

書評

A. E. Maxwell, *Analysing Qualitative Data*

London, Methuen & Co. 1961, 163 pp.

神田 祐一

本書は題名の示すとおり、変数が非連続的な性格をもつ場合のデータの分析方法を解説したものである。かかる定性的データの統計的分析方法としては、通常最も一般的なものとして、変量をもとにした二つの分類間の「独立性」あるいは「関連性」(association)の検定方法が用いられている。本書は、定性的なデータの分析をめぐって、このような検定の応用やこれに関連した種々のトピックを包括的に概説している。従って本書が分析の対象としているデータは、従来の経済分析においてはあまり分析の対象とされていないが、精神病理学や心理学あるいは社会調査の分野では、得られるデータが、質的要因によって分類された人員、世帯等の分布や、ある属性の所有非所有によって分類された標本である場合や、さらには何らかの形で定

義された scale にもとづいて分類された標本である場合が多い。そしてこれらの分野においては従来からさまざまな統計的分析手法が開拓されてきているが、本書ではこのような手法を、著者の関与している精神病理学及び心理学上の豊富なデータを計算例として駆使しながら、体系的に解説している。もともと本書は同じ著者による

*Experimental Design in Psychology and the Medical Sciences*, 1958.

の補稿としての意図をもつものであり(この書物では連続変量の解析のみが扱われている)、従って本書で使用されているデータが、我々の扱うデータとは異った分野のものであることは明らかであり、また展開されている統計手法が従来の経済分析にあまり用いられなかったものであるとしても、我々にとっても、本書の内容はかなり興味深いものである。その理由の第一は、本書がごく最近までに開拓されてきた検定の応用や、それに関連した手法をきわめて網羅的にまた体系的に展覧しているということ、次に分野こそ異なっているが本書で分析例として用いられているデータのかなりが、非実験的なデータであるという(我々の扱うデータとの)共通性をもっているということに在る。

本書は十一章よりなっているが、内容の構成上これを二つの部分にわけることができるであろう。最初の部分は、分割表における独立性あるいは関連性を詳論したI—VI章であり、後の部分は適合度検定や順位相関の検出、その他最後の item

analysis の初歩的解説までに至るVII—XI章である。以下各章ごとに内容を要約することにしよう。

## 二

I—VI章のうち、最初の二章は分割表における独立性の検定にかんし初歩的な解説が行われ、後の四章ではこの主題にかんし近時開拓されてきた目新しいトピックが紹介されている。まずI章では分類間の独立性あるいは関連性や $\chi^2$ 検定の意味内容を説明することから始まり $n \times n$ 分割表の分析が、またII章では自由度2以上の $m \times n$ 分割表の分析が行われている。この二章は、分析すべきデータの母集団からの抽出にあたっての注意を加えるなど、本書の序論ともいえるものであるが、帰無仮説——関連性ゼロ——の下での分割表の枠内期待値が5以下の場合 $\chi^2$ 値に偏りが生ずるという立言を、Yatesの補正值や、(I章では $2 \times 2$ 分割表、II章では $2 \times 3$ 分割表をもとにした) Fisherの exact test からの結果によって数値的に吟味しているのが、この二章の中心的内容をなすものであろう。その他I章においては分割表周辺度数が固定されており、枠内の値が正規分布に従う場合に、その枠内の母集団平均値の信頼区間を求める方法がCox (1938)<sup>(1)</sup>から紹介されている。またII章においては、自由度 $n$ が30以上で枠内期待度数がきわめて小さい場合における $\chi^2$ 値の平均値と分散を求める式がHaldane (1939)及びDawson (1934)<sup>(2)</sup>から与えられ、この条件の下で、自由度30以上のときに正規変量として通常用いられる

を使った場合の有意水準の偏りが吟味されている。

第1表：事業所数（産業×保有状況）

|     | 業      |        |        |     |
|-----|--------|--------|--------|-----|
|     | I      | II     | III    |     |
| 保有  | 58(57) | 78(83) | 93(82) | 222 |
| 非保有 | 19(20) | 34(29) | 16(27) | 76  |
|     | 77     | 112    | 109    | 298 |

( ) 内は帰無仮説の下での期待度数。

III章及びIV章では、分割表における $\chi^2$ 値の分割、換言すれば全自由度の分割がとりあげられている。まずIII章での論点を明らかにするため、次の例を挙げてみる。

例一 標本調査の結果、ある地域の事業所の産業(I、II、III)別自動車保有状況が第一表で与えられたものとする。この表によって産業と自動車保有との関連性を検定すれば

$$\chi^2 = 7.118 \quad (\text{自由度 } n=2)$$

となり、5%で有意である。

しかし表を一見すればわかるとおり、産業I、IIの間では保有率に有意な差がなく、上の結果が有意となったのは、産業IIIの保有率が異なるためであるとみられる。これを確認するために、上の $\chi^2$ 値を以下の成分に分割して、その各々を検定してみるとよい。すなわち

(i) 産業I、IIと産業IIIの保有率の差によってもたらされる変動部分…… $\chi^2(n_1=1)$

(ii) 産業ⅠとⅡの保有率の差によってもたらされる変動部分

$$\sum_{i=1}^m x_i^2 = x_1^2 + x_2^2 \quad (n_2=1)$$

この結果 $x_1^2$ が有意となり、 $x_2^2$ が有意とならなければ、さきの推論は確証されたことになる。

Ⅲ章では、この例のごとき $m \times n$ 分割表の $x^2$ 値の分割、すなわち一方の分類基準による $N$ 階層をいくつかのグループにまとめた場合のグループ間及びグループ内での関連性の検定方法が述べられ、ついで一般的な $m \times n$ 分割表における $x^2$ 値の分割方法が、 $m \times n$ の場合における分析例とともに、解説されている。

( $m \times n$ 分割表においては、 $x^2$ 値を(1)(1)(1)個の成分にわけることができる。すなわち、枠を適当にグループピングして(1)(1)(1)個の異った分割表を作成することができる。)

Ⅳ章では、 $x^2$ 値分割の今一つの方法として、分割表におけるトレンドの検出方法、すなわち $x^2$ 値を(線型)トレンドに依る部分とその残余部分とにわけける方法が概説されている。この方法は二つの分類基準のいずれもが大小関係によって表わすことのできる場合に可能である。(ただし $m \times n$ 分割表の場合、 $N$ 個の分類の方だけを大小関係で表わすことが可能なら、もう一方の分類の方は1及びゼロで表わすことができる。)分析方法の概略を述べると、まず各々の分類を数量化して $x$ 、 $y$ と表わし、 $x$ にたいする $y$ の線型回帰係数の推定値 $b_{yx}$ または $y$ にたいする $x$ のそれ $b_{xy}$ を求め、母集団回帰係数ゼロの仮説の下で分散 $V(b_{yx})$ 、 $V(b_{xy})$ を用いて $x^2$ 変量 $g_{yx}^2/V(b_{yx})$

または $g_{xy}^2/V(b_{xy})$ を求め、トレンドを検出する。結局、 $x^2_T$ は $x^2 = g_{yx}^2/V(b_{yx}) (= g_{xy}^2/V(b_{xy}))$  ( $n_1=1$ )、 $x^2 = x^2 - x^2_T$  ( $n_2=n-1$ )とに分割されることになる。

Ⅴ章では、同一母集団にかんして同じ分類基準による分割表の分析を数回くりかえした場合、これらを結合してより確定的な結論をみちびくための手法や、他の点では異った条件にある対象に、同一の分類基準で分割表を作成した場合の対象間の比較の問題がとりあげられている。前者の手法としては、同一の $m \times n$ 分割表をくりかえして得た場合、各々の $x^2$ 値の平方根の和を自由度合計で除した正規変量を利用して、一方の分類による二つの比率の大きさがどちらの方向に偏っているかを判定すること、またその際に分割表間で標本数が大幅に異っている場合の取扱いなどがある。総じて、本章の内容は目新しいものがなく、前章までの補注とみなしてよいであろう。

これまでに示されてきた方法は、二つの分類基準による分割表の分析にかんするものであったが、実際上は三つ以上の分類基準をとり入れる必要がしばしば生ずる。Ⅵ章では、このような場合の一つの分析方法として、Dyke-Patterson (1952)によって行われた $2^n$ 型分割表の分析に分散分析を応用する方法が解説されている。ここではこの章で扱われている問題の性質と接近方法の概略を、例一と同じ自動車保有にかんする例をもとにして、説明しよう。

例二 産業Ⅰ、Ⅱ、地域A、B、規模(i)、(ii)による三重層別によって、第二表のごとき事業所分布を得たものとする。これ

第2表: 事業所数及び自動車保有事業所数及び保有率(産業×地域×規模)

| 全 標 本  | $n_1 + \dots + n_8$      |                 |                  |               |                           |               |                  |                 |
|--------|--------------------------|-----------------|------------------|---------------|---------------------------|---------------|------------------|-----------------|
| 産 業    | I<br>$n_1 + \dots + n_4$ |                 |                  |               | II<br>$n_5 + \dots + n_8$ |               |                  |                 |
| 地 域    | A<br>$n_1 + n_2$         |                 | B<br>$n_3 + n_4$ |               | A<br>$n_5 + n_6$          |               | B<br>$n_7 + n_8$ |                 |
| 規 模    | (i)<br>$n_1$             | (ii)<br>$n_2$   | (i)<br>$n_3$     | (ii)<br>$n_4$ | (i)<br>$n_5$              | (ii)<br>$n_6$ | (i)<br>$n_7$     | (ii)<br>$n_8$   |
| 保有事業所数 | $m_1$                    | $m_2$           | $m_3$            | $m_4$         | $m_5$                     | $m_6$         | $m_7$            | $m_8$           |
| 保 有 率  | $p_1 = m_1/n_1$          | $p_2 = m_2/n_2$ | $p_3$            | $p_4$         | $p_5$                     | $p_6$         | $p_7$            | $p_8 = m_8/n_8$ |

ら事業所について自動車保有状況を調べたところ、第二表における事業所数 $n_i$ ( $i=1, \dots, 8$ )につき $m_i$ だけの保有事業所の存在ことが明らかとなった。これまでの分割表の分析の立場にたつて分析を行うとすれば、自動車保有非保有と三つの分類基準との関連性を検定することになるであろう。しかし、表を一見すれば、2<sup>3</sup>型分散分析を用いてより良好な分析が可能となることがわかる。すなわち産業、地域、規模の三要因によって

自動車保有率の変動を説明することが可能であり、また関連性の検定を行うよりも、より明確な結論の導出が可能となるように思われる。

本書で行われている計算の概略を以下に要約する。まず0/1 $p_i$ なる制約による分析上の困難を避けるため、通常行われているように $P_i$ をロジットに変換する。すなわち、

$$\text{logit } P_i = \log_e P_i / (1 - P_i) = Z_i'$$

を求め、さらに $2P_i - 1 = e^{Z_i'} / (1 + e^{Z_i'})$ とおき $Z$ 変量

$$Z = \frac{1}{2} \log_e \frac{(1+r)}{(1-r)}$$

を導出し、これを被説明変数として、2<sup>3</sup>型の分散分析によって各要因の主効果及び交互作用を検出する。次にこれら効果の推定値をもとにし、枠内標本数の不揃いを考慮した修正、あるいは加重推定の方法が行われる。これは各効果の最尤推定値を得るため行われる反復計算の手法であり、安定した推定値が得られるまでの反復計算が必要であるが、本書ではこの計算の第一ステップまでを、交互作用を無視した主効果の推定のみについて行っている。無論、計算の煩雑さを厭わなければ交互作用の推定も可能であるが、交互作用を無視したモデルの適合度検定によって交互作用の有意性如何をチェックすることもできる。

(1) Cox, D. R., 'The regression analysis for binary sequences', *J. R. S. S. Ser. B*, vol. 20, 1958, pp. 1—24.

(2) Haldane, J. B. S., 'The mean and variance of  $\chi^2$  when used as a test of homogeneity, when expectations are small', *Biometrika*, vol. 31, 1939, pp. 346—55.

Dawson, R. B., 'A simplified expression for the variance of the  $\chi^2$ -function on a contingency table,' *Biometrics*, vol 41, 1954, p. 280.

(c) たとえば  $2 \times 3$  型の場合にはもとの分割表を

|       |       |       |  |  |
|-------|-------|-------|--|--|
| $a_1$ | $a_2$ | $a_3$ |  |  |
| $b_1$ | $b_2$ | $b_3$ |  |  |
| $n_1$ | $n_2$ | $n_3$ |  |  |

の二通りの分割表を作成することが出来る。またこれら各々の  $\chi^2$  値を求める公式は次のように与えられる。

$$(i) \chi_1^2 = \frac{N(b_3(a_1+a_2)-a_3(b_1+b_2))^2}{ABn_3(n_1+n_2)}$$

$$(ii) \chi_2^2 = \frac{N^2(a_1b_2-a_2b_1)^2}{ABn_1n_2(n_1+n_2)}$$

これを例一の場合に適用すれば

$$(i) \chi_1^2 = 6.349 (2.5\% \text{ 有意}) (n_1=1)$$

$$(ii) \chi_2^2 = .769 (N.S.) (n_2=1)$$

$$\text{計 } \chi_T^2 = 7.118 (5\% \text{ 有意}) (n=2)$$

となり、予想した結果が確認される。なお、例一及び後の例二はいずれも筆者の考えた仮設例であり、本書の計算例が精神病理学及び心理学に限られていて、本稿にそのまま紹介することが適当でない」と判断したことによることを、お断りしておく。

(4) Dyke, G.V. & Petterson, H.D., 'Analysis of factorial arrangement when the data are proportions,'

*Biometrics*, vol. 8, 1952, pp. 1-12.

### III

I—VII 節では、専ら分割表における関連性の検定に視点が向けられていたが、VIII 章では適合度の検定が取扱われる。ここでは、まず抽出されたデータがある基準に従って分類した場合、このデータと何らかの関連をもつ母集団の（同一分類基準による）期待度数が明らかなきに、抽出データがその母集団から抽出されたものとみなしてよいかを検定する方法が述べられ、次に二項分布、ポワソン分布など理論的分布にかんする適合度検定の方法が解説されている。

VIII 章は、Sperman の  $\rho$  や Kendall の  $\tau$  をはじめとする順位相関の手法の解説にあてられている。この二つの方法の他には、 $n$  個の同一対象にたいして  $K$  組の順位を与えたとき、この  $K$  通りの順位がどの程度一致しているかを測る Kendall の一致係数 (coefficient of concordance)  $W$  がとりあげられている。また順位づけを行う対象の数が大きい場合に、これをいくつかの不完備型ブロックにわけて、各ブロック内での標本に対して同一方法で順位をつけておき、この順位がブロック間でランダムであるかどうかを判定する手法が詳細に説明されている。この方法は、順位相関法にたいする不完備型ブロック計画の一つの応用であるが、本章冒頭に述べられているごとく、標本数の大なる場合に順位づけにともなう困難を回避するための工夫として貴重なものである。

IV章では、系列データが1またはゼロいずれかの値をとる場合の、非パラメトリック検定の手法がいくつか紹介されている。これらは主として学習実験からのデータにおいて有効な分析ツールとなる。

最後の二章では、X章がいわゆる分類問題——複合母集団の要素を、何らかの最適な方法によって、その母集団内での分類のいずれに属するかを決める問題——を扱っており、XI章ではitem analysisの初歩的な解説がなされている。前者は決定理論の応用であり、後者はattitude measurementの問題であって、いずれも興味深いものであるが、我々の分野での実際の応用という面からは、いまだほど遠いものがある。

#### 四

本書はMaurice Bartlettの編集によるMethuen's Monographs on Applied Probability and Statistics中の一冊であり、このmonographは最近までに開拓されてきた確率論及び統計学の各分野の業績を拡張範囲にわたって手ぎわよく展望することを目的としている。従って本書がきわめて網羅的に主題の整理、要約を行っていることは当然であるとしても、本書の題名が示すごとく、一つ一つデータを例示して具体的な計算手順を明らかにすることに力点をおいているのが、本書の特徴であるように思われる。このことは、一つには本書の主題を解明していくのに最善の方法であるということにもよろうし、また本書が同じ著者による書物の補稿として、精神病理学における

統計手法の応用を解説するという意図をもつためであろう。半面、本書では、検定に用うべき統計量の導出方法が何ら明らかにされていないかったり、計算手順の意味内容についてふれていない場合がきわめて多い。この種の書物につきまとう一つの欠点かと思われる。しかしこの点に関しては、引用文献を各章ごとに掲げるといふ配慮はなされている。

最後に、本書で展開されている統計手法が、我々の分野における分析において、どの程度有効であるかという問題が残る。この種の手法は、一、二のものを除き、従来の経済分析にはほとんど用いられていない。一つには、経済分析で扱うデータが連続的性格のものであるということにもよるし、また本書の如く検定のみが力点がおかれるのであれば、経済分析の場合、さして有意義ではなく、力点は推定論の方に在る。ただ、本書のVI章で概説されている方法は、2<sup>n</sup>型分割表の取扱いの1方法という観点をはなれても、我々の分析にとってかなり参考となるものがある。この章の方法が、経済データにたいする分散分析の応用につきまとう一つの困難を回避する道への示唆を与えるからである。また、我々の関連分野、たとえば市場調査においては、主たる説明変数となる所得データなどが得られない場合、他のデータを適当にクロスすることによって、いわゆるfact findingsのみを行う目的のためには、本書の内容がかなり参考となるかも知れない。たとえば、例一、二のごとき適用例がそれを示している。