

44-S006

統計解析用予測モデルJUMPS

(標準アプリケーション・システム
の設計とプログラムの研究開発)

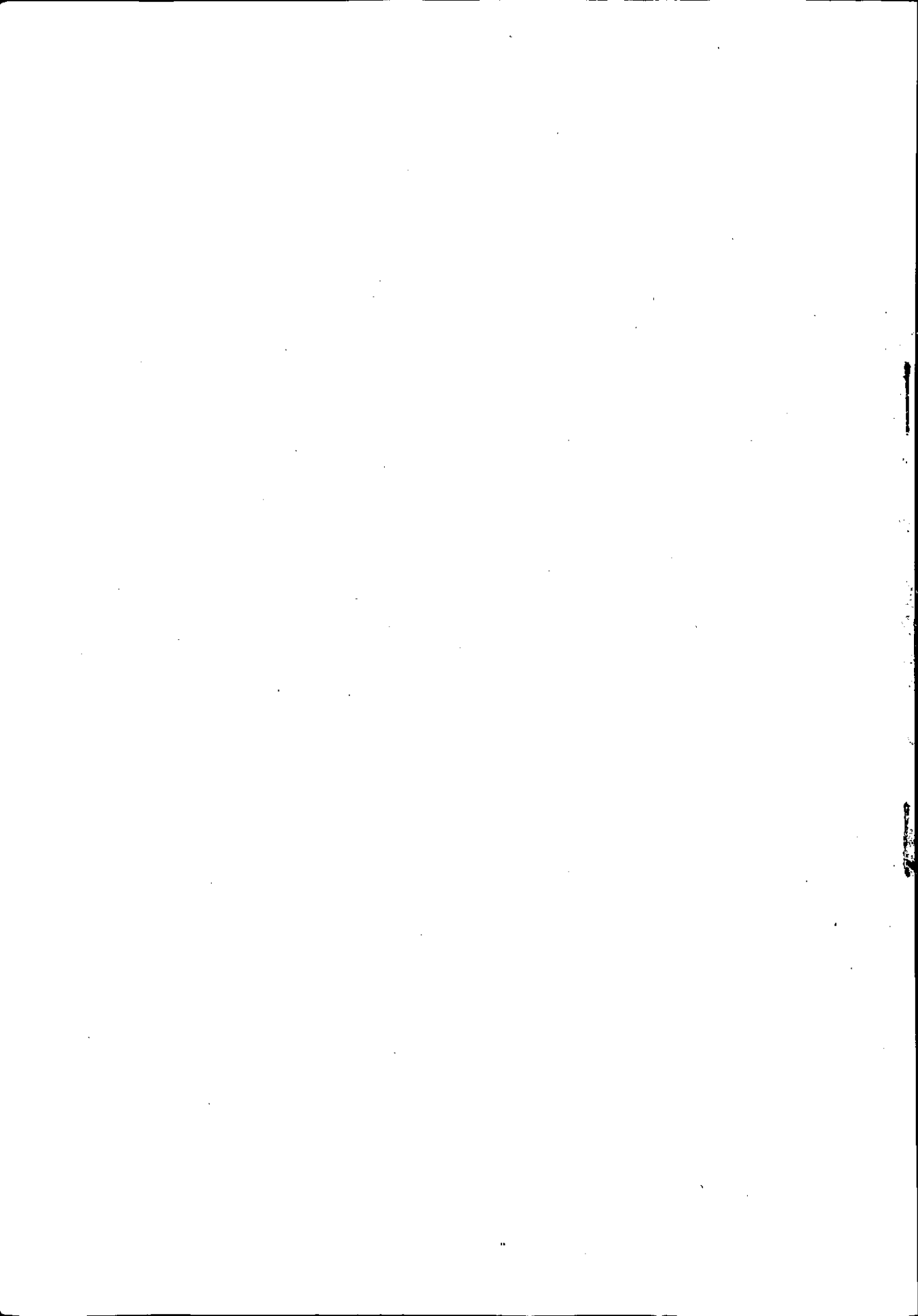
昭和45年6月

JIPDEC

財団法人 日本情報処理開発センター



この事業は、日本自動車振興会の機械工業振興資金による「昭和
44年度情報処理に関する調査・研究補助事業」のうち「標準アプ
リケーション・システム設計とプログラムの研究開発」の一部とし
て実施したものであります。



序 に 代 え て

最近、予測、統計、数理計画のためのアプリケーションプログラムとして、シミュレーションなどの手法を主体的にもりこんだ汎用プログラムの開発が盛んに行なわれておりますが、この報告書は統計解析、データ分析および予測の一連の数学的処理を目的とし、グラフィックディスプレイを入出力装置として使用することにより解の妥当性、最適性を試行錯誤によって会話的に求める統計解析用予測モデルJUMPS (J I P D E C ' S Universal Mathematical Programming System) を開発するための基礎研究をとりまとめたものであります。

ここに本研究実施にご尽力を賜った早稲田大学新沢雄一教授に心より感謝の意を表しますとともに、本報告書が各方面に利用され、わが国情報処理産業発展の一助として寄与できますようお願いいたします次第であります。

昭和 4 5 年 6 月

財団法人 日本情報処理開発センター

会長 難 波 捷 吾

はじめに

本報告書は、JUMPS 70 (JIPDEC'S Universal Mathematical Programming System in 1970) の基本設計書である。

今日、コンピュータのハードウェアの性能の向上はもとより、コンピュータ利用の高度化と相俟ってソフトウェアの開発・充実は目ざましいものがある。すでにいくつかのすぐれたシステム・プログラムや言語プログラム、アプリケーション・プログラムが開発され、情報検索、データ分析、モデル・ビルディング、予測、シミュレーション、最適化問題などはもとより経営、分析、設計、コントロールの各分野に広く用いられている。

だがすでに開発された各種のプログラムは、それぞれすぐれた特長をもつとはいえ、個別問題の処理に対して人間が望ましいと考える結論に到達するまでに処理手順を試行錯誤によって繰返し、多大の労力と時間を掛けたり、あるいはシステムとして一応まとまってはいても利用に柔軟性を欠いたり、あるいはまた利用者にとって取り扱いがきわめて難渋であったりして、現実の問題処理に必ずしも適しているとはいえない。

その主たる理由は次のように指摘することができるであろう。

1. コンピュータと人間が対話を通じて人間の要求する目標を総合的に解決するプログラム・システムが完成していないこと。

この問題は、すでに第3世代のコンピュータが普及しつつあるにもかかわらず、第3世代コンピュータの中心的な利用・開発の課題である、(1)通信システムとコンピュータとの結合、(2)会話形式による多重同時並行処理が言語プログラムやアプリケーション・プログラムに充分応用されていないことによるのである。

2. 何らかの目標をもってコンピュータを利用する利用者とプログラム・システム開発者との考え方が必ずしも一致していないこと。

この問題は、とくにプログラム・システム設計者が現実の問題処理の経験が少ないかあるいは全くないために、利用者の立場や何が問題であるかを理解していないことや、あるいはコンピュータの性能と問題処理とが一致していないためによるのである。

3. 人間とコンピュータの長所を生かすような工夫がなされていないこと。

この問題こそがあるいは最重要な問題であろう。いうまでもなくコンピュータの最大の特長は(1)記憶の大容量性、(2)演算の高速性、(3)処理の正確性にあり、人間の長長は判断、推理、創造性にある。コンピュータに委ねるべき領域と、人間が主導権を持つべき領域とが確定してはじめて真の意味の Man-Machine System が作られるのである。

JUMPS は上記の諸点に留意し、その名称の示す通り汎用性のある予測手法、数理計画法のプログラム・システムであると同時に、第3世代以後のコンピュータを前提として Jump 機能を full に

活用し、人間がきわめて自由に問題解決のために模型をつくり、予測、シミュレーションその他の処理を一貫して行なうことができるよう設計されている。

本報告書では第1章において、JUMP Sの基本設計の目標と特長および機器構成について述べ、第2章においてJUMP Sを構成している基本的な構成要素の概念をFlow, Module, Phase, 応答形式に関連して説明し、第3章において主題分析の考え方およびFileの構成について明らかにする。第4章のデータ抽出においては変数のType(内生, 外生, 未決定, ダミー変数)を分類し、Priority, Sector, Function, Time Lag, Period, Unitなどの諸定義および応答の段階に必要な各種テーブルの指定について述べる。第5章データ分析では、データの時系列に関する分析、方法と指定条件およびGPF Sの考え方および解析内容が書かれてある。第6章変数間の関係(Interrelation of Variable)では、定義形式と応答形式によって相関分析および因子分析が行なえるようにしている。第7章模型作成(Model Building)では直接最小自乗法、間接最小自乗法をはじめ逐次最小自乗法、二段階最小自乗法、三段階最小自乗法、操作変数法および最尤法の指定の仕方を説明し、またとくに方程式の選択と適度認定問題の自動化について述べている。第8章推定と評価(Estimation and Evaluation)では、上記の各種模型の推定に関し各テスト条件の指定の仕方を述べる。第9章予測(Forecasting)では内挿、外挿の指定と予測期間の指定について、定義形式と応答形式に分けて述べている。第10章シミュレーションでは模型、制御変数、比較条件および選択条件の指定に関して説明し、各種分布の乱数の指定について述べている。第11章最適化問題では連立方程式、線型計画法、二次計画法の定義形式による指定方法を説明している。

対話言語としてのJUMP Sはこの報告書の段階で完結したものではない。きわめて短期間のうちに基本設計を行ないまた現状では、必ずしも特定のコンピュータを想定して設計したわけではないので、具体的に詳細なプログラムを組む段階において補足、修正、変更が行なわれることはやむを得ない。しかし、その場合にもJUMP Sの基本理念は失わぬようにプログラム設計を行ない、我が国のコンピュータのソフトウェア開発の一助にもなりうれば本システムの基本設計を行なった者にとってこれに優る喜びはない。

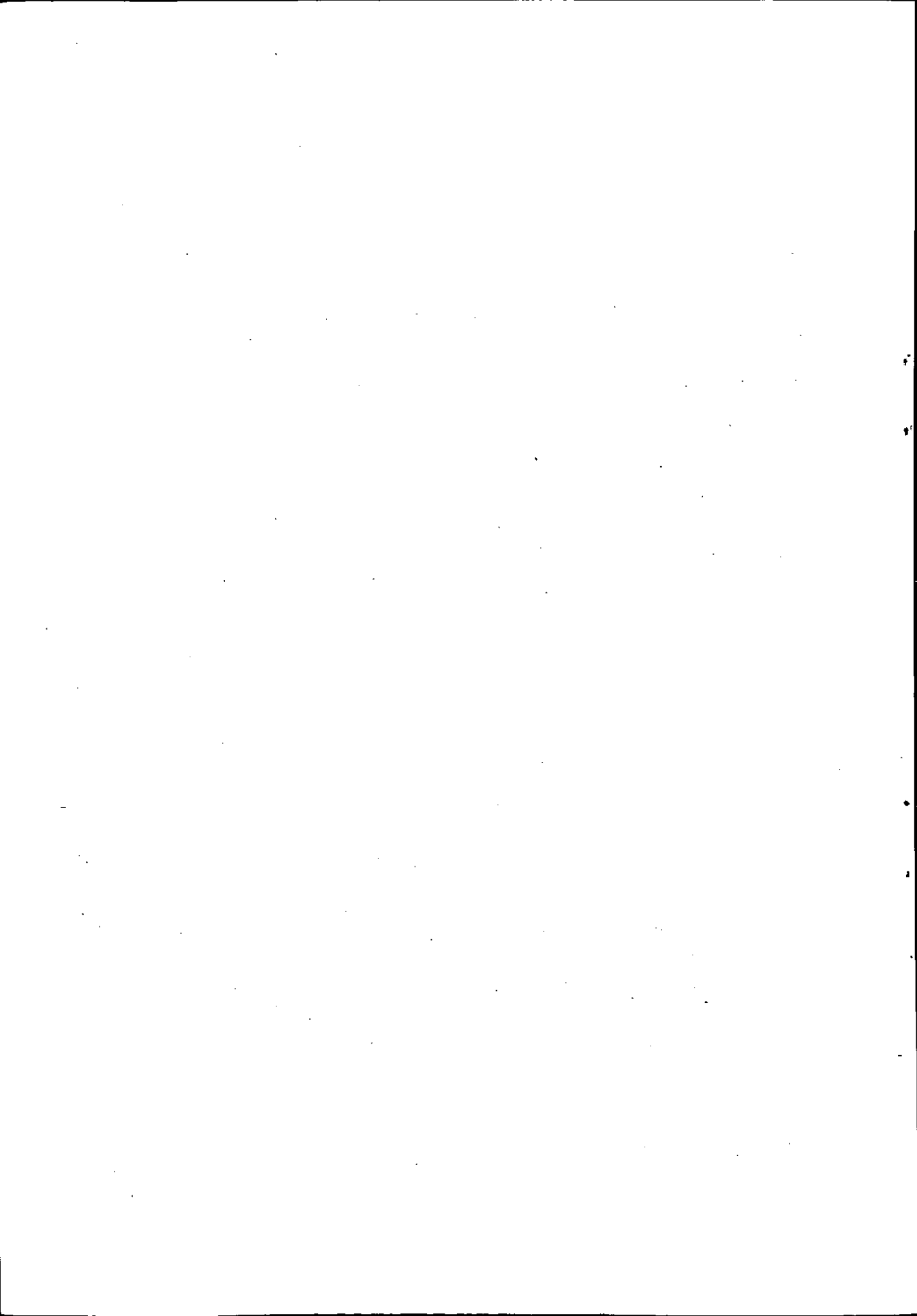
JUMP Sの基本設計において常に変らぬ立場で設計者の自由を許し、絶えず励ましの言葉を与えて下さった日本情報処理開発センターの開発課長山本欣子氏、同課の伊藤哲史氏および榎本晃氏に心から御礼申し上げるとともに、設計の段階でたゆみない討論と資料および原稿整理にもおしみな助力を与えてくれた早稲田大学・大学院理工学研究科・博士課程の星野供二氏に心から感謝の言葉を捧げたい。

昭和45年5月

新 沢 雄 一

目 次

1. JUMPS基本設計の目標と特長	1
1.1 目 標	1
1.2 特 長	1
1.3 システムおよび機器構成	2
2. JUMPSの構成要素	5
2.1 BUF(Basic Unit Flow)	5
2.2 BDF(Basic Dialogue Flow)	6
2.2.1 Orientation Phase	9
2.2.2 Execution Phase	11
2.2.3 JUMPS Phase	16
2.3 General STEP Flow	18
2.4 Learning Function	21
3. 主題分析(Theme Analysis)	25
3.1 JUMPSにおける主題の構造分析	25
3.2 JUMPSのInformation Retrieval	30
3.3 FILEの構成	32
4. データ抽出(Data Retrieval)	39
4.1 データの抽出における定義形式	40
4.2 データの抽出における応答形式	44
5. データ分析(Data Analysis)	47
5.1 データの分析における定義形式	49
5.2 データの分析における応答形式	55
6. 変数間の分析(Interrelation of Variables)	58
6.1 変数間の分析における定義形式	59
6.2 変数間の分析における応答形式	60
7. 模型作成(Model Building)	61
7.1 Model Buildingの内容	61
7.2 Model Buildingの定義形式	77
8. 推定と評価(Estimation and Evaluation)	78
9. 予 測(Forecasting)	96
10. シミュレーション(Simulation)	97
11. 最適化問題(Optimization)	108



1. JUMPS基本設計の目標と特長

1.1 目 標

JUMPSの基本設計を行なうにあたって、とくに次の諸点を目標に設計した。

1) 使い易さ

専門家でも非専門家でも、対話形式、定義形式およびJUMPを通じて自由に修正、加工、変更がユーザの理解に合わせて即座に行なえる。

2) 汎 用 法

予測手法、数理計画法その他について現状の水準を網羅し、単に経営あるいは経済の問題ばかりでなく、数学、物理、化学、生物学、工学など広い範囲に利用できる。

3) 対話言語 (Dialogue Language) の作成

コンピュータのソフトウェアの開発は個々のサブルーチンから始まって、アッセンブラ、コンパイラが開発され、また個有の問題に対してはシミュレーション手法、パッケージが作られているが、FORTRAN、GOBOLなどのコンパイラにしてもStatement Languageであって、コンピュータと対話しながら決定を加えてゆく対話言語 (Dialogue Language) ではない。

JUMPSはディスプレイ装置あるいは他の装置を駆使して対話を行ないながら問題の解決を行なえるよう設計されている。

1.2 特 長

JUMPSの設計にあたって上記の目標のもとに設計が行なわれたが、その特長は次の諸点である。

1) システムの拡張性と柔軟性

JUMPSはここに提出する内容に終るものではない。徹底したモジュール形式によりプログラムが作成され、JUMP機能と相俟ってJUMPシステムは必要に応じて拡張が可能であり、またコンピュータの性能に合わせてシステムの大きさを決めることができる。

2) 徹底した対話形式の活用

JUMPSは本論の各章・各節で説明しているように処理方法、処理手順のみならず選択、修正、変更などコンピュータと人間とが対話形式によって問題の分析を深めることができる。

3) 定義形式と応答形式の自由選択

JUMPSを利用する際に、変数、方程式あるいは方程式群、諸条件に対して予定した指定を専門家の立場から定義形式で与えることもできれば、分析方法を熟知していない利用者にとってもコンピュータとの対話を通じて分析を進めることができるように応答形式と定義形式を各Stageおよび各Stepで自由に選ぶことができる。

4) Jump機能の活用による融通性

JUMPSは各応答Stage, 応答Stepおよび応答SubstepにおいてRETURN, REPEAT, LEAP, CARRY-ONを指定することによって, すでに実行してきたStage, Step, Substepに返ったり, あるいはLOCKしたStage, Stepを繰り返したり, または不要なStage, Step, Substepを飛び越したり, あるいは順次実行することができ, 分析を深めることができる。

5) 演算処理および手順の自動化

JUMPSは各Stage, 各Step, 各Substepにおいてデータ, 変数, 方程式, 分析結果など前提条件が不足していたり欠除している場合にはWarningを利用者に与えて, 前提条件を満たすために自動的に必要なStage, Step, SubstepへAUTO-RETURNし, 前提条件が満たされると再びその処理段階へもどる。

さらに方程式と変数の斉合成や認定条件の適度性はOPTIMAL, SELECT, ALLなどの指定によって自動的に模型を作成する。模型が自動的に適度認定になるようなアルゴリズムが組まれているが, この例は他に例をみない。そのほか処理手順についてもJUMPS独得のアルゴリズムが組まれている。

6) 既成モデルとの自由結合

模型(モデル)を作成する場合に, 主題分析によってデータ抽出の段階から模型を作成してゆくことは当然であるが, すでに作成されている模型が存在する場合は, 応答形式あるいは定義形式によってこの既成の模型を利用できるように組まれている。

7) Learning機能による処理手順・範囲の自動的処理

一般のプログラムは殆んどLearning機能を備えていないため, 特定問題の専門家も非専門家も等しく同一手順で模型を作成しなければならない。このことは往々にして専門家にとっては不十分な結果を与え, 非専門家には理解することが困難な過剰な結果を与えることになるのである。このような弊害を取り除き, また要求に応じて分析を深めることができるように自動的にコンピュータが人間の専門程度を常に測定し, 処理手順や解答の範囲を処理するような機構が備えられている。

8) 主題分析におけるKeyword構成

主題分析によって検索する場合に, File形式を一般化し融通性と拡張性をもたせなければならない。JUMPSの現段階においては意味論的なDocumentationは行なってはいないが, 主題をKeynounとKeyverbに分け, その合成としてのKeywordの作り方に工夫がしてある。

1.3 システムおよび機器構成

JUMPSは, 可能なかぎり各社各種のコンピュータにかけられるように考えて, コーディング言語としてはFORTRAN JISレベル7000を使用することとする。しかしながら次の点は全くできないか, あるいは非能率的なことが起り得るので, この点についてはアセンブリ言語を使わざるを得ない。

全くできない機能

1. グラフィックディスプレイに対する機能。
2. 割込処理等のモニタあるいはエグゼクに対する連絡機能。
3. ダイレクト・アクセス・ファイルあるいはパーティションド・シーケンシャル・ファイルに対するアクセス

非能率的なことが起る場合

1. 各種テーブルが主として1文字あるいは1ビット単位で記憶装置を使いたい場合
2. ファイル設計上、シーケンシャル・アクセス以外のファイルが使えない事

上記の問題以外の各種処理はすべてFORTHANによってコーディングし、そのルーチンは各種のコンピュータで使用できるようにし、アセンブリ言語でコーディングしたプログラムは完全にサブルーチン化し該当部分さえ組みなおせば他のコンピュータにかけられるように作り上げる。

JUMPSのシステム設計を進めていく場合に、仮想のコンピュータを想定してシステムの開発を考えることも可能であるが、このようにすることは、応々にして机上の空論に終り、実際にプログラムを作成することができないものを作り上げてしまう場合がある。

このために、図1-1で示すように現存のHITAC-8400とHITAC-8811グラフィック・コンピュータに対してCGOS (Comprehensive Graphic Operating System) という汎用グラフィックオペレーティングシステムを開発中なので、そのシステムをそのまま利用して作り上げることを考えている。

CGOSの概説

CGOSにおけるコントロールプログラムシステムの特徴は、H-8400プロセッサ内のGJM (Graphic Job Monitor) と、H-8811プロセッサ内のGCH (Graphic Communication Handler) の両プログラムに代表される。

GJM及びGCHは、CGOS

のシステム・イニシエーション時にそれぞれ主記憶装置にローディングされ、DOS IIIエグゼクティブ、H-8811エグゼクティブ等の制御プログラムと協業して、CGOSの運用を管理する。

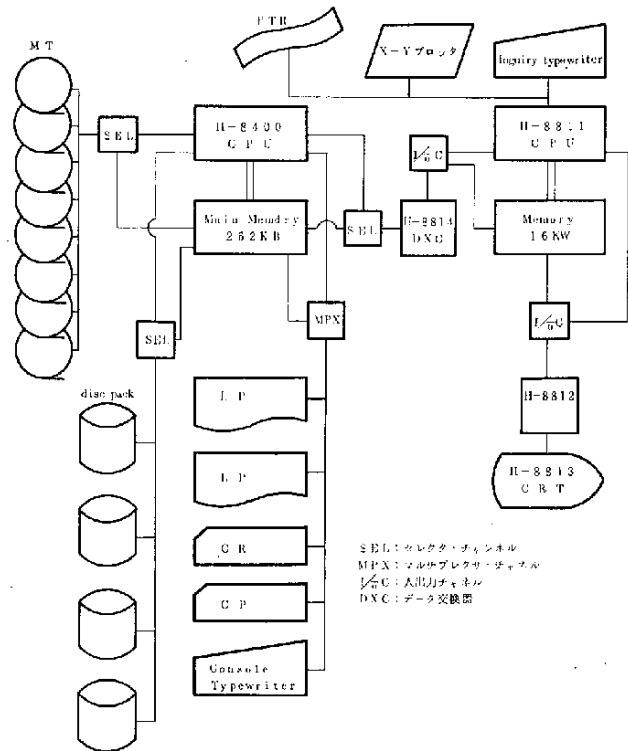


図1-1

CGOSのコントロール・プログラムシステムは従来H-8400付属のコンソールタイプライタで行っていた、DOSオペレーションのはずべてを、H-8813グラフィック・CRT・コンソールからのオペレーションで代行することができる。

GJMONは、H-8400プロセッサ内で、DOS IIIエグゼクティブのTで、スレーブモードスタとして常駐し、次の機能を果たす。

1. DXC経由による、H-8811プロセッサとの間の相互データ転送
2. DXC経由で受け取ったメッセージデータの解析
3. H-8813グラフィック・CRT・コンソールから要求のあったプログラムの起動、及びモニタリング
4. グラフィック・コマンドリクエスト
5. グラフィック・アプリケーションプログラムとの間の相互データ移送

GCHはH-8811プロセッサ内に常駐し、次の機能を果たす。

1. DXC経由による、H-8400 host computer との間の相互データ転送
2. DXC経由で受け取ったキャラクタ・メッセージの表示
3. DXC経由で受け取った濃縮図形データからディスプレイコマンド列への翻訳、表示
4. メッセージ・インディケータの表示と、オペレーション・メッセージ・ストリングの編集と転送
5. コマンド・コンパイルされた図形データに対する、ライトペンのピックアップの受付と、プリンキング処理
6. コマンド・コンパイルされた図形データのエッジ・シザリング処理
7. ファンクションキー(0~31)の割込み制御
8. ダイアル(1, 2)の割込み制御
9. トラッキング・クロスの表示と、ライトペンによる手書入力データの処理
10. CRT表示データのX-Yプロッタ上へのハード・コピー

CGOSとJUMPSの関係

CGOSは、概説に述べたように各種機能を持っているが、JUMPSはCGOSの下にあって、その機能を使いながら動くのである。使用する主な機能は次のとおりである。

1. JUMPSシステムのローディング
2. JUMPSシステムをオーバーレイして使用
3. 文字及びグラフの表示
4. 文字入力及び、文字列のピックアップ
5. Quit機能、Restart機能

これ等の機能は図1-2で示されるような形でJUMPSの各コマンドを通じて使用されることになる。

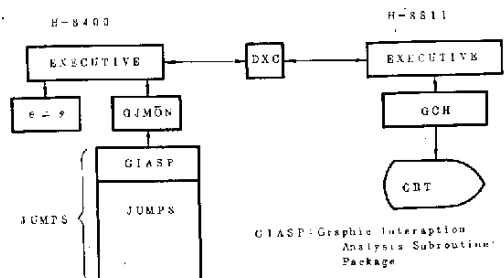


図1-2 CGOSとJUMPSの関係図

2. JUMPSの構成と要素

2.1 BUF (Basic Unit Flow)

JUMPSの特長の一つに、モジュール化されたBasic Unit Flowがある。BUFは図2-1のように次の3つの基本的なPhaseから成っている。

Phase I ORIENTATION Phase

Phase II EXECUTION Phase

Phase III JUMPS Phase

Phase IのORIENTATION Phaseでは、人間と機械との会話のうちにあるいはPhase I以前に発生した(または与えられた)条件を整理して、Phase IIのEXECUTION Phaseで何を対象に何を実行するかを決定する。

Phase IIのEXECUTION Phaseでは、Phase Iで与えられた条件にしたがって、数値解析、論理演算、統計処理などの各種演算とデータ蓄積、検索などのデータ処理およびI/O装置や補助記憶装置などの指定を行なう。このPhaseは問題によって下位レベルのPhase I, II, IIIすなわちORIENTATION, EXECUTION, JUMPに再分割される。

Phase IIIのJUMP Phaseは、すでに処理通過したStep, Stageへ戻るRETURNと、これから処理したいと考えるStep, Stageへ跳躍するLEAPとに分かれ、これらの指定を人間がシステム外から指定する場合、人間と機械との会話によって決定される場合、機械が一方向に選択して決定する場合とがある。JUMP Phaseは基本的にはRETURN機能とLEAP機能を持っているが、すでに実行したあるStageまたはStepから特定のStageまたはStepを全く同じ手順で繰り返す場合は、ProcessをLockしてREPEATを用いてもよい。REPEATはいわゆるおさらいであってLockされたProcessをたどるが、途中の判断を機械が人間に要求する場合は、JUMPによって任意のStageあるいはStepに飛ぶことができるとともに、CARRYとONを指定すればREPEATによって支配されたStageあるいはStepまでを自動的に進めることができる。

INITIALIZATIONはBUFの一部ではなく、機械がSTART命令によりBUFに入る前に、初期値、初期条件などの指定を行なう部分である。

ENDもBUFの一部ではなく、人間と機械の相互的な関係のすべての処理が終了し、機械が個々の最終段階を処理するPhaseである。ENDには、人間の与えた目的、判断、結果などについてのPrescriptionを含ましめる。

人間がJUMP (RETURN, LEAP, CARRY ON)あるいはREPEATの実行について判断できない場合演算、および処理を一時的に停止する命令としてHALTがある。HALTの場合は、JUMPあるいはREPEATの指定によって演算および処理が開始される。

演算を全く停止する場合はSTOPの指定を行なう。STOPによって機械は自動的にEND Phaseに入る。

2.2 BDF(Basic Dialogue Flow)

JUMPSにおける人間と機械との会話についての基本的なフローは図2-2に示してある。

右欄は人間の頭脳内の処理(判断, 指定, 選択)を表わし, 中央欄は人間と機械との会話を表わし, 左欄は機械の内部処理を表わしている。人間はこのシステム利用についての主題を与え, 対象をはじめ各Stage, Stepにおける具体的な内容の指定, 選択, 判断および訂正を行なう。機械は人間が提出した主題について分析を行ない, 可能性を示し, 診断を下すとともに各段階のExecutionを行なう。

INITIALIZATION Phase では人間(オペレータあるいは会話者)が機械に対してkeyboard上でmanualに

startの指令を行ない, 機械的Initializationの初期値あるいは諸条件をkeyboard, card, tape, 入力装置, その他の方法によって与える。機械はこれらの初期値あるいは初期条件によって会話モードに入る準備を行なう。

ORIENTATION Phaseは人間によるobjectの指定, 改訂および修正, 機械による分析, 可能性の提出および診断のDiagnosis Stepと, 機械によるAlternative Stages, Stepsの提出, および人間によるSelection Stepをふくむ。

EXECUTION PhaseはORIENTATION Phaseによって指定された目的にしたがって第i stage 第j stepの問題を処理するExecution Stepと, 機械によるDisplayの提出, 人間によるDisplayの選択, 結果の処理, 結果の提出というOutput Stepとに分けられる。

JUMP PhaseではReturn, Leap, Repeat, Carry on, Halt, Ssop の指定が行なわれる。

ORIENTATION Phaseにおけるobjectの内容は

1. このシステムを利用する利用者の主題

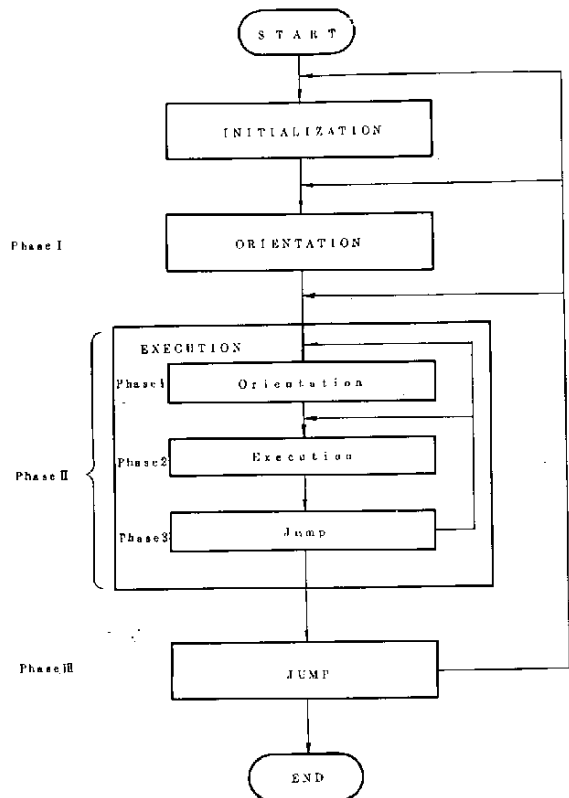


図 2 - 1

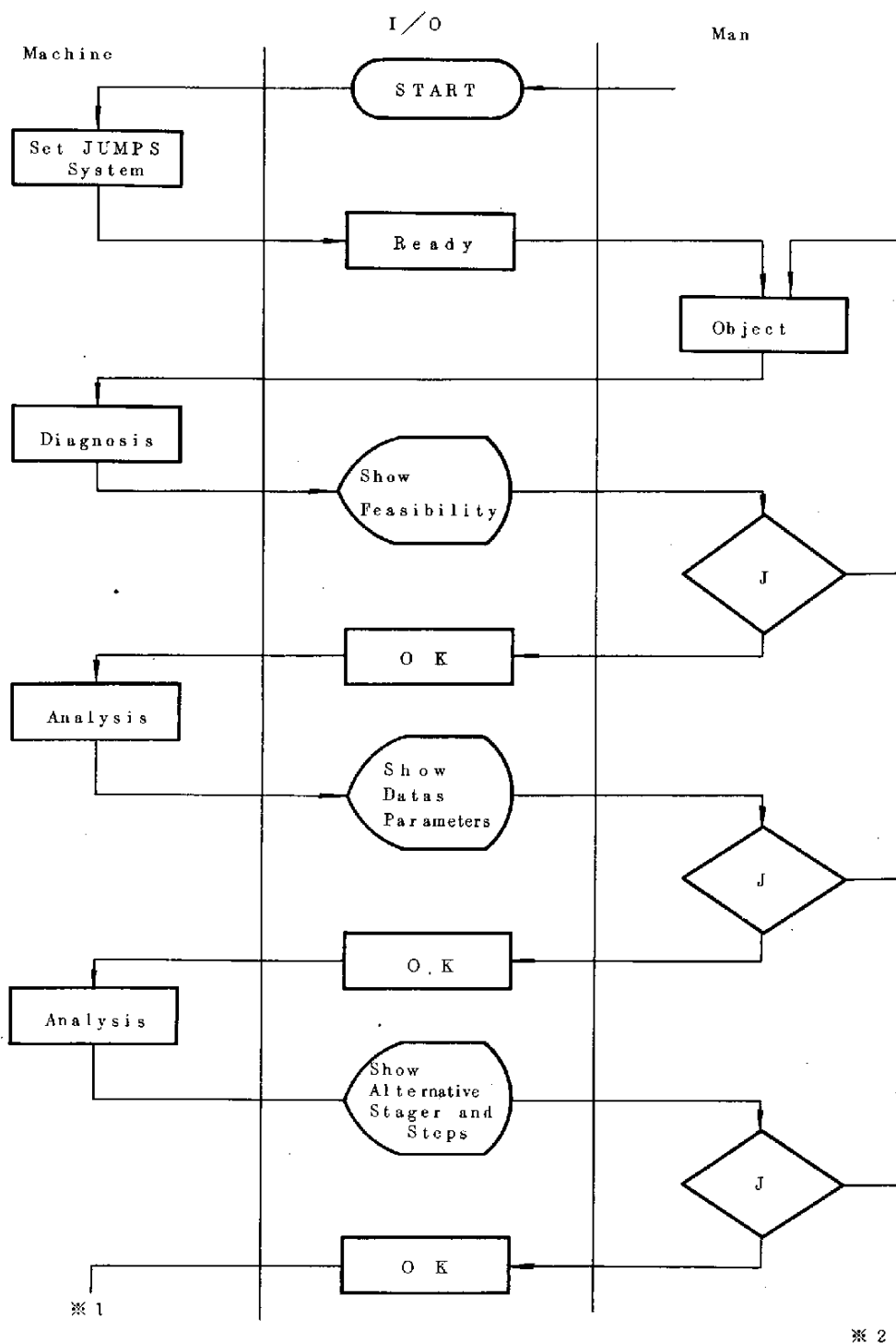


图 2-2(1)

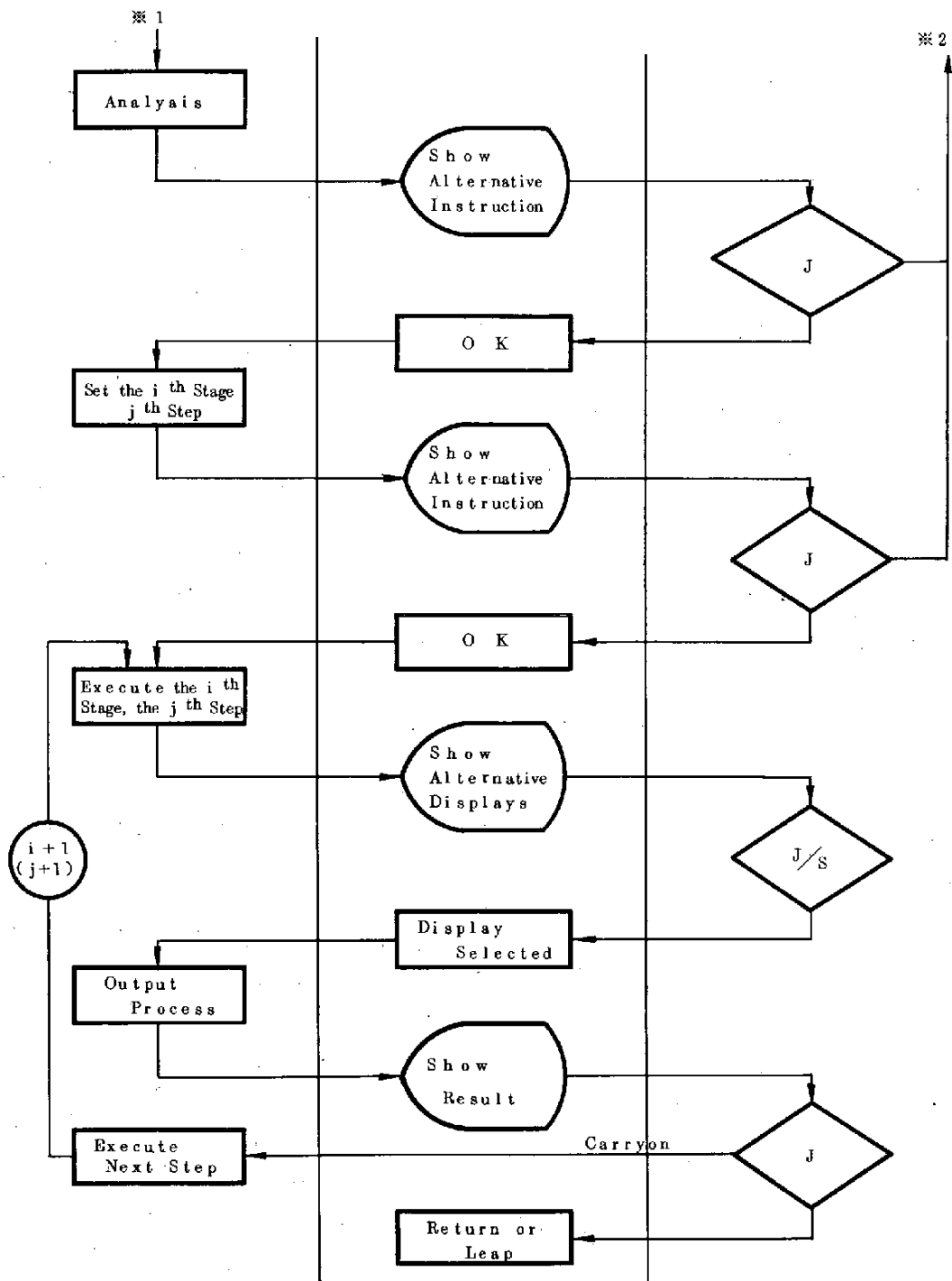


图 2-2(2)

2. 利用者の対象, データ, 方法などについての要求
3. 各 Stage あるいは Step において機械が質問する項目についての利用者の解答をふくんでいる。

object step のレベルは

1. 機械にあらかじめ組み込まれている objects を display し人間に選択させる。
2. ある程度の機械で処理できる objects の範囲を人間に予備知識として与えておき, 応答形式にて object を選択・決定する。
3. 全く Man-Machine の応答より object を選択・決定する。

前述のように JUMP S は Stage, Phase, Step からなる。Phase は Step の集合であり, Stage は Phase の集合である。

Step はその性質から, 次の 4 module に分けられる。

1. B D M: Basic Dialogue Module
2. T A I M: Theme Analysis Initial Module
3. I O D M: Input-Output Dialogue Module
4. J M: Jump Module

BDM は最も基本的な Module で Orientation Phase における Object Step の一部, Keynoun Step, Keyverb, Instruction Step, Execution Phase における Operation Step は全て BDM からなる。

BDM は図 2-3 に示すように人間と機械の基本的な対話の module であって, 人間の判断, 機械の diagnosis, ディスプレイ装置を介する input-output 部分から成っている。BDM 内部において人間が機械にある判断を与えると, 機械は分析・診断してその結果を示し, 人間は機械から表示された結果を見ながら Jump し, もし OK ならば次のボタンは機械に移り, 機械は変数, 式, 条件などを示すとともに人間にさらにデータ, 変数, 式, 条件などを附加する必要があるかどうかあるいは選択するかを尋ね, 次にボタンは人間に移り, 人間が OK という判断を行なうまで, データ, 変数, 式, 条件などの附加, 選択, 変更が行なわれ, 機械が分析・診断をつづける。

2.2.1 ORIENTATION Phase

TAIM は object step の一部で, 主題 (Theme) が与えられると Theme について量的属性をもつ Keynoun として適格であるかどうかを分析して結果を示す module である。これは BDM に似ているが, 人間が結果をみながら Theme の変更について Input するところが異っている。

Object Step の一部をなす BDM は, すでに TAIM で Theme の分析が行なわれるから Diagnosis 部分は必要とせず, 目的の質問については人間に質問する。人間は OK となるまで目的に修正を行なう。この Object Step の BDM においては, 条件の附加が機械から人間に対して質問され, 条件の変更がある場合は TAIM の Analysis へ戻る。

Keynoun step は一つの BDM からなる。この step では主題に適合する Data File, 変

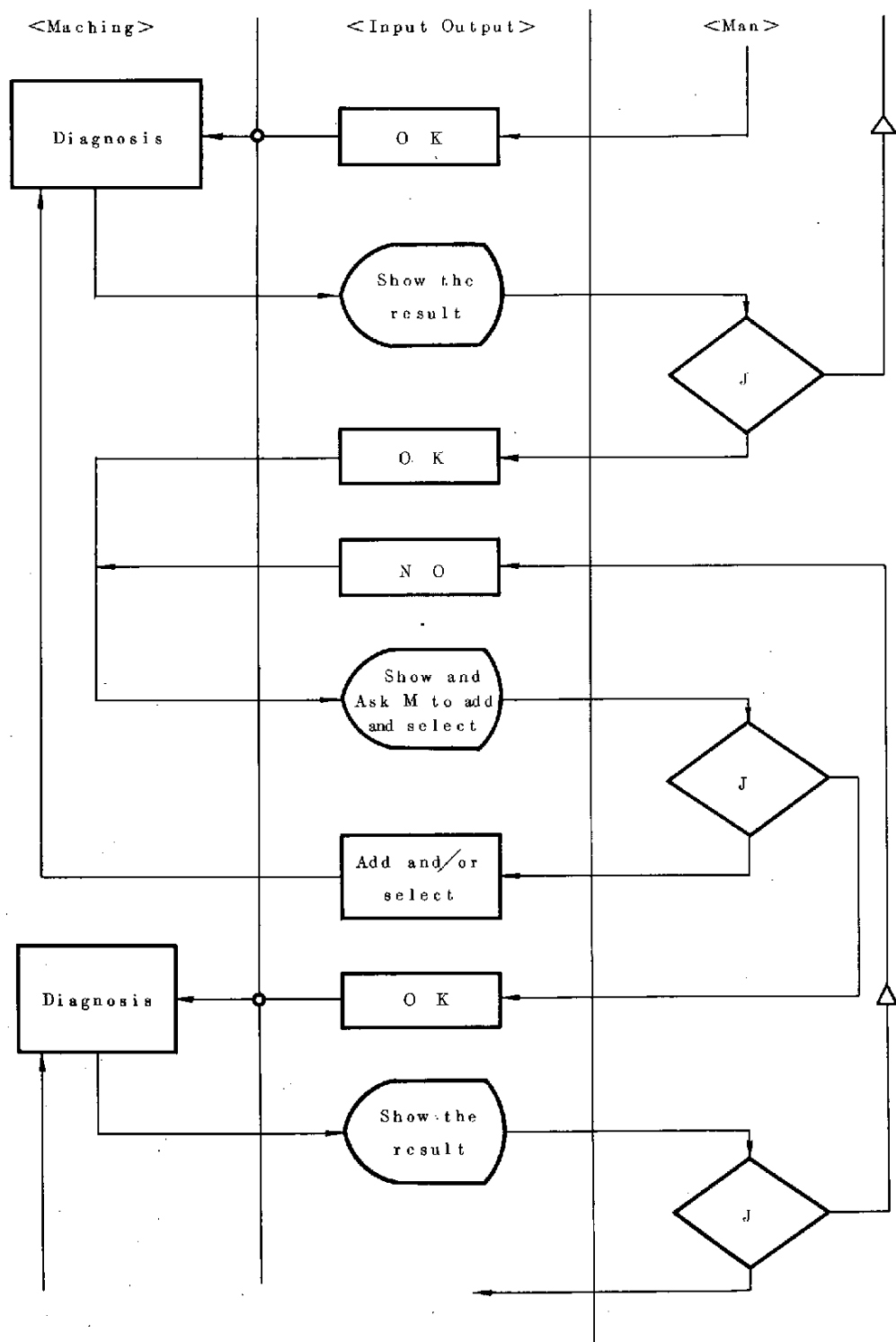


图 2-3 Basic Dialogue Module (BDM)

数Fileを検索するとともにdataおよびparameterの追加・修正が必要な場合はobject stepへJumpする。

Keyverb stepは一つのBDMからなる。このstepではobject stepで決定した主題にしたがってKeyverbに関するTableをsearchし、主題の目的の可能性について結果を示し、stageあるいはstepの附加と選択を機械が人間に対して行なう。

Instruction stepも一つのBDMからなる。このstepではobject stepにおいて決定した主題とKeynoun stepおよびKeyverb stepによる判定にしたがって主題のFeasibilityを知らせ、さらに機械は人間に分析条件のInstructionを質問する。指定されるInstructionは次に示すようである。

1. 期間(時系列)
2. 変数の優先度
3. 内生変数と外生変数の区分
4. 単位の指定
5. タイム・ラグの指定
6. 変数の組合せ
7. 定義式の指定

2.2.2 Execution Phase

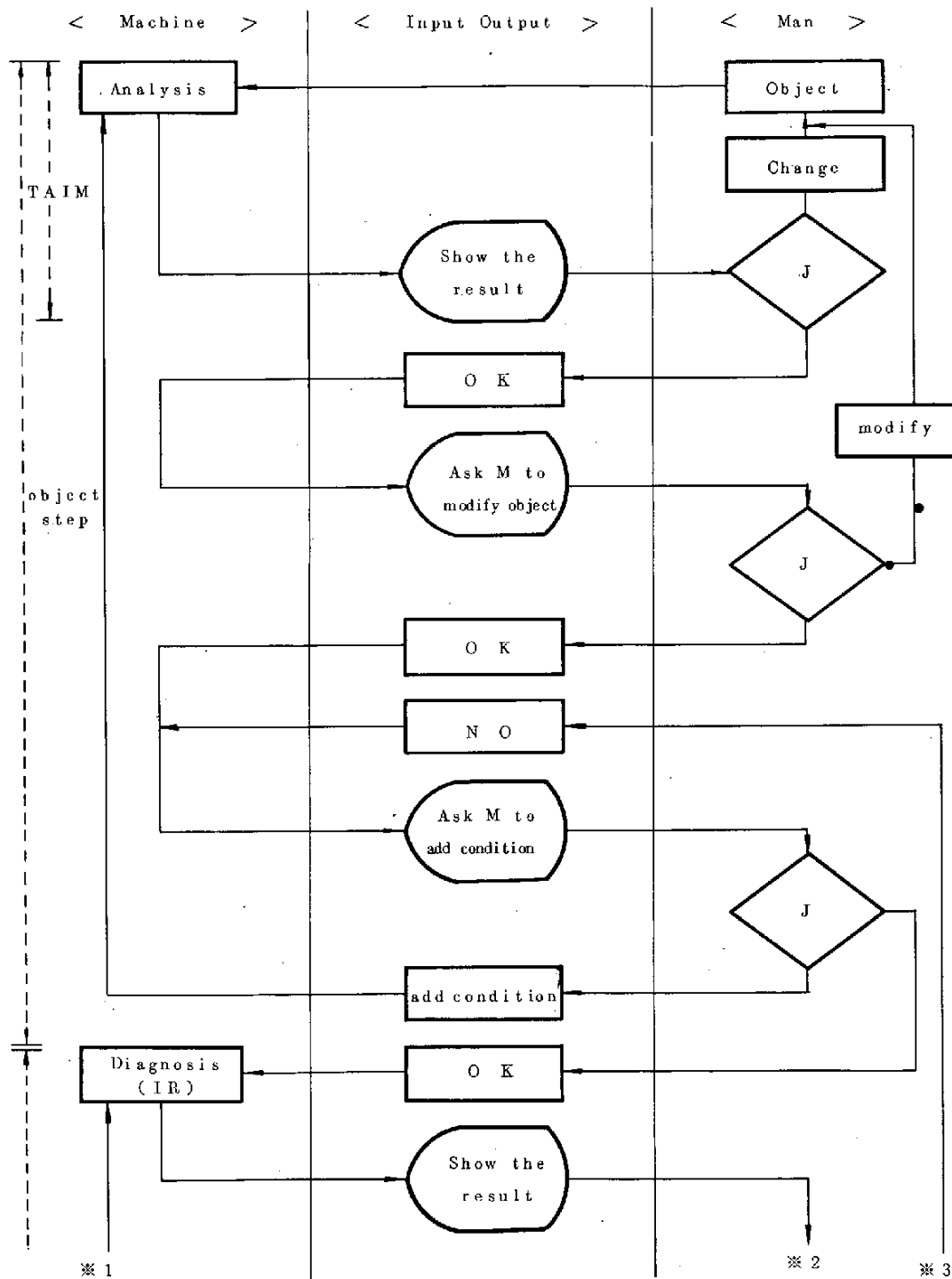
Execution PhaseはOperation stepとOutput Control stepからなる。

Operation stepは一つのBDMからなり、次に示すstagesの個有の性質によってそれぞれのデータ処理および演算を行なう。

1. Data Base
2. Interrelation between variables
3. Estimation of Parameters
4. Model Building
5. Forecasting
6. Simulation
7. Optimization
8. その他

Operation stepはデータ処理および演算を行なってただちにその結果をDisplayに表示するとともに、また機械は人間に対してデータ処理および演算に必要な諸条件の附加あるいは選択を質問する。

Output Control stepは一つのI O D Mから成っている。I O D MとはInput-Output Dialogue Moduleであって、Input装置あるいはOutput装置、Input, Output 様式などについて人間に質問し、人間がそれぞれの装置および様式を選択・指定する。



☒ 2-4(1) ORIENTATION Phase

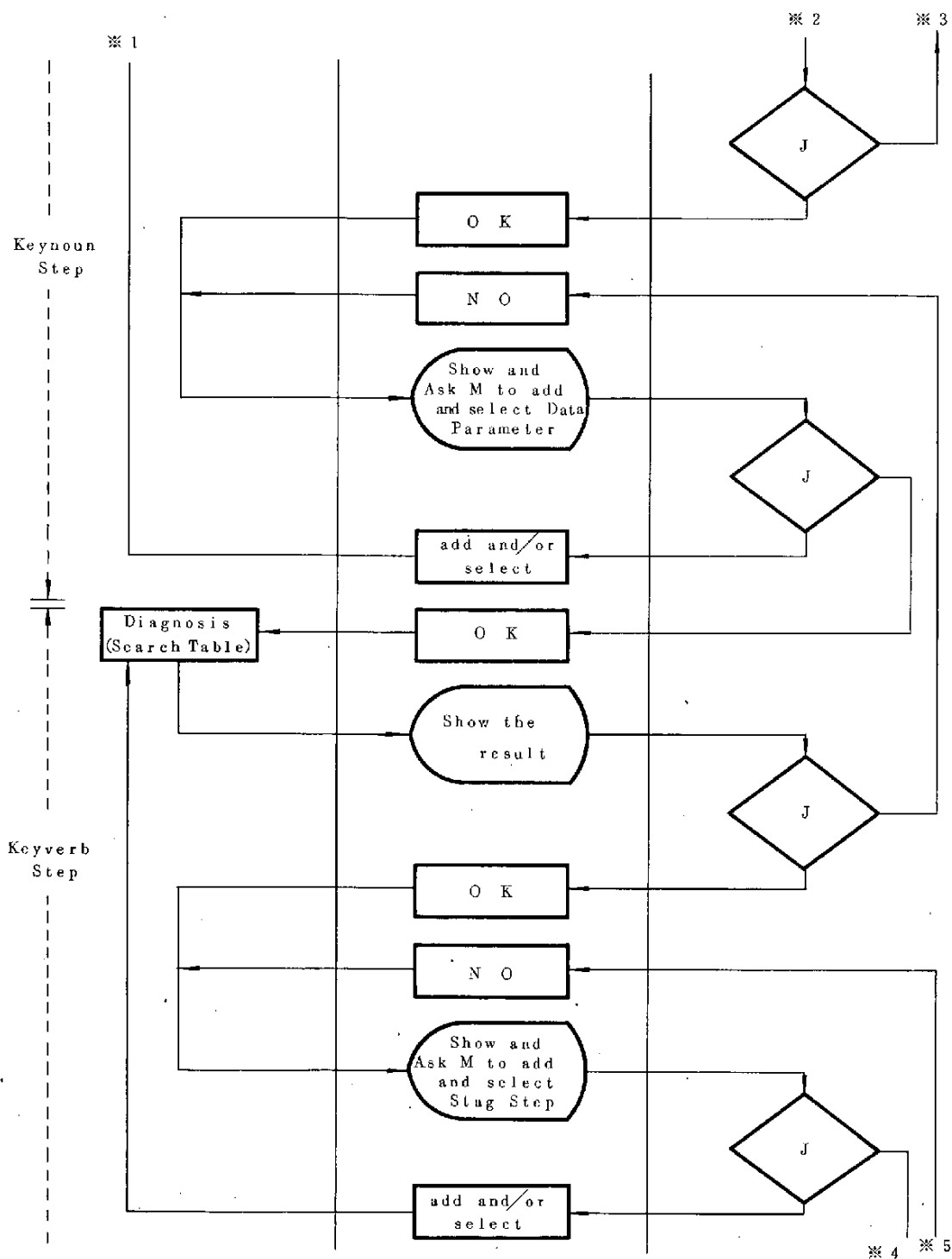


图 2-4(2)

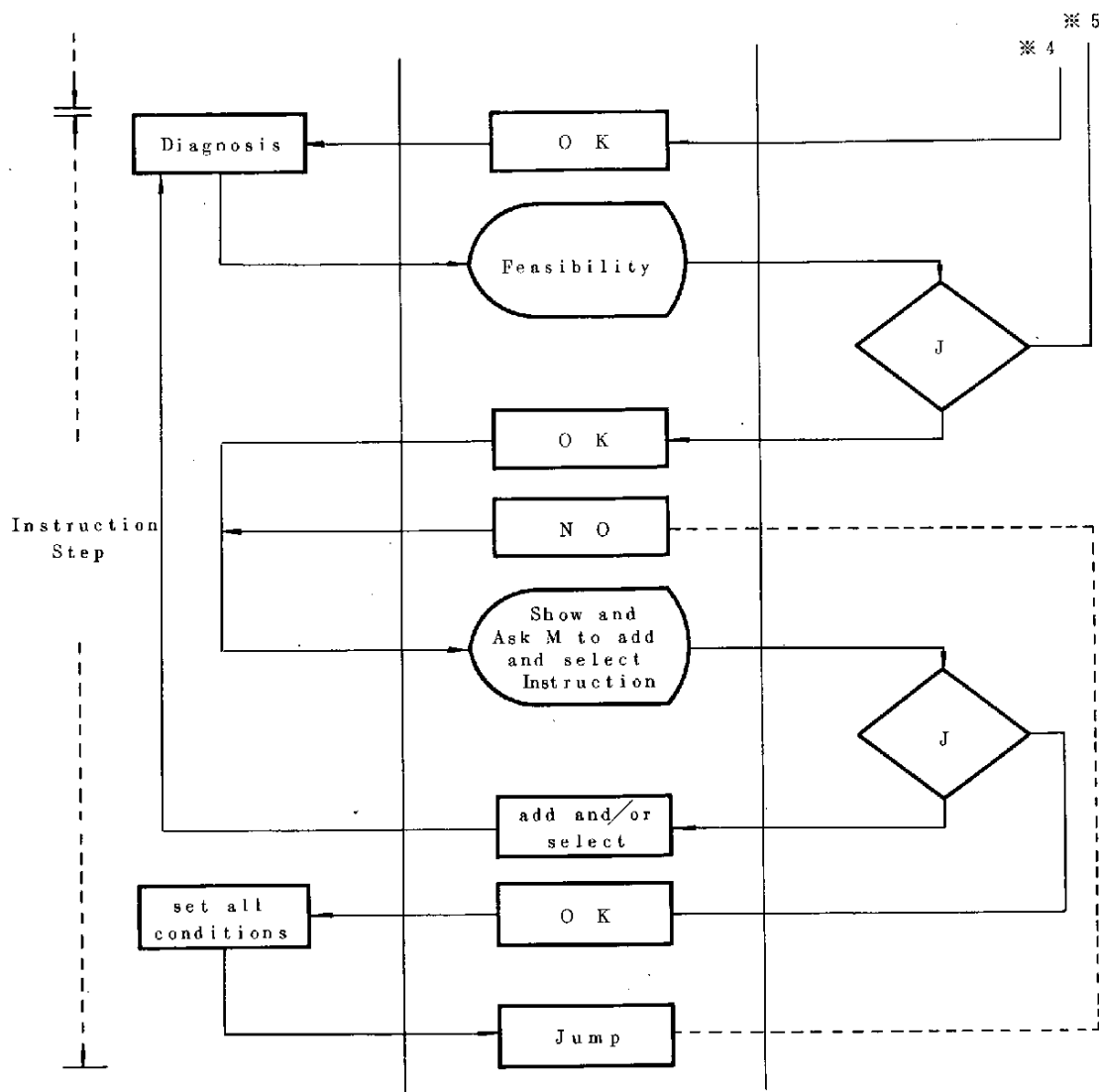


图 2-4(3)

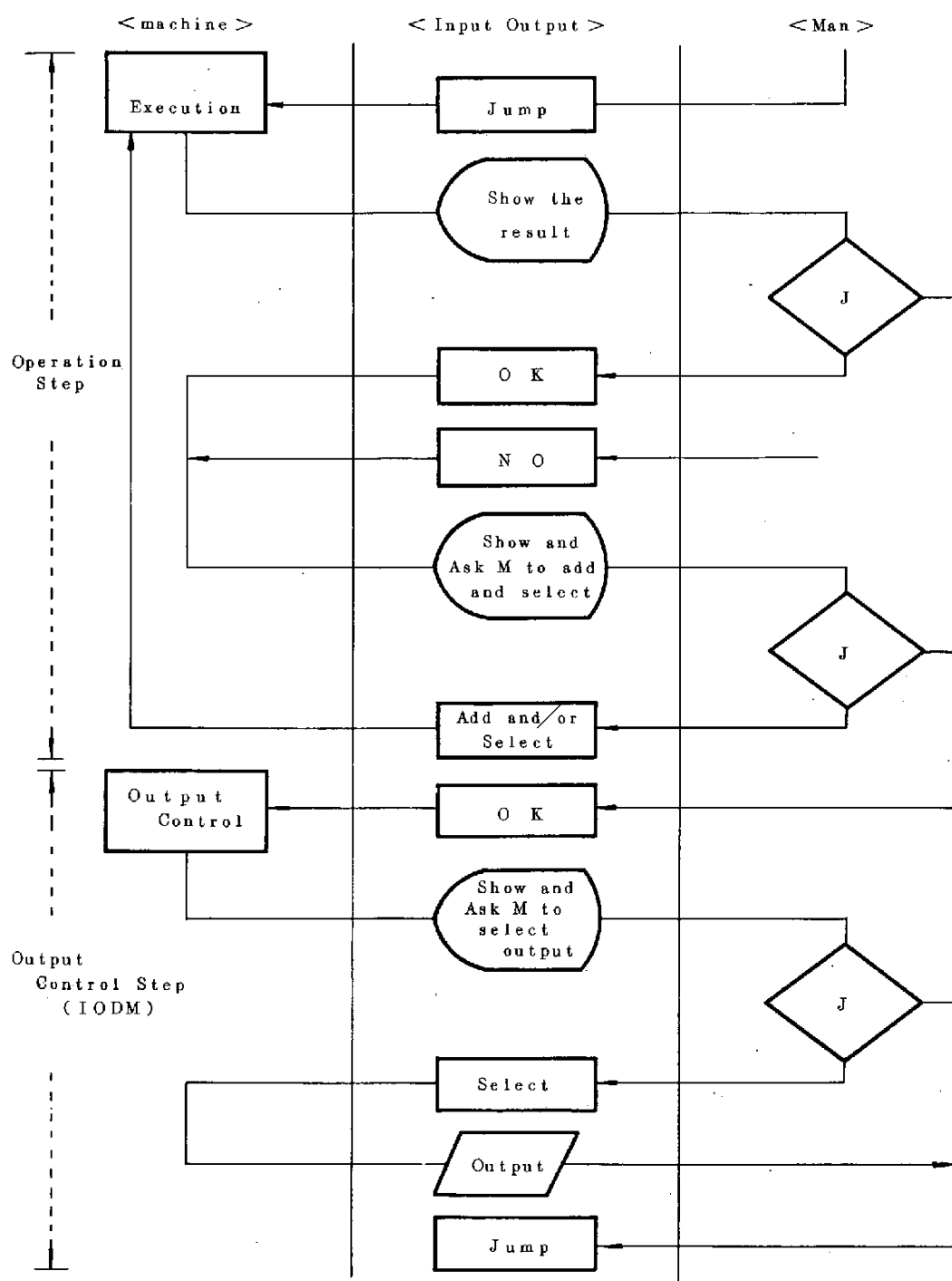


图 2-5 Execution Phase

2.2.3 JUMP Phase

各 stage, phase, step が終ると自動的に JUMP Phase に入るか、あるいは特定 Step において Display 表示が行なわれるときに JUMP Phase に飛ぶかを自由に指定することによって JUMP Phase に入る。

前者を Unvoluntary Jump (UJ) といひ後者を Voluntary Jump (VJ) という。UJ のときは特定の Jump の指定を人間が与えなければ、next stage, next phase, next step は自動的に機械が選んで先へ進むが、VJ のとき機械は自動的に next stage, next phase, next step を選ばない。したがって VJ は必ず次に実行される stage, phase, step を指定しなければならない。

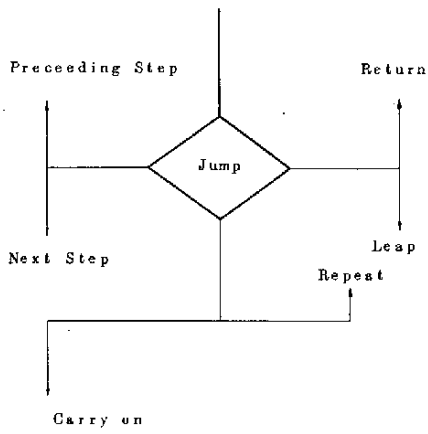
UJ あるいは VJ いずれにせよ JUMP Phase では、初めに RETURN, LEAP, STOP のいずれかの Type を人間に選択させる。

次に RETURN Type ならば Return, Repeat, Preceeding step を表示し、LEAP Type ならば Leap, Carry on, Next step を表示し、STOP Type ならば Stop-End, Halt を表示し、そのいずれかを人間に選択させる。さらに stage, phase, step の順に表示があり、それぞれ一つずつ選択される。

Repeat と Carry on は X_1 -stage, Y_1 -phase, Z_1 -step から X_2 -stage, Y_2 -phase, Z_2 -step までの from-to 形式をとるため、さらに stage, phase, step の順に表示があり、それぞれ一つずつ選択される。

Return (Leap) は任意の stage, phase, step へ戻る (進む) ことを意味し、Repeat (Carry on) はある stage, phase, step を繰り返す (続ける) ことを意味する。

Stop の場合は自動的に END Phase に入る。Halt の場合は一時演算処理を中断する。



Return {
Return
Repeat
Preceding Step

Leap {
Leap
Carry on
Next Step

Stop {
Stop-End
Halt

Stage

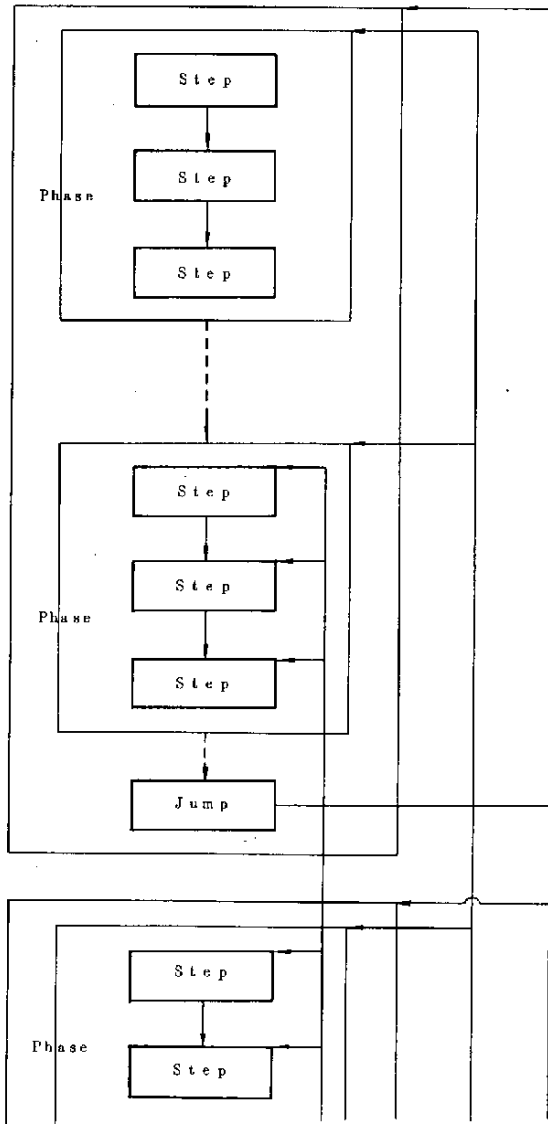
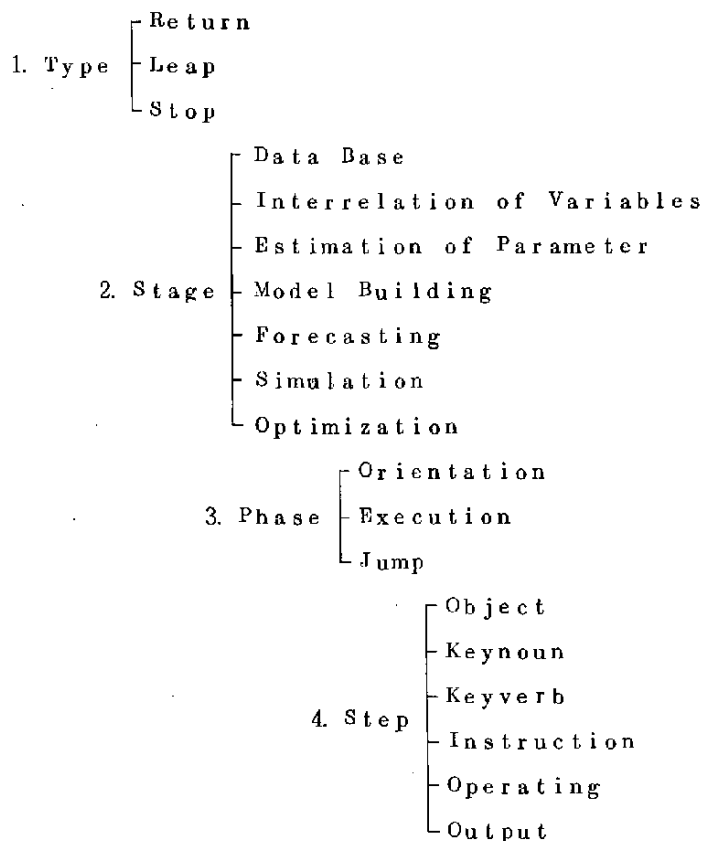


图 2-6 Jump phase

Jump phaseの選択Level



2.3 General Step Flow (GSF)

Orientation Phase, Execution PhaseおよびJump Phaseとを各step毎に接続させGeneral Step Flow (GSF)を書けば図2-7のようになる。図においてJUMPの性質を利用して、各stepあるいはphaseからその他のstepあるいはphaseへ自由にJumpすることができる。

JUMPSに用いられるmodule, TAIM, BDM, IODM, JMをstage, phase, step別に関係づけると、表2-1のようになるであろう。

Theme Analysis Stage (TAS)以外のstageを実行させる場合は、必ず前もってTheme Analysis Stageによって主題分析を行なわなければならない。またTAS以外のstageを実行する場合は、前もって各stage個有の条件を与える必要があり、特有のinstruction Stepを各stageごとに通らなければならない。

Execution phaseのOutput StepはIODMを利用するが、IODM自身はData, 変数, 条件などを附加する場合にInput指定のModuleとして用いられる。

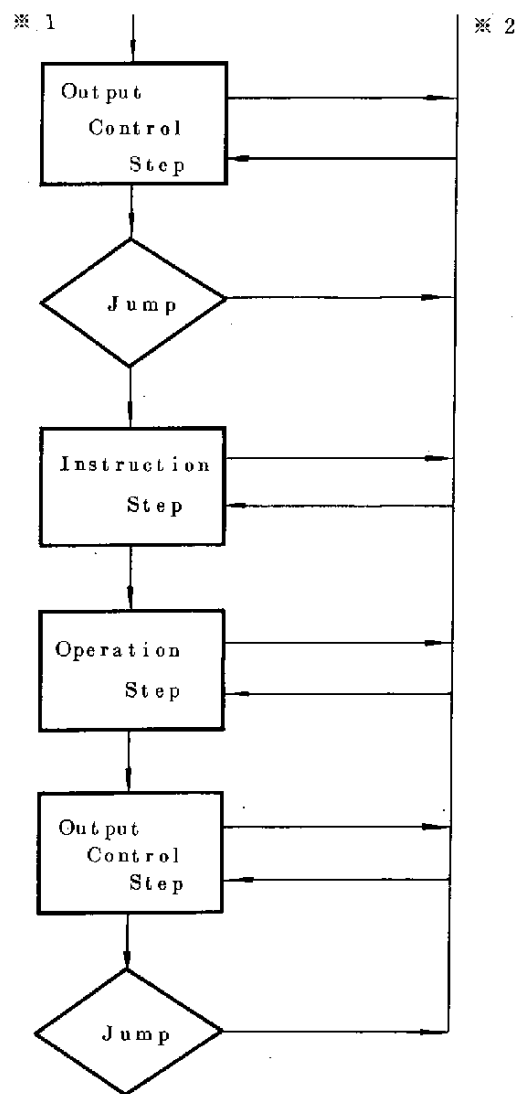
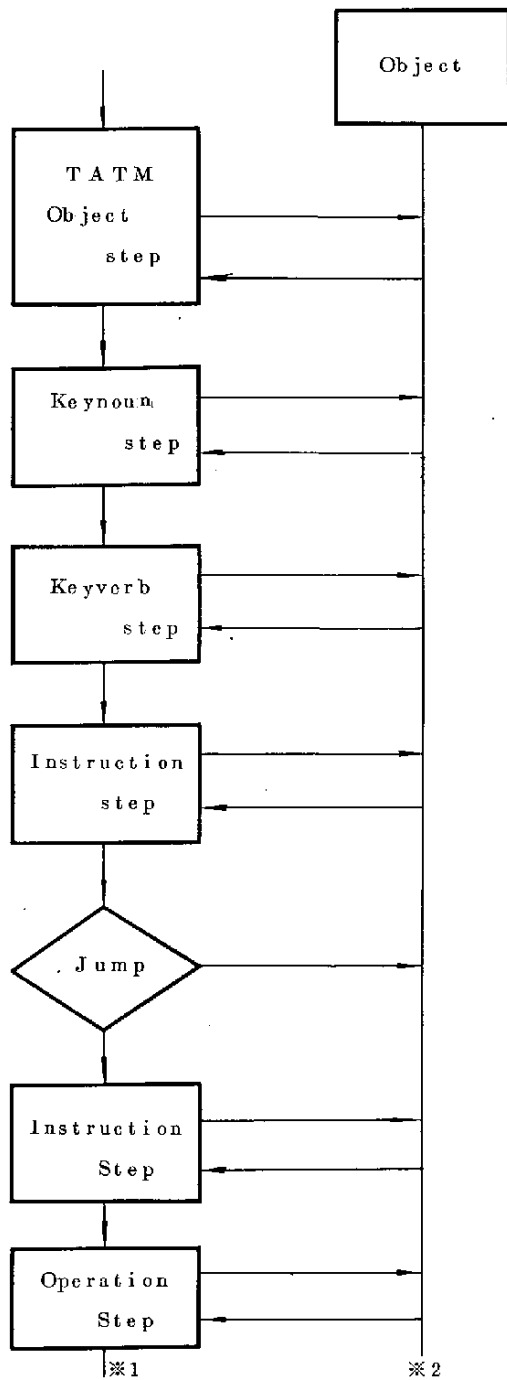


图 2-7 General Step Flow (GSF)

表2-1 Moduleの組合せ表

Phase Step Stage	ORIENTATION Phase						EXECUTION Phase			JUMP Phase
	Initial Step	Object Step	Keynoun Step	Keyverb Step	Instruction Step		Operation Step	I/O Control Step		Jump
Theme Analysis	TAIM	BDM	BDM	BDM	BDM					JM
Data Base					BDM		BDM	IODM		JM
Interrelation between Variables					BDM		BDM	IODM		JM
Estimation of Parameter					BDM		BDM	IODM		JM
Model Building					BDM		BDM	IODM		JM
Forecasting					BDM		BDM	IODM		JM
Simulation					BDM		BDM	IODM		JM
Optimization					BDM		BDM	IODM		JM

TAIM: Theme Analysis Initial Module

BDM: Basic Dialogue Module

IODM: Input-Output Control Dialogue Module

J M: Jump Module

応答形式

応答形式は基本的には大別して次の3つの形に分けることができる。

- (1) あくまでも computer が人間に対して質問 (Request) あるいは要求の形で問いかけ、人間が Yes, No あるいは computer が表示する幾つかの選択子の中から必要項目を選択する C→M 形式
- (2) 人間が computer に質問あるいは要求する M→C 形式
- (3) 上記二形式の混合形式
 1. C→M 形式
 - (1) computer による質問および要求と選択表あるいは選択子の表示
 - (2) 選択表 (Choice Table) あるいは選択子の中から人間が1つあるいはそれ以上の項目の

選択

(3) 問題によっては computer の質問に対して人間による Choice Table の項目,あるいは選択子と Input Table の項目との対応づけ

computer による人間に対する要求は次のような形で Display に表示される。

○ SELECT A INPUT FORM FROM FOLLOWING FORMS

- DIALOGUE FORM
- DEFINITION FORM

○ SELECT A INPUT UNIT FROM FOLLOWING UNITS

- CARD READER (UNIT NO.)
- TAPE READER (")
- CONSOLE TYPE WRITER(")
- DISPLAY (")
- MAGNETIC TAPE (")
- MAGNETIC DRUM (")
- MAGNETIC DISK (")

2. M→C形式

人間による computer に対する要求は次のような形をとる。Display には必ず Input, Output および Jump の表示があり, そのうちの1つを選べばそれぞれの Choice Table が Display に表示される。

DISPLAY
• INPUT
• OUTPUT
• JUMP

(a) INPUT の場合の Choice Table は C→M 形式と同じである。

(b) OUTPUT の場合の Choice Table は C→M 形式と同様 Unit の指定と GRAPH および TABLE などの表示方法と各種書式からなる。

(c) JUMP の場合は前述 (P.) のようである。

2.4 Learning Function

JUMPS の Learning 機構は, 使用者の理解度を示す確定的な状態として初期値を与える段階と, 応答によって試行錯誤しながら使用者の専門程度の状態を決めてゆく段階とに分かれる。

初期値として使用者の確定的な指定による内生変数 x_i ($i = 1, 2, \dots, l$) と外生変数 z_j ($j = 1, 2, \dots, m$) と不確定な変数 Y_k ($k = 1, 2, \dots, n$) の総計 $(l + m + n)$ に対する

n の比, また priority の判定の深さ $P_1, P_2, \dots, P_k$ と変数合計の総和の比は, それぞれ使用者がどの程度問題を理解しているかを表わす係数と考えることができる。

この性質を利用して次式の値

$$\frac{n}{k} = \frac{n}{l+m+n} \times \frac{l+m+n}{k}$$

を専門度係数と呼ぶことができるであろう。

一般に専門度係数を用いるときは

$$y = \log\left(\frac{n+1}{k}\right)$$

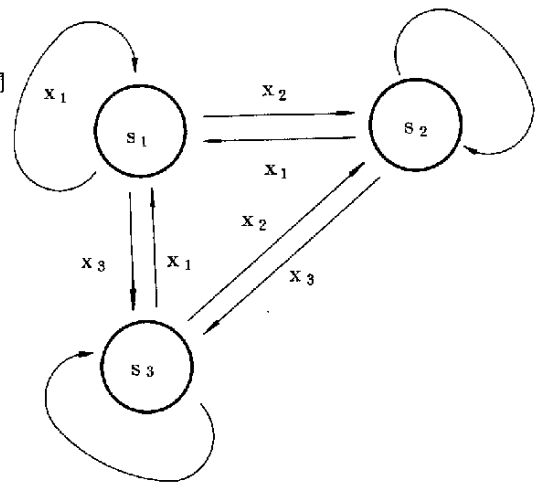
の形式をとる。

これらの専門度係数によって与えられる初期状態から出発して, 各 step stage における応答回数と応答時間の積の総和を判定して, 確定的な専門程度を表す状態からそれぞれの専門化の程度を表す状態へ移行することが可能である。その関係は次の図のように示される。

このように自動的に判定された専門程度に応じて, computer では計算条件および step, stage の選択の制限をゆるめたり, きつくしたりすることができる。

s_1, s_2, s_3 : 専門度を示す状態

x_1, x_2, x_3 : Σ 応答回数 $\times$ 応答時間



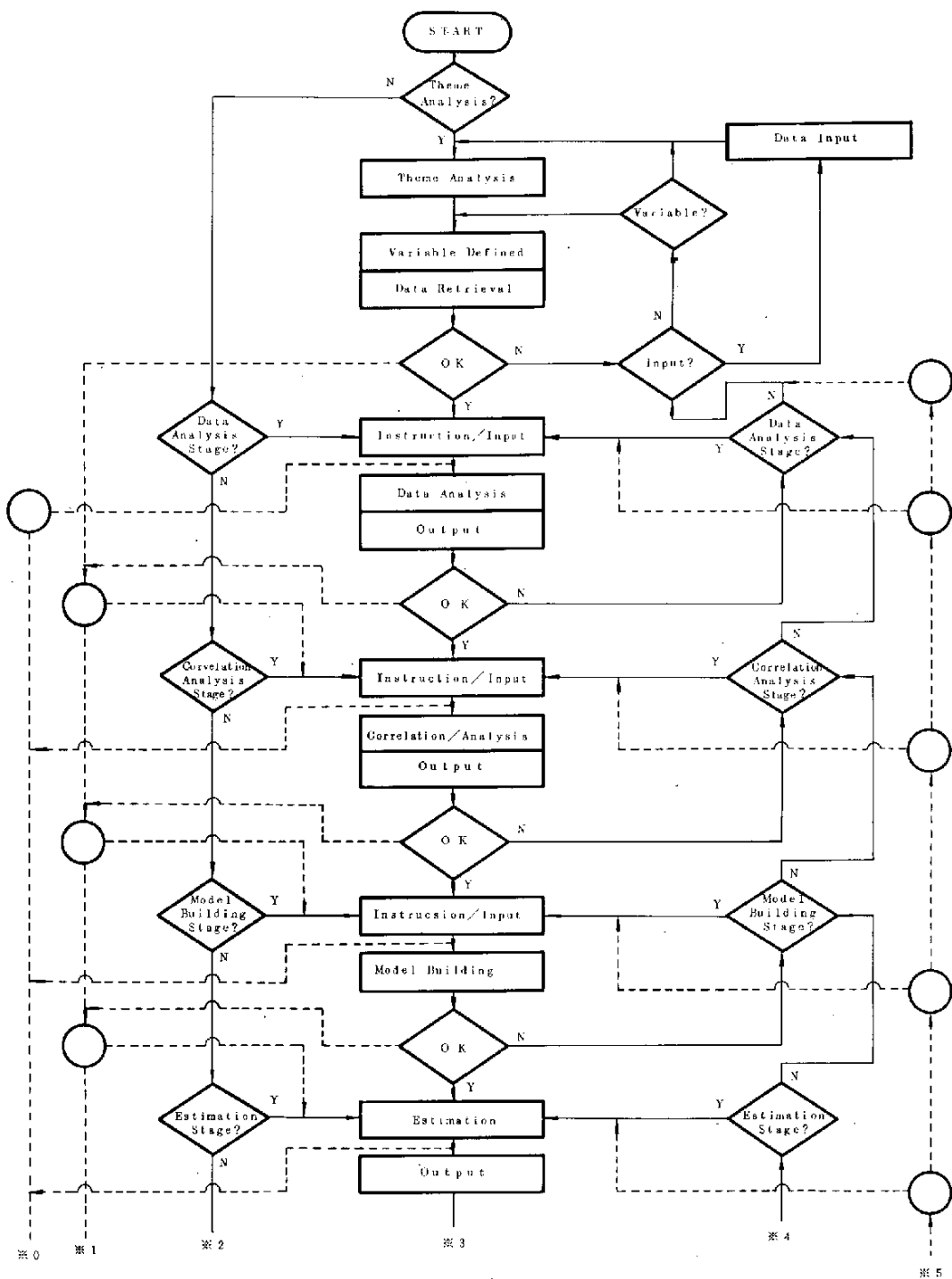


图 2-8(1)

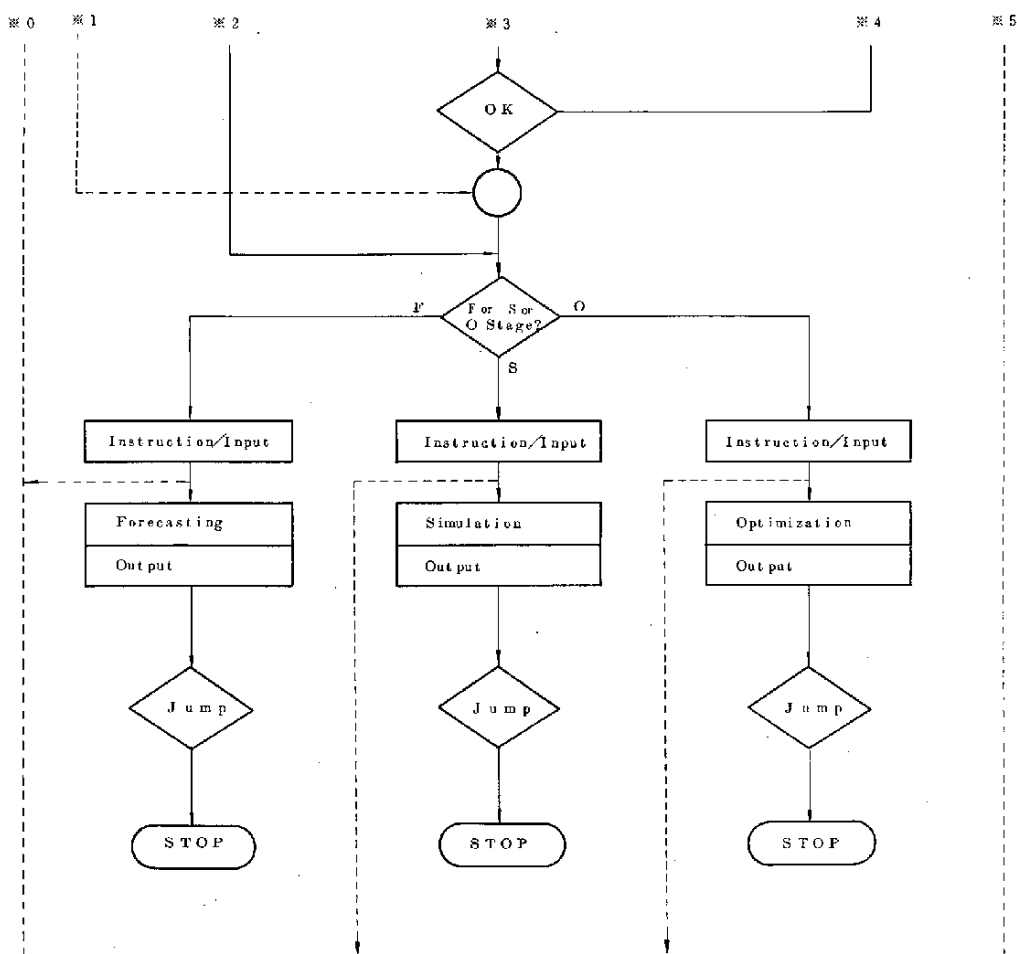


图 2 - 8(2)

3. 主題分析 (Theme Analysis Stage)

人間が提出する目的に対する主題分析におけるMan-Machineの応答は、主題分析以後のすべての分析、検索、演算に影響するがゆえに、非常に重要な、また最も難しいStageでもある。それは人間の多様で且つある程度抽象的な要求に対して、機械側では人間の要求を可能な限り満足するように、問題の領域を定義し具体的な処理手順を確定しなければならぬからである。

人間が提出する目的が明確であれば、問題の対象領域はより限定される一方、結果の厳密性が要求されるであろう。逆に、目的が明確でなければ、問題の対象領域は漠然としたものになる一方、結果の厳密性を要求することはできない。

いまX軸方向に目的の明確性を取り、Y軸に情報量をとれば、問題の対象領域の限定度と厳密性は図3-1のように示すことができる。

目的の明確性についてJUMPSは、使用者の目的の明確性に合せて変数の範囲、処理手順および演算精度が自動的に指定されるようLearning機構を組み入れてある。たとえば目的が明確でなく、結果の厳密性を要求しない場合や応答過程において試行錯誤の繰り返しが多いような場合には、機械は自動的により単純な組合せを選ぶ。

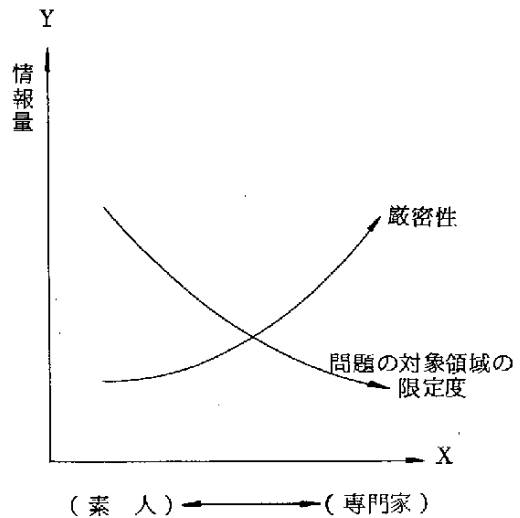


図 3-1

目的の明確性とLearning過程についてはLearning Functionの項で再述する。

3.1 JUMPSにおける主題の構造分析

人間がある目的を達成しようとする場合、

why	何故に
what	何を
when	何時
where	何処で
how	如何にして

JUMPSでは図3-2に示したように人間の要求と、その要求に対する機械の回答、機械の質問とその質問に対する人間の回答という対話を繰返すことによって人間の目的意識を明確にしてゆくが人間が心の中で自問自答

図 3-2 JUMPS における主題の構造分析

する5つの疑問も人間の思考形式のままでは機械に理解されることは困難である。人間がしばしば発する why (何故に) の過程は機械的処理がとくに困難であり、why は人間の問題意識の中で整理されることを前提にして、JUMPS では why の過程を分析対象として扱っていない。しかしながら JUMPS では、what (何を) について厳密に処理することによって、問題の対象領域と具体的な実行処理の対象領域を確定する。

主題分析はこのための段階であって、人間と機械との対話に共通の取り決めが必要となる。とくに人間（user）の目的意識をより明確に判定できるような分類体系をあらかじめシステム内に用意しなければならない。しかも対話の過程にあって問題意識がより具体化されることを考慮して、目的に対する変更、修正が柔軟に行なわれるように作らなければならない。

このような諸条件を考え、JUMPS には問題の対象領域を確定する場合の中心となる概念として **Keynoun** があり、また具体的な処理手順を選択する場合の中心となる概念として **Keyverb** がある。

一般に主題の構造は、noun あるいは noun の群による形容句 (Adj. ph.) や形容詞によって形作られ、それは図 3-3 の如く示される。主題の構造を分析する場合には対象とする領域が如何なる概念群に包含されているかを明らかにしなければならない。

- 1.) 主題の領域を確定するためのKeyとなる名詞をKeynounと呼ぶ。Keynoun は単一の noun または単一の概念としてのadj+nounあるいは複数のnoun または複数のadj

+nounからなる。

例

<noun>

<adj+noun>

<noun>+<noun>

<adj+noun>+<adj+noun>

<Income>, <Cat>, <Dog>

<National Income>

<Black Cat>

<White Cat>

<Income>+<Product>

<Cat>+<Dog>

<National Income>

+<National Product>

<Black Cat>+<white Dog>

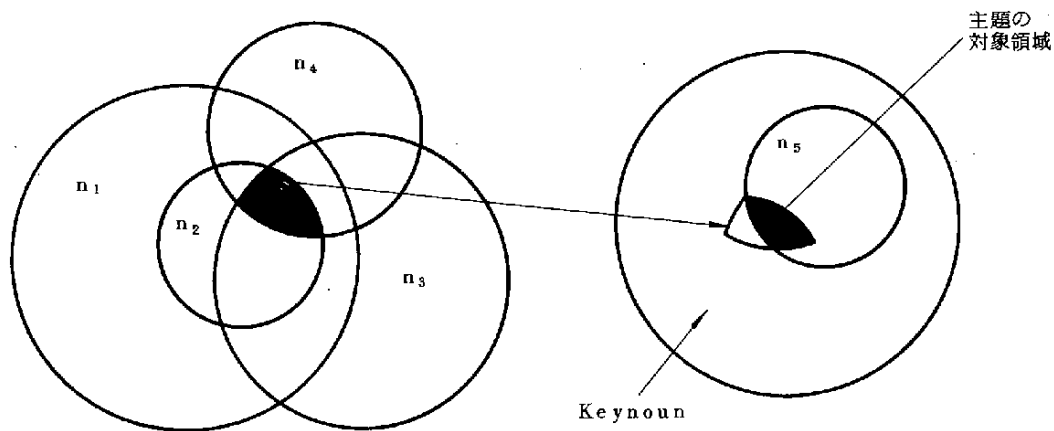
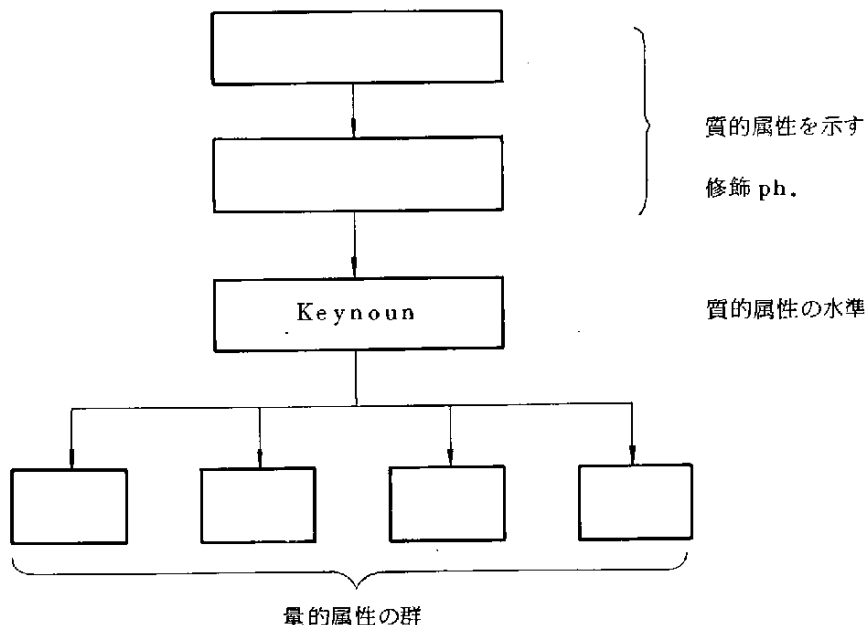


図 3-3

- 2.) Keynoun は質的属性を表わす形容詞あるいは形容詞によってmodifyされ、修飾された Keynoun は質的属性のレベルを規定すると同時に、それはまたいくつかの量的属性と対応している。



- 3.) Keynoun を修飾する形容句は preposition + noun の形をとる。

例

Sedan

of Toyota Autsmobile Mfg. Co.

in Autsmobile Industry

in Transportation Machinery Ind.

of Secondary Industry

in Japan

JUMPSにおいて用いられる前置詞は次の種類がある。

at + noun

in + noun

of + noun

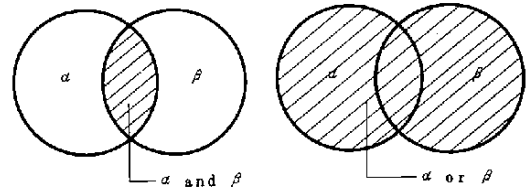
on + noun

at, in, of, on に導かれた名詞は、置かれた順序によって概念の上下レベルを決定する。

with, againstに導かれた名詞は、レベルとしては前の名詞と並列された概念として扱う。

- 4.) Keynoun および質的属性と量的属性は、それぞれAND, ORによって結合される。ただしANDの場合は、ある概念 α とある概念 β との共通概念を意味し、ORの場合はある概念とある概念との両者を含む概念を意味している。

ある概念 α とある概念 β とをそれぞれ独立して利用する場合は、それぞれ α のKeynounと β のKeynounとを分けて指定しさらに+で結合しなければならない。いま、



α = 太郎 (TARO)

β = 花子 (HANAKO)

とし、Keynoun を父 (FATHER) としよう。

FATHER of TARO AND HANAKO

は太郎と花子の共通の父であり、

FATHER of TARO OR HANAKO

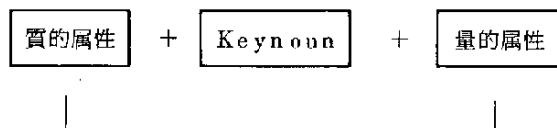
は太郎と花子の共通の父の場合も、またそれぞれ別の父の場合も含む。

FATHER of TARO + FATHER of HANAKO

の場合は太郎の父、花子の父をそれぞれ独立概念として取扱うことになる。

Keynoun の量的属性はたとえば、資本金、設備投資額、生産量、出荷額、労働者数……など具体的な量的データを指定するための属性を表わしている。

量的属性はKeynoun の下部並列概念として処理され、Keynounを中心として質的属性と量的属性が組合わさってKeywordをつくる。



Keyword

Keyword が確定すると問題の対象領域が確定し、JUMPSは自動的にデータのIndexに対応する検索のためのIndex Codingを行なう。

Keyverb

主題分析においてKeyword が確定すると、機械は人間に処理手順を確定するための指定を要求する。JUMPSにおいて処理手順の選択はKeyverb を中心に行なわれる。Keyverbの種類は、表3-1に示すとおりであるが、主題分析においていずれかのKeyverb が必ず入

間によって指定されなければならない。JUMPSのプログラムではJUMPSの利用を柔軟にするために、主題分析の段階で決定した処理手順の指定を、JUMPを利用している間に出てくる結果の如何によって人間がその都度新たに選択、訂正することができる。

表 3-1 Keyverb

Data Base
Interrelation of Variables
Model Building
Estimation of Parameters
Forecasting
Simulation
Optimization

分析条件の指定

期 間

優先度

内生変数と外生変数の区分

単位の指定

タイム・ラグの指定

変数の組合せ

定義式の指定

すなわち基本的にMan-MachineのOrientation Phase における対話には、以上のa, b, cを含んでいなければならない、それらをもう一度整理すれば

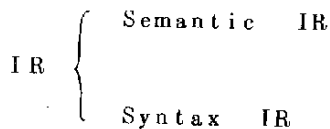
- (1) Data Retrieval のための質的Keynoun, 量的Keynounの指定あるいは選択
 - (2) Phase, Step, Stage決定のためのKeyverbの指定あるいは選択
 - (3) 分析条件の指定あるいは選択
- となる。

3.2 JUMPSのInformation Retrieval

主題分析の段階で人間と機械の対話を通じて、Keynoun を中心に人間の要求する問題の領域が明確にされ、Keyword が作られKeywordのcodingが行なわれる。次の段階は問題領域に対

応する変数のDataを検索することである。

1. JUMPSにおけるIRの考え方



Semantic IR における主題分析では、各Keywordと一対一の対応関係が決められている。Syntax IRにおいては、主題分析からのKeyword によって、それと対応関係にある情報の内容が取出されている。

いずれにせよ一般のIRでは、図3-4のような過程を経て個々のデータや情報が取出されるがこのデータ検索段階では時間的に効率よく、どのように情報を検索するかが問題となる。

JUMPSでは、主題分析によるKeywordが個別Dataあるいは個別情報に対応することはもちろん、群としてのDataあるいは情報にも対応している。さらに、それらの群内における個別データや情報が、主題に対してどのような関係にあり、またどのように関係づけることが最適であるかを求めている。

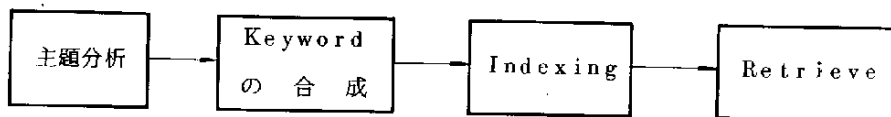


図 3 - 4

(i) FILE作成の方針

上で述べたように、JUMPSでは種類の異なる関連情報を群として取扱い、関連の重要さの程度がUserによって相対的に決定されるから、検索時間の効率を考えればUser OrientedなFileを各User側であらかじめ準備することが一番望ましい。しかし User が効率的なData Fileを用意すれば汎用性は無くなってしまう。この関係は次のように書けるであろう。

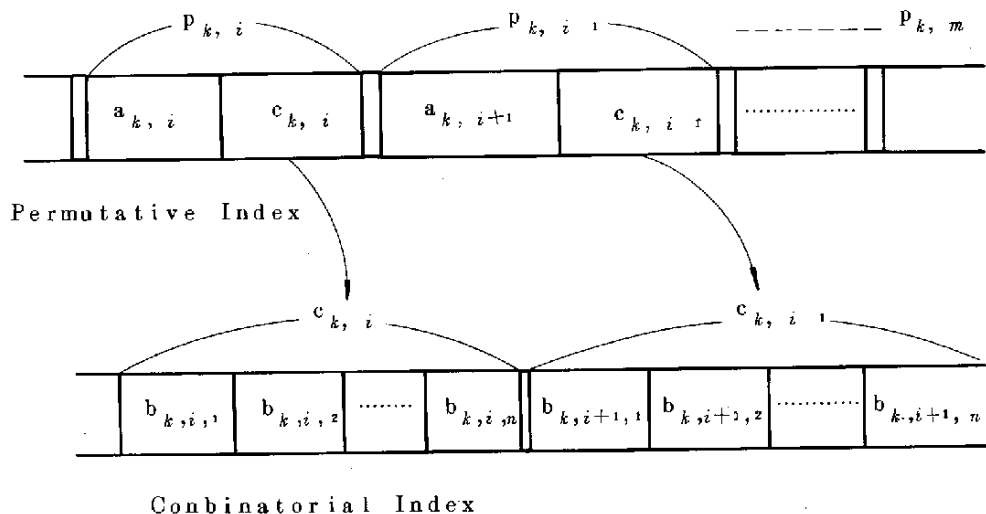
- | | |
|-------------|--|
| 汎
用
性 | ① JUMPS FILE として具体的なFILE構成とDataの内容を決定してしまう |
| | ② FILE構成だけ決定しておいてDataはUser側で整備する |
| | ③ FILE構成およびDataの種類をUser側で準備する |
| 効
率 | |

結果的にJUMPSでは折衷案②を選んだが、以下に述べるような効率と最的性を満足するgeneralなFILE構成を考えている。

3.3 FILEの構成

(1) Index Image

code化されたKeywordは、JUMPSの中で図3-5のようなIndex Imageにしたがって処理される。Index Imageとはcode化されたKeywordにしたがってData Fileの中にあるDataを検索するためのIndexing Tableの基本的な構成様式をいう。



- k ; 並列概念の分類を示す(質的属性)
- i ; 上下概念の順序を示す(質的属性)
- n ; 量的属性に関する並列概念の分類を示す

図 3-5 Index Image

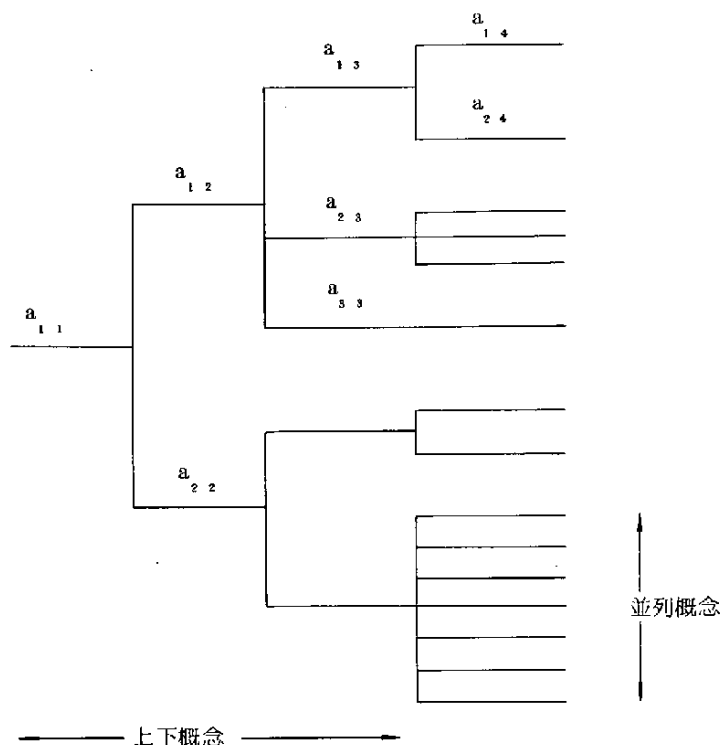
質的概念を示すインデックスをPermutative Indexと呼び、量的な概念の区分を示すインデックスをCombinatorial Indexと呼ぶ。

(2) Permutative Index

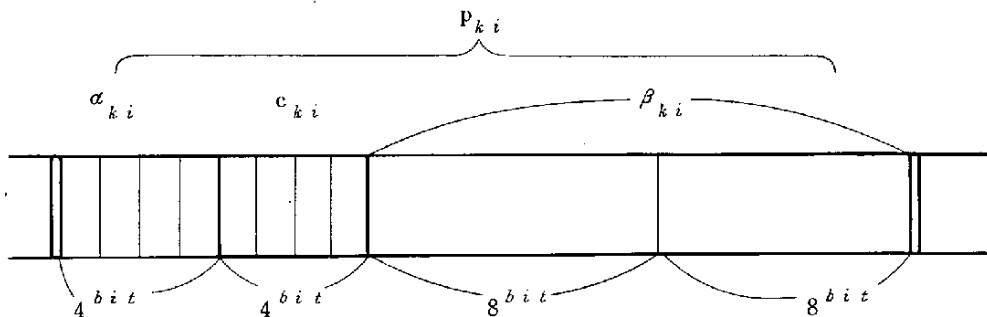
質的属性を $a_{k_1}, a_{k_2}, a_{k_3}, \dots, a_{k_m}$ とし、概念領域を $a_{k_1} > a_{k_2} > a_{k_3} > \dots > a_{k_m}$ とすれば、 $a_{k_l} > a_{k_m}$ ならば a_{k_l} は a_{k_m} の上位概念である。

量的属性のIndexを表わすcodeすなわち a_{k_1} に対応して c_{k_1} , a_{k_2} に対応して $c_{k_2}, \dots$, a_{k_m} に対応して c_{k_m} をもち、 a_{k_i} と c_{k_i} とが一セットとなって p_{k_i} を構成する。

しかしながら、一般に質的属性を示す概念は以下に示すようにTrec構造になっており、並列概念を示す k の値はvariableである。



例えば東京都の区のように23区ある場合に同一レベルで8分類したいときは、23区を適当に東、西、南、北、東北、西北、東南、西南地区などに分割することもできるし、あるいは目的によって23分類することもできる。前者の場合は2レベルに、後者の場合は1レベルに並列概念が入ることになる。JUMPSではこのように並列概念の数を柔軟性をもたせ自由に指定できるよう、Permutative Indexの中をさらに次のように利用する。



図で α_{ki} と β_{ki} とで a_{ki} を構成し、さらに c_{ki} と一緒になって p_{ki} を構成する。 α_{ki} は β_{ki} の数、すなわち並列概念の数を示す部分であり、 β_{ki} は並列概念のそれぞれを区別する code であり、以下の如く k に対応する。

β_{1i}	0 0 0 1
β_{2i}	0 0 1 0
β_{3i}	0 1 0 0
β_{4i}	1 0 0 0

a_{ki} に附随する c_{ki} は基本的には 4 bits = 1 nibble であって $c_{k1}, c_{k2}, \dots, c_{k16}$ 、すなわち 16 の質的属性の上下概念を規定することができる。

(3) Combinatorial Index

Permutative Index によって指定された各 c_{ki} ($i = 1, 2, \dots, 16$) は、さらに量的属性としてのいくつかの概念群に対応する。すなわちそれぞれの c_{ki} に対し b_{kij} ($j = 1, 2, \dots, n$) が対応する。

この j の数は、質的属性の上下概念、並列概念に対応する量的属性を示す概念が共通する場合と、共通でなくそれぞれ個有の量的属性である場合があり、それらは図 3 - 6 のように区別される。

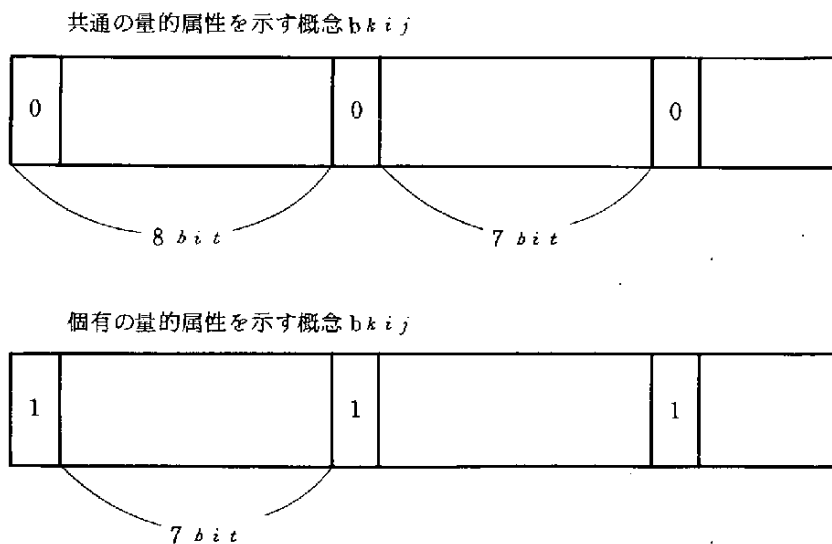


図 3 - 6

JUMPSでは主題分析によって決定した質的属性と量的属性からなるKeywordを Keyword codeに変換し、さらに各Keywordはそれに対応するIndex codeが与えられる。Keywordの質的属性は、アルファベットからなる1語あるいは2語、または3語によって表現される。

1語よりなる属性の長さは制限はないが、JUMPS内では最初のアルファベットを含んで最高6桁のアルファベットとして処理される。

2語あるいは3語よりなる量的属性は、必ず1つあるいは2つの語がハイフン(-)によって結びつけられていなければならない。

たとえばnational incomeやgross national productは

NATIONAL-INCOME (2語)

GROSS-NATIONAL-PRODUCT (3語)

のように表現される。

JUMPSではこのような合成語のハイフンを用いて表現するが、ハイフンが1箇の場合、すなわち2語よりなる合成語の場合は、第1語の3文字と第2語の3文字の頭が採られ次のように表わされる。

NATINC

ハイフンが2箇の場合、すなわち3語よりなる合成語は第1語の頭から2文字、第2語の頭から2文字、第3語の頭から2文字がとられ

GRNAPR

という語に置きかえられる。

Symbol Table

連番	Keyword	Keyword Code		Index Code			
				α	c	β	b
0001							
0002							
0023	Population of Tokyo	401823	569124	0001	0011	00100110	0010100
				4^{bit}	4^{bit}	8^{bit}	8^{bit}

6文字の語は10進数に交換され、さらに二乗採中法によって6桁の10進数よりなるKeyword codeに変換される。

$$\begin{array}{l} \text{NATINC} \quad 513953 \\ 513953^2 = 264147686209 \\ \quad \quad \quad \underbrace{\hspace{1.5cm}}_{6\text{桁}} \\ \text{GRNAPR} \quad 795179 \\ 795179^2 = 632309642041 \\ \quad \quad \quad \underbrace{\hspace{1.5cm}}_{6\text{桁}} \end{array}$$

Keyword code が重複する場合はK+1, K+1が重なる場合はK+2のようにcode化される。このようにしてすべてのKeywordは質的属性を示すKeywordのcodeと量的属性を示すKeywordのcodeに変換され、Symbol Tableに格納される。

Symbol Tableをあらかじめ連番をもって用意されている。Symbol TableにないSymbolが使用される場合、機械はJUMPS利用者に対して「UNDEFINED SYMBOL」という表示を行ない、DataのInputが必要かどうかを尋ねる。この場合はData Input Stageに戻ってDataをInputしなければならない。

Keyverb Routine Control Table (KRCT)

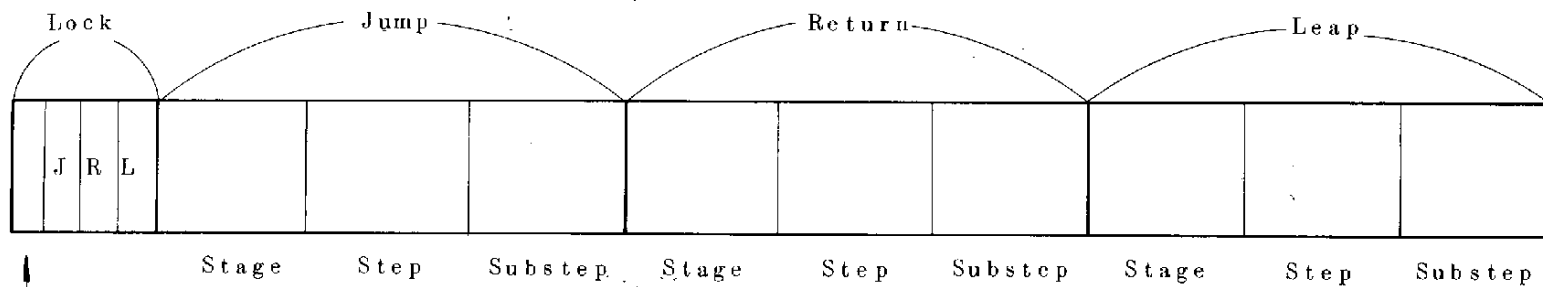
JUMPSはJumpを効率的に実行するために特殊なKeyverb Routine Control Table (KRCT)をもっている。KRCTはStage, Stepを示すcode番号と、それぞれのStageあるいはStepのrelocatableなEntry, Exitが書き込まれている。すなわちJUMPSのStart開始番地が指定メモリに格納されると、移動したbyte××××だけEntry, Exitの番地に移動が発生する。(ここで、JUMPSのStart開始番地はrelocatableで相対番地として〇〇〇〇バイトがプログラム上に与えられている。)

JUMPSではprocessの途中でJUMP, RETURN, LEAPを任意に指定することができるように、KRCTはさらに5byteからなるSequence Table (ST)をもっている。STの最初のnibble=4bitは、それぞれ利用者のInitialization段階における大まかな指定、JUMP, RETURN, LEAPの指定に対応してUNLOCK, LOCKを決めることができる。LOCKの場合は機械自身が最適な手順を予定してStage, Stepの連結したcode番号を格納している。UNLOCKのJUMPは使用者による次のStage, Stepの指定によってStage, Stepのcode番号が記憶される。RETURN, LEAPの場合も同様であるから、繰り返し同じprocessを利用する場合は、使用者はLOCKの指定を各Stage, Stepで行わなければならない。

Keyverb / Rtn control Table

Stage	Step	Entry	Exit	Sequence				
				1	2	3	4	5
A	A	1 0 2 8	1 0 8 6					
	B	1 0 8 7	1 1 0 3					
	C	1 1 0 4	1 4 0 1					
	D	1 4 0 5	1 7 9 8					
B								
C	A	3 8 0 5	3 9 3 8					
	B	3 9 3 9	4 0 8 3					

5 byte



0 大まかな処理手順の指定

1

4. データの抽出 (Data Retrieval)

JUMPS ORIGINAL File は各種の利用目的をもった利用者に利用されるために、種々のデータが格納されている。ある特殊な利用目的のために、このJUMPS ORIGINAL FILEを利用するためには、人間が主題を与え、computerがデータ領域を限定し、応答形式によって必要なデータを抽出するか、あるいは人間があらかじめ変数および方程式を定義し領域を限定してデータを抽出するかいずれかを行なって、人間が求める目的に対するデータを抽出しなければならない。

いずれにしても抽出されたデータはDisk Fileに導びかれるが、このFileを Selected Fileと呼ぶ。

以下では人間とcomputerの応答によってexecutionを導く形式を応答形式と呼び、人間があらかじめexecutionの内容を指定する方法を定義形式と呼ぶことにする。

もともとのJUMPS Original Fileは設定される主題に対してデータのIR及び、概念の上下関係を考えて構成されているために必ずしもEXECUTION時において処理しやすいデータ構成をなしていない。

このために、もしJUMPS Original FileをEXECUTION時にそのまま使うとすれば、データを使うたびに編集していかなければならない。

データ抽出は、このようなことをさけるためにEXECUTION時に必要とするデータを必要なフォーマットになおしてSelected Fileと称する1つのFileを作り上げる働きをいうものである。

会話形の特徴として必要とされるデータは、問題の解決の過程においてUserが必要とするときに指定され、一度指定されたとしても、必要に応じて再指定が行なわれ、その指定は何時なされるかは不定である。このために、データ抽出の働きを1つのstepとして独立させることは効率を考えれば不適當である。

このために変数が指定されるたびにJUMPS Original Fileから読み出して実際に使う形に編集してSelected File上に登録することとする。ここで作り上げるFile は再指定されればUpdateをし、必要な時に速くdataをとりだせる形をしていなければならない。

この条件を満たすFile形式としては次の3つが上げられる。

1. Direct Access File
2. Indexed Sequential File
3. Partitioned Sequential File

ここでもう一つ要求される特徴として、データ長が不定であることを考えて、しかも Update がおこなわれることを考えれば、Indexed sequential Fileがこのシステムでは、一番使いやすいこととなる。

4.1 データ抽出における定義形式

DEFINE VARIABLE

JUMPSでは応答形式によってInputのDATA、変数、方程式などを決めることができるが、利用者の便宜を考慮して、利用者自身があらかじめ定義したDATA、変数、方程式をInputすることができる。この場合も、確定したDATA、変数、方程式ばかりでなく、コンピュータと人間とがDATA、変数、方程式を選ぶことができるように工夫されている。

アルファベットによって書かれた文字の集合をNAMEという。NAMEは質的属性と量的属性をもつKeywordのこともあれば、DATA名のこともある。

変数にはXi Type, Zj Type, Yk Type, およびDl Type Variable の四種類がある。Xi Type Variableは内生変数であり、Zj Type Variable は外生変数である。Yk Type Variableは内生、外生が未決定の変数であり、Dl Type Variableはdummy変数である。

変数を定義するには〔DEFINE VARIABLE : 〕というLABELをつけ、等号をはさんで右辺にNAMEあるいは変数を書き、左辺にXi, Zj, Yk, DlそれぞれのTypeの変数を書く。ここでi, j, k, lはそれぞれ3桁まで許される。

あらかじめZ3のNAMEが定義されていれば

$$X5 = Z3$$

のように変数の再定義が可能である。

DEFINE VARIABLE ;

X1=NATIONAL-INCOME

X2=CONSUMPTION

X3=GNP

Xi=

Xm=CAPITAL

Z1=

Z2=

```

Z j =
Z n =
Y 1 = POPULATION
Y 2 =
Y k =
Y o =
D 1 =
D 2 =
D l =
D p =

```

VARIABLE

A1, A2, A3, , B1, B2, B3のようにXi, Zj, Yk, D1 以外の変数を定義することも可能であるが、それぞれ変数が内生 (ENDO genous Variable) であるか、外生 (EXO genous Variable) であるか、内生、外生が不確定 (UNDetermined Variable) であるか、またダミー (DUMmy変数) であるかを次のように定義しなければならない。この場合はLABELは [BARIABLE:]である。

VARIABLE ;

ENDO-VAR : A1, A2, A3,

B1, B2, B3,

EXO-VAR : C1, C2, C3,

D1, D2, D3,

UND-VAR : E1, E2, E3,

F1, F2, F3,

DUM-VAR : G1, G2, G3,

あらかじめ定義されたXi, Zj, Yk, D1 Typeの変数であっても、LABEL [VARIABLE:]を用いることによって再定義が可能である。

ENDO-VAR : Z1, Z2, Z3,

この場合のZ1, Z2, Z3, …… は外生変数ではなく内生変数である。

PRIORITY

変数は利用目的によっていくつかの変数群からなるSectorに分け、Sectorの中で変数に優先順位をつけることができる。

優先順位を指定する場合はLABELとして、[PRIORITY;]と書き、次に第iセクタの指定として[SECTOR i;]と書き、優先順位第j位の変数群を

Pj(X1, X2, …… , Z1, Z2, Y1, D1)

のように指定する。

変数に括弧をつければ、利用者が一応優先順位第i位にその変数をランクさせてはいるが、computerと人間との会話を通してそのほかのrankに移すことができるような変数である。

PjはP1からP16まで可能である。すべての機械的処理はPjの高い順(1, 2, 3, ……)から行なわれ、Pjをもたない変数は最後に処理される。

処理の途中で再定義によりPRIORITYの変更が可能である。

PRIORITYの用いられる変数はXi, Zj, Yl, Dk以外に一般の変数およびFunctionで定義された変数、LAGあるいはPERIODを指定された変数を含んでいる。

PRIORITY

SECTOR 1

P1(X1, X2, X5, …… , Z4)

P2(X3, X6, Z2, (Z3))

P3(X4, (X9), X12, Z7, …… , (Z9))

P1(X13, (X15), Z13, Y2, (Y5))

SECTOR 2

P1(X10, Z6, Y7, D2)

P2(X21, A1, B2, C3, ……)

DEFINE FUNCTION

JUMPSではFORTRANの使用を認めているから、次のようにFORTRAN IV形式で自由に方程式を書ける。

DEFINE FUNCTION;

L1=LOG(X1)

L2=LOG(X2)

L3=LOG(X3)

E1=EXP(X1)

$$E2 = EXP(X2)$$

$$E3 = EXP(X3)$$

$$G1 = (Y1 * Y2) / (D * D)$$

$$G2 = (X2 + X3) / Y1$$

$$G3 = X1(T) - X1(T-1)$$

DEFINE TIME LAG/PERIOD

時系列データを処理する場合、次のように変数にTime Lagおよびperiodの指定をすることができる。

Time LagはLABELとして〔LAG:〕を用い、Time Lagを指定すべき変数につづく等号の右辺に $T(\alpha, \beta)$ の形式を書く。

たとえば $\alpha = -1$, $\beta = 2$ とすれば

$$X3 = T(-1, 2)$$

は $X3_{t-1}$, $X3_t$, $X3_{t+1}$, $X3_{t+2}$ を意味する。

$$X1 = T(0, 0)$$

は $X1_t$ であって、この場合は $X1$ だけと同じである。

DEFINE TIME:

$$LAG: X1 = T(0, 0)$$

$$X2 = T(-2, 0)$$

$$X3 = T(-1, 2)$$

注) ただしDisplay表示あるいはOutputにおいてLagは変数の区別をするために

$$X1 \rightarrow X1(T+0)$$

$$X2 \rightarrow X2(T+0), X2(T-1), X2(T-2)$$

$$X3 \rightarrow X3(T+2), X3(T+1), X3(T+0),$$

$$X3(T-1)$$

と表す。

LABEL〔PERIOD:〕に導びかれた変数は等号に次く右辺に $P(\alpha, \beta, r)$ の α , β , r を指定することによって、データ期間と時系列データの性質(単位期間)を指定することができる。

α は利用されるデータの開始時期, β は終了時期, r は時系列データの性質を示す。たとえば1945年第1期からはじまり1970年第4期で終る4半期データならば

$$P(1945-1, 1970-4, QT)$$

と書けばよい。

データによってデータ期間が色々な場合があり、とくにTime Lagによって時系列を変化させる

場合は、すべての変数のデータ期間が必ずしも一致しない。そこでデータの始まりを指定し、すべての変数のデータ期間が一致する最大まで、あるいはデータの終了時を指定し、すべての変数のデータ期間が一致する最小時まで computer によって一致させるという意味から、次のように指定することができる。

PERIOD :

X1=P(1945, 1970, YR)

X2=P(1945-1, max, BI)

X3=P(min, 1970-4, QT)

X4=P(min, max, MT)

ここで YR ; 年
 BI ; 期
 QT ; 4半期
 MT ; 月
 WK ; 週
 DY ; 日

はそれぞれのデータの性質(集計単位)を示す。また2ヶ月単位のデータが必要なときには2MT, 5日単位のデータが必要なときには5DYとする。

4.2 データの抽出における応答形式

データ抽出の応答形式において内容的な処理は定義形式と同一であるが、定義形式のように変数の個数を人間が指定するのではなく、computerが主題分析と専門レベルの判定から変数の個数を定めて表示する。

応答の場合は選択表とINPUT表とがDisplayに表示される。選択表はさらに次の各構成要素の表からなる。

1. 変数選択表 Variable Table
2. 定義表 Definition Table
3. 選択子表 Option Table
4. 応答状況指定 Dialogue Status Instruction
5. 条件表 Condition Table

変数選択表(VT)には変数が並べられ各変数の下にある点(dot)にライトペンで指示を与えると、その変数が現在応答の対象として指定されたことを意味する。

定義表 (DT) は次の各項目からなる。

- ・ TYPE
- ・ SECTOR
- ・ PRIORITY
- ・ LAG
- ・ PERIOD
- ・ UNIT

それぞれの意味は定義形式と同じであり、各項目の左側の dot にライトペンで指定すれば、その項目について定義を行なうことができる。

定義はさらに選択子表 (OT) から TYPE の場合は X, Z, Y, D のいずれかを選び、数値の場合は SIGN から +, - のいずれかを選び、さらに NUMBER から 0, 1, 2, 3, 4, 5, 6, 7, 8, 9 のいずれかを選ぶことができる。数字の桁は次の項目の指定 (dot) によって定まる。

応答形式は必ず応答状況指定 (DSI) が表示されている。DSI は次のようである。

- ・ CONTINUE
- ・ NEXT
- ・ CHANGE
- ・ OK

同一変数について定義を繰り返す場合は CONTINUE あるいは CHANGE を指定し、変数を順次指定する場合は直接変数を指定するか NEXT を指定すればよい。

表示されている変数の処理が終わった場合あるいは次の Stage へ飛ぶ場合は OK を指定する。

1 つ 1 つの定義づけに対応して 1 箇所ずつ定義表の () の中がうめられ、最終的には INPUT 表に整理される。

Variable Table (VT)

VARIABLE (・X1, ・X2,)

Definition Table (DT)

- ・ TYPE (・)
- ・ SECTOR (・)
- ・ PRIORITY (・)
- ・ LAG (・ , ・)
- ・ PERIOD (・ , ・ , ・)
- ・ UNIT (・ , ・)

Option Table (OT)

TYPE X Z Y D

SIGN + -

NUMBER 0, 1, 2, 3, 4, 5, 6, 7, 8, 9

Dialogue Status Instruction

- CONTINUE
- NEXT
- CHANGE
- OK
- INPUT TABLE

5. データの分析 (Data Analysis)

応答形式にせよ定義形式にせよデータ抽出段階が終了すると、抽出された Selected File 上の各データは Priority, Sector, Time, の条件にしたがってデータ分析の段階に入る。

(1) データ分析の準備

JUMPS ORIGINAL FILE 作成の段階で各データは Validity check(p) を受け、さらに基本的統計量 (平均値, 分散) の計算を行なっている。データ分析段階では Selected File 上のデータの附加・修正が行なわれると、再び Validity Check および基本的統計量の計算を行なう。

個別のデータについて Time 条件の指定がなければ、データ分析の準備はこれで終了するが、Time Lag あるいは Period の指定があつて各 Sector 内のデータを揃える場合には、基本的統計量は再計算される。

データ期間を揃えるためには、

DEFINE PERIOD ; (α , β , r)

のように r に対して YR, BI, QT, MT, WK, DY を与えればよい。YR は BI に対して上位データといい、QT は BI に対しての下位データと呼ぶ。

DEFINE PERIOD ; (1945-1, 1970-4, QT)

を行なったにもかかわらず、データが QT の上位データ YR, BI の場合は、使用者が ORIGINAL FILE のデータ分析の結果をみて特別に指定する必要がある。もし指定がなければ機械的に 1 期間のデータの差分を $\frac{1}{2}$ あるいは $\frac{1}{4}$ したものに、期首のデータに下位データの期数を掛け加えて作る。すなわち、

$$YR_T + (t-1) * (YR_{T+1} - YR_T) / 4$$

ただし T : 上位データの期数

t : 下位データの期数

のように計算され、対象とするデータの概念がストックである場合は上式のままでよいが、フローの場合は YR_T の値は下位データの期数で割算して平均値を用いる。

またデータ分析の結果を使って指定する場合は

F (T) ; T (T), C (T), S (T)

T (T) ; 傾向修正値

C (T) ; 循環修正値

S (T) ; 季節修正値

が与えられることになる。

(2) データ分析の内容

データ分析の段階では個別データの

- i 傾向分析
- ii 周期解析
- iii 季節の調整
- iv その他

を行なう。

i) 傾向分析では最小二乗法を用いて次の関数型のパラメータが順次求められる。

- ・ 一次式
- ・ 二次式
- ・ $\vdots$
- ・ n 次式
- ・ 指数関数
- ・ 対数関数
- ・ ロジスティック（成長曲線）

ii) 周期解析は次の方法で行なうことができる。

- ・ ペリオドグラム
- ・ 自己相関母数推定
- ・ スペクトル解析

iii) 季節調整は次の方法で行なうことができる。

- ・ ハーバード法
- ・ 連環比率法
- ・ 移動平均法
- ・ 指数平滑法
- ・ センサスEPA法

(3) データの加工パターン

データの加工パターンを示すには Label PATTERNを書き、パターンを示した後、続いて変数あるいは変数群を指定する。

PATTERN: $F(T) = T(T) + C(T) + S(T) + D(T)$

$F(T)$; (v_i, s_j)

ただし $T(T)$; 傾向変動

$C(T)$; 循環変動

S (T) ; 季節変動
D (T) ; 不規則変動

5.1 データの分析における定義形式

データの分析は定義形式あるいは応答形式のいずれの形式でも行なうことができる。

定義形式の場合は必ず ANALYSIS という Label を書き、次にいかなる分析を実行させるか種類を書く。

データ分析の種類を指定する Label は

TREND 傾 向 値 (超勢変動)
CYCLE 周 期 解 析 (循環変動)
SEASON 調整指数解析 (季節変動)
GPFS GPFS 法
CENSUS センサス局法

などである。

この Label に導びかれてデータ分析を行なうために、さらに関数、方法を指定しなければならない。その場合

関 数 FUNCTION
方 法 METHOD

という Label を指定する。

FUNCTION の一般形式は、次のように表現される。

FUNCTION ; $f((n \text{ または } n'), v_i, s_j)$
 f 関数の型あるいは方法
 n 次 数
 n' 特 殊 指 定
 v_i 変 数 ($i = 1, 2, \dots, l$)
 s_j セクタあるいは変数群 ($j = 1, 2, \dots, m$)

関数の型には次のような形式がある。

LIN (v_i, s_j)
QUA (v_i, s_j)
POL ($(n), v_i, s_j$)
EXP ($(n), v_i, s_j$)

LOG (v_i, s_j)

LGT ((n'), v_i, s_j)

PSN ((n'), v_i, s_j)

変数 v_i はたとえば $X_1, X_2, \dots, X_\ell, Z_1, Z_2, \dots, Z_m$, セクタ S_j はたとえば SECTOR1, SECTOR2, $\dots$, SECTOR j ($X_1, X_2, \dots, X_k$) のように指定する。

n は次数を示し, ((1), $X_1, X_2, \dots, X_\ell$), ((3), X_2, X_5), ((4), $X_1, X_2, \dots$, SECTOR3), ((2), $X_1, X_2,$), ((3), $X_3, \dots, X_8,$ SECTOR4)) のいずれの形式をとってもよい。

((n), $X_1, X_2, \dots, X_\ell$) は (n) に導びかれたすべての変数 $X_1, X_2, \dots, X_\ell$ が n 次の多項式であり, ((n), X_1) の場合は X_1 だけが n 次の多項式である。

次数 (n) を (m, n) と指定すれば, m 次から n 次までの多項式についてパラメータを推定する。

なお ANALYSIS は推定の段階と検定の段階とを含む。ANALYSIS の変りに ESTIMATE を使えば推定を行ない, TEST を用いれば検定を行なう。ESTIMATE および TEST は各 Label の後に任意に指定することができる。

ANALYSIS ; TREND

① FUNCTION ; LIN (v_i, s_j)

この形式は直接最小二乗法によって, 時間 t に対して線型方程式 (linear function)

$$X_{i,t} = \alpha_{i,0} + \alpha_{i,1} t$$

のパラメータ $\alpha_{i,0}$ および $\alpha_{i,1}$ を推定する。

② FUNCTION ; QUA (v_i, s_j)

この形式は直接最小二乗法によって時間 t に対して二次方程式 (quadratic function)

$$X_{i,t} = \alpha_{i,0} + \alpha_{i,1} t + \alpha_{i,2} t^2$$

のパラメータ $\alpha_{i,0}$, $\alpha_{i,1}$ および $\alpha_{i,2}$ を推定する。

③ FUNCTION ; POL ((n), v_i, s_j)

この形式は直接最小二乗法によって時間 t に対して n 次方程式 (polynomial function)

$$X_{i,t} = \alpha_{i,0} + \alpha_{i,1} t + \alpha_{i,2} t^2 + \dots + \alpha_{i,n} t^n$$

のパラメータ $\alpha_{i,0}$, $\alpha_{i,1}$, $\alpha_{i,2}$, $\dots$, $\alpha_{i,n}$ を推定する。

FUNCTION ; POL ((1), v_i, s_j) は FUNCTION (v_i, s_j) と同じである。

④ FUNCTION ; EXP ((n), v_i, s_j)

この形式は指数関数

$$X_{i,t} = \alpha_0 \cdot e^{\alpha_1 t + \alpha_2 t^2 + \dots + \alpha_n t^n}$$

であり、機械内部では

$$\log X_{i,t} = \log \alpha_0 + \alpha_1 t + \alpha_2 t^2 + \dots + \alpha_n t^n$$

のパラメータを推定するから、 $X_{i,t} \leq 0$ であってはならない。もしデータが $X_{i,t} \leq 0$ ならば、この関数形式が使用不可能であることをコンピュータは警告する。

⑤ FUNCTION : LOG ((n'), v_i, s_j)

この形式は対数関数

$$X_{i,t} = \alpha_0 + \alpha_1 \log n't$$

であり、 n' は底の指定である。また $t > 0$ でなければならない。

⑥ FUNCTION : LGT ((n'), v_i, s_j)

この形式は Logistic 関数

$$X_{i,t} = C + \frac{K}{1 + \alpha_0 \cdot e^{\alpha_1 t}}$$

のパラメータ α_0 , α_1 を推定するが、 C および K は n' として (K, C) を指定しなければならない。もし (K, C) を指定しないときは n' に AUTO を書けば自動的に K および C を computer が選ぶ。

K, C の指定については応答形式を参照。

⑦ FUNCTION : PSN ((n'), v_i, s_j)

この形式は Pearson 型の分布関数を表わし、 n' を 1, 2, 3, 4, 5, 6, 7 と指定することによって pearson 型第 I, II, III, IV, V, VI, VII 型の分布関数のパラメータが推定される。

pearson 型については応答形式参照。

FUNCTION は上述のほか、使用者が FORTRAN 型式で具体的に関数の型を書くことができる。

たとえば次に示すようである。

FUNCTION : X1 = A0 + A1 * T + A2 * * T

循環変動あるいは季節変動を分析する場合は、Label にそれぞれ

ANALYSIS : CYCLE

ANALYSIS : SEASON

によって指定される。

この 2 つの分析にはそれぞれ固有の分析方法すなわち、循環変動は周期と振幅を、季節変動は周期内の指数を計算する固有の分析方法がある。さらにこの 2 つの分析に共通して用いられる方

法がある。

循環変動の分析を実行するLabelは

ANALYSIS ; CYCLE

である。周期解析の方法の固有の型には次の型式がある。

METHOD ; PGM((n'), v_i, s_j)

METHOD ; ACR((n'), v_i, s_j)

METHOD ; SRM((n'), v_i, s_j)

① METHOD ; PGM((n'), v_i, s_j)

この形式はPERIODGRAMによって周期を解析する。PGMのn'は(m, n)型であり(m < n), mとnの範囲で周期についての分散が計算され, 最小の分散をもつ周期が指定される。

② METHOD ; ACR((n'), v_i, s_j)

この形式はAuto Correlation Analysisを行ない, GORRELOGRAMによって周期解析を行なう。n'は(m, n)によって, 周期m期からn期までの変化を与える。自己相関係数をρとすれば

$$\rho_m = \frac{\text{cov}(y_t, y_{t+m})}{\sqrt{\text{var } y_t} \sqrt{\text{var } y_{t+m}}}$$

で与えられる。

③ METHOD ; SRM((n', c'), v_i, s_j)

この形式はSPECTRAM(調和解析)法を用いて, 周期および振幅について推定を行なう。n'は項数, c'は周期の指定であり, (n', c')の形をとる。このc'は周期の条件であり, c'の指定が1個で(λ)であれば, 真の周期が既知として関数

$$Y_{it} = \alpha_{i0} + \alpha_{i1} \cos \frac{2\pi}{\lambda} t + \alpha_{i2} \cos 2 \frac{2\pi}{\lambda} t + \dots + \alpha_{ik} \cos k \cdot \frac{2\pi}{\lambda} t + \dots + \beta_{i1} \sin \frac{2\pi}{\lambda} t + \dots + \beta_{i2} \sin 2 \cdot \frac{2\pi}{\lambda} t + \dots + \beta_{ik} \sin k \frac{2\pi}{\lambda} t + \dots$$

を用いて最小二乗法によりパラメータα_{i0}, α_{i1}, α_{i2}, …… α_{ik}, …… β_{i1}, β_{i2}, …… β_{ik}, …… が推定される。

c'を(μ₁, μ₂)のように指定した場合は, λは未知で, 関数

$$\alpha_{im} = \frac{2}{n} \sum Y_{it} \cos t \frac{2\pi}{\mu} t$$

$$\beta_{im} = \frac{2}{n} \sum Y_{it} \sin t \frac{2\pi}{\mu} t$$

を最小にするμがμ₁ ~ μ₂の範囲内で選ばれる。

季節変動を分析するLabelは

ANALYSIS ; SEASON

であって固有の方法には

METHOD ; HVD ((n'), v_i , s_j)

METHOD ; LRV ((n'), v_i , s_j)

METHOD ; MVA ((n'), v_i , s_j)

① METHOD ; HVD ((n'), v_i , s_j)

この形式はいわゆるハーバード法であり、 n' は繰り返し期間の指定である。 v_i は変数、 s_j は変数群の指定である。

② METHOD ; LRV ((n'), v_i , s_j)

この形式はLink Relative法であり、 n' は平均値の指定である。平均値には 算術平均AMと中位数MEとがある。

③ METHOD ; MAV ((n'), v_i , s_j)

この形式は移動平均法であり、 n' は (m , n) によって、平均をとる最小項数 m と最大項数 n の指定を行なう。 n' を指定しなければコンピュータは自動的に項数の値を $m=2$, 3 , ふやしていく。

TREND, CYCLE, SEASONに共通して用いられる方法に次の方法がある。

METHOD ; DIF ((n'), v_i , s_j)

METHOD ; ESM ((n'), v_i , s_j)

METHOD ; FTR ((n'), v_i , s_j)

① METHOD ; DIF ((n'), v_i , s_j)

この形式の方法は時系列について n 次の階差 (Difference) を求め、時系列の次数を打ち出す。 n' は (m , n) によって m は低次min, n は高次maxを指定する。

② METHOD ; ESM ((n'), v_i , s_j)

この形式の方法はExponential Smoothing法であり、

$$\text{新推定値} = \text{旧推定値} + \alpha (\text{新データ} - \text{旧データ})$$

の α を推定する。 n' は (m , n) によって、 m は最小項数を、 n は最大項数を指定する。

③ METHOD ; FTR ((n'), v_i , s_j)

この形式の方法は時系列についてFACTOR ANALYSISを行ない、 n' は (m , n) によって、 m は最小項数、 n は最大項数を指定する。

データ分析を行なう場合、各レベルあるいは変数または変数群に対してALL, OPTIMAL, SELECTIONを用いることができる。形式は

ANALYSIS ; ALL

FUNCTION ; ALL (SECTOR ALL)

FUNCTION ; OPTIMAL(v_i, s_j)

などである。

ALL機能とはたとえば

ANALYSIS ; ALL(SECTOR i)

と書けば、第 i セクタのすべてに対してTREND, CYCLE, SEASONの分析を行なうことを意味し、

ANALYSIS ; TREND

FUNCTION ; ALL(SECTOR ALL)

の場合はすべてのSECTORについてTREND分析のすべてを行なう。

OPTIMALの場合はALLの位置にOPTIMALを書けばよい。コンピュータは指定された条件にしたがって各種のテストの結果から最適な結果だけを提出する。

SELECTIONの場合はALLの位置にSELECTION(n')を書き、 n' は最適な結果を含めて n' で指定された個数を結果として打ち出す。

定義形式で用いることができる書式を以下に示す。

ANALYSIS ; TREND

FUNCTION ; LIN(v_i, s_j)

FUNCTION ; QUA(v_i, s_j)

FUNCTION ; POL((n) , v_i, s_j)

FUNCTION ; EXP((n) , v_i, s_j)

FUNCTION ; LOG((n') , v_i, s_j)

FUNCTION ; LGT((n') , v_i, s_j)

FUNCTION ; PSN((n') , v_i, s_j)

FUNCTION ; User指定

ANALYSIS ; CYCLE

METHOD ; PGM((n') , v_i, s_j)

METHOD ; ACR((n') , v_i, s_j)

METHOD ; SRM((n') , v_i, s_j)

ANALYSIS ; SEASON

METHOD ; HVD((n') , v_i, s_j)

METHOD ; LRV((n') , v_i, s_j)

METHOD ; MVA((n') , v_i, s_j)

共通して用いられる方法

METHOD ; DIF((n') , v_i, s_j)

METHOD : ESM((n'), v_i, s_j)

METHOD : FTR((n'), v_i, s_j)

その他の指定

ANALYSIS : ALL

FUNCTION : ALL(SECTOR ALL)

METHOD : OPTIMAL(v_i, s_j)

なおGPFS, CENSUS局法を用いるときは

ANALYSIS : GPFS(v_i, s_j)

ANALYSIS : CENSUS(v_i, s_j)

と指定する。

5.2 データの分析における応答形式

データ分析の応答形式は変数選択の段階で選ばれた個別変数およびSECTORの個別変数について分析が行なわれる。

選択表には次のように変数表, セクタ表が表示される。

VARIABLE (·X1, ·X2, X3,, Zn)

SECTOR (·S1, ·S2, S3,, Sk)

データ分析の応答形式は各変数毎にPATTERNを決めなければならない。

PATTERNは選択子表からPATTERN, ·T(T), ·C(T), ·S(T), ·D(T),
·ESM, GPFS, CENSUS, FTRおよびOPERATOR +, ·-, ·*, ·／を選ぶ。

結果はたとえば次のように順次表示される。

PATTERN(X1, T(T)+S(T))

PATTERN(X2, T(T)+C(T)+S(T))

PATTERN(X3, T(T)*S(T))

次にNEXTの指定をすれば, X1から分析が行なわれる。分析Step の応答は次のように行なわれる。

VARIABLE(X_i) ·DIF ·NORMAL

または SECTOR(S_j) ·DIF ·NORMAL

が表示され, 次に各選択項目を選択する。

TREND ·Y ·N

·LIN

```

      . QUA
      . POL (          )
      . EXP (          )
      . LOG (          )
      . LGT (          )
      . PSN (          )
      . USR (          )

OPTION  SIGN    + ,    -
        NUMBER  . 0, . 1, . 2, . 3, . 4, . 5, . 6, . 7,
                . 8, . 9
        BASE    . N ,    E
        DECIMAL .
CONDITION      . AUTO
                . SELECT (          )
                . OPTIMAL
                . ALL
DSI            . CONTINUE
                . NEXT
                . CHANGE
                . OK
                . EXECUTE

```

TRENDの・Y, ・NによってTRENDのSubstepをCANCELすることができる。何にも指示しなければ当然TREND分析を行なう。分析方法について複数の分析方法の指定をゆるし、分析方法の()内はOPTIONの指定によって埋められる。

VARIABLE TYPEの指定によって変数それ自身を扱うか、時間的な変化すなわち階差を与えることがどきる。両方を必要とするときはDIFとNORMALを同時に指定する。何にも指定しなければ、EXECUTEのStepでNORMALとして実行する。

CYCLEおよびSEASONの分析も、分析方法を除いてTRENDと同様である。

```

VARIABLE  (Xi)  . DIF  . NORMAL
SECTOR    (Sj)  . DIF  . NORMAL
CYCLE     . Y    . N
          . PGM (    ,    )
          . ACR (    ,    )

```

.SRM(, ,)
 OPTION SIGN + , -
 NUMBER 0 , 1 , 2 , 3 , 4 , 5 , 6 , 7 ,
 8 , 9 , 10
 OPERATOR + , - , * , /
 CONDITION . AUTO
 . SELECT()
 . OPTIMAL
 . ALL
 DSI . CONTINUE
 . NEXT
 . CHANGE
 . OK
 . EXECUTE
 VARIABLE (X_i) .DIF .NORMAL
 SECTOR (S_j) .DIF .NORMAL
 SEASON .Y .N
 .HVD(,)
 .LRV(,)
 .MVA(,)
 OPTION SIGN + , -
 NUMBER 0 , 1 , 2 , 3 , 4 , 5 , 6 ,
 7 , 8 , 9
 OPERATOR + , - , * , /
 MEAN .AM, .MEDIAN

なお、データ分析の1例として、頁98にGPTSをあげておく。

6. 変数間の分析 (Interrelation of Variables)

このステージの主たる目的は、主として数多い変数群からそれぞれの変数相互間の関係を調べることにより、次の Model Building のための重要な情報を提供することにある。ステップの主たる内容は、各変数間の相関分析、因子分析などである。

I) 単相関係数

いま2変量を y_i, y_j とすれば、 y_i, y_j の直接的な結びつきの程度を計る尺度である相関係数 r_{ij} は次のように与えられる。

$$r_{ij} = \frac{\sigma_{y_i y_j}}{\sigma_{y_i} \cdot \sigma_{y_j}}$$

ここで平均値を $\bar{y}_i, \bar{y}_j$, それぞれの分散を $\sigma_{y_i}^2, \sigma_{y_j}^2$, 共分散を $\sigma_{y_i y_j}$ とすると

$$\sigma_{y_i}^2 = \frac{1}{T} \sum_{t=1}^T (y_{it} - \bar{y}_i)^2$$

$$\sigma_{y_j}^2 = \frac{1}{T} \sum_{t=1}^T (y_{jt} - \bar{y}_j)^2$$

$$\sigma_{y_i y_j} = \frac{1}{T} \sum_{t=1}^T (y_{it} - \bar{y}_i)(y_{jt} - \bar{y}_j)$$

と表わせる。ただし t は期を示すもので $t = 1, 2, \dots, T$ である。

II) 因子分析

2変数 y_i, y_j が独立でなく、それぞれ第 t 番目の値 y_{it}, y_{jt} の関係を分析する場合に用いられる。この場合には機械内部で次のような相関表を準備する必要がある。

y_i と y_j の相関表

$y_j \backslash y_i$	p_1	-----	p_l	-----	p_m	f_{y_j}
q_1	f_{11}	-----	f_{l1}	-----	f_{m1}	$f_{\cdot 1}$
$\vdots$	$\vdots$		$\vdots$		$\vdots$	$\vdots$
q_k	f_{1k}	-----	f_{lk}	-----	f_{mk}	$f_{\cdot k}$
$\vdots$	$\vdots$		$\vdots$		$\vdots$	$\vdots$
q_n	f_{1n}	-----	f_{ln}	-----	f_{mn}	$f_{\cdot n}$
f_{y_i}	$f_{1\cdot}$	-----	$f_{l\cdot}$	-----	$f_{m\cdot}$	N

6.1 変数間の分析における定義形式

相関分析にはUnconditional分析とConditionalとがある。Unconditional 分析は抽出されたすべてのデータの組合せに対して分析を行ない、Conditional 分析は指定された変数あるいは変数群の組合せに対して分析を行なう。定義形式の一般型は次に示すようである。

ANALYSIS ; CORRELATION

UNCOND ; LEVEL(k)((n'))

COND ; LEVEL(k)((n'), v_i , s_j)

ここで(n')は(m, n)の形式によって、mは相関係数の高い順にm番目まで、nは相関係数のレベルを指定する。

たとえば次のような場合

COND ; LEVEL(k)((5, 0.98), SECTOR2)

セクタ2に含まれるすべての変数の組合せの中から、相関係数の値が0.98以上かあるいは高い順に5個までという、どちらかゆるい方の制限を満たすまで選ばれる。すなわち相関係数が0.98以上の組合せがすべて(5個以上)選ばれるが、あるいは数が少なければ0.98以下も含めて5個まで選ばれる。X1に関するOutput表示は

(X1, X2) = 0.9932

(X1, X3) = 0.9914

(X1, X6) = 0.9803

(X1, Z1) = 0.9800

(X1, Z4) = 0.9758

のようになる。

因子分析の定義形式は次のようである。

ANALYSIS ; CORRELATION

CRMRIX ; LEVEL(k)((n'₁), (n'₂), (n'₃), v_1 ,
 v_2)

(n'₁), (n'₂), (n'₃)はたとえば次のようである。

(V1₀, V2₀), (V1_m, V2_m), (C_{v1}, C_{v2})

ここでV1₀, V2₀は変数 v_1 , v_2 のそれぞれminimumとする値であり、V1_m, V2_mはそれぞれのmaximumを示す値である。C_{v1}, C_{v2}はそれぞれの変数のClass Intervalを示す値である。

6.2 変数間の分析における応答形式

相関分析 stage の応答形式では、はじめにコンピュータはすでに整理して選ばれた変数、および SECTOR を表示する。

各変数あるいは SECTOR について相関係数の指定を行なってレベルを示さなければならない。LEVEL の指定の方法は定義形式と同様である。すなわち k においてレベルの番号を、 m において相関係数の高い順に m 番目まで、 n において相関係数のレベルを指定する。

CRMRIX の場合は、定義形式と同様に 2 変数指定を行なわなければならない。

CORRELATION

·LEVEL () ()

·LEVEL () (,), (,), (,)

LEVEL ·CRMRIX

·UNCONDITIONAL

·CONDITIONAL

OPTION SIGN ·+ , ·-

NUMBER ·0 , ·1 , ·2 , ·3 , ·4 , ·5 , ·6 , ·7 , ·8 , ·9

DSI ·CONTINUE

·NEXT

·CHANGE

·OK

·EXECUTE

7 模型作成 (Model Building)

7.1 Model Building の内容

IDL S, SWLS, 2 SLS, 3 SLS, INSTVなどによってパラメータの推定を行なう場合、特別な指定のない限り必ずDL Sを通らなければならない。

また2 SLS, 3 SLS, INSTVなどによって、パラメータを推定する場合は必ずIDL Sを通らなければならない。もしこのルールが守られていなければcomputerは必ずDL S, またはIDL Sの推定過程を通ることを要求する。

しかしながらFORECAST, SIMULATION, OPTIMIZATIONなどのstageでEXECUTIONを行なった結果

JUMPでIDL S, 2 SLS, 3 SLS, INSTVおよびSWLSにもどるときは、すでにDL SあるいはIDL Sが行なわれているので必ずしもこのSubstep を通らなくてもよい。

1) 直接最小二乗法 (D. L. S. M) Direct Least Squares Method

二変量の相関々係を求める場合に、われわれはすでに直接最小二乗法によって、それぞれの変量 (X, Y) の回帰方程式のパラメータを計算した。この節では、さらに一般化して多変量の回帰方程式のパラメータを最小二乗法によって推定しよう。

いま、被説明変数、説明変数、観測誤差をそれぞれ次の記号によって示そう。

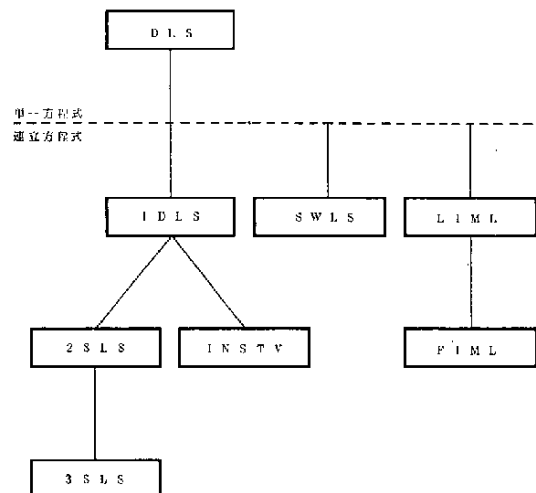
X_t 被説明変数

$Z_{1t}, Z_{2t}, \dots, Z_{kt}$ 説明変数

U_t 観測誤差

ここで、観測期間 $t = 1, 2, \dots, T$ において、 X_t, Z_{1t}, U_t の間に

$$\left. \begin{aligned} X_1 &= a_1 Z_{11} + a_2 Z_{21} + \dots + a_k Z_{k1} + U_1 \\ X_2 &= a_1 Z_{12} + a_2 Z_{22} + \dots + a_k Z_{k2} + U_2 \\ &\vdots \\ X_t &= a_1 Z_{1t} + a_2 Z_{2t} + \dots + a_k Z_{kt} + U_t \\ &\vdots \\ X_T &= a_1 Z_{1T} + a_2 Z_{2T} + \dots + a_k Z_{kT} + U_T \end{aligned} \right\} \text{---(1)}$$



のような関係が成り立つとする。たとえば、被説明変数 X_t を t 期の消費とし、法人所得を Z_{1t} 、個人業主所得を Z_{2t} 、勤労所得を Z_{3t} とし、 X_t と Z_{1t} 、 Z_{2t} 、 Z_{3t} の間に線型関係が認められれば、 X_t の理論値 X'_t に

$$X'_t = \alpha_1 Z_{1t} + \alpha_2 Z_{2t} + \alpha_3 Z_{3t}$$

のような関係が成り立つ。だが、現実の資料は理論値 X'_t ではなく、 X_t であるから、誤差項 U_t を含めて

$$X_t = \alpha_1 Z_{1t} + \alpha_2 Z_{2t} + \alpha_3 Z_{3t}$$

となる。前者を理論式 (Theoretical Equation) といい、後者を統計式 (Statistical Equation) という。これら両者は

$$\text{統計式} = \text{理論式} + \text{誤差項}$$

の関係を持つ。

直接最小二乗法とは、統計式において誤差項の2乗の総和を最小にするように、パラメータ α_1 、 α_2 、……、 α_k を求めることである。

計算の手順を単純化するために一般式(1)をベクトルおよび行列によって再定式化しよう。

$$\alpha = [\alpha_1 \quad \alpha_2 \quad \alpha_3 \quad \cdots \quad \alpha_k]$$

$$Z_t = [Z_{1t} \quad Z_{2t} \quad \cdots \quad Z_{kt}]$$

$$Z'_t : Z_t \text{ の転置行列}$$

とすれば、

$$X_t = \alpha_1 Z_{1t} + \alpha_2 Z_{2t} + \cdots + \alpha_k Z_{kt} + U_t$$

は、

$$X_t = \alpha Z'_t + U_t \quad (t=1, 2, \cdots, T) \quad \cdots \cdots \cdots (2)$$

したがって

$$U_t = X_t - \alpha Z'_t \quad \cdots \cdots \cdots (2)'$$

によって表わすことができる。

直接最小二乗法では U_t の総和 L を各パラメータ α_1 、 α_2 、……、 α_k について微分して0とおき、 α_1 、 α_2 、……、 α_k の推定値 $\hat{\alpha}_1$ 、 $\hat{\alpha}_2$ 、……、 $\hat{\alpha}_k$ を求めるのであるから、

$$L = \sum_{t=1}^T U_t^2 = \sum_{t=1}^T (X_t - \alpha Z'_t)^2 \quad (t=1, 2, \cdots, T)$$

ここで、 L を α_i について微分して0とおけば、

$$\begin{aligned} \frac{\partial L}{\partial \alpha_i} &= -2 \sum_{t=1}^T (X_t - \alpha Z'_t) Z_{it} \\ &= 0 \end{aligned}$$

パラメータ α_i は α_1 、 α_2 、……、 α_k の K 個あるから、それぞれについて微分し0とおけば、次の K 個の連立方程式を導くことができる。

$$\left. \begin{aligned} \sum_{t=1}^T X_t Z_{1t} &= \alpha_1 \sum_{t=1}^T Z_{1t} \cdot Z_{1t} + \alpha_2 \sum_{t=1}^T Z_{2t} \cdot Z_{1t} + \cdots + \alpha_k \sum_{t=1}^T Z_{kt} \cdot Z_{1t} \\ \sum_{t=1}^T X_t Z_{2t} &= \alpha_1 \sum_{t=1}^T Z_{1t} \cdot Z_{2t} + \alpha_2 \sum_{t=1}^T Z_{2t} \cdot Z_{2t} + \cdots + \alpha_k \sum_{t=1}^T Z_{kt} \cdot Z_{2t} \\ &\vdots \\ \sum_{t=1}^T X_t Z_{kt} &= \alpha_1 \sum_{t=1}^T Z_{1t} \cdot Z_{kt} + \alpha_2 \sum_{t=1}^T Z_{2t} \cdot Z_{kt} + \cdots + \alpha_k \sum_{t=1}^T Z_{kt} \cdot Z_{kt} \end{aligned} \right\} (3)$$

各方程式の両辺に $\frac{1}{T}$ を掛けると、

$$\left. \begin{aligned} \frac{1}{T} \sum_{t=1}^T X_t Z_{1t} &= \alpha_1 \frac{1}{T} \sum_{t=1}^T Z_{1t} \cdot Z_{1t} + \alpha_2 \frac{1}{T} \sum_{t=1}^T Z_{2t} \cdot Z_{1t} + \cdots + \alpha_k \frac{1}{T} \sum_{t=1}^T Z_{kt} \cdot Z_{1t} \\ \frac{1}{T} \sum_{t=1}^T X_t Z_{2t} &= \alpha_1 \frac{1}{T} \sum_{t=1}^T Z_{1t} \cdot Z_{2t} + \alpha_2 \frac{1}{T} \sum_{t=1}^T Z_{2t} \cdot Z_{2t} + \cdots + \alpha_k \frac{1}{T} \sum_{t=1}^T Z_{kt} \cdot Z_{2t} \\ &\vdots \\ \frac{1}{T} \sum_{t=1}^T X_t Z_{kt} &= \alpha_1 \frac{1}{T} \sum_{t=1}^T Z_{1t} \cdot Z_{kt} + \alpha_2 \frac{1}{T} \sum_{t=1}^T Z_{2t} \cdot Z_{kt} + \cdots + \alpha_k \frac{1}{T} \sum_{t=1}^T Z_{kt} \cdot Z_{kt} \end{aligned} \right\} (4)$$

方程式の各項はパラメータ $1, \alpha_1, \alpha_2, \cdots, \alpha_k$ をもつ二次の積率で表わすことができる。

ここで、

$$\begin{aligned} m_{XX} &= \frac{1}{T} \sum_{t=1}^T X_t \cdot X_t \\ m_{XZ_i} &= \frac{1}{T} \sum_{t=1}^T X_t \cdot Z_{it} \\ m_{Z_i Z_j} &= \frac{1}{T} \sum_{t=1}^T Z_{it} \cdot Z_{jt} \end{aligned}$$

と定義すれば、第(4)式は

$$\left. \begin{aligned} m_{XZ_1} &= \alpha_1 m_{Z_1 Z_1} + \alpha_2 m_{Z_2 Z_1} + \cdots + \alpha_k m_{Z_k Z_1} \\ m_{XZ_2} &= \alpha_1 m_{Z_1 Z_2} + \alpha_2 m_{Z_2 Z_2} + \cdots + \alpha_k m_{Z_k Z_2} \\ &\vdots \\ m_{XZ_k} &= \alpha_1 m_{Z_1 Z_k} + \alpha_2 m_{Z_2 Z_k} + \cdots + \alpha_k m_{Z_k Z_k} \end{aligned} \right\} (5)$$

となる。

さらに、

$$\begin{aligned} M_{XZ} &\equiv [m_{XZ_1} \ m_{XZ_2} \ \cdots \ m_{XZ_k}] \\ M_{XZ}' &\equiv M_{ZX} \\ \alpha &\equiv [\alpha_1 \ \alpha_2 \ \cdots \ \alpha_k] \\ 0 &\equiv [0 \ 0 \ \cdots \ 0] \\ M_{ZZ} &\equiv \begin{bmatrix} m_{Z_1 Z_1} & m_{Z_2 Z_1} & \cdots & m_{Z_k Z_1} \\ m_{Z_1 Z_2} & m_{Z_2 Z_2} & \cdots & m_{Z_k Z_2} \\ \vdots & \vdots & \ddots & \vdots \\ m_{Z_1 Z_k} & m_{Z_2 Z_k} & \cdots & m_{Z_k Z_k} \end{bmatrix} \end{aligned}$$

と定義すれば、第(5)式は

$$M_{ZX} = M_{ZZ} \alpha'$$

または、

$$M_{ZZ} \alpha' = M_{ZX}$$

であるから、 α' の推定値 $\hat{\alpha}' \equiv [\hat{\alpha}_1 \quad \hat{\alpha}_2 \quad \dots \hat{\alpha}_k]$ は

$$\hat{\alpha}' = M_{ZZ}^{-1} M_{ZX} \dots \dots \dots (6)$$

このようにして、 $\alpha_1, \alpha_2, \dots, \alpha_k$ の最小二乗推定値を求めることができる。

次に、誤差の分散を求めてみよう。

この推定値をもとにして、観察値 Z_{it} を与えれば、被説明変数の理論値 X'_t が導かれるが、最小二乗法によって求められた理論式がどの程度現実を表現しているかを知るために、理論値と観察値との関係について調べてみなければならない。すなわち、観察値 X_t と理論値 X'_t との関係を調べるには $(X_t - X'_t = U_t)$ 誤差についてその分散を調べてみればよい。

誤差の分散を求めるには、 K 個の未知パラメータおよび T 個の観察値が与えられると、前述の分散分析で示したように、自由度を修正しなければならない。すなわち自由度は観察値の個数 T から未知パラメータの個数を引いた $T - K$ である。いま、誤差の分散を σ_u^2 とし、自由度修正済誤差の分散を S^2 とすれば、

$$\sigma_u^2 = \frac{1}{T} \sum_{t=1}^T U_t^2$$

$$S^2 = \frac{1}{T-K} \sum_{t=1}^T U_t^2 = \frac{1}{T-K} \sum_{t=1}^T (X_t - \hat{\alpha}' Z'_t)^2$$

ここで、 $\sum_{t=1}^T (X_t - \hat{\alpha}' Z'_t)^2$ の $\hat{\alpha}$ は一定であるから、

$$\begin{aligned} \sum_{t=1}^T (X_t - \hat{\alpha}' Z'_t)^2 &= \sum_{t=1}^T (X_t - \hat{\alpha}_1 Z_{1t} - \hat{\alpha}_2 Z_{2t} \dots - \hat{\alpha}_k Z_{kt})^2 \\ &= \sum_{t=1}^T \left\{ \begin{bmatrix} 1 & -\hat{\alpha}_1 & -\hat{\alpha}_2 & \dots & -\hat{\alpha}_k \end{bmatrix} \begin{bmatrix} X_t \\ Z_{1t} \\ Z_{2t} \\ \vdots \\ Z_{kt} \end{bmatrix} \right\}^2 \\ &= \begin{bmatrix} 1 & -\hat{\alpha}_1 & -\hat{\alpha}_2 & \dots & -\hat{\alpha}_k \end{bmatrix} \sum_{t=1}^T \begin{bmatrix} X_t \\ Z_{1t} \\ Z_{2t} \\ \vdots \\ Z_{kt} \end{bmatrix} \begin{bmatrix} X_t & Z_{1t} & Z_{2t} & \dots & Z_{kt} \end{bmatrix} \begin{bmatrix} 1 \\ -\hat{\alpha}_1 \\ -\hat{\alpha}_2 \\ \vdots \\ -\hat{\alpha}_k \end{bmatrix} \end{aligned}$$

$$\begin{aligned}
&= (1 - \hat{\alpha}) \sum_{t=1}^T \begin{bmatrix} X_t X_t & X_t Z_{1t} & X_t Z_{2t} & \cdots & X_t Z_{kt} \\ X_t Z_{1t} & Z_{1t} Z_{1t} & Z_{1t} Z_{2t} & \cdots & Z_{1t} Z_{kt} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ X_t Z_{kt} & Z_{kt} Z_{1t} & Z_{kt} Z_{2t} & \cdots & Z_{kt} Z_{kt} \end{bmatrix} \begin{bmatrix} 1 \\ -\hat{\alpha}_1 \\ -\hat{\alpha}_2 \\ \vdots \\ -\hat{\alpha}_k \end{bmatrix} \\
&= (1 - \hat{\alpha}) T \begin{bmatrix} \frac{1}{T} \sum X_t X_t & \frac{1}{T} \sum X_t Z_{1t} & \cdots & \frac{1}{T} \sum X_t Z_{kt} \\ \frac{1}{T} \sum X_t Z_{1t} & \frac{1}{T} \sum Z_{1t} Z_{1t} & \cdots & \frac{1}{T} \sum Z_{1t} Z_{kt} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{1}{T} \sum X_t Z_{kt} & \frac{1}{T} \sum Z_{kt} Z_{1t} & \cdots & \frac{1}{T} \sum Z_{kt} Z_{kt} \end{bmatrix} \begin{bmatrix} 1 \\ -\hat{\alpha}_1 \\ \vdots \\ -\hat{\alpha}_k \end{bmatrix} \\
&= (1 - \hat{\alpha}) T \begin{bmatrix} m_{XX} & M_{XZ} \\ M_{ZX} & M_{ZZ} \end{bmatrix} \begin{bmatrix} 1 \\ -\hat{\alpha}' \end{bmatrix} \\
&= T \{ m_{XX} - \hat{\alpha} M_{ZX} M_{XZ} - \hat{\alpha} M_{ZZ} \} \begin{bmatrix} 1 \\ -\hat{\alpha}' \end{bmatrix} \\
&= T \{ m_{XX} - \hat{\alpha} M_{ZX} - (M_{XZ} - \hat{\alpha} M_{ZZ}) \hat{\alpha}' \} \\
&= T \{ m_{XX} - \hat{\alpha} M_{ZX} - M_{XZ} \hat{\alpha}' + \hat{\alpha} M_{ZZ} \hat{\alpha}' \} \\
&\quad (\text{ここで, } \hat{\alpha}' = M_{ZZ}^{-1} M_{ZX} \text{ であるから}) \\
&= T \{ m_{XX} - \hat{\alpha} M_{ZX} - M_{XZ} M_{ZZ}^{-1} M_{ZX} + \hat{\alpha} M_{ZZ} M_{ZZ}^{-1} M_{ZX} \} \\
&= T \{ m_{XX} - M_{XZ} M_{ZZ}^{-1} M_{ZX} \}
\end{aligned}$$

ゆえに,

$$S^2 = \frac{T}{T-K} [m_{XX} - M_{XZ} M_{ZZ}^{-1} M_{ZX}]$$

このようにして、自由度修正済の誤差の分散を求めることができるが、ここで注意せねばならぬことは M_{ZZ} の逆行列 M_{ZZ}^{-1} において、行列式 $\det M_{ZZ}$ が 0 であってはならない。もし、 $\det M_{ZZ} = 0$ ならば逆行列 M_{ZZ}^{-1} を解くことができない。すなわち、ある二つの行、あるいは二つの列に同じ値の数字が排列されている場合は、

$$\begin{aligned}
M_{ZZ} &= \begin{bmatrix} m_{Z1} & Z_1 & m_{Z2} & Z_2 & \cdots & m_{Zk} & Z_k \\ m_{Z1} & Z_1 & m_{Z2} & Z_2 & \cdots & m_{Zk} & Z_k \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ m_{Z1} & Z_1 & m_{Z2} & Z_2 & \cdots & m_{Zk} & Z_k \end{bmatrix} \\
&= 0
\end{aligned}$$

であるが、このような現象は説明変数の特定の二変数が全く同一な変数であったり、相関関数がきわめて強い類似の変数である場合に生ずる。これを多重共変性 (Multi-Collinearity) と呼ぶ。

次に直接最小二乗法によってパラメータ α を推定した場合に、もう一つの重要な問題はパラメータの分散である。

前にのべたように、

$$\alpha = [\alpha_1 \quad \alpha_2 \quad \dots \alpha_k]$$

$$\hat{\alpha} = [\hat{\alpha}_1 \quad \hat{\alpha}_2 \quad \dots \hat{\alpha}_k]$$

$$M_{ZX} = \begin{bmatrix} m_{XZ_1} \\ m_{XZ_2} \\ \vdots \\ m_{XZ_k} \end{bmatrix} = \begin{bmatrix} \frac{1}{T} \sum_t X_t \cdot Z_{1t} \\ \frac{1}{T} \sum_t X_t \cdot Z_{2t} \\ \vdots \\ \frac{1}{T} \sum_t X_t \cdot Z_{kt} \end{bmatrix}$$

とすれば、第(2)式 ($X_t = \alpha Z' + U_t$) により、

$$\begin{aligned} M_{ZX} &= \begin{bmatrix} \frac{1}{T} \sum_t (\alpha Z' + U_t) Z_{1t} \\ \frac{1}{T} \sum_t (\alpha Z' + U_t) Z_{2t} \\ \vdots \\ \frac{1}{T} \sum_t (\alpha Z' + U_t) Z_{kt} \end{bmatrix} \\ &= M_{ZZ} \alpha' + \begin{bmatrix} \frac{1}{T} \sum_t Z_{1t} U_t \\ \frac{1}{T} \sum_t Z_{2t} U_t \\ \vdots \\ \frac{1}{T} \sum_t Z_{kt} U_t \end{bmatrix} \end{aligned}$$

ここで右辺第2項を

$$M_{ZU} = \begin{bmatrix} \frac{1}{T} \sum_t Z_{1t} U_t \\ \frac{1}{T} \sum_t Z_{2t} U_t \\ \vdots \\ \frac{1}{T} \sum_t Z_{kt} U_t \end{bmatrix}$$

と定義すれば、

$$M_{ZX} = M_{ZZ} \alpha' + M_{ZU}$$

したがって第6)式より

$$\begin{aligned}\hat{\alpha}' &= M_{ZZ}^{-1} M_{ZX} \\ &\quad - M_{ZZ}^{-1} (M_{ZZ} \alpha' + M_{ZU}) \\ &= \alpha' + M_{ZZ}^{-1} M_{ZU}\end{aligned}$$

$\hat{\alpha}'$ の平均値を求めれば、

$$E(\hat{\alpha}') = \alpha' + E(M_{ZZ}^{-1} M_{ZU})$$

いま、 $E(M_{ZZ}^{-1} M_{ZU})$ が0ならば、 $E(\hat{\alpha}') = \alpha'$ であるから、標本抽出回数を大きくすれば、パラメータの推定値の平均値が理論値に近づいてくる。すなわち、最小二乗法によるパラメータの推定値は不偏推定値である。したがってパラメータの推定値は不偏性をもつことがわかる。

つぎに、推定されたパラメータ $\hat{\alpha}$ の分散共分散行列を $\Sigma \hat{\alpha}$ という記号で表わせば、

$$\Sigma \hat{\alpha} = \begin{bmatrix} \text{Var}(\hat{\alpha}_1) & \text{Cov}(\hat{\alpha}_1, \hat{\alpha}_2) & \cdots & \text{Cov}(\hat{\alpha}_1, \hat{\alpha}_k) \\ \text{Cov}(\hat{\alpha}_2, \hat{\alpha}_1) & \text{Var}(\hat{\alpha}_2) & \cdots & \text{Cov}(\hat{\alpha}_2, \hat{\alpha}_k) \\ \vdots & \vdots & \ddots & \vdots \\ \text{Cov}(\hat{\alpha}_k, \hat{\alpha}_1) & \text{Cov}(\hat{\alpha}_k, \hat{\alpha}_2) & \cdots & \text{Var}(\hat{\alpha}_k) \end{bmatrix}$$

である。ここで $\text{Var}(\hat{\alpha}_i)$ は $\hat{\alpha}_i$ の分散、 $\text{Cov}(\hat{\alpha}_i, \hat{\alpha}_j)$ は $\hat{\alpha}_i, \hat{\alpha}_j$ の共分散である。

$\Sigma \hat{\alpha}$ はさらに次のように展開することができる。

$$\Sigma \hat{\alpha} = E\{[\hat{\alpha}' - E(\hat{\alpha}')] \cdot [\hat{\alpha}' - E(\hat{\alpha}')]'\} \cdots \cdots (7)$$

この証明は次のように明らかである。

$$\begin{aligned}\Sigma \hat{\alpha} &= E \left\{ \begin{bmatrix} \hat{\alpha}_1 - E(\hat{\alpha}_1) \\ \hat{\alpha}_2 - E(\hat{\alpha}_2) \\ \vdots \\ \hat{\alpha}_k - E(\hat{\alpha}_k) \end{bmatrix} \cdot [\hat{\alpha}_1 - E(\hat{\alpha}_1), \hat{\alpha}_2 - E(\hat{\alpha}_2), \cdots, \hat{\alpha}_k - E(\hat{\alpha}_k)] \right\} \\ &= E \left\{ \begin{bmatrix} (\hat{\alpha}_1 - E(\hat{\alpha}_1))^2 & (\hat{\alpha}_1 - E(\hat{\alpha}_1))(\hat{\alpha}_2 - E(\hat{\alpha}_2)) & \cdots & (\hat{\alpha}_1 - E(\hat{\alpha}_1))(\hat{\alpha}_k - E(\hat{\alpha}_k)) \\ (\hat{\alpha}_2 - E(\hat{\alpha}_2))(\hat{\alpha}_1 - E(\hat{\alpha}_1)) & (\hat{\alpha}_2 - E(\hat{\alpha}_2))^2 & \cdots & (\hat{\alpha}_2 - E(\hat{\alpha}_2))(\hat{\alpha}_k - E(\hat{\alpha}_k)) \\ \vdots & \vdots & \ddots & \vdots \\ (\hat{\alpha}_k - E(\hat{\alpha}_k))(\hat{\alpha}_1 - E(\hat{\alpha}_1)) & (\hat{\alpha}_k - E(\hat{\alpha}_k))(\hat{\alpha}_2 - E(\hat{\alpha}_2)) & \cdots & (\hat{\alpha}_k - E(\hat{\alpha}_k))^2 \end{bmatrix} \right\}\end{aligned}$$

第 ℓ 標本のパラメータ α'_ℓ の推定値 $\hat{\alpha}'_\ell$ は

$$\hat{\alpha}'_\ell = E(\alpha'_\ell) + M_{ZZ}^{-1} M_{ZU} \ell$$

であるから

$$\hat{\alpha}'_\ell - E(\alpha'_\ell) = (M_{ZZ}^{-1} M_{ZU})_\ell$$

従って

$$\begin{aligned}\Sigma \hat{\alpha} &= E\{[M_{ZZ}^{-1} M_{ZU}][M_{ZZ}^{-1} M_{ZU}]'\} \\ &= E\{M_{ZZ}^{-1} M_{ZU} M_{UZ} M_{ZZ}^{-1}\}\end{aligned}$$

ここで,

$$M_{ZU} = \begin{pmatrix} \frac{1}{T} \sum_{t=1}^T Z_{1t} U_t \\ \frac{1}{T} \sum_{t=1}^T Z_{2t} U_t \\ \vdots \\ \frac{1}{T} \sum_{t=1}^T Z_{kt} U_t \end{pmatrix} = \begin{pmatrix} \frac{1}{T} Z_{11} & \frac{1}{T} Z_{12} & \frac{1}{T} Z_{13} & \cdots & \frac{1}{T} Z_{1T} \\ \frac{1}{T} Z_{21} & \frac{1}{T} Z_{22} & \frac{1}{T} Z_{23} & \cdots & \frac{1}{T} Z_{2T} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \frac{1}{T} Z_{k1} & \frac{1}{T} Z_{k2} & \frac{1}{T} Z_{k3} & \cdots & \frac{1}{T} Z_{kT} \end{pmatrix} \begin{pmatrix} U_1 \\ U_2 \\ \vdots \\ U_T \end{pmatrix}$$

さらに,

$$M_{Zk} = \begin{pmatrix} \frac{1}{T} Z_{11} & \frac{1}{T} Z_{12} & \cdots & \frac{1}{T} Z_{1T} \\ \frac{1}{T} Z_{21} & \frac{1}{T} Z_{22} & \cdots & \frac{1}{T} Z_{2T} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{1}{T} Z_{k1} & \frac{1}{T} Z_{k2} & \cdots & \frac{1}{T} Z_{kT} \end{pmatrix}$$

$$U = [U_1 \quad U_2 \quad \cdots \quad U_T]$$

とすれば,

$$M_{ZU} = M_{Zk} U'$$

$$M_{ZU} = U M'_{Zk}$$

したがって,

$$\sum \hat{\alpha} = E \{ M_{ZZ}^{-1} M_{Zk} U' U M'_{Zk} M_{ZZ}^{-1} \}$$

U は誤差項であるから, U の平均値 $E(U) = 0$, 分散 $E(U_t)^2 = \sigma_{uu}$, 共分散 $E(U_s U_t) = \sigma_{st} = 0$ とすれば, 誤差項の分散共分散行列

$$E(U' U) = E \begin{pmatrix} U_1 U_1 & U_1 U_2 & \cdots & U_1 U_T \\ U_2 U_1 & U_2 U_2 & \cdots & U_2 U_T \\ \vdots & \vdots & \ddots & \vdots \\ U_T U_1 & U_T U_2 & \cdots & U_T U_T \end{pmatrix}$$

において, 対角行列すなわち分散を残して 0 であるから

$$E(U' U) = \begin{bmatrix} \sigma_{11} & 0 & \cdots & 0 \\ 0 & \sigma_{22} & \cdots & 0 \\ \vdots & & \ddots & \\ 0 & 0 & \cdots & \sigma_{TT} \end{bmatrix} = \sigma_u^2 I$$

ここに,

$$I = \begin{bmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 \\ \vdots & & \ddots & \\ 0 & 0 & \cdots & 1 \end{bmatrix}$$

したがって,

$$\begin{aligned} \sum \hat{\alpha} &= M_{ZZ}^{-1} M_{Zkt} E(U' U) M'_{Zkt} M_{ZZ}^{-1} \\ &= \sigma_u^2 M_{ZZ}^{-1} M_{Zkt} M'_{Zkt} M_{ZZ}^{-1} \end{aligned}$$

ここで

$$M_{Zkt} M'_{Zkt} = \frac{1}{T} M_{ZZ}$$

であるから,

$$\sum \hat{\alpha} = \frac{\sigma_u^2}{T} M_{ZZ}^{-1}$$

ゆえに, パラメータの分散共分散行列の推定値は $\sum \hat{\alpha} = \frac{\sigma_u^2}{T} M_{ZZ}^{-1}$ である。実際には, 資料の数に制約されるから, 分散は自由度修正済の分散 S^2 を利用して,

$$\sum \hat{\alpha} = \frac{S^2}{T} M_{ZZ}^{-1}$$

により計算を行う。この式の右辺の行列の対角要素がパラメータ α_i の推定値 $\hat{\alpha}_i$ の分散である。

<直接最小二乗法の諸条件>

直接最小二乗法によるパラメータの推定はどの理論模型, あるいはデータによっても求められるものではない。直接最小二乗法は次にのべるような厳しい条件にのみに適用されるパラメータの推定方法である。

いま, 推定すべきパラメータをもつ方程式を

$$X_t = \alpha_1 Z_{1t} + \alpha_2 Z_{2t} + \cdots + \alpha_k Z_{kt} + U_t \quad (t=1, 2, \cdots, T)$$

とすれば, 誤差項 U_t は次の2つの条件のいずれかに従っている。

第1に, 誤差項 U_t は平均値 $E(U_t) = 0$ と分散 $E(U_t)^2 = \sigma_{uu}$ $E(U_{\ell m}) = 0$ をもち, U_t が同一の正規分布に属するという。

正規分布と独立性の仮定。もし, U_{ℓ} と U_m とが独立性を仮定できず, $U_{\ell} = g(U_m)$ ならば, 直接最小二乗法によって推定されるパラメータをもつ原方程式の変数の選び方が不適当であることを意味する。ここにいう誤差項 U_t の正規分布の仮定は極めて厳重な一つの仮定である。

第2に、誤差項 U_t は平均値 $E(U_t) = 0$ と一定の分散 $E(U_t^2) = \sigma_{uu}$ をもつ同一の分布に属し、 $E(U_{tm}) = 0$ の仮定である。この仮定では分布の型は特に定められていない。さらに、直接最小二乗法で推定されるパラメータをもつ方程式の変数の性質によって、次の三つに分類することができる。

A: すべての説明変数 ($Z_{1t}, Z_{2t}, \dots, Z_{kt}$) が外生変数である場合。

B: 説明変数 ($Z_{1t}, Z_{2t}, \dots, Z_{kt}$) が外生変数とラグ付き内生変数で構成されている場合。

C: 説明変数が先決変数ばかりでなく、同時従属変数を含む場合。

上に分類した誤差項の二つの条件と変数についての三つの状態の組み合わせによって、直接最小二乗法による推定値の特性は変化する。

ある模型の未知のパラメータ $\alpha = [\alpha_{j1}, \alpha_{j2}, \dots, \alpha_{jk}]$ の推定値は標本のとり方によって種々の値を示すが、ある標本 j による推定値を

$$\hat{\alpha}_j = [\hat{\alpha}_{j1}, \hat{\alpha}_{j2}, \dots, \hat{\alpha}_{jk}]$$

とすれば、推定値の平均値 $E(\hat{\alpha}') = \frac{1}{L} \sum_{j=1}^L (\hat{\alpha}'_j)$

$$E(\hat{\alpha}') = \alpha' + E(M_{ZZ}^{-1} \cdot M_{ZU})$$

である。

いま標本数 L を増加させることによって $E(M_{ZZ}^{-1} M_{ZU}) \rightarrow 0$ ならば、最小二乗法による推定値 $\hat{\alpha}'$ は不偏性をもっていることがわかる。また、標本に含まれる観測数 $T \rightarrow \infty$ のときに、

$M_{ZZ}^{-1} M_{ZU} \rightarrow 0$ に確率収束するならば、 $\hat{\alpha}'$ は一貫性をもっている。さらに、 M_{ZZ}^{-1} が 0 ならば、式 $\hat{\alpha}' = M_{ZZ}^{-1} M_{ZU}$ において解は存在しない。すなわち、Multicollinearity の問題が生じる。このように、最小二乗法においては、 $M_{ZZ}^{-1} M_{ZU}$ の項がきわめて重要な役割を果たしている。

また、 $M_{ZZ}^{-1} M_{ZU}$ が 0 であるためには、 $M_{ZZ}^{-1} \neq 0$ であるなら、 $M_{ZU} = 0$ でなければならない。すなわち、

$$M_{ZU} = \begin{pmatrix} \frac{1}{T} \sum_t Z_{1t} U_t \\ \frac{1}{T} \sum_t Z_{2t} U_t \\ \vdots \\ \frac{1}{T} \sum_t Z_{kt} U_t \end{pmatrix} = \begin{pmatrix} \frac{1}{T} Z_{11}, \frac{1}{T} Z_{12}, \dots, \frac{1}{T} Z_{1T} \\ \frac{1}{T} Z_{21}, \frac{1}{T} Z_{22}, \dots, \frac{1}{T} Z_{2T} \\ \vdots \\ \frac{1}{T} Z_{k1}, \frac{1}{T} Z_{k2}, \dots, \frac{1}{T} Z_{kT} \end{pmatrix} \begin{pmatrix} U_1 \\ U_2 \\ \vdots \\ U_T \end{pmatrix}$$

$$= \begin{pmatrix} 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}$$

であって、 $\frac{1}{T} Z_{it} \neq 0$ であるから、 $U_t (t=1, 2, \dots, T) = 0$ でなければならない。したがって

$M_{ZU}=0$ となるためには、 Z_{it} のいかんにかかわらず、 $U_t=0$ であって、 Z_{it} と U_t は相互に独立でなければならない。もし、説明変数 Z_{it} と U_t との間に何等かの関数関係が存在すれば、 $M_{ZU} \neq 0$ である。

誤差項の二つの条件と説明変数の三つの条件を組み合わせて、直接最小二乗法による推定値の持つ特性をつぎのように分類することができるであろう。

説明変数 \ 誤差項	$E(U_t)=0$ σ_{uu} 一定 正規分布, 時間に独立		$E(U_t)=0$ σ_{uu} 一定 同一分布, 時間的に独立	
	小標本特性値	大標本特性値	小標本特性値	大標本特性値
A 外性変数 だけ	不偏性, 有効性 最良不偏性	一致性 漸近的正規性 漸近の有効性	不偏性, 有効性 最良不偏性	一致性
B 外生変数 とラグ付 内生変数		一致性 漸近的正規性 漸近の有効性		一致性
C 同時従属 変数を含む				

このようにして、直接最小二乗法によるパラメータの推定において、説明変数の性質によってはたとえば説明変数に同時従属変数が含まれていれば、 X の特性値は小標本特性値（不偏性、有効性、最良不偏性）および大標本特性値（一致性、漸近的正規性、漸近の有効性）のいずれをももたないから、このような方程式のパラメータの推定に直接最小二乗法を用いることは推定理論の上から望ましくない。

ここで、すべての同時従属変数を被説明変数とし、すべての説明変数を外生変数とするような試みが必要である。このような考え方から間接最小二乗法によるパラメータの推定が行われる。

2) 間接最小二乗法 Indirect Least Squares Method (I. L. S. M)

ヘンリー・シュルツ (Henry Schultz) が示してから、しばしば用いられる例として、物価と数量との関係から需要曲線と供給曲線を推定する方法がある。いま、図7-1のように、価格を Y 軸にとり、その財の取引量を X 軸方向にとれば、需要曲線は $X-Y$ 平面上で右上りの曲線となり、供給曲線は右下りの曲線となる。

われわれが何等かの目的をもって、これらの両曲線の形を決めるときに、具体的な統計数値を集めて曲線のパラメータを推定する。

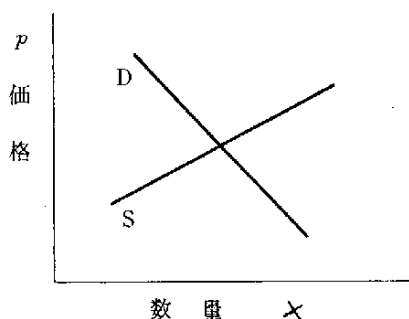


図 7-1

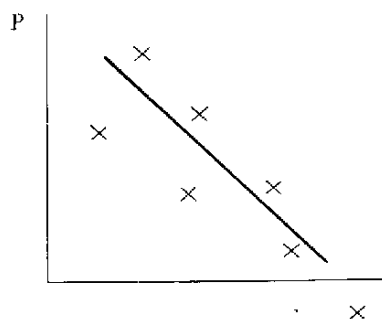


図 7-2

だが、図 7-2 に示したように、集められた統計数値を $X-Y$ 平面上にプロットして、直接最小二乗法によって曲線のパラメータを推定することにはどのような意味があるであろうか。

この図に示された曲線は需要曲線でもなければ、また供給曲線でもない。いま、需要曲線と供給曲線とが線型であるとすれば、ある財の需要量 X^D と供給量 X^S とは、

$$X^D = \alpha_0 - \alpha_1 P$$

$$X^S = \beta_0 + \beta_1 P$$

で表わすことができる。この式において X^D , X^S はそれぞれ需要量、供給量、 P は価格、 α_0 , α_1 , β_0 , β_1 はパラメータである。われわれが手にする統計数値からでは、これらの方程式のパラメータ、 α_0 , α_1 , β_0 , β_1 を推定することはできない。なぜならば純粹に需要方程式だけについて価格と需要量とを実験的に求めることができず、従って、統計数値そのものは図 7-3 に示

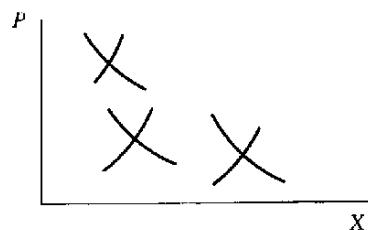


図 7-3

したように、それぞれが需要と供給との均衡した点であるからである。

いま、 X^D を λ 倍して X^S に加えると、 $X^D = X^S = X$ として、

$$\lambda X^D = \lambda (\alpha_0 - \alpha_1 P)$$

$$+) X^S = \beta_0 + \beta_1 P$$

$$X(1+\lambda) = (\lambda\alpha_0 + \beta_0) + (\beta_1 - \lambda\alpha_1)P$$

$$X = \frac{\lambda\alpha_0 + \beta_0}{1+\lambda} + \frac{(\beta_1 - \lambda\alpha_1)P}{1+\lambda}$$

この式は需要方程式であろうか、供給方程式であろうか、右辺の第二項の分子を置きかえると、

$$X = \frac{\lambda\alpha_0 + \beta_0}{1+\lambda} - \frac{(\lambda\alpha_1 - \beta_1)P}{1+\lambda}$$

したがって、需要方程式でも供給方程式でもない。

いま、ある財の需要量 X^D は価格 P と所得水準によって影響をうけ、その財の供給量 X^S は価格 P と生産性 W によって影響をうけるとすればどうであろうか。

$$X^D = \alpha_0 - \alpha_1 P + \alpha_2 Y$$

$$X^S = \beta_0 + \beta_1 P + \beta_2 W$$

均衡において、 $X^D = X^S$

したがって、この二つの方程式から

$$P = \frac{\alpha_0 - \beta_0}{\alpha_1 + \beta_1} + \frac{\alpha_2}{\alpha_1 + \beta_1} Y - \frac{\beta_2}{\alpha_1 + \beta_1} W \quad \dots\dots\dots(1)$$

同様にして

$$X = \frac{\alpha_0 \beta_1 + \alpha_1 \beta_0}{\alpha_1 + \beta_1} + \frac{\alpha_2 \beta_1}{\alpha_1 + \beta_1} Y + \frac{\alpha_1 \beta_2}{\alpha_1 + \beta_1} W \quad \dots\dots\dots(2)$$

この二式を基本的な方程式から導き出された誘導型 (Reduced form) と呼び、これらの誘導型の方程式にはそれぞれ同時従属変数を含まず、ある期の所得水準と生産性とは与えられれば、それぞれ個別の直接最小二乗法によって解くことができる。

すなわち、

$$\left. \begin{aligned} A &= \frac{\alpha_0 - \beta_0}{\alpha_1 + \beta_1}, & B &= \frac{\alpha_2}{\alpha_1 + \beta_1}, & C &= \frac{\beta_2}{\alpha_1 + \beta_1} \\ D &= \frac{\alpha_0 \beta_1 + \alpha_1 \beta_0}{\alpha_1 + \beta_1}, & E &= \frac{\alpha_2 \beta_1}{\alpha_1 + \beta_1}, & F &= \frac{\alpha_1 \beta_2}{\alpha_1 + \beta_1} \end{aligned} \right\} \dots\dots\dots(3)$$

とすれば、(1)および(2)式は、

$$P = A + BY - CW$$

$$X = D + EY + FW$$

であって、各方程式のパラメータ A, B, C および D, E, F を直接最小二乗法で推定することになる。いま、それぞれのパラメータの推定値を $\hat{A}, \hat{B}, \hat{C}, \hat{D}, \hat{E}, \hat{F}$ とすれば、これを(3)式に代入して $\alpha_0, \alpha_1, \alpha_2, \beta_0, \beta_1, \beta_2$ の推定値を求めることができる。このような方法によって方程式のパラメータを推定するのが間接最小二乗法の手順である。

直接最小二乗法で行ったと同様に、一般化して考えてみよう。いま、 G 個の同時従属変数 ($X_{1t}, X_{2t}, \dots, X_{Gt}$) と K 個の先決変数 ($Z_{1t}, Z_{2t}, \dots, Z_{Kt}$) をもつ G 個の方程式からなる理論模型を考えてみよう。

この理論模型は次のように表わされる。

$$\beta_{11} X_{1t} + \beta_{12} X_{2t} + \dots + \beta_{1G} X_{Gt} + r_{11} Z_{1t} + r_{12} Z_{2t} + \dots + r_{1K} Z_{Kt} = U_{1t}$$

$$\beta_{21} X_{1t} + \beta_{22} X_{2t} + \dots + \beta_{2G} X_{Gt} + r_{21} Z_{1t} + r_{22} Z_{2t} + \dots + r_{2K} Z_{Kt} = U_{2t}$$

$\vdots$

$$\beta_{G1} X_{1t} + \beta_{G2} X_{2t} + \dots + \beta_{GG} X_{Gt} + r_{G1} Z_{1t} + r_{G2} Z_{2t} + \dots + r_{GK} Z_{Kt} = U_{Gt}$$

この理論模型において、

$$X_t = [X_{1t} \ X_{2t} \ \dots \ X_{Gt}] \quad \dots\dots\dots \text{同時従属変数のベクトル}$$

$$Z_t = [Z_{1t} \ Z_{2t} \ \dots \ Z_{Kt}] \quad \dots\dots\dots \text{先決変数のベクトル}$$

$X'_t \ Z'_t \ \dots\dots\dots$ それぞれの転置行列

$U_t \equiv [U_{1t}, U_{2t} \ \dots\dots U_{Gt}] \ \dots\dots$ 誤差項のベクトル

$B \equiv \begin{pmatrix} \beta_{11} & \beta_{12} & \dots\dots \beta_{1G} \\ \beta_{21} & \beta_{22} & \dots\dots \beta_{2G} \\ \vdots & & \\ \beta_{G1} & \beta_{G2} & \dots\dots \beta_{GG} \end{pmatrix} \ \dots\dots$ 同時従属変数のパラメータの行列

$\Gamma \equiv \begin{pmatrix} \gamma_{11} & \gamma_{12} & \dots\dots \gamma_{1k} \\ \gamma_{21} & \gamma_{22} & \dots\dots \gamma_{2k} \\ \vdots & & \\ \gamma_{G1} & \gamma_{G2} & \dots\dots \gamma_{Gk} \end{pmatrix} \ \dots\dots$ 先決変数のパラメータの行列

とすれば、理論模型はベクトル、行列によって次のように表わすことができる。

$$BX'_t + \Gamma Z'_t = U'_t$$

この理論模型の誘導型は、同時従属変数がすべて先決変数によって説明されるのであるから、

$$X_{1t} = \lambda_{11} Z_{1t} + \lambda_{12} Z_{2t} + \dots\dots + \lambda_{1k} Z_{kt} + v_{1t}$$

$$X_{2t} = \lambda_{21} Z_{1t} + \lambda_{22} Z_{2t} + \dots\dots + \lambda_{2k} Z_{kt} + v_{2t}$$

$\vdots$

$$X_{Gt} = \lambda_{G1} Z_{1t} + \lambda_{G2} Z_{2t} + \dots\dots + \lambda_{Gk} Z_{kt} + v_{Gt}$$

すなわち、

$\pi \equiv \begin{pmatrix} \lambda_{11} & \lambda_{12} & \dots\dots \lambda_{1k} \\ \lambda_{21} & \lambda_{22} & \dots\dots \lambda_{2k} \\ \vdots & & \\ \lambda_{G1} & \lambda_{G2} & \dots\dots \lambda_{Gk} \end{pmatrix} \ \dots\dots$ 誘導型の先決変数のパラメータの行列

$V_t \equiv [v_{1t}, v_{2t} \ \dots\dots v_{Gt}] \ \dots\dots$ 誘導型の誤差のベクトル

とすれば、誘導型は

$$X'_t = \pi Z'_t + V'_t$$

である。理論模型は

$$BX'_t + \Gamma Z'_t = U'_t$$

であるから

$$X'_t = -B^{-1} \Gamma Z'_t + B^{-1} U'_t$$

すなわち、

$$\pi = -B^{-1} \Gamma$$

$$V'_t = B^{-1} U'_t$$

である。

間接最小二乗法による基本形のパラメータの推定においては、誘導型

$$X'_t = \pi Z'_t + V'_t$$

によって、パラメータ π の推定値 $\hat{\pi}$ を直接最小二乗法によって推定し、

$$\pi = -B^{-1} \Gamma$$

を解いて、 B および Γ の推定値 $\hat{B}$ 、 $\hat{\Gamma}$ を求める。

3) 認定の問題

間接最小二乗法の考え方と解き方についての手順は上述のとおりであるが、誘導型のパラメータを推定したからといって基本形のパラメータが必ず導かれるということはない。

いま、理論模型の基本形に同時従属変数が一個増加すると方程式の数を一本増加し、それぞれの方程式に増加した同時従属変数が一個ずつ増えるから、推定されるべきパラメータは

$1 \times G + 1 (G + 1 + K)$ 個増加する。(図7-4参照)

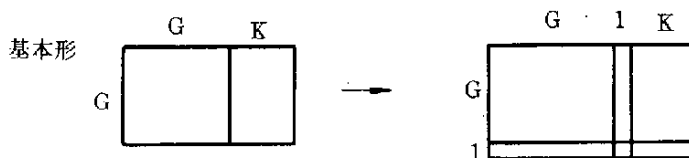


図7-4

一般に理論模型の基本形において同時従属変数が ℓ 個増加すれば、初めの完全模型の推定されるべきパラメータは変数の数 $(G + K)$ に方程式の数 G を掛けた $G(G + K)$ から、 ℓ 個の新変数の数 $(G + \ell + K)$ に方程式の数 $(G + \ell)$ を掛けた $(G + \ell + K)(G + \ell)$ になるから、
 $(G + \ell + K)(G + \ell) - (G + K)G = \ell(2G + K + \ell)$ 個増加する。この場合の誘導型のパラメータは図7-5のように、 GK 個から $(G + \ell)K$ すなわち、 $(G + \ell)K - GK = \ell K$ 個増加する。

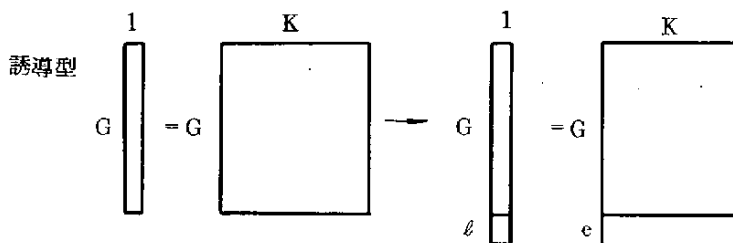


図7-5

このように完全模型で同時従属変数を ℓ 個増加させることは $\ell(2G + K + \ell)$ の推定されるべき未知のパラメータを増加させることであり、誘導型では ℓK 個のパラメータの増加となり、

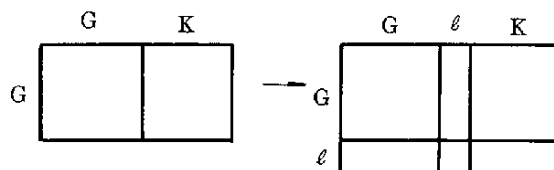


図7-6

差し引き

$$\ell(2G+K+\ell) - \ell K = 2G\ell + \ell^2$$

だけ完全理論模型の未知パラメータを決定することができない。(図7-6)

次に、同時従属変数をそのままにして、
先決変数を m 個増加させた場合はどうであ
ろうか。

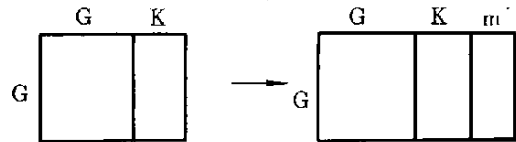


図7-7

完全理論模型の未知パラメータの個数は
 $G(G+K)$ から $G(G+K+m)$ すなわ
ち、 $G(G+K+m) - G(G+K) = Gm$
個未知パラメータが増加する。(図7-7)

他方、誘導型においては GK 個から、 $G(K+m)$ 個すなわち、

$$G(K+m) - GK = Gm$$

Gm 個推定さるべきパラメータが増加する。したがって先決変数の個数が m 個増加すれば、完全理
論模型も誘導型も等しく、未知パラメータが Gm 個増加している。このように、変数の増減があれ
ば、完全理論模型の未知パラメータと誘導型のパラメータの数には一致しない場合があるから、た
とえ誘導型の方程式のパラメータを直接最小二乗法で推定したとしても、完全理論模型のパラメー
タを決定するわけにはゆかない。

そこで、いま完全理論模型と誘導型の方程式が同時に成り立つとすれば、基本形の第 i 番目の式

$$\beta_{i1} X_{1t} + \beta_{i2} X_{2t} + \cdots + \beta_{iG} X_{Gt} + r_{i1} Z_{1t} + r_{i2} Z_{2t} + \cdots + r_{ik} Z_{kt} = U_{it}$$

に誘導型を代入して、

$$\begin{aligned} & \beta_{i1} (\lambda_{11} Z_{1t} + \lambda_{12} Z_{2t} + \cdots + \lambda_{1k} Z_{kt}) + \beta_{i2} (\lambda_{21} Z_{1t} + \lambda_{22} Z_{2t} + \cdots + \lambda_{2k} Z_{kt}) \\ & + \cdots + \beta_{iG} (\lambda_{G1} Z_{1t} + \lambda_{G2} Z_{2t} + \cdots + \lambda_{Gk} Z_{kt}) \\ & + r_{i1} Z_{1t} + r_{i2} Z_{2t} + \cdots + r_{ik} Z_{kt} \\ & = U_{it} - \beta_{i1} V_{1t} - \beta_{i2} V_{2t} - \cdots - \beta_{iG} V_{Gt} \end{aligned}$$

両辺を $t = 1, 2, \dots, T$ で合計すれば、右辺は誤差の合計で0であるから、いま、左辺だけを
考えて、 Z_{jt} についてまとめると

$$\begin{aligned} & (\beta_{i1} \lambda_{11} + \beta_{i2} \lambda_{21} + \cdots + \beta_{iG} \lambda_{G1} + r_{i1}) Z_{1t} + \cdots \\ & + (\beta_{i1} \lambda_{1k} + \beta_{i2} \lambda_{2k} + \cdots + \beta_{iG} \lambda_{Gk} + r_{ik}) Z_{kt} = 0 \end{aligned}$$

先決変数 Z_{jt} は0でないから

$$\begin{aligned} & \beta_{i1} \lambda_{11} + \beta_{i2} \lambda_{21} + \cdots + \beta_{iG} \lambda_{G1} + r_{i1} = 0 \\ & \vdots \\ & \beta_{i1} \lambda_{1k} + \beta_{i2} \lambda_{2k} + \cdots + \beta_{iG} \lambda_{Gk} + r_{ik} = 0 \end{aligned}$$

である。

ここに K 個の方程式があるが、未知パラメータは $\beta_{i1}, \beta_{i2}, \dots, \beta_{iG}$ の G 個あるいは β_{i1}
の基準化によって1とし、全体を割れば、未知パラメータは $\beta_{i2}^*, \beta_{i3}^*, \dots, \beta_{iG}^*$ の $G-1$ 個と

$r_{i1}^*, r_{i2}^*, \dots, \beta_{ik}^*$ の K 個あるから、 $G-1$ 個は決定することができない。そこで、完全理論模型のある特定の方程式の変数を $G-1$ 個減少させれば、 K 個の方程式によって K 個の未知パラメータを決定させることができる。

完全理論模型の第 i 番目の方程式において採用されている同時従属変数の数を G^d 個とし、現われてこない同時従属変数の数を $G^{d,d}$ とすれば、 $X_{1t}, X_{2t}, \dots, X_{G^d t}$ が、この方程式にあらわれてくる。また、 $X_{G^{d,d}+1t}, X_{G^{d,d}+2t}, \dots, X_{Gt}$ は現われてこない。さらに、この方程式に現われる先決変数の数を K^* 、現われてこない先決変数の数を K^{**} とすれば、 $Z_{1t}, Z_{2t}, \dots, Z_{K^*t}$ はこの方程式に現われ、 $Z_{K^*+1t}, Z_{K^*+2t}, \dots, Z_{Kt}$ は現われてこない。

従って、このようにすれば、基本型に誘導型を代入した式は

$$\left. \begin{aligned} \beta_{i1} \lambda_{11} + \beta_{i2} \lambda_{21} + \dots + \beta_{iG^d} \lambda_{G^d 1} + r_{i1} &= 0 \\ \vdots \\ \beta_{i1} \lambda_{1k^*} + \beta_{i2} \lambda_{2k^*} + \dots + \beta_{iG^d} \lambda_{G^d k^*} + r_{ik^*} &= 0 \end{aligned} \right\} K^* \text{個} \quad \dots\dots\dots (1)$$

と

$$\left. \begin{aligned} \beta_{i1} \lambda_{1, k^*+1} + \beta_{i2} \lambda_{2, k^*+1} + \dots + \beta_{iG^d} \lambda_{G^d, k^*+1} + 0 &= 0 \\ \vdots \\ \beta_{i1} \lambda_{1k} + \beta_{i2} \lambda_{2k} + \dots + \beta_{iG^d} \lambda_{G^d k} + 0 &= 0 \end{aligned} \right\} K^{**} \text{個} \quad \dots\dots\dots (2)$$

ここで、(1)式においては前提条件からして問題はないが、(2)式の条件を考慮しなければならない。

いま(2)式と $\beta_{i1} = 1$ で基準化すれば、

$$\left. \begin{aligned} \lambda_{1, k^*+1} + \beta_{i2}^* \lambda_{2, k^*+1} + \dots + \beta_{iG^d}^* \lambda_{G^d, k^*+1} + 0 &= 0 \\ \vdots \\ \lambda_{1, k} + \beta_{i2}^* \lambda_{2, k} + \dots + \beta_{iG^d}^* \lambda_{G^d k} + 0 &= 0 \end{aligned} \right\}$$

において、方程式の個数は $K^{**} - K - K^*$ 個、未知数は $G^d - 1$ 個ゆえに、方程式が未知数より少ないければ、この未知数を決定することができないから、

$$K^{**} \geq G^d - 1$$

でなければならない。この条件を構造認定の Rank Condition (位階条件) という。

7.2 Model Building の定義形式

DEFINE: EQUATION

$EQ(k) = (v_i, s_j)$

$EQ(k) = (X1, X2, \dots, Z5, SECTOR2)$

ESTIMATE: DLS

$EQ(k): TEST((n'))$)

ESTIMATE: IDLS

$EQ(1): TEST((n'))$)

$EQ(2): TEST((n'))$)

$\vdots$

$EQ(i): TEST((n'))$)

8 推定と評価 (Estimation and Evaluation)

JUMPSでは変数, 方程式, 方程式体系などについてパラメータを計算すると, 検定に必要な特性値 (平均値, 分散, 共分散, 非対称度, 鋭度など) の統計量を自動的に計算する。

統計量の検定には $TEST(m')(n', v_i, s_j, EQ_k)$ を用いる。 $TEST(m')(n', v_i, s_j, EQ_k)$ において (m') は検定方法の指定である。検定方法には次の指定方法がある。

CS		χ^2 - 検定
F		F - 検定
Z		Z - 検定
T		T - 検定
NM		Neumann 比検定
DW		Durbin Watson 比検定
ER		誤差率
SG		符号条件の検定

また n' は検定基準を与える。

ある変数群, 方程式群, 方程式体系群がある場合, それらの群の中から指定個だけの変数, 方程式, 方程式体系を選ぶために, $SELECT$ および $OPTIMAL$ を指定することができる。

$SELECT(m', (n'), v_i, s_j, EQ_k, GR_\ell)$

$OPTIMAL(m', (n'), v_i, s_j, EQ_k, GR_\ell)$

$SELECT$ の m' によって選択方法 $CS(n')$, $Z(n')$, $T(n')$, $VAR(n')$, $COR(n')$, $NM(n')$, $ER(n')$, $DW(n)$ などを指定する。たとえば $ER(n')$ の n' によって誤差率の小さな変数, 方程式, 方程式体系の個数 m と誤差率の許容水準 n が指定される。

$n' = (m, n)$

m は個数, n は誤差率, DW , χ^2 , Z , T , および COR の許容水準である。 m と n の条件はゆるい方が採用される。

$SELECT$ に続く $OPTIMAL$ の m' によって, $SELECT$ のすべての条件を満足する変数, 方程式, 方程式体系が選出され, その中で最も重要な検定条件が指定される。 $OPTIMAL$ が単独で用いられる場合も m' によって指定された検定条件に最も適合する変数, 方程式, 方程式体系が選ばれる。

応答形式

Model Building Stage の Instruction Step は確認の Substep, 指定の Substep とがあり, 確認の Substep では初めに前段階までの分析結果を表わす変数分析表と相関表が表示される。

変数分析表は変数名、変数、優先度および分散、 $D-W$ 係数、誤差率、変動係数が次のように表示される。

	VARIABLE	PRIORITY	VAR	$D-W$	ER	VR
NATIONAL INCOME	X1	2				
	X2	3				

Model Building を行なう人は、この表を見ながら分析が行なわれていない変数、あるいは相関関係についてそれぞれの stage にもどって分析することができる。自動リターン (AUTO-RETURN) の指定を行なえば、指摘した分析の stage に computer は自動的にもどり、その分析が OK になったときに再び Model Building Stage へ帰ってくる。特別な分析を必要とするときは JUMP 命令によって任意の stage あるいは step へ飛ぶことができる。

変数を選び終ると相関分析の結果表が提出される。分析不足のデータについては AUTO-RETURN あるいは JUMP によって相関分析 stage で再び相関係数についての分析が行なわれ、OK によって再び Model Building の stage にもどる。

Model Building Stage の Instruction Step の確認の Substep が OK ならば、次に指定の Substep へと作業は移る。

この Substep では Sector をふくめた変数表が表示され、定義表としては方程式、被説明変数、推定方法および対数、指数がある。さらに条件表、選択子表、応答状況指定がある。

Model Building を行なうために次のように方程式の指定を行なう。

$$\bullet EQ(k) = (X1, X2, \dots, Y1, \dots, Z1, \dots, S1)$$

次に各方程式の被説明変数が何であることを示すために、

• EXPLAINED

にライトペンで指定を行ない、さらに $EQ(k)$ を指定し、次いで変数表の変数に指定を行なう。

推定方法はこの段階では直接最小自乗法と間接最小自乗法の二方法のうちいずれかあるいは両者を選ぶ。

• ESTIMATE • DLS

• IDLS

応答状況指定においては • CONTINUE, • NEXT, • CHANGE, • OK,

• EXECUTE のほかに • DELETE が加えられている。変数表あるいは SECTOR, あるいは方程式において変数, SECTOR, 方程式の変更および削除が直接的に行なわれる。

VARIABLE • X1, • X2, • X3, ……

• Y1, • Y2, • Y3, ……

• Z1, • Z2, • Z3, ……

• D1, • D2, • D3, ……

SECTOR • S1, • S2, • S3, ……

• EQ • LOG • EXP

• EXPLAINED	
• ESTIMATE	• DLS • IDLS
DSI	• CONTINUE
	• NEXT
	• CHANGE
	• DELETE
	• OK
	• EXECUTE
OPTION SIGN	• +, • -
OPERATOR	• +, • -, • *, • 1
NUMBER	• 0, • 1, • 2, • 3, • 4, • 5, • 6, • 7, • 8, • 9,
CONDITION	• ALL • OPTIMAL
	• SELECT
	• TEST
TEST	• CS • F • Z
	• T • NM • D-W
	• ER • SG

ALL, OPTIMALおよびSELECTは前のStageと同様であるが、CONDITIONでTESTを指定すると、TESTの方法が表示されて、ライトペンによってそのうちのいくつかを選ばなければならない。TESTの優先度は指定した順であって、順位についてはCHANGEによって変更ができる。

表示形式は次のようである。

• $EQ1 = (X1, X2, \dots, S1, S2, \dots)$
 • EXPLAINED • $EQ1(X1),$ • $EQ2(X2), \dots$

IDLSを指定した場合は前もってDLSを指定しなければ、Computerは自動的にDLSのSubstepを処理し、IDLSの確認Substepに移る。したがってIDLSに入るときに自動的処理を行なわせないようにするにはDLSも指定しておかなければならない。このようにすると使用者の考えによる独特の推定を行なうことができる。

DLSがすでに実行されているか、定義形式などによってDLSのSubstepの処理が必要でない場合には直接IDLSに入ることができる。

IDLSの内部的な処理は原則として自動的に行なわれるが、諸条件の変更を要する場合には、たとえばpriorityの変更、内生・外生の再定義、誤差率など推定の水準の指定などは処理の途中で応答することができる。

応答に要する条件は方程式体系に関する条件であって次のように表現される。

EQUATION NO. VARIABLE A-P R R E R DW

EQ 1 X1 + X2

EQ 2 X1 + X3

⋮

EQ n X2 + X3

OPTION SIGN +, -

OPERATOR +, -, *, /

NUMBER 0, 1, 2, 3, 4, 5, 6, 7, 8, 9

単一方程式のパラメータの推定と方程式の自動選択

連立方程式モデルを想定して単一方程式のパラメータを推定する場合、利用者は各変数をあらかじめ

内生変数 (X1, X2, ..., Xi, ..., Xg')

外生変数 (Z1, Z2, ..., Zj, ..., Zk')

および内生あるいは外生に決められない変数

(Y1, Y2, ..., Yl, ..., Yl')

の3つの変数群に分けて指定しなければならない。

いくつかの類似変数を含む単一方程式のパラメータの推定の場合、たとえばX1 = 国民所得、X2 = 可処分所得、X3 = 個人業種所得などはその動向がきわめて類似しているので、推定の計算途中でmulticollinearityが発生する可能性がある。このような類似変数を2つ以上ふくむ方程式をとくに作成する場合を除いて

((X1, X2, X3), X4, X5, ...)

のように指定すると括弧内(X1, X2, X3)の中から以下の手順を経てもっともよい変数が1つ選ばれる。

1.0 コンピュータは内生変数のpriorityの高い特定変数を被説明変数とし、そのほかの内生変数を1個づつ選び、2変数模型として各方程式のパラメータを推定する。

・類似変数のない場合

$X1 = \alpha_1 + \alpha_2 X2$

$X1 = \alpha_1 + \alpha_3 X3$

$X1 = \alpha_1 + \alpha_4 X4$

⋮

$X1 = \alpha_1 + \alpha_i X_i$

⋮

$X1 = \alpha_1 + \alpha_g X_g'$

・類似変数のある場合 (X 1, X 2, X 3)

$$\begin{array}{rcl}
 X_1 & = & \alpha_1 + \alpha_4 X_4 \\
 & \vdots & \\
 X_1 & = & \alpha_1 + \alpha_i X_i \\
 & \vdots & \\
 X_1 & = & \alpha_1 + \alpha_{G'} X_{G'} \\
 X_2 & = & \alpha_1 + \alpha_4 X_4 \\
 & \vdots & \\
 X_2 & = & \alpha_1 + \alpha_i X_i \\
 & \vdots & \\
 X_2 & = & \alpha_1 + \alpha_{G'} X_{G'} \\
 X_3 & = & \alpha_1 + \alpha_4 X_4 \\
 & \vdots & \\
 X_3 & = & \alpha_1 + \alpha_i X_i \\
 & \vdots & \\
 X_3 & = & \alpha_1 + \alpha_{G'} X_{G'}
 \end{array}$$

- 1.1 これらの方程式群から符号条件の合致しない方程式は除かれる。
- 1.2 次に方程式群の最小誤差率をもつ方程式が選ばれる。
- 2.0 先に選ばれた変数 (X 1) 以外の変数たとえば X 2 が被説明変数に選ばれ、X 1 を除くその他の変数を説明変数とする 2 変数模型が次のように作られる。
- 2.1 この方程式群から 1.1, 1.2 の手順を経て単一方程式が 1 つ選ばれる。
- 2.2 1.0 - 1.2 まで選ばれた方程式と 2.1 で選ばれた方程式について、誤差分散を比較して良い方程式あるいは良い方から指定された数だけの基本方程式が数本残される。
- 3.0 次に、出てきた変数たとえば X 1, X 2 を除いて 2.0, 2.1, 2.2 の手順を経て誤差の分散の最も小さな方程式が選ばれる。
- 4.0 この手順が最終の G' 変数まで繰り返される。このようにして第 1 の基本方程式

$$X_s = \alpha_s + \alpha_t X_t$$

あるいは指定された数本の基本方程式が選ばれる。

- 5.0 第 1 の基本方程式 $X_s = \alpha_s + \alpha_t X_t$ に対して外生変数 Z 1, Z 2, ..., Z_j, ..., Z_k の中から priority の高いあるいは priority が等しいときには X_s との相関が最も (相関係数) 高い Z_j が選ばれ

$$X_s = \alpha_s + \alpha_t X_t + \beta_j Z_j$$

のパラメータ $\alpha_s, \alpha_t, \beta_j$ が推定され、パラメータの各種条件が reserve される。

- 5.1 次に priority の高い外生変数あるいは同じ priority の場合は相関係数の高い外生

変数が選ばれ、5.1は $j = 1, 2, \dots, k'$ まで繰返され結果は reserve される。

ここで $j = 1, 2, \dots, k'$ までの方程式群に対して符号条件が検定され、不合格の方程式は取り除かれる。

5.2 次に誤差率の最も小さい外生変数 $Z_{\ell'}$ を持つ方程式

$$X_s = \alpha_s + \alpha_t X_t + \beta_{\ell'} Z_{\ell'}$$

を第2基本方程式として誤差率が小さくなるように影響を与えた変数が1つ加えられ、

$$X_s = \alpha_s + \alpha_t X_t + \beta_{\ell'} Z_{\ell'} + \beta_p Z_p$$

$$(p = 1, 2, \dots, k'', p \neq \ell')$$

$p = 1, 2, \dots, k''$ までの方程式群に対して5.1の過程が繰返され、符号条件が指定され不合格の方程式が除かれる。

第2基本方程式

$$X_s = \alpha_s + \alpha_t X_t + \beta_{\ell'} Z_{\ell'}$$

の分散と

$$X_s = \alpha_s + \alpha_t X_t + \beta_{\ell'} Z_{\ell'} + \beta_p Z_p$$

の分散が比較され、第2基本方程式の分散よりも小さくなれば、その方程式が現行の基本方程式として残される。分散が小さくなる場合はさらに5.2の過程が繰返される。分散が大きくなった場合は、第1基本方程式の誤差率と現行の基本方程式の誤差率とが比較される。

6.0 第1基本方程式がよければ、現行の基本方程式は捨てられ内生変数が加えられる。第1基本方程式に priority の高い変数あるいは priority が同一水準の場合、被説明変数に対して相関係数の高い変数が選ばれる。

$$X_s = \alpha_s + \alpha_t X_t + \alpha_u X_u \quad (u = 1, 2, \dots, G'')$$

誤差率が比較され、最も誤差率の低い u をもつ X_s が選ばれる。

$$X_s = \alpha_s + \alpha_t X_t \text{ の分散 } Var1$$

$$X_s = \alpha_s + \alpha_t X_t + \alpha_u X_u \text{ の分散 } Var2$$

ここで、 $Var1 > Var2$ のときは

$$X_s = \alpha_s + \alpha_t X_t + \alpha_u X_u + \alpha_v X_v \quad (v = 1, 2, \dots, G', v \neq t, v \neq u)$$

とし、分散が比較され、ある変数が加えられて分散が大きくなった場合は、その変数はとり除かれ現行基本方程式となる。

$Var1 < Var2$ のときは現行の Z_p をもつ基本方程式の Z_p ($p = 1, 2, \dots, k''$) が加えられ、5.2の過程が繰返される。

I パラメータが推定されて作られたいくつかの単一方程式は元来相互に関係はない。同時連立方程式体系を作るために間接最小二乗法による推定を行なう場合は、

J \ 条件	Sign	α_t	β_j
1			
2			
$\vdots$			
$\vdots$			
k'			

I-1 方程式群を構成する単一方程式は内生変数の数と、方程式の数とが一致しなければならない。

I-2 各方程式のパラメータを推定する場合に、方程式群のどの単一方程式も内生変数と外生変数として現われてこない変数の数 N が、方程式群の方程式の数 G と $G = N - 1$ ，すなわち Just Identifiable (適度認定) でなければならない。

II 方程式群の方程式の数(m)と内生変数の数(n)との一致

II-1 第1に方程式群の内生変数をそろえるために各基本方程式に含まれている内生変数が整理され確認される。全く重複する変数をもつ方程式は単一方程式としてみなされる。

第1方程式	X 1		X 3						
2	X 1	X 2							
3	X 1		X 3	X 4					
4		X 2		X 4					
5		X 2		X 4	X 5				
6	X 1	X 2		X 4					
7							X 6		
8		X 2	X 3	X 4					
$m = 8$	X 1	X 2	X 3	X 4	X 5	X 6		$n = 6$	

II-2 m と n が比較され

$m > n$ ならば II-3 へ

$m < n$ ならば II-4 へ

$m = n$ ならば II-5 へ

II-3 $m > n$ ならば方程式の数を $m - n$ 本減らさなければならない。この手順は次のように行なわれる。

II 3.1 各方程式の共通内生変数の個数が数えられ、同一内生変数の存在しない1個の単独内生変数をもつ方程式は残される。2個以上の共通内生変数の場合は II 3.4 へ。

上例では 第5式 X 2, X 4, X 5 および

第7式 X 6

が残される。

II 3.2 残された方程式の単独内生変数以外の内生変数と共通の内生変数をもつ方程式が選ばれる。

上例では次の4方程式が選ばれる。

				誤差率
第2式	X 1	X 2		0.12
第4式		X 2	X 4	0.38
第6式	X 1	X 2	X 4	0.08

第8式 X 2 X 3 X 4 0.0 6

II 3.3 選ばれた方程式の誤差率が比較され、最も高い誤差率をもつ方程式が一本はずされる。
誤差率は各方程式の内生変数が1つつ被説明変数になって内生変数の数だけ計算され、最も小さな誤差率をもつ内生変数が被説明変数となる。上例では第4式がはずされる。

II 3.4 priority係数

第1式	X 1		X 3				2
2	X 1	X 2					1.5
3	X 1		X 3	X 4			2.7
4							
5		X 2		X 4	X 5		
6	X 1	X 2		X 4			2.3
7						X 6	
8		X 2	X 3	X 4			3.0

各方程式の内生変数の priority の総和を変数の数で除した priority 係数が比較され最も高い係数を含む方程式が除かれる。上例では第8方程式の係数が3.0で最も高くこの方程式が除かれる。

priority 係数が高く等しい場合は、その方程式の誤差率が比較され誤差率の大きい方程式が除かれる。なお誤差率が等しい場合には誤差の分散が比較され、分散の大きい方程式が除かれる。さらに誤差の分散が等しい場合には $D-W$ 比が比較され2より遠い $D-W$ 比をもつ方程式が除かれる。

これらの選定基準によって結論が出ない場合は Warning を表示する。

II-2へとぶ。

II-4 $m < n$ ならば変数を $n - m$ 箇減らさなければならない。その手順は次のように行なわれる。

II 4.1 各方程式の内生変数の priority が比較され、最も低い priority の変数をもつ方程式が選ばれる。

第1式	X 1		X 3	X 4				
2		X 2	X 3		X 5			
3	X 1			X 4		X 6		
4		X 2			X 5	X 6		
5			X 3	X 4		X 6	X 7	
$m = 5$	X 1	X 2	X 3	X 4	X 5	X 6	X 7	$n = 7$

上例でX 5, X 6, X 7が最も低い priority をもつとすれば、第3, 第5および第5式が選ばれる。

II 4.2 選ばれた各方程式のうち1つしかでてこない変数は取り除かれる。例えば第5式のX 7は取り除かれる。次に誤差率が計算され、X 6を取り除いた方程式の誤差率の総和の平均値と

X 5 を取り除いた誤差率の総和の平均値とが比較され、大きな誤差率の平均値をもつ変数が取り除かれる。この例ではたとえば X 6 が除かれる。II-2 へとぶ。

第1式	X 1		X 3	X 4	
2		X 2	X 3		X 5
3	X 1			X 4	
4		X 2			X 5
5			X 3	X 4	

II-5 $m=n$ ならば各方程式の外生変数が整理された $G=m=n$ と各方程式に現われてこない内生変数および外生変数の総和 N が求められ、 $G-1$ と N とが比較される。

II 5.1 $N < G-1$ の場合は、II 5.2 へ

$N > G-1$ の場合は、II 5.3 へ

$N = G-1$ の場合は、II 5.4 へ

II 5.2 $N < G-1$ の場合は各方程式に共通な外生変数が数個存在すれば、 $N = G-1$ になるまで priority の低い変数がはずされる。priority が一致していれば、誤差率を計算して高い誤差率をもたらし外生変数がとりのぞかれる。なお $N < G-1$ ならば他の方程式の外生変数をふやすために II 5.5 へ。

II 5.3 $N > G-1$ の場合は他の方程式に現われている外生変数で当該方程式に現われてない外生変数であってすでにこの方程式を作成する場合に除かれている変数すなわち予定変数の中から priority の高い外生変数を選ばれ変数が附加される。ただし附加する変数の priority が同レベルにある場合には誤差率を計算して低い誤差率をもたらし変数が加えられる。予定変数のない場合は共通の変数のない方程式の外生変数が一つ取り除かれる。II 5.1 へとぶ。

II 5.4 $N = G-1$ ならば、すべての方程式は Just Identifiable である。

II 5.5 その方程式以外の方程式に予定変数 Y_k を一個加える。II 5.1 へとぶ。

Execution phase

ここで間接最小二乗法によってパラメータが推定される。

$$B X + \Gamma Z = u$$

$$B^{-1} B X = -B^{-1} \Gamma Z + B^{-1} u$$

$$X = -B^{-1} \Gamma Z + B^{-1} u$$

$$= \pi Z + V$$

π が直接最小二乗法によって推定され

$$-B^{-1} \Gamma = \hat{\pi}$$

が解かれる。

ただしこの stage において、SWLS, 2 SLS, 3 SLS, INSTV, LIML および FIML による推定方法は応答形式にしたがう。

間接最小二乗法によるパラメータの推定

間接最小二乗法によってパラメータを推定するには、誘導型の係数行列 π を直接最小二乗法によって推定し

$$B^{-1} F = \hat{\pi}$$

を解かなければならない。

上式を解くためには $B^{-1} F$ の係数行列において B の対角行列 $\beta_{11}, \beta_{22}, \dots, \beta_{ii}, \dots, \beta_{cc}$ が基準化され、且つ存在しなければならない。このために B の正方行列条件と $BX + F'Z$ の認定条件を満足した B の対角行列が存在し、基準化できるようにするには、次のアルゴリズムにしたがわねばならない。

- i) B の係数行列 β_{ij} が 0 のベクトルをはずす。
- ii) Houthacker の Northwest Rule のように B の対角行列に 0 以外のベクトルをもってくるのであるが、 B 行列は正方行列であって各変数には少なくとも必ず一個以上の 0 以外の係数が存在するので、変数の置き換えは行なわない。
- iii) 変数 X_1 のパラメータをもつ方程式が選ばれ各方程式において X_1 のパラメータを含む変数のパラメータの数が比較され、最もパラメータの個数の少ない方程式のパラメータが第 1 列に並べられる。
- iv) 第 1 列に並べられた方程式以外の変数 X_2 のパラメータをもつ方程式が選ばれ、各方程式において X_2 のパラメータを含む変数のパラメータの数が比較され、最もパラメータの個数の少ない方程式のパラメータが第 2 列に並べられる。
- v) 4 の過程が変数 X_i について繰り返される。
- vi) ある方程式にのみ 2 個以上の個有の変数が存在する場合は、対角要素が 0 になるから、対角要素が 0 になった場合は、 C 箇の変数の中で同一方程式に 2 箇以上の個有の変数があるかどうかを調べる。もし 2 箇以上の個有の変数があれば、そのほかの方程式を選んでその個有の変数を加える過程に入る。

もし 2 個以上の個有の変数が存在しなければ、対角要素が 0 の行の中に現われてくる方程式が選ばれ、さらに対角要素が 0 である方程式の変数と同一の変数をもつかどうかを調べ、同一の変数をもつ方程式と対角要素が 0 である方程式と置き換えられる。

もし同一変数をもつ方程式がなければ、7 の過程に入る。

- vii) 対角要素が 0 である行の中から、そのほ

かの方程式が選ばれ Northeast のコーナーから $X_C, X_{C-1}, \dots, X_1$ の順序で対角行列にパラメータが並べられ 6 の過程を繰り返す。

- viii) 方程式の自動選択の過程に入る。

変数 式番	X 1	X 2	X 3	X 4	X 5	X 6
1	○			○	○	
2		○			○	
3						
4		○		○		
5	○				○	
6		○	●			●

変数 式番	X 1	X 2	X 3	X 4	X 5	X 6
1	○			○		
2		○			○	
3			○		○	
4				●		●
5		○	○		○	
6		○		●		

逐次最小二乗法

ある理論模型 $BX_t' + \Gamma Z_t' = U_t'$ の方程式の順序を入れ換えることによって、同時従属変数のパラメータの行列 B が

$$B = \begin{bmatrix} \beta_{11} \\ \beta_{21} & \beta_{22} \\ \vdots \\ \beta_{G1} & \beta_{G2} & \beta_{G3} & \cdots & \beta_{GG} \end{bmatrix}$$

のように三角形に並ぶときに逐次最小二乗法を用いることができる。

いま,

$$X_t = [X_{1t} \ X_{2t} \ \cdots \ X_{Gt}] \quad : \text{同時従属変数}$$

$$\Gamma = \begin{bmatrix} \gamma_{11} & \gamma_{12} & \cdots & \gamma_{1k} \\ \gamma_{21} & \gamma_{22} & \cdots & \gamma_{2k} \\ \vdots \\ \gamma_{G1} & \gamma_{G2} & \cdots & \gamma_{Gk} \end{bmatrix} \quad : \text{先決変数のパラメータ}$$

$$Z_t = [Z_{1t} \ Z_{2t} \ \cdots \ Z_{kt}] \quad : \text{先決変数}$$

$$U_t = [U_{1t} \ U_{2t} \ \cdots \ U_{Gt}]$$

とすれば,

$$\begin{bmatrix} \beta_{11} \\ \beta_{21} & \beta_{22} \\ \vdots \\ \beta_{G1} & \beta_{G2} & \beta_{G3} & \cdots & \beta_{GG} \end{bmatrix} \begin{bmatrix} X_{1t} \\ X_{2t} \\ \vdots \\ X_{Gt} \end{bmatrix} + \begin{bmatrix} \gamma_{11} & \cdots & \gamma_{1k} \\ \vdots \\ \gamma_{G1} & \cdots & \gamma_{Gk} \end{bmatrix} \begin{bmatrix} Z_{1t} \\ Z_{2t} \\ \vdots \\ Z_{kt} \end{bmatrix} = \begin{bmatrix} U_{1t} \\ U_{2t} \\ \vdots \\ U_{Gt} \end{bmatrix}$$

は

$$BX_t' + \Gamma Z_t' = U_t'$$

である。

この三角模型のパラメータを推定するには、同時従属変数が 1 個でそのほかの変数が先決変数である。第一番の式

$$\beta_{11} X_{1t} + r_{11} Z_{1t} + \dots + r_{1k} Z_{kt} = U_{1t}$$

から基準化して

$$X_{1t} = \frac{-r_{11}}{\beta_{11}} Z_{1t} + \frac{r_{12}}{\beta_{11}} Z_{2t} + \dots + \frac{-r_{1k}}{\beta_{11}} Z_{kt} + U_{1t}$$

この式の先決変数 Z_{1t} のパラメータを直接最小二乗法によって推定しさらに、この推定値をもとに Z_{1t} の観測値を与えると X_{1t} が得られるから、この X_{1t} の推定値 $\hat{X}_{1t}$ を第2の式に代入して、さらに、

$$\beta_{22} X_{2t} = -(\beta_{21} \hat{X}_{1t} + r_{21} Z_{1t} + \dots + r_{2k} Z_{kt}) + U_{2t}$$

このようにして、同時従属変数の少ない方程式のパラメータから推定し、その結果を次の方程式に逐次代入し、最小二乗法によって次の方程式のパラメータを推定する方法を逐次最小二乗法という。この推定値は一致性をもつ。

応答形式

DLSが終了した段階で $EQ(k)$, $k = 1, 2, \dots, n$ が表示される。

この表示された方程式群をみて、三角模型に改善できるかどうかを人間が判断し

$$\begin{aligned} \cdot EQ1 &= (\cdot X1 && \cdot Z1 & \cdot Z2) \\ \cdot EQ2 &= (\cdot X1 & \cdot X2 && \cdot Z1 & \cdot Z2) \\ \cdot EQ3 &= (\cdot X1 & \cdot X2 & \cdot X3 && \cdot Z1 & \cdot Z2) \\ \cdot EQ4 &= (\cdot X1 & \cdot X2 & \cdot X3 & \cdot X4 && \cdot Z1 & \cdot Z2) \\ &\vdots \\ \cdot EQn &= (\cdot X1 & \cdot X2 & \cdot X3 & \cdot X4 & \dots & \cdot Xn & \cdot Z1 & \cdot Z2) \end{aligned}$$

に改善する可能性があれば、STEPWISE Substep を指定することができる。不揃いの場合は各方程式についてDLSの段階にAUTO-RETURNして対角要素に被説明変数がかかるようにパラメータを推定しなおす。

OKならば次にEXECUTIONの指定を行なうとComputer は自動的にStepwiseのアルゴリズムにしたがって各方程式のパラメータを推定する。

二段階最小二乗法(Two Stage Least Squares)

認別問題(Identification Problem)で述べたように、過剰認定(Over identifiable)な理論模型には間接最小二乗法を適用することができない。H. Theil() は、このような模型のパラメータの推定に対して、二段階最小二乗法を提唱している。

いま、完全理論模型を

$$\begin{aligned} \beta_{11} X_{1t} + \dots + \beta_{1G} X_{Gt} + r_{11} Z_{1t} + \dots + r_{1k} Z_{kt} &= U_{1t} \\ \beta_{21} X_{1t} + \dots + \beta_{2G} X_{Gt} + r_{21} Z_{1t} + \dots + r_{2k} Z_{kt} &= U_{2t} \\ &\vdots \\ \beta_{G1} X_{1t} + \dots + \beta_{GG} X_{Gt} + r_{G1} Z_{1t} + \dots + r_{Gk} Z_{kt} &= U_{Gt} \end{aligned}$$

とし、

$$B \equiv \begin{pmatrix} \beta_{11} & \beta_{12} & \cdots & \beta_{1G} \\ \beta_{21} & \beta_{22} & \cdots & \beta_{2G} \\ \vdots & \vdots & \ddots & \vdots \\ \beta_{G1} & \beta_{G2} & \cdots & \beta_{GG} \end{pmatrix}$$

$$X_t = [X_{1t}, X_{2t}, \cdots, X_{Gt}]$$

$$\Gamma \equiv \begin{pmatrix} \gamma_{11} & \gamma_{12} & \cdots & \gamma_{1k} \\ \gamma_{21} & \gamma_{22} & \cdots & \gamma_{2k} \\ \vdots & \vdots & \ddots & \vdots \\ \gamma_{G1} & \gamma_{G2} & \cdots & \gamma_{Gk} \end{pmatrix}$$

$$Z_t = [Z_{1t}, Z_{2t}, \cdots, Z_{kt}]$$

$$U_t = [U_{1t}, U_{2t}, \cdots, U_{kt}]$$

とすれば、完全理論模型は

$$\beta X_t' + \Gamma Z_t' = U_t'$$

である。

この模型の誘導型は

$$X_t' = -B^{-1} \Gamma Z_t' + B^{-1} U_t'$$

ここで、

$$\pi = -B^{-1} \Gamma$$

$$V_t' = B^{-1} U_t'$$

とすれば

$$X_t' = \pi Z_t' + V_t'$$

である。

いま、完全理論模型 $BX_t' + \Gamma Z_t' = U_t'$ において、第1番目の方程式の中で係数0の項を除いて

$$\beta_d^{(1)} \equiv [\beta_{12} \cdots \beta_{1G}]$$

$$X_d^{(1)} \equiv [X_{2t} \cdots X_{Gt}]$$

$$r_* \equiv [\gamma_{11} \cdots \gamma_{1k}]$$

$$Z_{*t} \equiv [Z_{1t} \cdots Z_{kt}]$$

とすると、

$$\begin{aligned} X_{1t} &= \beta_{12} X_{2t} + \beta_{13} X_{3t} + \cdots + \beta_{1G} X_{Gt} + \gamma_{11} Z_{1t} + \cdots + \gamma_{1k} Z_{kt} + U_{1t} \\ &= \beta_d^{(1)} X_d^{(1)'} + r_* Z_{*t}' + U_{1t} \end{aligned}$$

この式の中の同時従属変数

$$X_d^{(1)} \equiv [X_{2t}, X_{3t}, \cdots, X_{Gt}]$$

だけの誘導型をみちびくと

$$X_d^{(1)'} \equiv \pi_d^{(1)} Z_{*t}' + V_d^{(1)'}$$

ただし, $\pi_d^{(1)}$ は π の 1 行と (G_d) 行以下の行を除いた行列である。

いま, 次のように理論模型があつて,

$$\beta_{12} X_{2t} + \cdots + \beta_{1G} \Delta X_{G_d t} + r_{11} Z_{1t} + \cdots + r_{1k} Z_{kt} + U_{1t} = 0$$

$$\beta_{22} X_{2t} + \cdots + \beta_{2G} \Delta X_{G_d t} + r_{21} Z_{1t} + \cdots + r_{2k} Z_{kt} + U_{2t} = 0$$

⋮

$$\beta_{G_d 2} X_{2t} + \cdots + \beta_{G_d G} \Delta X_{G_d t} + r_{G_d 1} Z_{1t} + \cdots + r_{G_d k} Z_{kt} + U_{G_d t} = 0$$

第 1 番目の式の係数を除いた β_{ij} の行列を(1)で表わし,

$$B_d^{(1)} = \begin{pmatrix} \beta_{22} & \beta_{23} & \cdots & \beta_{2G} \Delta \\ \vdots & & & \\ \beta_{G_d 2} & \beta_{G_d 3} & \cdots & \beta_{G_d G} \Delta \end{pmatrix}$$

とし,

$$X_{d t}^{(1)} = [X_{2t}, X_{3t}, \cdots, X_{G_d t}]$$

$$r_d^{(1)} = \begin{pmatrix} r_{21} & r_{22} & \cdots & r_{2k} \\ \vdots & & & \\ r_{G_d 1} & r_{G_d 2} & \cdots & r_{G_d k} \end{pmatrix}$$

$$Z_t = [Z_{1t}, Z_{2t}, \cdots, Z_{kt}]$$

$$U_d^{(1)} = [U_{2t}, \cdots, U_{G_d t}]$$

とすれば,

$$B_d^{(1)} X_{d t}^{(1)'} + r_d^{(1)} Z_t' = -U_d^{(1)}'$$

この式の誘導型は

$$X_{d t}^{(1)'} = -\{B_d^{(1)}\}^{-1} r_d^{(1)} Z_t' - \{B_d^{(1)}\}^{-1} U_d^{(1)}'$$

ここで,

$$\pi_d^{(1)} = -B_d^{(1)-1} r_d^{(1)}$$

$$V_d^{(1)'} = -B_d^{(1)-1} U_d^{(1)'}$$

とすれば,

$$X_{d t}^{(1)'} = \pi_d^{(1)} Z_t' + V_d^{(1)'} \cdots \cdots \cdots (1)$$

$\pi_d^{(1)}$ の最小二乗推定値を $\hat{\pi}_d^{(1)}$ とすれば,

$$\hat{\pi}_d^{(1)} = M_{x_d z}^{(1)} M_{z z}^{-1} \cdots \cdots \cdots (2)$$

ここで,

$$M_{x_d z}^{(1)} = \begin{pmatrix} m_{x_d z 1} & \cdots & m_{x_d z k} \\ \vdots & & \\ m_{x_d z 1} & \cdots & m_{x_d z k} \end{pmatrix}$$

$$M_{zz} = \begin{pmatrix} m_{z1z1} & \cdots & m_{z1zk} \\ \vdots & & \\ m_{zkz1} & \cdots & m_{zkzk} \end{pmatrix}$$

(1)から

$$X_{dt}(1) - V_{dt}(1)' = \pi_{dt}' Z_t' \quad \dots\dots\dots (3)$$

であるから、(2)式から求められた推定値を(3)に代入すれば

$$X_{dt}(1) - V_{dt}(1)' = M_{xdt} M_{zz}^{-1} Z_t'$$

ここで、第1番目の方程式を推定するためには、第1番目の方程式は

$$X_{dt} = \beta_{dt}^{(1)} X_{dt}(1)' + r_* Z_*' t + U_{dt}$$

であるから、第(3)式の結果を利用して、

$$\begin{aligned} X_{dt} &= \beta_{dt}^{(1)} X_{dt}(1)' - \beta_{dt}^{(1)} V_{dt}(1)' + \beta_{dt}^{(1)} V_{dt}(1)' + r_* Z_*' t + U_{dt} \\ &= \beta_{dt}^{(1)} (X_{dt}(1)' - V_{dt}(1)') + r_* Z_*' t + (U_{dt} + \beta_{dt}^{(1)} V_{dt}(1)') \end{aligned}$$

U_{dt} は平均値0、その分散は有限一定である分布から得られた確率変数であるから、 U_{dt} の一次結合 $(U_{dt} + \beta_{dt}^{(1)} V_{dt}(1)')$ も平均値0、有限の分散をもつ確率変数であると考えることができる。すなわち二段階最小二乗法では推定さるべきパラメータをもつある方程式が過剰認定であれば、その方程式（上の例では第1番目の方程式）をのぞいたその他の方程式によって誘導型を導き、誘導型のパラメータを直接最小二乗法によって推定し、その結果

$$\begin{aligned} \bar{X}_{dt}(1)' &= X_{dt}(1)' - V_{dt}(1)' \\ &= M_{xdt} M_{zz}^{-1} Z_t' \end{aligned}$$

を過剰認定の式（第1番目の方程式）の同時従属変数として、あたえると同時に、 $Z_*' t$ を観察資料から与えることによって、それぞれのパラメータ $\beta_{dt}^{(1)}, r_*$ を直接最小二乗法によってさらに推定する。このように、最小二乗法によって求めた値をもとにして、さらに最小二乗推定を行うという意味で、二段階最小二乗法という。二段階最小二乗法は大標本特性値のうち、一致性をもつ。

応答形式

2SLSは基本的に応答形式である。とくにIDLSで、ある方程式がOver Identificationである場合、その方程式の変数の中で1つだけ2SLSによって解かれる変数を決め、そのほかの変数はIDLSによってJust Identifiableな模型として推定され、その結果を先に決めた変数の推定に用いる。したがって2SLSによって指定されるパラメータの変数を応答の中で決めなければならない。

$$\begin{aligned} EQ1 &= X1 \quad X2 \quad X3 \\ EQ2 &= X1 \quad \quad \quad X4 \\ EQ3 &= \quad \quad X2 \quad X3 \quad \quad X5 \\ &\vdots \\ EQn &= X1 \quad \quad \quad X5 \end{aligned}$$

- ・ 2 S L S
- ・ 3 S L S
- ・ U N E X P L A I N E D

第1に方程式を選び、その方程式の中から2 S L S で最終的に決定するパラメータの変数を指定する。

3 S L S は2 S L S に準ずる。

操作変数法 (Instrumental Variable)

今までの最小二乗法は直接、間接、二段階、逐次によらず推定さるべき方程式に関する資料をもとにパラメータの推定を行ってきたが、操作変数法では理論模型全体から必要な一部の情報と推定さるべき方程式とを関連せしめてパラメータを推定する。

二段階最小二乗法で行ったように、理論模型の第1番目の方程式を

$$\beta X_{it}' + \gamma Z_{it}' = U_{it}$$

とする。

ここで、

$$\beta = [\beta_{11}, \beta_{12}, \dots, \beta_{1G}, \underbrace{0, \dots, 0}_{G^d \text{ 個}}]$$

$$\gamma = [\gamma_{11}, \gamma_{12}, \dots, \gamma_{1k}, \underbrace{0, \dots, 0}_{k^{**} \text{ 個}}]$$

いま、この方程式において、正規化して $\beta_{11} = 1$ とすれば

$$X_{it} = -\beta_{12} X_{2t} - \beta_{13} X_{3t} \dots - \beta_{1G} X_{Gt} - \gamma_{11} Z_{1t} \dots - \gamma_{1k} Z_{kt} + U_{it}$$

である。ここで、この方程式には含まれていない先決変数で模型全体に含まれている先決変数のうちから、 $G^d - 1$ 個の変数を適当に選び

$$Z_{k^*+1t}, Z_{k^*+2t}, \dots, Z_{k^*+G^d-1t}$$

とし、この方程式が適度に識別可能であれば、いま選ばれた先決変数は次数条件によって選出可能である。すなわち、すべての内生変数 $G = G^d + G^d$ 、すべての外生変数 $K = K^* + K^{**}$ 、この方程式にあらわれてこない変数 $N = G^d + K^{**}$ であって、もし、この方程式が適度に識別可能ならば、 $N = G - 1$ の階数条件によって

$$K^{**} = K - K^* = G^d - 1$$

が保証されるから、前述の先決変数は選出可能である。

このように、この方程式が適度に識別可能であれば、先決変数は一意的に選ばれるが、この方程式が過剰に識別可能であれば、先決変数の選び方は $N^{C_{G^d-1}}$ 通りある。操作変数法というのは、この方程式に含まれている先決変数と、この方程式に現われてこないで、その他の方程式から選ばれた先決変数を操作変数としてパラメータを推定する方法である。

いま、推定さるべき方程式の両辺に先決変数 Z_i ($i = 1, 2, \dots, k^*, k^*+1, \dots, k^*+G^d-1$) を乗じ、 i について合計すると、

$$\sum_{t=1}^T \begin{pmatrix} Z_{1t} \\ Z_{2t} \\ \vdots \\ Z_{k*+G^d-1t} \end{pmatrix} X_{it} = \begin{pmatrix} Z_{1t} \\ Z_{2t} \\ \vdots \\ Z_{k*+G^d-1t} \end{pmatrix} \quad [-\beta_{12} X_{2t} \cdots \beta_{1G^d} X_{G^d t} - r_{11} Z_{1t} \\ \cdots r_{1k} * Z_{k*} + U_{1t}]$$

両辺を T で割ってモーメントで表現すれば、

$$M_{xiz} = -\beta_d^{(1)} M_{xz} - r_* M_{z*z} + M_{uz}$$

この場合

$$M_{zxi} = [m_{z1xi}, m_{z2xi}, \cdots, m_{zk*+G^d-1xi}]$$

$$M_{xiz} = M_{zxi}'$$

$$\beta_d^{(1)} = [\beta_{12}, \cdots, \beta_{1G^d}]$$

$$M_{xz}^{(1)} = \begin{pmatrix} m_{x2z1} \cdots m_{x2zk*+G^d-1} \\ \vdots \\ m_{xG^dz1} \cdots m_{xG^dzk*+G^d-1} \end{pmatrix}$$

$$r_* = [r_{11} \cdots r_{1k*}]$$

$$M_{z*z} = \begin{pmatrix} m_{z1z1} \cdots m_{z1zk*+G^d-1} \\ \vdots \\ m_{zk*z1} \cdots m_{zk*zk*+G^d-1} \end{pmatrix}$$

$$M_{zu} = [m_{z1u1}, m_{z2u1}, \cdots, m_{zk*+G^d-1u1}]$$

である。

この式から

$$[\beta_d^{(1)} r_*] \begin{pmatrix} M_{xz}^{(1)} \\ M_{z*z} \end{pmatrix} = -M_{xiz} + M_{uz}$$

が直ちに求められる。

したがって、

$$[\beta_d^{(1)} r_*] = -M_{xiz} \begin{pmatrix} M_{xz}^{(1)} \\ M_{z*z} \end{pmatrix}^{-1} + M_{uz} \begin{pmatrix} M_{xz}^{(1)} \\ M_{z*z} \end{pmatrix}^{-1}$$

この式の右辺第二項は誤差項の一次結合であって、平均値 0、分散一定の確率変数のベクトルであるから、合計した結果は 0 である。

したがって、 $[\beta_d^{(1)} r_*]$ の推定値を $[\hat{\beta}_d^{(1)} \hat{r}_*]$ とすれば、

$$[\hat{\beta}_d^{(1)} \hat{r}_*] = -M_{xiz} \begin{pmatrix} M_{xz}^{(1)} \\ M_{z*z} \end{pmatrix}^{-1}$$

このようにして、操作変数法によって、パラメータを推定することができるが、この推定値は大標本特性値の一致性をもっている。

INSTV の応答形式

INSTV では過剰認定の方程式に対して、この方程式の先決変数とこの方程式に現われてこないで他の方程式から選ばれた先決変数を操作変数として、 NC_{G^d-1} 通りの先決変数の選び方の組合せ

がある。Computerは方程式 $EQ(k)$ を指定すれば自動的にJust Identifiableのように先決変数を順次選り、IDLSでパラメータを推定する。この場合次のような指定をしなければならない。

CONDITION	・ AUTO		
	・ SELECT		
	・ OPTIMAL		
	・ TEST		
TEST	・ CS	・ F	・ Z
	・ T	・ NM	・ D-W
	・ ER	・ SG	

9. 予 測 (Forecasting)

Forecasting Stage は Interpolation と Exterpolation に分けられる。Interpolation では観察値の存在する範囲内において推定されたパラメータをもとに、観察値と理論値の比較を行ない、Exterpolation では観察値をもとに、観察値を越えて将来の理論値と変動の範囲を与えてくれる。

Forecasting Stage の定義形式は次のようである。

FORECAST :

INPOL (m' , (n'), v_j , S_i , EQ_k , GR_l)

EXPOL (m' , (n'), v_i , S_j , EQ_k , GR_l)

ここで、 m' は推定方法であり、 n' は予測期間の指定である。予測期間 n' の指定は、

(n_1' , n_2' , n_3')

の形をとる。

n_1' は予測の最初の日、期、年を表わし、 n_2' は予測期の最終の日、期、年を表わし、 n_3' において日、期、年などの単位の指定を必ずつける。

Output 形式は Display および表示方法は使用者が行ない、Output の内容は自動的に表示される。

応答形式

Forecasting Stage では、主観分析によるか、あるいは定義形式の INPUT によるかして変数、方程式体系が表示され、INPOL, EXPOL のいずれかを選択するとともに予測期間の指定を行なわなければならない。

ただし、Forecasting を行なうためには、模型があらかじめ作成されていなければならない。

もし作成されていなければ確認 Substep で自動的に AUTO-RETURN で処理される。

VARIABLE	· X1 · X2 · X3
	· Z1 · Z2 · Z3
SECTOR	· S1 · S2 · S3
EQUATION	· EQ1 · EQ2 · EQ3
GROUP	· GR1 · GR2 · GR3
PERIOD	(, ,)
OPTION	· PERIOD · YR · BI · QT
	· MT · WK · DY
	· NUMBER · 0 · 1 · 2 · 3 · 4 · 5 · 6 · 7 · 8 · 9
	· HYPHON · -

10. シミュレーション (Simulation)

Simulation Stage に入るためには、あらかじめ模型が作成されていなければならない。Simulation Stage はある意味では、Forecasting Stage と Optimization Stage と模型の形および手法は類似しているが、変数、関数、条件を Operational に変化し、現実の姿をこの模型によって Simulate する。

Simulation Stage の定義形式は次のように行なわれる。

SIMULATION:

MODEL : ($V_i(n')$, $EQ_k(n')$, $GR_l(n')$)

OPERATE : ($V_i(n')$, $EQ_k(n')$, $GR_l(n')$)

COMPARE : (V_i , EQ_k , GR_l)

SELECT : ($V_i(n')$, $EQ_k(n')$, $GR_l(n')$)

Simulationを行なう場合は、MODEL: ($V_i(n')$, $EQ_k(n')$, $GR_l(n')$) によって最初にMODELの指定を行なう。もしMODELが与えられていない場合は、(n')によって V_i , EQ_k , GR_l の推定方法を指定する。 n' には、GPFS, DLS, IDLSなどがある。

MODELが決まっています、そのMODELの中でSimulationを行なうため、変数 v_i , 方程式 EQ_k , 方程式群 GR_l の何をOPERATIONALな手とするかについて、OPERATEのラベルに続いて変数、方程式、方程式群を指定しなければならない。

n' はOPERATEするための方法の指定であり、たとえば、

B - DIST 二項分布
P - DIST ポアソン分布
N - DIST 正規分布
E - DIST 指数分布
ER - DIST アーラン分布
U - DIST 一様分布

などが指定できる。

Simulationの結果を比較するために、ラベルCOMPAREを用いて、指定された変数、方程式、方程式群が比較される。

SELECTおよびOPTIMIZEは、Simulationの結果について前者は n' で指定された個数をOptimalな状態の結果から n' 個取出し、後者は最適な解だけを与える。

Two Persons Gameを行なうときは、WEとTHEYとに意志決定のStageが分かれ、THEYがNatureであるか選択主体であるかに分類され、さらに確率条件が加えられるかどうか質問される。

Multi-StageのBusiness Gameを行なう場合には、意志決定の項目と、その数およびルール並びに結果表の指定を行なわなければならない。

G P F S

G P F Sは、コンパイルの容易さと弾力的に拡張できるという2つの目的をもち、いわゆるモジュラー形式(コントロール・プログラムとサブルーティンの組みの)で作成される。G P F Sはつぎの6段階から成る。

- 第1段階 研究対象の時系列をデータとして計算機に読み込むか、あるいはデータを発生させる。
- 第2段階 予測を行なう。
- 第3段階 予測誤差を計算する。
- 第4段階 誤差計算の結果を印刷する。
- 第5段階 予測プロジェクションを行なう。
- 第6段階 点描写する。

(シミュレータの論理的な流れについては図10-1参照のこと。)

第1段階では、シミュレータの最初の選択が行なわれる。すなわち、この段階では、現実の時系列がデータとして読み込まれるか、あるいは、現実のデータがない場合には、擬正規乱数発生装置を利用し、時系列を発生させる。時系列が発生すると、利用者は、発生した時系列の平均、分散、自己相関の程度、平均的な直線傾向、二次項および、季節変動を指定することができる。時系列が読み込まれる場合は、その時系列が印刷される。時系列発生装置は、つぎの原則に従って作動する。0と1の間の一様確率分布から、12個の観察値が得られる。(これは、乱数表から12個の数値を拾い出すか、あるいは、内部乱数発生装置が利用される。)

X_i を上述のようにして得られた第*i*番目の一数とすれば X_i の期待値 $E(X_i) = \frac{1}{2}$ 、分散 $\text{Var}(X_i) = \frac{1}{12}$ である。 $S = \sum_{i=1}^{12} X_i$ とすれば、 $E(S) = 6$ 、 $\text{Var}(S) = 12$ $\text{Var}(X_i) = 1$ 、つぎに S から6を引けば、(これを U と呼ぶ) $E(U) = 0$ 、 $\text{Var}(U) = 1$ 、中心極限定理によって、 U は平均値0、分散1をもつ正規分布をなす。この時系列のつぎの値を発生させるために、さらに12個の観察値が選ばれ、このようにして上述の手順が繰り返される。

毎回12個の異なった観察値を得ることによって、自己相関のない時系列が得られる。自己相関関係を導入する場合は、G P F Sでは、つぎの U の計算に、それまでの数値を利用するにすぎない。ある U から、つぎの U へ計算を行なうときに、共通した数値の個数が C に等しければ、ある一期のラグに対する自己相関係数は $C/12$ である。 $C=12$ の場合、したがって分散が0の場合をのぞいて、このようにして、平均値0、分散1、および、自己相関係数 $C/12$ という1期のラグをもった時系列を得る。定数 K_0 を掛け、定数 K_1 を得ることによって、この時系列を平均 K_1 と分散 K_2^2 、($C=12$ の場合、分散は0であるから、これを除く)をもつ一つの時系列に変えることができる。時間の一次関数であるような一数を加えるか差引くことによって、(あるいは、 K_1 を時間の一次関数とすることによって)、直線の傾向が導入される。また二次式は、時間の二次関数を加えるか引くことによって得られる。季節性は、全観察期間にわたって繰り返される数列を掛けることによって導入される。

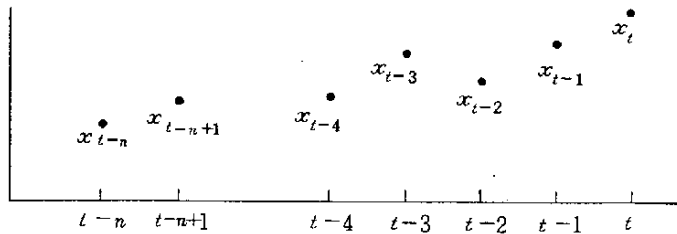


図10-2

予 測 公 式

第2段階において、実際の予測が行なわれる。GPFSは、つぎのような9つの異なった予測方式をもっている。

1. 単純移動平均法
2. 最小自乗法の傾向値をともなった単純移動平均法
3. 二重移動平均法
4. 最小自乗法の傾向値と指数季節変動をともなった単純移動平均法
5. 単純指数平滑法
6. 傾向値および、季節指数をともなった単純指数平滑法
7. 二次指数平滑法
8. 二重指数平滑法
9. 三重指数平滑法

この9つの予測方式の内容は、つぎのように要約することができる。方程式において用いられている変数およびパラメータをつぎのように約束する。

α ……平滑化定数 $0 \leq \alpha \leq 1$

β ……平滑化定数 $0 \leq \beta \leq 1$

r ……平滑化定数 $0 \leq r \leq 1$

$\bar{F}_i$ ……第*i*期の季節係数の推定値 ($i = 1, 2, \dots, m$, この場合*m*は一季節周期の期の数を表わす)

n ……単純移動平均における期の数

$\bar{T}_t$ ……*t*期における傾向線の推定値

X_t ……*t*期における時系列の観察値

$\bar{X}_t$ ……*t*期における時系列の期待値(平均値)の推定値

$\hat{X}_{t+L}$ ……*t+L*期の予測値(もし*t*が今期であるならば $\hat{X}_{t+L}$ は、将来の第*L*番目の期の予測値を与える。)

Z……初期条件設定期の長さ

1. 単純移動平均ルーティン (SMA)

$$\widetilde{X}_t = \frac{1}{n} (X_t + X_{t-1} + \cdots + X_{t-n+1}) = \frac{1}{n} \sum_{i=t-n+1}^t X_i \quad t=n \quad (1)$$

$$\widetilde{X}_{t-1} = \frac{1}{n} (X_{t-1} + X_{t-2} + \cdots + X_{t-n}) = \frac{1}{n} \sum_{i=t-n}^{t-1} X_i$$

$$\begin{aligned} \widetilde{X}_t &= \frac{1}{n} (X_t + X_{t-1} + \cdots + X_{t-n+1} + X_{t-n} - X_{t-n}) \\ &= \frac{1}{n} (X_t + X_{t-n}) + \frac{1}{n} (X_{t-1} + \cdots + X_{t-n}) \\ &= \widetilde{X}_{t-1} + \frac{1}{n} (X_t - X_{t-n}) \quad t \geq n+1 \quad (2) \end{aligned}$$

$$\widetilde{X}_{t+L} = \widetilde{X}_t \quad t \geq z \geq n \quad (3)$$

2. 最小自乗傾向値をもった単純移動平均法 (SMAT)

$$\widetilde{M}_t = \frac{1}{n} \sum_{i=t-n+1}^t X_i \quad t=n \quad (4)$$

$$\widetilde{M}_t = \widetilde{M}_{t-1} + \frac{1}{n} (X_t - X_{t-n}) \quad t \geq n+1 \quad (5)$$

X_t の t への回帰係数を $\widetilde{T}_t$ とすれば,

$$\widetilde{T}_t = \frac{\sum_{i=1}^n (X_i - \bar{X}) (t - \bar{t})}{\sum_{i=1}^n (t - \bar{t})^2}$$

$$t = \frac{n(1+n)}{2} \times \frac{1}{n} = \frac{(1+n)}{2}$$

$$\begin{aligned} \sum_{i=1}^n (t - \bar{t})^2 &= \sum_{t=1}^n t^2 - n(\bar{t})^2 \\ &= \frac{n(n+1)(2n+1)}{6} - n\left(\frac{1+n}{2}\right)^2 \\ &= \frac{n(n^2-1)}{12} \end{aligned}$$

であるから,

$$\begin{aligned}\widetilde{T}_t = & \frac{12}{n(n^2-1)} \left[\frac{(n-1)}{2} X_t + \frac{(n-3)}{2} X_{t-1} \right. \\ & \left. + \dots + \frac{n-(2n-1)}{2} X_{t-n+1} \right] \quad t = n \quad (6)\end{aligned}$$

同様に

$$\begin{aligned}\widetilde{T}_{t-1} = & \frac{12}{n(n^2-1)} \left[\frac{(n-1)}{2} X_{t-1} + \frac{(n-3)}{2} X_{t-2} \right. \\ & \left. + \dots + \frac{n-2(n-1)}{2} X_{t-n} \right]\end{aligned}$$

$$\begin{aligned}\widetilde{T}_t = & \frac{12}{n(n^2-1)} \left[\frac{(n-1)}{2} X_t + \frac{(n-3)}{2} X_{t-1} \right. \\ & \left. + \dots + \frac{n-(2n-1)}{2} X_{t-n+1} + \frac{n-(2n-1)}{2} X_{t-n} - \frac{n-(2n-1)}{2} X_{t-n} \right] \\ = & \widetilde{T}_{t-1} + \frac{2}{n(n^2-1)} \left[\frac{n-1}{2} X_t + \frac{n+1}{2} X_{t-n} - n\bar{M}_{t-1} \right] \quad t \geq n+1 \quad (7)\end{aligned}$$

$$\widetilde{X}_t = \bar{M}_t + \frac{(n-1)}{2} \widetilde{T}_t \quad t \geq n \quad (8)$$

$$\hat{X}_{t+L} = \widetilde{X}_t + \widetilde{T}_{tL} \quad t \geq z \geq n \quad (9)$$

この場合 $\widetilde{T}_t$ は最小自乗法による傾斜，すなわち X_t の t への回帰係数であるが， $\bar{M}_t$ （単純移動平均）は， n 期間の中央における推定値である。

3. 二重移動平均

この方法は，(4)，(5)，(8)，と同じような考え方から導くことができる。

$$\widetilde{M}_t = \frac{1}{n} \sum_{i=t}^{t+n-1} X_i \quad t = n \quad (10)$$

とすれば

$$\widetilde{M}_t = \widetilde{M}_{t-1} + \frac{1}{n} (X_t - X_{t-n}) \quad t \geq n+1 \quad (11)$$

さらに

$$\widetilde{M}_t^{(2)} = \frac{1}{n} \sum_{i=t}^{t+n-1} \widetilde{M}_i \quad t = 2n-1 \quad (12)$$

とすれば

$$\widetilde{M}_t^{(2)} = \widetilde{M}_{t-1}^{(2)} + \frac{1}{n} (\widetilde{M}_t - \widetilde{M}_{t-n}) \quad t \geq 2n \quad (13)$$

$$\tilde{M}_t = \tilde{M}_t^{(2)} + \frac{(n-1)}{2} \tilde{T}_t \quad \text{であるから}$$

$$\tilde{T}_t = \frac{2}{n-1} (\tilde{M}_t - \tilde{M}_t^{(2)}) \quad t \geq n-1 \quad (14)$$

$$\tilde{X}_t = 2\tilde{M}_t - \tilde{M}_t^{(2)} \quad t \geq n-1 \quad (15)$$

$$\hat{X}_{t+L} = \tilde{X}_t + \tilde{T}_t L \quad t \geq z \geq n-1 \quad (16)$$

4. 最小自乗傾向値と季節指数を伴った単純移動平均ルーティン (SMATS)

このルーティンは、季節的影響度について、若干の調節をする以外は、上述の § 2 で与えられた傾向値を伴った単純移動平均法と考え方は同じである。季節指数が、はじめに推定され、移動平均を作るときに用いられる観測値は、この季節係数によって割られ、季節性が除去される。このようにして、(4), (5), (6), (7), (8), および(9)の方程式によって与えられたと同様な手順が用いられる。予測値は、固有の季節係数を生ずることによって得られる。

$\tilde{F}_t$ を第 t 期の季節係数の推定値 (どの時期においても m 個の係数がある) としよう。下記の方程式の中で $t \leq m$ の場合は、(25) から (27) までの方程式によって得られた下の初期値が使われる。

$$\begin{aligned} \tilde{D}_t &= \text{季節性を除去した移動平均} \\ &= \frac{1}{n} \sum_{i=t}^{t-n+1} \frac{X_i}{\tilde{F}_{i-m}} \quad t=n \quad (17) \end{aligned}$$

$$\tilde{D}_t = \tilde{D}_{t-1} + \frac{X_t(\tilde{F}_{t-m} - X_{t-n}/\tilde{F}_{t-n-m})}{n} \quad t \geq n+1 \quad (18)$$

$$\begin{aligned} \tilde{T}_t &= \frac{12}{n(n^2-1)} \left[\frac{n-1}{2} X_t / \tilde{F}_{t-m} + \frac{n-3}{2} X_{t-n} / \tilde{F}_{t-m-1} + \dots \right. \\ &\quad \left. + \frac{n-(2n-1)}{2} X_{t-n-1} / \tilde{F}_{t-m-n+1} \right] \quad t=n \quad (19) \end{aligned}$$

$$\begin{aligned} \tilde{T}_t &= \tilde{T}_{t-1} + \frac{12}{n(n^2-1)} \left[\frac{n-1}{2} X_t / \tilde{F}_{t-m} \right. \\ &\quad \left. + \frac{n+1}{2} X_{t-n} / \tilde{F}_{t-m-n} - n\tilde{D}_{t-1} \right] \quad t \geq n+1 \quad (20) \end{aligned}$$

$$\tilde{X}_t^d = \tilde{D}_t + \frac{(n-1)}{2} \tilde{T}_t \quad (21)$$

季節性を除去した期待値とすれば、

$$\widetilde{F}_t = \beta X_t / \widetilde{X}_t^d + (1-\beta) \widetilde{F}_{t-m} \quad t \geq \max[m+1, n] \quad (22)$$

$$\widetilde{X}_t = [\widetilde{X}_t^d] \widetilde{F}_{t-m} \quad t \geq n \quad (23)$$

$$\hat{X}_{t+L} = (\widetilde{X}_t^d + \widetilde{T}_t L) \widetilde{F}_{t-m+L}, \quad t \geq z \geq \max[n, m], \quad L=1, 2, \dots, m \quad (24)$$

z は m の倍数であり、また $L > m$ であるから、個々の $\widetilde{F}$ を再び利用することによって、予測値を得ることができる。 $\widetilde{F}$ の初期値に対して、

$$\widetilde{F}_t = \frac{X_t}{\bar{X}_i - (\frac{m+1}{2} - j) \widetilde{T}_1} \quad t \leq z \quad (25)$$

を得る。この場合、 $\widetilde{T}_1$ は (35) 式によって与えられる初期傾向値の推定値であり、 $\bar{X}_i$ は t 全体の周期における X の平均値であり、また j は第 1 期、第 2 期などから第 m 期までの周期内の期を表わす。初期条件を与えるときに、各期に対応する $\widetilde{F}_t$ が平均化される。すなわち、

$$\widetilde{F}_j = \frac{\sum_{i=0}^{(z/m)-1} \widetilde{F}_{j+im}}{Z/m} \quad j = 1, 2, \dots, m \quad (26)$$

とすれば $\widetilde{F}_j$ は初期値 $\widetilde{F}$ を与えるために基準化される。

$$\widetilde{F}_j = \bar{F}_j \left(\frac{m}{\sum_{i=1}^m \bar{F}_i} \right) \quad j = 1, 2, \dots, m \quad (27)$$

また $\widetilde{F}_t = \widetilde{F}_j$ であって、 $t \leq m$ 、もし $m < n$ ならば、方程式 (22) が用いられるときは、 $t = n$ まで $\widetilde{F}_j$ が繰り返される。もし $m \geq n$ ならば方程式 (22) は $t \geq m+1$ として用いられる。

5. 単純指数平滑ルーティン (SES)

指数平滑法というのは、

$$\text{新推定値} = \text{旧推定値} + \alpha (\text{新データ} - \text{旧推定値})$$

あるいは、

$$\text{新推定値} = \alpha (\text{新データ}) + (1-\alpha) (\text{旧推定値})$$

である。

$$\widetilde{S}_t = \alpha X_t + (1-\alpha) \widetilde{S}_{t-1} \quad t \geq 3 \quad (28)$$

$$\widetilde{S}_2 = (X_2 + X_1) / 2 \quad (29)$$

$$X_{t+L} = \widetilde{X}_t = \widetilde{S}_t \quad t \geq z \geq 3 \quad (30)$$

6. 傾向値と季節指数を伴う単純指数平滑法 (SESTS)

$$\tilde{S}_t = \alpha X_t + \tilde{F}_{t-m} + (1-\alpha)(\tilde{S}_{t-1} + \tilde{T}_{t-1}) \quad t \geq 2 \quad (31)$$

$$\tilde{S}_1 = \frac{\sum_{i=1}^m X_i}{m} \quad (32)$$

$$\tilde{F}_t = \beta(X_t / \tilde{S}_t) + (1-\beta) \tilde{F}_{t-m} \quad t \geq m \quad (33)$$

$t \leq m$ ならば, (25) 式から (27) 式までに与えられている $\tilde{F}$ の初期推定値を利用する。

$$\tilde{T}_t = r(\tilde{S}_t - \tilde{S}_{t-1}) + (1-r)\tilde{T}_{t-1} \quad t \geq 2 \quad (34)$$

$$\tilde{T}_1 = \frac{\sum_{i=2-m}^z X_i}{m} - \frac{\sum_{i=1}^m X_i}{m} \quad (35)$$

$$\tilde{X}_t = \tilde{S}_t \tilde{F}_{t-m} \quad t \geq 2 \quad (36)$$

$$\hat{X}_{t+L} = [\tilde{S}_t + L\tilde{T}_t] \tilde{F}_{t-m+L} \quad t \geq z \geq 2 \text{ および } L=1, 2, \dots, m \quad (37)$$

$L > m$ の場合は, 予測値は, 個々の $\tilde{F}$ を再び利用することによって得られる。

7. 二次指数平滑ルーティン (SOES)

$$\tilde{T}_t = \tilde{S}_t - \tilde{S}_{t-1} \quad (38)$$

$$\tilde{S}_t = \tilde{X}_t = \alpha \tilde{X}_t + 2(1-\alpha)\tilde{S}_{t-1} - (1-\alpha)\tilde{S}_{t-2} \quad t \geq 4 \quad (39)$$

$$\tilde{S}_2 = (X_2 + X_1) / 2 \quad (40)$$

$$\tilde{S}_3 = (X_3 + X_2) / 2 \quad (41)$$

$$\hat{X}_{t+L} = \tilde{X}_t + L\tilde{T}_t \quad t \geq z \geq 4 \quad (42)$$

方程式 (39) は, $\tilde{S}_t = \alpha X_t + (1-\alpha)(\tilde{S}_{t-1} + \tilde{T}_{t-1})$ と書きなおすことができる。

8. 二重指数平滑法 (DES)

$$\tilde{S}_t = \alpha X_t + (1-\alpha)\tilde{S}_{t-1} \quad t \geq 3 \quad (43)$$

$$\tilde{S}_2 = (X_2 + X_1) / 2 \quad (44)$$

$$\tilde{S}_t^{(2)} = \alpha \tilde{S}_t + (1-\alpha)\tilde{S}_{t-1}^{(2)} \quad t \geq 3 \quad (45)$$

$$\tilde{S}_2^{(2)} = \tilde{S}_2 \quad t \geq 3 \quad (46)$$

$$\tilde{T}_t = [\alpha / (1-\alpha)] (\tilde{S}_t - \tilde{S}_{t-1}^{(2)}) \quad (47)$$

$$\tilde{X}_t = 2\tilde{S}_t - \tilde{S}_t^{(2)} \quad t \geq 3 \quad (48)$$

$$\hat{X}_{t+L} = \tilde{X}_t + L\tilde{T}_t \quad t \geq z \geq 3 \quad (49)$$

9. 三重指数平滑法 (TES)

$$\tilde{S}_t = \alpha X_t + (1-\alpha) \tilde{S}_{t-1} \quad t \geq 3 \quad (50)$$

$$\tilde{S}_2 = (X_2 + X_1) / 2 \quad (51)$$

$$\tilde{S}_t^{(2)} = \alpha S_t + (1-\alpha) \tilde{S}_{t-1}^{(2)} \quad t \geq 3 \quad (52)$$

$$\tilde{S}_2^{(2)} = \tilde{S}_2 \quad (53)$$

$$\tilde{S}_t^{(3)} = \alpha S_t^{(2)} + (1-\alpha) \tilde{S}_{t-1}^{(3)} \quad t \geq 3 \quad (54)$$

$$\tilde{S}_2^{(3)} = \tilde{S}_2 \quad (55)$$

$$\tilde{T}_t = [\alpha/2(1-\alpha)^2] \{ (6-5\alpha) \tilde{S}_t - 2(5-4\alpha) \tilde{S}_t^{(2)} + (4-3\alpha) \tilde{S}_t^{(3)} \} \quad t \geq z \geq 3 \quad (56)$$

$$\tilde{Q}_t = [\alpha^2/(1-\alpha)^2] \{ \tilde{S}_t - 2\tilde{S}_t^{(2)} + \tilde{S}_t^{(3)} \} \quad t \geq 3 \quad (57)$$

とすれば,

この場合 $\tilde{Q}_t$ は, どのような二次成分形式をも説明している。

$$\tilde{X}_t = 3\tilde{S}_t - 3\tilde{S}_t^{(2)} + \tilde{S}_t^{(3)} \quad t \geq 3 \quad (58)$$

$$X_{t+L} = \tilde{X}_t + L\tilde{T}_t + (L^2/2)\tilde{Q}_t \quad t \geq z \geq 3 \quad (59)$$

いかなる作業に対しても, 初期条件設定期の Z は, つねに, これから始めようとする予測方式が必要とする時間よりも, 大きくなければならない。二重移動平均方式は, 通常, $t=2n-1$ になるまで待たなければならぬから, 開始がもっとも遅い。だが, 季節指数をもった方式では, 季節周期が非常に長くなれば, 遅くなることもある。このようにして, 利用者は, Z が $(2n-1, m)$ の最大よりも大きくなるように, 二重移動平均法において使用される最長の n と, 季節性を利用する方式においては, 周期の長さ m とを考慮しなければならない。また Z は, m の倍数でなければならない。

インプット

GPFSの利用者は, LEADTIME (リード・タイム) という1つの変数の範囲を指定する。

LEADTIME とは, 予測が行なわれる将来の期間の数である。たとえば, LEADTIME を1から8の範囲と指定すれば, このシミュレータは, 各時期において, 予測時点の各時期において, 1つはつぎの期, 1つは次期ともう1つつぎの期など, 将来の8期までの予測, すなわち, 8個の予測を行ない, 各期についで, 予測値と現実値との間の誤差を記録する。これが第3段階である。

誤差が, 各期に対して記録される。さらに, GPFSによって予測が行なわれるが, 誤差が記録されないようなく (利用者によって指定される) 初期条件設定期が1個ある。この指定は, 誤差計算の初期条件をクリアするために行なわれるのである。各期の計算が進み, 予測値が得られると, 誤差が蓄えられ, 最終期の計算が終了した後に予測誤差の標準偏差が, 各リード・タイムに対して計算される。

さらに、予測誤差を計算する場合、予測計画のための販売資料を作成するのであるか、在庫管理のための需要を予測するのであるかによって、誤差計算の方法が本質的に異なるのであるが、GPFSは、使用者がSALE、あるいはINVENTORYのいずれかを指定することによって、いずれのシミュレーションにも利用することができる。すなわち、SALEの場合は、ある期の予測値と現実値との間の差が問題となるが、INVENTORYの場合は、ある一定のリード・タイム全体にわたって利用できる在庫量を予測しなければならない。したがって、在庫予測は、累積量あるいは、予測時点までの合計した数値に対して行なわれなければならない。たとえば、変数LEADTIME=3とすれば、われわれが関心をもっている数値というのは、1期の予測値、つぎに1期を含む2期間の予測値、さらに、2期までを含む3期間の予測値である。この場合の誤差は、これら3期間の現実値の合計と、この予測値の合計との差である。このようにして、LEADTIMEは、販売の問題の場合と、在庫型の場合とでは、異なった意味をもっている。販売型では、LEADTIMEが k に等しいということは、将来の第 k 期の予測を1つすればよいが、在庫型の場合には、第 k 期までの累積量を予測することを意味しているのである。いずれの場合にも、LEADTIMEの各数値に対して、1組の標準偏差を得ることができる。

11. 最適化問題 (Optimization)

JUMPS の現段階では、Optimization Stage は次の5つの Step からなるものとする。

- (1) Simultaneous Equation Step
 - (2) Linear Programming Step
 - (3) Quadratic Programming Step
 - (4) Dynamic Programming Step
 - (5) GO-NO-Network
 - (6) etc.
- (1) SE Step (Simultaneous Equation Step)

この Step では、式の形が

$$AX = R$$

の連立方程式を解して X の値を求める。ラベルは、

OPTIMIZATION;

SE: (EQ k, GR l)

であり、EQ k の場合は変数と方程式の数とが一致しなければならない。不一致があれば、応答形式を通じて修正される。

マトリックス A およびベクトル R を変化させて Simulation を行なう場合は、Simulation の指定条件

OPERATE, COMPARE, SELECT, OPTIMIZE

を使用することができる。マトリックス A および R が時間 t の関数であれば、 A および R の各要素について時間 t の関数を指定することができる。たとえば、

$$EQ1(T) = A_{10}(T) + A_{11}(T)X1 + A_{12}(T)X2 + A_{13}X3$$

$$A_{10}(T) = B_{10} + B_{11} * T$$

$$A_{11}(T) = B_{20} + B_{21} * T + B_{22} ** T$$

$$A_{12}(T) = B_{30} + B_{31} * T + B_{32} ** T$$

であって、

$$A(T(n'))X = R(T(n'))$$

ならば、 n' において指定された $n' = (n_1, n_2, n_3)$ の n_1 によって初期、 n_2 によって終期、 n_3 によって日、月、期、年などの単位期間の指定を行ない、 T の変化に応じた X_i を求めることができる。

METHOD の指定が行なわれなければ、Sweep-out によって連立方程式は解かれる。Gauss-Seidel あるいは Gramér など特別の指定が行なわれれば、それらの方法によっ

て連立方程式が解かれる。

JUMPS の内部では、 A_{ij} ($i=1, 2, \dots, m, j=1, 2, \dots, n$) について $m=n$ であるかどうか自動的に検討され、 $m=n$ でなければ Warning Sign が出される。

(2) L P Step (Linear Programming Step)

この Step では初めにLPの形を次のように指定する。

$AX \leq R$ 型

OPTIMIZATION

LP; $AX + U = R$

MAX; CX

$AX \geq R$ 型

LP; $AX - U = R$

MIN; CX

次にAの配列の大きさを次のように指定する。

MATRIX: $A(m, n)$

i) マトリックスAが実数ならば、連続関数として処理し、マトリックスAが整数ならば、

Integer Programming として処理される。

また、A, X, RについてREAL, INTEGERを指定することができる。

ii) Degenerate が発生すれば、演算は一時停止し、

① 自動的に続行

② 人間の選択

③ Criss-Cross 法

の3種類の中から1つの方法の選択・指定を機械側で要求する。

① 自動的に続行する場合は、機械は Degenerate の発生した方程式あるいはプロセスを乱数によって1つ選択する。

② 人間の選択の場合は、応答形式にしたがう。

③ Criss-Cross 法の場合は、機械が自動的に Primal-Dualを反転しながら解に到達する。

iii) 解は X_j および CX_j の値である。ただし、Primal の状態においてX, RはColumn Vector, C は row vector と約束する。

(3) Q P Step (Quadratic Programming Step)

この Step では初めにQPの形を次のように指定する。

標準型

OPTIMIZATION

QP; $MX = E$

$SX = 1$

一般型

OPTIMIZATION

$$QP: AX=R$$

また、目的関数について MIN と MAX について次のように指定する。

$$\text{MIN: } V = X C X$$

$$\text{MAX: } U = L E - X C X$$

現段階における QP の処理方法は Markowitz の Critical Method と Sharppe の Diagonal Method を用意する。指定の方法はたとえば次のようである。

METHOD : CRITICAL

MATRIX : $A(m, n, V_{ij})$

VECTOR : $X(n, V_j)$

VECTOR : $R(m, V_i)$

または METHOD : DIAGONAL

VECTOR : $M(n, v_i)$

VECTOR : $X(n, v_i)$

EQ 1 : $I = A(N+1) + C(N+1)$

EQ 2 : $R = A(I) + B(I)I + C(I)$

収益の平均と分散

マーコヴィッツが提出した問題を要約すれば、高い予想収益をのぞましいと考え、また収益の不確実性をのぞましくないと考えている投資家が、いくつかの証券を保有することによって（あるいは、いくつかの金融資産を保有することによって）、将来の収益を予想して、一定金額の貨幣をこれらの資産にどのように配分すれば、最大の期待効用を得ることができるであろうかという問題である。

いま単純化のために n 種の証券に投資をする場合を考えてみよう。マーコヴィッツは、ある証券が、ある年度あるいは期間 t にもたらした収益を、つぎのように定義する。

P_{it-1} 第 i 番目の証券の $t-1$ 期末の価格

P_{it} 第 i 番目の証券の t 期末の価格

D_{it} 第 i 番目の証券の t 期中の配当

R_{it} 第 i 番目の証券の t 期中の収益率

とすれば、

$$R_{it} = \frac{P_{it} - P_{it-1} + D_{it}}{P_{it-1}} \quad (1)$$

R_{it} は t 期中に投資された単位貨幣量（たとえば 1 ドル）についての収益あるいは損失を意味する。いま、いくつかの証券（ $i = 1, 2, 3, \dots, n$ ）があって、これらの証券の過去の業績だけを考え、それぞれの証券の数量を異にした組合せを求めると、複数（ n 種まで）の証券によってもたら

された収益率を考えることができる。

複数以上の証券をもつ場合、その組合せは無数にあり、組合せに応じて収益を計算することができる。複数以上の証券、あるいはその他の金融資産に拡張して、複数以上の金融資産を組み合わせることを、ポートフォリオといい、この無数に考えられるポートフォリオ（すなわち組み合わせられた金融資産）の中から、危険が最も少なく、しかも収益の最も高いポートフォリオを選択することが、ポートフォリオ・セレクションの問題である。

1つのポートフォリオの平均収益は、このポートフォリオに含まれる証券の最高収益より低く、最低収益より高いどこかの点にある。収益だけを考えてみれば、複数以上のポートフォリオを考えず、平均収益のもっとも高い証券を単独に選択することがあるいはもっとも望ましいかもしれない。だが、各資産について毎期の収益の変動をみると、Aは平均収益が低いが安定的で損失も小さく、たとえば最大25%の損失があり、Bは平均収益が高いが、不安定で、たとえば最大45%の損失があるとすれば、たとえ平均収益が最高であっても単独に保有することは損失を招くことが大きいからいずれかの証券をどのように組み合わせるべきか保有したならばよいであ

ろうか。A、Bからなる1つのポートフォリオをつくと、各証券の収益は時間の経過にともなって変動するが、同方向に変動すればするほど、互いに相殺し合う程度は小さく、したがって、ポートフォリオの収益の変動は大きい。逆にいうと、証券の変動が互いに関係する度合いが小さければ小さいほど、このポートフォリオの平均収益は安定的であり、しかも損失は小さい。ある証券を選択する場合、収益が高く、バラツキが小さいほど望ましいが、このバラツキを測

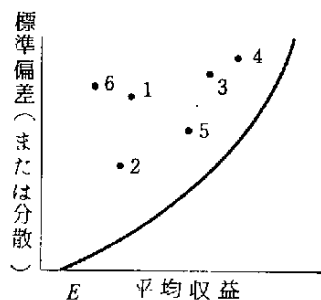


図 11-1

定するには一般に用いられているように、標準偏差の考え方をを用いることができる。あるポートフォリオにおいても、標準偏差の考え方を利用すると、証券間の変動の相関が高ければ、各証券の変動が相殺されないから、全体としての標準偏差は大きくなる。したがって、ポートフォリオの標準偏差は、

- a 各証券の標準偏差、
- b 各一對の証券の相関関係
- c 各証券に投資された金額

によって決定される。

いま、 $X-Y$ 平面上において、 X 軸方向に n 種の個々の証券の平均収益をとり、 Y 軸方向に、それぞれの平均収益に対応する証券の標準偏差（あるいは分散）をとり、つぎに、これらの証券の任意の組合せからなるポートフォリオの平均収益と標準偏差（または分散）を無数にとれば、個々の証券の平均収益と標準偏差（または分散）によって表わされる n 個の点と、1本の曲線に囲まれた領域とを決めることができる。この曲線は後述するように各証券の一対ずつの相関関係から出てくるものであるが、図 11-1 からわかるように、ポートフォリオを2つのグループに分けてみると、

- 1 平均収益および収益の標準偏差（または分散）が曲線上に表わされるものと、

2 平均収益および収益の標準偏差（または分散）が曲線の上に表わされるものと2分類される。しかも、いかなるポートフォリオも、この曲線以下にはない。ある平均収益をもたらす点をたとえば、Eで固定すると、曲線から上方に1つの証券あるいは複数の証券からなるポートフォリオの標準偏差（あるいは分散）があれば、バラツキが大きく、もっとも小さなバラツキをしめす標準偏差（あるいは分散）は曲線上にある。このことは、いくつかの証券を組み合わせることによって、バラツキを少なくすることができることを意味している。なぜいくつかの証券の組合せによるポートフォリオの標準偏差（または分散）が単独の証券の分散より小さくなるかは、つぎのように簡単に証明することができる。

先述したように、A, B 2種の証券があって、それぞれのt期の収益が R_{At}, R_{Bt} であり、 $t=1, 2, \dots, T$ 期全体の平均収益をそれぞれ $\bar{R}_A, \bar{R}_B$ それぞれの分散を σ_A^2, σ_B^2 , 共分散を σ_{AB} とすると、

$$\begin{aligned}\sigma_A^2 &= \frac{1}{T} \sum_{t=1}^T (R_{At} - \bar{R}_A)^2 \\ \sigma_B^2 &= \frac{1}{T} \sum_{t=1}^T (R_{Bt} - \bar{R}_B)^2 \\ \sigma_{AB} &= \frac{1}{T} \sum_{t=1}^T (R_{At} - \bar{R}_A)(R_{Bt} - \bar{R}_B)\end{aligned} \quad \dots\dots\dots (8)$$

A, B からなるポートフォリオの分散を σ^2 とし、A, B をそれぞれ $p\%$ $q\%$ 保有したとすれば、

$$\begin{aligned}\sigma^2 &= \frac{1}{T} \sum_{t=1}^T [p(R_{At} - \bar{R}_A) + q(R_{Bt} - \bar{R}_B)]^2 \\ &= p^2 \sigma_A^2 + q^2 \sigma_B^2 + 2pq\sigma_{AB}\end{aligned} \quad \dots\dots\dots (9)$$

いま $\sigma_A^2 > \sigma_B^2$ の場合を考えてみると、

(9) 式は

$$\begin{aligned}\sigma^2 &\leq p^2 \sigma_A^2 + q^2 \sigma_A^2 + 2pq\sigma_{AB} = \\ &\sigma_A^2 - 2pq\sigma_{AB} \left(\frac{\sigma_A}{\sigma_B} - \rho_{AB} \right) \frac{\sigma_A \sigma_B}{\sigma_B}\end{aligned} \quad \dots\dots\dots (10)$$

(10) 式の右辺第2項の σ_A / σ_B は $\sigma_A^2 > \sigma_B^2$ であれば1より大であり、 $\rho_{AB} = \frac{\sigma_{AB}}{\sigma_A \sigma_B}$

すなわち相関係数であって、 $-1 \leq \rho_{AB} \leq 1$ であるから、右辺第2項の括弧の中は、必ず正である。

$$\therefore \sigma_A^2 \geq \sigma^2 \quad \dots\dots\dots (11)$$

(11) を (9) に代入すれば

$$\sigma^2 \geq \sigma_B^2 \quad \dots\dots\dots (12)$$

がみちびかれ、したがって

$$\sigma_A^2 \geq \sigma^2 \geq \sigma_B^2 \quad (13)$$

このことは、投資者が、いくつかの異なった分散をもつ証券をもつことによって、全体の分散を小さくすることができることを意味し、したがって危険を回避することができるのである。ポートフォリオの分析は、このようにして、いろいろな予想収益の水準 E のおおのに対して、分散が最小になるようなポートフォリオを選択することである。選択者には、それゆえに欲している収益水準と危険水準のどのような組合せを決定するかが残されている。

効率的ポートフォリオ

マーコヴィッツは、間一の収益をもつあらゆる可能なポートフォリオの中で、収益の分散が最小であるポートフォリオを効率的ポートフォリオ (efficient portfolio) と定義しているが、図11-1の曲線部分がこの効率的なポートフォリオにあたっている。これを別の表現でいえば、もし収益の分散が同一であらゆる可能なポートフォリオの中で、最大の予想収益をもつポートフォリオは効率的であると定義することもできる。

マーコヴィッツは、ポートフォリオ・セレクションの分析が、

第1に、いくつかの効率的なポートフォリオの集合をつくり、

第2に、危険と収益の最適な組合せを投資家に与える1つのポートフォリオを選ぶことからなるといっている。効率的ポートフォリオの集合から、最適なポートフォリオを選ぶために、マーコヴィッツは、はじめにつきのように3次計画 (Quadratic Programming) の問題として定式化を行なう。

X_i を証券 i に配分される資産の割合とし、(先の例では p, q にあたる)、 μ_i および σ_{ii} をそれぞれ、第 i 証券からの収益の期待値および分散とし、 σ_{ij} を第 i 証券と第 j 証券からの収益間の共分散、 E と V とをそれぞれ、そのポートフォリオからの収益の期待値および分散とすれば、

$$\sum_{i=1}^n X_i \mu_i = E \quad (14)$$

$$\sum_{i=1}^n X_i = 1 \quad (15)$$

$$X_i \geq 0 \quad (16)$$

の条件の下で、

$$\min V = \sum_{i=1}^n \sum_{j=1}^n X_i X_j \sigma_{ij} \quad (17)$$

という2次計画の問題を解けば、 E の各可能な値に対して、効率的ポートフォリオの集合を求めることができる。

一般的なポートフォリオが実行可能であるための制限条件式は、

$$\left. \begin{array}{l} A_1 X_1 + A_2 X_2 + A_3 X_3 + \dots + A_N X_N = B \\ \text{または} \\ A_1 X_1 + A_2 X_2 + A_3 X_3 + \dots + A_N X_N \geq B \\ \text{または} \\ A_1 X_1 + A_2 X_2 + A_3 X_3 + \dots + A_N X_N \leq B \end{array} \right\} \dots \dots \dots (18)$$

の3つの形のいずれかあるいは複数をとっている。

目 的 関 数

目的関数 (17) 式の意味はどうであろうか、証券のポートフォリオ分散は、

$$V = X_1^2 \sigma_{11} + X_2^2 \sigma_{22} + X_3^2 \sigma_{33} + 2 X_1 X_2 \sigma_{12} \\ + 2 X_1 X_3 \sigma_{13} + 2 X_2 X_3 \sigma_{23}$$

$X_3 = 1 - X_1 - X_2$ をこれに代入して整理すると、

$$V = X_1^2 \{ \sigma_{11} - 2 \sigma_{13} + \sigma_{33} \} + X_2^2 \{ \sigma_{22} - 2 \sigma_{23} + \sigma_{33} \} \\ + 2 X_1 X_2 \{ \sigma_{12} - \sigma_{13} - \sigma_{23} + \sigma_{33} \} \\ + 2 X_1 \{ \sigma_{13} - \sigma_{33} \} + 2 X_2 \{ \sigma_{23} - \sigma_{33} \} + \sigma_{33}$$

各証券の分散および、各証券間の共分散がつぎのように与えられたとすると、

$$\sigma_{11} = \sigma_{22} = 0.01, \\ \sigma_{33} = 0.04, \sigma_{12} = 0.003, \\ \sigma_{13} = \sigma_{23} = 0 \\ V = 0.05 X_1^2 + 0.05 X_2^2 \\ + 0.09 X_1 X_2 - 0.08 X_1 \\ - 0.08 X_2 + 0.04$$

目的関数の条件を満足する V に任意の値を与えると、楕円を無数に描くことができる。この楕円は V を一定にした等分散曲線 (Iso-variance curve) である。 V が小さければ、投資者は危険をなるべく小さく考え、すなわち、危険を回避しているのであり、 V が大きければ危険にはあまりこだわっていないといえることができる。前者を投資家といい、後者を投機家ということもある。 V を小さくしていくとポートフォリオが実行可能であるか可能でないかにかかわらず楕円の中心 C に到達する。

目的関数 V を一般的に考えてみると、

$$\begin{aligned}
V &= \sum_{i=1}^n \sum_{j=1}^n X_i X_j \sigma_{ij} \\
&= \{X_1^2 \sigma_{11} + X_2^2 \sigma_{22} + \dots + X_n^2 \sigma_{nn} \\
&\quad + 2(X_1 X_2 \sigma_{12} + X_1 X_3 \sigma_{13} + \dots + X_{n-1} X_n \sigma_{n-1n})\} \dots \dots \dots (19)
\end{aligned}$$

であるから、もしも、投資対象である証券がある特定の産業にのみかよった場合のように、証券相互間の収益の相関が非常に高い場合は目的関数は、必ずしも2次計画にはならない。すなわち、このような例で、相関係数 ρ_{ij} が1を想定できるような場合は、

$$\rho_{ij} = \frac{\sigma_{ij}}{\sigma_i \sigma_j} = 1$$

であって、(ただし、 σ_i, σ_j は第 i 証券および第 j 証券それぞれの標準偏差) $\sigma_{ij} = \sigma_i \sigma_j$ とおくことができるから、目的関数 V は、

$$\begin{aligned}
V &\leq \{X_1^2 \sigma_{11} + X_2^2 \sigma_{22} + \dots + X_n^2 \sigma_{nn} \\
&\quad + 2(X_1 X_2 \sigma_1 \sigma_2 + \dots + X_{n-1} X_n \sigma_{n-1} \sigma_n)\} \\
&= (X_1 \sigma_1 + X_2 \sigma_2 + \dots + X_n \sigma_n)^2 \dots \dots \dots (20)
\end{aligned}$$

ゆえに目的関数 V は、

$$V^{1/2} \leq X_1 \sigma_1 + X_2 \sigma_2 + \dots + X_n \sigma_n \dots \dots \dots (21)$$

によって置き換えることができる。すべての相関係数 $\rho_{ij}=1$ ならば完全に等式が成立し、目的関数は線型表示となり、したがって、線型計画と同質の問題となる。他方もし証券相互間の収益の変動が独立であれば、 $\rho_{ij}=0$ であり、したがって $\sigma_{12}=\sigma_{13}=\dots=\sigma_{n-1n}=0$ であるから、目的関数は、

$$V = X_1^2 \sigma_{11} + X_2^2 \sigma_{22} + \dots + X_n^2 \sigma_{nn} \dots \dots \dots (22)$$

である。

(14)、(15)、(16)式を制限条件として目的関数(17)を最小にすることとはどういう意味であろうか。前例に引き続いて3証券のポートフォリオを考えてみれば、ちょうど重ね合わせたような図を描くことができる。ポートフォリオの収益には、 $E_1, E_2, E_3, \dots$ など無数にあり、したがって、平行な無数の等平均直線があり、またその分散も $V_1, V_2, V_3, \dots$ など無数にありしたがって同一の中心をもつ楕円が無数にあるが、いま仮に特定の値 E_1 を入れた場合 $E=E_1$ は E_1 と分散の関係はどうなっているであろうか、 E_1 の等平均直線上では同一の収益を期待することができるが、1点でこの等平均直線に接する分散の楕円すなわち等分散曲線は1つしかない。もし等平均直線が2点で等分散曲線と交われば、同一収益ではあっても先の分散よりも大きく、したがって危険が大きい。もし等平均直線に等分散曲線が接しなければ、分散は小さいが、 E_1 の収益を得ると

とはできない。このようにしてある E の値に対して必ず1つのもっとも小さい分散がなければならない。無数の等平均直線と接する無数の等分散曲線との無数の点の軌跡をつくと直線が得られる。この直線は、同一の期待収益をもつポートフォリオの中で、最小の分散をもつすべての点の軌跡であり、これを臨界線 (critical line) という。逆にいうと臨界線上に、ある期待収益の最小の分散が存在することになる。したがって、臨界線は必ず等分散曲線の中心 C を通るが、それは実行可能ポートフォリオの集合を横切る場合もあれば横切らない場合もある。われわれのポートフォリオが効率的であるためには、ポートフォリオの点 P が下記の条件を満足しなければならない。

- 1 P は実行可能ポートフォリオである。
- 2 もし実行可能ポートフォリオが、より大きな期待収益をもつならば、それはポートフォリオ P よりもより大きな分散をもっている。
- 3 もし実行可能ポートフォリオがより小さな期待収益の分散をもつならば、それは、ポートフォリオ P よりも小さな期待収益をもっている。

あるポートフォリオが1の条件を満たしていなければ、そのポートフォリオは、効率的でもなければ、非効率的でもなく、単に実行可能ではないだけである。条件1, 2, 3, のうち2, 3, をそれぞれ1つ欠除するポートフォリオもあるが、これらは効率的ではない。点 C が実行可能ポートフォリオであれば、これら3つの条件を満足するから、効率的ポートフォリオであって、このポートフォリオの分散より小さい分散をもつポートフォリオはない。したがって、いかなるポートフォリオも、 C の期待収益より小さな期待収益をもつポートフォリオは効率的ではない。分散の視点から考えると、分散を一定にしての期待収益の最大値は、臨界線上にある。

C が実行可能ポートフォリオでなければ、 C の方向に向かって臨界線上の実行可能な点までが効率的ポートフォリオである。

以上で理解されるようにさらに、証券の種類を増加し4種の証券のポートフォリオをとり上げれば、3次元空間で説明することができる。

4証券からなるポートフォリオを X_1, X_2, X_3, X_4 として、制限条件 $X_1 \geq 0, X_2 \geq 0, X_3 \geq 0, X_4 \geq 0$

$$X_1 + X_2 + X_3 + X_4 = 1$$

の意味を考えてみよう。 X_1, X_2, X_3 を軸に $X_1 = 1, X_2 = 1, X_3 = 1$ のそれぞれの場合を考えるとすれば、 a, b, c, d を四隅とする四面体を描くことができる、いま $X_3 = 0, X_1 + X_2 + X_3 + X_4 = 1$ の条件にしたがうすべての点 (ポートフォリオ) からなる部分空間を S_{124} , $X_2 = 0, X_1 + X_2 + X_3 + X_4 = 1$ の条件にしたがうすべての点からなる部分空間を S_{134} とし、 $X_2 = 0, X_3 = 0, X_1 + X_2 + X_3 + X_4 = 1$ の条件を満足するすべての点の部分空間を S_{14} などすると、

$$S_{1234}, S_{123}, S_{124}, S_{134}, S_{234}, S_{12}, S_{13}, S_{14}, S_{23}$$

$$S_{24}, S_{34}, S_1, S_2, S_3, S_4$$

という15の部分空間を考えることができる。これらの部分空間に対応して3証券で考えたような臨

界集合

$$\ell_{234}, \ell_{123}, \dots, \ell_{12}, \dots, \ell_1$$

を考えることができる。 ℓ_{1234} は、ポートフォリオPが S_{1234} (すなわち $X_1 + X_2 + X_3 + X_4 = 1$) にあって、同一の期待収益をもつ S_{1234} のあらゆるポートフォリオの中でPが最小の分散をもつという条件を満足するPのあらゆる点の軌跡である。 ℓ_{123} については、ポートフォリオの収益

$E = \mu_1 X_1 + \mu_2 X_2 + \mu_3 X_3 + \mu_4 X_4$ において $\mu_1 = \mu_2 = \mu_3$ ならば、 ℓ_{123} は最小の分散をもつポートフォリオの点であり、 $\mu_1 \neq \mu_2 \neq \mu_3$ ならば、 ℓ_{123} は直線である。以下同じように解釈することができる。3証券の例でも述べたように、効率的ポートフォリオは臨界線上にあるが、臨界線上にあるすべてのポートフォリオは必ずしも効率的ではない。いま、すべてのポートフォリオの中で最小の分散をもつ実行可能なポートフォリオをXとし、臨界線(ℓ_{1234})に沿って期待収益を増加させ、他の臨界線(ℓ_{123})との交点で、交わった臨界線に乗り換え、さらに期待収益を増加させ、さらにつぎの臨界線(ℓ_{23})との交点で、新しい臨界線に乗り換えてゆけば最後に、最大の期待収益をもつXに到達することができる。このように臨界線の交点をコーナー・ポートフォリオといいコーナー・ポートフォリオを探しながら期待収益が最大であって分散のより小さなポートフォリオを探す方法をマーコヴィッツの臨界線法という。 n 種の証券からなるポートフォリオを分析する場合、 $n-1$ 次の空間が必要となり、図示できない。 n 種の効率的ポートフォリオを表わすためには、 $X-Y$ 平面において、 X 軸方向にポートフォリオの期待値 E 、 Y 他方向にその分散(または標準偏差)をとり($E-V$ 平面)、 E の最大値から E を固定して分散(あるいは標準偏差)の最小のポートフォリオの点を求め、次第に E を減少させ、それに対応する分散の点の軌跡を求めると、図11-1のような曲線($E-V$ 曲線)を描くことができる。ポートフォリオの問題というのは、個々の証券(あるいは資産)の期待収益と分散が与えられて、ポートフォリオの収益がより大きくそして、分散がより小さい、効率的なポートフォリオを求めるにはどうしたらよいかという問題である。

定 式 化

マーコヴィッツの模型において、先述したように、第 i 証券に投資される投資比率を、 X_i 、 $i=1, 2, \dots, n$ 、とすれば、 $X_1, X_2, \dots, X_n$ からなる1つのポートフォリオXをつぎのように定義することができる。

$$X' = [X_1, X_2, \dots, X_n] \quad (23)$$

は転置行列を表わす、ポートフォリオが実行可能であるためには、

$$\begin{aligned} a_{11}X_1 + a_{12}X_2 + \dots + a_{1n}X_n &= b_1 \\ a_{21}X_1 + a_{22}X_2 + \dots + a_{2n}X_n &= b_2 \\ \vdots &\vdots \\ a_{m1}X_1 + a_{m2}X_2 + \dots + a_{mn}X_n &= b_m \end{aligned} \quad (24)$$

$$X_1 \geq 0, X_2 \geq 0, \dots, X_n \geq 0 \quad (25)$$

の条件を満足せねばならない。ここで、

$$A = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{bmatrix} \quad \cdots \cdots \cdots (26)$$

$$b' = [b_1, b_2, \cdots, b_m] \quad \cdots \cdots \cdots (27)$$

$$0' = [\bar{0}, 0, \cdots, 0] \quad \cdots \cdots \cdots (28)$$

と定義すれば(24)および(25)は

$$AX = b \quad \cdots \cdots \cdots (29)$$

$$X \geq 0 \quad \cdots \cdots \cdots (30)$$

(29)式において、 $m=1$, $a_{11}=a_{12}=a_{13}=\cdots=a_{1n}=1$ とすれば、ポートフォリオの標準型を求めることができる。

ポートフォリオの期待収益 E は μ_i を第 i 証券の期待収益とすれば (11) 式から $E = \sum_{i=1}^n X_i \mu_i$ であるから、

$$\mu' = [\mu_1, \mu_2, \cdots, \mu_n] \quad \cdots \cdots \cdots (31)$$

と定義すれば (14) 式は、

$$E = \mu' X \quad \cdots \cdots \cdots (32)$$

σ_{ij} を第 i 証券および第 j 証券の期待収益の分散あるいは共分散とし、

$$C = \begin{bmatrix} \sigma_{11} & \cdots & \sigma_{1n} \\ \vdots & \ddots & \vdots \\ \sigma_{n1} & \cdots & \sigma_{nn} \end{bmatrix} \quad \cdots \cdots \cdots (33)$$

とすれば、分散共分散行列 C は対称行列であって、正の半正定値である。目的関数であるポートフォリオの収益の分散を表わす (17) 式

$$V = \sum_i \sum_j X_i X_j \sigma_{ij} \quad \text{は}$$

$$V = [X_1, X_2, \cdots, X_n] \begin{bmatrix} \sigma_{11} & \sigma_{12} & \cdots & \sigma_{1n} \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_{n1} & \sigma_{n2} & \cdots & \sigma_{nn} \end{bmatrix} \begin{bmatrix} X_1 \\ X_2 \\ \vdots \\ X_n \end{bmatrix} \quad \cdots \cdots \cdots (34)$$

と書きなおすことができるから

$$V = X' C X \quad \cdots \cdots \cdots (35)$$

である。このようにして (14) ~ (17) で述べたポートフォリオの問題は、

$$E = \mu' X$$

$$AX = b$$

$$X \geq 0$$

の制限条件の下で目的関数

$$V = X' C X$$

を最小にすることと要約することができる。

臨 界 線 法

効率的ポートフォリオの集合から、最適なポートフォリオを選ぶために、投資者は、彼の効用関数の形を決めなければならない。マーコヴィッツ⁽¹⁰⁾⁽¹¹⁾は最適ポートフォリオの選択について、特定の判定基準を主張してはいないが、投資者の効用関数と、その選択基準との関係についていくつかの研究を行なっている。

ポートフォリオ・セレクションを実行する人の期待効用 U は、あるポートフォリオからの収益 R の関数であると考えられる。

$$U = U(R) \quad \dots\dots\dots (36)$$

だが、一般にこの収益 R は確定的ではなく、ばらつくことが予想され、すなわち、確率分布をなしているから、効用関数 U は、各資産のもたらす平均収益 μ_i によって計算されるポートフォリオからの収益の期待値 E と、分散 V (すなわち危険)の関数である。

$$U(R) = f(E, V) \quad \dots\dots\dots (37)$$

ここで効用を最大にするには、マーコヴィッツの模型のように、すべての効率的ポートフォリオの集合を発生させるように考えるか、あるいは、後述のファラー(Farrar D.E)⁽²⁾のように、投資者に対して最適のポートフォリオを選び出すように計画するか、いずれかしなければならない。

マーコヴィッツは、効率的ポートフォリオの集合を計算するために、非模型計画法に関するクーン、タッカー(Kuhn, Tucker)^{(6),(7)}の証明を用いた。これはまた経済学の効用理論で、所得額 M が一定のとき、その範囲内で、どのように所得を配分し財貨を購入したならば、最大の効用を得ることができるという問題を解く場合に、普通行なわれている方法に似ている。すなわち効用理論ではこの問題をつぎのように解決する。いま x_i 、 p_i をそれぞれ第 i 財の数量と価格とすれば、貨幣所得額 M は、

$$M = \sum_{i=1}^n p_i x_i$$

であり、 U を効用とし、 p_i ($i=1, 2, \dots\dots\dots n$)が与えられるとすれば、所得制限方程式にラグランジ乗数 λ を掛け、効用関数と結合し、財貨数量 x_i と λ について、この結合された方程式

$$\Phi = U + \lambda (M - \sum_{i=1}^n p_i x_i)$$

を微分し、それぞれ($i=1, 2, \dots\dots\dots n$)を0とおくことによって、 U が最大の場合の x_i を求めることができる。この場合の目的関数は U であり、一次の所得制限方程式に制約された目的関数を最大にする問題であった。ポートフォリオ・セクションの目的関数は収益の分数からなる二次式である。この場合も前例と話しような考え方でとくことができる。すなわち、クーン、タッカーは、非線

型計算法の解法で、

$$AX = b$$

$$X \geq 0$$

の制限条件の下で目的関数

$$f = X' C X + \nu' X \dots\dots\dots (38)$$

を最小にするように X を選ぶことという問題をつぎのように解いている。上式において ν は、 ν_i を要素とする n 次の列ベクトルであり、 C 、 X 、 A は、すでに定義したベクトル、およびマトリックスである。ここで制限条件 AX にラグランジ乗数 2λ を掛けて、

$$\phi(X, \lambda) = f + 2\lambda' AX \dots\dots\dots (39)$$

をつくる。ラグランジ乗数 λ を 2 倍したのは、目的関数 f が 2 次式であり、 f を X_i について微分すると係数 2 が出てくるから、あらかじめ全体として整理できるような形にしたためである。このような意味で、目的関数の右辺第二項の ν' は $2\nu_i$ の要素からなると考えてよい。また λ は、第 K 番目の要素が λ_k であるようなラグランジ乗数の m 次のベクトルである。 $\phi(X, \lambda)$ を変数 X_i で微分すれば次式を得る。

$$\eta_i = \frac{1}{2} \frac{\partial \phi(X, \lambda)}{\partial X_i} = \sum_j \sigma_{ij} X_j + \nu_i + \sum_k \lambda_k a_{ki} \dots\dots\dots (40)$$

クーン、タッカーの証明においては、

$$AX = b$$

$$X \geq 0$$

であって、

$$\begin{aligned} \eta_i &\geq 0 & i=1, \dots\dots\dots n \\ \eta_i X_i &= 0 & i=1, \dots\dots\dots n \end{aligned} \dots\dots\dots (41)$$

のようなベクトル λ が存在する場合にのみ、 $AX = b \geq 0$ を条件に X は f を最小化することができる。条件 $\eta_i X_i = 0$ ということは、 $\eta_i > 0$ ならば $X_i = 0$ 、 $X_i > 0$ ならば $\eta_i = 0$ を意味する。この場合ラグランジ乗数は、不等式ではなく等式に関係しているから λ の記号には、いかなる制限もない。

この証明からマーコヴィッツは、ポートフォリオの問題を多変数関数の条件つき最大化の問題として、臨界線法を提案した。すなわち目的関数である効用関数 U をポートフォリオ E によって制約された関数と考え、

$$U = \lambda_E E - V \dots\dots\dots (42)$$

を最大にする問題としてとらえた。この目的関数 U において、ラグランジ乗数 λ_E はどういう意味をもっているであろうか。 λ_E は 0 から ∞ までの値をとりうるが、図 11-2 で示したように、期待収益と危険を表わす分数 V 、すなわち $E-V$ 平面上において、効率的ポートフォリオの曲線に接する直線の傾斜である。 λ_E が 0 ならば、直線 U は期待収益 E の軸に平行し、 λ_E が ∞ ならば、直線 U は分散 V の軸に平行する。またポートフォリオ・セレクションの実行者の行動として λ_E はどのような意味をも

つだろうか。すなわち λ_E が 0 に近ければポートフォリオ・セレクションの実行者の行動は期待収益が小さくても危険が少なくなることを望んでおり、危険を回避して安全性をねらう投資家 (investor) の行動であり、 λ_E が ∞ に近ければ、危険があっても期待収益の大きいことを望んでおり、危険をあまり顧みない投機家 (speculator) の行動である。このようにして λ_E は、投資家と投機家とを区別しているのであるが、マーコヴィッツの臨界線法では、 $\lambda_E = \infty$ 、すなわち危険をまったく恐れない投機家の行動から、ポートフォリオを解いて、臨界線法をたどって、次第に危険回避的な投資家の行動 ($\lambda_E \rightarrow 0$) に近づいていく。

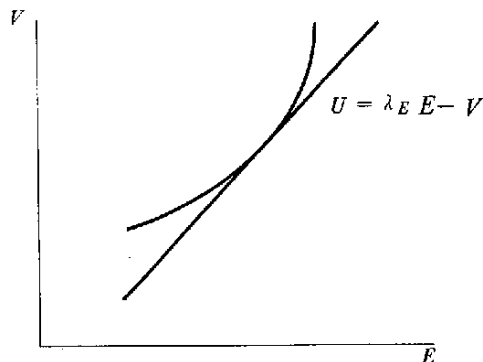


図 11-3

目的関数 $u = \lambda_E E - V$ において、 $\lambda_E = 0$ とすれば、投機家が E の最大だけをねらって危険については顧みないのであるから、 V には無関係であり、したがって、マーコヴィッツの模型では、

$$AX = b$$

$$X \geq 0$$

の制限条件の下に、目的関数 $E = \mu' X$ を最大にすること、すなわち、線型計画法とまったく同一の問題となる。

λ_E が ∞ でなければ、制限条件 $AX = b$ 、 $X \geq 0$ に制約された効用関数 $u = \lambda_E E - V$ すなわち

$$u = \lambda_E \mu' X - X' C X \quad \dots\dots\dots (43)$$

を最大にする問題となるから、クーン、タッカーにしたがって、

$$\phi(X, \lambda) = u + 2\lambda' (AX - b) \quad \dots\dots\dots (44)$$

をつくり、 X_i ($i=1, 2, \dots\dots\dots n$)、 λ_k ($k=1, 2, \dots\dots\dots m$)、について微分し、その微係数を 0 とおくことによって最大化の問題を求めることができる。すなわち、

$$\phi(X, \lambda) = \lambda_E \mu' X - X' C X + 2\lambda' [AX - b] \quad \dots\dots\dots (45)$$

を X_i と λ_E で微分して 0 とおけば、

$$\eta_i = \frac{1}{2} \frac{\partial \phi(X, \lambda)}{\partial X_i} = - \sum_{j=1}^n \sigma_{ij} X_j + \lambda_E \mu_i + \sum_{k=1}^m \lambda_k a_{ki} = 0 \quad \dots\dots\dots (46)$$

$$\eta_k = \frac{1}{2} \frac{\partial \phi(X, \lambda)}{\partial \lambda_k} = \sum_{i=1}^n a_{ki} X_i - b_k = 0$$

であり、また、 $i=1, 2, \dots\dots\dots n$ 、 $k=1, 2, \dots\dots\dots m$ であるから、この演算の結果は次のようになる。

$$\begin{bmatrix} \sigma_{11} & \dots & \sigma_{1n} & a_{11} & \dots & a_{m1} \\ \vdots & & \vdots & \vdots & & \vdots \\ \sigma_{n1} & \dots & \sigma_{nn} & a_{n1} & \dots & a_{mn} \\ \vdots & & \vdots & \vdots & & \vdots \\ a_{1i} & \dots & a_{in} & 0 & \dots & 0 \\ \vdots & & \vdots & \vdots & & \vdots \\ a_{mi} & \dots & a_{mn} & 0 & \dots & 0 \end{bmatrix} \begin{bmatrix} X_1 \\ \vdots \\ X_n \\ -\lambda_1 \\ \vdots \\ -\lambda_m \end{bmatrix} = \begin{bmatrix} 0 \\ \vdots \\ 0 \\ b_1 \\ \vdots \\ b_m \end{bmatrix} + \begin{bmatrix} \mu_1 \\ \vdots \\ \mu_n \\ 0 \\ \vdots \\ 0 \end{bmatrix} \lambda_E \quad (47)$$

ここで σ , A , A' , b , μ , をそれぞれマトリックス, およびベクトルとし, 新たに,

$$M = \begin{bmatrix} \sigma & A' \\ A & 0 \end{bmatrix}; \quad R = \begin{bmatrix} 0 \\ b \end{bmatrix}, \quad S = \begin{bmatrix} \mu \\ 0 \end{bmatrix} \quad (48)$$

を定義すれば, 上式は

$$M \begin{bmatrix} X \\ \lambda \end{bmatrix} = R + S \lambda_E \quad (49)$$

と書くことができる。この式において, λ_E が与えられれば, λ_E の水準によって効率的ポートフォリオを求めることができる。先にのべたようにマーコヴィッツ模型では $\lambda_E = \infty$ から解いて次第に λ_E を 0 に近づけながら効率的ポートフォリオ全体の集合を得るのであるが, λ_E を次第に変化させることによって, 全効率的ポートフォリオを求めることは実際上容易ではない。

だが, あるポートフォリオの中に採用されている (in) 資産によって効率的ポートフォリオの集合があり, このポートフォリオの中から 1 つの資産をはずし, 採用されていない (out) 資産の中から 1 つの資産を採用すると, 新たな効率的ポートフォリオの集合をつくることができる。前の効率的ポートフォリオの集合と新たな効率的ポートフォリオの集合の接点にコーナー・ポートフォリオが存在する。したがってコーナー・ポートフォリオにおいて資産を 1 つ入れ替えることができる。このように資産の入れ替えによってコーナー・ポートフォリオを求め, 各コーナー・ポートフォリオを結ぶ臨界線をたどってゆけば容易に効率的ポートフォリオの全集合を把握することができる。

マーコヴィッツはコーナー・ポートフォリオの性質を利用し, 前述したように, はじめに $\lambda_E = \infty$ において線型計画法によってコーナー・ポートフォリオを求め, この結果を利用して資産の入れ替えによりつぎのコーナー・ポートフォリオを求め λ_E の水準をつぎのように次第に小さくしている。

$$M \begin{bmatrix} X \\ \lambda \end{bmatrix} = R + S \lambda_E \quad (49)$$

において, あるポートフォリオに採用されていた資産 X_i は, この式に現われてこないから, M の第 i 行第 i 列の交叉している要素を 1 にして, 第 i 行第 i 列の i のほかの要素を 0 とおいたマトリックス

$\widetilde{M}$ をつくることができる。第 i 行第 i 列をユニット・クロスという。

たとえば、

$$M = \begin{pmatrix} \sigma_{11} & \sigma_{12} & \sigma_{13} & a_1 \\ \sigma_{21} & \sigma_{22} & \sigma_{23} & a_2 \\ \sigma_{31} & \sigma_{32} & \sigma_{33} & a_3 \\ a_1 & a_2 & a_3 & 0 \end{pmatrix}$$

の第2行第2列のユニット・クロス $\widetilde{M}$ は、

$$\widetilde{M} = \begin{pmatrix} \sigma_{11} & 0 & \sigma_{13} & a_1 \\ 0 & 1 & 0 & 0 \\ \sigma_{31} & 0 & \sigma_{33} & a_3 \\ a_1 & 0 & a_3 & 0 \end{pmatrix} \dots\dots\dots (50)$$

である。ここでユニット・クロスを用いたのは、 $\widetilde{M}$ の逆行列を求めるためである。 R についても同じように、採用されていない資産に対応する係数を0とおいた列ベクトルを $\widetilde{R}$ とすることができる。このようにすると、あるポートフォリオの臨界線は、

$$\widetilde{M} \begin{bmatrix} X \\ \lambda \end{bmatrix} = R + \widetilde{S} \lambda_E \dots\dots\dots (51)$$

であり、したがって、 λ_E に値を与えると、

$$\begin{bmatrix} X \\ \lambda \end{bmatrix} = \widetilde{M}^{-1} R + \widetilde{M}^{-1} \widetilde{S} \lambda_E \dots\dots\dots (52)$$

によって X 、 λ を解くことができる。だが、この式の、 $\widetilde{M}^{-1}$ はこのポートフォリオに採用されていない資産にあたる列と行の交叉する要素がユニット・クロスになっているから逆行列を求めた後はユニット・クロスをゼロにおきかえたゼロ・クロスにしなければならない。いま第 l 番目の臨画線上の X 、 λ を求めるとすれば $\widetilde{M}^{-1}$ をゼロ・クロスにした $N(l)$ によって書きかえれば、 $\widetilde{S} = S$ でよいから

$$\begin{bmatrix} X \\ \lambda \end{bmatrix} = N(l) R + N(l) S \lambda_E \dots\dots\dots (53)$$

である。つきに上(52)式を(46)式に代入して各行において $\lambda_E^{(j)}$ ($j=1, 2, \dots\dots\dots n$)をもとめ、臨界線は $\lambda_E^{(j)}$ のうちの最大のものがつぎの臨界線と交叉するから(49)式に $\lambda_E^{(j)}$ を代入して、この l 段階におけるポートフォリオを求める。このようにしてつぎの $l+1$ 段階の λ_E がきまって新しく採用される資産とはずされる資産がきまるから $l+1$ のポートフォリオをきめることができる。この手順を順次くりかえし λ_E の値を小さくしていくのである。

対角線法

マールコヴィッツは臨界線法でのべたように、すべての効率的ポートフォリオの集合を発生させるように考えたが、それを選択する場合の特定の判定基準を主張しなかった。ただいくつかの効用関数と選択基準の間の関係を研究したにすぎない。たとえば、 $U = c + aR + bR^2$ の形の効用関数を考えてみれば、投資家が、最大にしようとする期待効用は、 $E(U) = c + a\mu + b\mu^2 + b\sigma^2$ である。この場合、 U は効用、 R はあらゆるポートフォリオからの収益、 a, b, c は定数であって、投資家は、ポートフォリオを R の期待値 μ と標準偏差 σ に基づいてのみ選ぶことを意味する。またもし効用関数が、

$$U(R) = \begin{cases} 1, & R \geq d \\ 0, & R < d \end{cases} \quad (54)$$

であれば、投資家は、ある最低の受けいれることのできる収益率 d 以下になる R の確率を最小にするポートフォリオを選ぶように行動する。すなわち、

$$E(U) = P(R \geq d) = 1 - P(R < d) \quad (55)$$

投資家の最適ポートフォリオを選ぶために、効用関数と最適ポートフォリオとを結合したのはファラーであった。

ファラーは、効用関数 U を収益の期待値の回りで、デラー展開し、

$$U(R) = U(E) + U'(E)(R-E) + \frac{U''(E)}{2} \times (R-E)^2 + \dots \quad (56)$$

$U(R)$ の期待値をつぎのように求めている。すなわち

$$\begin{aligned} E\{U(R)\} &= U(R) = U(E) + U'(E)\varepsilon(R-E) \\ &+ \frac{U''}{2}(E)\varepsilon(R-E)^2 + \dots \quad (57) \end{aligned}$$

$\varepsilon(R-E)$ は平均値を原点とする二次の積率であるから $\varepsilon(R-E) = 0$ であり、 $\varepsilon(R-E)^2$ は二次の積率 V であるから、この結果を利用して、 $U(R)$ の期待値 $\bar{U}(R)$ 式の第 4 項以下を近似的に省略すれば、

$$\bar{U}(R) = U(E) + \frac{U''}{2}(E)V \quad (58)$$

ここで限界効用通減の法則を仮定すれば、

$$U'' < 0 \text{ であるから}$$

$$0 > \frac{U''(E)}{2}$$

$$A = -\frac{U''(E)}{2} \text{ とおけば、}$$

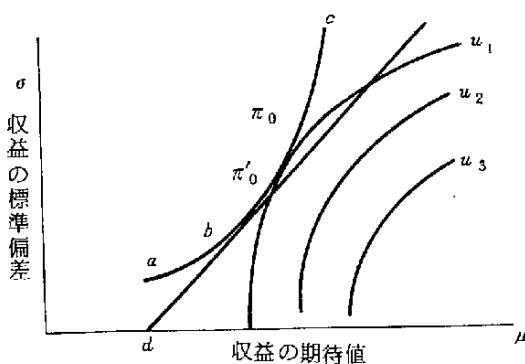


図 11-4

$$\bar{U}(R) = E - AV \dots\dots\dots (59)$$

ファラーは、このようにして求めた目的関数をマーコヴィッツの模型に導入することによって最適ポートフォリオを決定しようとした。形式的には、ファラーの模型は、

$$\sum X_i = 1$$

$$X_i \geq 0$$

$$\text{最大: } E - AV \dots\dots\dots (60)$$

である。この場合 A は危険回避 (risk aversion) の一定係数であり、そのほかの変数は今まで定義したとおりである。もし $E-V$ 平面上に、目的関数 $E-AV$ を書けば、右側にあればくるほどより高次の効用を表わす無数の無差別曲線の群を画くことができる。図11-3においては、この曲線群を U_1, U_2, U_3 として示してある。最適ポートフォリオ π_0 は、これらの無差別曲線の集合と、効率的ポートフォリオの集合との間の接点を見つけ出すことによって得られるのである。ファラーは、この模型を投資信託のポートフォリオ・セレクションの問題に応用し、危険回避係数 A を決定するために利用した。さらにシャープ (William E. Sharpe) ^{(13), (14), (15)} は、投資家の危険回避を経験的に証明するために1954-64年間の34種類のオープン・エンド投資信託の平均収益率を分析した。その結果、年間収益が非常に変動するよ

うな投資信託は、変動の少ない収益をもたらす投資信託よりも、平均的に、より大きな収益をもたらすことを発見し、この危険回避説を実証した。ファラーは、危険回避係数 A はすべて正であって、その値は、個人の資金の投資目的物に対する一般的な協定によってきまるということを発見している。

このような分析に対して、ロイ (A. D. Roy) は、現実の資産選択者は、収益と分散とをいかに組み合わせるかということよりも、むしろ、財政的災難を避けようとして行動すると考えた。すなわち、ポートフォリオからの収益 R が、ある危険水準 d 以下に下がる確率を最少にするように行動する投資家を考えると、もし、 R が平均的値 μ と分散 σ^2

をもって正規分布 $[N(\mu, \sigma)]$ しているとすれば、この投資者は、 $P(R \leq d)$ を最少にするように行動するから

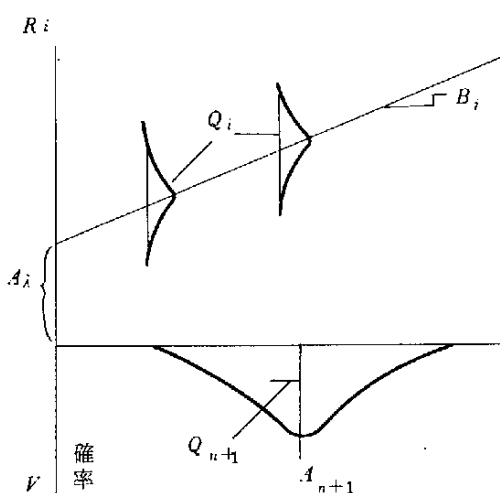


図11-5

$$\begin{aligned} P(R \leq d) &= P\left(\frac{R - \mu}{\sigma} < \frac{d - \mu}{\sigma}\right) = \Phi\left(\frac{d - \mu}{\sigma}\right) \\ &= 1 - \Phi\left(\frac{\mu - d}{\sigma}\right) \dots\dots\dots (61) \end{aligned}$$

ゆえに $P(R \leq d)$ は $\frac{\sigma}{\mu-d}$ が最少のとき、最少となるように、ポートフォリオを選ぶ、図的にいうと、最適ポートフォリオ π'_0 は、図11-3のように効率的ポートフォリオの集合と、 d からの延長線との接点を見出すことによって決定される。

対角線法

マーコヴィッツの模型を基礎にした臨界線法によって、ポートフォリオ分析を完成するための計算量は非常にばう大なものとなる。すなわち上述のような臨界線法では、つぎの計算手順をふまなければならないからである。

- 1) 効率的ポートフォリオに含まれる各資産の量と λ の値との関係で計算される。
- 2) (1)において計算された関係を用いてあるポートフォリオにとり入れられた資産から、1つの資産をはずしとり入れられていなかった資産を新たに組み入れるために、 λ の値を決定するために、各証券が決められる。
- 3) このポートフォリオにとり入れられるかあるいは迫出される1つの資産を決定するために、最高の λ のつぎの λ が決定される。このようにして、つぎのコーナー・ポートフォリオを決める。
- 4) 新しいコーナー・ポートフォリオを計算する。

以上の4つの段階の手順を繰り返し計算を行なうのであるが、

- 1) 資産の種類が第1段階、第2段階に影響をおよぼし、
- 2) コーナー・ポートフォリオの個数が、第1段階から第4段階までの繰り返し演算回数に影響し、
- 3) 第2段階では、各コーナー・ポートフォリオごとに分散・共分散行列の逆行列の計算を必要とし、その計算がきわめて複雑である。

また、これを電子計算機で計算する場合、記憶容量は、主としてこの分散・共分散行列の大きさによってきまる。すなわち、ポートフォリオの標準型の場合、もし N 種類の資産が分析の対象であれば、この行列は $1/2(N^2 + N)$ 個の要素をもっているのである。ポートフォリオ分析では非常に多くの比較計算をしなければならぬので、資産の種類が多くなれば極端に計算量が増加する。シャープは、このような比較計算量を軽減するために、マーコヴィッツの模型を発展させて、対角線模型(Diagonal Model)を開発した。計算量を軽減するだけの問題ならば、目的関数
$$V = \sum_{i=1}^n \sum_{j=1}^n X_i X_j \sigma_{ij}$$

において、すべての資産収益の変動が独立であると仮定し、共分散が

$$\sigma_{12} = \sigma_{13} = \dots = \sigma_{n-1n} = 0$$

ならば、

$$V = \sum_{i=1}^n X_i^2 \sigma_{ii}$$

であって、分散・共分散行列の対角行列のみが残り、計算量は非常に減少する。シャープはこの点に着目したが、この仮定はどの場合にも当てはめることができない。また、証券投資がある特定産業だ

けに備った場合、あるいは一般に産業活動の変動とともに、各証券の収益が同時に同方向に変動するような傾向がみとめられ、すなわち各証券間の相関が非常の場合、あるいは極限を考えてすべての相関係数が1を仮定できる場合には、目的関数が V ではなく、 $V^{1/2}$ となる。だが、このような仮定を、直接的に行なうことは一般に困難である。シャープは、直接的に証券相互間に高い相関関係が存在するということを仮定せず、しかも大部分のこのような相関関係をとらえ、計算量を減少させるための工夫をした。対角線模型の主な特長は、各種の証券（あるいは資産）の収益が、ある基本的要因との共通的な関係を通してのみ関係しているという仮定である。したがって、考え方としては、直接的に証券相互間の共分散を0とするのでもなければ、また直接的に証券相互間の相関係数を1にしているわけではない。どの証券の収益も、確率要素と、この単一の外部要因とによってのみ決定される。いま R_i をある証券からの収益とし、 A_i, B_i をパラメータ、 C_i を期待値0、分散 Q_i をもつ確率変数とし、 I をある指標の水準とすれば、

$$R_i = A_i + B_i I + C_i \quad \dots\dots\dots (62)$$

ここで指標 I は、各証券からの収益に最も莫大な単一の影響を与えると考えられる全体としての証券市場の水準、総国民生産物、ある価格指数あるいは、そのほかの要因である。 I の将来の水準は、部分的には、確率要素によって決定される。すなわち、

$$I = A_{n+1} + C_{n+1} \quad \dots\dots\dots (63)$$

この場合、 A_{n+1} はパラメータ、 C_{n+1} は、期待値0、分散 Q_{n+1} をもつ1つの確率変数である。

ここですべての i および j ($i \neq j$)に対して C_i と C_j の間の共分散が0であると仮定する。図11-4は、この模型を図示したものである。 A_i と B_i は、 I の水準に期待値 R_i を関係づける直線のパラメータであり Q_i は、期待値 R_i の分散である。（この分散はこの直線に沿って各点において同一の分散である）。 A_{n+1} は、 I の期待値と、 Q_{n+1} すなわちその期待値の周りの分散である。

対角線模型においては、証券アナリストから、つぎのような予測値を必要とする。

- 1) N 証券のおおのに対して A_i, B_i, Q_i の値
- 2) 指標 I に対して A_{n+1} と Q_{n+1} の値

このようにすると、証券アナリストからの推定値の数は、100種類の証券の分析では、5,150個から302個へ、200種の証券の分析では、2,003,000から6,002個へと非常に減少する。対角線模型のパラメータが定められると、標準的ポートフォリオ分析は必要なすべてのインプットが導かれる。その関係は、

$$E_i = A_i + B_i (A_{n+1}) \quad \dots\dots\dots (64)$$

$$V_i = (B_i)^2 (Q_{n+1}) + Q_i \quad \dots\dots\dots (65)$$

$$C = (B_i)(B_j)(Q_{n+1}) \quad \dots\dots\dots (66)$$

あるポートフォリオからの収益 R_p は、その中に含まれる証券からの収益の加重平均であって、

$$R_p = \sum_{i=1}^n X_i R_i \quad \dots\dots\dots (67)$$

対角線模型の仮定の下では (62) 式を (67) 式に代入して

$$R_p = \sum_{i=1}^N X_i (A_i + B_i I + C_i) = \sum_{i=1}^N X_i (A_i + C_i) + \left[\sum_{i=1}^N X_i B_i \right] I \dots\dots (68)$$

すなわち、 N 種の基本的証券への投資収益と、指標 I に対する投資収益との和である。

X_{n+1} を I の水準に対する R_p の加重平均反応性と定義すれば、

$$X_{n+1} \equiv \sum_{i=1}^N X_i B_i \dots\dots\dots (69)$$

で、この変数と (63) 式を (68) 式に代入すれば、

$$\begin{aligned} R_p &= \sum_{i=1}^N X_i (A_i + C_i) + X_{n+1} (A_{n+1} + C_{n+1}) \\ &= \sum_{i=1}^{N+1} X_i (A_i + C_i) \dots\dots\dots (70) \end{aligned}$$

あるポートフォリオの期待値は、このようにして

$$E = E(R_p) \dots\dots\dots (71)$$

$$E = \sum_{i=1}^{N+1} X_i A_i \dots\dots\dots (72)$$

であり、分散は、

$$V = \sum_{i=1}^{N+1} X_i^2 Q_i \dots\dots\dots (73)$$

この方程式は、 I の将来の値の期待値と分散を描写するために、なぜパラメータ A_{n+1} と Q_{n+1} を利用するかという理由を示している。このようにして、 N 証券を考える場合に、その分散・共分散行列は、上に定義された $N+1$ 番目の証券を含むことによって、対角線上にのみ、非常要素をもつ対角行列として表現される。このようにするとポートフォリオ分析の問題を解く場合にとくに、分散・共分散行列の逆行列をつくるときに臨界線法の第1の段階に必要な演算回数を非常に減少させることができる。対角線法によってポートフォリオの問題を定式化すれば

$$\begin{aligned} X_i &\geq 0 \quad (i=1, 2, \dots\dots\dots n), \quad \sum_{i=1}^N X_i = 1, \\ \sum_{i=1}^N X_i B_i &= X_{n+1} \end{aligned}$$

の制限条件の下で、目的関数

$$\lambda \sum_{i=1}^{N+1} X_i A_i - \sum_{i=1}^{N+1} X_i^2 Q_i \dots\dots\dots (74)$$

を最大にすることである。

決定理論への応用

ファラー、ロイおよびシャープの論文は、マーコヴィッツの模型を理解するのに非常に役に立つが、しかしながら、これらの論文では、マーコヴィッツの模型を決定理論と結びつけて考えていなかった。現代の統計的決定理論とマーコヴィッツの模型を結合し、発展させたのはマオとサーンドル (C. T. Mao, Carl Erik Sarndal)⁽⁵⁾である。マオとサーンドルは、ポートフォリオ・セクションの問題をきわめて狭義に技術的に解釈し、証券分析は科学ではなく *art* であるという立場からポートフォリオ・セクションの分析を展開する。これらの議論を要約すればつぎのようになるであろう。

証券から得られる収益の期待値と分散は、必然的に、将来の商況、金融財政政策等の維持および政治・経済の発展というようなある種の環境要素の予測にもとづいた主観的な推定である。通常これらの要因のすべてが証券投資に同時にしかも決定的な影響を与えるものではない。だが、ある一定の環境において、決定的な要因あるいはそのほかの要因がなんであっても、ポートフォリオ・セクションの模型に、これらの要因に対して証券アナリストがもっている見解を明確に組み入れられれば、より現実的なポートフォリオ・セクションの模型をつくることができるであろう。

いま経済の領域に不確実性の要素が入って、一般的な企業活動の水準をきめているような環境を考えてみよう。投資家は、確率 p をもって商況が有利 (favorable) であり、 q をもって不利 (unfavorable) であると信じているとする。ここでは2種類の自然の状態を想定したが、もちろん n 種類の自然の状態をこの2分類から容易に一般化することができる。さらに投資家は、 n 種の証券表のおのおのからの収益について、つぎのような確信 (beliefs) をもつ。すなわち、

有利な商況	不利な商況
μ_i	μ_i'
σ_{ii}	σ_{ii}'
σ_{ij}	σ_{ij}'

この場合、すべての記号は、すでに定義したものである。商況を一定に保つことは、証券収益のパラツキを減少せしめるが、まったく消滅させるわけではないから、したがって収益率は、このような場合でも確率変数と考えることができる。

さらに、自然の真の状態 θ (この場合は商況) に投資家が光を当て実験を行なうと仮定する。投資家はこの実験の結果を観察し、この結果を z とする。つぎに、自然の状態について投資家は基本確率を改訂する。これらの新しい確率は、実験の結果が与えられた場合の、商況に対する条件的確率である。基本確率 p および q と改訂した後の確率 p^* と q^* とを区別するために、前者を事前確率、後者を事後確率と呼べば、ベイズの定理を利用して、 z が与えられたときのそれぞれの θ_1, θ_2 の生起する確率 p^*, q^* をつぎのように求めることができる。

$$p^* = P(\theta = \theta_1 | z) = \frac{P(z | \theta_1) P(\theta_1)}{P(z | \theta_1) P(\theta_1) + P(z | \theta_2) P(\theta_2)} \dots\dots\dots (74)$$

$$q^* = P(\theta = \theta_2 | z) = \frac{P(z | \theta_2) P(\theta_2)}{P(z | \theta_1) P(\theta_1) + P(z | \theta_2) P(\theta_2)} \quad (75)$$

すなわち、 θ_1 、 θ_2 の生起する確率をそれぞれ $P(\theta_1)$ 、 $P(\theta_2)$ とし、 θ_1 と z をともに生起する確率および θ_2 と z をともに生起する確率をそれぞれ $P(\theta_1 \cap z)$ 、 $P(\theta_2 \cap z)$ とし、 θ_1 を条件にして z の生起する確率を $P(z | \theta_1)$ 、 θ_2 を条件にして z の生起する確率を $P(z | \theta_2)$ 、 z を条件にして θ_1 あるいは θ_2 の生起する確率を $p^* = P(\theta_1 | z)$ 、 $q^* = P(\theta_2 | z)$ とすれば、

$$P(\theta_1 \cap z) = P(\theta_1) \cdot P(z | \theta_1) \quad (76)$$

$$P(\theta_2 \cap z) = P(z) P(\theta_1 | z) \quad (77)$$

$$\therefore P(z | \theta_1) = \frac{P(z) \cdot P(\theta_1 | z)}{P(\theta_1)} \quad (78)$$

いまある実験の結果 θ_1 、 θ_2 のいずれかの事象が生起し、しかもこれらの事象が互いに排反事象であれば、

$$\theta_1 \cap \theta_2 = 1 \quad (i = j) \quad (79)$$

また、 θ_1 あるいは θ_2 のいずれかが生起するとすれば

$$\theta_1 \cup \theta_2 = 1 \quad (80)$$

$$z = (z \cap \theta_1) \cup (z \cap \theta_2) \quad (81)$$

$(z \cap \theta_1)$ と $(z \cap \theta_2)$ とは互いに排反事象だから

$$P(z) = P(z \cap \theta_1) + P(z \cap \theta_2) \quad (82)$$

また $P(z \cap \theta_i) = P(\theta_i) P(z | \theta_i)$ であるから

$$P(z) = P(\theta_1) P(z | \theta_1) + P(\theta_2) P(z | \theta_2) \quad (83)$$

$$\therefore P(\theta_1 | z) = \frac{P(\theta_1) P(z | \theta_1)}{P(z)} \quad (84)$$

$$= \frac{P(\theta_1) P(z | \theta_1)}{P(\theta_1) P(z | \theta_1) + P(\theta_2) P(z | \theta_2)}$$

よって(74)式が導かれた。同様に(75)式を求めることができる。

この場合 $P(z | \theta_1)$ と $P(z | \theta_2)$ とは、自然の状態がそれぞれ θ_1 および θ_2 の下における実験結果 z を観察する確率であり、 $P(\theta_1)$ と $P(\theta_2)$ は自然の状態の事前確率である。

つぎに、決定理論をポートフォリオ・セレクションに応用する場合に、ベイオフ・マトリックスが必要である。

1つ1つの可能なポートフォリオ π に対して、1つ1つの可能な自然の状態のもとに、その収益の期待値と分散とを計算し、これらの計算結果を、投資者の目的関数に代入することによって、すべての状態の下における1つ1つのポートフォリオのベイオフを決定することができる。

表 11-1

自然の 状 態	ポ ー ト フ ォ リ オ					
	π_1	π_2	π_3		π_m	
θ_1	$u(\pi_1, \theta_1)$	$u(\pi_2, \theta_1)$	$u(\pi_3, \theta_1)$		$u(\pi_m, \theta_1)$	
θ_2	$u(\pi_1, \theta_2)$	$u(\pi_2, \theta_2)$	$u(\pi_3, \theta_2)$		$u(\pi_m, \theta_2)$	

表 11-1において、 $u(\pi_1, \theta_1)$ は、自然の状態が θ_1 のときのポートフォリオ π_1 の効用を表わし、 $u(\pi_1, \theta_2)$ は自然の状態が θ_2 のときのポートフォリオ π_2 の効用を表わしている。この他にも同じように解釈される。この表には、投資家が、これ以上追求する必要のないいくつかのポートフォリオがある。いまこのポートフォリオ π_n と π_m をとりあげて、もしすべての状態に対して、 $u(\pi_n, \theta) \leq u(\pi_m, \theta)$ ならば、ポートフォリオ m はポートフォリオ n より優位であるという。どのような投資家も優位でないポートフォリオを選ぶことはない。優位なこれらのポートフォリオはこの容認できるものとして選択されれば、最適ポートフォリオはこの容認できる集合の中にある。だが投資家が、容認できるポートフォリオの中で選択を行なう場合に、どのような選択原理を適用するだろうか、マオとサードルは、ベイジアン戦略によって、ウェイトとして各種の可能な自然の状態の事後確率を用いて、ベイズの加重平均を最大にするポートフォリオを考え、マーコヴィッツの模型をつぎのように再定式化した。もちろんこの模型はポートフォリオの標準型であるが、一般化したポートフォリオの問題としてさらに拡張してもよい。

$$\sum_{i=1}^n X_i = 1$$

$$X_i \geq 0$$

$$\begin{aligned} \text{Max } p * \left(\sum_{i=1}^n X_i \mu_i - A \sum_{i=1}^n \sum_{j=1}^n X_i X_j \sigma_{ij} \right) \\ + q * \left(\sum_{i=1}^n X_i \mu'_i - A \sum_{i=1}^n \sum_{j=1}^n X_i X_j \sigma_{ij}' \right) \dots\dots\dots (85) \end{aligned}$$

この式に含まれる μ_i , μ'_i , σ_{ij} , σ_{ij}' は上述のとおりである。目的関数は、

$$\text{Max } \sum_{i=1}^n X_i (p * \mu_i + q * \mu'_i) - A \sum_{i=1}^n \sum_{j=1}^n X_i X_j (p * \sigma_{ij} + q * \sigma_{ij}') \dots\dots\dots (86)$$

と整理することができる。

ベイジアン解が、一般的に、標準ポートフォリオ・セレクションの問題に還元することは明かである。マオとサードルの模型は、ファラーの模型と異なり、目的関数の中に明確に2つの自然の状態を説明の便宜のため θ_1 と θ_2 とし、ポートフォリオ π は点 (u, u_1) によって表わされるとすれば、 $u_1 = u(\pi, \theta_1)$ と $u_2 = u(\pi, \theta_2)$ であるから、図 11-5 のようにこの模型を図示することができるであろう。

すべてこのような点の集合は、陰影をつけた領域にあり、またそのうちのいくつかのポートフォリ

オは優位でないから、したがって容認できない。
容認できるポートフォリオの集合は、円弧 efg の内側にあり、またこれは陰影をつけた領域内の上方にある。この模型にあってはベリオフの加重平均が目的関数であり、これを最大にすることによって、最適ポートフォリオを求めることができるから、図で説明すると、最適ポートフォリオは、 efg と p^*/q^* の傾斜をもつ直線との接点を見つけることによって決定される。

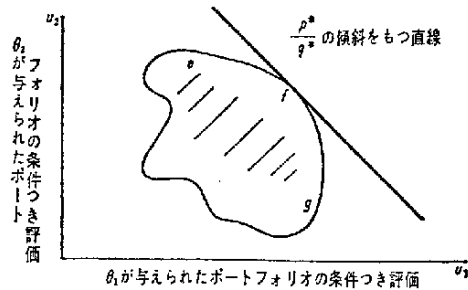


図 11-6

参考文献

1. Cord, Joel (1964): "A Method for Allocating Funds to Investment Projects When Returns are Subject to Uncertainty," *Management Science*, 10, 335~341.
2. Farrar, Donald E. (1962): *The Investment Decision Under Uncertainty*, Englewood Cliffs, New Jersey.
3. Dantzig, George B. (1963): *Linear Programming and Extensions*, Princeton Univ. Press, Princeton, New Jersey.
4. Dantzig, George B., Alex Orden and Philip Wolfe (1955): "The Generalized Simplex Method for Minimizing a Linear Form under Linear Inequality Constraints," *Pacific Journal of Mathematics*, 5, 183~195.
5. Houthakker, H. (1960): "The Capacity Method of Quadratic Programming," *Econometrica*, 28, 62~87.
6. Kuhn, H. W. and A. W. Tucker "Non Linear Programming" in *Proceedings of the Second Berkeley Symposium on Mathematical Statistics and Probability* edited by Jerzy Neyman, Univ. of Calif. Press, 481~492.
7. Künzi, Hans P. and Wilhelm Krelle (1966): *Nonlinear Programming Translated by Frank Levin*, Blaisdell Publishing Co.
8. Martin, A. D. (1955): "Mathematical Programming of Portfolio Selections," *Management Science*, 1, 152~166.
9. Markowitz, Harry, (1952): "Portfolio Selection," *Journal of Finance*, 7, 77~91.
10. Markowitz, Harry (1959): *Portfolio Selection, Efficient Diversification of Investments*, New York, N. Y.
11. Mao, James C. T. and Särndal, Carl Erik (1966): "A Decision Theory Approach to Portfolio Selection," *Management Science*, 12, 8.
12. Roy, A. D. (1952): "Safety First and Holding of Assets," *Econometrica*, 20, 431~449.
13. Sharpe, William F. (1963): "A Simplified Model for Portfolio Analysis," *Management Science*, 9, 277~293.
14. Sharpe, William F. (1963): "Communication to the Editor," *Management Science*, 9, 498.
15. Sharpe, William F. (1963): "Risk-Aversion in the Stock Market; Some Empirical Evidence," *Journal of Finance*, 20, 416~422.
16. Wolfe, Philip (1959): "The Simplex Method for Quadratic Programming," *Econometrica*, 27, 382~398.
17. 新沢雄一 "資産選択理論の解説" オペレーションズ・リサーチ Vol. 1. 12-1, 2, 3, 1967 日科技連
18. 新沢雄一 "汎用予測シュミレータ" オペレーションズ・リサーチ

—— 禁 無 断 転 載 ——

昭和 45 年 6 月 発 行

発行所 財団法人 日本情報処理開発センター

東京都港区芝公園21号地1番5

機 械 振 興 会 館 内

TEL (434) 8211 (代表)

印刷所 三協印刷株式会社

東京都渋谷区渋谷3-11-11

TEL (407) 7316

