

経済と統計の間で¹

美添 泰人*

In the Realm between Economics and Statistics

Yasuto Yoshizoe*

経済分野における統計を適切に分析しようとする際、いくつかの課題に直面する。一般に、検証されるべき仮説は経済学の領域における理論ないし経験から導かれるもので、統計学に固有の対象とはいえないが、適切な資料と分析手法の探求は統計学の本来の課題でもある。このような課題のうち、ベイズ統計学、頑健統計学およびデータ解析、公的統計の利用を対象として、筆者の現状認識と今後の展望を記す。

We encounter with several problems whenever we intend to apply adequate statistical analyses in the field of economics. In general, hypotheses concerning economic issues are derived from economic theory or experience, which are not necessarily proper statistical problems. Yet, the process of seeking appropriate data as well as investigating analytical methods belong to the inherent field of statistics. In the lecture, I presented my personal recognition of current issues of economic statistics, specifically, on Bayesian methods, robust statistics and exploratory data analysis, and challenges to effective use of official statistics.

キーワード: ベイズ統計, 頑健統計学, データ解析, 経済統計, 公的統計

1. はじめに

経済と統計の関係については、経済には直接的な関係が深いとは思えない R. A. Fisher も、その著書 *Statistical Methods for Research Workers* (1963) で次のように指摘している。「統計学とは (i) 集団の研究, (ii) 変動の研究, (iii) データの簡約方法に関する研究とみなすことができる」とし、さらに「社会科学においても統計的方法は本質的に大切であって、社会科学が科学の域に達したのは主として統計的方法の賜である。統計的方法に対する社会科学のこの特殊な依存関係によって、統計学は経済学の 1 部門とみなすべきであるという不幸な誤解が生じた。実は経済学上のデータの処理に適した方法は、今までのところでは、たいていは生物学その他の科学の研究においてのみ発達してきたのである。」などと述べた上で、経済の分野には大量の統計資料があることを指摘している。ま

¹ 本稿の内容は統計関連学会連合大会における会長就任講演 (2009 年 9 月 8 日) に基づくものである。

* 青山学院大学経済学部, 〒150-8366 東京都渋谷区渋谷 4-4-25.

た、別の機会に「農業の統計データに比較して経済の分野では豊富な統計資料があり、統計的手法を適用するのに適した分野である」という趣旨の指摘もあり、経済分野における統計資料に対してある種の羨望が窺える。

ところで、経済統計を適切に分析しようとする際、一般的な統計手法の選択に限らず、いくつかの課題に直面する。基本的に、統計資料を利用して検証されるべき仮説は経済学の領域における理論ないし経験から導かれるもので、統計学に固有の対象とはいえないが、具体的な課題に対応した統計資料の整備・利用と、適切な分析手法の探求は統計学の本来の課題といえる。

経済統計という用語については、官庁統計シンポジウムにおける西村清彦氏（日本銀行副総裁）の発言にあるように、経済を対象とする「統計」として調査・作成方法まで含める視点と、逆に「経済」の問題から出発して、それを分析するために統計を利用するという、やや異なる視点がある。前者は、経済構造を把握する目的があって、そのために統計を作成し利用するという姿勢であり、後者は経済分析の視点が先にあって、その目的で使える統計を探すという姿勢である。

経済の実証分析についても、理論的に導かれた経済モデルを検証する立場と、統計データの分析を通じて新たな知見を得たり、経済構造の変化を発見する立場には違いがあるし、さらに多くのエコノミストが関わっているように、市場の動向を示す指標を提供することも経済統計の重要な課題である。

広い意味での経済と統計の関わりに対する筆者の現状認識と、今後の展望を紹介することが本稿の目的である。以下ではまず、経済統計の分析に関わる統計的手法として、ベイズ統計学と頑健統計学およびデータ解析の視点に触れ、続いて統計情報の作成・提供・利用に関わる制度的な問題として公的統計に触れる。

2. ベイズ統計学

近年、ベイズ統計の応用が活性化し、さまざまな水準の解説書も急増している。経済分析でもベイズ統計の応用が主流になりつつあると言っても良いほどの活況であり、個人的には歓迎すべきことと考えている。しかしその一方で、ベイズ統計の意味を十分に理解しているとはいいがたい分析も増加しているように見受けられる。

その発展過程をたどれば明らかなように、ベイズ統計学が合理的な結論を導く理由が手法としての論理整合性にあることは、L. J. Savage, D. V. Lindley, J. W. Pratt などが指摘してきたところであり、現代的なベイズ統計学の手法を適切に適用するためには、理論的定式化の理解が不可欠である。

1960年代にはベイズ統計の手法、特に主観確率は一般的に受け入れられていたとは言えない。Lindley (1965) では、正規分布や二項分布の母数に関するベイズ統計の推論は古典

的な手法と類似の結果を与えることを紹介し、利用者に安心感を与えることが意図されていた。それに対して Box and Tiao (1973) では、古典的な統計学では解決が困難とされたいくつかの問題に対して、ベイズ統計を利用すれば説得的な解が得られることが示された。また、計量経済学の分野でも Shiller (1973) による分布ラグの推定手法は、ベイズの手法が大きな成功を収めた例である。その後、数値計算の発展とともに、ベイズ統計が急速に普及するまで過程を美添 (1983, 1989, 1993–96, 1996) などに基づいて概観する。

2.1 ベイズ統計学の構造

一般に、標本空間 $x \in X$, 母数空間 $\theta \in \Theta$, および 密度関数 $p(x|\theta)$ からなるモデル (X, Θ, P) の下で、統計的推論が定式化される。なお、本稿では確率分布を表わすのに一般的に $p(x|\theta)$, $p(y)$, $p(\theta)$ などの記号を用いる。Wald (1950) による統計的決定問題では、これに加えて、行動の集合である決定空間 $d \in D$ と母数 θ に対して損失関数 $L(d, \theta)$ が与えられ、最適な行動 (decision rule) $d = \delta(x)$ を定める。

Lehmann (1959, 1986) などに詳細な解説があるとおおり、たとえばリスク (損失の期待値) $R(\theta, \delta) = E[L(\delta(x), \theta) | \theta] = \int_X L(\delta(x), \theta) p(x|\theta) dx$ を用いるとしても、これは未知の θ を含むため、ミニマックスなどの適当な基準を用いて決定関数 δ が選ばれる。特別な場合として、母数の推定問題や、仮説検定の問題が導かれることも良く知られている。ここで、これらは便宜的に導入された基準であり、その背景に理論的な前提を設けて導出されたものではない点に注意が必要である。

Savage (1954) 以降の正統的なベイズ統計では、以上の統計的モデルに加えて事前分布 $p(\theta)$ が導入される。また、損失関数と類似の概念として効用関数 $U(d, \theta)$ が用いられる。そこでは、ベイズの定理によって事後分布 $p(\theta|x) = p(x|\theta)p(\theta)/p(x)$ を求め、行動 d は事後期待効用 $\int_{\Theta} U(d, \theta) p(\theta|x) d\theta$ を最大にするように決定される。ここで選択される行動は、統一的な原理から整合的に導かれる基準であるという点について以下に検討する。

2.2 論理整合性 (coherence)

ベイズ統計の正当化に関してはいくつかの議論があるが、ここで紹介するものは意思決定者の合理的な行動を前提として導かれる論理整合性 (coherence) であり、筆者が最も重要な根拠と考えるものである。合理的行動に関する公理から、事前分布 $p(\theta)$ の存在、効用関数 $U(d, \theta)$ の存在、さらに行動 d は事後期待効用 $E(U|x)$ を最大にするように選ばれるべきであること、の3つが整合的に導かれる。L. J. Savage, J. W. Pratt, M. H. DeGroot などによっていくつかの公理体系が紹介されているが、ある公理体系では効用は確率の一種であり、そのときは事後期待効用も確率の一種となる。有限母数空間については弱い公理の下で事前分布の存在は示されるし、実用上はそれで十分であろう。

また決定と統計的推測 (statistical inference) を分離する場合にも、Lindley (1972) のよ

うな議論が可能である。すなわち統計的推測とは、発生し得る（ある効用関数 U による）任意の決定問題に対して必要な情報をあらかじめ整理しておくことである。ここで、最適な行動の基準である事後期待効用の評価には効用関数 U の他に事後分布 $p(\theta|x)$ だけが必要だから、ベイズ統計では事後分布を評価することが推測そのものである。ただし A. P. Dempster のように、ベイズ統計 Bayesian Statistics とベイズ決定理論 Bayesian Decision Theory は異なるという主張もある。

以上の論理的整合性が最大の特長であり、結論に矛盾が生じないことが、ベイズ統計を正当化する重要な根拠を与えている。ところで標本理論の手法は、特定の効用関数と事前分布を仮定したベイズ統計による解と一致することが多い。ベイズ推定量などと呼ばれる解 $d = \delta(x)$ は効用関数に対応して変化することから、標本理論の手法は恣意的な方法でありながら、それを明示していないものと言える。

2.3 事前分布の設定

事前確率の存在は正当化できても、具体的な事前分布の定め方は、また別な問題である。まず、決定問題に関わる個人の主観的判断としての事前分布を心理実験を通じて定める方法があり、その変種として専門家集団の見解を集約して母数に関する判断とする提案もある。前者の例としては M. Freedman と L. J. Savage の有名な実験があり、後者の例は Press (1989) にある。また、Bernardo and Smith (1994) に紹介されている交換可能性 (exchangeability) も、ある種の問題に対して説得力のある事前分布を提供する枠組みである。

精密測定の原理 L. J. Savage らによる精密測定 (precise measurement または安定的推定) の原理と呼ばれる定理は、任意の事前分布 $p(\theta) > 0$ に対して、標本サイズ n が大きくなると事後分布 $p(\theta|x_1, \dots, x_n)$ は母数の真の値に集中することを主張する。このように、事前分布の多少の差は結論に影響しないことも、ベイズ統計の信頼性を保証する。なおここで重要なのは $p(\theta) > 0$ という条件であり、真の値について $p(\theta_0) = 0$ なら、常に $p(\theta_0|x) = 0$ となり観測値は事後分布に反映されない。不適切なモデルの設定が重大な情報の損失を与えることは当然であるが、古典的な統計学では明確な意識をせずに狭いモデルを扱う場合が少なくない。このことは $p(\theta_0) = 0$ と想定していることを意味している。ベイズ統計でもその事情は同様とはいえ、モデルの作り方次第では、母数の制約条件を緩めることが容易に行える点でも、ベイズ統計の実用性は高い。

情報のない事前分布 整合性の考え方とは異なって、データのもつ情報を十分に引き出すために、なるべく個人的な主観の入り込む余地のない事前分布を採用し、また効用関数は特定化しないという立場がある。この立場は、Savage (1954) によるベイジアン革命以前から Bayes-Laplace の流れをくんでいた Jeffreys (1961) に代表されるものであるが、Box and Tiao (1973) も同じような考え方を採用している。情報を持たない事前分布 (non-

informative prior) のひとつである Jeffreys の invariant prior では、観測値の分布の Fisher 情報量行列 $I(\theta) = E[(\partial \log p / \partial \theta)(\partial \log p / \partial \theta)' | \theta]$ の行列式を用いて $p(\theta) \propto |I(\theta)|^{1/2}$ とする. また Box-Tiao の data translated likelihood に基づく事前分布では、 $\phi = \phi(\theta)$ と変換した母数に関する尤度関数 $\ell(\theta) = p(x|\phi)$ において ϕ が (近似的な) 位置尺度とみなせる場合に、 ϕ の事前分布として一様分布を仮定する. この考え方による事前分布は Jeffreys の事前分布と、ほぼ一致する. ところで θ そのものが位置母数で $p(\theta)$ が一様分布なら、事後分布 $p(\theta|x)$ は尤度関数 $p(x|\theta)$ に比例する. したがって尤度関数を利用する手法は、ベイズ統計の特殊な場合と言える.

データのもつ情報を重視する立場では、狭いモデルに限定してデータの意味を見逃すという過ちを避けるというデータ解析学派 (data analysis school) と結びついている. この立場に「情報のある事前分布」を追加したものがきわめて実用的であり、ベイズ統計の成功例の多くは、情報のある事前分布を効果的に利用したことによる.

共役事前分布 指数分布族 $p(x|\theta) = a(\theta)b(x) \exp \sum_{i=1}^m \theta_i t_i(x)$ を対象とする問題では、 $\alpha_0, \dots, \alpha_m$ を超母数 (hyper parameter) として、事前分布を $p(\theta) \propto \{a(\theta)\}^{\alpha_0} \exp\{\sum_{i=1}^m \alpha_i \theta_i\}$ とおくと、事後分布では超母数だけが変化する. これを元の分布に対して共役 (conjugate) な事前分布と呼ぶ. 精密測定 of the principle from, 事前情報を近似的に表現する事前分布として用いられる.

2.4 標本理論の非整合性

以上のようにベイズ統計の手法が公理体系から整合的に導かれるのに対して、標本理論には標本分布を手法の評価に利用するという以外に明確な原理がないため、問題に応じて、不偏性、一致性、最小分散性など、多くの原理が導入されてきた. しかし、一貫した原理を持たないことから、論理的な矛盾が発生することは当然である.

標本理論のいくつかの原理 標本理論で用いられる代表的な原理を確認しよう.

まず、弱い尤度原理は「確率分布 $p(x|\theta)$ に従う 2 組の観測値 (ベクトル) x_1, x_2 について、尤度関数が $p(x_1|\theta) \propto p(x_2|\theta)$ という比例関係を満たすとき x_1 と x_2 は同じ推論を導く」というものである. これは最小十分統計量を用いる推論と同等であり、最も基本的で一般に認められている.

一方、強い尤度原理は「異なった確率分布に従う 2 組の観測値 (ベクトル) x, y の尤度関数が $p(x|\theta) \propto p(y|\theta)$ を満たすとき x と y は同一の推論を導く」というものである. この原理は、以下に示すように標本理論では受け入れられない.

例 1 $x \sim N(\theta, 1)$ に対して仮説 $H_0: \theta = 0$ を検定するとき、標本平均が $|\bar{x}| > c/\sqrt{n}$ を満たすまで標本を取り続ける逐次検定では、実際に $\theta = 0$ の場合にも確実に仮説が棄却される. 尤度は停止規則には依存しないから、強い尤度原理に従えば、この結論を受け入れ

なければならない。

次に、条件性の原理は、ある変数 z の分布が θ と無関係のとき「 θ に関する推論は z の値を固定した条件付分布で評価する」ことを要請する。例として、ある物体の特性 θ を測定するのに、精度の異なる器具を乱数 $z \sim U(0, 1)$ を用いて無作為に選び、 $z < 1/2$ のとき $x \sim N(\theta, \sigma_1^2)$ 、 $z > 1/2$ のとき $x \sim N(\theta, \sigma_2^2)$ とする場合、使った器具を知って推論を行うのであれば当然の原理である。このような確率変数を補助統計量とよぶ。

以下に記す例の問題は、ベイズ統計では矛盾なく解決されるものである。

不偏性 統計量 $T(x)$ が不偏という条件 $E[T(X) | \theta] = \theta$ は標本空間 X に依存する。例えば、成功の確率を θ とする n 回の実験で r 回の成功が観測された場合、 n を固定していれば r は二項分布で $X = \{0, 1, \dots, n\}$ 、不偏推定量は r/n である。一方、 r を固定したときは n は負の二項分布で $X = \{r, r+1, \dots\}$ 、不偏推定量は $(r-1)/(n-1)$ と違う結果となる。これも強い尤度原理を否定する例である。

例 2 はかりを 1 度だけ使用してふたつの物体 A, B の重さの不偏推定量 $\hat{A}, \hat{B}$ を作成する。そのためには、コインを投げて表が出たら A を測定し、その結果 x から $\hat{A} = 2x$ 、 $\hat{B} = 0$ とする。裏が出たら B を測定し、 $\hat{A} = 0$ 、 $\hat{B} = 2x$ とすればよい。しかし、現実に観測値 x が与えられたときには、この推定量は明らかに誤った推論を与える。この例では、不偏推定量は条件性の原理を満たしていないが、標本理論では現実の観測値ではなく、可能な結果をすべて考慮に入れるため、不偏性はこのような奇妙な結果を与えることがある。なお観測値を固定する点については図 1 でも触れる。

仮説検定 母数 θ に関する信頼係数は、同じ手順を繰り返すとき「信頼区間が θ を含む割合」という標本理論の評価を表すものである。この方法が、繰返しを前提としない場合に無意味になることは、広く知られている。すなわち、観測値 x から計算された母数 θ に関する信頼係数 $1 - \alpha$ の信頼領域 $R(x)$ は、すべての θ について $\Pr\{\theta \in R(x) | \theta\} = 1 - \alpha$ が成立するということであり、これは $\Pr\{\theta \in R(x) | x\} = 1 - \alpha$ を意味しない。したがって、特定の観測値 x が与えられたときには何も結論できないという極論 (J. Neyman) もある。

その他、さまざまな原理 (不変性, 一致性, UMVUE, etc.) が提唱されているが、個々の問題に応じて適用される原理の数を増すことによって、標本理論の本質的な欠陥を取り除くことはできないのは当然である。

2.5 ベイズ統計学と標本理論の関係

図 1 の楕円群は、事前分布と尤度関数から求められる同時分布 $p(x, \theta) = p(x|\theta)p(\theta)$ の等高線を表わしている。ベイズ統計では観測値 x が与えられた後は x を固定して図の水平な直線上の事後分布 $p(\theta|x)$ にのみ関心があるのに対して、標本理論では、ある $\theta = \theta_0$ を

固定したときの図の垂直な直線上の分布 $p(x|\theta_0)$ から統計量の挙動を評価することになる．
 このように両者の違いは， θ を固定して繰り返し標本 x について評価するか， x を固定して θ について評価するかにある．

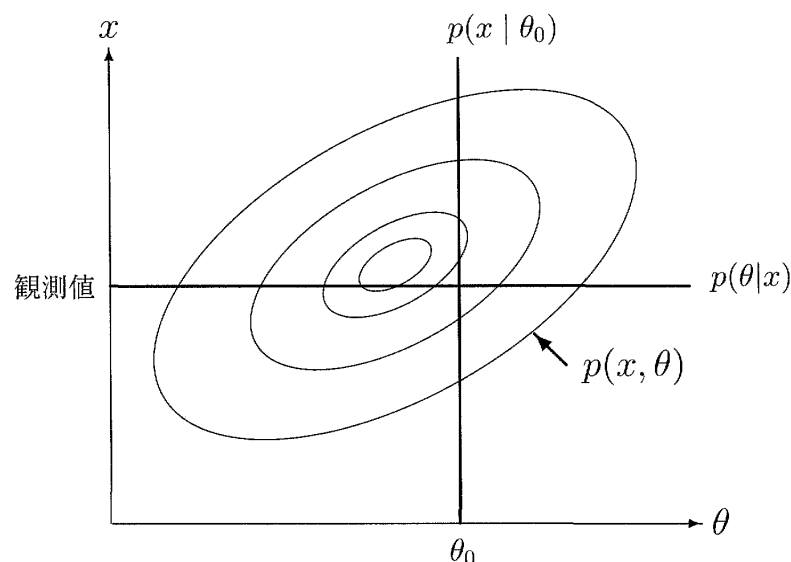


図1 同時分布・標本理論・事後分布

Jeffreys の議論，事後オッズ，区間推定 標本理論において強い尤度原理が認められないとした例1 (165 ページ) について，ベイズ統計と標本理論を比較することで，両者の違いが明らかになる．Jeffreys (1961) は，この問題に関する事前分布として $P\{\theta = 0\} = \pi_0$, $P\{\theta \neq 0\} = 1 - \pi_0$, かつ $\Theta - \{0\}$ の範囲で連続な $p(\theta)$ を想定している．標本平均 $\bar{x}$ が与えられたときの事後確率は容易に評価できるが，仮説 $H_0 : \theta = 0$ が任意の有意水準でちょうど棄却される値 $\bar{x} = c/\sqrt{n}$ を代入して $n \rightarrow \infty$ とすると， $p(\theta)$ が有界なら $P\{\theta = 0 | x\} \rightarrow 1$ が導かれ，データは $\theta = 0$ を明確に示している．これも，有意水準 α の検定が整合的でないことの一つの表現である．

なお，この議論には注意すべき点がある．標準的なベイズ統計の区間推定 (Credible interval) では最大事後密度 (HPD: highest posterior density) 領域を用いるが，例1の正規分布 $x \sim N(\theta, 1)$ に関しては，標準的な事前分布を $\theta \sim N(m, 1/\tau_0)$ とすると ($\tau_0 \rightarrow 0$ のとき)，次の区間が得られる．

$$\Pr\left\{\bar{x} - 1.96 \frac{1}{\sqrt{n}} < \theta < \bar{x} + 1.96 \frac{1}{\sqrt{n}} \mid x\right\} = 0.95$$

一方， $\theta = 0$ が成立するかどうかを考えると，Jeffreys の議論では事後オッズ (posterior odds) は次のようになる．ただし $\tau_1 = \tau_0 + n$ とする．

$$\frac{P\{\theta \neq 0 | x\}}{P\{\theta = 0 | x\}} = \frac{1 - \pi_0}{\pi_0} \frac{\sqrt{\tau_0}}{\sqrt{\tau_1}} \exp\left\{-\frac{n}{2} \left[\frac{\tau_0}{\tau_1} (\bar{x} - m)^2 - \bar{x}^2\right]\right\}$$

いま、数値例として $\bar{x} = 0.015$, $n = 40000$, $m = 0$, $\tau_0 = 0.01$, $\pi_0 = 0.5$ とすると、推定区間 (Credible interval) は $0.0052 < \theta < 0.0248$ と $\theta = 0$ を含まないのに対して、事後確率は $P\{\theta = 0 | x\} = 0.9569$ と大きく、当然、事後オッズも $P\{\theta \neq 0 | x\} / P\{\theta = 0 | x\} = 0.0450$ と非常に小さい。

このように事後オッズと区間推定の間で一見したところ矛盾する結論が得られることから、ベイズ統計の理論的な不備を指摘する見解もあるが、それは事前分布が違うことを理解していない議論である。この例では区間推定に関しては情報のない事前分布として $\theta \sim N(m, 1/\tau_0)$ ($\tau_0 \rightarrow 0$) を想定した。当然、小さな $h > 0$ に対して $|\theta| < h$ となる確率は小さい。これに対して Jeffreys の事前分布では測度ゼロの集合 $\{\theta = 0\}$ に $\pi_0 > 0$ を付与している。このような大きな違いがあれば、区間推定と事後オッズで結論が異なるのは当然である。

なお、筆者は1点の確率を正と想定することは不適当と考える。すなわち仮説 $H_0: \theta = \theta_0$ の現実的な意味は、合理的な範囲で小さな $h > 0$ に対して $|\theta - \theta_0| < h$ が満たされるということである。例1は標本理論が整合的でないことを示しているが、ベイズ統計に関しては矛盾はない。

標本理論とベイズ統計学の比較 多くの矛盾はあるものの、標本理論が実用的な成果をあげてきたことは事実である。しかしそれは、ある種の問題については、標本理論の手法はベイズ統計の近似になっているという偶然によるものである。例えば、不偏推定量、信頼区間、仮説検定などが、ベイズ統計の解として構成できることが Pratt (1965) によって示されている。実際的に重要な結果は、正規線形回帰モデルにおける推論がベイズ統計の近似的な解となることで、多くの実用的な問題はこの手法に基づいている。

しかし、複雑な問題に対しては整合性が保証されていないし、Birnbaum (1972) による「弱い尤度原理と条件性の原理は強い尤度原理の十分条件である」という結論は、標本理論の非整合性を端的に表わすものである。

例3 最も簡単な二項分布 $B(n, \theta)$ の推論でも、標本理論とベイズ統計を比較することができる。ベイズ統計では共役事前分布はベータ分布 $Be(a_0, b_0)$ で与えられ、事後分布 $p(\theta|x)$ は $Be(a_1, b_1)$ (ただし $a_1 = a_0 + x$, $b_1 = b_0 + n - x$) となる。ベータ分布 $Be(a, b)$ の上側 $100\alpha\%$ 点を $B(\alpha; a, b)$ と表すとき、 θ_1, θ_2 を $p(\theta_1|x) = p(\theta_2|x)$, $\alpha_1 + \alpha_2 = \alpha$, $\theta_1 = B(1 - \alpha_1; a_1, b_1)$, $\theta_2 = B(\alpha_2; a_1, b_1)$ を満たすように定めれば、 $(1 - \alpha)$ HPD 区間は $\theta_1 < \theta < \theta_2$ と与えられる。事後分布 $p(\theta|x)$ が非対称なら $\alpha_1 \neq \alpha_2$ である。

一方、正規近似を利用しない場合の標本理論は次のようになる。 α_1 と $\alpha_2 = \alpha - \alpha_1$ に対応して、 $\Pr\{x \leq m_1 | \theta\} \geq \alpha_2$ となる最小の整数を $m_1(\theta)$, $\Pr\{x \geq m_2 | \theta\} \geq \alpha_1$ となる最大の整数を $m_2(\theta)$ とすると、 $\Pr\{m_1(\theta) \leq x \leq m_2(\theta) | \theta\} \geq 1 - \alpha$ となるから、不等式 $m_1(\theta) \leq x \leq m_2(\theta)$ を θ について解けばいい。そのためには $x \sim B(n, \theta)$,

$u \sim Be(m+1, n-m)$ のとき $\Pr\{x \leq m\} = \Pr\{u > \theta\}$ が成り立つことを利用すると、 $B(1-\alpha_1; x, n-x+1) < \theta < B(\alpha_2; x+1, n-x)$ という信頼区間が導かれる。

事前分布を $a_0 = b_0 = 1$ と想定し、標本理論の α_1, α_2 を信頼区間が最短になるように定めると、HPD 区間と信頼区間はほぼ一致する。この例でも、論理の一貫性ととも、ベイズ統計の手法の簡明さは明らかである。

2.6 ベイズ統計の有効な適用

現実には、母数に関して何らかの先験的情報を持っている場合が多い。標本理論では情報の確からしさを表現することが難しく十分なデータ分析が妨げられるのに対して、ベイズ統計では事前分布を工夫することで新たな展開が開ける。特に、線形モデルにおける線形制約の問題は汎用的な手法を与え、多次元母数のモデルに利用されている。その応用として滑らかな応答局面を持つ回帰モデルの推定、計量経済分析における Shiller ラグ、季節調整法 Decomp (BAYSEA) などがある。以下、線形制約を取り上げる。

行列表記で線形回帰モデル $y = X\beta + u$, $u \sim N(0, \sigma^2 I)$ を想定し、既知の R, r による線形制約を $R\beta = r$ とする。このとき、仮説 $H_0: R\beta = r$ が受容されれば線形制約つき最小二乗法の解 $\hat{\beta}_R$ を求め、棄却されれば通常解 $\hat{\beta} = (X'X)^{-1}X'y$ を求めるという手順では、線形制約は無視するか、正しいとするかの選択しかない。

一方、ベイズ統計では、線形制約を表現する事前分布として $R\beta - r \sim N(0, \phi^{-1}I)$ を想定できる(簡単のために σ を固定する)。ここで ϕ は情報の信頼性を表わす母数である。これ以外には情報がないときは β の事前分布は $p(\beta|\sigma^2) \propto \exp[-(\phi/2)(R\beta - r)'(R\beta - r)]$ となる。この事前分布は一般に可積分ではない(improper prior)が、事後分布 $p(\beta|y, \sigma^2)$ ではその困難は解消され、 $N[\bar{\beta}, \sigma^2(X'X + kR'R)^{-1}]$ となる。ただし $\bar{\beta} = (X'X + kR'R)^{-1}(X'y + kR'r)$, $k = \sigma^2\phi$ である。ここで k は事前情報と観測値の精度の比であり、 $k \rightarrow 0$ は情報のない状態、 $k \rightarrow \infty$ は完全な情報を表わす。実際、 $k \rightarrow 0$ とすると $\bar{\beta} \rightarrow \hat{\beta}$ (最小二乗推定量) となり、 $k \rightarrow \infty$ とすると $\bar{\beta} \rightarrow \hat{\beta}_R$ (制約付き最小二乗推定量) となる。母数 σ^2 については共役事前分布を想定するのが最も簡単であるが、詳細は省略する。いずれにしても、 $0 < \phi < \infty$ となるような中間的な値を用いるところに、データ解析的な手法としてのベイズ統計の長所がある。

数値計算の手法 指数分布族に含まれない t 分布などを含む複雑なモデル、特に母数 $(\theta_1, \dots, \theta_k)$ が高次元となる問題では事後分布を解析的に求めることは困難であったが、今日では、マルコフ連鎖・モンテカルロ法(MCMC)などの数値計算法によって事後分布を評価することができる。MCMCの基本的な考え方は次のとおりである。事後分布 $p(\theta|y)$ に対して推移確率を $p(\theta, \theta^*)$ とする θ のマルコフ連鎖で $p(\theta|y)$ が定常分布となるようなものを選ぶ。マルコフ連鎖の選び方としては Metropolis 法, Gibbs Sampler など、いく

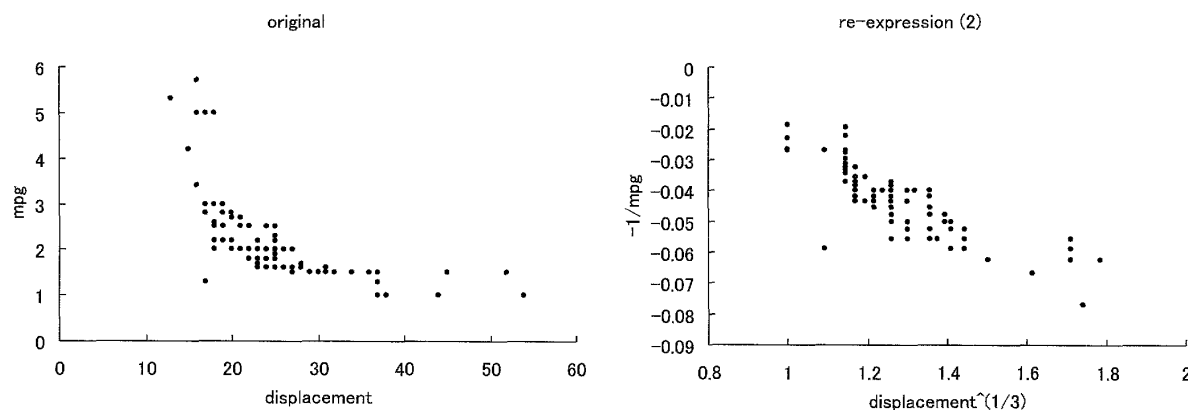


図2 排気量 (x) と燃費 (y) の関係 (左: 原単位, 右: $x^{1/3}$ と $-1/y$)

つかの方法が提案されている. 適当な初期値 $\theta^{(0)}$ から, このようなマルコフ連鎖に従って $\theta^{(1)}, \theta^{(2)}, \dots, \theta^{(m)}$ を生成すると, 十分に大きい m に対して $\theta^{(m)}$ の分布は事後分布 $p(\theta|y)$ に収束するから, 収束後に生成される乱数を用いれば事後分布を評価したり, θ のさまざまな関数 $g(\theta)$ の期待値を数値的に求めることができる. 収束および精度の評価については基本的なマルコフ連鎖の理論が利用できるほか, 多くの方法が開発されている. そのため, 標本理論では解決が困難とされた複雑な問題でも, 数値計算によって解答を求めることができる. このように MCMC は一般的な乱数発生の手法であり, ベイズ統計の事後分布の評価において多大な貢献をしている.

しかし, 不適切な統計モデルからはどのような数値計算を用いても不適切な結果しか導けない. MCMC が使えるという理由で, 標本理論に比較してベイズ統計が有用であるという主張が増えているが, これはベイズ統計の本質を無視した皮相的な見解である.

2.7 ベイジアン統計学はいつでも有用か

統計分析に際してはデータ処理に関する十分な考察が必要である. Berger (1995) は, アメリカにおける乗用車のデータに関して $y_{ijk} = \beta'x_{(ijk)} + \alpha'x_{(ij)}^* + \epsilon_{ijk}$ という, 複雑なモデルの推定を論じている. ここで y_{ijk} は第 i メーカー・第 j モデル・第 k 観測値の燃費 (対数), $x_{(ijk)}$ は乗用車の特性を表すベクトル, $x_{(ij)}^*$ はメーカーとモデルの特性, β と α はそれぞれ固定効果, 変動効果を表すベクトルである (' は転置記号). さらに β に関しては事前の知識から $\beta_{10} > 0, \beta_4 \leq \beta_5 \leq \beta_6$ などが知られており, $\alpha_{ij} \sim N(\mu_i, V_i), (j = 1, \dots, J_i)$ かつ μ_i は AR(1) 過程に従うなどと仮定している. これほど複雑なモデルになると, 伝統的な標本理論の手法では実質的に解くことは困難となり, 数値計算は確かに有効である.

しかし, 本当にこれほど複雑なモデルが必要かどうかは, 検討の余地がある. このデータの問題点は図2 (左) のような強度の非線形性が観察されるところにある. Berger は被説明変数の対数を利用しているが, その程度の変換では不十分である. EDA (探索的デー

タ解析) の手法を利用すれば、自然に選択される変換は図 2 (右) のように、 y の逆数と x の 3 乗根である。このような状況であれば比較的簡単な線形回帰モデルによっても十分な分析ができる。以上のような準備作業の上で、なおかつ複雑なモデルを推定する必要があるれば、そのときこそベイズ統計の強力な分析手法がその力を発揮することになるだろう。

経済分野でベイズ統計の手法が真に有効なのは、経済的な問題に固有の知識を統計モデルの中に取り入れるように事前分布を工夫することを通じてである。ごく一般的な、ただし相当程度複雑なモデルを作り、情報のない事前分布を想定し、MCMC などによって事後分布を評価するだけなら、ベイズ統計を用いる必然性はないし、練習問題以上の価値はない。ベイズの手法は悪いモデルに対しても利用可能であるが、統計家のデータを見る能力に代替可能なものではないことは明確に認識されるべきである。

3. 頑健統計学およびデータ解析の視点

経済統計には外れ値が少なからず存在する。しかし家計の所得や資産保有額などに現れる外れ値の多くは「正しい」外れ値であり、経済変数の分布が強い正の歪みを持つことから必然的に発生するものである。そのため、公的統計における標本調査の実施や統計表の集計にあたっては、頑健統計学の視点の重要性が次第に認識され、手法が明示されるようになってきたことは、当然の成り行きである。

その一方で、計量分析の中には、頑健でない手法を適用しているにもかかわらず、いくとおりかの条件の下で分析結果に大きな差がないことをもって、ロバストな推定結果が得られたとする主張も散見される。形式的な頑健性の議論は経済分析に適當ではないという意見は尊重するものの、特に経済統計の分析に当たっては、頑健性とデータ解析の視点が重要である。以下では、そのような例を紹介したい。

3.1 頑健統計学およびデータ解析の視点—現状と評価

日常用語とは言えない「頑健性」という表現は、経済分析の分野では明確な定義を意識することなく使われているように見受けられる。ある期間の経済分析を、その前後の期間に適用してみて、その結果が大きく異なるかどうかを確認することは分析の姿勢として重要である。しかし、最小二乗法で推定したいくつかの結果を比較すると、いずれも外れ値の影響を同じように受けるため、似たような結果が得られることがある。すなわち頑健でない手法から得られた結果を単純に比較するだけでは不十分である。

経済モデルについての議論は観測結果を前提として行われるが、同時にデータ発生に関しても十分な注意を払わなければならない。最も単純な場合では、無作為性 (randomness) や独立性は暗黙に仮定されているし、分布に関する仮定や未知の母数に関するある種の事前情報も存在する。一方で、これらの仮定は厳密な意味では正しくない点に注意が必要で

あり、仮定のほとんどは漠然とした知識や経験に基づく判断を数理的に処理しやすい形で定式化したものと理解すべきである。このようなモデルの仮定が、実際の経済データ発生過程からわずかに異なる場合に、推定や検定などの統計的推論に大きな影響がない場合に限り、数理的な処理が容易な方法を選ぶことは正当化できる。このような意味で安定的な手法が頑健な手法である。

頑健統計学を体系化した Huber (1981 など) がときどき指摘している通り、統計的手法に関する「頑健 (robust)」および「母数によらない (nonparametric)」という性質は関係が深いものの、意味は異なる。たとえばピアソンの相関係数とケンドールの順位相関係数を比較すると、通常は前者は 2 変量正規分布の下でその確率分布を評価して推論が行われるのに対して、後者は連続性だけを仮定した分布に基づいた推論が行われる。この意味で前者はパラメトリックな推定量、後者は母数によらない推定量と呼ばれる。このように、母数によらない手法とは広い分布のクラス（連続かつ対称など）に対して同一の確率的な命題が成立することを指すのが一般的な用法であり、厳密な定義があるわけではない。「分布によらない (distribution-free)」という表現も同じ内容を指している。実用的な手法としては、順位に基づく検定や、それから導出される推定が広く利用されている。

しかし、母数（モデル）に基づいた推論と母数によらない（分布の仮定に依存しない）推論の区別は自明とは言えない。たとえば、算術平均 $\bar{x}$ は頑健でない推定量、中央値 m は頑健な推定量とされるが、モデルに依存するかどうかという基準を用いると、両者とも特定の分布を想定した場合の母集団平均 μ に対する最尤推定量 (MLE) であるから、パラメトリックな推定量とも言える。具体的には、母集団分布として正規分布 $N(\mu, \sigma^2)$ を仮定した場合が算術平均 $\bar{x}$ であり、母集団分布として $f(x) = (1/2\sigma)e^{-|x-\mu|/\sigma}$ という密度関数で表される両側指数分布を仮定した場合が中央値 m である。

一方、これらの推定量は母集団分布とは無関係な基準から導くことができるから、母数によらない推定量でもある。実際、観測値 $x_1, \dots, x_n$ が与えられたとき、 $\sum_i (x_i - a)^2$ を最小にすると $a = \bar{x}$ 、 $\sum_i |x_i - a|$ を最小にすると $a = m$ が得られる。

頑健な手法 頑健性 (robustness) という表現も、必ずしも厳密に定義されたものとは言えず、仮定がわずかに違っていた場合でも利用できる手法という程度の意味で用いられることが多い。Huber (1981) でも「あらゆる仮定からのわずかな相違 (deviations from the assumptions) に対する非過敏性 (insensitivity)」を頑健性と呼んでいる。

最も研究が進んでいる「分布に関する頑健性 (distributional robustness)」では、真の確率分布が、正規分布などの想定されたモデルからわずかに異なる場合の問題を扱う。頑健性については、分布の他にも独立性や回帰モデルにおける線形性など、さまざまな側面があり、これらの標準的な仮定が成立しない場合の影響の大きさを評価することも重要であるが、分布に関する頑健性に比較して、解決すべき多くの課題が残されている。

ほとんどの古典的な手法について、分布に関する頑健性に関しては、想定された分布（通常は正規分布）よりも真の分布の裾が広い (long-tailed) 場合に、推定や検定の結果が大きな影響を受ける傾向がある。他方で、真の分布の裾が短い場合には結論は大きな影響を受けない。このことから分布に関して頑健な手法と、外れ値に対して安定的な手法が同じ意味で用いられることがある。

実用的には、外れ値を除去した後で古典的な手法を用いることで、極端に悪い結果を避けることはできるが、以下の理由から、特に経済分析に関しては、問題はそれほど簡単ではない。

まず世帯の所得や企業の出荷額など、真の分布が正規分布と大きく異なる場合には、分布を反映した観測値を外れ値として除去しては正しい分析ができなくなる。このことは、適切な変数変換（EDA の再表現）を利用して分布を対称化する前の段階では、形式的に外れ値を判断することの危険性を示している。

次に、通常的手法を用いて外れ値を検出することは容易ではない。回帰分析における最小二乗法では、少数の外れ値があると、多くの観測値について残差が大きくなる一方で外れ値に対する残差はそれほど大きくないことがあり、通常に残差分析では外れ値を検出することは難しい。このような場合、経験的にも、頑健な手法を直接適用する方が、外れ値を除去してから古典的な手法を適用するよりも安全かつ効率的な手法である。

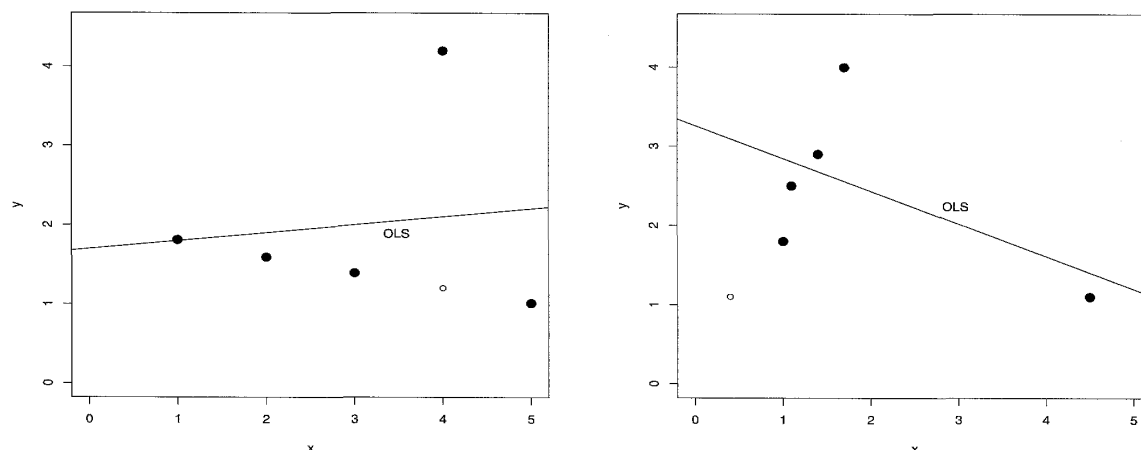
3.2 具体的な手法の例

講演では頑健統計学と経済統計との関わりをいくつかの手法を通じて紹介したが、ここには、最も一般的な手法の例を収録する。

回帰分析 回帰モデル $y = f(x) + u$ において回帰式を線形 $f(x) = \alpha + \beta x$ とし、誤差項 u に正規分布を想定する場合、仮定が正しければ α, β の推定量としては最小二乗法が最も優れている。

分布に関する頑健性の問題としては、誤差項 u が正規分布ではない場合、特に裾の広い分布の場合への対策がある。すなわち残差を $r_i = y_i - \alpha - \beta x_i$ と定義するとき、最小二乗法は $\sum r_i^2$ を最小にする解として α, β を定めるが、これは外れ値に敏感であり、頑健ではない。それに対して、分布に関して頑健な推定量は適当な ρ を用いて $\sum \rho(r_i)$ を最小とすることで得られる。位置の母数と同様、 $\rho' = \psi$ が有界（かつ単調）であれば頑健となる。

頑健な推定法の中で Huber の ψ として知られているものは刈込み平均に対応する ρ または ψ を用いたもので、minimax 問題の解として、優れた性質を持っている。Huber の推定法では図 3（左）に示す y 方向の外れ値に対しては適切な処理がなされるが、図 3（右）のような x 方向の外れ値に対しては、これだけでは対処が不十分である。いずれも白い丸印 $\circ$ が、誤りを取り除いた観測値を示している。図には通常最小二乗法による結果だけ

図3 y 方向の外れ値 (左) と x 方向の外れ値 (右)

を掲載している。

比較的簡単で高い頑健性を持つ手法として知られているものに最小刈込み二乗法 (least trimmed squares, LTS) がある。これは n 個の残差の 2 乗を大きさの順に並べて $(r^2)_{1:n} \leq \dots \leq (r^2)_{n:n}$ としたときに、適当な h ($h \leq n$) に対して $\sum_{i=1}^h (r^2)_{i:n}$ を最小にするように回帰係数を定めるものである。 $h = n$ なら最小二乗法と一致するが、 $h < n$ とすれば大きな r_i を無視できるため x 方向の外れ値に対処することが可能となる。 $h \approx n/2$ のとき、中央値と類似の頑健性を実現できる。

この手法は最近の統計解析プログラムで提供されており、比較的容易に実行できる。この他にも残差の 2 乗のメディアンを最少にする LMS (least median of squares) などが提案されているが、いずれも収束が遅く、効率も高くないこともあり、Huber (2009) は根本的な問題の検討が必要と指摘している。

結論として、経済分析では最小二乗法と頑健な手法を併用し、大きな差が見出された場合には元の観測値に戻って慎重な検討を実施することが現実的、かつ効果的である。

カーネル法の頑健化 回帰モデル $y = f(x) + u$ において任意の「なめらかな」関数 $f(x)$ を表現する手法として、加重関数 $K(x)$ を適当に選んで x の近くの観測値について y を加重平均することによって f を推定するのがカーネル法であり、応答関数 f の推定は

$$\hat{f}(x) = \frac{1}{nh} \sum K\left(\frac{x - x_i}{h}\right) y_i$$

と表される。 K としては標準正規分布や一様分布など (期待値を 0 とする) 密度関数が提案されている。

この手法はノンパラメトリック回帰と呼ばれ、頑健とされることが多いが、それは正しい表現ではない。たしかに適当に K を選べば $\hat{f}(x)$ を局所的に推定できるため、 x の値が離れた観測値の影響は小さい。しかし、加重算術平均である $\hat{f}(x)$ は、 x の値が近い観測値

については y 変数の外れ値の影響を強く受けるため、 y の分布に関する頑健性を持っていないことは明らかである。このことは、移動平均 (MA) が頑健でない一方、Tukey (1977) が提案した移動メディアン (Running Median) が頑健であるという関係と類似している。

このカーネル推定量を頑健化することは難しい問題ではない。そのためには上記の推定法は $\sum (y_i - \hat{f})^2 K((x - x_i)/h)$ を最小化していること、すなわち最小二乗法に対応することに注意すればよい。したがって適当な ρ を用いて $\sum \rho((y_i - \hat{f})/\sigma) K((x - x_i)/h)$ を最小化すれば、頑健な推定法が得られる。ここで σ は誤差項の標準偏差に相当する尺度である。このように考えれば、通常のカーネル推定量では $\rho(x) = x^2$ としているため、頑健でないことは明白である。

3.3 若干の留意点

最近では、各府省が公表する経済統計においても、伝統的な分類表や平均、合計、比率だけでなく、分位点などの分布情報が次第に公表されるようになってきた。中には箱ヒゲ図の利用など、公的な統計においても定着した手法もある。また、複雑な手法を利用することから、その理論的背景の確認や、新たな手法の検討も従来以上に求められるようになってきている。そのため、次節に記すように、統計学の専門家集団である学会としても、この分野で果たすべき役割が大きくなることは明らかである。

最後に、経済分野における頑健統計学の応用例を紹介して、一般的な課題を記そう。内閣府は経済状況の把握のために景気動向指数を作成、公表している。政府が景気を判断するほか、政策評価に際しても参考にされる指標である。従来は、景気動向指数のうちで DI (diffusion index) と呼ばれる指標が中心であったが、2008 年 4 月以降は CI (Composite Index) と呼ばれる指標を中心とする公表形態に変更された。統計的手法の視点からは、DI、CI のいずれも主要な経済変数の成長率を観測値とする位置の尺度であり、より具体的には、DI は成長率の中央値、CI は過去 5 年間の分布の平均と標準偏差で標準化した成長率の平均とみなすことができる。ただし、これは近似的な説明であり、詳細は内閣府の解説を参照されたい。

定義からわかるように CI は外れ値に対して脆弱であり、消費税の導入時など、商業販売額が一時的に大きく変動した場合に指数全体に影響を与えることが知られていた。そのためあって、以前は安定的な DI を重視してきたのが、海外の動向も参考にしながら CI への移行が図られた際、刈込み平均を利用するように変更された。ところで、時系列データに関する頑健推定量の構成は、それほど簡単ではなく、特に構造的変化への対応には課題が残されている。

実際、2008 年秋のリーマンブラザーズ破綻に端を発した金融危機の拡大局面では、ほとんどすべての経済変数が、従来の基準では外れ値となるほどの変化を示した。このような

時期には、過去の分布を参照して頑健化された CI は、経済全体の変化を捉えるためには過度に保守的になる。内閣府 CI の頑健化に際して原型を提示した筆者としても、早い機会に手法を見直す必要があるものと考えている。

この例のように、統計学の成果を反映した手法が作成され、実施に移された後でも、手法の適否に関する継続的な評価が必要であり、そのためにも、統計学の専門家による適切な助言が求められる。これも、次節に記す課題の一例である。

4. 公的統計における専門家の役割

2007 年に 60 年ぶりに改正された統計法は、日本の公的統計に関わる基本的な制度を規定するものであり、公的統計の体系的な整備に関連して、日本統計学会の会員にも密接に関わる点が少なくない。

新統計法では重要な基幹統計に関する基準などを定めている。これは利用価値の高い統計を作成するためには回答者の協力を得て正確な情報を収集することが前提となるためであり、調査票情報の適正管理義務、守秘義務などの規定も整備されている。

さらに新統計法では公的統計の二次的利用を促進することとしている。この点に関しては、専門家は利用者として質の高い統計を要求するだけでなく、統計作成過程に関しても意見を表明し、さらに技術的な協力を提供することが必要である。二次的利用の具体的な形として、委託集計や匿名標本データの提供が計画されているが、これらの実現のためには、分析に必要な情報を確保しながら実効性のある秘匿措置を講じなければならず、数理統計の専門的な知識が必要となる。

よりよい経済統計を国民の共有財産として利用可能とするために、統計学会の会員が公的統計の分野へ積極的に参入し、効果的な手法を提示することを期待したい。

4.1 新統計法の成立まで

舟岡・美添他 (2005) に記したように、統計制度の改革が必要とされた理由は、戦後日本の統計作成機関の特徴であった高度に分散的な統計制度が、時代の変化とともに利点より欠点が目立ってきたことにつくる。終戦直後の食糧確保や経済復興のために必要な統計は当時の農林省や通商産業省を中心として作成され、この時期の統計が日本経済の復興に果たした役割は大きい。しかし、最近のように第三次産業の比率が急激に増大している中でサービス産業の調査が不十分であるなど、産業構造の変化に対して統計整備の遅れが目立ってきた。それは、各省が政策的に必要とする統計を、省ごとに人的資源および予算を投入して作成していたのに対して、サービス業のようにその範囲が複数の省にまたがる統計に関しては、各省に分散する統計組織間の調整が十分になされなかったことが原因である。統計を国全体として体系的に整備するために統計基準を担当する組織が設けられてい

るものの、2001年の省庁再編に伴ってその権限が縮小されたこともあって、各省にまたがる統計調査に関しては不十分な調整しか実現できていないのが実情である。

4.2 統計データの利用促進

新統計法には、公的統計の体系的整備、統計データの利用促進と秘密の保護、統計委員会の設置、という3本の柱があるが、ここでは統計データの利用促進の問題を検討したい。

旧統計法には、指定統計を作成するために集められた調査票を統計の目的以外に利用してはならないという規定があり、総務大臣の承認を得て使用の目的を公示したものについてのみ、研究等での目的外利用が認められてきた。正確な統計を作成することを主な目的として収集された統計資料を別な形で利用することは想定されていなかったことは、コンピュータのない時代には自然である。統計調査に対して報告義務を課される世帯や事業所の信頼を得るために調査情報の厳重な管理が必要とされたことも当然である。

このことから旧統計法の下では、研究目的による統計利用であっても、統計作成機関との共同研究、委託研究などを除いて例外的にしか認められなかった。さらに許可を得る際にも、研究者の所属機関が国立大学か私立大学かによって手続きの煩雑さに格差があった。

その後、松田芳郎一橋大学教授（当時）を代表者とする科学研究費重点領域研究に多数の経済統計関係者が参加し、ミクロデータに関わるさまざまな課題を明らかにし、解決方向が提示されたこともあって、2000年頃までには科研費など公的な研究資金による研究であれば、目的に公益性があるものと認められ、利用が拡大するようになった。その経緯は松田（1999）、松田他（2000-2003）などに詳しく記されている。それでもなお、目的外利用の実現までには長期間かかる上、利用期間が短いなどの問題があり、諸外国と比べても例外的に保守的な状況であった。

これに対して、新統計法における基本理念では、公的統計を体系的に整備するために、「統計の作成に用いられた秘密は保護されなければならない」と情報の管理を規定した上で、公的統計は「広く国民が容易に入手し、効果的に利用できるものとして提供されなければならない」と定めている。また具体的な提供方法として、調査票情報（いわゆるミクロデータ）の提供、委託による統計の作成、および匿名データの提供を明示している。ミクロデータを利用するための資格要件が学術研究の発展に資すると認める場合と限定されているが、従来の目的外利用に比べて、制度として大幅な改善が実現されている。

委託集計 旧統計法では規定がなかった方法であるが、委託集計には個別情報を秘匿して、目的に応じた集計表を作成する機能があるため、統計作成機関としても個別情報の十分な管理ができる。利用者からの具体的な要望を想定しつつ、各府省の統計作成部局において然るべきシステムを構築する作業が必要であるため、この方法での提供が開始され定着するまでには時間がかかることが予想されるが、大きな可能性を持った制度であり、統計の

有用性も飛躍的に高まる。

匿名マイクロデータ 匿名データとは、調査票の回答者（事業所、世帯など）が特定されないように加工されたマイクロデータを指す名称であり、新統計法では、この形式によるマイクロデータを比較的広い範囲の研究目的に提供することを想定している。

この方法については、個別情報に関する秘匿措置の研究が必要であり、その整備は簡単なものではない。まず、大企業を含む企業や事業所のデータでは秘匿は不可能である。ほとんどの統計調査において、大規模な企業は全数調査の対象として含まれているため、資本金や従業員数などの公開情報を利用することによって、企業を特定することが可能となる。したがって、経営上の重要な情報が含まれるマイクロデータを公開することは、企業の名称が削除されたとしても認められない。

他方、世帯の場合には比較的同質なので、マイクロデータを安全に公開する工夫の余地がある。収入が極めて大きな世帯や世帯人数が大きな世帯は特定される危険性があるが、世帯主年齢や所得を階級表示にするなどの加工によって、個別情報を秘匿したマイクロデータを提供することが可能となる。

この分野は Statistical Disclosure Control と呼ばれ、日本も含めて理論的な研究が進んでいるが、日本においては公的統計の作成機関にこの分野の理論家を育てる余裕がなかったため、当面は学会としての貢献が期待される。日本統計学会としても、この分野の研究を推進し、研究者にデータを提供することは有益、かつ安全であることを実績を持って示すことが必要である。

匿名データの利用実績はこれから蓄積していくことになるが、すでに各国のデータを収録して分析に提供している Luxembourg Income Study のように、研究者がプログラムを送付し、結果を受け取る方法を開発することも可能である。その際、プログラムの誤りや不適切な分析を避けるためには、疑似データによって事前の分析を実行できることが望ましい。簡易なマイクロデータで分析した結果を踏まえて、さらに詳細な分析を必要とする研究者に向けた使いやすい仕組みを構築する必要がある。

4.3 統計の二次利用に関する検討課題

二次利用の具体的な運用は、各省および総務省（統計基準担当）の今後の対応によって定まる。その動向を注視して、必要に応じて学会からの要望を提示することが望ましい。一方で、研究目的でマイクロデータを利用する研究者の側にも、秘密の保護に関して規律が必要である。ASA など海外の学会に倣って、日本統計学会としても倫理規定を設けることを検討すべき時期である。

統計データを提供する費用の問題もある。新統計法では、手数料は実費を勘案して定めることになっているが、委託集計でも匿名データの提供でも、各府省とも基本的なシステム

を構築する段階である。その開発への協力など、学会としてさまざまな形で運用に関わっていくことが望ましい。なお、現時点の情報によれば、各府省とも 2009 年度中には二次利用のための予算・定員は組まれていない。また、ここでは論じないが美添 (2008) に記すように、データの効果的な提供体制のためには、統計データアーカイブを設立することが必須である。そのためにも担当者の手配が必要となり、予算の獲得など、実現は容易ではない。

今後、研究目的で多くの利用者が出れば、マイクロデータ提供の単価が下がるうえ、統計の二次的提供を担当する職員が配置される可能性が高まるなどの効果があり、結果としてデータ作成費用が軽減される。以上のように、技術的な問題に対する学会の支援体制を整え、統計作成部局と密接に連携することがますます重要になる。会員諸氏の大きな貢献を期待したい。

参 考 文 献

- Berger, J. O. (1995) Recent Developments and Applications of Bayesian Analysis, *ISI*, **IP1-1**, 3–14.
- Bernardo, J. M. and Smith, A. F. M. (1994) *Bayesian Theory*, Wiley.
- Birnbaum, A. (1962) On the foundations of statistical inference, *J. Am. Statis. Assoc.*, **57**, 269–326.
- Box, G. E. P. and Tiao, G. C. (1973) *Bayesian Inference in Statistical Analysis*, Addison-Wesley.
- 舟岡史雄・美添泰人他 (2005) 「ゼミナール 日本の統計改革」(8 月 4 日から 9 月 16 日まで連載) 日本経済新聞
- Huber, P. J. (1981, 2009) *Robust Statistics*, New York, Wiley (Second ed., with E. M. Ronchetti).
- Jeffreys, H. (1961) *Theory of Probability, Third edition*, Oxford.
- Lehmann, E. L. (1986) *Testing Statistical Hypotheses, Second edition*, Springer-Verlag (First ed., 1959).
- Lindley, D. V. (1965) *Introduction to Probability and Statistics from a Bayesian Viewpoint*, Cambridge University Press.
- Lindley, D. V. (1972) *Bayesian Statistics, a Review*, Philadelphia, Society for Industrial and Applied Mathematics.
- 松田芳郎 (1999) 『マイクロ統計データの描く社会経済像』, 日本評論社.
- 松田芳郎他 (編) (2000–2003) 『講座 ミクロ統計分析』, 日本評論社, (第 1 巻 統計調査制度とミクロ統計の開示, 第 2 巻 ミクロ統計の集計解析と技法, 第 3 巻 地域社会経済の構造, 第 4 巻 企業行動の変容) .
- Pratt, J. W. (1965) Bayesian Interpretation of Standard Inference Statements, *J. R. Statis. Soc. Ser. B*, **27-2**, 169–203.
- Pratt, J. W., Raiffa, R., Schlaifer, R. (1989) *Introduction to Statistical Decision Theory*, MIT Press.
- Press, S. J. (1989) *Bayesian Statistics: Principles, Models, and Applications*, Wiley.
- Savage, L. J. (1954, 1972) *The Foundations of Statistics*, Wiley (revised ed., Dover Publications).
- Shiller, R. J. (1973) A distributed lag estimator derived from smoothness priors, *Econometrica*, **41**, 775–778. Reprinted in S. E. Fienberg and A. Zellner, *Studies in Bayesian Econometrics and Statistics*, North-Holland, 1974.
- Tukey, J. W. (1977) *Explanatory Data Analysis*, Addison-Wesley.
- Wald, A. (1950) *Statistical Decision Functions*, Wiley.
- 美添泰人 (1983) 「ベイズの手法による統計分析」, 竹内 啓 (編) 『計量経済学の新展開』, 東京大学出版会, 第 6 章.
- 美添泰人 (1989) 「多重共線性へのベイズ・アプローチ」, 鈴木雪夫・國友直人 (編) 『ベイズ統計学とその応用』, 東京大学出版会, 第 7 章.
- 美添泰人 (1993–96) 「ベイズ統計入門 (その 1) から (その 19)」, ESTRELA.
- 美添泰人 (1996) 「ベイズ統計学はいつでも有用か」『統計』第 47 巻第 2 号
- 美添泰人 (2008) 「統計改革の残された課題」, 21 世紀の統計科学・第 1 巻『社会・経済の統計科学』, 東京大学出版会, 171–196.