

推定と予測における直観的な論点の吟味

— 回帰分析の場合 —

竹 内 啓

1 問題の所在

経済分析の統計的方法の体系としての計量経済学の方法は、同時方程式体系の母数推定論を中心として、きわめて精緻なものとして完成されていることは、いまさいうまでもない。

しかしながら、現実の経済分析においては、精密な確率的モデルを前提とする厳密な方法を、現実のデータや、それが反映している経済構造のあり方に照らして見るとき、あまりにも sophisticated であって、形式的な厳密性が非現実的な基礎の上に立っているという感じを抱かせることは、決して少なくない。

勿論私は、計量経済学的方法を根本から否定しようとする考え方¹⁾には同調できないし、また想定されたモデルと現実との間に多かれ少なかれギャップが存在することは、モデルが一種の“理想化された関係の体系”を表現するものである以上、いわば最初から前提されているのであって、そのことは直ちにモデル分析の有効性を損うものではないと思っている。しかしそれにもかかわらず、モデルの極度の精密さが、果してその有効性を高める上に真に必要なものであるかという疑問は残るであろう。

もう一つの問題は、精密なモデル化が困難であることが明白な場合、しかもそこにおいて何らかの定量的分析が要求される場合も少なくないことである。このような問題を、計量経済学の対象の外にあるものとして、それを直観的な、純然たるアド・ホックな方法に委ねてしまえば、計量経済学

の体系のいわば理論的純粋性は保たれるかもしれないが、しかしそれは現実の経済分析の要求によく応えるものではなくなるであろう。

ここでいうまでもないことではあるが、モデルと現実とのギャップは、モデルの仮定をゆるめ、ただモデルを“一般化”することによって埋めることはできないことを強調しておきたい。そのようにして表面的にモデルを現実に対応させることができて、むやみに一般的なモデルは分析の道具としての有効性を失ってしまう。

ここで問題は、ある程度一般的な条件の下で、ある程度の有効性を保つような方法を見出すこと、すなわち広い意味で robust な方法を考えることである。robust な方法ということばはふつう確率分布の型が“標準型”(多くの場合正規分布が想定される)から離れた場合にも、効率を失わない方法という意味に用いられるが、ここではより一般に標準的な状況からのいろいろな方向へのズレに対して考えねばならない。

robust な方法を求める際の問題状況の特徴はかつて Allan Birnbaum が指摘したことがあるが²⁾、その fuzzyness、すなわちその限界が不明確であるという点にある。すなわち標準的な状況からの乖離の可能性について、その大体の方向と、程度とはわかっており、従って考えられる程度のズレに対してあまり効率を失わないような方法を求めることが問題になるが、その方向や限界について、明確な定義はできないというのが、その標準的な場合なのである。問題の要点はむしろ漠然としたものであっても、ある程度の方向や範囲が想定されているという点にあり、従って例えば分布型についての仮定を全く置かないノンパラメト

1) 計量経済学に対するマルクス経済学の立場からの批判はすでに1955年に出た広田純・山田耕之介「計量経済学批判」『近代経済学批判講座』第3巻(日本評論社)に基本的な論点は展開されている。

2) personal discussion による。

リック法とは、異なる考え方に立って問題が論ぜられなければならない。

実はわれわれが現実のデータ解析にあたって事前に持っている情報は、多かれ少かれ fuzzy なものである。われわれがモデルを構成するにあたって、その情報の一部は確固としたものとみなしてモデルの中に組み入れ、また一部は不確かなものとして一応除外し、またもしベイズ的方法を用いるならば一部は母数の事前分布という形で定式化する。このようにして fuzzy なものに一応明確な区切りをつけたものがモデルなものである。もしこのことが適切に行い得ないとされる場合には、モデルを構成することなく、fuzzy な情報をそのまま用いて直観的な方法が用いられることになる。

しかしこの2つの場合の区別は絶対的ではない。どのような場合においても、モデルは全く疑問の余地がないということはないし、またモデルの中に組みこまれなかった情報が全くないということもあり得ない。逆にどのように fuzzy な情報の下であっても、全くモデル化が不可能ということもない。多くの場合その差は程度の問題にすぎない。従って一見あいまいな情報をなるべく数学的に客観的な論理の中にとり込むことが必要である。

このような状況を理論的にとり扱うには、モデルについて一種の二重構造を想定しなければならないことが多い。すなわち“真の構造”についてのある程度一般的な、しかし精密には規定されていない前提と、それについて仮定される“モデル”すなわち必ずしも正確ではないが、一連の精密な仮定とである。そうして後者から導びかれる方法が、より一般的な前提の下でどのような性質を持つか、そうしてそのことと基準として、モデル自体の適切さをどのようにはかるかを考えねばならない。この問題の一つの解は赤池氏の情報量基準によって与えられる³⁾。それはあるモデルの適切さを

AIC = -対数尤度 + 推定される母数の数
で計るもので、これを最小にするモデルが、最も望ましいとされる。それはモデルのデータに対す

る「あてはまり」のよさと、母数の数によって表わされるモデルの複雑さの間のバランスをはかろうとするものである。この基準の合理化、或いはそれ以上立ち入った議論にはここではふれないが、ここで“モデル”と真の分布との違いが最初から考慮に入れられていることに注目したい。

以下では経済分析においても最もふつうに使われる手法である回帰分析について、問題を論じよう。

2 回帰分析における Mallows の C_p 統計量

いま、ある経済変数 Y の変化を、いくつかの説明変数を用いて“説明”することを考えよう。

いま被説明変数 Y 、および考えられる説明変数の候補 $x_1, x_2, \dots$ についての時系列データ

$$Y_t, x_{1t}, x_{2t}, \dots, t=1, 2, \dots, T$$

が与えられたものとしよう。

分析の目的は、これらのデータを用いて、適当な説明変数をえらんで推定式

$$\hat{Y} = \alpha + \beta_1 x_{j1} + \dots + \beta_p x_{jp}$$

を求めることである。

このときふつうに行われる方法は、まづいくつかのモデル

$M_i: Y_t = \alpha_i + \beta_{i1} x_{1t} + \dots + \beta_{ip} x_{ipt} + u_t \quad t=1 \dots T$
 $i=1, 2, \dots$ を考え、その中から最も適当なものをえらび、その母数を推定して、 Y の推定式を定めるという方法である。

この方法は実は論理的な難点をふくんでいる。というのはそれぞれのモデルの下での母数の推定は(当然のことながら)そのモデルが正しいという前提の下で行われる。しかし実はいくつかのモデルを比較する場合には、正しいモデルは2つ以上はあり得ないはずである。更に厳密に言えば真に“正しい”モデルは一つも存在せず、ただより真に近いモデルとそうでないモデルとが存在するだけであると考えねばならないであろう。そこで現実には多くのモデルについての計算結果をいろいろ比較して、適当と思われるものをえらぶことになる。そうしてモデルをえらんだ後には、それが多くの中からえらばれたものであるということとは不問にして、あたかもそれが最初から仮定された

3) 『数理科学』特集情報量基準 1976年3月号。

ものであるかのように見なされるのがふつうである。

このようなモデル選択においては、決定係数を基準とするほか、係数推定値の符号や大きさの尤もらしさ等から直観的に判断がなされるが、その論理は必ずしも明確でないことが多い。そこでこの点をもう少し論理的に扱うことにしよう。そのためにまず一般的な構造として、

$$Y_t = \eta_t + u_t \quad E(u_t) = 0 \quad t=1 \cdots T \quad (1)$$

と想定し、次に η_t を

$$\eta_t = \alpha + \beta_{11}x_{1t} + \cdots + \beta_{ip}x_{ipt} + \zeta_{it} \quad (2)$$

$$\sum_i x_{ijt}\zeta_{it} = 0 \quad j=1 \cdots p$$

と表現することにすれば、モデル M_i の“偏り”は

$$b_i^2 = \sum \zeta_{it}^2$$

で表わすことができる。この値の小さいモデルが望ましいモデルである。

ここで $E(u_t^2) = \sigma^2$ $E(u_t u_{t'}) = 0 \quad t \neq t'$ (3) とし、形式的にモデル M_i をあてはめて推定を行えば、 $\beta_{11} \cdots \beta_{ip}$ の最小 2 乗推定量 $\hat{\beta}_{11} \cdots \hat{\beta}_{ip}$ は、やはり不偏推定量になり、またその分散共分散行列は、偏り ζ_i が存在しない場合と同じになる。また、

$$\hat{Y}_{it} = \hat{\alpha} + \hat{\beta}_{11}x_{1it} + \cdots + \hat{\beta}_{ip}x_{ipt}$$

とおくと、

$$\sum E(Y_t - \hat{Y}_{it})^2 = (T-p-1)\sigma^2 + b_i^2$$

となるから、 $Q_i = \sum (Y_t - \hat{Y}_{it})^2$ とするとき

$$\hat{\sigma}_i^2 = Q_i / (T-p-1)$$

は過大な推定量になる。

ところで、一般に説明変数を増せば、偏り b_i^2 は減少するであろう。しかし逆にそれによって Y_t の“理論値” $\hat{Y}_{it}$ の分散は大きくなるであろう。そこでこの 2 つの傾向のつり合いを考えねばならない。

Mallows はこの点に関連して次のような概念を導入した⁴⁾。いま Y_t^0 を Y_t と同じ期待値を持ち、かつ $Y_1 \cdots Y_T$ とは独立な分布に従う変数と

する。すなわち、

$$Y_t^0 = \eta_t + u_t' \quad t=1, \cdots, T$$

であって u_t' $t'=1 \cdots T$ は u_t $t=1 \cdots T$ とは互いに独立であると仮定する。

そうして Y_t^0 の値を $\hat{Y}_{it}$ を用いて予測するときの平均 2 乗誤差を考えると、その和は、

$$W_i = \sum E(Y_t^0 - \hat{Y}_{it})^2 = (n+p+1)\sigma^2 + b_i^2$$

となる。従って、

$$W_i = E(Q_i) + 2(p+1)\sigma^2$$

となる。そこで W_i が小さいようなモデルが望ましいと考えることができる。そこで W_i の推定量として

$$\hat{W}_i = Q_i + 2(p+1)\hat{\sigma}^2$$

或いは、 $C_p = Q_i / \hat{\sigma}^2 + 2(p+1)$

を用いる。 $\hat{\sigma}^2$ は何らかの形で求めた σ^2 の推定量である。 C_p は Mallows の C_p 統計量といわれている。

いくつかのモデル M_i $i=1, 2, \cdots$ の中から最もよいモデルをえらぼうとすると、それぞれについて C_p 統計量を計算して、それが最小になるようなモデルをえらぶことが考えられる。

この方式については、いろいろな技術的な吟味もなされており、また母数の数 $p+1$ の係数 2 が最も適当であるかどうかについても議論の余地がある⁵⁾。しかし少なくとも第一次的近似としては、この方法はモデル選択に一つの合理的な基準を与えるといつてよい。

ここで一方における説明変数の組が、他方のそれの部分集合となっているような 2 つのモデル、

$$Y = \alpha + \beta_1 x_1 + \cdots + \beta_p x_p + u$$

および

$$Y = \alpha + \beta_1 x_1 + \cdots + \beta_p x_p + \beta_{p+1} x_{p+1} + \cdots + \beta_q x_q + u$$

の比較において、問題を仮説

$$\beta_{q+1} = \cdots = \beta_q = 0$$

の検定という観点から考えることもできよう。このとき、ふつうに用いられる F 統計量は、

$$F = \frac{(Q_p - Q_q)}{(q-p)\hat{\sigma}^2} \quad \hat{\sigma}^2 = \frac{Q_q}{n-q-1}$$

と表わされるから、Mallows の C_p 統計量を用い

4) Mallows, C. L., "Some comments on C_p ," *Technometrics*, 15(1973). ただし彼は C_p 統計量をこれより数年早く発表しており、この前からそれを用いた他の人の論文が現われている。

5) 竹内啓「回帰分析における変数選択基準について」『経済学論集』42 巻 2 号, 1976 年。

る方法は(σ^2 を上記の値に等しくおくと) $F \geq 2$ に
応じて仮説をすてたり, 受容したりすることに対
応する。

2 という限界は F 検定のふつうに用いられる
水準に対応する棄却限界よりも小さい。このよう
ないわゆる検定—推定型の多段階推測(或いは推
測過程)の問題において, 仮説の棄却限界はふつ
うの 5% 水準よりも小さく, 10~30% 水準に対
応する値にとらねばならないことは, 以前から知
られていたことであり, また実際 F 検定の限界
を 2 に定めるということも, これまで経験的に採
用されて来たところであるが, 上の議論はその一
つの根拠を与えている。

ただしここで注意すべきことは, Mallows の考
えを適用する場合, 候補となるモデルすべてにつ
いて同時に比較を行わねばならないこと, また
 σ^2 の推定量 $\hat{\sigma}^2$ はほぼ偏りのないものであること,
すなわち十分多数の説明変数を用いたモデルの残
差から計算されねばならないことである。従って
説明変数を逐次的に 1 つずつ増したり減らしたり
する場合に, F 検定を機械的に適用し, 2 という
値を基準として用いることは適当でない。特に一
つずつ説明変数をふやして行く, いわゆる前進法
は説明変数の少なすぎるモデルで止まってしまう
危険性がある⁶⁾。

3 回帰式の意味の再吟味

しかしこのような理論にもなお疑問の余地が残
る。

先にのべた構造においては, Y_i の期待値 η_i と,
説明変数 $x_{1i}, x_{2i}, \dots$ とを結びつける式(2)は, 全
く形式的に定義されており, そこには特に“因果
関係”は仮定されていないし, 説明変数 x_i の係
数 β_i が, その変数の Y におよぼす影響の強さを
直接表現しているわけでもない。このことは, 説
明変数の係数の定義が, 説明変数と“偏り”とが
直交するという条件(2)に依存していることから
明らかである。

従って係数 β_i の値はそのモデルにおいて当

変数 x_i のほかにどのような変数がふくまれるか
によって変わるのである。

もし特定の変数の直接の影響の大きさををはか
ることが必要であり, 従ってその変数の係数 β_i が
独立の意味を持つならば, 直交条件(2)はなり立
つとは限らない。従ってまた推定量 $\hat{\beta}_i$ は偏りを
持つことになる。しかしその偏りはまたモデルに
ふくまれる他の変数によって影響されるから, そ
の意味を解釈することが困難になる。

β_i の意味に立ち入ることをさけて, Y の“予
測”のみを問題としたのは, 論理的な困難をさけ
る巧妙な方法ということができるかもしれないが,
しかしなお問題は残されている。というのは, 予
測される時点において, 説明変数 x_{ii} の値が, デ
ータにおける値と全く一致するという仮定の不自
然さは, 便宜的なものとして容認するとしても,
説明変数の値の組が一致すれば, その期待値も一
致すると考えてよいか否かには疑問があるからで
ある。上記の議論では, 予測時点において, すべ
ての説明変数の値がデータ時点の値に一致すると
したのであるが, このような仮定の下では Y の
期待値もまた両者において等しくなると想定して
もよいかもしれない。しかし, 一度特定のモデル
を選択してしまえば, その式はその中に現われな
かった説明変数の値は考慮せずに用いられるであ
ろう。その式をあてはめるための条件として, そ
の中に現われない説明変数の値を前提することは,
不自然でもあり, またできる限り簡単なモデルを
えらぶという要求に反することになるであろう。

従って一部の説明変数の値を定めれば, 他の説
明変数の値が一義的に決定されてしまうような場
合(例えば, 多項式モデルにおいて, 各次数の項,
或いはそれと同等な直交多項式のそれぞれを説明
変数とみなすとすれば, それらの間には関数関係
が成立する)を除けば, 上のような接近法はなお
不十分であるといわねばならない。

このように考えると, 結局説明変数の間には何
らかの形で一定の関係が存在し, そうしてその関
係は予測期においてもなり立つと想定しなければ,
予測に関連して変数選択の問題を考えることがで
きないことがわかる。実際極端な場合として, デ

6) 5)に同じ。

ータ期間においては説明変数の間に完全な線型関係が成立する一方、予測期間においてその関係がもはや成り立たず、しかも Y はこれらすべての変数によって影響されるとすれば、適切な予測量を求めることはそもそも最初から不可能になる。

4 説明変数の種類区分

ここで更に立ち入って論ずるためには説明変数の種類をわけて考えることが必要になる。まず基本的に制御可能な変数と制御不能な変数との区別が重要である⁷⁾。制御可能な変数の中では、自由に動かし得る変数の数、すなわちその自由度は一般に最初から与えられている。そうしてまたそれらの変数の一部をモデルから除くとしても、それはそれらの変数を動かしても Y の値に影響を与えないであろうということを想定するだけで、それらを全く考慮しないわけにはいかない。それらの値もとにかく特定の値に定めておかねばならない。 Y の値を特定の方向に動かそうとして制御変数の値を変えたとするならば、データに照らしてあまり影響がないように思われる制御変数の値は、データにおける平均値に等しく取っておくのが最も安全であろう。これに対してもしそのような変数を何らかの理由によってデータ期間における値から大きく動かさねばならないとき、その結果を正確に予測することは非常に困難になる。

一方説明変数が制御不能であるときには、回帰式にもとづいて予測を行うということは、説明変数の値を知ることを前提として Y の値を予測することを意味する。そうしてモデルの中にふくまれる変数を制限するということは、説明変数の一部を知って予測を行おうとすることにほかならない。しかしその場合には他の変数がどうなっているかについての情報は得られないであろう。従って予測に用いられる変数以外の変数が大きく変動して、それによって Y の値に予期しない変動が生ずるというようなことはしないものとしなければならない。

$x_1, x_2, \dots$ が制御可能でない変数の場合、それら

を更に Y に対して明確な因果関係を持つと思われる変数と、 Y に影響を与えるような状況を直接或いは間接に表現していると思われる変数とに分けて考える必要がある。前者についてはその変動が Y に影響をおよぼさないということが、すなわち因果関係があるという想定が誤りであったことが明確に示されない限り、モデルから除くべきではない。後者については、実はそれは直接 Y に影響をおよぼすのではなく、 Y に影響をおよぼすような、直接には観測されない変数によって影響され、それを反映しているものとみなすことができるであろう。

以上のことをまとめると、モデルは一般に次のように表わされる。

$$E(Y_t) = \eta_t = \alpha + \beta_1 x_{1t} + \dots + \beta_r x_{rt} \\ + \beta_{r+1} \bar{x}_{1t} + \dots + \beta_{r+s} \bar{x}_{st} \\ + \gamma_1 \xi_{1t} + \dots + \gamma_k \xi_{kt}$$

ここに $x_1, \dots, x_r$ は制御可能な変数、 $\bar{x}_1, \dots, \bar{x}_s$ は因果関係を表わす変数、 $\xi_1, \dots, \xi_k$ は未知の構造を表わす変数である。

そうして $\xi_1, \dots, \xi_k$ に対しては、他の“説明変数” $x_1' \dots x_q'$ が存在して、

$$x_{jt}' = c_{j1} \xi_{1t} + \dots + c_{jk} \xi_{kt} + v_{jt} \quad j=1 \dots q$$

と表わされるものとしよう。

一般にこのような、モデルにおいて r, s, k はあまり大きくないものとしなければならない。しかも更に $\beta_1 \dots \beta_{r+s}$ のいくつかは絶対値が 0 に近いので、その項は無視してもよい程度であると想定されよう。

そうすると結局問題は $x_1 \dots x_r$ の巾の若干を固定し、 $\bar{x}_1 \dots \bar{x}_s$ の中の一部を除き、また $x_1' \dots x_q'$ の中からいくつかを $\xi_1 \dots \xi_p$ の代わりに用いて、

$$\hat{Y} = \hat{\alpha} + \hat{\beta}_1 x_{1t} + \dots + \hat{\beta}_r x_{rt} \\ + \hat{\beta}_{r+1} \bar{x}_{1t} + \dots + \hat{\beta}_{r+s} \bar{x}_{st} \\ + \hat{\gamma}_{11} x_{1t}' + \dots + \hat{\gamma}_{ip} x_{ip}'$$

という形の式を構成することである(ただしここで $\hat{\beta}_1 \dots \hat{\beta}_{r+s}$ のうち一部は最初から 0 とおく)。

そうすると説明変数の一部を除いたことの影響は、制御可能な変数についてはその値をデータにおける平均に等しく固定しておく限り問題はない。また因果関係を表わす変数については、それを無

7) この点に関しては竹内啓『計量経済学の研究』東洋経済新報社、1973年第、12章参照。

視することは偏りを生ずるが、その大きさを理論的に考慮する上で特別な点はない。

問題は最後の部分である。議論を簡単にするために $x_1 \cdots x_{r+s}$ を無視すると

$$\begin{aligned} Y_i &= \alpha + \gamma_1 \xi_{1i} + \cdots + \gamma_k \xi_{ki} + u_i \\ x_{ji}' &= c_{j1} \xi_{1i} + \cdots + c_{jk} \xi_{ki} + v_{ji} \end{aligned} \quad (4)$$

となる。ここで u_i と v_{ji} とは互いに独立であると考えてよいであろう。また v_{ji} $j=1 \cdots q$ も互いに独立と仮定しておく。このことは Y, x_j の相関を表わす部分は、すべて ξ で表現されるものと仮定すれば一般性を損わない。

いま ξ は未知であり、その個数 k も未知であるとする、上記の構造は因子分析のモデルに一致する。

ここで2つの方法が考えられる。一つはまず $x_1' \cdots x_q'$ についてふつうの因子分析の手法⁸⁾を用いて因子スコアを求めてそれを ξ_{it} とし、次に Y のこれらの値に対する回帰式を計算し、

$$\hat{Y} = \hat{\alpha} + \hat{\gamma}_1 \xi_{1i} + \cdots + \hat{\gamma}_k \xi_{ki}$$

とする。更に ξ_{it} の $x_1' \cdots x_q'$ に対する回帰式を

$$\xi_{it} = d_{i1} x_{1i}' + \cdots + d_{iq} x_{qi}'$$

とし、これを先の推定式に代入すれば、

$$\hat{Y}_i = \hat{\alpha} + \hat{\beta}_1' x_{1i}' + \cdots + \hat{\beta}_q' x_{qi}'$$

という形の“予測式”が得られる。

この方法は C. R. Rao が推称するところであって⁹⁾、彼は ξ に時々用いられるような主成分分析法による主成分ではなく、因子スコアを用いることを主張しているが、その理由は上記のような構造から導かれるであろう。ただしここで因子スコアをもう一度 x' で表現しておくことは、予測に用いるために必要である。またこの過程の途中で、適当に“座標軸の回転”を行って、最後に得られた $\hat{Y}$ の式においてなるべく多くの変数の係数を0に近くすることが望ましいであろう。

第2の方法は x, Y をともに一度に座標変換して、 Y の因子荷重から $\hat{\gamma}_1 \cdots \hat{\gamma}_k$ を推定することである。そうして次に上記と同様に ξ を x' の一次式で表わして代入し、 Y の予測式を求めることが

できる。どちらがよりよいかは一概にはいえないようである。

より一般に、制御可能な変数、および明確な因果関係を持つ変数が存在する場合には、まず Y のこれらの変数に対する回帰式を計算し、その残差部分について上記の方法を適用すればよい。

5 多重共線性と“縮小”推定量

現実のデータ解析において、回帰分析に際し多重共線性が生ずることは極めて多い。またそれに関していくつかの方法が提案されている。特に最近ではいわゆる Ridge estimator などという、推定量の絶対値を小さくする手法が論ぜられている¹⁰⁾。それは最小2乗法における正規方程式

$$M_{xx} \beta = M_{xy}$$

の代わりに、この係数行列を拡大して、

$$(M_{xx} + kI) \hat{\beta}^* = M_{xy}$$

の解を推定量とするものである。多重共線性のために x 変数のモーメント行列 M_{xx} が特異に近くなるとき、最小2乗解は大きく変動するが、ridge 推定量 $\hat{\beta}^*$ はあまり大きな値をとることがなくなる。

この方法についても Mallows の C_p 統計量を用いて、 k の適当な大きさを定めたりすることも可能である。

しかしこの場合にも説明変数の性質が明確でないならば、このようなことを行うことの意味がはっきりしない危険性がある。

また Mallows の場合と同じ基準を用いて、最小2乗推定量以外の推定量 $\hat{\beta}_i^*$ を用いて

$$Y_i^* = \bar{Y} + \hat{\beta}_1^* (x_{1i} - \bar{x}_1) + \cdots + \hat{\beta}_p^* (x_{pi} - \bar{x}_p)$$

とし、予測誤差の平均平方の和を計算すれば、

$$\begin{aligned} \sum E(Y_i^0 - \hat{Y}_i^*)^2 &= (\beta^* - \beta)' M_{xx} (\beta^* - \beta) + b^2 \\ &\quad + \text{trace } M_{xx} \Sigma_{\beta^*} + (n+1)\sigma^2 \end{aligned}$$

という形に表わされる。ただし β^* は $\hat{\beta}_i^*$ の期待値、 β は β_i をそれぞれ成分とするベクトル、 Σ_{β^*} は $\hat{\beta}^*$ の分散共分散行列である。また b^2 はモデル自体の偏りの2乗和である。

8) 竹内啓・柳井晴夫『多変量解析の基礎』東洋経済新報社、1972年参照。

9) 1975年来日の際の講義。

10) Hoerl, A. E. and Leslie, R. N., "Ridge regression: biased estimators in non-orthogonal problems," *Technometrics*, 12, 1970.

説明変数を直交変換して $M_{xx}=I$ としておけば、

$$\begin{aligned} \sum (Y_i^0 - \hat{Y}_i^*)^2 &= \sum (\beta_i^* - \beta_i)^2 + \sum V(\hat{\beta}_i^*) \\ &\quad + b^2 + (n+1)\sigma^2 \\ &= \sum E(\hat{\beta}_i^* - \beta_i)^2 + b^2 + (n+1)\sigma^2 \end{aligned}$$

となる。そこで $\sum E(\hat{\beta}_i^* - \beta_i)^2$ をなるべく小さくする推定量が望ましい推定量である。

ここでもし誤差 u が正規分布に従うものと仮定すれば、 $M_{xx}=I$ のとき最小 2 乗推定量 $\hat{\beta}_i$ は互いに独立に平均 β_i 分散 σ^2 の正規分布に従い、そうして $p > 2$ ならば、Stein の有名な推定量¹¹⁾

$$\hat{\beta}_i^* = \max \left\{ 0, 1 - \frac{(p-2)f\hat{\sigma}^2}{(f+2)\sum \hat{\beta}_i^2} \right\} \hat{\beta}_i$$

を用いれば、つねに

$$\sum E(\hat{\beta}_i^* - \beta_i)^2 < p\sigma^2 = \sum E(\hat{\beta}_i - \beta_i)^2$$

となることが示される。ただしここで $\hat{\sigma}^2$ は σ^2 の不偏推定量、 f はその自由度である。仮説 $\beta_i \equiv 0$ を検定する場合の F 統計量を F にすれば、上の推定量は

$$\hat{\beta}_i^* = \max \left(0, 1 - \frac{(p-2)f}{(f+2)F} \right) \hat{\beta}_i$$

となる。Stein 推定量はいろいろな形に変形することができるが、いずれにしても最小 2 乗推定量を縮小する (shrink) ならば、その比率が $\hat{\beta}_i$ の値に無関係に定められる ridge 推定量よりも、 F 比に依存して定める Stein 型の推定量の方がより適当であると思われる。

Stein 型の推定量はモデルの中の変数の一部の係数にのみ適用することもできる。たとえば一組のダミー変数 (季節変動を表わす項など) の係数の推定量を $\hat{\gamma}_1 \cdots \hat{\gamma}_k$ とするとき、もしこれらが互いに独立ならば、これを次のように “縮小” して用いるのが適当である。

$$\hat{\gamma}_i^* = \max \left(0, 1 - \frac{(k-2)f\hat{\sigma}^2}{(f+2)\sum \hat{\gamma}_i^2} \right) \hat{\gamma}_i$$

しかし Stein 型の推定量にしても、性質の異なる説明変数をすべて一まとめにして適用することは好ましくない。

多重共線性の処理にあたっては、それぞれの変

数の特質を考慮することが重要であって、ただ数理形式的方法によってのみそれを扱うことは適当でないと思う。

なお多重共線性が生ずる場合、しばしば “おかしい” 推定値が現われる。しかしある値がおかしいと判定されることは、係数に対してある程度事前の情報が存在していることを意味する。すなわちわれわれは係数の存在範囲についてある程度の知識は持っていることが多い。もしそのような情報、或いは知識を事前分布という形で表現することができるならば、Bayes の方法が適用できるであろう。いま β の事前分布が平均ベクトル b 分散行列 Σ を持つとすれば、その事後平均はほぼ

$$(M_{xx}\hat{\sigma}^2 + \Sigma)\hat{\beta}^* = M_{xy}\hat{\sigma}^2 + \Sigma b$$

或いは

$$(M_{xx} + \hat{\sigma}^{-2}\Sigma)\hat{\beta}^* = M_{xy} + \hat{\sigma}^{-2}\Sigma b$$

をみたす。従って $b=0$ とすれば、これは ridge type の推定量になる。或いはモデルを書変えて、

$$y_i - b'x_i = \alpha + (\beta - b)'x_i + u_i$$

という形に変形し、これに ridge 推定量の考え方を適用すると考えてもよい。ridge 推定量はむしろ Bayes 推定量の一つと理解する方がその意味が明確になるかもしれない。

或いはまた上記のように変換したモデルに Stein の方法と適用することも考えられる。

Bayes 法に対する批判は、主として事前分布の想定に主観性が入るということにある。しかしながら ridge 法などは、それ以上に恣意的な面をふくんでいるのであって、たまたまあるデータに適用して尤もらしい結果が得られたとしても、その判定は Bayes 法の場合以上に主観的といわざるを得ない。私は Bayesian ではないが、もし事前の判断を推定の中に取り入れるとするならば、論理の明確な Bayes 法による方が、Ridge 法などのあまりにも ad hoc な方法を用いるよりも適当であると思う。

6 変数変換の問題

ところで、被説明変数 Y に、予測するに当って、適当な変数変換を行ってから回帰分析を適用することもある。変換を主要な理由としては、非線形

11) 回帰分析の場合については C. Stein “Multiple regression” in *Contribution to Probability and Statistics*, ed. by Olkin et al., Stanford, 1960.

関係の線形化, 誤差分散の均一化の2つが考えられる。対数変換などは, この2つの目的に同時に合致すると考えられることが多いが, それがつねに同時に解決されるとは限らない。また第3の理由として, 誤差分布の正規分布からの極端な隔り, とくに非対称性を除くという目的も考えられる。正の値のみをとる変数を対数変換する目的の一つはこの点にもある。しかしこの目的も上記の目的と同時に達成できるとは限らない。

例えば Y がポアソン分布に従うと考えられる離散量であるときには, $\log Y$ について線形モデルを考えるのが母数の構造としては最も適當であるが, 分散を均一にするためには $\sqrt{Y}$ を, 正規分布に近づけるためには $Y^{2/3}$ を用いなければならない。

ここでいろいろな変数変換の適切さを比較することを考えよう。そのためには先の構造を僅かに一般化して(3)の代りに,

$$E(u_i^2) = \sigma_i^2 \quad E(u_i u_j) = 0 \quad (5)$$

と仮定する。もし観測値が互いに独立ならば, (1), (2), (5)で表わされるモデルはどのように変換を行った値についてもあてはまることはいうまでもない。

いま簡単のために説明変数がすべて互いに直交するとしよう。そうするとモデル M_i をあてはめて最小2乗法により推定を行ったときの係数の推定量は,

$$\hat{\beta}_{ij} = \sum x_{ijt} Y_t / \sum x_{ijt}^2$$

となる。これはやはり β_{ij} の不偏推定量になるが, その分散は今度は

$$(\sum x_{ijt}^2 \sigma_i^2) / (\sum x_{ijt}^2)^2$$

になって, 簡単にならない。そうしてまた

$$E\{\sum (Y_t - \hat{Y}_t)^2\} = \sum \sigma_i^2 + b_i^2 - \sum_j (\sum x_{ijt}^2 \sigma_i^2) / (\sum x_{ijt}^2)$$

となるから, これを, $\sum \sigma_i^2 + b_i^2 - \sum \bar{\sigma}_{ij}^2$ と表わそう。次に先の場合と同様にして Y_t^0 を Y_t と等しい平均分散を持ち, かつそれを独立な量とすると,

$$E(\sum (Y_t^0 - \hat{Y}_t)^2) = \sum \sigma_i^2 + b_i^2 + \sum \bar{\sigma}_{ij}^2$$

となることがわかる。従って

$$E\{\sum (Y_t^0 - \hat{Y}_t)^2\} = E\{\sum (Y_t - \hat{Y}_t)^2\} + 2 \sum_j \bar{\sigma}_{ij}^2$$

となることがわかる。ここで $\bar{\sigma}_{ij}^2$ は σ_i^2 の加重平均であるが, その値は x_{ijt}^2 と σ_i^2 との関係に依存して定まるから, 一般にはそれについて何もいうことはできない。従って Mallows の C_p 統計量にあたるものを計算することはできない。勿論分散比が既知であれば, その逆比を加重として加重最小2乗法を適用することができる。そのときには,

$$E(\sum (Y_t^0 - \hat{Y}_t)^2 / \sigma_i^2) = E(\sum (Y_t - \hat{Y}_t)^2 / \sigma_i^2) + 2(p+1)$$

が成立するから, Mallows の考え方を僅かに修正して適用することができる。しかも $\sum_j \bar{\sigma}_{ij}^2$ が未知の場合その代わりにいいかげんな推定量を用いると, 偏った結論に導かれる危険性が大きい。

分散の構造についての2つのモデル

$$1 \quad E(u_i^2) = \sigma^2, \quad 2 \quad E(u_i^2) = c_i \sigma^2 \quad (c_i \text{ は既知})$$

を比較するには, もはや平均2乗誤差は使えないから対数尤度を比較しなければならない。正規を仮定するとモデル2の下での対数尤度は,

$$2 \log \lambda^{(2)} = \text{const} - n \log \sigma^{*2} - \sum \log c_i$$

となる。ただし σ^{*2} は加重最小2乗法を適用した場合の σ^2 の推定量である。モデル1の下での対数尤度は,

$$2 \log \lambda^{(1)} = \text{const} - n \log \sigma^2$$

となる。従って

$$(\prod c_i)^{1/n} \sigma^{*2} > \sigma^2$$

ならば, 分散の不均一性を考慮しないモデル1の方がよく, 逆の不等号が成立すれば, 分散比についての仮定2を入れた方がよいことになる。

分散の不均一性は, その構造について明確な仮定を想定しない限り, データからは検出されないことが多く, そうしてそれを見落とすと, モデルの適切さについて誤った判断を下す危険性が大きい。従ってデータの性質についての考慮から, その構造についてのあり得べきパターンを注意深く考慮しなければならない。そうしていくつかの可能性についてモデルを作って比較する必要がある。

変数変換において, 分散の不均一性には大きな注意を払わねばならない。多くの場合変数変換に

ついて1)分散の均一性を重視して非線形モデルを構成するか、2)線形モデルにして分散の不均一性を許容するかの2つの方法がある。前者の場合は非線形最小2乗法を、後者の場合は加重最小2乗法を用いなければならない。コンピューターがあれば計算の複雑さは本質的な困難とはならない。しかし結果の評価は、多くの場合1)の考え方をとった方が容易になる。

被説明変数 Y に直接回帰モデルをあてはめた場合と、 $Y' = \phi(Y)$ と変換して、 Y' に線形モデルをあてはめた場合との2つを比較する一つの方法は、その対数尤度を比較することである。すなわち Y についてのモデルの対数尤度を $\log \lambda^{(1)}$ 、 Y' についてのモデルの対数尤度を $\log \lambda^{(2)}$ とするならば、 Y 自体についての第2のモデルの下での尤度はヤコビアン因子を加えて、

$$\log \lambda^{(2)} + \sum \log |\phi'(Y_i)|$$

となる。そこでこの値と $\log \lambda^{(1)}$ を比較してどちらがより適当であるかを判定することができる。回帰モデルの場合、

$$\log \lambda = \text{const} - n \log \hat{\sigma}$$

となるから、このような計算は容易である。

ただしこの場合2つのモデルにおいては当然に誤差分散の構造についての想定が異って来ることに注意を要する。誤差分散についての別個のチェックなしに上記のような方法だけでモデルを選択することは危険である。

なお Y 自体についてのモデルと、変換した値についてのモデルにおいて推定される母数の数に差があるときには、赤池情報量(AIC)との考え方を応用して、それぞれの対数尤度に母数の数を加えて比較を行えばよい。

変数変換を行うこと自体の評価は行われることが少ないが、AICなどを用いる評価は、少なくとも第一次的接近として有効であろう。

7 非正規性と最小2乗法の是非

ところでこれまで推定法はほとんど最小2乗法によるものとして来た。ところで最小2乗法は誤差分布の正規性を前提とし導びかれるものであるから、このことは間接的に誤差の正規性を想定し

ていることを意味する。“縮小法”を論じたときにも正規性の仮定の下で議論を進めた。

しかし回帰式の誤差項が正規分布に従うという仮定は、あまり根拠のないものであるといわざるを得ないであろう。この点を強調すると推定法としての最小2乗法の妥当性も疑わしいと思われるかもしれない。もし誤差分布が正規分布より著しくスソが長い(確率密度が0に近づく近づく方がおそい)と想定されるならば、最小絶対偏差法、すなわち、

$$d = \sum |Y_i - \alpha - \beta_{1p}x_{1ip} - \dots - \beta_{kp}x_{kip}|/n$$

を最小にする値を推定値とする方法を適用することも考えられる。そうして最小2乗法と最小偏差法のどちらをとるべきかについては比

$$\sqrt{\frac{2}{\pi}} \hat{\sigma}/d$$

を基準とすることができる。正規分布の場合この値がほぼ1になるから、これが1より著しく大きい、例えば2以上ならば、平均偏差法を用いる方がよい。

このようなタイプのいわゆる robust な推定量を求める問題は、P. Huber らによって論ぜられているが¹²⁾、しかし問題の形式的な処理よりも、残差分析を用いる方が实际的である。それには

$$\hat{u}_i = Y_i - \hat{Y}_i$$

或いは、それを基準化した、

$$\hat{u}_i' = \hat{u}_i/a_i \text{ ただし } a_i = E(\hat{u}_i^2)/\sigma^2$$

を用いる方がよい。 $\hat{u}_i'$ の中で著しく大きいものがあれば、それを“異常値”とみなすことができる。その棄却限界としては、ふつうの Thompson-Smirnoff の棄却限界の値を $\sqrt{1-1/n}$ 倍したものをいれれば近似的には十分である。このような“異常値”はそれが本当に他の値にあてはまる式からズレた値であるか、または分布が正規分布でないかの、2つの理由によって生じ得る。しかしいずれの場合にもその値を除いて最小2乗推定を行えば、精度のよい推定量が得られるであろう。しかしこのようにして予測値を計算したとき、その信頼性は見かけより落ちることに注意しなければ

12) P. Huber, "Robust estimation," Lecture note, Yale University, 1971.

ばならない。すなわち予測される値が“異常値”になるかもしれないからである。“異常値”による検定は、コーシー分布のような非常にスロが長い分布に対しては(標本数があまり小さくない限り)検出力がかなり高くなるので、このような方法はかなり有力である。すなわち“異常値”に注意することにより(そうしてこのことはデータにつきまとうミスを発見するためにも怠ってはならないことである)、非正規性から生ずる最小2乗推定量の効率の低下は、かなり防ぐことができる。

現実の問題においては、データから明らかに見られる非正規性を見落さない限り、非正規性の問題はそれほど重大でないといってよい。

8 む す び

以上にのべたような問題点は、より複雑な問題、例えば同時方程式体系や、時系列の場合に生ずる

ものと、本質的には共通であると考えられる。そのような場合にも、より簡単なモデルとより複雑なモデルとの間の選択、変数変換の是非、非正規性の問題等が生ずる。そこでやはり C_p 統計量や尤度比、或いはそれを拡張した AIC 統計量などの考え方を利用することができるであろう。またその際、あれこれの手法を形式的に適用するだけでなく、各変数の意味と性質、分析や予測の目的を具体的に考慮して、適切(relevant)な結論を導くことが大切である。このことは統計的推測のすべての場面を通じていわれるべきことであるが、その際何らかの意味で指針となるような数学的理論を構成することが大切であると思う。ここでのべたことはそのための一つの手がかりとなり得るであろう。

(東京大学経済学部)

農 業 経 済 研 究

第 48 巻 第 3 号

(発 売 中)

今村奈良臣：農地賃貸借の構造変化

新谷正彦：戦前養蚕部門における夏秋蚕の普及と生産弾性値の変化

朽木昭文：農民層分化に関する確率論的アプローチ(明治初期を中心として)

鈴木敏正：不足払制度下における「酪農危機」の生成メカニズムについて
——農産物過剰論的接近——

長谷川宏二：農村社会研究の展開と課題——農業経済学との関連において——
神谷一夫

《書 評》

玉城昌幸：足羽進三郎著『農業協同組合の研究』

吉田忠：御園喜博著『現代農業経済論』

常盤政治：暉峻衆三編著『政治革新と世界の農業問題』

B 5 判・50 頁・500 円

日本農業経済学会編集・発行／岩波書店発売