	5	3	2	2000
13	12	4	4	2	3	4	2	2	4	4	3	3	4	2010
14	13	3	4	3	2	3	3	4	4	4	4	2	5	2010
15	14	4	3	5	4	2	3	3	2	2	4	3	3	2010
16	15	1	4	3	2	3	2	4	1	3	2	3	2	2010
17	16	4	4	3	2	3	4	4	4	3	3	3	3	2000
18	17	3	2	2	4	1	4	4	4	4	3	2	2	2000
19	18	4	4	4	4	4	3	3	4	2	2	1	3	2010
20	19	3	3	3	3	3	4	2	3	4	4	3	3	2010
21	20	4	2	5	2	2	2	4	2	4	4	3	5	2010

```

> #共分散構造分析の実行
> #lavaanパッケージの読み込み
> library(lavaan)

```

```

> # モデルの設定
> model.1 <- '
+ # latent variable definitions
+ f1 =~ x1 + x2 + x3
+ f2 =~ x4 + x5 + x6
+ f3 =~ x7 + x8 + x9
+ f4 =~ x10 + x11 + x12
+
+ # regression
+ f3 =~ f1
+ f4 =~ f1 + f2
+
+ # variances and covariances
+ ,f1 =~ f2
+ '
>

```

```

> # lavaan関数を使って計算
> fit.model.1 <- lavaan(model.1, data=d1, model.type="sem", group="era",
+ fixed.x=F, meanstructure=F,
+ auto.var=TRUE, auto.fix.first=TRUE)

```

```

> summary(fit.model.1, fit.measures=TRUE, standardized=TRUE)

```

lavaan (0.5-10) converged normally after 142 iterations

```

Number of observations per group
2010                                136
2000                                140

```

```

Estimator                        ML
Minimum Function Chi-square      140.293
Degrees of freedom                100
P-value                          0.005

```

Chi-square for each group:

2010	84.267
2000	56.026

Chi-square test baseline model:

Minimum Function Chi-square	353.995
Degrees of freedom	132
P-value	0.000

Full model versus baseline model:

Comparative Fit Index (CFI)	0.818
Tucker-Lewis Index (TLI)	0.760

Loglikelihood and Information Criteria:

Loglikelihood user model (H0)	-4611.761
Loglikelihood unrestricted model (H1)	-4541.615
Number of free parameters	56
Akaike (AIC)	9335.523
Bayesian (BIC)	9538.265
Sample-size adjusted Bayesian (BIC)	9360.699

Root Mean Square Error of Approximation:

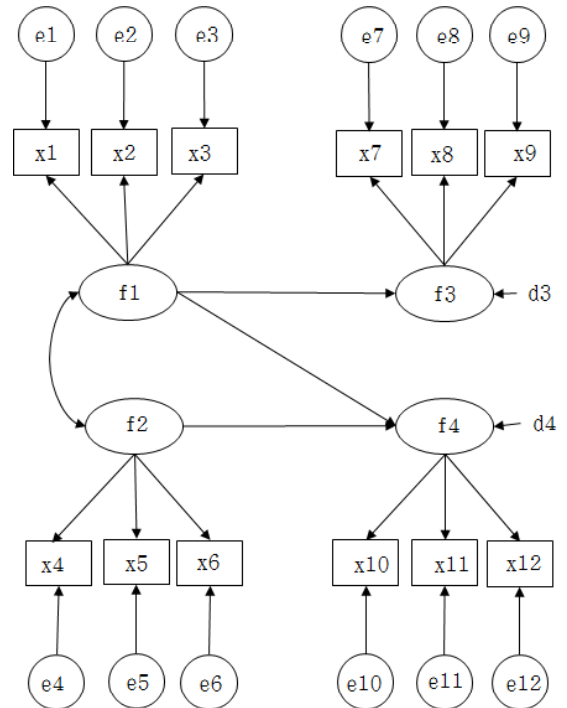
RMSEA	0.054
90 Percent Confidence Interval	0.031 0.074
P-value RMSEA ≤ 0.05	0.362

Standardized Root Mean Square Residual:

SRMR	0.070
------	-------

Parameter estimates:

Information	Expected
Standard Errors	Standard



Group 1 [2010]:

	Estimate	Std. err	Z-value	P(> z)	Std. lv	Std. all
Latent variables:						
f1 =						
x1	1.000				0.423	0.428
x2	1.105	0.443	2.496	0.013	0.467	0.500
x3	0.805	0.365	2.207	0.027	0.341	0.349
f2 =						
x4	1.000				0.924	0.903
x5	0.164	0.209	0.781	0.435	0.151	0.156
x6	0.272	0.327	0.831	0.406	0.251	0.240
f3 =						
x7	1.000				0.648	0.617
x8	0.458	0.275	1.665	0.096	0.297	0.293
x9	0.457	0.273	1.673	0.094	0.296	0.298
f4 =						
x10	1.000				0.711	0.731
x11	0.323	0.201	1.612	0.107	0.230	0.237
x12	0.461	0.251	1.835	0.066	0.327	0.323
Regressions:						
f3						
f1	0.845	0.387	2.184	0.029	0.551	0.551
f4						
f1	0.837	0.388	2.161	0.031	0.499	0.499
f2	0.106	0.163	0.649	0.516	0.137	0.137
Covariances:						
f1						
f2	0.087	0.059	1.465	0.143	0.223	0.223

Variances:

x1	0.799	0.124	0.799	0.817
x2	0.655	0.117	0.655	0.750
x3	0.838	0.117	0.838	0.878
x4	0.194	0.988	0.194	0.185
x5	0.911	0.114	0.911	0.976
x6	1.030	0.145	1.030	0.942
x7	0.684	0.249	0.684	0.619
x8	0.934	0.128	0.934	0.914
x9	0.904	0.124	0.904	0.911
x10	0.439	0.260	0.439	0.465
x11	0.887	0.113	0.887	0.944
x12	0.922	0.126	0.922	0.896
f1	0.179	0.103	1.000	1.000
f2	0.853	0.996	1.000	1.000
f3	0.292	0.233	0.696	0.696
f4	0.355	0.258	0.702	0.702

Group 2 [2000]:

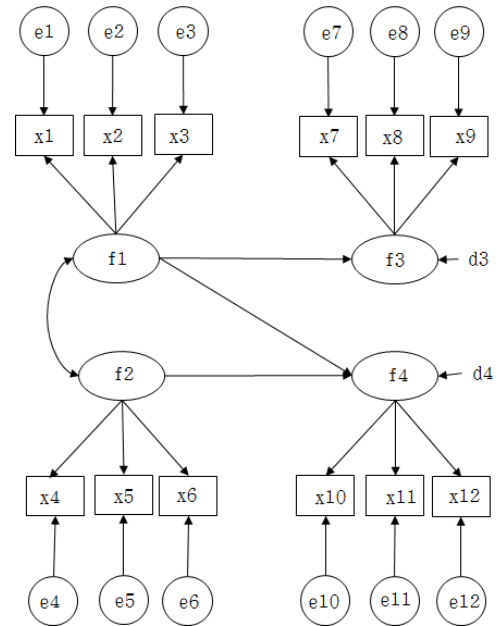
	Estimate	Std. err	Z-value	P(> z)	Std. lv	Std. all
Latent variables:						
f1 = ~						
x1	1.000				0.454	0.469
x2	1.387	0.430	3.225	0.001	0.629	0.563
x3	1.109	0.354	3.136	0.002	0.503	0.500
f2 = ~						
x4	1.000				0.374	0.368
x5	1.419	0.497	2.854	0.004	0.531	0.516
x6	1.902	0.687	2.768	0.006	0.711	0.687
f3 = ~						
x7	1.000				0.547	0.594
x8	0.841	0.270	3.117	0.002	0.460	0.457
x9	0.976	0.302	3.237	0.001	0.534	0.521
f4 = ~						
x10	1.000				0.361	0.337
x11	1.641	0.626	2.621	0.009	0.592	0.582
x12	1.537	0.587	2.616	0.009	0.554	0.557
Regressions:						
f3 ~						
f1	0.715	0.263	2.715	0.007	0.593	0.593
f4 ~						
f1	0.276	0.165	1.674	0.094	0.347	0.347
f2	0.434	0.235	1.843	0.065	0.449	0.449
Covariances:						
f1 ~						
f2	0.048	0.030	1.567	0.117	0.281	0.281
Variances:						
x1	0.729	0.109			0.729	0.780
x2	0.854	0.152			0.854	0.683
x3	0.760	0.119			0.760	0.750
x4	0.895	0.119			0.895	0.865
x5	0.774	0.128			0.774	0.733
x6	0.566	0.166			0.566	0.528
x7	0.550	0.111			0.550	0.648
x8	0.799	0.119			0.799	0.791
x9	0.764	0.127			0.764	0.728
x10	1.017	0.134			1.017	0.887
x11	0.685	0.133			0.685	0.661
x12	0.684	0.123			0.684	0.690
f1	0.206	0.096			1.000	1.000
f2	0.140	0.082			1.000	1.000
f3	0.194	0.094			0.648	0.648
f4	0.077	0.054			0.590	0.590

等値制約 — semパッケージを使う方法

```
seq.1 <- specifyEquations()
x1 = 1*f1
x2 = b2.1*f1
x3 = b3.1*f1
x4 = 1*f2
x5 = b5.1*f2
x6 = b6.1*f2
x7 = 1*f3
x8 = b8.1*f3
x9 = b9.1*f3
x10 = 1*f4
x11 = b11.1*f4
x12 = b12*f4
f3 = a31.1*f1
f4 = a41.1*f1 + a42.1*f2
V(x1) = ev1.1
V(x2) = ev2.1
V(x3) = ev3.1
V(x4) = ev4.1
V(x5) = ev5.1
V(x6) = ev6.1
V(x7) = ev7.1
V(x8) = ev8.1
V(x9) = ev9.1
V(x10) = ev10.1
V(x11) = ev11.1
V(x12) = ev12.1
V(f3) = evf3.1
V(f4) = evf4.1
V(f1) = vf1
V(f2) = vf2.1
C(f1, f2) = cf12.1
```

#b12の値は2つの母集団で共通

#f1の分散の値は2つの母集団で共通



```
seq.2 <- specifyEquations()
x1 = 1*f1
x2 = b2.2*f1
x3 = b3.2*f1
x4 = 1*f2
x5 = b5.2*f2
x6 = b6.2*f2
x7 = 1*f3
x8 = b8.2*f3
x9 = b9.2*f3
x10 = 1*f4
x11 = b11.2*f4
x12 = b12*f4
f3 = a31.2*f1
f4 = a41.2*f1 + a42.2*f2
V(x1) = ev1.2
V(x2) = ev2.2
V(x3) = ev3.2
V(x4) = ev4.2
V(x5) = ev5.2
V(x6) = ev6.2
V(x7) = ev7.2
V(x8) = ev8.2
V(x9) = ev9.2
V(x10) = ev10.2
V(x11) = ev11.2
V(x12) = ev12.2
V(f3) = evf3.2
V(f4) = evf4.2
V(f1) = vf1
V(f2) = vf2.2
C(f1, f2) = cf12.2
```

#b12の値は2つの母集団で共通

#f1の分散の値は2つの母集団で共通

sem パッケージでは、多母集団分析において、母集団間の等値制約はおけるが同一母集団内で等値制約はおけない（多母集団分析でなければ、単一母集団内で等値制約はおける）

```

> setwd("d:¥¥")
> d1 <- read.table("多母集団分析_データ.csv", header=TRUE, sep=",")
> head(d1)
  id x1 x2 x3 x4 x5 x6 x7 x8 x9 x10 x11 x12 era
1  1  4  3  2  3  4  4  2  4  2  4  4  3 2010
2  2  2  3  2  5  3  4  5  3  3  4  4  4 2010
3  3  1  4  2  2  3  3  4  3  2  3  4  4 2000
4  4  1  5  2  2  3  2  3  2  2  5  4  3 2000
5  5  2  3  1  4  2  1  3  3  2  1  1  1 2000
6  6  2  4  3  3  5  4  4  3  4  2  4  1 2010
>
> # id 列の削除
> d1 <- d1[,c(-1)]
>
>
> #データフレームの行数 (標本サイズ), 平均, 標準偏差, 共分散, 相関
> n.d1 <- nrow(d1)
> mean.d1 <- apply(d1, 2, mean)
> sd.d1 <- apply(d1, 2, sd)
> cov.d1 <- cov(d1)
> cor.d1 <- cor(d1)
> round(data.frame(n.d1, mean.d1, sd.d1, cov.d1, cor.d1), 2)

  n.d1 mean.d1 sd.d1    x1    x2    x3    x4    x5    x6    x7    x8    x9    x10    x11    x12
x1   276    3.01  0.98  0.96  0.28  0.19  0.27 -0.01  0.03  0.11  0.07  0.08  0.06  0.07  0.07
x2   276    3.01  1.03  0.28  1.07  0.22 -0.01  0.03  0.04  0.28  0.06  0.09  0.08  0.09  0.11
x3   276    3.00  0.99  0.19  0.22  0.99 -0.01  0.03  0.01  0.12  0.08  0.09  0.21  0.01  0.07
x4   276    2.99  1.02  0.27 -0.01 -0.01  1.05  0.20  0.24  0.03  0.04  0.02  0.09  0.08  0.09
x5   276    3.00  1.00 -0.01  0.03  0.03  0.20  1.01  0.22  0.05  0.23  0.02  0.08  0.08  0.11
x6   276    3.03  1.04  0.03  0.04  0.01  0.24  0.22  1.09  0.09  0.05  0.07  0.09  0.30  0.12
x7   276    3.03  0.99  0.11  0.28  0.12  0.03  0.05  0.09  0.98  0.21  0.23  0.07  0.05  0.07
x8   276    3.01  1.01  0.07  0.06  0.08  0.04  0.23  0.05  0.21  1.03  0.24  0.09  0.03  0.11
x9   276    3.00  1.01  0.08  0.09  0.09  0.02  0.02  0.07  0.23  0.24  1.03  0.06  0.04  0.22
x10  276    2.99  1.03  0.06  0.08  0.21  0.09  0.08  0.09  0.07  0.09  0.06  1.05  0.17  0.23
x11  276    2.99  1.00  0.07  0.09  0.01  0.08  0.08  0.30  0.05  0.03  0.04  0.17  0.99  0.24
x12  276    2.99  1.01  0.07  0.11  0.07  0.09  0.11  0.12  0.07  0.11  0.22  0.23  0.24  1.01
era  276 2004.93  5.01 -0.04  0.11  0.22 -0.26  0.45  0.13  0.26  0.36  0.11 -0.29 -0.09  0.13

  era x1.1 x2.1 x3.1 x4.1 x5.1 x6.1 x7.1 x8.1 x9.1 x10.1 x11.1 x12.1 era.1
x1 -0.04 1.00 0.28 0.19 0.27 -0.01 0.03 0.11 0.07 0.08 0.06 0.07 0.07 -0.01
x2  0.11 0.28 1.00 0.22 -0.01 0.02 0.04 0.28 0.06 0.09 0.08 0.09 0.10 0.02
x3  0.22 0.19 0.22 1.00 -0.01 0.03 0.01 0.12 0.08 0.09 0.21 0.01 0.07 0.04
x4 -0.26 0.27 -0.01 -0.01 1.00 0.20 0.23 0.03 0.04 0.02 0.09 0.08 0.08 -0.05
x5  0.45 -0.01 0.02 0.03 0.20 1.00 0.21 0.05 0.23 0.02 0.08 0.08 0.10 0.09
x6  0.13 0.03 0.04 0.01 0.23 0.21 1.00 0.09 0.04 0.07 0.09 0.29 0.11 0.02
x7  0.26 0.11 0.28 0.12 0.03 0.05 0.09 1.00 0.21 0.23 0.06 0.05 0.07 0.05
x8  0.36 0.07 0.06 0.08 0.04 0.23 0.04 0.21 1.00 0.24 0.09 0.03 0.11 0.07
x9  0.11 0.08 0.09 0.09 0.02 0.02 0.07 0.23 0.24 1.00 0.06 0.04 0.22 0.02
x10 -0.29 0.06 0.08 0.21 0.09 0.08 0.09 0.06 0.09 0.06 1.00 0.17 0.22 -0.06
x11 -0.09 0.07 0.09 0.01 0.08 0.08 0.29 0.05 0.03 0.04 0.17 1.00 0.24 -0.02
x12 0.13 0.07 0.10 0.07 0.08 0.10 0.11 0.07 0.11 0.22 0.22 0.24 1.00 0.03
era 25.09 -0.01 0.02 0.04 -0.05 0.09 0.02 0.05 0.07 0.02 -0.06 -0.02 0.03 1.00
>
>
>
>
> #多母集団パス解析
> #semパッケージの読み込み
> library(sem)
> opt <- options(fit.indices = c("GFI", "AGFI", "RMSEA", "NFI", "NNFI", "CFI", "RNI", "IFI",
+ "SRMR", "AIC", "AICc", "BIC", "CAIC"))
>
> seq.1 <- specifyEquations()
1: x1 = 1*f1
2: x2 = b2.1*f1
3: x3 = b3.1*f1
4: x4 = 1*f2
5: x5 = b5.1*f2
6: x6 = b6.1*f2
7: x7 = 1*f3
8: x8 = b8.1*f3

```

```

9: x9 = b9.1*f3
10: x10 = 1*f4
11: x11 = b11.1*f4
12: x12 = b12*f4          # b12は共通
13: f3 = a31.1*f1
14: f4 = a41.1*f1 + a42.1*f2
15: V(x1) = ev1.1
16: V(x2) = ev2.1
17: V(x3) = ev3.1
18: V(x4) = ev4.1
19: V(x5) = ev5.1
20: V(x6) = ev6.1
21: V(x7) = ev7.1
22: V(x8) = ev8.1
23: V(x9) = ev9.1
24: V(x10) = ev10.1
25: V(x11) = ev11.1
26: V(x12) = ev12.1
27: V(f3) = evf3.1
28: V(f4) = evf4.1
29: V(f1) = vf1          # vf1は共通
30: V(f2) = vf2.1
31: C(f1, f2) = cf12.1
32:
Read 31 items

```

	A	B	C	D	E	F	G	H	I	J	K	L	M	N
1	id	x1	x2	x3	x4	x5	x6	x7	x8	x9	x10	x11	x12	era
2	1	4	3	2	3	4	4	2	4	2	4	4	3	2010
3	2	2	3	2	5	3	4	5	3	3	4	4	4	2010
4	3	1	4	2	2	3	3	4	3	2	3	4	4	2000
5	4	1	5	2	2	3	2	3	2	2	5	4	3	2000
6	5	2	3	1	4	2	1	3	3	2	1	1	1	2000
7	6	2	4	3	3	5	4	4	3	4	2	4	1	2010
8	7	4	4	4	4	4	3	2	1	2	2	2	1	2000
9	8	3	4	3	3	3	4	3	4	3	4	4	3	2000
10	9	3	4	3	2	3	5	4	4	3	2	5	4	2010
11	10	3	5	2	2	4	4	4	4	5	4	1	3	2010
12	11	4	2	3	4	4	5	3	4	2	5	3	2	2000
13	12	4	4	2	3	4	2	2	4	4	3	3	4	2010
14	13	3	4	3	2	3	3	4	4	4	4	2	5	2010
15	14	4	3	5	4	2	3	3	2	2	4	3	3	2010
16	15	1	4	3	2	3	2	4	1	3	2	3	2	2010
17	16	4	4	3	2	3	4	4	4	3	3	3	3	2000
18	17	3	2	2	4	1	4	4	4	4	3	2	2	2000
19	18	4	4	4	4	4	3	3	4	2	2	1	3	2010
20	19	3	3	3	3	3	4	2	3	4	4	3	3	2010
21	20	4	2	5	2	2	2	4	2	4	4	3	5	2010

```

> seq.2 <- specifyEquations()
1: x1 = 1*f1
2: x2 = b2.2*f1
3: x3 = b3.2*f1
4: x4 = 1*f2
5: x5 = b5.2*f2
6: x6 = b6.2*f2
7: x7 = 1*f3
8: x8 = b8.2*f3
9: x9 = b9.2*f3
10: x10 = 1*f4
11: x11 = b11.2*f4
12: x12 = b12*f4          # b12は共通
13: f3 = a31.2*f1
14: f4 = a41.2*f1 + a42.2*f2
15: V(x1) = ev1.2
16: V(x2) = ev2.2
17: V(x3) = ev3.2
18: V(x4) = ev4.2
19: V(x5) = ev5.2
20: V(x6) = ev6.2
21: V(x7) = ev7.2
22: V(x8) = ev8.2
23: V(x9) = ev9.2
24: V(x10) = ev10.2
25: V(x11) = ev11.2
26: V(x12) = ev12.2
27: V(f3) = evf3.2
28: V(f4) = evf4.2
29: V(f1) = vf1          # vf1は共通
30: V(f2) = vf2.2
31: C(f1, f2) = cf12.2
32:
Read 31 items
>
> mg.seq <- multigroupModel(seq.1, seq.2, groups=c("2000", "2010"))
> d1$era <- factor(d1$era)
> sem.mg.seq <- sem(mg.seq, data=d1, group="era")
> summary(sem.mg.seq)

```

```

Model Chisquare = 142.18 Df = 102 Pr(>Chisq) = 0.0052915
Chisquare (null model) = 351.43 Df = 132
Goodness-of-fit index = 0.92313
Adjusted goodness-of-fit index = 0.89147
RMSEA index = 0.053624 90% CI: (0.030182, 0.073589)

```

Bentler-Bonnett NFI = 0.59542
 Tucker-Lewis NNFI = 0.76302
 Bentler CFI = 0.81688
 SRMR = 0.071584
 AIC = 250.18
 AICc = 169.06
 BIC = 445.68

Iterations: initial fits, 55 57 final fit, 139

era: 2000

Model Chisquare = 56.452 Df = 50 Pr(>Chisq) = 0.24655
 Chisquare (null model) = 201.42 Df = 66
 Goodness-of-fit index = 0.93918
 Adjusted goodness-of-fit index = 0.90513
 RMSEA index = 0.030468 90% CI: (NA, 0.064737)
 Bentler-Bonnett NFI = 0.71973
 Tucker-Lewis NNFI = 0.93711
 Bentler CFI = 0.95236
 SRMR = 0.060757
 AIC = 112.45
 AICc = 71.082
 BIC = 194.82
 CAIC = -240.63

Normalized Residuals

Min.	1st Qu.	Median	Mean	3rd Qu.	Max.
-1.5600	-0.2540	0.0294	0.0521	0.4330	2.5100

R-square for Endogenous Variables

x1	x2	x3	x4	x5	x6	f3	x7	x8	x9	f4
0.2035	0.3143	0.2536	0.1291	0.2539	0.4952	0.3492	0.3526	0.2094	0.2704	0.4270
x10	x11	x12								
0.1615	0.3376	0.2497								

Parameter Estimates

	Estimate	Std Error	z value	Pr(> z)	
b2.1	1.442812	0.389261	3.7065	2.1011e-04	x2 <--- f1
b3.1	1.166645	0.328216	3.5545	3.7870e-04	x3 <--- f1
b5.1	1.416354	0.503401	2.8136	4.8995e-03	x5 <--- f2
b6.1	1.992752	0.734282	2.7139	6.6500e-03	x6 <--- f2
b8.1	0.840692	0.271591	3.0954	1.9653e-03	x8 <--- f3
b9.1	0.973363	0.303320	3.2090	1.3319e-03	x9 <--- f3
b11.1	1.351874	0.401026	3.3710	7.4886e-04	x11 <--- f4
b12	1.122060	0.206539	5.4327	5.5518e-08	x12 <--- f4
a31.1	0.744994	0.252750	2.9476	3.2030e-03	f3 <--- f1
a41.1	0.352728	0.184365	1.9132	5.5722e-02	f4 <--- f1
a42.1	0.554111	0.263727	2.1011	3.5634e-02	f4 <--- f2
ev1.1	0.742763	0.104995	7.0743	1.5026e-12	x1 <--> x1
ev2.1	0.861751	0.154395	5.5815	2.3851e-08	x2 <--> x2
ev3.1	0.760041	0.120954	6.2837	3.3057e-10	x3 <--> x3
ev4.1	0.907271	0.120576	7.5245	5.2925e-14	x4 <--> x4
ev5.1	0.793223	0.128037	6.1952	5.8193e-10	x5 <--> x5
ev6.1	0.544595	0.173776	3.1339	1.7251e-03	x6 <--> x6
ev7.1	0.553551	0.112276	4.9303	8.2118e-07	x7 <--> x7
ev8.1	0.804780	0.120195	6.6956	2.1477e-11	x8 <--> x8
ev9.1	0.770869	0.128744	5.9876	2.1293e-09	x9 <--> x9
ev10.1	1.000172	0.135317	7.3913	1.4536e-13	x10 <--> x10
ev11.1	0.690611	0.134516	5.1340	2.8361e-07	x11 <--> x11
ev12.1	0.728699	0.111881	6.5131	7.3593e-11	x12 <--> x12
evf3.1	0.196204	0.095259	2.0597	3.9427e-02	f3 <--> f3
evf4.1	0.110379	0.056731	1.9456	5.1697e-02	f4 <--> f4
vf1	0.189725	0.066529	2.8518	4.3479e-03	f1 <--> f1
vf2.1	0.134550	0.080425	1.6730	9.4327e-02	f2 <--> f2
cf12.1	0.044349	0.027926	1.5881	1.1226e-01	f2 <--> f1

era: 2010

Model Chisquare = 85.731 Df = 50 Pr(>Chisq) = 0.0012428
 Chisquare (null model) = 150.02 Df = 66

Goodness-of-fit index = 0.9066
Adjusted goodness-of-fit index = 0.85429
RMSEA index = 0.072756 90% CI: (0.045455, 0.098419)
Bentler-Bonnett NFI = 0.42853
Tucker-Lewis NNFI = 0.43864
Bentler CFI = 0.57472
SRMR = 0.082729
AIC = 141.73
AICc = 100.91
BIC = 223.29
CAIC = -209.9

Normalized Residuals

Min.	1st Qu.	Median	Mean	3rd Qu.	Max.
-2.43000	-0.62200	-0.07270	-0.00711	0.33300	2.84000

R-square for Endogenous Variables

x1	x2	x3	x4	x5	x6	f3	x7	x8	x9	f4
0.1906	0.3249	0.0736	0.4507	0.0404	0.1081	0.3338	0.4514	0.0669	0.0745	0.3129
x10	x11	x12								
0.2209	0.0980	0.2405								

Parameter Estimates

	Estimate	Std Error	z value	Pr(> z)		
b2.2	1.228827	0.446818	2.75017	5.9564e-03	x2 <--- f1	
b3.2	0.610847	0.108573	5.62615	1.8428e-08	x3 <--- f1	
b5.2	0.282537	0.498890	0.56633	5.7117e-01	x5 <--- f2	
b6.2	0.500396	0.319582	1.56578	1.1740e-01	x6 <--- f2	
b8.2	0.370129	0.235223	1.57352	1.1560e-01	x8 <--- f3	
b9.2	0.385046	0.361949	1.06381	2.8741e-01	x9 <--- f3	
b11.2	0.674508	0.229227	2.94254	3.2554e-03	x11 <--- f4	
b12	1.122060	0.206539	5.43267	5.5518e-08	x12 <--- f4	
a31.2	0.940164	0.318145	2.95514	3.1253e-03	f3 <--- f1	
a41.2	0.386055	0.397133	0.97210	3.3100e-01	f4 <--- f1	
a42.2	0.224145	0.249560	0.89816	3.6910e-01	f4 <--- f2	
ev1.2	0.805920	0.199240	4.04497	5.2330e-05	x1 <--> x1	
ev2.2	0.595305	0.127854	4.65613	3.2222e-06	x2 <--> x2	
ev3.2	0.890622	0.134377	6.62781	3.4071e-11	x3 <--> x3	
ev4.2	0.579531	0.117089	4.94948	7.4413e-07	x4 <--> x4	
ev5.2	0.902362	0.343941	2.62360	8.7006e-03	x5 <--> x5	
ev6.2	0.982018	0.115522	8.50068	1.8848e-17	x6 <--> x6	
ev7.2	0.610496	0.148560	4.10943	3.9665e-05	x7 <--> x7	
ev8.2	0.960500	0.276821	3.46975	5.2093e-04	x8 <--> x8	
ev9.2	0.925167	0.126300	7.32514	2.3865e-13	x9 <--> x9	
ev10.2	0.719727	0.121768	5.91065	3.4075e-09	x10 <--> x10	
ev11.2	0.854564	0.119675	7.14070	9.2857e-13	x11 <--> x11	
ev12.2	0.811409	0.118450	6.85022	7.3736e-12	x12 <--> x12	
evf3.2	0.334711	0.156121	2.14392	3.2039e-02	f3 <--> f3	
evf4.2	0.140189	0.269106	0.52094	6.0241e-01	f4 <--> f4	
vf1	0.189725	0.066529	2.85175	4.3479e-03	f1 <--> f1	
vf2.2	0.475457	0.353317	1.34570	1.7840e-01	f2 <--> f2	
cf12.2	0.067428	0.055899	1.20625	2.2772e-01	f2 <--> f1	

> stdCoef(sem.mg.seq)

Group: 2000

	Std. Estimate		
1	0.4510666	x1 <--- f1	
2	b2.1 0.5606032	x2 <--- f1	
3	b3.1 0.5035817	x3 <--- f1	
4	0.3593734	x4 <--- f2	
5	b5.1 0.5038711	x5 <--- f2	
6	b6.1 0.7037276	x6 <--- f2	
7	0.5938131	x7 <--- f3	
8	b8.1 0.4575488	x8 <--- f3	
9	b9.1 0.5199742	x9 <--- f3	
10	0.4018642	x10 <--- f4	
11	b11.1 0.5810708	x11 <--- f4	
12	b12 0.4997119	x12 <--- f4	
13	a31.1 0.5909734	f3 <--- f1	
14	a41.1 0.3500562	f4 <--- f1	

```

15 a42.1      0.4631008    f4 <--- f2
16 ev1.1      0.7965389    x1 <--> x1
17 ev2.1      0.6857240    x2 <--> x2
18 ev3.1      0.7464055    x3 <--> x3
19 ev4.1      0.8708508    x4 <--> x4
20 ev5.1      0.7461139    x5 <--> x5
21 ev6.1      0.5047674    x6 <--> x6
22 ev7.1      0.6473860    x7 <--> x7
23 ev8.1      0.7906491    x8 <--> x8
24 ev9.1      0.7296269    x9 <--> x9
25 ev10.1     0.8385052   x10 <--> x10
26 ev11.1     0.6623567   x11 <--> x11
27 ev12.1     0.7502880   x12 <--> x12
28 evf3.1     0.6507504    f3 <--> f3
29 evf4.1     0.5730035    f4 <--> f4
30   vf1      1.0000000    f1 <--> f1
31   vf2.1    1.0000000    f2 <--> f2
32 cf12.1     0.2775711    f2 <--> f1

```

Group: 2010

```

      Std. Estimate
1      0.4365258    x1 <--- f1
2   b2.2      0.5699934    x2 <--- f1
3   b3.2      0.2713559    x3 <--- f1
4      0.6713234    x4 <--- f2
5   b5.2      0.2009067    x5 <--- f2
6   b6.2      0.3288227    x6 <--- f2
7      0.6718929    x7 <--- f3
8   b8.2      0.2585861    x8 <--- f3
9   b9.2      0.2729710    x9 <--- f3
10     0.4699609   x10 <--- f4
11  b11.2     0.3130131   x11 <--- f4
12   b12      0.4903575   x12 <--- f4
13 a31.2     0.5777449    f3 <--- f1
14 a41.2     0.3722827    f4 <--- f1
15 a42.2     0.3421730    f4 <--- f2
16 ev1.2     0.8094452    x1 <--> x1
17 ev2.2     0.6751075    x2 <--> x2
18 ev3.2     0.9263660    x3 <--> x3
19 ev4.2     0.5493249    x4 <--> x4
20 ev5.2     0.9596365    x5 <--> x5
21 ev6.2     0.8918756    x6 <--> x6
22 ev7.2     0.5485599    x7 <--> x7
23 ev8.2     0.9331332    x8 <--> x8
24 ev9.2     0.9254868    x9 <--> x9
25 ev10.2    0.7791368   x10 <--> x10
26 ev11.2    0.9020228   x11 <--> x11
27 ev12.2    0.7595495   x12 <--> x12
28 evf3.2    0.6662108    f3 <--> f3
29 evf4.2    0.6871266    f4 <--> f4
30   vf1      1.0000000    f1 <--> f1
31   vf2.2    1.0000000    f2 <--> f2
32 cf12.2    0.2245030    f2 <--> f1
>
>

```

等値制約 — lavaanパッケージを使う方法

lavaanパッケージを使う方法の例（2群の多母集団分析）

```

model.1 <- '
# latent variable definitions
f1 =~ x1 + c("b12","")*x2 + c("b131","b132")*x3
f2 =~ x4 + equal(c("b12","")*x5 + c(NA,0)*x6
f3 =~ x7 + x8 + equal(c("b131","b132"))*x9
f4 =~ x10 + x11 + c("b412","b412")*x12

# regression
f3 ~ f1
f4 ~ f1 + f2

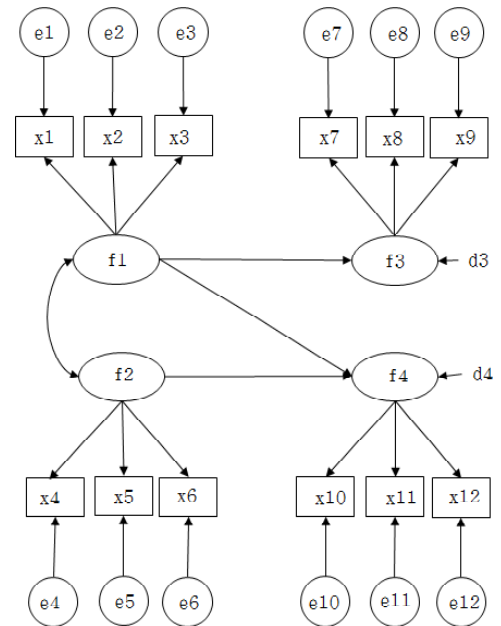
# variances and covariances
f1 ~~ c("vf1","vf1")*f1
f1 ~~ f2

```

b12 : 第1群において、 $x_2 \leftarrow f_1$ と $x_5 \leftarrow f_2$ は同じ値
 b131 : 第1群において、 $x_3 \leftarrow f_1$ と $x_9 \leftarrow f_3$ は同じ値
 b132 : 第2群において、 $x_3 \leftarrow f_1$ と $x_9 \leftarrow f_3$ は同じ値
 b412 : 第1群と第2群の $x_{12} \leftarrow f_4$ は同じ値
 c(NA, 0) : $x_6 \leftarrow f_2$ のパス係数を、第1群は推定するが、
 第2群は0に固定
 vf1 : 第1群と第2群の f_1 の分散は同じ値

equal()を使うように説明されているが、c()に同じ系数名を書きだけでも同じ値で推定されるようだ。

semパッケージを使うときも、係数に同じ名前をつければ、等値制約となる。



```

> setwd("d:¥¥")
> d1 <- read.table("多母集団分析_データ.csv", header=TRUE, sep=",")
>
> # id 列の削除
> d1 <- d1[,c(-1)]

> #データフレームの行数（標本の大きさ）、平均、標準偏差、共分散、相関
> n.d1 <- nrow(d1)
> mean.d1 <- apply(d1, 2, mean)
> sd.d1 <- apply(d1, 2, sd)
> cov.d1 <- cov(d1)
> cor.d1 <- cor(d1)
> round(data.frame(n.d1, mean.d1, sd.d1, cov.d1, cor.d1), 2)
>

> # 群ごとのデータフレーム作成と記述統計量
> ds1 <- d1[d1$era==2000,]
> n.ds1 <- nrow(ds1)
> mean.ds1 <- apply(ds1, 2, mean)
> sd.ds1 <- apply(ds1, 2, sd)
> cov.ds1 <- cov(ds1)
> cor.ds1 <- cor(ds1)
>

> ds2 <- d1[d1$era==2010,]
> n.ds2 <- nrow(ds2)
> mean.ds2 <- apply(ds2, 2, mean)
> sd.ds2 <- apply(ds2, 2, sd)
> cov.ds2 <- cov(ds2)
> cor.ds2 <- cor(ds2)
>

```


SRMR

0.087

Parameter estimates:

Information Standard Errors			Expected Standard				
Group 1 [2010]:							
		Estimate	Std. err	Z-value	P(> z)	Std. lv	Std. all
Latent variables:							
f1 =		1.000				0.478	0.479
x1		1.000				0.478	0.479
x2	(b12)	0.848	0.261	3.253	0.001	0.405	0.441
x3	(b131)	0.552	0.205	2.688	0.007	0.264	0.270
f2 =		1.000				0.423	0.420
x4		1.000				0.423	0.420
x5	(b12)	0.848	0.261	3.253	0.001	0.359	0.365
x6		0.935	0.514	1.818	0.069	0.396	0.378
f3 =		1.000				0.600	0.571
x7		1.000				0.600	0.571
x8		0.537	0.283	1.899	0.058	0.322	0.318
x9	(b131)	0.552	0.205	2.688	0.007	0.331	0.332
f4 =		1.000				0.420	0.440
x10		1.000				0.420	0.440
x11		0.734	0.353	2.081	0.037	0.308	0.318
x12	(b412)	1.241	0.355	3.499	0.000	0.522	0.506
Regressions:							
f3							
f1		0.720	0.286	2.515	0.012	0.574	0.574
f4							
f1		0.353	0.207	1.705	0.088	0.401	0.401
f2		0.386	0.282	1.372	0.170	0.389	0.389
Covariances:							
f1							
f2		0.044	0.046	0.954	0.340	0.217	0.217
Variances:							
f1	(vf1)	0.228	0.069			1.000	1.000
x1		0.766	0.120			0.766	0.770
x2		0.681	0.110			0.681	0.806
x3		0.882	0.115			0.882	0.927
x4		0.835	0.139			0.835	0.823
x5		0.838	0.124			0.838	0.867
x6		0.936	0.156			0.936	0.857
x7		0.742	0.188			0.742	0.674
x8		0.918	0.129			0.918	0.899
x9		0.884	0.121			0.884	0.890
x10		0.736	0.115			0.736	0.807
x11		0.845	0.118			0.845	0.899
x12		0.790	0.144			0.790	0.744
f2		0.179	0.105			1.000	1.000
f3		0.241	0.168			0.670	0.670
f4		0.109	0.072			0.620	0.620

Group 2 [2000]:

		Estimate	Std. err	Z-value	P(> z)	Std. lv	Std. all
Latent variables:							
f1 =		1.000					
x1		1.000				0.478	0.490
x2		1.343	0.335	4.009	0.000	0.642	0.573
x3	(b132)	1.014	0.208	4.878	0.000	0.485	0.482
f2 =		1.000					
x4		1.000				0.331	0.325
x5		2.514	2.450	1.026	0.305	0.831	0.809
x6		0.000				0.000	0.000
f3 =							

x7		1.000				0.538	0.585
x8		0.856	0.255	3.356	0.001	0.461	0.458
x9	(b132)	1.014	0.208	4.878	0.000	0.546	0.532
f4 = ~							
x10		1.000				0.448	0.411
x11		1.199	0.408	2.940	0.003	0.538	0.528
x12	(b412)	1.241	0.355	3.499	0.000	0.556	0.566
Regressions:							
f3							
f1		0.682	0.193	3.545	0.000	0.606	0.606
f4							
f1		0.378	0.182	2.083	0.037	0.404	0.404
f2		0.311	0.230	1.353	0.176	0.229	0.229
Covariances:							
f1							
f2		0.036	0.040	0.907	0.364	0.229	0.229
Variances:							
f1	(vf1)	0.228	0.069			1.000	1.000
x1		0.725	0.105			0.725	0.760
x2		0.842	0.152			0.842	0.672
x3		0.775	0.114			0.775	0.767
x4		0.925	0.152			0.925	0.894
x5		0.365	0.663			0.365	0.346
x6		1.071	0.128			1.071	1.000
x7		0.558	0.100			0.558	0.658
x8		0.800	0.118			0.800	0.790
x9		0.757	0.120			0.757	0.717
x10		0.987	0.138			0.987	0.831
x11		0.747	0.134			0.747	0.721
x12		0.656	0.126			0.656	0.679
f2		0.109	0.118			1.000	1.000
f3		0.183	0.080			0.633	0.633
f4		0.149	0.078			0.742	0.742
>							
>							