



Title	汎用統計パッケージSPSS ^ Xのサブプログラム
Author(s)	小野寺, 義孝; 竹村, 和久; 吉村, 英 他
Citation	大阪大学大型計算機センターニュース. 1988, 70, p. 71-159
Version Type	VoR
URL	https://hdl.handle.net/11094/65791
rights	
Note	

The University of Osaka Institutional Knowledge Archive : OUKA

<https://ir.library.osaka-u.ac.jp/>

The University of Osaka

汎用統計パッケージSPSS[×]のサブプログラム

1983年、SPSSは大幅に改定・強化され、SPSS[×]としてリリースされた。本センターにおいてもサービスされており、その機能の豊富さと使いやすさから、利用されている方も多いと思われる。既に日本語マニュアル「新版SPSS[×]—I基礎編」¹⁾が出版されており、基本的な機能については解説されている。しかし、多変量解析を中心として、多くのサブプログラムを紹介したものが、英文マニュアルのほかになく、利用者は不便を強いられているようである。またそのため、SPSS[×]の利用に二の足を踏まれる場合も多いようである。

そこで今回、SPSS[×]利用者の便を図るため、上記の日本語マニュアルに紹介されていないサブプログラムについて、その用法を中心に紹介することとした。基本的な機能についての詳細は上記マニュアルを参照されたい。また、その他の機能の詳細については、追って日本語マニュアルの続編も出版される予定であるが、それまでは英文マニュアル^{2), 3), 4)}を参照されたい。

(1) 紹介するサブプログラム

本解説でその用法を紹介するサブプログラムは次のとおりである。なおTABLESは、利用者が限られると思われるので今回は省略する。

ANOVA	多元配置の分散共分散分析
ONEWAY	一元配置の分散分析
MANOVA	分散共分散分析
HILOGLINEAR	階層的対数線形モデル解析
LOGLINEAR	対数線形モデル解析
PROBIT	プロビット
PLOT	プロット（各種2次元図の作成）
PEARSON CORR	ピアソンの相関係数
PARTIAL CORR	偏相関係数
REGRESSION	重回帰分析
DISCRIMINAT	判別分析
FACTOR	因子分析
PROXIMITIES	距離行列作成
ALSCAL	多次元尺度構成法
CLUSTER	クラスター分析
QUICK CLUSTER	高速クラス分析
NONPAR CORR	順位相関係数

NPAR TESTS	ノンパラメトリック検定
BOX-JENKINS	ボックス・ジェンキンス
RELIABILITY	信頼性解析
SURVIVAL	生存（表）分析

(2) サブコマンドなどの書式における表記法

- ・ `[]` はその中に書かれた事柄が省略可能であることを示す。
- ・ `{}` あるいは縦に並んだ `{}` は、その中からいずれかを選択すること示す。
- ・ 下線は標準値（省略時に自動的に採用される値）を示す。

[例] `LEVEL = { ORDINAL ([ UNTIE ] [ SIMILAR ])`
 `{ INTERVAL ( { 1 または d } )`
 `{ RATIO    ( { 1 または d } )`
 `{ NOMINAL`

この例では、`LEVEL=` のパラメータとして `ORDINAL`, `INTERVAL`, `RATIO`, `NOMINAL` の中から一つ選択し、`ORDINAL` にはさらに、その後に `()` でくくって `UNTIE` または `SIMILAR` が付加でき、これは省略可能であることを示している。また、`ORDINAL` が選ばれ `UNTIE` と `SIMILAR` が省略されたときは、`UNTIE` が標準値として選ばれることを示している。

1 ANOVA：多元配置分散共分散分析

ANOVA は多元配置分散共分散分析を行うとともに、MCA（多重分類分析）表を作成するプログラムである。ただし独立変数が 1 つの場合には `ONEWAY` を、また従属変数が複数（多重従属変数）の場合、繰返しのあるデザイン、入れ子型のデザイン、および要因・共変量間の交互作用を処理する場合は `MANOVA` を利用する。

1.1 一般書式

ANOVA 従属変数リスト BY 独立変数リスト（最小値，最大値）...
 [WITH 共変数リスト]/...

BY の左側に従属変数を、右側に独立変数を、WITH の右側に共変数を書く。スラッシュ `/` を付けることにより、さらに分析リストを続けることができる。

1.2 オプション

- 1 欠損値定義を無視し、データを除外しない。このオプションを指定しなければ、欠損値を含むケースは分析から除外され、使用されない。
- 2 変数ラベルおよび変数値ラベルの印刷を省略する。
- 3 すべての要因間交互作用を無視する。
- 4 3 要因間以上の交互作用を無視する。
- 5 4 要因間以上の交互作用を無視する。

- 6 5 要因間以上の交互作用を無視する。
- 7 共変量と要因を同時に投入する。
- 8 共変量を要因の後に投入する。
- 9 一括投入型回帰分析的処理法を採用する。
- 10 階層的なステップダウン処理法を採用する。
- 11 80文字の幅で印刷する。

1.3 追加統計

- 1 MCA（多重分類分析）表を作成する（オプション9の場合を除く）。
- 2 共変量にたいする標準化されていない偏回帰係数を表に印刷する。
- 3 平均、度数の表を印刷する（オプション9の場合を除く）。

1.4 コマンドの編成例

```
ANOVA  PRESTIGE BY REGION(1,9) SEX,RACE(1,2) WITH EDUC
STATISTICS 2
OPTIONS 10,11
```

2 ONEWAY：一元配置分散分析

ONEWAYは一元配置の分散分析を行うサブプログラムである。もちろんANOVAやMANOVAを使用しても一元配置の分散分析は可能であるが、ONEWAYには多くの追加テストを含むという特徴がある。すなわち、傾向(TREND)テスト、アプリアリな対比(CONTRAST)テスト、さまざまな多範囲(RANGE)検定などの機能をもっている。

2.1 一般書式

ONEWAY 従属変数リスト BY 独立変数リスト (最小値、最大値)

[/POLYNOMIAL=最高次数][/CONTRAST=対比係数リスト]

[/CONTRAST=...]

```
[/RANGES={LSD          }({{0.05}})]
          {DUNCAN       }{alpha}
          {SNK
          {TUKEYB
          {TUKEY
          {LSDMOD
          {SCHEFFE
          {レンジの値のリスト}
```

[/RANGES=...]

2.2 サブコマンド

POLYNOMIALサブコマンドによって、主効果(級間変動)を最高5次までのトレンド成分に展開できる。また、CONTRASTサブコマンドにより、任意のグループ間の平均値の差をt検定することがで

きる。RANGESサブコマンドを使用すると、グループ平均のあらゆる組合せについて、差の検定を行うことができる。

2.3 オプション

- 1 欠損値指定を無視する。
- 2 欠損値をリスト単位で排除する。
- 3 変数ラベルの印刷を省略する。
- 4 各グループの度数、平均、標準偏差をマトリックス出力する。
- 6 独立変数の変数値ラベルの先頭8文字をグループラベルとする。
- 7 オプション4の出力をマトリックス入力する。
- 8 度数、平均、プールされた分散、自由度をマトリックス入力する。
- 10 レンジテストにおける全グループに対して調和平均を使用する。

2.4 追加統計

- 1 各グループごとに、度数、平均、標準偏差、標準誤差、最小値、最大値、平均に対する95%の上側、下側信頼限界の印刷を行なう。
- 2 母数模型に対しては、標準偏差、標準誤差、平均の95%の信頼水準における両側信頼限界を、変量模型に対しては、標準誤差、平均の95%の信頼水準における両側信頼限界、および推定級間分散 (estimate of between-component variance) を印刷する。
- 3 グループ間の分散の均一性についての統計量 (コ克蘭のC、バートレット・ボックスのF、ハートレイの最大分散と最小分散の比) を印刷する。

2.5 コマンドの編成例

```
ONEWAY  WELL BY EDUC6(1,6)
POLYNOMIAL = 2/  CONTRAST = 2*-1,2*1/
CONTRAST = 2*0,2*-1,2*1/  CONTRAST = 2*-1,2*0,2*1/
RANGES = SNK/  RANGES = SCHEFFE (.01)
STATISTICS ALL
```

3 PLOT: プロット

PLOTは、2次元の図を作成するプログラムである。散布図(scatterplot)、等値線図(contour plot)、重ね書き図(overlay plot)、回帰図(regression plot)などが作成できる。プロット位置を示す記号も、オプションによって選択できる。また縦軸、横軸の長さや縮尺も、自由に設定することができる。

3.1 一般書式

```
PLOT  [HSIZEサブコマンド] [/VSIZEサブコマンド]
```

[/CUTPOINTサブコマンド] [/SYMBOLSサブコマンド]

[/MISSINGサブコマンド] [/FORMATサブコマンド]

[/TITLEサブコマンド] [/HORIZONTALサブコマンド]

[/VERTICALサブコマンド]

/PLOTサブコマンド[/PLOTサブコマンド...]

PLOTでは、サブコマンドによって図の明細を指示する。かならず指定しなければならないのは、PLOTサブコマンドだけであり、その他はすべてオプションである。

HSIZE、VSIZE、CUTPOINT、SYMBOLS、MISSINGなどは共通サブコマンドと呼ばれ、それぞれ1回しか使用できず、かならずPLOTサブコマンドより前に置かなければならない。

FORMAT、TITLE、HORIZONTAL、VERTICALなどはローカルサブコマンドと呼ばれ、PLOTサブコマンドとともに何度でも使用できる。ただしローカルサブコマンドの指定は、直後のPLOTサブコマンドに対してのみ有効である。PLOTサブコマンドは、かならず最後に置く。各サブコマンドは、必ず/(スラッシュ)で区切っておく。

3.2 サブコマンド

(1) PLOTサブコマンド

[書式] PLOT = 変数リスト WITH 変数リスト[(PAIR)]

[BY 変数名][;変数リスト ...]

キーワードWITHの左側の変数を縦軸、右側の変数を横軸とする図が作成される。キーワードPAIRを使用すると、WITHの左側のリストの第1番目の変数と右側の第1番目の変数、左側の第2番目の変数と右側の第2番目の変数、…というようにリスト順のペアが作られる。PAIRを使用しない場合には、左側の変数群と右側の変数群の全ての組合せについて、図が作成される。キーワードBYの後ろにはコントロール変数を指定する。この場合、コントロール変数の変数値ラベルの1番目の文字がプロット位置を示すシンボルとして使用される。変数値ラベルがない場合には、変数値の1番最初の文字または数字が使用される。また等値線図のコンター変数もキーワードBYの後に置く。

(2) TITLEサブコマンド

[書式] TITLE = '文字列'

60字以内の文字列で、図に表題(タイトル)をつけることができる。

(3) VERTICAL、HORIZONTALサブコマンド

[書式] (VERTICAL または HORIZONTAL) = ['文字列']

[STANDARDIZE] [REFERENCE(変数値リスト)]

[MIN(最小値)] [MAX(最大値)] [UNIFORM]

'文字列' 縦軸・横軸に40字以内でラベルをつけることができる。

MIN(最小値) 各軸の最小値を指定する。指定された値より小さいデータは図に含まれない。

標準値は、データの最小値。

MAX (最大値) 各軸の最大値を指定する。指定された値より大きいデータは図に含まれない。

標準値は、データの最大値。

UNIFORM 各軸の尺度をそれぞれ統一する。

REFERENCE (変数値リスト) 各軸に対しそれぞれ10本まで、参照のための垂線をひくことができる。

STANDARDIZE 変数値を平均0、標準偏差1となるように標準化する。

(4) FORMATサブコマンド

[書式] FORMAT = {DEFAULT
 {CONTOUR[{10}]}
 {n}
 {OVERLAY
 {REGRESSION

DEFAULT FORMATサブコマンドを省略するか、キーワードDEFAULTを指定すると、2変数の散布図が作成される。

CONTOUR 最高35段階までの等値線図を作成する。キーワードCONTOURの後の()の中に、何段階に分割するかを数字で指定する(標準値は10)。コンター変数は連続変数でなければならない。PLOTサブコマンドのキーワードBYの後に指定する。

OVERLAY 1つの枠の中に最高20個までの図を重ねて書くことができる。ただし重ねることができるのは単純な散布図や回帰図だけであり、コントロール変数やコンター変数を持つ図は使用できない。重ねる図はPLOTサブコマンドの中で指定する。

REGRESSION 縦軸変数の横軸変数への回帰を計算し印刷を行なう。また散布図の枠と回帰直線との交点をRという文字を使って示す。この2つのRを結んだものが回帰直線となる。

(5) HSIZE、VSIZEサブコマンド

[書式] HSIZE = {80}/VSIZE = {40}
 {n} {n}

縦軸の長さ、および横軸の長さを指定する。デフォルト値は縦軸の長さ40、横軸の長さ80である。

(6) SYMBOLSサブコマンド

[書式] SYMBOLS = {ALPHANUMERIC
 {NUMERIC
 {'記号列' [, '重複点の記号列']
 {'X' 16進記号列' [, '重複点の16進記号列']
 {DEFAULT

プロット位置を示す記号を、最高36種類まで指定することができる。ただしコントロール変数を用いる図では、SYMBOLSサブコマンドは適用されない。

ALPHANUMERIC 1～9までの数字とA～Zまでの英字、および*印が記号として用いられる。

したがって*印は、その位置に36以上のケースが重なっていることを示す。このセットがデフォルトとなる。

NUMERIC 1～9までの数字と*印が記号として用いられる。したがって*印は、その位置に10以上のケースが重なっていることを示す。

‘記号列’[‘重複点の記号’] アポストロフィーで囲むことによって、任意の記号のリストを指定することができる。コンマの後に第2のリストを指定すると、第1のリストの記号と第2のリストの記号を重ね打ち（合成）した記号を使用する。第2のリストの記号には、16進数またはキーボードの文字が使用できる。

X'16進記号列'[‘重複点の16進記号列’] Xの後にアポストロフィーで囲んだ16進数のリストを置くことによって16進数の記号を指定することができる。コンマの後に第2のリストを指定すると、第1のリストの記号と第2のリストの記号を重ね打ち（合成）した記号を使用する。第2のリストの記号には、16進数またはキーボードの文字が使用できる。

(7) CUTPOINTサブコマンド

[書式] CUTPOINT = {EVERY(1またはn)}
 {変数値リスト}

このサブコマンドを省略すると、ケース数が1増加するごとに違った記号が割当てられる。1ケースごとではなく、あるまとまったケース数ごとに異なった記号を割当てたいときには、CUTPOINTサブコマンドを使用する。ただし、重ね図(overlay plot)や等値線図(contour plot)、またコントロール変数を持つ図には、適用されない。

EVERY(n) nケースごとに異なる記号が使用される。

変数値リスト 指定された変数値ごとに異なる記号が使用される。

(8) MISSINGサブコマンド

[書式] MISSING = [{PLOTWISE}] [INCLUDE]
 {LISTWISE}

PLOTWISE 欠損値のプロット単位(plotwise)の除去。ある図（プロット）を構成する変数のどれかに欠損値がある場合、そのケースはその図から除去される。MISSINGサブコマンドの指定が無い場合は、この欠損値処理方式が選択される。

LISTWISE 欠損値のリスト単位(listwise)の除去。PLOTサブコマンドで指定された変数のどれか1つにでも欠損値があれば、そのケースはすべての図（プロット）から除去される。

INCLUDE 欠損値定義を無視しすべてのデータを計算に含める。

3.3 コマンドの編成例

```
PLOT HSIZE = 38/  
      VSIZE = 25/  
      SYMBOLS = ',-+oxXNPX' 0#'/
```

```

TITLE='PRESSURE WITH WEIGHT' /
HORIZONTAL='WEIGHT IN POUNDS' /
PLOT = PRESSURE WITH WEIGHT/
FORMAT = CONTOUR/
TITLE='WEIGHT WITH AGE BY PRESSURE' /
HORIZONTAL='YEARS OF AGE' MIN(15) MAX(85)/
PLOT = WEIGHT WITH AGE BY PRESSURE/

```

4 HILOGLINEAR

HILOGLINEARは、度数分布表 (frequency tables) に対する階層的对数線形モデル (hierarchical log-linear model) のパラメーターを推定するプログラムである。HILOGLINEARは、繰返し比例当てはめ法 (iterative proportional fitting algorithm) を使用するため、LOGLINEARよりこのようなモデルにたいして効果的である。

4.1 一般書式

```

HILOGLINEAR  変数リスト (最小値, 最大値) [変数リスト...]
              [/MISSINGサブコマンド] [/CWEIGHTサブコマンド]
              [/PRINTサブコマンド]  [/PLOTサブコマンド]
              [/CRITERIAサブコマンド] [/METHODサブコマンド]
              [/MAXORDERサブコマンド]
              [/DESIGNサブコマンド [/DESIGNサブコマンド/...]]

```

HILOGLINEARコマンドの変数名は、必ず指定しなければならない。サブコマンドは8種類あり、すべてオプションである。ある1つのDESIGNサブコマンドに関係するサブコマンドは必ずDESIGNサブコマンドの前に置かなければならない。

4.2 サブコマンド

(1) DESIGNサブコマンド

[書式] DESIGN [=効果名 効果名*効果名 ...]

階層モデルにおいては、高次の交互作用が存在することは、同じ変数についてより低次の交互作用や主効果も必ず存在するということを意味している。したがって、DESIGNサブコマンドである交互作用を指定したとすれば、それはより低次の交互作用と主効果も同時に指定したことになる。たとえば、

DESIGN = A*B*C D

という指定は、ABCの3次の交互作用、AB、AC、BCの2次の交互作用、およびA、B、C、Dの各主効果を含んだモデルであるということを意味する。

DESIGNサブコマンドの後に何も指定しなければ、飽和モデル（全ての交互作用と主効果を含んだモデル）を選択したことになる。1つのHILOGLINEARコマンドの中で、複数のDESIGNサブコマンドの指定が可能である。

（２）CWEIGHTサブコマンド

〔書式〕 CWEIGHT={変数名
 {(マトリックス)}

CWEIGHTサブコマンドは、あるモデルの各セルに対し、ウェイトづけを行う。あるモデルに対し構造的な0 (structural zeros)を与えたり、新しい枠(margin)にあうように表を調整したりする時に用いる。なお、アグリゲートデータに対するウェイトづけにはWEIGHTコマンドを用い、CWEIGHTサブコマンドは使用しない。各セルに対するウェイトづけの指定は、CWEIGHTサブコマンドの後にウェイト変数を指定するか、括弧で囲まれたウェイトマトリックスを指定することによって行う。

（３）METHODサブコマンド

〔書式〕 METHOD[=BACKWARD]

サブコマンドは、指定されたモデルから項減少法(backward elimination of terms)によって、“最適(best)”なモデルを選び出す。このサブコマンドは、直後のデザインにのみ有効である。METHODサブコマンドを省略すると、DESIGNサブコマンドで指定されたモデルがテストされるだけで、モデル選択は行われない。

（４）MAXORDERサブコマンド

〔書式〕 MAXORDER=k

モデルにおける最高次元数の指定に用いる。k次元以下の全ての交互作用と主効果を含んだモデルを指定したことを意味する。

（５）PRINTサブコマンド

〔書式〕 PRINT=[DEFAULT] [ASSOCIATION]
 [FREQ] [RESID] [ESTIM] [ALL] [NONE]

PRINTサブコマンドを使用することによって、標準出力に含まれない統計値や図表を出力させたり、標準出力の中の不要なものを省略することができる。

DEFAULT 標準出力。飽和モデル以外のモデルに対しては、FREQおよびRESIDキーワードによる出力が行われる。飽和モデルに対しては、FREQ、RESID、およびESTIMキーワードによる出力が行われる。飽和モデルにおいては観測度数と期待度数は一致し、残差は0となる。

FREQ 観測度数および期待度数の出力。

RESID 残差および標準化残差の出力。

ESTIM 飽和モデルに対するパラメータ推定値の出力。

ASSOCIATION 飽和モデルに対して、各効果の部分的な結合を要求する。

ALL 利用可能な全ての出力。

NONE デザインに関する情報とあてはまりのよさに関する統計値のみが出力される。

(6) PLOTサブコマンド

[書式] PLOT=[DEFAULT] [RESID]
 [NORMPLOT] [NONE]]

PLOTサブコマンドがないと、図は何も出力されない。キーワードがすべて省略された場合は、キーワードDEFAULTを指定した場合と同じ図が出力される。

DEFAULT RESIDおよびNORMPLOTキーワードを指定した場合と同じ出力になる。

RESID 標準化残差を縦軸、観測度数および期待度数をそれぞれ横軸とする図が出力される。

NORMPLOT expected normal value および偏差をそれぞれ縦軸とし、標準化残差を横軸とする図が出力される。

NONE 何も出力されない。すでに使用されているPLOTサブコマンドの指定を、無効にするために使用する。

(7) CRITERIAサブコマンド

[書式] CRITERIA=[CONVERGE($\left\{ \begin{smallmatrix} 0.25 \\ n \end{smallmatrix} \right\}$)] [ITERATE($\left\{ \begin{smallmatrix} 20 \\ n \end{smallmatrix} \right\}$)]
 [P($\left\{ \begin{smallmatrix} 0.05 \\ \text{prob} \end{smallmatrix} \right\}$)] [MAXSTEPS($\left\{ \begin{smallmatrix} 10 \\ n \end{smallmatrix} \right\}$)] [DEFAULT]

繰返し比例当てはめ法(iterative proportional fitting) やモデル選択を行うさいに使用する基準値を指定する。

DEFAULT パラメータを標準値にもどす。標準値以外の値を指定したあとで、パラメータを標準値にもどしたい時に用いる。

CONVERGE(n) 収束の基準値。標準値は0.25。

ITERATE(n) 最大繰返し回数。標準値は20。

P(n) χ^2 値の有意水準。標準値は0.05。

MAXSTEPS(n) モデル選択における最大ステップ数の指定。標準値は10。

(8) MISSINGサブコマンド

[書式] MISSING={LISTWISE
 {INCLUDE }
 {DEFAULT } }

LISTWISE 欠損値のリスト単位の除去。

INCLUDE 欠損値定義を無視し全てのデータを計算に含める。

DEFAULT 欠損値の標準処理。

4.3 コマンドの編成例

HILOGLINEAR WATSOFT (1,3) BRANDPRF PREVUSE TEMP (1,2) /

```
PRINT = ALL/  
PLOT = DEFAULT /  
METHOD = BACKWARD /  
DESIGN /
```

5 ALSCAL：多次元尺度法

ALSCALはリリース2.1から利用できるもので、Takane, Young and deLeeuw (1977)により提唱され、後に Young, Takane and Lewycky (1978)により改良された交互最小2乗法 (alternating least-squares approach) を用いた多次元尺度法と展開法のプログラムである。5つの異なるモデルの中の1つを採用し、多次元空間上に刺激座標や重み付けを位置づける。適当なサブコマンドオプションで多くの多次元尺度法や展開法モデルを指定できる。また、刺激座標、重み付け、変換等を含んだ様々なプロットも可能である。ALSCALは極めて柔軟なプログラムで計量多次元尺度法 (MDS) と非計量MDS、古典的MDS、反復MDS、重み付けMDS、展開法 (MDU) を一つのプログラムに統合したものである。現在ではALSCALが個人差多次元尺度モデル (McGee, 1968; Carrol and Chang, 1970) と展開法モデル (Coombs, 1964; Young, 1972)、個人差を扱わない多次元尺度モデル (Torgerson, 1952; Sheperd, 1962; Kruscal, 1964; Guttman, 1968)、さらに外的基準のある多次元尺度モデルや外部展開法モデル (Carrol, 1972) を統合した唯一の尺度構成プログラムである。ALSCALはあらゆる種類の2元データ (例えば単一の行列) や3元データ (例えば、幾人かの被験者おのの行列) に適用できる。データは条件付きでも (conditional or unconditional)、また繰り返しがあるものでもないものでも、さらに欠損値の有無にかかわらず、処理が可能である。順序データの同順位データはそのまま同順位にしておくこともできるし、順位関係をつけることもできる。

5.1 一般書式

ALSCALの一般書式は次の通りで、最低指定しなくてはならないのは VARIABLES サブコマンドである。これにより尺度化する変数を指定する。

```
ALSCAL VARIABLES=変数リスト[/サブコマンド/...]
```

他のサブコマンドはすべてオプションである。デフォルトとしてALSCALは2次元非計量ユークリッド多次元尺度法解を求め、その際データは順序尺度で1つ以上の正方対称行列からなっていると仮定されている。プリント出力は自動的に各行列の欠損値の数、反復手続きによる Young の S-STRESS の値の変化、各入力行列に対する2つの当てはまりの測度 (Kruskal の STRESS と相関の平方、RSQ)、モデルが適切である場合の布置と重み付け、を含むものになる。

5.2 サブコマンド

入力、モデル、出力に関するオプションにおいて、様々なサブコマンドの指定が可能である。ALSCAL のコマンドの後であればサブコマンドはどのような順序でも良い。サブコマンドはサブコマン

ドのキーワードで始まり、続いてイコール記号や指定内容を書かねばならない。サブコマンドを複数指定する場合にはそれぞれをスラッシュで区切る。

(1) FILEサブコマンド

〔書式〕 FILE=handle名

FILEサブコマンドにより付加的データ、例えば、尺度化された行列に対する重み付け、あるいはまた刺激座標に対する初期布置、確定(fixed)布置を含んだファイルを読み込むことができる。FILEサブコマンドを用いる場合にはFILE HANDLEを指定しなくてはならない。

読み込まれる行列や行列の値のタイプに関してオプションを後につけることが可能である。初期布置にはINITIAL、確定布置にはFIXEDを指定する。連続的なALSCAL次元と対応する布置、重み付けファイルの変数はDIM1, DIM2, ..., DIMrの変数名を持つことになる。ここでrはALSCAL次元の最大数である。ファイルはまた全ての行に数値タイプを示すための短い文字型変数(string variable) TYPE_を含む。刺激座標値はCONFIG、行刺激座標はROWCONF、列刺激座標はCOLCONF、そして被験者重み付けと刺激重み付けはそれぞれSUBJWGHTとSTIMWGHTとして示すこと。ALSCALはCONFIGとROWCONFを区別せず受け入れる。CONFIG, ROWCONF, COLCONF, SUBJWGHT, STIMWGHTをそのまま書いてもよいし、おのおのはじめの3文字だけでもよい。

ファイル上の座標や重み付けのセットはCON, ROW, COL, SUB, STIの順でなくてはならない。それらが適切な順序でファイル上に現われる限りALSCALは不用なタイプは読み飛ばす。

CONFIG[{INITIAL}] 刺激布置の読み込み。布置／重み付けファイルが初期刺激座標

{FIXED}を含んでいる。この種の入力にはSHAPE=SYMMETRICかSHAPE=ASYMMETRIC、あるいは行列中の変数の数がALSCALコマンド上の変数の数に等しい場合に適切である。

ROWCONF[{INITIAL}] 行刺激布置の読み込み。布置／重み付けファイルが初期行刺激座標を含

{FIXED}んでいる。SHAPE=RECTANGULARの場合にはこの指定が必要である。観測値の数はINPUTサブコマンドで指定した行数と等しいか、INPUTを指定しない場合には距離行列ファイル(proximity file)におけるケース数と等しいこと。

COLCONF[{INITIAL}] 列刺激布置の読み込み。布置／重み付けファイルが初期列刺激座標を含

{FIXED}でいる。この種のファイルはSHAPE=RECTANGULARで行列中の観測値数がALSCALのコマンドにおける変数の数と等しい場合に入力となり得る。

SUBJWGHT[{INITIAL}] 被験者重み付け(行列)の読み込み。布置／重み付けファイルは被験者の

{FIXED}重み付けを含んでいる。被験者一重み付け行列における観測値の数は距離行列ファイルに於ける行列の数に等しくなくてはならない。被験者重み付けの入力は、MODEL=INDSCAL、MODEL=AINDSあるいはMODEL=GEMSCALの時のみ許される。

STIMWGHT[(INITIAL)] 刺激重み付けの読み込み。布置／重み付けファイルが刺激重み付けである。

{ FIXED } 布置／重み付けファイルに於ける観測値数が距離行列ファイルに於ける行列数と等しいときにのみこのオプションを指定できる。この種の行列入力にはMODEL=AINDS あるいは MODEL=ASCAL のときのみ許される。

(2) INPUTサブコマンド

[書式] INPUT = ROWS { ALL
 n }

INPUT サブコマンドは矩形行列（行の数と列の数が異なる行列）の読み込みかたをALSCALに指定するものである。一般にALSCALは実行ファイルにおける各ケース（これはデータ行列における一行を表わしている）を一行単位に読み込むことになる。もし矩形行列を読み込む際に INPUT サブコマンドを用いない場合には入力行列の一行が各ケースに対応していると仮定される。データ行列が複数のデータ行列から成るのではなくそれ自体で一つのデータを構成しているなら ALL を指定し、ファイルに複数の入力矩形行列がある場合には各行列の行数を n で指定する。全行数が n で割られ自動的にそのデータ行列に含まれる行列の数が求められる。

(3) SHAPEサブコマンド

[書式] SHAPE = { SYMMETRIC
 ASYMMETRIC
 RECTANGULAR }

SHAPEサブコマンドは尺度化する入力行列の形式を指定するものである。指定する形式は次のように、対称行列、非対称行列、矩形行列のいずれかである。

SYMMETRIC 対称行列。これはデフォルトである。

ASYMMETRIC 非対称行列。行列の上三角部分と下三角部分の対応成分が異なる行列。対角成分は無視される。

RECTANGULAR 矩形行列。行と列は異なるアイテムを表わしている。

(4) LEVELサブコマンド

[書式] LEVEL = { ORDINAL ({ UNTIE } [SIMILAR])
 INTERVAL ({ 1 または d })
 RATIO ({ 1 または d })
 NOMINAL }

LEVELサブコマンドは行列内の測度の水準、即ち、順位尺度、間隔尺度、比尺度、名義尺度の水準を指定するものである。オプションとして間隔尺度と比尺度には多項式変換 (polynomial transformation) の次数を指定できる。次数は LEVEL=INTERVAL あるいは LEVEL=RATIO の後の括弧の中に (d) の形で、数字で指定する。d には 1 から 4 までの整数を指定できる。デフォルト値は 1 であり、これは線型変換に対応している。

ORDINAL[({ UNTIE } [SIMILAR])] 順位水準データ。デフォルトである。データは Kruskal (1964) の最小 2 乗単調変換 (least-squares monotonic transformation) により順位として扱

われる。この分析は非計量（ノンメトリック）である。かっこ内のパラメータを省略すると、分析を通して同順位データを同順位のままとし、データは非類似性として扱われる。

順位データを連続量と見なしたいのであれば SIMILAR オプションを指定すればよい。

INTERVAL[(d)] 間隔水準データ。古典的な回帰技法を用いて計量的分析が行われる。

RATIO[(d)] 比水準データ。間隔水準と同様、計量的分析が行われる。

NOMINAL 名義水準データ。Takane, Young, and deLeeuw (1977)の最小2乗カテゴリカル変換を用い、データを名義水準として扱う。名義水準データの非計量分析にこの指定を行う。

(5) CONDITIONサブコマンド

[書式] CONDITION={ MATRIX
 ROW
 UNCONDITIONAL }

条件性(Conditionality)とはデータセット内の数字が比較できない状態を言う。CONDITION サブコマンドは尺度化するデータ行列が条件付きであることを指定するものである。

MATRIX 数字の意味が被験者に関して条件付きである。これはデフォルトである。

ROW 数字の意味が行に関して条件付きである。各行列の行内の数字間でのみ比較が意味を持つ。この指定は非対称データあるいは矩形行列データに対してのみ適当で、MODE L=ASCAL あるいは MODEL=AINDSの場合は指定できない。

UNCONDITIONAL 数字の意味が条件付きではない。入力行列に於ける全ての値の比較が意味を持つ。

(6) MODEL (METHOD) サブコマンド

[書式] {MODEL または METHOD}={EUCLID
 INDSCAL
 ASCAL
 AINDS
 GEMSCAL}

MODELサブコマンドは分析のための尺度化モデルを指定するためのものである（キーワード METHOD を MODEL の代わりに用いることもできる）。5つの尺度化、展開モデルのタイプからいずれかを指定できる。

EUCLID ユークリッド距離モデル。このモデルはデフォルトである。

INDSCAL 個人差（重み付け）ユークリッド距離モデル。ALSCALはCarrol and Chang (1970)が提案した重み付き個人差ユークリッド距離モデルを使ってデータを尺度化する。

ASCAL 非対称ユークリッド距離モデル。このオプションは Young (1975b)が提案したユークリッドモデルである。

AINDS 非対称個人差ユークリッド距離モデル。このオプションはYoung (1975a)の非対称ユークリッドモデルと Carrol and Chang(1970) が提案した個人差モデルを組み合わせたもの

である。

GEMSCAL 一般化ユークリッド計量個人差モデル。

(7) CRITERIAサブコマンド

```
[書式]  CRITERIA=[NEGATIVE] [CUTOFF{0}] [CONVERGE{.001}]  
                                     {c}          {c}  
        [ITER{30}] [STRESSMIN{.005}] [NOULB]  
          {ni}          {s}  
        [DIMENS{ 2 }]  
          {min[,max]}] [DIRECTIONS(r)]  
        [TIESTORE(n)]
```

CRITERIA サブコマンドは尺度化するモデルの諸点を指定し、尺度化の解の収束基準を設定するものである。ALSCALはS-STRESSと呼ばれる適合基準を最小化しようデザインされた反復アルゴリズム (iterative algorithm) である。ALSCALはS-STRESSの改善が CONVERGE の値以下になるか、反復回数が最大になるか、あるいはS-STRESSの値が最小値に達するまで反復計算を繰り返す。

CONVERGE(c) CONVERGE を c に設定する。これは一回の繰り返しでS-STRESSがどれほど改善されたかの基準である。デフォルトではCONVERGE=0.001 である。CONVERGE=0 の場合には ITER オプションで違う値を指定しない限り、30回の反復計算が行われる。

ITER(ni) 反復の最大回数をniに設定する。デフォルトは30である。STRESSMIN(s)最小ストレスをsに設定する。デフォルトではS-STRESSの値が 0.005 以下になると反復計算は自動的に終了する。STRESSMIN には0から1までの値を設定できる。

NEGATIVE 個人差モデルで負の重み付けを認める。デフォルトでは重み付けが負の値をとることは許されない。

CUTOFF(c) 欠損値として距離を処理する際の打ち切り値をcに設定する。

デフォルトでは負の類似性 (あるいは非類似性) は全て欠損値として処理され、0と正の類似性は欠損値として扱われない (CUTOFF = 0)。CUTOFF 値を変えることでその値以上の類似性は欠損値として扱わないように指定できる。ユーザー指定の欠損値とシステム欠損値は CUTOFF 指定に関係なく欠損値と見なされる。

NOULB 欠損値の上界と下界 (upper and lower bounds) の推定を行なわない。デフォルトでは初期布置を計算するために上界と下界の推定が行われる。

DIMENS(min[,max]) 尺度解 (scaling solution) の最小次元数と最大次元数を設定する。デフォルトでは2次元の尺度解が計算される。最小数と最大数は2から6までの整数値であること。括弧の中に一つだけ整数値を指定してもよい。

DIRECTIONS(r) 一般化ユークリッドモデルにおける主方向 (principal direction) 数をrに設定する。このオプションは MODEL=GEMSCAL 以外では効果を持たない。主方向数は1か

ら DIMENS オプションで指定した次元数までのいずれか正の整数となる。デフォルトでは方向数は次元数に等しく設定される。

TIESTORE(n) 同順位処理に必要な記憶容量 (storage) を n に設定する。このオプションは順位データにおける同順位処理に必要な記憶容量を指定するものである。デフォルトでは n は 1000 か行列のセル数、いずれか少ない方に設定される。これで不十分な場合にはより多くの容量が必要であるというメッセージを出して ALSCAL は終了する。

(8) PRINT サブコマンド

[書式] PRINT = [DATA] [HEADER] [INTERMED]

標準以外の出力を指定するためには PRINT サブコマンドを利用する。

DATA 入力データのプリント。

INTERMED 尺度化プロセスにおける中間ステップの印刷。素データ、欠損値パターン、欠損値の推定がなされたデータ、規準化されたデータや平均、推定された加算定数を持つ平方データ、各個人のスカラー積、など印刷される。

HEADER ヘッダーページのプリント。この指定をすると、以下に述べるその分析に於けるデータ、モデル、出力、アルゴリズムそれぞれのオプションリストがプリントされる。

Data Options 行と列の数、行列の数、尺度のレベル、データ行列の形、データのタイプ (類似性が非類似性か)、同順位は同順位のままか、あるいは同順位でないようにするか、条件性、データ打ち切り値、が含まれる。

Model Options 指定したモデルタイプ (EUCLID, INDSCAL, ASCAL, AINDS, GEMSCAL)、次元の最小値と最大値、負の重み付けが許されているか、が含まれる。

Output Options 出力オプションにはプリント出力がジョブオプションヘッダ (job option header) とデータ行列を含むか否か、布置と変換のプロットをしたか、出力データセットを作ったか、初期刺激座標、初期列刺激座標、初期被験者重み付け、初期刺激重み付けなどを計算したか、が含まれる。

Algorithmic Options 最大反復回数、収束基準、最大 S-STRESS 値、上界と下界により欠損値が推定されているか否か、順位データにおける同順位処理に割り当てられている記憶容量が含まれる。

(9) PLOT サブコマンド

[書式] PLOT=[DEFAULT] [ALL]

多次元尺度解をプロットする。

DEFAULT 標準ののプロットを行なう。刺激座標のプロット、行列重み付けのプロット (MODEL=INDSCAL, MODEL=AINDS, MODEL=GEMSCAL の場合。MODEL サブコマンド参照)、刺激重み付

けのプロット (MODEL=AINDS, MODEL=ASCAL の場合)、がある。さらにデフォルトではデータとモデルの間の線型適合度プロット、データのタイプによっては非線型適合度やデータ変換のプロットも含まれる。ALSCALは $d*(d-1)/2$ ページに刺激空間をプロットし、要求された場合には被験者の各重み付け空間も同じスペースでプロットする。ここで d は解に於ける次元数である。

ALL 各行列に対する刺激空間のプロット。PLOT=ALL を指定することにより各被験者のデータ変換のプロット (CONDITION=MATRIX あるいはデータが名義尺度、順位尺度の場合)、各行 (CONDITION=ROW の場合) に関するプロット、重み付けモデルに関しては各被験者の重み付け空間のプロット、が得られる。このオプションでは、特に CONDITION=ROW の場合、出力が数千ページに及ぶこともある。

(10) OUTFILEサブコマンド

[書式] OUTFILE=handle名

handle名で指示したシステムファイルに、座標と重み付け行列を保存する。システムファイルは TYPE_ で名付けられた英数字 (短い文字型) 変数を持ち、各行の値の種類を示している。数字変数 DIMENS は次元の数を、数字変数 MATNUM は各座標セットが対応する被験者 (行列) を、また変数 DIM1, DIM2... DIMr はモデルの r 次元に対応する変数を示している。分割ファイル(split-file)の変数値もまたこのファイルに含まれる。

以下のリストは7種のALSCALのファイル出力である。各識別子(identifier)のはじめの3文字だけがシステムファイルの変数 TYPE_ に書き込まれる。例えば、CONFIG なら CON になる。

CONFIG	SHAPE=SYMMETRIC あるいは SHAPE=ASYMMETRIC に対する刺激布置座標
ROWCONF	SHAPE=RECTANGULAR に対する行刺激布置座標
COLCONF	SHAPE=RECTANGULAR に対する列刺激布置座標
SUBJWGHT	MODEL=INDSCAL, MODEL=AINDS, MODEL=GEMSCAL に対する被験者 (行列) 重み付け
FLATWGHT	MODEL=INDSCAL, MODEL=AINDS, MODEL=GEMSCAL における平坦化被験者 (行列) 重み付け (flattened subject(matrix) weights)
GEMWGHT	MODEL=GEMSCAL に於ける一般化重み付け
STIMWGHT	MODEL=ASCAL, MODEL=AINDS に於ける刺激重み付け

5.3 制限事項

- 1 次元数は最大6まで。
- 2 データに重み付けを行うことはできない。
- 3 一課題(a partition)で扱えるデータ数は32,767まで。

5.4 コマンドの編成例

[例1] ALSCAL VARIABLES = 変数リスト

```

/INPUT = ROWS(8)
/SHAPE = RECTANGULAR
/CONDITION = ROW
/LEVEL = INTERVAL(2)
/MODEL = EUCLID

```

この例では1行が1ケースに対応しており、行列が8行からなることを示している。データ行列は矩形で行内でのデータ比較のみが意味を持つ。この場合は次数が2なので、1次と2次の項を持つ多項式変換を意味している。モデルとしてはユークリッド距離モデルを採用している。

[例2] ALSCAL VARIABLES = 変数リスト

```

/SHAPE = ASYMMETRIC
/LEVEL = ORDINAL
/MODEL = ASCAL
/FILE = FIRST CONFIG(FIXED) STIMWGHT(INITIAL)

```

この例では刺激座標が固定値として、刺激重み付けは純粋に初期値として読み込まれる。

[例3] ALSCAL VARIABLES = 変数リスト

```

/SHAPE = RECTANGULAR
/INPUT = ROWS
/CONDITION = ROW
/FILE = OUTSET ROWCONF

```

上のコマンドでは OUTSET が行刺激布置を含んでいることを示している。距離行列ファイルは矩形であり、各ケースはデータ行列一行に対応している。そして行列の行内での値の比較のみが意味を持っている。

6 PROXIMITIES : 距離行列作成

PROXIMITIESはデータを距離行列、類似性データ行列、非類似性データ行列に変換するサブプログラムである。ケース間、あるいは変数間の距離測度をユークリッド距離をはじめとする様々な尺度に変換、規準化できる。扱える原データは連続量データ、度数データ、2値データいずれでもよい。出力は一般的な行列の形の他、MDSで必要とされることが多い入力データの三角行列とすることも可能である。出力結果はそのままCLUSTER、FACTOR、ALSCALの入力データとして用いることができる。特にさまざまなMDSプログラムと等価なアルゴリズムを実現できるALSCALの使用者には有益なサブプログラムと思われる。

6.1 一般書式

PROXIMITIES 変数リスト/MISSINGサブコマンド/

READサブコマンド/STANDARDIZEサブコマンド/
 VIEWサブコマンド/MEASUREサブコマンド/IDサブコマンド/
 PRINTサブコマンド/WRITEサブコマンド/
 OUTFILEサブコマンド

6.2 サブコマンド

(1) MISSINGサブコマンド

[書式] MISSING={INCLUDE}
 {LISTWISE}

LISTWISEなら欠損値を持つケースを表単位に処理から除く（デフォルト）。INCLUDEだとユーザ定義の欠損値を処理に含める。このオプションでは、システムの欠損値のみが表単位で除去される。

(2) READサブコマンド

[書式] READ = [SIMILAR] [{SQUARE }]
 {TRIANGULAR }
 {SUBDIAGONAL}

READサブコマンドは距離行列作成のための行列の入力を指示するものである。類似性データ、非類似性データのラベルを与え、行列の形式を指定する。SIMILAR指定で類似性データであることを指示し、TRIANGULARは、入力行列が対角要素を含む下三角行列であることを、またSUBDIAGONALは、対角要素を含まない下三角行列であることを指示する。SQUAREはデフォルトで、入力行列は正方行列と見なされる。TRIANGULAR 行列は下三角行列であること。対角要素が含まれる。

(3) STANDARDIZEサブコマンド

[書式] STANDARDIZE = [{VARIABLE}] [{Z_____}]
 {CASE } {SD }
 {RANGE }
 {MAX }
 {MEAN }
 {RESCALE }
 {NONE }

STANDARDIZEサブコマンドはデータの規準化方法を指定するものである。各変数ごとに規準化するか、各ケースごとに規準化を行なうかをはじめに指定し、次に規準の仕方を指定する。デフォルトでは変数ごとに平均0、標準偏差1のZ得点になるように規準化を行う。ただし、STANDARDIZEサブコマンド自体を指定していない場合には NONEがデフォルトになり、規準化は行なわれない。

デフォルト以外の規準化方法としては

SD 標準偏差が1になるように規準化を行う。
 RANGE レンジが1になるように規準化を行う。
 MAX 最大値が1になるように規準化を行う。
 MEAN 平均が1になるように規準化を行う。
 RESCALE データが0から1の範囲になるように規準化を行う。

がある。いずれも割り算の処理がなされるので、もし0で割るような場合には注意が必要である。0で割るということが生じた場合には各方法により対処法が異なる。詳細はの英文マニュアルを参照すること。

(4) VIEWサブコマンド

[書式] VIEW = {CASE
 {VARIABLE}}

VIEWサブコマンドは、ケース間の距離を求めるのか、変数間の距離を求めるのかを選択するためのものである。省略された場合にはケース間の距離(CASE)が選択される。

(5) MEASUREサブコマンド

[書式] MEASURE = [{測度}] [ABSOLUTE] [REVERSE] [RESCALE]

MEASUREサブコマンドではどの測度を用いるか(測度)、および出力データの変換を指定する。変換方法次の3つがあり、並べた順にしたがってデータは変換される。

ABSOLUTE 距離の絶対値を採用する。

REVERSE 類似性を非類似性に、あるいはその逆に非類似性を類似性にする変換を行う。

RESCALE 0から1の範囲に距離の値を再スケールする。RESCALEでは始めにデータから最小値を引き、次にレンジで割ることで距離を標準化する。

連続量に対する測度

NONE 距離測度の計算を行わない。READサブコマンドを用いて既存の距離行列を入力するときのみ使う。

EUCLID ユークリッド距離。デフォルトである。

SEUCLID 平方ユークリッド距離(Squared Euclidean distance)。

CORRELATION 相関。

COSINE 方向余弦(Cosine of vectors of values)。類似性の測度である。

CHEBYCHEV チェビシェフの距離。

BLOCK 市街距離。

MINKOWSKI(p) ミンコフスキーのパワー距離。

度数データに対する測度

CHISQ 2つの度数の等しさに関する χ^2 検定

PH2 度数間の ϕ^2 検定

2値データに対する測度

オプションとして2つの整数値パラメーター、p(present)とnp(not present)が指定できる。もし両方のパラメーターを指定すると始めの値を特性が存在している指標として、第2番目の値を特性が無い指標として見なされる。それ以外の数値は無視される。

RR[(p[np])]	ラッセルとラオ(Russel and Rao)の類似性測度。
SM[(p[np])]	単純対応類似測度(Simple matching similarity measure)。
JACCARD[(p[np])]	JACCARD類似性測度。
DICE[(p[np])]	DICE(or Czekanowski or Sorenson)類似性測度。
SS1[(p[np])]	Sokal and Sneathの類似性測度1。
RT[(p[np])]	Rogers and Tanimoto の類似性測度
SS2[(p[np])]	Sokal and Sneathの類似性測度2。
K1[(p[np])]	Kulczynskiの類似性測度1。
SS3[(p[np])]	Sokal and Sneathの類似性測度3。

条件付き確率

K2[(p[np])]	Kulczynskiの類似性測度2。
SS4[(p[np])]	Sokal and Sneath の類似性測度4。
HAMANN[(p[np])]	Hamann 類似性測度。

予測性測度 (Predictability Measures)

LAMBDA[(p[np])]	グッドマンとクラスカルのラムダ測度 ; Goodman and Kruscal lambda。
D[(p[, np])]	Anderberg の D。
Y[(p[, np])]	Yule の 結合(colligation) Y 係数。
Q[(p[, np])]	Yule の Q 。

その他の2値測度

残りの2値測度は連続変数に対しての2値測度と等価な連関測度か、あるいはアイテム間の関連に関する特殊な特性を持った測度である。

OCHIAI[(p[, np])]	: Ochiai の類似性測度。
SS5[(p[, np])]	: Sokal and Sneath の類似性測度5。
PHI[(p[, np])]	: 四分点相関係数
BEUCLID[(p[, np])]	: 2値ユークリッド距離。
BSEUCLID[(p[, np])]	: 2値平方ユークリッド距離
SIZE[(p[, np])]	: サイズ差測度。
PATTERN[(p[, np])]	: パターン差測度。
BSHAPE[(p[, np])]	: 2値形状差測度(Binary shape difference)。
DISPER[(p[, np])]	: 散らばり類似性測度(Dispersion similarity measure) 。
VARIANCE[(p[, np])]	: 分散非類似性測度(Variance dissimilarity measure)。
BLWMN[(p[, np])]	: Lance-and-Williams のノンメトリック非類似度。Bray-Curtis ノンメトリック係数としても知られている非類似性測度である。

(6) IDサブコマンド

[書式] ID=変数名

IDサブコマンドはケース識別用の文字型変数を指定するために用いられる。識別子(identifier)として、ファイル中のどの文字型変数も指定できる。PROXIMITIESはこの変数の値によって、出力におけるケースを識別する。省略時は、ケースナンバーだけでケースを識別する。

(7) PRINTサブコマンド

[書式] PRINT={PROXIMITIES
 {NONE}}

常に計算される測度名とケース数のプリントに加えて、PRINT サブコマンドで距離行列のプリントを指定できる。PROXIMITIES アイテム間の距離行列をプリントする(デフォルト)。NONE 距離行列をプリントしない。

(8) WRITEサブコマンド

[書式] WRITE=PROXIMITIES

WRITEサブコマンドを用いて、計算された距離行列を出力ファイルに保存できる。

(8) OUTFILEサブコマンド

[書式] OUTFILE={*
 {handle}}

計算された距離行列を出力するファイルを指定する。SPSS*システムファイル、あるいは実行ファイルにセーブできる。パラメータになにも指定しないかアスタリスクを指定すると、距離行列と分割ファイル変数(split-file variables)を実行ファイルに書き出す。ALSCALではこの新しい実行ファイルを距離行列の入力として、直接に用いることができる。別のSPSS*システムファイルに書き出すためには、OUTFILE サブコマンドでファイルを指定しなくてはならない。もし、VIEW=VARIABLE (VIEWサブコマンド参照) ならば、新たなファイルの変数名は元の変数の変数名とラベルを持つことになる。また、VIEW=CASE (標準値) ならば新たなファイルの変数名は CASE1, CASE2,CASEn となる。ここで n はファイルのケース数、あるいは分割されたファイル上のグループのケース数である。新たなファイルでは分割ファイル変数の変数名と値は無視される。

6.3 制限事項

- 1 係数を計算するアイテム(ケースや変数)の数、ケース数が増大するに従い、必要記憶容量は急激に増加する。PROXIMITIESは素データをメモリー内のカレント分割ファイル(current split-file group in memory)に蓄えている。距離行列を計算する場合にはケースと下三角行列を保存している。80Kバイトの作業領域では、150のケースをわずかな変数を用いてしか距離計算できないかもしれない。ケースや変数が多いときは、相当に大きい作業領域が必要である。
- 2 PROXIMITIESは係数を計算する際にケースに対する重み付けを無視する。

7 QUICK CLUSTER: クイック・クラスター

大規模な数のケースを指定したグループ数に効率よくクラスター化するためにQUICK CLUSTERを利用できる。QUICK CLUSTERのアルゴリズムではクラスター変数の値を基にクラスターの中心(centers)を見出し、最も近いケースから中心に割り当てていく。QUICK CLUSTERは2.0版から利用できる。ケースをクラスター化した後で、QUICK CLUSTERでは最終クラスターの中心と各クラスターのケース数を印刷する。またQUICK CLUSTERでは各ケースがどのクラスターに所属するか(cluster membership)も表示する。引き続いての分析や報告のために、最終クラスター中心、各ケースの分類、仮の(interim)分類クラスター中心からの距離を保存して置くことができる。

QUICK CLUSTERアルゴリズムのデフォルトである NOINITIAL オプションを指定すると McQueen の k-平均クラスター法と等価なアルゴリズムになる。

7.1 一般書式

QUICK CLUSTERの書式は次のとおりである。

```
QUICK CLUSTER 変数リスト[/MISSINGサブコマンド/  
READサブコマンド/INITIALサブコマンド/CRITERIAサブコマンド/  
PRINTサブコマンド/WRITEサブコマンド/SAVEサブコマンド]
```

変数リストのみ必須で、他にサブコマンドにより、7つのオプションを好きな順に指定できる。

7.2 サブコマンド

(1) MISSINGサブコマンド

```
[書式] MISSING={LISTWISE } [INCLUDE]  
              {PAIRWISE }  
              {DEFAULT }
```

欠損値を持つケースの扱いを変えることができる。

LISTWISE 欠損値を持つケースを表単位に除く。DEFAULT を指定することでこのデフォルト処理を明示できる。

PAIRWISE 欠損値を変数の組合せ単位(pairwise)に除く。

INCLUDE 使用者定義の欠損値を計算に含める。

(2) READサブコマンド

```
[書式] READ=[INITIAL]
```

初期クラスター中心を行列入力ファイルから読み込むには READ サブコマンドを用いる。オプションとしては READ のすぐ後にキーワード INITIAL を指定できる。もし、固定書式の行列入力を行なう場合には、5F16.5の書式で皆同じく書かれているものと見なされる。

(3) INITIALサブコマンド

```
[書式] INITIAL = 数値のリスト
```

プログラム中で初期クラスター中心をセットするには INITIAL サブコマンドを用いることがで

きる。要求するクラスターの数だけ各クラスター変数についての値を与えてはならない。

[例] QUICK CLUSTER A B C D

/CRITERIA = CLUSTER(3)

$$/INITIAL = (13 \ 24 \ 1 \ 8 \ 7 \ 12 \ 5 \ 9 \ 10 \ 18 \ 17 \ 16)$$

この例では、4つのクラスター変数を指定し、3つのクラスターを要求しているので12の数値を与えなくてはならない。初めの初期クラスター中心は変数Aについては13、Bについては24、CとDについてはそれぞれ1と8になる。

(4) CRITERIAサブコマンド

[書式] CRITERIA=[CLUSTER ($\begin{matrix} \{2\} \\ \{k\} \end{matrix}$)] [NOINITIAL] [NOUPDATE]

CRITERIA サブコマンドを用いて、クラスターの数やオプションの数を指定できる。

CLUSTER(k) クラスターの数。kによってケース数が割り当てられるクラスター数を指定する。

デフォルトは2つのクラスターである。

NOINITIAL 初期クラスター中心を選択しない。このキーワードでは初期中心として欠損値なしの初めのkケースを用いることを指定するものである。

NOUPDATE クラスタ中心を更新しない。このオプションにより、迅速なクラスタ化が可能になる。しかし、結果はデフォルトによるものほど良いものではない。以前にクラスタ化した結果を基にケースを分類するような場合にはこのオプションが使える。

(5) PRINTサブコマンド

```
[書式] PRINT=[INITIAL] [CLUSTER] [ID (変数名)]
        [DISTANCE] [ANOVA] [NONE]
```

印刷出力に対して以下の指定を行なう。

INITIAL 初期クラスター中心の印刷。デフォルトである。

ID(変数名) ケースに対する同定変数(Identifying variable for cases)。ファイル上の変数を識別子(identifier)として用いることもできる。QUICK CLUSTERではこの変数の値を出力に於けるケースを識別するために用いる。デフォルトではケース番号をケースの識別に用いる。

CLUSTER クラスターの所属(Cluster membership)の印刷。

DISTANCE 最終クラスター中心間の距離の印刷。

ANOVA 各クラスター変数に関する単一F検定の印刷。このF検定は単に記述的なものである。クラスター間に差がないとする帰無仮説を検定するためのものとして結果の確率を解釈すべきではない。ケースはクラスター化する変数について差が最大になるように組織的にクラスター化されているのである。

NONE 条件付きではない出力(Unconditional output)。SPSSXの2.1版から利用できるようになった。分類と最クラスター中心、各クラスターのケース数、のみを印刷する。他の指定は全て無視される。

(6) WRITEサブコマンド

〔書式〕 WRITE [=FINAL]

WRITE サブコマンドを用いることにより、出力ファイルに最終クラスターの中心行列を保存できる。オプションとして WRITE の後に、キーワード FINAL を指定することもできる。出力書式は5F16.5となる。

(7) SAVEサブコマンド

〔書式〕 SAVE=[CLUSTER (変数名)] [DISTANCE (変数名)]

ケースのクラスター所属と分類クラスター中心からの距離を、新しい変数として実行ファイルに保存する。

CLUSTER(変数名) クラスター所属の保存。保存された変数は各ケースが所属するクラスターを示す整数値になっている。値は1からクラスター数までである。

DISTANCE(変数名) 最も近い分類クラスター中心からの距離を保存する。

8 SURVIVAL：生存（表）分析

SURVIVALは生存表(life table)の産出、生存関数のプロット、サブグループの比較を行なうものである。グラフィック装置や他のプログラムで使えるように出力ファイルを作成できる。生存分析は従属変数が初期事象(initial event)と終末事象(termination event)の間の時間間隔(time interval)を表わしているときに有用である。このタイプの分析は医学研究で最も頻繁に行われる。病気の発見が初期事象であり、患者の死亡が終末事象である（従い、文字どおりの生存）。結婚から第一子誕生までの時間や同じ会社での採用年数のような他の事象を研究するのに利用できる。中途打ち切り観測値(Censored Observations)（終末事象が分析時点までに生じていないとか、他の原因での死亡や連絡を取り続けられなくなったとかで研究から抜けてしまうことで生じる）も欠損値としては扱わず、中途打ち切りではない観測値の両方を生存時間の計算に利用する。生存得点(Survival Score)、統計量 D (statistic D)（D統計量が大きいほど、サブグループは異なる生存分布からのものである可能性が高い）、Dの有意水準、リストされた各変数について生存表(life table)、生存関数のプロットと特定のサブグループの比較、を出力できる。また、他のプログラムで利用できるように生存表をファイルに書き込むことができる。

8.1 書式

SURVIVAL TABLESサブコマンド/INTERVALSサブコマンド/

[/PLOTSサブコマンド][/COMPAREサブコマンド]

(1) TABLESサブコマンド

「BY コントロール変数リスト (最小値、最大値) ...」

〔例〕 TABLES = ONSSURV, RECSURV BY TREATMNT (1,3) BY SEX(1,2)

(2) INTERVALSサブコマンド

INTERVALS サブコマンドは調べる期間を決定し、分析に於ける時間をいかにグループ化するかを定めるためのものである。

INTERVALS = THRU 240 BY 12 /

(3) STATUSサブコマンド

〔書式〕 STATUS = 状態変数 {最小値、最大値} FOR { ALL
数値 生存変数リスト }

大阪大学大型計算機センターニュース

名を指定する。値の範囲は終末事象が生じたことを示すコードを認識するためのものである。

```
[例] TABLES=ONSSURV BY TREATMNT (1,3) /  
      INTERVALS=THRU 50 BY 5, THRU 100 BY 10  
      STATUS=OUTCOME (3,4) FOR ONSSURV
```

この例では、OUTCOME の3あるいは4のコードが生存変数 ONSSURV に終末事象が生じたことを示している。もし、OUTCOME のコードが“1”生存、“2”研究から脱落、“3”無関係の原因で死亡、“4”病気による死亡、ならば、原因がなんであれ死んだ患者はすべて終末事象に至ったと見なされる。値の範囲にコードが当てはまらない観測値はすべて中途打ち切りケースと分類される。

(4) PLOTSサブコマンド

```
[書式] PLOTS = { ALL  
                {LOGSURV}  
                {SURVIVAL}  
                {HAZARD}  
                {DENSITY} } = { ALL  
                                {生存変数リスト} } BY { ALL  
                                                        {独立変数リスト}  
                ... BY { ALL  
                        {コントロール変数}
```

生存表に加え、SURVIVALでは3つの生存関数をプロットできる。PLOTS サブコマンドではそのあとに求めるプロットに応じて5つのキーワードの一つを指定する。デフォルトでは要求された各関数は各生存変数についてプロットされる。一次のコントロール変数の個々の値はプロット上ではコードによって示される。もし、二次のコントロール変数が使われた場合には二次のコントロール変数の各値について別々なプロットが出力される。TABLES サブコマンドに指定された変数のみがリストでき、生存変数、一次コントロール変数、二次のコントロール変数それぞれの役割は変更できない。変数のグループを指定する際に T0 を用いることができる。

(5) COMPAREサブコマンド

```
[書式] COMPARE = { ALL  
                 {生存変数リスト} } BY { ALL  
                                         {独立変数リスト}  
                 ... BY { ALL  
                         {コントロール変数リスト}
```

コントロール変数によって定義されたサブグループ間の生存を比較するためには、“COMPARE” サブコマンドを用いる。変数を指定せずに COMPARE を指定すると、“TABLES” 変数リストを使ったデフォルトの比較セットが出力される。比較が可能になるためには最低一つの生存変数と一つの一次コントロール変数が必要である。

8.3 オプション

- 1 欠損値を計算に含める。利用者定義の欠損値が分析に含まれる。
- 2 欠損値を表単位(listwise)に除く。
- 3 比較のみを行なう。TABLES サブコマンドで指定した生存表は計算されず、プロットの要求も無視される。
- 4 生存表の印刷の省略。プロットと比較結果のみの出力。

- 5 メモリーが不十分なときには近似比較を計算する。
- 6 近似比較のみを行う。
- 7 対単位の比較の実行。
- 8 生存表データレコードの書き出し。全ての生存表統計量がファイルに書き出される。
- 9 生存表データとラベルレコードの書き出し。変数名、変数ラベル、数値ラベルが生存表統計量と共に書き出される。

ファイルへの書き出しにあたってはオプション8と9を用いるが5種類の書式が決められている。詳細は英文マニュアル参照のこと。

8.4 制限事項

- 1 生存変数の最大数は20
- 2 一次と二次のコントロール変数リスト上の変数はあわせて100まで
- 3 INTERVALS サブコマンドに指定できる THRU BY の最大数は20まで
- 4 プロットできる値の数は35まで

9 NONPAR CORR：順位相関係数

NONPAR CORRは2つの順位相関係数、スピアマンの ρ (rho)とケンドールの $\tau - b$ (tau-b)を計算し、有意水準を求めるプログラムである。デフォルトではスピアマンの順位相関係数が出力される。ケンドールの順位相関係数のみを出力させることも、両方の順位相関係数を出力させることも可能である。有意水準に関するデフォルトでは各係数の有意水準は片側検定を基にしている。オプションで両側検定によって計算された有意水準を要求することもできる。

9.1 一般書式

NONPAR CORR 変数リスト WITH 変数リスト [/ 変数リスト……]

一般書式は上記のとおりで、NONPAR CORR の後に順位相関を求める変数のリストを続ける。いくつかの変数リストをスラッシュで区切って複数の順位相関係数を求めることもできるし、キーワード WITH を用いて、ある変数セットとある変数セットの順位相関係数を求めることもできる。変数リストの指定の仕方により、下三角行列、あるいは正方行列のいずれかをプリントすることができる。いずれのプリントの指定においても連続する変数にはキーワード T0 を用いることができる。変数は数値型でなくてはならない。変数リストで単に変数を並べただけの場合にはリスト上の各変数と他の全ての変数との順位相関係数を下三角行列の形で印刷する。変数それ自身との相関（対角成分）、及び重複する上三角行列の要素はプリントされない。

9.2 オプション

- 1 ユーザーの指定した欠損値を計算に含める。
- 2 欠損値を表単位に削除する。

- 3 両側検定による有意水準を計算する。順位相関係数の後に示されるアスタリスク“*”が有意水準を示す。“*”では0.01%以下の水準を、“**”では0.001%水準以下の有意性を示す。
- 4 行列をファイルに書き出す。kendallカスピアマンの順位相関係数、あるいはその両方と各順位相関係数を計算するのに用いられたケース数を書き出す。変数リストは単に変数を並べたものであること(WITHは使えない)。オプション4を指定した場合にはFILE HANDLE コマンドとPROCEDURE OUTPUT コマンドによって出力するファイルを指定しなければならない。
- 5 kendallの τ_b 。kendallの順位相関のみが求められる。
- 6 kendallとスピアマンの順位相関係数。両順位相関係数が求められる
- 7 無作為抽出。もし、ケースを保存するスペースが無い場合、無作為にケースを出力する。
- 8 ケース数と有意水準の印刷を省略する。
- 9 連続フォーマットを指定する。行列のはじめの行から順に連続配列の形式で相関係数をプリントしていく。1ページに6列分の順位相関係数、計算された変数名がプリントされる。

9.3 制限事項

- 1 変数リストの数は25まで
- 2 一度のNONPAR CORRで処理できる変数の数は100まで

9.4 コマンド編成の例

〔例1〕(下三角行列の形で出力する場合)

```
NONPAR CORR    PRESTIGE SPPRES PAPRES16 DEGREE RADEG MADEG
```

正方向行列を出力するには2つの変数リストをキーワード WITH で区切らなくてはならない。その場合、はじめの変数リストと2番目の変数リストの相関行列が正方向行列の形で出力される。

〔例2〕(正方向行列の形で出力する場合)

```
NONPAR CORR    PRESTIGE SPPRES PAPRES16 WITH DEGREE RADEG MADEG
```

この例では9つの順位相関係数が出力される。WITHの前にリストされた変数は列に、後の変数は行に表示される。両変数リストに同じ変数が現われない限り、相関が1になるものは現われない。もし、他の分析手続きやプログラムで用いるために相関行列を書き出す場合にはWITHを用いてはならない。

10 RELIABILITY：信頼性解析

RELIABILITYは一般に利用される信頼性係数を計算し、加算尺度の成分に関して項目分析を行うものである。項目平均、標準偏差、項目間共分散行列と項目間相関行列(inter-item correlation matrices)、尺度平均、項目間相関(item-to-item correlations)、などを含めた基本統計量の概要が印刷される。繰り返し要因がある分散分析(repeated measures design analysis of variance)、各セルに観測値が一つの場合の二元配置分散分析(two-way factorial analysis of variance wi

th one observation per cell)、加算性に関する Tukey の検定、繰り返しのある実験計画に於いて平均の差を検定する Hotelling の T2 検定、順位データに対する Friedman の二元配置分散分析を行うこともできる。データをケースとしても、相関行列、共分散行列の形としてでも入力できる。下三角行列やベクトルとして入力された行列を含め、様々な形式で行列を読み込むことができる。更に多様な形式で行列を書き出すこともできる。

10. 1 一般書式

RELIABILITY VARIABLESサブコマンド/SCALEサブコマンド

[/FORMATサブコマンド/MODELサブコマンド] [...]

上記のように、コマンドRELIABILITY の後にサブコマンドを並べる。VARIABLESとSCALEサブコマンドは必須である。スラッシュで区切って複数の分析を一度に行なうこともできる。

10. 2 サブコマンド

(1) VARIABLESサブコマンド

[書式] VARIABLES=変数リスト

分析に使用する全変数を上げる。

(2) SCALEサブコマンド

[書式] SCALE (尺度名) =変数リスト

SCALE サブコマンドはテストされる尺度を指定するためのものである。SCALE サブコマンドでは括弧内に任意の尺度名、その後に尺度を構成する変数のセットを書く。尺度名は表示上で分析を明示するためのものである。尺度名は最大8文字でAからZまでの文字と0から9までの数字から成っていないてはならない。SCALE サブコマンドで指定する変数は全て前もって VARIABLES サブコマンドで指定されている必要がある。SCALE サブコマンドでは隣合う変数のセットに対してキーワード T0 を用いることができる。変数の順序は VARIABLES サブコマンドで並べた変数の順序に従う。VARIABLES サブコマンド上に挙げた変数全てを指定する場合にはキーワード ALL を用いることもできる。

(3) FORMATサブコマンド

[書式] FORMAT={INPUT } ([{ 4 }] / [{ 5 }])
{OUTPUT } { m } { c }

RELIABILITYでは標準の入力書式で行列の読み込み書き込みができる。その場合ケース数と標準偏差は8F10.4、相関行列は8F10.7、平均ベクトルは5D16.9、共分散行列の書式は4D20.13で入力される。行列が自由領域書式の場合にはこれらの書式は適用されない。幾つか他の書式を利用することもできる。その場合には FORMAT サブコマンドを用いる。キーワード INPUT により行列の読み込みに際して別な書式を指定でき、キーワード OUTPUT により行列の書き出しをコントロールできる。キーワード INPUT あるいは OUTPUT の後に続く括弧内は書式の選択番号でスラッシュで区切

る。最初の書式番号は平均についての書式（もし、相関行列が入力、あるいは出力されている場合には標準偏差に当たる）を指示している。2番目の書式番号は相関行列、あるいは共分散行列についての書式を指示している。各書式番号は特定の書式に対応している。5つの書式が利用できる。詳しくは英文マニュアル参照のこと。

(4) MODELサブコマンド

[書式] MODEL= {ALPHA
 {SPLIT [(n)]}
 {GUTTMAN
 {PARALLEL
 {STRICTPARALLEL}}

MODEL サブコマンドにより信頼性分析のタイプを指定する。キーワード MODELにつづいて以下に挙げるキーワードの内の一つを指定する。

ALPHA Cronbach の α と標準化項目 α (standardized item α)。

SPLIT(n) 折半信頼性係数 (Split-half coefficients)。n を指定せずに SPLIT を指定すると SCALE サブコマンドで挙げてある順序で項目は半分に分けられる。はじめに挙げてある $n/2$ 個の項目と後は残りの項目である。もし尺度が奇数の数の変数から成っている場合にははじめの部分が多い項目数となる。GUTTMAN 真の信頼性係数に対するガットマンの下限推定値 (Guttman's lower bounds for true reliability)。

PARALLEL 平行性の仮定の基での信頼性係数の最尤推定値。

STRICTPARALLEL 強平行性の仮定の基での信頼性係数の最尤推定値。

10. 3 オプション

- 1 欠損値を計算に含める。
- 3 変数ラベルの省略
- 4 正方共分散行列の読み込み。
- 5 正方相関行列とその後にくる標準偏差ベクトルの読み込み。
- 6 三角相関行列あるいは三角共分散行列の読み込み。
- 7 行列の前に平均ベクトルがくる場合の読み込み。
- 8 三角共分散行列の書き出し。
- 9 VARIABLES サブコマンドで指定された変数順序ではなく SPSS*変数定義情報 (SPSS* variable dictionary) を正方行列の情報指標 (index) とする。このオプションはデータが三角行列や単一ベクトルとして読み込まれる場合には用いることができない。
- 10 行列が書き出された後、計算を行わない。全ての尺度は無視 (bypassed) される。
- 11 平均の書き出し。
- 12 単一ベクトルとしてフォーマットされた三角行列の読み込み。

13 単一ベクトルとしてフォーマットされた共分散行列の書き出し。

14 解法 2 を選択する。計算には共分散行列を用いる。以下のいずれかを望む場合に指定することになる。

1) 尺度上に分散が 0 の項目が生じた場合それを除く。

この処理は α 係数、Tukey の検定 (STATISTICS 11)、分散分析 (STATISTICS 11) の結果に影響を与える。

2) 重相関係数の 2 乗。

これらは解法 2 が用いられた時には全項目統計量 (STATISTICS 9) に含まれるが、解法 1 が用いられた時には含まれない。

3) 標準化項目 α (Standardized item α)。

15 Friedman の χ^2 統計量と Kendall の一致係数。

16 Cochran の Q

17 行列の後にケース数がかかる場合。

10. 4 追加統計

STATISTICS コマンドにより以下の追加統計を印刷できる。

1 項目平均と標準偏差

2 項目間分散共分散行列

3 項目間相関

4 尺度平均と分散

5 項目平均に関する統計量の概要。項目平均の平均 (the average item mean over the number of items)、項目平均の分散、項目平均の最大値と最小値、項目平均の範囲 (レンジ)、項目平均の最小値に対する最大値の比 (最大値/最小値)。

6 項目分散に関する統計量の概要。この統計量の出力は STATISTICS 5 によるものと同一の形式になるが、項目平均ではなく項目分散を基に計算が行われる。

7 項目間共分散に関する統計量の概要。この出力形式は STATISTICS 5、STATISTICS 6 のものと同じであるが共分散について計算が行われる。

8 項目間相関に関する統計量の概要。この出力形式は STATISTICS 5、6、7 に同じであるが計算は相関を基に行われる。

9 項目に関する全統計量 (item total statistics)。この中には各項目と項目セットの間の関係を扱っている 5 つの統計量が含まれる。各項目に関して以下のものが計算される。

1) 項目が除かれたときの尺度平均。これは尺度から特定の項目が除かれたときの尺度得点の平均である。

2) 項目が除かれたときの尺度分散。これは尺度から特定の項目が除かれたときの尺度得点の分

散である。

- 3) 修正化項目尺度間相関(Corrected item total correlations)。これは特定の項目が尺度から除かれたときの、その項目と(それ以外の項目セットからなる)尺度得点の間の相関係数である。
- 4) 重相関係数の2乗。各項目について尺度を構成する残りの項目セットにより回帰させ、重相関係数の2乗を計算する(これは解法1が用いられた場合には計算されない)。
- 5) 項目が除かれた場合の α 係数。各項目について尺度を構成する他の項目だけによる Cronbach の α が計算される。

10 分散分析表

11 Tukey の加算性に関する検定

12 Hotelling の T^2

10. 5 制限事項

- 1 VARIABLES サブコマンドの最大数は10まで。
- 2 SCALE サブコマンドの最大数は50まで。
- 3 VARIABLES サブコマンドで挙げられている変数を全て合計した数の最大数は500まで。同じ変数であっても繰り返し挙げられてあれば数に入り、500までの制限に従う。
- 4 SCALE サブコマンドで指定できる変数の最大数は500まで。
- 5 SCALE サブコマンドで挙げられている変数を全て合計した数の最大数は1000まで。同じ変数であっても繰り返し挙げられてあれば数に入り、1000までの制限に従う。
- 6 複数の VARIABLES サブコマンドを処理するに十分な作業容量が無い場合には、割り当てられた作業容量が十分になるまで後ろから VARIABLES サブコマンドが除かれることになる。

10. 6 コマンドの編成例

[例1] RELIABILITY FORMAT = OUTPUT (1/2) /

VARIABLES = Q1 TO Q25 /

SCALE (TESTS) = ALL

OPTIONS 8 10 11

最初の書式番号は平均についての書式(もし、相関行列が入力、あるいは出力されている場合には標準偏差に当たる)を指示している。2番目の書式番号は相関行列、あるいは共分散行列についての書式を指示している。各書式番号は特定の書式に対応している。5つの書式が利用できる。もし、標準の書式の一つだけを修正したいとか、標準の書式の両方は使わないという場合には対応する番号を省略できる。但し、スラッシュは必要である。

[例2]

RELIABILITY VARIABLES = VAR1 TO VAR10 /

```
SCALE (VERBAL) = VAR1 VAR3 VAR5 VAR7 VAR9 /
```

```
SCALE (MATH) = VAR2 VAR4 VAR6 VAR8 VAR10
```

この例では VARIABLES サブコマンドで計算される行列を使って2つの分析を指定している。欠損データがある場合には、このように分析を指定することは2つの VARIABLES サブコマンドを用いた場合と異なる結果になることがある。

[例3]

```
RELIABILITY VARIABLES = VAR1 VAR3 VAR5 VAR7 VAR9 /
```

```
SCALE (VERBAL) = VAR1 TO VAR9 /
```

```
VARIABLES = VAR2 VAR4 VAR6 VAR8 VAR10 /
```

```
SCALE (MATH) = VAR2 TO VAR10
```

この場合では、はじめの VARIABLES サブコマンドで VAR1, VAR3, VAR5, VAR7, VAR9 に関してのみ欠損値が表単位に除かれる。作られる共分散行列はその他の変数の欠損値に影響されない。第2の VARIABLES サブコマンドは VAR2, VAR4, VAR6, VAR8, VAR10 のみを用いた新たな行列を作り出す。

[例4]

```
RELIABILITY VARIABLES = VAR1 VAR2 VAR3 VAR4 VAR5 /
```

```
SCALE (RATING) = VAR1 TO VAR5 /
```

```
MODEL = SPLIT
```

この場合、折半法による信頼性係数分析が行われる。

[例5]

```
FILE HANDLE TESTSCOR / file specifications
```

```
GET FILE = TESTSCOR
```

```
FILE HANDLE RELMAT / file specifications
```

```
PROCEDURE OUTPUT OUTFILE = RELMAT
```

```
RELIABILITY VARIABLES = Q1 TO Q25 /
```

```
SCALE (TESTS) = ALL
```

```
OPTIONS 8 10 11
```

この例は RELMAT というファイルに共分散行列と平均を書き出す場合を示している。

11 PEARSON CORR: ピアソン相関係数

PEARSON CORRは、ピアソン積率相関係数をその有意水準とともに算出する。また、オプションとして、単変量の統計や共分散なども算出する。

11. 1 一般書式

```
PEARSON CORR 変数リスト [WITH 変数リスト] /変数リスト...
```

PEARSON CORRの実行のためには、変数リストを指定する必要がある。変数リストの指定の方法には、以下の2種類がある。第1の方法は単一リストによる指示で、リスト上にある全変数相互間の相関行列が算出される。第2の方法は、WITHを用いるリストによる指示で、キーワードWITHで結びつけられた前のリストと後のリストの変数がひとつずつ組み合わせられる。他の手続やプログラムで相関行列を書き出す際には、WITHによる書式を用いることはできない。なお、ひとつの指示リストが終ると、その後にスラッシュ [/] をおいて次のリストを必要なだけ続けることができる。

```
[例] PEARSON CORR X1 TO X10/  
      A1 TO A5 WITH A6 TO A9
```

11. 2 オプション

OPTIONS コマンドにより、欠損値の処理、出力形式、有意水準の検定法などを、指定できる。このサブプログラムの標準の欠損値処理方式は、ペア単位の除去である。この方法では、ペアを構成するどちらかに欠損値のあるケースは、そのペアに限り計算から排除される。もし次のオプション1または2を選択しないと、この方式が採用される。以下オプション番号毎に簡単に説明する。

- 1 欠損値定義を無視して、欠損値としてあらかじめ定義している値も含めて計算を行う。
- 2 リスト単位で欠損値を除去して計算する。この方式によると、変数リストにある変数のどれかひとつにでも欠損値がみつかり、そのケースはそのリストに基づく全計算から排除される。
- 3 有意性の両側検定を行う。デフォルトでは、片側検定を行う。
- 4 計算に用いたケース数とともに、相関行列をファイルに書き出す。その際、PROCEDURE OUTPUT コマンドを用いる必要がある。
- 5 ケース数と有意水準の印刷を省略する。
除される。
- 6 相関行列の右上半分のみを追込み式で印刷する。標準処理では、行列形式で印刷する。
- 7 計算に用いたケース数を書き出さず、相関行列のみをファイルに書き出す。その際、PROCEDURE OUTPUT コマンドを用いる必要がある。

11. 3 追加統計

相関係数、ケース数、有意水準は自動的に出力されるが、このほか STATISTICS コマンドを通して、次の統計を追加出力できる。

- 1 各変数の平均、標準偏差、欠損値を除いたケース数。
- 2 相関係数を求める変数のペア毎の、偏差の積和と共分散。

11. 4 制限事項

- 1 指定できる変数リストの最大数は40である。
- 2 変数の総数は500までである。

3 要素(individual element)の数は250までである。

12 PARTIAL CORR: 偏相関係数

PARTIAL CORRは、種々のオーダーの偏相関係数を計算する。プログラムへの入力は素データでもよいし、相関行列でもよい。この相関行列は、PEARSON CORR, REGRESSION, DISCRIMINANT, FACTORの出力でもよい。

12. 1 一般書式

PARTIAL CORR 変数リスト [WITH 変数リスト]
BY コントロールリスト (オーダー値) [/...]

変数リストを指定する必要がある。変数リストは、次の3部分からなる。第1は相関リストで、ひとつまたは複数の変数のペアを指定する。第2は相関リストの変数を統制するための変数を指示するコントロールリストである。最後の第3は、相関リストとコントロールリストから生じる偏相関の次数を示すオーダー値である。これらの変数リストの指定がすむと、その後にスラッシュ [/] をおいて次のリストを続けることができる。

[例] PARTIAL CORR A1 TO A10 BY X1(1)/
B1 TO B10 BY X2, X3(2)

(1) 相関リスト

相関リストは、コントロールリストに現れる変数によって統制されるゼロオーダーの単純相関のための変数ペアを指定するリストである。この単純相関の計算は、PEARSON CORRと同じ作業であって、このリストの機能と書式はPEARSON CORRのそれと等しい。すなわち、その書式には、単一リストによる形式とWITHをもちいるリストによる形式とがある。

(2) コントロールリスト

相関リストの最後の変数名の直後に、キーワードBYをおき、BYに続けてコントロールリストを指定する。このリストは、相関リスト上の変数が組むペアを統制する変数の単なるリストにすぎず、統制変数がいかに変数ペアを統制するかは、次のオーダー値によって決まる。

(3) オーダー値

オーダー値は正の整数値で、コントロールリストにおける最後の変数の後に、かっこで包んでおく。この数値は、偏相関のオーダーを示す。

[例] PARTIAL CORR A1 WITH A2 BY X1, X2(1)

この例では、X1とX2をそれぞれ単独で統制したA1とA2の間の2個の第1オーダー偏相関係数が算出される。

[例] PARTIAL CORR A1 WITH A2 BY X1, X2(2)

この例では、X1とX2を同時に統制したA1とA2の間の1個の第2オーダー偏相関係数が算出される。

[例] PARTIAL CORR A1 WITH A2 BY X1,X2(1,2)

この例では、上記の2種類の偏相関係数がともに算出される。

前述X1とX2をそれぞれ単独で統制したA1とA2の間の2個の第1オーダー偏相関係数が算出される。

12. 2 オプション

OPTIONSコマンドにより、欠損値の処理、出力形式、有意水準の検定法などを、指定できる。欠損値の標準処理方式は、リスト単位の除去である。もし次のオプション1または2を選択しないと、この方式が採用される。以下オプション番号毎に簡単に説明する。

- 1 欠損値定義を無視して、欠損値としてあらかじめ定義している値も含めて計算を行う。
- 2 ペア単位で欠損値を除去して計算する。この方式によると、ペアを構成するどちらかに欠損値のあるケースは、そのペアに限り計算から排除される。
- 3 有意性の両側検定を行なう。デフォルトでは、片側検定を行なう。
- 4 相関行列データを入力する。この場合、NUMERICコマンドで変数を定義し、FILE HANDLEやINPUT MATRIXコマンドでファイルの定義をする必要がある。
- 5 リスト上の全変数間の相関行列とケース数のファイルへの出力を行なう。この場合、FILE HANDLEやPROCEDURE OUTPUTコマンドでファイルの定義をする必要がある。
- 6 相関行列にPARTIAL CORRで用いるより多くの変数が含まれている場合、オプション4に加えてこのオプションを用い、行列を指示する。
- 7 自由度と有意水準を省略して印刷する。
- 8 相関行列の右上半分のみを追込み式で印刷する。デフォルトでは、行列形式で印刷する。

12. 3 追加統計

偏相関係数、ケース数、有意水準は自動的に出力されるが、このほかSTATISTICSコマンドを通して、次の統計を追加出力できる。

- 1 偏相関の計算に用いた単純相関係数とその自由度および有意水準
- 2 各変数の平均、標準偏差、欠損値を除いたケース数。
- 3 単純相関を計算できない変数の組合せがひとつでもあらわれると、ペリオド(.)を与え、相関係数を出力する。

12. 4 制限事項

- 1 指定できる変数リストの最大数は25である。
- 2 変数の総数は400までである。
- 3 統制変数の最大数は100までである。
- 4 各リストの偏相関のオーダーは5種類までである。オーダー値の最大は100までである。

13 REGRESSION：重回帰分析

REGRESSIONは、ひとつの従属変数と複数の独立変数との関係を分析する重回帰分析を行い、それに関連する統計やプロットを算出する。

13. 1 一般書式

```
REGRESSION [READサブコマンド][WIDTHサブコマンド]
           [/SELECTサブコマンド][MISSINGサブコマンド]
           [/DESCRIPTIVESサブコマンド][WRITEサブコマンド]
           /VARIABLESサブコマンド
           [{NOORIGIN, ORIGINサブコマンド}][METHODサブコマンド]
           [/RESIDUALSサブコマンド][CASEWISEサブコマンド]
           [/SCATTERPLOTサブコマンド][PARTIALPLOTサブコマンド]
           [/SAVEサブコマンド]
```

13. 2 サブコマンド

13. 2.1 必要最小限のサブコマンド

(1) VARIABLESサブコマンド

[書式] VARIABLES={変数リスト}
 {(COLLECT)}
 {ALL}

分析で用いる変数を指定する。

ALL 実行ファイルにあるすべての変数を計算の対象に含める。

(COLLECT) DEPENDENTサブコマンドとMETHODサブコマンドで指定したすべての変数を計算の対象に含める。

[例] REGRESSION VARIABLES=(COLLECT)/
 DEPENDENT=SAVINGS/
 METHOD=ENTER POP15 POP75 INCOME GROWTH/

(2) DEPENDENTサブコマンド

[書式] DEPENDENT=変数リスト

従属変数を指定する。ひとつのREGRESSIONコマンドに、複数のDEPENDENTサブコマンドを与えてもよい。

[例] REGRESSION VARIABLES=X1 TO X5 Y1 TO Y3/
 DEPENDENT=Y1 TO Y2/
 STEPWISE/
 DEPENDENT=Y3/ENTER

(3) METHODサブコマンド

[書式]

```
[METHOD=] {STEPWISE [変数リスト] } [...] [/...]  
           {FORWARD  [変数リスト] }  
           {BACKWARD [変数リスト] }  
           {ENTER     [変数リスト] }  
           {REMOVE    [変数リスト] }  
           {TEST(変数リスト)(変数リスト)...}
```

変数選択の方法を指定する。以下の6種類の選択法を指定できるが、METHOD=というキーワード自体は省略してもよい。また、METHODサブコマンドは、ひとつの回帰方程式に複数与えることもできる。なお、回帰方程式作成のためにREMOVEまたはTESTを指定した場合、あるいはVARIABLES=(COLLECT)を指定した場合は、METHODサブコマンド上に必ず変数リストを与えなければならない。

FORWARD 前進的投入法。有意確率水準に達するまで回帰方程式に変数をひとつずつ加えてゆく方法である。

BACKWARD 後退的除去法。独立変数全部を含めた回帰分析を行ない、それから変数を有意確率水準に達するまでひとつずつ除去してゆく方法である。

STEPWISE ステップワイズ選択法。有意確率水準に達するまで後退的除去と前進的投入とを交互に逐次的に繰返す方法である。

ENTER 強制的投入法。減少トレランスの順に変数をひとつずつ加えてゆく方法である。ENTERに変数リストを添えない場合、トレランス基準に合格したすべての変数が投入される。

REMOVE 強制的除去。指定した変数を回帰方程式から除去する。

TEST 独立変数サブセットの検定。モデルから指定した独立変数のサブセットを除去したことによる R^2 の変化値の有意性検定を行なう。指定の際、各サブセットを括弧でくくっておく必要がある。

```
[例] REGRESSION VARIABLES=X1 TO X5 Y/  
      DEPENDENT=Y/  
      STEPWISE/ENTER
```

13. 2.2 方程式のコントロールのためのサブコマンド

VARIABLESサブコマンドとDEPENDENTサブコマンドとの間に以下の3つのサブコマンドを付加的に与えることができる。

(1) CRITERIAサブコマンド

```
[書式] CRITERIA=[DEFAULTS]  
        [TOLERANCE({0.01})]  
        {値 }  
        [MAXSTEPS(n)]  
        [PIN({0.05})][POUT({0.10})]  
        {値 } {値 }
```

[FIN({3.84})][FOUT({2.71})]
 {値} {値}

回帰方程式作成のための統計的基準をコントロールするサブコマンドである。CRITERIAサブコマンドを指定しない場合、DEFAULTSを指定したのと同じになり、標準値が採用される。

PIN F-to-enterのための有意水準。その変数が新たに追加されることによって寄与率が有意に改良されるかどうかの検定統計量Fの有意水準を指定する。標準値は0.05である。

POUT F-to-removeのための有意水準。その変数の寄与率がある新しい変数の追加後に下がった場合、その変数が構成から除外されるための検定統計量Fの有意水準を指定する。標準値は0.10である。

FIN F-to-enterの合格水準で、Fの絶対値を括弧内に指定する。標準値は3.84である。

FOUT F-to-removeの合格水準で、Fの絶対値を括弧内に指定する。標準値は2.71である。

TOLERANCE トレランス。標準値は、0.01である。

MAXSTEPS 反復計算の打ち切り上限回数。標準値は、変数選択法がSTEPWISEの際が独立変数の数の2倍で、FORWARDまたはBACKWARDの際がPIN, POUT, またはFIN, FOUTの基準に合格した変数の数である。

[例] REGRESSION VAR=X1 TO X5 Y/
 CRITERIA=PIN(.19) TOL(.0001)/ DEP=Y/ FORWARD/
 CRITERIA=DEFAULTS/ DEP=Y/ FORWARD

(2) STATISTICSサブコマンド

[書式] STATISTICS=[DEFAULTS*] [R] [COEF] [ANOVA]
 [OUTS] [ZPP] [LABEL] [CHA] [CI] [F]
 [BCOV] [SES] [LINE] [HISTORY] [XTX]
 [COND] [END] [TOL] [ALL]]

回帰方程式や独立変数の統計の結果の印刷出力をコントロールする。このサブコマンドは、DEPENDENTサブコマンドに先行して与えなければならない。STATISTICSのキーワードは、全体的な指定、回帰方程式の概括統計の指定、独立変数の概括統計の指定、ステップの概括統計の指定の4種類に大別できる。

全体的な指定

DEFAULT R, ANOVA, COEF, OUTS, STATISTICSサブコマンドを指定しない場合、標準出力としてこれらの統計が印刷される。

ALL LABEL, F, LINE, ENDを除く他のすべての概括統計を印刷する。

回帰方程式の概括統計の指定。

R 重相関係数 R 、 R^2 、修正された R^2 、決定係数、推定標準誤差を印刷する。

ANOVA 分散分析表、重相関係数 R に対する F 値とその有意水準を印刷する。

CHA ステップ間の R^2 の変化、変化に対する F 値とその有意水準を印刷する。

BCOV 標準化されていない回帰係数に対する分散-共分散行列を印刷する。

XTX 掃き出し行列を印刷する。

COND 掃き出し行列の部分行列の条件数に対する上限と下限を印刷する。

独立変数の概括統計の指定

COEFF 標準化されていない回帰係数 B 、 B の標準誤差、標準化回帰係数 β 、 B の t 値、 t 値の両側検定の有意水準を印刷する。

OUTS 次のステップでその変数が回帰方程式に投入される予定の場合の β 、 B の t 値、 t 値の両側検定の有意水準、方程式のすべての変数を統制した場合の従属変数との偏相関係数を印刷する。

ZPP 各独立変数と従属変数との相関係数、他の独立変数を統制した場合の従属変数との偏相関係数、部分相関係数を印刷する。

CI 標準化されていない回帰係数の 95% 信頼区間を印刷する。

SES β の近似標準誤差を印刷する。

TOL トレランスと最小トレランスを印刷する。

LABEL 変数ラベルを印刷する。

F B の F 値とその有意水準を印刷する。

ステップの概括統計の指定

LINE 実行された各ステップの結果の単一の概括ラインを印刷する。

HISTORY ステップの最終的概括統計を印刷する。

END STEPWISE, FORWARD, BACKWARD については各ステップにひとつのラインを印刷し、ENTER, REMOVE については、各ブロックにひとつのラインを印刷する。

(3) ORIGINサブコマンドとNOORIGINサブコマンド

[書式] {NOORIGIN または ORIGIN}

ORIGINは、定数項を削除して回帰方程式が原点を通るようにするサブコマンドである。NOORIGINは、標準処理であり、定数項を含んだ回帰方程式を作成する。ORIGINサブコマンドまたはNOORIGINサブコマンドは、DEPENDENTサブコマンドとMETHODサブコマンドの前に指定しなければならない。

[例] REGRESSION VARIABLES=GNP TO M1/
ORIGIN/DEPENDENT=GNP/FORWARD

13. 2.3 残差分析のためのサブコマンド

(1) 一時的変数

残差の分析には、後述するサブコマンドのRESIDUALS, CASEWISE, SCATTERPLOT, PARTIALPLOT, SAVEを使用する。各分析において、REGRESSIONは、以下の12個の一時的変数を計算できる。

PRED 標準化されていない予測値。
RESID 標準化されていない残差。
DRESID 除去された残差。
ADJPRED 修正された予測値。
ZPRED 標準化された予測値。
ZRESID 標準化された残差。
SRESID スチューデント化された残差。
SDRESID スチューデント化され除去された残差。
SEPREP 予測値の標準誤差。
MAHAL マハラノビスの距離。
COOK クックの距離。
LEVER 影響値(leverage values)。

(2) RESIDUALSサブコマンド

[書式] RESIDUALS=[DEFAULTS][ID(変数名)][DURBIN]

 [{SEPARATE}]
 [POOLED]

 [HISTGRAM({ZRESID
 {一時的変数リスト} })]

 [OUTLIERS({ZRESID
 {一時的変数リスト} })]

 [NORMPROB({ZRESID
 {一時的変数リスト} })]

 [SIZE({LARGE})]
 [{SMALL}]

以下のキーワードにより、外れ値の情報、一時的変数のヒストグラムや正規確率プロットなどを得ることができる。

DEFAULTS SIZE(LARGE), DURBIN, NORMPROB(ZRESID), HISTOGRAM(ZRESID), OUTLIERS(ZRESID)

SIZE(プロットサイズ) プロットのサイズ。SMALLまたはLARGEで与えることができる。標準値は、LARGEである。

HISTGRAM(変数リスト) 一時的変数や命名された変数のヒストグラム。無指定時の変数は、ZRESIDである。

NORMPROB(変数リスト) 標準化変数の正規確率プロット。無指定時の変数は、ZRESIDである。

OUTLIER(変数リスト) 指定した変数の値に基づく10個の外れ値。無指定時の変数は、ZRESIDである。

DURBIN ダービン-ワトソンの統計量。デフォルト。

ID(変数名) ケース単位のプロットまたは外れ値のプロットのラベリングのために指定した変数の値を使用する。

POOLED プールされたプロットと選択されたケースと選択されなかったケースに対する統計を表示する。

(3) CASEWISEサブコマンド

```
[書式] CASEWISE=[DEFAULTS][{OUTLIERS({3})}]]  
                                     {値}  
                                     {ALL }  
                                     [PLOT({ZRESID  
                                     {一時的変数名}})]  
                                     [{DEPENDENT PRED RESID}]  
                                     {一時的変数リスト}]
```

ケース単位のプロットのための変数を同定し、ケース単位の印刷出力のための変数を指定し、ケースの選択をコントロールする。

DEFAULTS OUTLIER(3), PLOT(ZRESID), DEPENDENT, PRED, PRESIDを指定したのと同じ。

OUTLIER(値) 括弧内に値を指定することによって外れ値のプロットを制限する。デフォルト値は、3である。

ALL ケース単位のプロットにおいてすべてのケースを含める。

PLOT(変数名) ケース単位のプロットにおいて、括弧内に指定した一時的変数の値をプロットする。無指定時の変数は、ZRESIDである。

変数リスト 指定した変数の値を表示する。無指定時の変数は、従属変数、PRED, RESIDである。

(4) WIDTHサブコマンド

```
[書式] WIDTH={132}  
          {n }
```

REGRESSIONの結果の出力幅をコントロールする。WIDTH=の後に、72から132までの値を指定する。標準値は132である。

(5) SCATTERPLOTサブコマンド

```
[書式] SCATTERPLOT=[SIZE({SMALL})](  
                                     {変数名, 変数名}...  
                                     {LARGE})
```

散布図のための変数のペアを指定し、そのプロットの大きさをコントロールする。

SIZE(プロットサイズ) プロット大きさを、SMALLまたはLARGEで指定する。標準値はSMALLである。

(変数名、変数名) プロットする変数を指定する。第1の変数が垂直軸、第2の変数が水平

軸になる。なお、一時的変数を指定する場合には、その変数の前に*印を付ける。

[例] SCATTER PLOT (*RES,*PRE)(*RES,Y)/

(6) PARTIALPLOTサブコマンド

[書式] PARTIALPLOT={ALL
 {変数リスト}}[SIZE({SMALL})]
 {LARGE}]

偏回帰プロットを行ない、そのプロットの大きさをコントロールするためのサブコマンドである。偏回帰プロットは、従属変数の残差と指定した独立変数の残差との散布図である。このプロットを行なうには、少なくとも2つの独立変数が回帰方程式中になければならない。

SIZE(プロットサイズ) プロット大きさを、SMALLまたはLARGEで指定する。標準値はSMALLである。

変数名リスト プロットする変数リストを指定する。標準値はALLで、すべての独立変数がプロットされる。

(7) SAVEサブコマンド

[書式] SAVE=一時的変数名(新しい名) 一時的変数名(新しい名)...

1 2 個の一時的変数(本項(1)一時的変数を参照)を新しい変数名で保存する。保存された変数は次例のように、後続のステップで使用できる。

[例] REGRESSION VARIABLES=Y TO X10/

DEPENDENT=Y/ENTER/

SAVE PRED(PREDSR)/

PLOT PLOT PREDSR WITH Y

13. 2.4 欠損値処理およびケース選択のためのサブコマンド

(1) MISSINGサブコマンド

[書式] MISSING={LISTWISE
 {PAIRWISE
 {MEANSUBSTITUTION}}}[INCLUDE]]

欠損値の扱いについての指示を与える。

LISTWISE リスト単位で欠損値を除去して計算する(標準処理)。

PAIRWISE ペア単位で欠損値を除去して計算する。

MEANSUBSTITUTION 欠損値をその変数の平均値で置き換える。

INCLUDE 欠損値を含めて計算する。

(2) SELECTサブコマンド

[書式] SELECT={({ALL})
 {変数名 関係演算子 値}}]

回帰方程式の算出に使用する変数を選択する。なお、関係演算子は、EQ, NE, LT, LE, GT, GEである。

13. 2.5 記述統計出力のサブコマンド

(1) DESCRIPTIVEサブコマンド

[書式] DESCRIPTIVES=[DEFAULTS] [MEAN] [STDDEV] [CORR]
[COV] [VARIANCE] [XPROD] [SIG] [N] [BADCORR]
[ALL] [NONE]

通常REGRESSIONは、記述統計を印刷出力しない。DESCRIPTIVESサブコマンドを用いると、種々の記述統計を印刷出力できる。

NONE 記述統計を印刷出力しない。DESCRIPTIVESサブコマンドを省略した場合もこの扱いとなる。

DEFAULTS MEAN, STDDEV, CORRを指定したのと同じになる。サブコマンド名の場合のみは、この扱いとなる。

ALL MEAN, STDDEV, VARIANCE, CORRを指定したのと同じになる。

MEAN 変数の平均値。

STDDEV 変数の標準偏差値。

VARIANCE 変数の分散。

COV 共分散行列。

XPROD 平均偏差の積和。

CORR 相関行列。

SIG 相関係数の片側検定。

BADCORR 相関係数が計算不能になった場合、相関行列を出力する。

N 相関係数の計算に用いられたケース数。

13. 2.6 行列の入出力のためのサブコマンド

REGRESSIONでは、行列の入出力を行なうことができる。たとえば、PEARSON CORRのようなサブプログラムによって出力された結果を入力して計算することができる。

(1) READサブコマンド

[書式] READ=[DEFAULTS] [MEAN] [STDDEV] [CORR] [N]
[VARIANCE] [COV] [INDEX]]

相関行列または共分散行列を入力して、分析を行なう。

DEFAULTS MEAN, STDDEV, CORR, N指定に相当。サブコマンド名READの場合には、この扱いとなる。

MEAN 変数の平均値を行列に先行して入力する。

STDDEV 変数の標準偏差値を行列に先行して入力する。

VARIANCE 変数の分散を行列に先行して入力する。

CORR 相関行列を入力する。
 COV 共分散行列を入力する。
 N 相関係数の計算に用いられたケース数が行列の後に来る。
 INDEX DATA LISTコマンドの変数の数と順序によって行列に番号を付ける。

なお、行列入力の際は、NUMERICコマンドで変数を定義し、INPUT MATRIXコマンドでその行列を含んだファイルを定義する必要がある。

(2) WRITEサブコマンド

[書式] WRITE=[DEFAULTS] [MEAN] [STDDEV] [CORR]
 [VARIANCE] [COV] [N] [NONE]]

相関行列や共分散行列などを外部ファイルに出力する。キーワードのうちDEFAULTSからNまでは、READサブコマンドのものと同様である。NONEは、それまでのWRITEサブコマンドによる指定を解除する。なお、行列の外部出力の際、FILE HANDLEコマンドやPROCEDURE OUTPUTコマンドを用いて、その行列を含んだファイルを定義する必要がある。

14 DISCRIMINANT：判別分析

DISCRIMINANTは、一般にグループ数が2つ以上の場合について線形判別関数を求める。また、求めた判別関数を用いて、ケースを分類して判別の成功度を確かめることもできる。

14. 1 一般書式

DISCRIMINANT GROUPSサブコマンド /VARIABLESサブコマンド
 [/SELECTサブコマンド][ANALYSISサブコマンド]
 [/METHODサブコマンド][TORERANCEサブコマンド]
 [/MAXSTEPSサブコマンド]
 [/FINサブコマンド][FOUTサブコマンド]
 [/PINサブコマンド][POUTサブコマンド]
 [/FUNCTIONSサブコマンド][PRIORSサブコマンド]
 [/SAVEサブコマンド][ANALYSISサブコマンド]

GROUPSサブコマンドとVARIABLESサブコマンドは、不可欠なサブコマンドである。

14. 2 サブコマンド

(1) GROUPSサブコマンド

[書式] GROUPS=変数名(最小値、最大値)

判別の対象グループを指定する。指定の際、特定の変数(グループ変数)の変数値の種類によってグループを区別し、変数値の最小値と最大値をその順にかっこで与える。

[例] DIDCRIMINANT GROUPS=WORLD(1,3)/

VARIABLES=VAR1 TO VAR9

本サブコマンドでの変数の指定は、1変数に限られる。また、グループ変数の変数値は整数でなければならない。変数値の最小値から最大値までのすべての整数がそれぞれグループとみなされるが、もしも該当するケースがひとつでもない整数があれば、それは無効とされ、グループの数から除いて計算を実行する。

(2) VARIABLESサブコマンド

[書式] VARIABLES=変数リスト

分析で用いられるすべての変数（グループ変数を除く）を指定する。指定の仕方は、通常の変数リストの指定と同じである。

(3) ANALYSISサブコマンド

[書式] ANALYSIS=変数リスト（水準） [変数リスト. . .]

1回のDISCRIMINANTの手続の中で変数の構成を変えて重複して分析することが可能であるが、具体的にこの判別変数群を指定するのがANALYSISサブコマンドである。DISCRIMINANTでは、段階的に変数を加える変数選択方式と全変数を同時にあつかう直接方式とがあるが、そのいずれをとるかにより変数の表記法が異なる。まず、以下に直接方式のプログラム例を示す。

[例1] DIDCRIMINANT GROUPS=WORLD(1,3)/

VARIABLES=V1 TO V10/

ANALYSIS=V1 TO V5/

ANALYSIS=V4 TO V10/

変数選択方式では、インクルージョンレベルを変数リストの後のかっこの中に指定する。インクルージョンレベルは、その値が大きいほど早い段階で変数が判別変数に加えられることを意味している。また、レベル値が偶数の時は同じレベルの変数が一括して投入されることを表わし、レベル値が奇数の時は同じレベルの変数がMETHODサブコマンドで与える統計的基準に関して最適な変数から順に投入されることを表わしている。レベル値が1以外の変数は、ひとたび投入されたら判別変数から除かれることはないが、レベル値が1の変数は投入された後でも他の変数の追加によって除外されることが有り得る。レベル値は0から99まで許され、省略時の標準値は1である。なお、変数選択方式では、METHODなどのサブコマンドによって具体的な基準を与える必要がある。

[例2] DIDCRIMINANT GROUPS=CAT(1,3)/

VARIABLES=V1 TO V8/

ANALYSIS=V1 TO V3(2)

V4 TO V6(1) V7 V8(3)/

METHOD=WILKS/

(4) METHODサブコマンド

[書式] METHOD={DIRECT
 {WILKS
 {MAHAL
 {MAXMINF
 {MINRESID
 {RAO

変数選択の過程における統計的基準の種類を指示するためのサブコマンドである。このサブコマンドは、ANALYSISサブコマンドの後に続けて与える必要がある。このサブコマンドのキーワードには以下のものがある。

DIRECT 変数選択を行なわない直接方式。デフォルト。

WILKS ウィルクスのラムダ統計量を最小にする変数。

MAHAL 2 グループ間のマハラノビス汎距離の最小値が最大となる変数。

MAXMINF 2 グループ間のマハラノビス汎距離をF値に変換して、その最小F値が最大となる変数。

MINRESID グループ間の非説明変動和を最小にする変数。

RAO ラオのV統計量の増加が最大となる変数。

このサブコマンドで指定する基準で追加変数を選ぶ時は、偏F値の大きさがひとつの基準となる。偏F値は、その変数が新たに追加されることによって判別対象グループの分離が有意に改良されるかどうかの検定統計量である。

(5) MAXSTEPSサブコマンド

[書式] MAXSTEPS={2v
 {m }

変数の追加と除外を反復する場合、同じ変数が交互に追加されたり除外されたりするような事態を防ぐためのサブコマンドで、反復の打ち切り上限回数の指定に使われる。省略時の標準値は、インクルージョンレベルが1の変数の数の2倍とインクルージョンレベルが1以上の変数の数の総和である。

(6) 変数選択に用いられるその他(上記(4),(5)以外)のサブコマンド

変数選択のためのその他のサブコマンドとして以下のものがある。これらのサブコマンドは、METHODサブコマンドの後に続ける必要がある。

TOLERANCEサブコマンド トレランス水準を指定する。TOLERANCE=の後に0から1までの値を与える。省略時の標準値は、0 0 1である。

FINサブコマンド F-to-enterの合格水準。その変数が新たに追加されることによって判別対象グループの分離が有意に改良されるかどうかの検定統計量(偏F値)の合格水準を指定する。省略時の標準値は1. 0であり、FIN=の後にFの絶対値を記入する。

FOUTサブコマンド F-to-removeの合格水準。その変数の偏F値がある新しい変数の追加後に

下がった場合、その変数が構成から除外されるための棄却水準を指定する。省略時の標準値は 1.0 であり、FOUT= の後に F の絶対値を記入する。

PINサブコマンド F-to-enterのための有意確率水準をPIN=の後に指定する。

POUTサブコマンド F-to-removeのための有意確率水準をPOUT=の後に指定する。なお、PINまたはPOUTは、それを指定していれば、FINまたはFOUTの指定に優先して適用される。

VINサブコマンド ラオのV-to-enterの合格水準。METHOD=RAOを指定した時は、偏F値のチェックに加えて、変数の追加に伴うV統計量の変化量の有意性検定が行われ、合格しない変数は候補から除かれる。VIN=の後には、この合格水準に対応する値を指定する。省略時の標準値は、0.0である。

(7) FUNCTIONSサブコマンド

```
[書式] FUNCTIONS={g-1, 100.0, 1.0}
                {nf, cp, sig}
```

一般に判別分析で算出可能な判別関数の数は、グループ数を g 、変数の数を n とすると、 $(g-1)$ と n のいずれか小さい方の数である。DISCRIMINANTでは、この数の範囲内で実際に算出する判別関数の数をFUNCTIONSサブコマンドを用いて選択することができる。このサブコマンドには以下の3つの指定値がある。

nf 算出した判別関数の数の最大値。

cp 固有値の累積の割合の上限値。

sig 判別関数の追加に関する検定の有意水準。

これら3種の指定値のいずれか1種が達成された時、判別関数の追加算出は打切られる。nf, cp, sigの標準値は、それぞれ、(g-1)とnのいずれか小さい方の数、100.0, 1.0であり、FUNCTIONSの指定を省略するとこれらの値が使われる。標準値の1部でも変えたい時は、FUNCTIONSサブコマンドで3種の値をその順序に従って指定しなければならない。

```
[例] DISCRIMINANT GROUPS=CLASS(1,4)/
      VARIABLES=S1 TO S10/
      FUNCTIONS=4,100,.80/
```

(8) SELECTサブコマンド

[書式] SELECT=変数名 (値)

SELECTサブコマンドは、特定のケースを選んで判別関数の諸係数を算出し、その結果を用いて、判別得点の算出、ケースのプロット、ケースの分類を行なえる。このサブコマンドは、ANALYSISサブコマンドに先行しなければならない。SELECTサブコマンドの後にANALYSISサブコマンドを2つ以上重ねると、そのすべてに対して、同じSELECT機能が働く。SELECT=の後に書かれた変数の値が、括弧の中の値に該当するケースだけが、判別関数の諸係数の算出に使われる。

```
[例] DISCRIMINANT GROUPS=TYPE(1,3)/
      VARIABLES=X1 TO X10/
      SELECT=YEAR(88)
```

(9) PRIORSサブコマンド

[書式] PRIORS={
 EQUAL
 SIZE
 確率リスト

本サブコマンドにより母集団構成比にあたる事前確率を指定し、ケースの分類を行うことができる。パラメータには以下の3種類がある。

EQUAL 判別グループの事前確率をすべて等しいとみなす。

SIZE 計算に使われた各グループのケース数の割合をその事前確率とみなす。

確率リスト 事前確率のリストを記入する。確率の合計値は1でなければならない。

```
[例] DISCRIMINANT GROUPS=TYPE(1,5)/
      VARIABLES=X1 TO X10/
      PRIORS=4*.15,.40/
```

(10) SAVEサブコマンド

[書式] SAVE=[CLASS=変数名] [PROBS=ルート名]
 [SCORES=ルート名]

実行ファイルに保存される情報のタイプや各情報に対応する変数名を指定する。次のキーワードを用いることによって、3種類の変数が保存される。

CLASS 予測されたグループ値を含む変数を保存する。

PROBS 各ケースに対するグループ所属の確率を保存する。ルート名をつける必要がある。

SCORES 判別得点を保存する。ルート名をつける必要がある。

```
[例] DISCRIMINANT GROUPS=WORLD(1,3)/
      VARIABLES=X1 TO X10/
      SAVE=CLASS=PRDCLAS
      SCORES=SCORE PROBS=PRB
```

14. 3 オプション

- 1 欠損値定義を無視して、欠損値としてあらかじめ定義している値も含めて計算を行う。
- 2 マトリックス出力。統計7、8またはオプション11を指定すると、グループ別の共分散行列が出力され、それ以外の場合は、グループ内共分散行列が出力される。な、おこの出力を再入力して次の分析を行うこともできる。
- 3 マトリックス入力。この場合、INPUT MATRIXコマンドが必要である。また、行列に含まれる

変数が定義されていない時は、NUMERICコマンドで定義をする必要がある。

- 4 ステップごとの印刷出力を省略する。
- 5 概括表の印刷出力を省略する。
- 6 パターン行列をバリマックス回転する。
- 7 構造行列をバリマックス回転する。
- 8 欠損値を平均値で置き換える。
- 9 SELECTサブコマンドによって選択されなかったケースのみを分類する。
- 10 GROUPSサブコマンドによって指定されなかったケースのみを分類する。
- 11 ケース分類の際、グループごとの共分散行列を使用する。

14. 4 追加統計

DISCRIMINANTでは、通常、判別関数に関して、固有値、分散説明率、累積寄与率、正準相関係数、ウィルクスのラムダ、カイ二乗値、その自由度および有意水準が出力され、変数選択のステップにおいては、ウィルクスのラムダ、F 値、自由度、有意水準、トレランスが出力され、最終的な結果として、判別関数の標準化された判別係数、構造行列などが出力される。なお、STATISTICSコマンドを通して、次の統計を追加出力できる。

- 1 全体およびグループごとの変数別平均値。
- 2 全体およびグループごとの変数別標準偏差値。
- 3 グループ内共分散行列。
- 4 グループ内相関行列。
- 5 2 グループ間のマハラノビス汎距離に関する F 値の行列。
- 6 変数別の F 値。
- 7 グループ別共分散行列の同等性に関するボックスの M 検定。
- 8 グループ別の共分散行列。
- 9 全体の共分散行列。
- 10 判別領域図。
- 11 標準化されていない正準判別関数。
- 12 分類関数の係数。
- 13 ケース分類結果の表。
- 14 ケースごとの分類情報と判別得点。
- 15 全グループのケースのプロット。
- 16 グループ別のケースのプロット。

14. 5 制限事項

- 1 GROUPS, SELECT, VARIABLESサブコマンドは、それぞれ1回しか用いることができない。

- 2 欠損値のペア単位の処理はできない。

15 FACTOR：因子分析

FACTORは、主成分分析ならびに因子分析を行う。FACTORは、4つのブロックを持っている。それらは、変数選択ブロック、抽出ブロック、回転ブロック、SAVEブロックである。

15. 1 一般書式

FACTOR VARIABLESサブコマンド

```
[/MISSINGサブコマンド] [/READサブコマンド]
[/WRITEサブコマンド] [/WIDTHサブコマンド]
[/ANALYSISサブコマンド[/ANALYSISサブコマンド/...]]
[/PRINTサブコマンド] [FORMATサブコマンド]
[/CRITERIAサブコマンド]
[/EXTRACTIONサブコマンド] [/ROTATIONサブコマンド]
[/EXTRACTIONサブコマンド] [/ROTATIONサブコマンド]
[.....]
[/DIAGONALサブコマンド] [/PLOTサブコマンド]
[/SAVEサブコマンド]
```

15. 2 サブコマンド

15. 2.1 変数選択ブロックのサブコマンド

(1) VARIABLESサブコマンド

[書式] VARIABLES=変数リスト

変数選択のサブコマンドは、唯一常に必要なサブコマンドである。このサブコマンドは、分析に用いるためのメインリストを作成するために使用される。たとえば、10個の変数VAR1からVAR10について主成分分析を行なう場合、以下のような指示で簡単に実行できる。

[例1] FACTOR VARIABLES=VAR1 TO VAR10/

変数リストには、その課題で用いる全変数を書く。変数リストはスラッシュ[/]でしめくくる。

(2) MISSINGサブコマンド

[書式] MISSING= {LISTWISE}
 {PAIRWISE}
 {MEANSUB}
 {INCLUDE}
 {DEFAULT }

欠損値の扱いについての指示を与えるサブコマンドであり、計算の第1段階である相関行列の計算にかかわる。このサブコマンドは、同時に2つ以上重ねて与えることができる。

LISTWISE リスト単位で欠損値を除去して計算する。MISSINGサブコマンドを省略した場合、標

準処理として、この処理が行われる。

PAIRWISE ペア単位で欠損値を除去して計算する。

MEANSUB 欠損値をその変数の平均値で置き換えて計算する。

INCLUDE 欠損値を含めて計算する。

DEFAULT LISTWISEと同じ。

(3) WIDTHサブコマンド

[書式] WIDTH={132}
 {n}

印刷出力結果のページ幅を72桁から132桁の範囲で指定する。なお、標準値は、最大の132である。

15. 2.2 抽出ブロックのサブコマンド

(1) ANALYSISサブコマンド

[書式] ANALYSIS=変数リスト

VARIABLESサブコマンドで指定した変数リストの中から特定の変数を計算用として指定する。このサブコマンドは、同時に2つ以上重ねて与えることができる。

[例2] FACTOR VARIABLES=VAR1 TO VAR10/

ANALYSIS=VAR1 TO VAR10/

ANALYSIS=VAR1 TO VAR8

(2) EXTRACTIONサブコマンド

[書式] EXTRACTION={PC
 {PAF
 {ALPHA
 {IMAGE
 {ULS
 {GLS
 {ML
 {PA1
 {PA2
 {DEFAULT }

次のキーワードにより、因子の抽出法を指定する。このサブコマンドは、同時に2つ以上重ねて与えることができる。

PC 主成分分析（本サブコマンド省略時の標準値）。

PA1 主成分分析。PCと同じである。

PAF 主因子法による因子分析。

PA2 主因子法による因子分析。PAFと同じである。

ALPHA アルファ因子分析。

IMAGE イメージ因子分析。

ULS 非加重式最小二乗法による因子分析。

GLS 一般化最小二乗法による因子分析

ML 最尤法による因子分析。

DEFAULT PCを指定したのと同じ。

[例3] FACTOR VARIABLES=VAR1 TO VAR10/

 EXTRACTION=ULS/

 EXTRACTION=ML

(3) PRINTサブコマンド

[書式] PRINT=[DEFAULT] [INITIAL] [EXTRACTION]

 [ROTATION] [UNIVARIATE]

 [CORRELATION] [DET] [INV] [REPR]

 [AIC] [KMO] [FSCORE] [SIG] [ALL]

次のキーワードを上記の書式に従って与えることにより、結果出力の項目を選択できる。

UNIVARIATE 平均値と標準偏差値。

INITIAL 初期共通性、相関行列の固有値、説明分散率。

CORRELATION 相関行列。

SIG 相関の有意水準。

DET 相関行列の行列式。

INV 相関行列の逆行列。

AIC 反イメージ共分散行列と反イメージ相関行列。

KMO Kaiser-Meyer-Olkin測度と Bartlettテスト。

EXTRACTION 共通性、固有値、回転前の因子負荷行列。

REPR 再現相関と残差相関。

ROTATION 回転後の因子パターンと構造行列、変換行列、因子間相関行列。

FSCORE 因子得点係数行列。

ALL 利用できるすべての統計。

DEFAULT INITIAL, EXTRACTION, ROTATIONの3キーワードを指定したのと同じ。

(4) FORMATサブコマンド

[書式] FORMAT=[SORT][BLANK(n)][DEFAULT]

因子の解釈を容易にするために、因子負荷行列や因子構造行列の再編成を行なうサブコマンドである。以下のキーワードを指定することによって、出力結果の印刷の仕方を選択できる。

SORT 因子負荷量の大きさの順に変数を並びかえて表示する。

BLANK(n) n以下の絶対値の係数を削除して表示する。

DEFAULT 上記の2つの機能を働かせない。

[例4] FACTOR VARIABLES=VAR1 TO VAR10/

```
MISSING=MEANSUB/
WIDTH=100/
ANALYSIS=VAR1 TO VAR10/
FORMAT=SORT BLANK(.3)/
EXTRACTION=ULS/
```

(5) PLOTサブコマンド

[書式] PLOT=[EIGEN] [ROTATION(n1, n2)]

スクリープロットや回転後の因子空間への変数のプロットが可能になる。以下のキーワードを指定することによって、プロットの仕方を選択できる。

EIGEN スクリープロット。

ROTATION(n1 n2) 因子空間への変数のプロット。n1とn2にプロットする番号を指定する。

(6) CRITERIAサブコマンド

[書式] CRITERIA=[DEFAULT] [FACSCORES(n)]
 [MINEIGEN({1.0})] [ITERATE({25})]
 {eig} {ni}
 [RCONVERGE({0.0001})] [DELTA({0})]
 {r1} {d}
 [{KAISER }] [ECONVERGE({0.001})]
 {NOKAISER} {el }

各因子抽出や各因子回転に対して、以下のキーワードによりその基準を指定する。

FACTORS(nf) 抽出する因子数を指定。標準値は、MINEIGENで指定された値より固有値の大きい因子の数。

MINEIGEN(eg) 抽出する因子のための最小固有値をegに変更する。標準値は、1.0である。

ITERATE(ni) 因子解を求めるための繰返し計算の回数を指定する。標準値は25である。

ECONVERGE(e1) 因子抽出のための繰返しを止める収束判定基準を指定する。標準値は、0.001である。

RCONVERGE(r1) 因子回転のための繰返しを止める収束判定基準を指定する。標準値は、0.0001である。

KAISER Kaiserの正規化を行なう。標準処理なので、この指定がなくても実行される。

NOKAISER Kaiserの正規化を行わない。

DELTA(d) 斜交回転の際の回転の度を指定。標準値は0である。

DEFAULT 上記のすべての指定を標準値に戻す。

(7) DIAGONALサブコマンド

[書式] DIAGONAL=値リスト

主因子法を用いる場合に、相関行列の対角要素を予め与えるためのサブコマンドである。なお、省略時の標準値は、各対角要素に対してすべて1.00である。

```
[例5] FACTOR VARIABLES=VAR1 TO VAR6/  
      DIAGONAL=.55 .45 .35 .40 .60 .70/  
      EXTRACTION=PAF/  
      ROTATION=VARIMAX
```

15. 2.3 回転ブロック

(1) ROTATIONサブコマンド

```
[書式] ROTATION={VARIMAX  
                  {EQUAMAX  
                  {QUARTIMAX  
                  {OBLIMIN  
                  {NOROTATE  
                  {DEFAULT
```

因子の回転とその方法を指示する。EXTRACTIONサブコマンドを使用していない場合、標準の因子回転法はVARIMAXとなる。EXTRACTIONサブコマンドを使用して、ROTATIONサブコマンドを使用しない場合、SAVEサブコマンドがある場合を除いて、因子の回転は行なわれない。SAVEサブコマンドを指定した場合、ROTATIONサブコマンドを使用しない場合でも、標準処理としてバリマックス回転が行なわれる。回転を行ないたくない場合は、ROTATIONサブコマンドを用いてキーワードNOROTATEを指定する必要がある。ROTATIONサブコマンドには、以下のキーワードがある。

VARIMAX バリマックス回転。
EQUAMAX エカマックス回転。
QUARTIMAX コーティマックス回転。
OBLIMIN 直接オブリミン法による斜交回転。
NOROTATE 回転を行なわない。
DEFAULT VARIMAXと同じ。

なお、このサブコマンドは、1つの抽出に対して2つ以上重ねて与えることができる。

```
[例6] FACTOR VARIABLES=VAR1 TO VAR10/  
      EXTRACTION=ML/  
      ROTATION=VARIMAX/  
      ROTATION=OBLIMIN
```

15. 2.4 SAVEブロックのサブコマンド

(1) SAVEサブコマンド

因子得点を計算し、新たな変数として実行ファイルに保存する。

```
[書式] SAVE={ {REG_____} ( {ALL} ルート名 )  
             {BART_____} {n } }
```

{AR
{DEFAULT}

まず、以下のキーワードで因子得点計算の方法を指定する。

REG 回帰分析による方法。

BART Bartlettの方法

AR Anderson-Rubinの方法

DEFAULT REGと同じ。

つづいてその後に、求めたい因子得点の数 n を指定する。求められる最大のは、因子解の個数に等しい。キーワードALLを用いると、可能なすべての因子得点を計算することができる。最後に、7字以内で因子得点保存ための変数の命名を行なう。

〔例7〕 FACTOR VARIABLES=VAR1 TO VAR10/

MISSING=MEANSUB/ WIDTH=100/

ANALYSIS=VAR1 TO VAR10/

PRINT=CORRELATION DET REPR/

CRITERIA=FACTORS(3)/

FORMAT=SORT BLANK(.3)/

CRITERIA=FACTORS(2)/

PLOT=EIGEN ROTATION(1 2)/

EXTRACTION=ULS/

ROTATION=VARIMAX/

SAVE AR (ALL FSULS)/

上の例の場合、Anderson-Rubinの方法で、FSULS1とFSULS2と命名された2つの因子得点が計算され、実行ファイルにその結果が保存されている。なお、SAVEサブコマンドを同時に2つ以上重ねて与え、異なる方法で計算した2種以上の因子得点を保存することもできる。

15. 2.5 その他のサブコマンド

(1) READサブコマンド

〔書式〕 READ={CORRELATION[TRIANGLE]}
 {FACTOR(n)}
 {DEFAULT}

以下のように、相関行列または因子負荷行列を読み込んでFACTORを実行させる。

CORRELATION 相関行列を読み込む。

FACTOR(nf) 因子負荷行列を読み込む。nfは、分析に用いる因子数である。

DEFAULT CORRELATIONと同じ。

キーワードCORRELATIONの後にキーワードTRIANGLEを続けることによって、対角要素を含んだ下

三角行列として入力することができる。

```
[例8] FACTOR READ=CORRELATION TRIANGLE/  
      VARIABLES=VAR1 TO VAR10/
```

なお、相関行列や因子負荷行列を読み込んで計算を行う場合、NUMERIC コマンドで変数を定義し、INPUT MATRIXコマンドでその行列を含んだファイルを定義する必要がある。

(2) WRITEサブコマンド

```
[書式] WRITE={CORRELATION[TRIANGLE]}  
           {FACTOR(n)}  
           {DEFAULT}
```

相関行列または因子負荷行列を特定のファイルに出力する。その内容は、次のキーワードにより指定する。

CORRELATION 相関行列を書き込む。

FACTOR 因子負荷行列を書き込む。

DEFAULT CORRELATIONと同じ。

なお、相関行列や因子負荷行列を書き込む場合、FILE HANDLEコマンドやPROOCEDURE OUTPUTコマンドを用いて、その行列を含んだファイルを定義する必要がある。

16 CLUSTER：クラスター分析

CLUSTERは、クラスター分析を行なう。クラスター分析では、ケース間の類似性に基づいて、その分類を行う。クラスター分析における分類作業は、大別すると、集合的手法と分割的手法があり、それぞれの中で階層的手法と非階層的手法とに分けられる。集合的手法は、類似性の高いもの同志を少しづつ集めてゆくもので、分割的手法は、逆に全体をいくつかに分割するものである。階層的手法は、階層構造のあるクラスターを求めるよう分割あるいは結合を進めるもので、一方、非階層的手法は、そのようなクラスターの階層構造を求めない。CLUSTERは、集合的階層的クラスター分析を行う。

16. 1 一般書式

CLUSTER 変数リスト

```
[/MISSINGサブコマンド] [/READサブコマンド]  
[/WRITEサブコマンド] [/MEASUREサブコマンド]  
[/METHODサブコマンド] [/SAVEサブコマンド]  
[/IDサブコマンド] [/PRINTサブコマンド]  
[/PLOTサブコマンド]
```

16. 2 サブコマンド

(1) 変数リストの指定

[書式] 変数リスト

分析に使用する変数を指定する。

[例1] CLUSTER A B C

上の例は、変数A, B, Cの値に基づいて、ユークリッド距離の2乗値による平均法のクラスター分析を行う。

(2) METHODサブコマンド

[書式] METHOD={BAVERAGE}[(ルート名)][...]
 {WABERAGE}
 {SINGLE}
 {COMPLETE}
 {CENTROID}
 {MEDIAN}
 {WARD}

クラスターの結合法を指定する。

BAVERAGE 群間平均法によるクラスター分析。

WABERAGE 群内平均法によるクラスター分析。

SINGLE 最近隣法によるクラスター分析。

COMPLETE 最遠隣法によるクラスター分析。

CENTROID 重心法によるクラスター分析。

MEDIAN メディアン法によるクラスター分析。

WARD ワード法によるクラスター分析。

(3) MEASUREサブコマンド

クラスター分析に用いる近接性(proximity)の測度を算出する。以下のキーワードにより、近接性の測度を選択できる。

SEUCLID ユークリッド距離の2乗値。デフォルト。

EUCLID ユークリッド距離。

COSINE 変数ベクトル間の交角の余弦。

CHEBYCHEV チェビシェフ距離。

BLOCK 市街距離。

POWER(p,r) パワー距離。pとrは、パワー距離のパラメーター。

(4) SAVEサブコマンド

[書式] SAVE=CLUSTER({水準
 {最小値、最大値}})

特定のクラスターの水準でのクラスターメンバーシップを実行ファイルに新しい変数として保存する。SAVEの指定には、CLUSTER(最小値, 最大値)またはCLUSTER(n)がある。ここでの最小値と最大値は、保存したいクラスターの個数の最小値と最大値であり、nはその単一の値である。メン

バーシッパが保存されるクラスターの手続毎に、METHODサブコマンドにおいて、クラスター名を指定しなければならない。

```
[例2] CLUSTER A B C
      /METHOD=BAVERAGE (CLUSMEM)
      /SAVE=CLUSTER(3,5)
```

(5) IDサブコマンド

[書式] ID=変数名

ケースの同定をするための変数を指定する。

(6) PRINTサブコマンド

[書式] PRINT=[CLUSTER({水準
 {最小値、最大値}})] [DISTANCE]

[SCHEDULE] [NONE]]

CLUSTERは通常、用いたクラスタリングの技法の数や近接度の指標を印刷する。本コマンドにより、以下のように印刷出力の仕様を変更できる。

SCHEDULE クラスタ集合化のスケジュール。

CLUSTER (最小値, 最大値) クラスタメンバーシップ。クラスターの個数の最小値と最大値を指定することができる。また、単一の水準で値を入れることも可能である。

DISTANCE 項目間の距離行列を出力する。

NONE プロットを行わない。

(7) PLOTサブコマンド

[書式] PLOT=[VICICLE[(最小値 [, 最大値 [, 増分]])]]

[HICICLE[(最小値 [, 最大値 [, 増分]])]]

[DENDROGRAM] [NONE]]

クラスター解のプロットや樹状図の指定を行う。

VICICLE[(最小値 [, 最大値 [, 増分]])] 垂直の氷柱状のプロット。オプションとしてクラスターの水準の最小値、最大値、最小値から最大値への増分を指定できる。たとえば、最小値を1、最大値を9、増分を2とすると、1、3、5、7、9がプロットされる。

HICICLE[(最小値 [, 最大値 [, 増分]])] 水平の氷柱状のプロット。最小値、最大値、増分についてはVICICLEと同じ。

DENDROGRAM 樹状図の作成。

NONE プロットを行わない。

(8) READサブコマンド

[書式] READ=[SIMILAR] [{TRIANGLE}
 {LOWER
 }]

近接性行列を行列形式で入力するときに使用する。

SIMILAR 入力する行列が類似性に関する指標である場合。このキーワードを省略すると、入力された行列を距離行列または非類似性行列とみなす。

TRIANGLE 対角要素を含んだ下三角行列で与える場合。

LOWER 対角要素を含まない下三角行列で与える場合。

(9) WRITEサブコマンド

[書式] WRITE[=DISTANCE]

近接性行列を行列形式で出力する。

キーワードDISTANCEにより、距離行列として出力される。DISTANCEを省略すると、類似度行列として出力される。

(10) MISSINGサブコマンド

[書式] MISSING={LISTWISE}
{INCLUDE}

欠損値の処理に関する指定を行う。

LISTWISE リスト単位で欠損値を持つケースを除外して計算する。

INCLUDE 欠損値を持つケースも含めて計算する。

17 PROBIT

PROBITは、反応が2分割されるような従属変数（例えば、生存か死亡か、雇用中か失業かといった変数）に対する、1つ以上の独立変数の影響を推定するプログラムである。このプログラムは投薬反応分析 (dose-response analyses) などにおいて最も効果的であるが、ロジスティック回帰モデル(logistic regression models)の推定にも利用できる。

17.1 一般書式

PROBIT 反応度数変数名 OF 観察度数変数名 [WITH 変数リスト]

[BY 変数名(最小値, 最大値)]

[/MISSINGサブコマンド][MODELサブコマンド]

[/LOGサブコマンド][PRINTサブコマンド]

[/CRITERIAサブコマンド][NATRESサブコマンド]

PROBITは、指定されたモデルのパラメーターの最尤推定値を計算する。**PROBIT**では、変数リストとサブコマンドの指定を行う。変数リストには必ず反応度数変数、観察度数変数、および1つ以上の予測変量 (predictor)を指定しなければならない。グループ変数の指定、および各サブコマンドの指定はオプションである。**PROBIT**では一般に、個々の観察データを直接入力することはできない。このような場合は、**AGGREGATE**コマンドを使用し、前もって観察度数、および反応度数を計算

(5) PRINTサブコマンド

[書式] PRINT=[ALL][CI][FREQ][RMP][PARALL]
[NONE][DEFAULT]

標準出力に含まれない統計値や図表を出力させる。

DEFAULT 下記のFREQ、CI、およびRMPが出力される。

ALL 利用可能な全ての出力。

FREQ 観測された度数と予測された度数が、残差とともに出力される。

CI 特定の反応率を生起させる予測変量の値の95%信頼区間。

RMP グループ変数によって定義されたグループ間のRMP（相対的中央値ポテンシー、Relative median potency）を表示する。

PARALL 各グループの回帰直線の平行性のテスト(test of the parallelism)。

NONE 次の情報のみが無条件に出力される。ケースとモデルに関する情報。さらに予測変量が1つの場合には、予測変量を対数変換したものを横軸とし、反応率をプロビット変換したものを縦軸とする分布図。PLOGITモデルに対しては、パラメータ推定値と共分散。

(6) MISSINGサブコマンド

[書式] MISSING={LISTWISE}
{INCLUDE}
{DEFAULT }

LISTWISE 欠損値のリスト単位の除去。変数のいずれかに欠損値をもつケースはすべて除去される。

INCLUDE 欠損値定義を無視し、すべてのデータを計算に含める。

DEFAULT 標準処理。LISTWISEに同じ。

17. 3 コマンド編成例

PROBIT R OF N BY ROOT(1,2) WITH X

/MODEL = BOTH

/MATRES

/PRINT = ALL

18 LOGLINEAR

LOGLINEARは、カテゴリー変数を持つモデルにたいして、モデルの当てはめ(model fitting)、仮説検定、パラメータの推定などをおこなう一般的なプロシジャである。したがってLOGLINEARには、多重コンティンジェンシー表(multi-way contingency tables)の一般モデル、ログジットモデル、カテゴリー変数にたいするロジスティック回帰(logistic regression)、quasi-independence モデルなど、さまざまな関連する手法が含まれている。

18. 1 一般書式

LOGLINEAR 従属変数リスト (最小値, 最大値)

[BY] 独立変数リスト (最小値, 最大値) [WITH 共変量リスト]

[/WIDTHサブコマンド][CWEIGHTサブコマンド]

[/GRESIDサブコマンド][PRINTサブコマンド]

[/NOPRINTサブコマンド][PLOTサブコマンド]

[/CONTRASTサブコマンド][CRITERIAサブコマンド]

[/DESIGNサブコマンド[/DESIGNサブコマンド...]]

変数名の指定を、まず最初に行なう必要がある。キーワードBYを使用する場合は、キーワードBYの前のカテゴリ変数が従属変数となり、後のカテゴリ変数が独立変数となる。キーワードWITHの後のセル共変量(Cell covariates)は、連続的変量でなければならない。DESIGNサブコマンドによって、当てはめるモデルの指定を行う。検討したいモデルごとにDESIGNサブコマンドが必要である。サブコマンドは何度でも使用できる。またDESIGNサブコマンドを除いて、サブコマンドで指定した内容は新しく変更されるまで、モデルからモデルへと引きつがれる。あるDESIGNサブコマンドに関係するサブコマンドは、必ずそのDESIGNサブコマンドの前に置かなければならない。もし最後のDESIGNサブコマンドの後にサブコマンドがあれば、飽和モデルとみなされる。

18. 2 サブコマンド

(1) DESIGNサブコマンド

[書式] DESIGN=効果 効果 ... 効果 BY 効果 ... [/DESIGN...]

当てはめるモデルの指定を行なう。このサブコマンドがない場合は、飽和モデルとみなされる。飽和モデルとはすべての主効果と交互作用効果とを含むモデルである。1つのモデルにたいして必ず1つのDESIGNサブコマンドが必要である。主効果のみのモデルの場合は、変数名だけを指定する。交互作用効果を指定する場合には、キーワードBYを使用する。つぎの例は、変数A、Bの2要因の主効果と交互作用効果を含むモデルを指定している。

[例1] LOGLINEAR A(1,4) B(1,5)/

DESIGN = A,B,A BY B/

括弧で囲まれた整数が後についている変数は、自由度1の分割(single-degree-of-freedom partition)によるものであることを示している。

(2) CWEIGHTサブコマンド

[書式] CWEIGHT={変数名}
{matrix}

セルに対する重みづけを行う。変数を指定する方法と、マトリックスを利用する方法とがある。変数を用いる場合、重みづけ変数として指定できる変数名は1つだけであり、数値変数でなければ

ならない。マトリックスを利用する場合には、CWEIGHTサブコマンドの等号の後に、括弧で囲まれた重みづけのマトリックスを指定する。マトリックスの中には、必ずセル数と同じ数だけの要素がなくてはならない。同じ数値が連続する時は、アスタリスク(*)を利用することができる。例えば、

```
[例2]  LOGLINEAR A(1,2) BY B(1,3) C(1,2)/  
        CWEIGHT=(0 3*1 0 3*1 0 3*1)/
```

という指定は、

```
[例3]  LOGLINEAR A(1,2) BY B(1,3) C(1,2)/  
        CWEIGHT=(0 1 1 1 0 1 1 1 0 1 1 1)/
```

とまったく同じである。マトリックスを利用する場合には、CWEIGHTサブコマンドの指定は何度でも行なえる。CWEIGHTサブコマンドの指定内容は、新しく変更されるまでずっと引継がれる。

(3) GRESIDサブコマンド

GRESIDサブコマンドは、各セルの観察度数、期待度数、調整済み残差(adjusted residuals)などの1次結合を計算する。変数を指定する方法と、マトリックスを利用する方法とがある。マトリックスの中には、必ずセル数と同じ数だけの要素がなくてはならない。例えば

```
[例4]  LOGLINEAR MONTH(1,18) WITH Z/  
        GRESID=(6*1,12*0)/ GRESID=(6*0,6*1,6*0)/  
        GRESID=(12*0,6*1)/ DESIGN=Z
```

では、1番はじめのGRESIDサブコマンドによって、最初の6ヵ月が結合され”初期”効果となっている。2番目のGRESIDサブコマンドは”中期”効果として、次の6ヵ月を結合し、3番目のGRESIDサブコマンドは”後期”効果として、最後の6ヵ月を結合している。

(4) PRINT,NOPRINTサブコマンド

[書式] (PRINTまたはNOPRINT)=(キーワード群)

PRINTサブコマンドを使用することによって、標準出力に含まれない統計値や図表を出力させることができる。またNOPRINTサブコマンドによって、結果の出力を制限することができる。以下のキーワードを使用する。

FREQ 各セルの観察度数と期待度数、およびそれぞれの%を出力する。

RESID 残差、標準化残差(standardized residuals)、調整済み残差(adjusted residuals)を出力する。

DESIGN モデルのデザインのマトリックスを出力する。

ESTIM モデルのパラメータ推定値を出力する。

COR パラメータ推定値の相関マトリックスを出力する。

DEFAULT FREQおよびRESIDの出力。DESIGNサブコマンドが省略されている場合には、ESTIMも出力される。

NONE デザインに関する情報と、当てはまりの良さに関する統計値のみ出力される。NOPRINT
サブコマンドでは使用できない。

(5) PLOTサブコマンド

[書式] PLOT={キーワード群}

オプションとなっている図を表示させる。キーワードを省略すると、図は何も表示されない。

RESID 調整済み残差と、観測度数および期待度数を軸とする散布図。

NORMPROB 調整済残差と標準化値(normal value)および標準化値からの偏差の2種類の散布図。

NONE 何も出力されない。

DEFAULT RESIDおよびNORMPROBを指定したことになる。

(6) CONTRASTサブコマンド

[書式] CONTRAST (変数名)={キーワード群}

要因についての対比(contrast)のタイプを指定する。要因となるのは、カテゴリカルな従属変数または独立変数である。CONTRASTサブコマンドの後に変数名を括弧で囲んで指定し、その後に等号と選択した対比のキーワードを続ける。LOGLINEARにおいては、MANOVAの場合と違い、対比係数の和が0にならなくてもよいし、また直交している必要もない。

使用できるキーワードは次のとおりである。

DEVIATION(refcat) 全体的な効果(overall effect)からの偏差。2.1版からこれがCONTRASTサブコマンドがない場合の標準値となっている。Refcatとはパラメータ推定値が表示されないカテゴリーのことである。Refcatの指定がない場合は、変数の最後のカテゴリーがデフォルト値として用いられる。

DIFFERENCE それより前にある全ての水準の効果の平均との対比。HELMERT対比の逆になる。

HELMERT それより後にある全ての水準の効果の平均との対比。

SIMPLE(refcat) 各水準の要因と最終水準のそれとの対比。2.1版より前の版では、SIMPLEがCONTRASTサブコマンドが省略された場合のデフォルト値となる。キーワードSIMPLEの後に括弧で囲んで参照カテゴリーを指定することができる。参照カテゴリーが省略された場合は、最後のカテゴリーがデフォルト値として用いられる。

REPEATED 隣接した水準(level)間の対比。

POLYNOMIAL(metric) 直交多項式による対比。標準値は等間隔(equal spacing)。括弧の中に線形多項式の係数を指定することにより、水準間の線形比較ができる。

[BASIS]SPECIAL(matrix) ユーザーの定義による対比。必ずカテゴリー数の2乗の数だけの要素を指定する必要がある。2.1版からは、キーワードBASISの使用が可能となった。SPECIALの前にBASISが指定されると、特定の対比に対して基本マトリックスが生成される。BASISがない場合は、指定されたマトリックスが、基本マトリックス(basis

matrix)となる

(7) CRITERIAサブコマンド

[書式] CRITERIA=[CONVERGE($\frac{0.001}{\text{eps}}$)] [ITERATE($\frac{20}{n}$)]
[DELTA($\frac{0.5}{d}$)] [DEFAULT]

CONVERGE(eps) 収束の基準値eps。

ITERATE(n) 最大繰返し回数n。

DELTA(d) セルデルタ値d。分析の前に、デルタ値を各セルの度数に加える。

DEFAULT 各パラメータを標準値（上記書式中の下線部分）にリセットする。

(8) WIDTHサブコマンド

[書式] WIDTH={ $\frac{132}{72}$ }

WIDTHによって、表示幅を変更することができる。ただし1度に使用できる表示幅は1種類だけであり、途中から変更はできない。また、狭い表示幅を選択した時は、表示の1部が省略される。

18. 3 オプション

1 欠損値の定義を無視する。

18. 4 コマンド編成例

[例1] 一般対数線形モデル(general log-linear model)の例

```
LOGLINEAR DPREF(2,3) RACE ORIGIN CAMP(1,2)/  
PRINT=DEFAULT ESTIM/  
DESIGN=DPREF,RACE,ORIGIN,CAMP,  
DPREF BY RACE,DPREF BY ORIGIN,DPREF BY CAMP,  
RACE BY CAMP,RACE BY ORIGIN,ORIGIN BY CAMP,  
RACE BY ORIGIN BY CAMP,  
DPREF BY ORIGIN BY CAMP
```

[例2] 多項式ロジットモデル(multinomial logit model)の例

```
LOGLINEAR PREF(1,5) BY RACE ORIGIN CAMP(1,2)/  
PRINT=DEFAULT ESTIM/  
CONTRAST(PREF)=SPECIAL(5*1,1 1 1 1 -4.3 -1 -1 -1 0,  
0 1 1 -2 0,0 1 -1 0 0)/  
DESIGN=PREF,PREF BY RACE,PREF BY ORIGIN,PREF BY CAMP,  
PREF BY RACE BY ORIGIN,PREF BY RACE BY CAMP,  
PREF BY ORIGIN BY CAMP,PREF BY RACE BY ORIGIN BY CAMP
```

19 MANOVA：多変量分散分析

MANOVAは、多変量分散共分散分析に関する汎用プログラムである。MANOVAを使用することによって、分散分析および共分散分析を行うことができる。乱塊法(randomized block)、分割プロット(split-plot)、入れ子型のデザイン、繰返しのあるデザインなども分析できる。また重回帰係数の推定も可能であるし、主成分、判別関数の係数、正準相関係数など、一般線形モデル(general linear model)に関する様々な統計値を得ることができる。

19. 1 一般書式

MANOVA 従属変数リスト [BY 要因リスト (最小値, 最大値) ...]

[WITH 共変量リスト ...]

[/WSFACTORSサブコマンド] [/TRANSFORMサブコマンド]

[/WSDSIGNサブコマンド] [/MEASUREサブコマンド]

[/RENAMEサブコマンド] [/PRINT, NOPRINTサブコマンド]

[/PLOTサブコマンド] [/METHODサブコマンド]

[/READサブコマンド] [/WRITEサブコマンド]

[/ANALYSISサブコマンド] [/PARTITION]

[/CONTRASTサブコマンド] [/SETCONSTサブコマンド]

[/ERRORサブコマンド] [/DESIGNサブコマンド]

MANOVAでは、まず最初に分析に用いる変数名の指定を行う必要がある。MANOVAでは3種類の変数を用いる。従属変数および共変量は連続変量(continuous variables)であり、要因(factors)はカテゴリカル変量である。

従属変数リストは、MANOVAコマンドの直後に続ける。普通の場合、MANOVAはリスト上の従属変数を合成変量として扱う。つまり、多変量デザイン(multivariate design)を想定する。しかしながらANALYSISサブコマンドを使用することによって、変数の種類や組合せを変更することができる。

要因リストは、キーワードBYの後に続ける。要因名の後には、レベルの最小値と最大値を括弧で囲んで指定する。この範囲外のケースは分析から除外される。レベルの値が連続した整数値でない場合は、RECODEコマンドを使って修正しておく必要がある。

共変量リストは、キーワードWITHのあとに続ける。

19. 2 サブコマンド

MANOVAには、非常に多くのサブコマンドがある。これらは次の4つのグループに分けることができる。

デザインに関するサブコマンド WSFACORSサブコマンドは、繰返しのあるデザインにおける被験者内要因(within-subjects factor)を指定する。WSDSIGNサブコマンドは、被験者内要因

印刷出力に関するサブコマンドPRINT サブコマンドは、オプションの出力を行う。PLOTサブコマンドは、オプションの分布図を出力する。

マトリックス入出力に関するサブコマンド WRITEサブコマンドによって、マトリックス出力を行うことができる。この出力はREADサブコマンドによって読むことができ、後の分析で使用する事が可能である。

[書式] ANALYSIS [(REPEATED {CONDITIONAL })]
{UNCONDITIONAL}
=従属変数リスト [WITH 共変量リスト]

```
[例 1] MANOVA A,B,C,D,E,F BY FAC(1,4)/  
ANALYSYS = (A,B/ C/ D WITH E ) WITH F/
```

```
[例2] MANOVA A,B,C,D,E,F BY FAC(1,4)/
      ANALYSIS = A B WITH F/ DESIGN/
```

```
ANALYSYS = C WITH F/ DESIGN/
```

```
ANALYSYS = D WITH E F/
```

という指定と同じである。またANALYSYSサブコマンドの後に、括弧で囲んでキーワードCONDITIONALを続けると、後のリストに対しそれ以前の全ての従属変数を共変量に含める指定をしたことになる。例えば、

[例3] MANOVA A, B, C, D, E, F BY FAC(1, 4)/

```
ANALYSYS(CONDITIONAL) = (A B C/ D E) WITH F/
```

という指定は、

[例4] MANOVA A, B, C, D, E, F BY FAC(1, 4)/

```
ANALYSYS = A B C WITH F/ DESIGN/
```

```
ANALYSYS = D E WITH A B C F/
```

という指定をしたことと同じになる。繰返しのあるデザインの場合には、ANALYSYSサブコマンドの後に、キーワードREPEATEDを括弧で囲んで続ける。

(2) DESIGNサブコマンド

[書式] DESIGN={ [CONSTANT ...]
 { [効果 効果 ...]
 { [CONTIN(変数リスト) ...]
 { [効果 効果 ... BY 効果 ...]
 { [効果 効果 ... {WITHIN} 効果 効果 ...]
 {W
 }
 { [効果 + 効果 ...]
 { [要因(水準) ... [WITHIN 要因 (分割の指示) ...]
 { [CONPLUS ...]
 { [MWITHIN ...]
 { { [被検項群] {AGAINST} {WITHIN または W}
 { [被検項=n] {VS} {RESIDUL または R}
 {WR または RW}
 {n

被験者間モデルの指定を行う。このサブコマンドは、何個でも使用できる。ある1つのモデルに対する指定の中でこのサブコマンドが1番最後に来る。これ以外の全てのサブコマンドは、それに続くDESIGNサブコマンドに対して有効である。本サブコマンド省略時に採用される標準のモデルは完全要因計画モデル (full factorial model) である。本サブコマンドでは、次のキーワードを使用してモデルにおける効果のリストを指定する。主効果のみのモデルは、要因名を並べ、交互作用の指定はキーワードBYでおこなう。例えば、A, B, Cの3次の交互作用は、A BY B BY C と書けばよい。

CONTIN いくつかの連続変数をプールして単一の効果とする。ANALYSISサブコマンドに含まれていない連続変数と要因との間の交互作用も指定することができる。ただし共変量間の交互作用は指定できない。

WITHINまたはW 入れ子型のデザインであることを示す。WITHINの左側の項が右側の項の中に入

れ子となっていることを示す。

プラス記号(+) 効果をプールすることを示す。

CONSPLUS 従属変数の総平均(grand means)とパラメータ値の合計からなる推定値を得る。

誤差項を示すキーワード 次の3種類が用意されている。

WITHIN または W ... セル内誤差項

RESIDUAL または R ... 残差誤差項

WR または RW ... セル内誤差項と残差誤差項との結合

誤差項の指定は、まず最初にテストしたい項の名前を、続いてキーワードVSまたはAGAINSTのいずれかを、最後に誤差項のキーワードを書くことによって行う。例えば、

[例5] DESIGN = AGE BY SEX AGAINST RESIDUAL

という指定は、AGEとSEXの2次の交互作用項について、残差を誤差項としてテストすることを意味する。誤差項としてはこの3種類のほかに、ユーザーによって最高10個まで定義して用いることができる。項(term) = n と指定することにより誤差項として定義できる。ここで n は1から10までの整数である。例えば、

[例6] DESIGN = AGE BY TREATMENT = 1

という指定は、AGEとTREATMENTの2次の交互作用項を、誤差項1として定義するということを意味する。誤差項の指定は、キーワードの場合と同じく、項(term) AGAINST n または 項(term) VS n という形で行なう。

CONSTANT 定数項(constant term)をモデルの中に含める。一般にMANOVAは、自動的に定数項をモデルに含める。ただしMETHODサブコマンドのなかでNOCONSTANTが指定された場合、DESIGNサブコマンドのなかでCONSTANTを指定しない限り、定数項はモデルに含まれない。

MWITHIN DESIGNおよびWSDSIGNサブコマンドのなかで使用され、繰返しのあるデザインでの単純効果(simple effect)をテストするのに用いられる。

(3) WSFACTORSサブコマンド

[書式] WSFACTORS=被験者内要因名(レベル数)...

被験者内要因の要因名と水準(level)を指定する。例えば、DRUG1からDRUG4に対し被験者内要因名としてTRIALと名付けたい場合には、

[例7] MANOVA DRUG1 TO DRUG4/

WSFACTORS = TRIAL(4)

と指定する。また被験者内要因が2つあり、最初が2水準で2番目が3水準の場合には、

[例8] MANOVA Y1 TO Y6 BY GROUP(1,2)/

WSFACTORS = DRUG(2),DOSE(3)/

というように指定する。MANOVAコマンドのなかの従属変数の並び方は、WSFACTORSサブコマンドの被験者内要因の順序と水準数できまる並び方と対応していなければならない。またMANOVAコマンドのなかの従属変数の数は、被験者内要因の水準数の積の整数倍となっていなければならない。WSFACTORSサブコマンドは、サブコマンドのなかで1番最初に置かなければならない。

(4) WSDSIGNサブコマンド

〔書式〕 WSDSIGN=効果 効果 ...

被験者内(within-subjects)モデルの指定を行なう。指定のしかたはDESIGNサブコマンドとほぼ同様であるが、キーワードで使用できないものがあるなど、いくつかの点で異なる。

繰返しのあるデザインで使用する

(5) ANALYSISサブコマンド

〔書式〕 ANALYSIS [(REPEATED) {CONDITIONAL }])
{UNCONDITIONAL}

=従属変数リスト [WITH 共変量リスト]

キーワードREPEATEDをともなったANALYSISサブコマンドは、WSDSIGNおよびWSFACTORSサブコマンドで定義された被験者内要因の効果を生成する。このどちらかのサブコマンドがない場合、ANALYSIS(REPEATED)を指定するとエラーメッセージが出る。ANALYSIS(REPEATED)サブコマンドのなかで、変数の指定を行うこともできる。

(6) MEASUREサブコマンド

〔書式〕 MEASURE=新しい名前 ...

MANOVAでは、多変量で繰返しのあるデザインを分析することができる。このときPRINTサブコマンドでキーワードSIGNIF(AVERF)を指定すると、SPSS[®]は多変量の場合の結果とともに、従属変数がプールされたもしくは”平均”された場合の結果も印刷する。MEASUREサブコマンドは、このプールされた場合の結果にラベル(名前)をつけるために使用される。

(7) TRANSFORMサブコマンド

〔書式〕 TRANSFORM[(変数リスト)]=[キーワード群]

従属変数や共変量の1次変換の指定を行う。変数リストはスラッシュで区切り、全体を括弧で囲む。各リストに含まれる変数の数は同じでなければならない。MANOVAは、指定された変換を各リストにたいして行なう。標準処理では、MANOVAは全ての従属変数と共変量を変換する。TRANSFORMサブコマンドは、1つのジョブのなかでいくつ使ってもよい。7種類の変換が利用できる。変換の種類を指定する前に、キーワードCONTRAST、BASIS、またはORTHONORMのいずれかを指定しなければならない。

以下の変換キーワードが使用できる。

DEVIATIONS(refcat) 従属変数を、リストに含まれる従属変数の平均と比較する。

DIFFERENCE 従属変数を、リストのなかでその従属変数より前に位置する従属変数の平均と比較する。HELMERTの逆になる。

HELMERT 従属変数を、リストのなかでその従属変数より後に位置する従属変数の平均と比較する

SIMPLE(refcat) 各従属変数を最後の従属変数と比較する。最後の従属変数以外の変数を比較の基準にしたい場合は、その変数の番号を括弧で囲んでキーワードのあとにつける。

REPEATED 隣どうしの2つの従属変数を比較する。したがって隣りあう2つの変数の差が新しい変数となる。

POLYNOMIAL(metric) 変換リスト中の変数に、直交多項式(orthogonal polynomials)を当てはめる。平均が1番目の変数と置きかわる。以後1次の成分が2番目の変数に、さらに2次の成分が3番目の変数にというように置き換わっていく。

SPECIAL(matrix) 任意の変換マトリックスを指定することができる。ただし行および列の数が変数の数と同じである正方行列でなければならない。

(8) RENAMEサブコマンド

[書式] RENAME={新しい名前} ...
 *

TRANSFORMまたはWSDSIGNサブコマンドを使用して従属変数や共変量の変換を行った場合には、RENAMEサブコマンドを利用して変数名を変更することができる。RENAMEサブコマンドを使用しなければ、以前の変数名がそのまま使用される。RENAMEサブコマンドを使用する場合は、全ての変数について指定を行わなければならない。このとき変数名を変更する必要の無いものに対しては、アステリスク(*)を指定すれば良い。

(9) METHODサブコマンド

[書式] METHOD=[MODELTYPE({MEANS })]
 {OBSERVATIONS})

 [ESTIMATION({QR }) {NOLASTRES} {NOBALANCED} {CONSTANT})]
 {CHOLESKY} {LASTRES} {BALANCED} {NOCONSTANT}]

 [SSTYPE({UNIQUE })]
 {SEQUENTIAL}]

MODELTYPE パラメータ推定のためのモデルを指定する。セル平均(cell means)モデルを使用する場合はMEANSを、観測(observations)モデルを使用する場合はOBSERVATIONSを指定する。

ESTIMATION パラメータ推定の方法に関する指定を行なう。選択肢のペアが4組あり、それぞれどちらか1方を選択する。

QRまたはCHOLESKY QRを使用すると一般に非常に正確な推定値が得られる。コレスキ一法を用いる場合にはCHOLESKYを指定する。

NOBALANCEDまたはBALANCED MANOVAは普通、バランスのとれていない（セル内のケース数が等しくない）デザインを想定している。つまりNOBALANCEDが標準である。しかしながら直交性をもつバランスのとれたデータを分析する場合には、キーワードBALANCEDを指定することによって、処理時間を大幅に節約することができる。

NOLASTRESまたはLASTRES MANOVAでは、最高次の交互作用のパラメータ推定値を必要としない場合には、これを省略し計算時間を大幅に節約することができる。省略した場合でも、交互作用の有意性の検定は行なうことができる。モデルの最後にこの交互作用を指定すれば、MANOVAは減算によりこの効果に対する平方和を計算することができるからである。キーワードLASTRESを指定することによって、MANOVAはこの計算を行なう。

CONSTANTまたはNOCONSTANT CONSTANTを選択すると、たとえDESIGNサブコマンドのなかではっきりと指定しなくとも、定数項がモデルのなかに含まれる。NOCONSTANTを選択すると、定数項はモデルのなかに含まれない。

SSTYPE MANOVAには、平方和を分解する手法が2通り用意されている。2.1版からは回帰的手法が標準となっている。この場合ある項の効果は、他の全ての項の効果について調整される。キーワードUNIQUEを指定したばあいもこの手法が用いられる。

平方和を階層的に分解する場合には、キーワードSEQUENTIALを指定する。この場合、各項はDESIGNサブコマンドのリスト上で先行する項に対してのみ調整される。これは直交分解であり、モデル中の各平方和を加えていくと全体の平方和と等しくなる。2.1版以前のSPSS[®]ではこの手法が標準となっている。

(10) PARTITIONサブコマンド

[書式] PARTITION (要因名) [= ({1,1,...}) / {自由度, 自由度, ...}]

要因に関する自由度を分割する。サブコマンド名の後に要因名を括弧で囲んで続ける。さらに等号を続け、その後に分割された自由度を示す整数のリストを括弧で囲んで続ける。例えば、水準が12（したがって自由度は11）である要因TREATMNTを、自由度1に分割したい場合には、

[例9] PARTITION(TREATMENT) = (1,1,1,1,1,1,1,1,1,1,1,1)/

と指定する。同じ整数が続くときはアスタリスク(*)を使うこともできる。次の

[例10] PARTITION(TREATMENT) = (11*1)/

という指定は、最初の指定とまったく同じである。どの要因も標準として最初は自由度1に分割さ

(11) CONTRASTサブコマンド

[書式] CONTRAST (要因名) = {
 DEVIATION[{(refact)}]
 SIMPLE[{(refact)}]
 DIFFERENCE
 HELMERT
 REPEATED
 POLYNOMIAL[{(1,2,3,...)}]}
 }

$$\left\{ \begin{array}{c} \text{SPECIAL(マトリックス)} \\ \text{\{metric\}} \end{array} \right\}$$

ある要因の水準間の対比を指定する。次の7種類の対比指定用のキーワードがある。

DEVIATION(refcat) 総平均からの偏差。MANOVAはふつう最後のカテゴリについての偏差を省略する。しかし偏差の合計は0にならなければならないから、他の偏差の合計がわかれば、省略された偏差も簡単に求めることができる。カテゴリの数字を括弧で囲むことによって、省略されるカテゴリを指定することもできる。対比マトリックスの一般形は次のようになる。ただし k は水準数である。

$$\begin{array}{l} \text{mean} \quad (\quad 1/k \quad 1/k \quad \cdots \quad 1/k \quad 1/k) \\ \text{df}(1) \quad (1-1/k \quad -1/k \quad \cdots \quad -1/k \quad -1/k) \\ \text{df}(2) \quad (-1/k \quad 1-1/k \quad \cdots \quad -1/k \quad -1/k) \\ \vdots \\ \vdots \end{array}$$

DIFFERENCE ある水準と、それより前に位置する水準の平均との対比を行う。ヘルマート対比の逆になる。対比マトリックスの一般形は次のようになる。

$$\begin{array}{l} \text{mean} \quad (\quad 1/k \quad 1/k \quad 1/k \quad \cdots \quad 1/k) \\ \text{df}(1) \quad (\quad -1 \quad 1 \quad 0 \quad \cdots \quad 0) \\ \text{df}(2) \quad (\quad -1/2 \quad -1/2 \quad 1 \quad \cdots \quad 0) \\ \vdots \\ \vdots \\ \text{df}(K-1) \quad (-1/(k-1) \quad -1/(k-1) \quad -1/(k-1) \quad \cdots \quad 1) \end{array}$$

SIMPLE(refcat) 各水準と最後の水準との対比を行う。最後の水準以外を比較カテゴリとしたい場合は、その水準の番号を括弧で囲んで指定する。対比マトリックスの一般形は次のようになる。

$$\begin{array}{l} \text{mean} \quad (\quad 1/k \quad 1/k \quad \cdots \quad 1/k \quad 1/k) \\ \text{df}(1) \quad (\quad 1 \quad 0 \quad \cdots \quad 0 \quad -1) \\ \text{df}(2) \quad (\quad 0 \quad 1 \quad \cdots \quad 0 \quad -1) \\ \vdots \\ \vdots \\ \text{df}(K-1) \quad (\quad 0 \quad 0 \quad \cdots \quad 1 \quad -1) \end{array}$$

HELMERT ヘルマート対比を行なう。ある水準と、それより後に位置する水準の平均との対比を行う。対比マトリックスの一般形は次のようになる。

$$\begin{array}{l} \text{mean} \quad (\quad 1/k \quad 1/k \quad \cdots \quad 1/k \quad 1/k) \\ \text{df}(1) \quad (\quad 1 \quad -1/(k-1) \quad \cdots \quad -1/(k-1) \quad -1/(k-1)) \\ \text{df}(2) \quad (\quad 0 \quad 1 \quad \cdots \quad -1/(k-2) \quad -1/(k-2)) \\ \vdots \\ \vdots \end{array}$$

$$\begin{aligned} \text{df}(K-2) & \left(\begin{array}{ccccc} 0 & 0 & 1 & -1/2 & -1/2 \end{array} \right) \\ \text{df}(K-3) & \left(\begin{array}{ccccc} 0 & 0 & \cdots & 1 & -1 \end{array} \right) \end{aligned}$$

POLYNOMIAL 直交多項式対比 (orthogonal polynomial contrasts) を行う。水準間の間隔を指定することができる。等間隔の場合には、水準数を k とすると、1 から k までの連続する整数を指定する。水準間の間隔が異なる場合には、それに対応した数値を指定する。例えば、第2のグループおよび第3のグループに対する投薬量が、それぞれ第1のグループの4倍および7倍である場合には、

CONTRAST(DRUG) = POLYNOMIAL(1, 4, 7)/

と指定すればよい。

REPEATED 隣あう2つの水準間の比較を行う。対比マトリックスの一般形は次の様になる。

$$\begin{aligned} \text{mean} & \left(\begin{array}{cccccc} 1/k & 1/k & 1/k & \cdots & 1/k & 1/k \end{array} \right) \\ \text{df}(1) & \left(\begin{array}{cccccc} 1 & -1 & 0 & \cdots & 0 & 0 \end{array} \right) \\ \text{df}(2) & \left(\begin{array}{cccccc} 0 & 1 & -1 & \cdots & 0 & 0 \end{array} \right) \\ & \vdots \\ \text{df}(K-1) & \left(\begin{array}{cccccc} 0 & 0 & 0 & \cdots & 1 & -1 \end{array} \right) \end{aligned}$$

SPECIAL 利用者の定義する対比を行う。行と列の数は水準の数と等しくなければならない。

次の例は4水準の要因 TREATMNT にたいするものである。

[例 1 1] CONTRAST(TREATMNT) = SPECIAL(1 1 1 1
3 -1 -1 -1
0 2 -1 -1
0 0 1 -1)/

(1 1) SETCONST サブコマンド

[書式] SETCONST=[ZETA(数値)] [EPS(数値)]

MONOVA で使用される重要な定数の設定を行う。

ZETA(zeta) パラメータ推定のための基本マトリックスを構成するときや、印刷時に使用するゼロの絶対的な値を設定する。標準値は、 10^{-8} である。

EPS(eps) コレスキー分解などを行なう時の、マトリックスの対角線要素をチェックするために使用するゼロの相対的な値を設定する。標準値は、 10^{-8} である。

(1 2) ERROR サブコマンド

[書式] ERROR= { WITHIN } または { W }
 { RESIDUAL } { R }
 { WITHIN+RESIDUAL } { WR }
 { n }

被験者間の各効果に対する標準設定の誤差項を指定する。

WITHIN セル内の誤差項をもちいる。Wと省記できる。

RESIDUAL 残差誤差項をもちいる。Rと省記できる。

WITHIN+RESIDUAL セル内誤差項と残差誤差項をプールしてもちいる。WRまたはRWと省記できる。

n DESIGNサブコマンドのなかで指定された項(model term)を使用する。

なおMANOVAは、以下の基準にしたがって誤差項を選択する。

- ・もしセル内誤差項があれば、これを使用する。
- ・セル内誤差項がなければ、残差誤差項を使用する。
- ・オブザベーションモデル(observations model)を使って処理されたデザインにたいしては、セル内誤差項と残差誤差項のプールされた誤差項がデフォルトとなる。

(13) PRINT, NOPRINTサブコマンド

[書式] {PRINT }={キーワード群}
 {NOPRINT}

MANOVAによる各種の印刷出力をコントロールする。PRINTサブコマンドによって指定された項目が印刷される。反対にNOPRINTサブコマンドは、指定する項目の印刷出力を停止する。キーワードの指定方法はどちらのサブコマンドも同じである。例えば、

[例12] MANOVA Y BY A(1,3) WITH X/
 PRINT=CELLINFO(MEANS)

という指定によって、各セルの平均が印刷される。

以下のキーワードが使用できる。

CELLINFO([MEANS] [SSCP] [COV] [COR]) セルに関する情報が印刷される。

MEANS 各セルの平均、標準偏差、および頻度(count)

SSCP 積和行列

COV 分散共分散行列

COR 相関行列

HOMOGENEITY([BARTLET] [COCHRAN] [BOXM]) 分散の同質性についての統計量を印刷する。

BARTLETT バートレット・ボックスのF。

COCHRAN コ克蘭のC。

BOXM ボックスのM(多変量の場合のみ)。多変量で繰り返しのあるデザインの場合には特に有効である。

DESIGN([ONEWAY] [OVERALL] [DECOMP] [BIAS] [SOLUTION]) デザインに関する情報が印刷される。

ONEWAY 各要因についての一元配置の基礎統計量。
 OVERALL デザインマトリックス。
 DECOMP デザインにたいするQR／コレスキー分解。
 BIAS デザイン中のバイアスを示す係数。
 SOLUTION 効果の検定に用いられたセル平均の、1次結合の係数。
 PRINCOMPS([COR] [NCOMP(n)] [MINEIGEN(eigencut)] [COV] [ROTATE(rottype)]) 主成分に関する統計量を印刷する。
 COR 誤差相関(error correlation)行列についての主成分分析。
 COV 誤差分散共分散行列についての主成分分析。
 ROTATE(rottype) 回転の種類を指定する。括弧のなかに、VARIMAX、EQUAMAX、QUARTMAX、またはNOROTATEのいずれかを指定する。
 NCOMP(n) 回転する主成分の数を指定する。
 MINEIGEN(eigencut) 主成分の抽出を打ち切る基準として、最小固有値eigencutを指定する。
 ERROR([SSCP] [COV] [COR] [STDDEV]) 誤差に関する統計量を印刷する。
 かっこ内のキーワードはキーワードCELLINFOの項を参照。STDDEVは標準偏差。SSCP
 SIGNIF([MULTIV] [EIGEN] [DIMENR] [UNIV] [HYPOTH] [STEPDOWN]
 [AVERF] [BRIEF] [AVONLY] [SINGLEDF]) 有意性の検定に関する統計量を印刷する。
 MULTIV グループ間の差についての多変量Fテスト。標準出力。
 EIGEN $S_n S_n^{-1}$ の固有値。標準出力。
 DIMENR dimension-reduction分析★。標準出力。
 UNIV 単変量Fテスト。標準出力。
 HYPOTH 仮説によるSSCPマトリックス。
 STEPDOWN ロイ・バークマンのステップダウンFテスト。
 AVERF 繰り返し要因にたいする平均Fテスト(averaged F test)。
 BRIEF 多変量に関する出力を縮小して印刷する。
 AVONLY 平均された結果のみを出力。繰り返しのある要因にたいしてもちいる。
 SINGLEDF 自由度1の効果のリスト
 DISCRIM([RAW] [STAN] [ESTIM] [COR] [ROTATE(rottype)] [ALPHA(alpha)]) 多変量分析における従属変数や独立変数のセットに対し正準分析を行う。もし予測値(predictor)のセットが連続変数であるならば、MANOVAは正準相関分析の結果を印刷する。もし予測値のセ

Vol. 18 No.2 1988-8

CELLPLOTS セルに関する統計量を用いた図として、セル平均－セル分散の分布図、セル平均－セル標準偏差の分布図、セル平均のヒストグラムを作成する。

BOXPLOT 間隔尺度の各変数について、箱型の図を印刷する。

NORMAL 各連続変数について、normal および detrended normal図を印刷する。

STEMLEAF 間隔尺度の各変数について、幹葉表示(stem-and-leaf display)の印刷をおこなう。

ZCORR 多変量分析における従属変数間の、セル内相関(within-cells correlations)の図。

PMEANS 各従属変数について、推定値の平均、調整された平均、観測値の平均のそれぞれが、グループごとに図示される。また、平均残差(mean residuals)も、各従属変数についてグループごとに図示される。

POBS 以下の図が印刷される。観測値－標準化残差の図、予測値－標準化残差の図、標準化残差に対する、normal および detrended normal probabilityの図。

SIZE 縦軸および横軸の長さを変更する。縦軸の長さ25行、横軸の長さ40桁が標準値となっている。

(15) WRITEサブコマンド

[書式] WRITE

各種のデータを外部ファイルに出力する。外部ファイルの指定は、PROCEDURE OUTPUTコマンドで行なう。MANOVAは、6種類のレコードを外部ファイルに出力する。タイプ1には、デザインに関する情報が含まれている。タイプ2および3には、セルのコードとセルごとの平均に関する情報が含まれている。タイプ4にはセルごとの度数が含まれている。タイプ5には、セル内の誤差相関行列が含まれている。タイプ6には、セル内の標準偏差が含まれている。

(16) READサブコマンド

[書式] READ

WRITEサブコマンドで出力された各種のデータを読み込む。

19.3 オプション

- 1 標準処理では、変数のいずれか1つにでも欠損値のあるケースは計算から除外されるが、本オプションにより、欠損値をもつケースも計算に含められる。

20 BOX-JENKINS：時系列データの解析

このサブプログラムはBox and Jenkins(1976)の手法によって一変数時系列データへの統計モデルのあてはめとそのモデルにもとづく予測を行うものである。時系列解析であるため、ケースの順序に意味があり、データは等間隔で並んでおらねばならず、しかも欠損値は許されない。対数変換、累乗変換、季節差分および非季節差分などのデータ変換機能が準備されている。時系列データのモデル化の作業は一時的モデルの同定、パラメータの推定、および予測という三段階で構成され

ており、適切なモデルが確定するまで最後の段階は実行できない。また、バッチモードでのこの作業を容易にするため、十分な情報が容易に高率よく入手できるよう配慮されている。

20. 1 一般書式

BOX-JENKINS VARIABLESサブコマンド/

[/系列変換サブコマンド][モデル指定サブコマンド]

[/アルゴリズム指定サブコマンド][初期値指定サブコマンド]

[/予測指定サブコマンド][表示指定サブコマンド]

[/IDENTIFY][/ESTIMATE][/FORECAST]

上記のように、変数および解析段階の指定のほかに、各段階ごとに必要な情報を入力するためと出力要求のためのサブコマンドが用意されている。ここで使用されるキーワードは先頭三文字のみ有効で、四文字目以降は省略できる。上記の一般書式で注意すべきことは、解析段階の指定は他のすべてのサブコマンドの後でなくてはならないことと、少なくとも一つは解析段階を指定しなくては成らないことである。

20. 2 解析段階の指定

以下の3キーワードのうち少なくとも一つを指定する。この指定は一つの処理についての指示の終りと見なされるので、必ず他のすべてのサブコマンドの後に置かなくてはならない。

IDENTITY モデルの同定。

ESTIMATE パラメターの推定。

FORECAST モデルにもとづく予測。

20. 3 サブコマンド

20. 3.1 VARIABLESサブコマンド

[書式] VARIABLES=変数リスト

等間隔で観測された時系列データの変数を指定する。変数は最大10個まで指定できるが、それぞれ個別に分析されるだけで、変数間の関係やラグなどの解析はされない。

20. 3.2 系列変換サブコマンド

データの分散を全期間にわたって一定にするための対数または累乗変換、季節的および非季節的な非定常値を除くための差分、およびこれらに付随した季節を作る期間および自己相関係数などの計算のためのおくれの指定を含む。

(1) LOGサブコマンド

[書式] LOG=[$\left\{ \begin{array}{l} \text{定数} \\ 0 \end{array} \right\}$]

新変数値 = $\ln(\text{旧変数値} + \text{定数})$ により自然対数に変換する。

(2) POWERサブコマンド

[書式] POWER=(n[, {定数}])
0

新変数値=(旧変数値+定数)ⁿ により累乗変換する。

(3) DEFERENCEサブコマンド

[書式] DIFERENCE=階数1[THRU 階数2 [BY {増分}]]
1

非季節的差分を行う。THRUを使用することにより複数のモデルを指定できる。一般には二次ないし三次の差分で定常化することができる。推定または予測の段階でTHRUを使用すると、それぞれのモデルごとに各種の計算結果が出力されることになるので、同定の段階で限定しておくが望ましい。

(4) SDIFERENCEサブコマンド

季節的差分を行う。書式はDEFERENCEサブコマンドと同様であるが、これに続いて、PERIODサブコマンドが必要である。

(5) PERIODサブコマンド

[書式] PERIOD={期間}
1

季節的差分を行うための季節を構成する期間を指定する。単位は観測値のそれとする。季節的差分を行う時は必須である。

(6) LAGサブコマンド

[書式] LAG={おくれ}
25

自己相関係数、自己共分散、偏自己相関係数、残差の自己相関係数の計算で使用するおくれの最大値を指定する。

20. 3.3 モデル指定のサブコマンド

(1) P, Q, SP, SQサブコマンド

[書式] {P }={パラメータ 1 } [THRU パラメータ 2]
{Q } 0
{SP}
{SQ}

自己回帰(Pサブコマンド)、移動平均(Qサブコマンド)、季節的自己回帰(SPサブコマンド)、季節的移动平均(SQサブコマンド)のパラメータの数を指定する。THRUにより幅で指定し、複数の指定を同時に行うことができるが、推定または予測の段階(20.2参照)でこれを行うと、大量の出力があるので注意が必要である。

(2) MALAG, ARLAGサブコマンド

[書式] {MALAG}=パラメータ1, パラメータ2, ...
{ARLAG}

移動平均(MALAG)、自己回帰(ARLAG)パラメータを非連続的に指定する。P, Qサブコマンドの

代用として用いることができ、より柔軟な指定ができる。

次の2例はまったく同じ効果となる。

[例1] Q=3

[例2] MALAG=1, 2, 3

(3) CONSTANT, NCONSTANTサブコマンド

[書式] {CONSTANT}
 {NCONSTANT}

定数項のあてはめを行う (CONSTANT) か、行わない (NCONSTANT) かを指定する。

(4) CENTER, NCENTERサブコマンド

[書式] {CENTER}
 {NCENTER}

NCENTERでは、推定に先立って変数値を平均値の周囲に集中させないようにする。CENTERとすると、全体の定数項の影響を考慮しないことになる。

20. 3.4 アルゴリズム指定サブコマンド

パラメータ推定のための計算における収束判定条件、最大反復数、後方予測数、計算途中での検定など、アルゴリズムを指定するもので、大部分は特に指定しなくとも十分な結果が得られる。

(1) ITERATEサブコマンド

[書式] ITERATE=反復数

推定値の計算ルーチンが収束するまでに使用する最大反復数で、省略時は40となる。

(2) BFRサブコマンド

[書式] BFR=後方予測数

推定値の計算で使用する後方予測値で、省略時は0となる。

(3) TEST, NTESTサブコマンド

[書式] {TEST}
 {NTEST}

TESTでは、収束のための繰り返しごとに推定値の検定を行い、NTESTではこれを行わない。

(4) FPRサブコマンド

[書式] FPR=経過出力間隔

収束のための反復計算の途中経過を印刷する間隔で、省略時は5回の反復ごとに印刷される。

(5) PCON, PP, PQ, PSP, PSQサブコマンド

[書式] サブコマンド名=(進み幅1, 進み幅2, ...)

最終推定値の計算ルーチンにおいて使用する進み幅を指定するもので、省略時はすべて0.1となる。この指定はそれぞれ、P, Q, SPおよびSQサブコマンドで指定したパラメータの数と対応していないといけない。

[書式] サブコマンド名=(収束条件1[, 収束条件2,...])

20. 3.5 初期値推定サブコマンド

(1) ICON, IP, IQ, ISP, ISQサブコマンド

20. 3.6 予測指定サブコマンド

(1) ORIGINサブコマンド

(2) LEADサブコマンド

(3) CINサブコマンド

(4) FCON, FP, FQ, FSP, FSQ

20. 3.7 表示指定サブコマンド

大阪大学大型計算機センターニュース

統計を指示する。

(1) PRINTサブコマンド

[書式] PRINT=[ACF] [PACF] [ACVF] [SER] [TSER] [DSER]
[RESID] [RACF]

予測値のほかに追加して印刷する数表を選択する。

- ACF 自己相関関数。同定段階で印刷される。
PACF 偏自己相関関数。同定段階で印刷される。
ACVF 自己共分散関数。同定段階で印刷される。
SER 時系列（原系列）。同定段階で印刷される。
TSER 定数または累乗変換された時系列。同定段階で印刷される。
DSER 差分された時系列。同定段階で印刷される。
RESID 残差系列（モデル推定時のみ）。推定段階で印刷される。
RACF 残差の自己相関関数。推定段階で印刷される。

(2) PLOTサブコマンド

[書式] PRINT=[ACF] [PACF] [SER] [TSER] [DSER]
[RESID] [RACF] [FCF] [FLF] [CIN]

作成するグラフを選択する。キーワードのうち、ACF～RACFはPRINTサブコマンドと同じである。

- FCF 予測関数。予測段階で作成される。
FLF 固定先行予測。予測段階で作成される。
CIN 信頼区間。予測段階で作成される。

21 NPAR TESTS

変数の分布に関する種々のノンパラメトリック・テストを行なう。例えば、ある変数が正規分布であるか、二つの変数が同じ母集団からの標本であるか、などのテストを行う。NPAR TEST では次のようなサブコマンドが14種類のテストに対応している。

CHI-SQUARE 等分布または指定した分布に対する適合度のテストで、名義尺度以上の変数が対象。

K-S(その1) Kolmogorov-Smirnov one sample テストで、等分布、正規分布、ポアソン分布への適合度をテストする。間隔尺度以上が対象。

RUNS 連テスト。2 値変数または順位尺度以上が対象。

MCNEMER 2 変数間の独立性の検定で、2 値変数が対象。

SIGN 2 変数が同一母集団からのものかの検定で、順位尺度以上が対象。

WILCOXON 同上。

COCHRAN 多変数間の独立性の検定で、2 値変数が対象。

FRIEDMAN k個の変数が同じ母集団からのものかの検定で、順位尺度以上が対象。

MEDIAN k個のグループが同じちゅうおうちを持つ母集団からのものかの検定で、順位尺度以上が対象。

M-W Mann-WhitneyのUテスト。二つのグループが同一母集団からのものかの検定で、対象は順位尺度以上。

W-W Wald-Wolfowitzの連テスト。同上。

K-S(その2) Kolmogorov-Smirnov two sample テスト。同上。

MOSES 二つのグループの幅の違いをテストするもので、順位尺度以上が対象。

K-W Kruskal-Wallisの一元配置分散分析で、k個のグループが同じ母集団からのものかを検定する。対象は順位尺度以上。

21. 1 一般書式

NPART TESTS [CHI-SQUAREサブコマンド]/[EXPECTEDサブコマンド]]
 [/K-Sサブコマンド]/[RUNSサブコマンド]/[BINOMIALサブコマンド]
 [/MCNEMARサブコマンド]/[SIGNサブコマンド]/[WILCOXONサブコマンド]
 [/COCHRANサブコマンド]/[FRIEDMANサブコマンド]/[KENDALLサブコマンド]
 [/MEDIANサブコマンド]/[M-Wサブコマンド]/[K-Sサブコマンド]
 [/W-Wサブコマンド]/[MOSESサブコマンド]/[K-Wサブコマンド]

サブコマンドを複数並べて同時にいくつかの検定を行なっても、また同じサブコマンドを2回以上使用してもよい。

21. 2 サブコマンド

(1) CHI-SQUARE, EXPECTEDサブコマンド

[書式] CHI-SQUARE=変数リスト [(最小, 最大)]/
 [EXPECTED={EQUAL
 {f1, f2, ..., fn}}]

最小, 最大は整数値を用いて変数の範囲を指定する。変数が整数でないものに用いると、すべての値は整数値に丸められる(小数点以下を切捨てる)。EXPECTEDにより、テストすべき分布を相対頻度の形で指定する。これを省略またはEQUALを選んだ時は、一様分布が仮定される。

[例1] CHI-SQUARE=A(1, 5)/EXPECTED=12, 3*16, 18

この例では、変数Aのカテゴリー1～5の相対頻度が12, 16, 16, 16, 18の仮説が検定される。

(2) K-Sサブコマンド (1 グループ)

[書式] K-S({UNIFORM} [, パラメータ])=変数リスト
 {NORMAL}
 {POISSON}

変数リストは分布をテストされる変数名。仮定する分布はキーワードUNIFORM（一様分布）、NORMAL（正規分布）、POISSON（ポアソン分布）の中から一つを選ぶ。各キーワードにつづくパラメータは次のとおりである。パラメータを省略すると標本値が使用される。

UNIFORM 最小と最大

NORMAL 平均と標準偏差

POISSON 平均

（３）RUNSサブコマンド

〔書式〕 RUNS({MEAN })=変数リスト
 {MEDIAN}
 {MODE }
 {値 }

MEAN, MEDIAN, MODEまたは指示する値を連テストの分岐点とする。

（４）MCNEMAR, SIGN, WILCOXONサブコマンド

〔書式〕 {MCNEMAR }=変数リスト [WITH 変数リスト]
 {SIGN }
 {WILCOXON}

使用する変数とその組合せを指示する。

〔例２〕 MCNEMAR=A B C (A, B) (A, C) (B C) の組合せでテストされる。

（５）COCHRAN, FRIEDMANサブコマンド

〔書式〕 {COCHRAN }=変数リスト
 {FRIEDMAN}

（６）MEDIANサブコマンド

〔書式〕 MEDIAN[{値}]=従属変数リスト BY 独立変数(値1, 値2)

独立変数によりグループ分けをする。値1のグループから値2までの値2-値1+1のグループに分け、従属変数のそれぞれについてテストする。MEDIANの後の値は仮定するMEDIANで、省略すると標本値を使用する。

（７）M-W, K-S(2グループ), W-Wサブコマンド

〔書式〕 {M-W}=従属変数リスト BY 独立変数(値1, 値2)
 {K-S}
 {W-W}

MEDIANサブコマンドとほぼ同じであるが、グループは値1と値2の二つだけとなる。

（８）MOSESサブコマンド

〔書式〕 MOSES(n)=従属変数リスト BY 独立変数(値1, 値2)

値1は制御グループを識別する値、値2は被検グループを識別する値。nは、値の小さいものと大きいものを削除するときの各ケース数を指示するもので、省略するとそれぞれ5%が削除される。

（９）K-Wサブコマンド

〔書式〕 K-W=従属変数リスト BY 独立変数(値1, 値2)

MEDIANサブコマンドと同様で、値2-値1+1のグループに分ける。

21. 3 オプション

- 1 欠損値の解除。
- 2 使用する変数のうち一つでも欠損値があればそのケースを除外する。
- 3 テストされる変数の組合せの特殊な扱い。

VARA VARB VARC ... に対しては、(VARA, VARB) (VARB, VARC)...、VARA VARB WITH VARC
VARD に対しては、(VARA, VARC) (CARB, VARD)のように組合せる。

- 4 ケース数が多くて作業領域が不足するときサンプリングを行う。ただし、RUNSテストではこのオプションは無視される。
- 5 結果出力の印刷幅を75桁とする。

21. 4 追加統計

- 1 記述統計（平均、最小、最大など）を出力する。

参考文献

- 1) 新版SPSS* I 基礎編、東洋経済新報社、1966。
- 2) SPSS* User's Guide, McGraw-Hill.
- 3) SPSS* Introductory Statistics Guide, McGraw-Hill.
- 4) SPSS* Advanced Statistics Guide, McGraw-Hill.

著者

小野寺義孝	東海女子短期大学初等教育専攻
竹村 和久	同志社大学文学部
吉村 英	光華女子短期大学家政科
山本嘉一郎	光華女子短期大学家政科