

北海道大学 学生員 ○田高 淳
北海道大学 正 員 山形 耕一
北海道大学 学生員 田村 亨

1. はじめに

近年交通計画では、個人交通行動や意志決定過程のモデル化、交通サービスに対する利用者のニーズや評価の把握などが要請され、意識量が変量として取り上げられる機会が多い。しかしながら、意識量は、多くの場合、比例尺度や間隔尺度で定量的に計測し得る変量ではなく、変量間において大小や等値関係のみを定義し得る順序尺度に従う変量である。また、その調査過程においては、「満足・普通・不満足」といったカテゴリ値としてしか扱え得ない順序カテゴリデータである。従来、このようなデータの分析方法としては、Pearsonの χ^2 検定や評点法などが挙げられるが、前者は順序関係を無視しており、後者は評点の与え方に恣意性があるなど、有効な分析手法と成り得ていない。本研究では、順序カテゴリデータの分析に有効とされている累積法を取り上げ、順序カテゴリデータに対する検出力という見地から検討を行う。

2. 累積 χ^2 統計量

いま、N個の観測値がA、B2つの要因について表1のように分類されているとする。ここで、A要因とB要因による事象が独立かどうかを検定する。B要因の水準に順序関係がある場合、 m_{il} の累積データをもとに水準ごとの χ^2 統計量を求め、これを合計したものが累積 χ^2 統計量Qである。

$$Q = \sum_{l=1}^{L-1} \sum_{i=1}^r \left\{ \frac{(m_{il} - n_{il}\hat{p}_l)^2}{n_{il}\hat{p}_l} + \frac{\{(n_{il} - m_{il}) - n_{il}(1 - \hat{p}_l)\}^2}{n_{il}(1 - \hat{p}_l)} \right\} \quad \dots (1)$$

$$= \sum_{l=1}^{L-1} \sum_{i=1}^r \frac{(m_{il} - n_{il}\hat{p}_l)^2}{n_{il}} \left\{ \frac{1}{\hat{p}_l} + \frac{1}{1 - \hat{p}_l} \right\} = \frac{1}{\hat{p}_l(1 - \hat{p}_l)} \sum_{l=1}^{L-1} \sum_{i=1}^r \frac{(m_{il} - n_{il}\hat{p}_l)^2}{n_{il}}$$

一方、従来の χ^2 統計量は、次式で表わされる。

$$\chi^2 = \sum_{l=1}^L \sum_{i=1}^r \left(n_{il} - \frac{n_{i.}n_{.l}}{N} \right)^2 / \left(\frac{n_{i.}n_{.l}}{N} \right) \quad \dots (2)$$

3. シミュレーションによる分析

累積法の評価を検出力によって行う。検出力は、2つの標本から、それぞれの標本が抽出された母集団の確率分布に違いがあるかどうかを統計的な意味で識別する能力である。本研究では、2つの順序カテゴリ型母確率分布を仮定し、それぞれの母確率分布から乱数を用いて抽出された2つの標本から母確率分布の異同を検定する試行を多数回繰返すことにより、累積 χ^2 検定及び χ^2 検定の検出力をシミュレーション的に求めた。母確率分布を定めるために、意識量などを表わす潜在変量として正規分布を仮定し、この正規分布の平均・分散を規則的に変化させることにより母確率分布

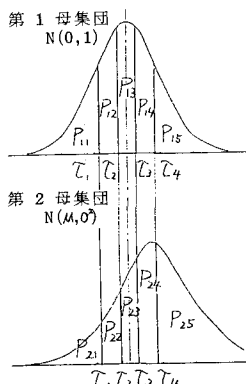


図 1. 2 標本の潜在分布モデル

表 1. 2 元分類表

A \ B	1	2	...	l	...	C	計
1							
2							
...							
l							
...							
r							
出現確率	$\hat{p}_l$	$1 - \hat{p}_l$					N

但し、 $l = 1, 2, \dots, C-1$
 $m_{il} = \sum_{j=1}^l n_{ij}$, $\hat{p}_l = \frac{m_{il}}{N}$

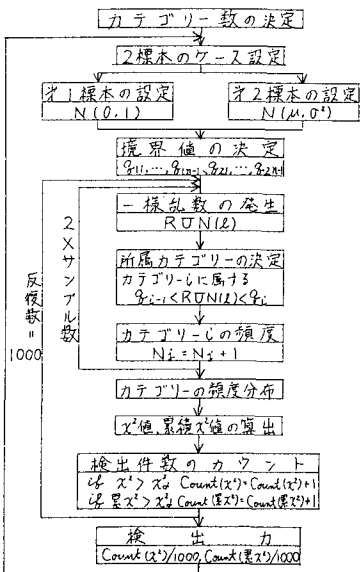


図 2. シミュレーションのフローチャート

の違いを表現している。すなわち、基準となる母確率分布（第1母集団と呼ぶ）としては、潜在変量の分布に正規分布 $N(0,1)$ をとった。そして、各カテゴリーの母確率が等しくなるようにカテゴリー境界値 τ_i を定めた。5カテゴリーの場合を考え、各カテゴリーの母確率を P_{1j} と表わすと、 $P_{11} = \dots = P_{15}$ より図1の如く境界値 τ_i が定まる。一方、比較対象する母確率分布（第2母集団と呼ぶ）の潜在変量の分布は $N(\mu, \sigma^2)$ とする。このとき、第2母集団の母確率 P_{2j} は、

$$P_{2j} = \int_{\tau_{j-1}}^{\tau_j} \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(z-\mu)^2}{2\sigma^2}} dz \dots (3) \quad (\text{但し、}\tau_0 = -\infty, \tau_5 = \infty)$$

となる。表2は、それぞれの μ 、 σ^2 における P_{2j} の値を示している。検出力に影響を及ぼす要因として、以下のものを取り上げ、図2のプロセスに従ってシミュレーション分析を行った。

表 2. 各第 2 母集団の確率分布

σ^2	μ	P_{21}	P_{22}	P_{23}	P_{24}	P_{25}
0.5	0.0	0.117	0.257	0.252	0.257	0.117
	0.2	0.070	0.202	0.243	0.302	0.182
	0.4	0.040	0.148	0.216	0.330	0.266
	0.6	0.021	0.100	0.178	0.335	0.366
	0.8	0.010	0.060	0.136	0.314	0.477
1.0	1.0	0.005	0.037	0.096	0.274	0.589
	0.2	0.149	0.186	0.177	0.228	0.261
	0.4	0.107	0.158	0.166	0.239	0.329
	0.6	0.075	0.129	0.151	0.241	0.405
	0.8	0.050	0.102	0.131	0.233	0.483
1.5	1.0	0.033	0.077	0.110	0.217	0.563
	0.0	0.246	0.180	0.147	0.180	0.246
	0.2	0.198	0.166	0.145	0.191	0.300
	0.4	0.155	0.149	0.140	0.197	0.359
	0.6	0.120	0.130	0.131	0.198	0.422
2.0	0.8	0.090	0.111	0.119	0.193	0.486
	1.0	0.066	0.092	0.106	0.184	0.551
	0.0	0.276	0.160	0.128	0.160	0.276
	0.2	0.231	0.151	0.127	0.167	0.325
	0.4	0.190	0.139	0.123	0.171	0.377
3.0	0.6	0.154	0.125	0.117	0.172	0.432
	0.8	0.123	0.111	0.109	0.169	0.488
	1.0	0.096	0.096	0.100	0.163	0.545
	0.0	0.314	0.134	0.105	0.134	0.314
	0.2	0.274	0.129	0.104	0.138	0.356
	0.4	0.237	0.122	0.102	0.140	0.399
	0.6	0.203	0.114	0.098	0.141	0.445
	0.8	0.172	0.105	0.094	0.139	0.490
	1.0	0.144	0.095	0.089	0.136	0.536

(i) 第2母集団 $N(\mu, \sigma^2)$ の μ 、 σ^2 の変化

$\mu = 0.1 \sim 1.0$ (0.2刻み)、 $\sigma^2 = 0.5, 1.0, 1.5, 2.0, 3.0$

(ii) 各標本の標本規模の変化 $n_1 = n_2 = 40, 60, 80, 100, 200$

(iii) 順序尺度データのカテゴリー数の変化 カテゴリー数 = 3, 4, 5

(iv) 有意水準の変化 $\alpha = 0.10, 0.05, 0.01$

4. 研究の成果と今後の課題

(1) 図3は、潜在変量の平均値の差と標本規模に対する（累積 z^2 検定- z^2 検定）の検出力の相対的な差を表わしている。累積 z^2 統計量と z^2 統計量を比較すると、潜在変量の平均値が異なる場合、すなわち、1次モーメントが異なる2つの確率分布の識別には累積 z^2 統計量が優れている。一方、潜在変量の平均値が同じで分散が異なる場合、すなわち、2次モーメントが異なる2つの確率分布の識別にはむしろ従来の z^2 統計量が勝っている。

(2) 潜在変量の平均値の差、標本規模及び検出力の関係を各有意水準ごとに図4の如く求めた。標本数の増加に伴い各順序カテゴリーへの多項分布が安定し、検出力が高まることが理解される。この図を用いることにより、個々の検定の信頼性、あるいは調査における標本数の決定に関して情報を得ることができる。

(3) 累積 z^2 統計量の検出力は、同じ平均値差に対して従来の z^2 検定の標本規模が約20ほど大きい標本の検出力に相当している。つまり、同程度の検出力を期待するのであれば累積 z^2 検定の採用により標本規模を20程度節約できる。

(4) 調査設計時における標本の規模決定は、検出力が平均値の差等により異なるため一概には言えないが、検出力0.5を目安とすると、カテゴリー数3~5では100程度は必要と考えられる。

今後は、人間がアンケート調査等で限られた順序カテゴリーデータに回答するときの思考過程を、実際の調査等を積み重ね経験的に把握し、これをモデルに組込んでより現実的な分析にしていく必要がある。

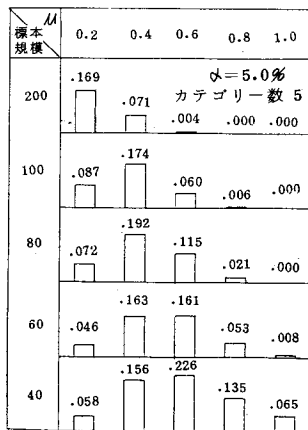


図 3. $\sigma^2=1.0$ の場合の (累積 z^2 検出力- z^2 検出力)

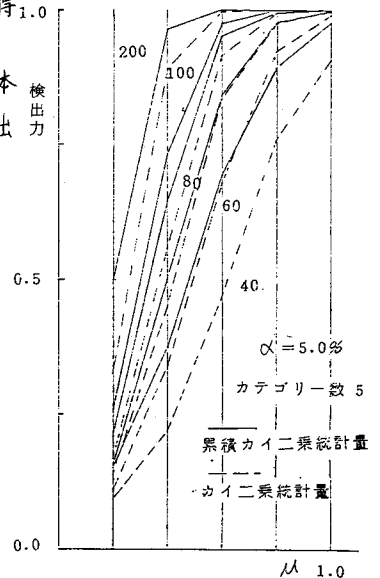


図 4. $\sigma^2=1.0$ の場合の μ と 標本規模の変化に伴う検出力