

分割表形式の誤診データの解析法について

名大 工学部 吉 村 功
中京大 体育学部 田 中 豊穂

1. 問題の所在

ある論文の中に、次頁に示す一節があった。ここでは、
“疾病 A が $\bar{A}$ と誤診される割合が約 15%, A でない疾病 $\bar{A}$ が A と誤診される割合が、同じく約 15%, したがって両方の誤診が相殺される”という結論が出されている。統計家にとっては自明なことであるが、このような結論のおろし方は、一般に正しくない。 A と $\bar{A}$ で病理解剖に付される個体の割合が違ふからである。そこで問題になることは、この種のデータをどう解析すべきかである。以下、これについて論じる。

2. 定式化

記号を表 2 および表 3 のように定める。表 2 において実際に観測できるのは、 $y_0, y_1, y_2, \dots, y_a$ である。 A_i と B_i は同じ疾病名とする。 A_0 すなわち B_0 を注目すべき疾病名とす

(1) 再不貧は死亡率の高い疾患であるだけに死亡統計による解析は疫学研究上重要である。一方死亡時診断に基礎をおく死亡統計は信頼度についてとかくの批判の、あるところである。

死は総合的な肉体的軽微であり、死因の総合もあって判断が容易でないことも事実である。この点を明らかにするために剖検例における臨床診断と病理診断の一致率を検討してみた。表1は1966～1970の5年間の日本剖検検体庫に掲載された剖検例について、臨床診断で再不貧、汎骨髄病、純赤血球性貧血とされたもの、および病理診断で再不貧と同程度の症例で、骨髄低(3)形成、骨髄病(萎縮)、脂肪腫とされたものをえらび、臨床診断と病理診断の一致率をみたものである。5年間の症例は946例、うち性不明9例をのぞく891(男468, 女423)である。臨床診断で再不貧群とされ

た男403例中、病理診断でも再不貧群は335(83.1%)、女では375例で323(86.1%)である。逆に病理診断で骨髄低形成、脂肪腫をふくめて再不貧群とすると、臨床診断との一致率は男400例中335(83.7%)、女371例中323(87.1%)となり、全体で85%は一致したといえてよい。このことは質的に85%の信頼性を示すとともに、誤診率がともに15%前後で相殺されるので、死亡統計の精度はそのままの故で再不貧死亡を示すことを物語っている。誤診の内容も臨床診断、病理診断がともに類似している。なお、臨床診断、病理診断がともに再不貧群であったものと、誤診をふくめた母集団の性、年齢別分布を比較したがほとんど同じであった。したがって、再不貧に関しては死亡統計はかなりの信頼性をもつものと言うことができよう。

表1 臨床診断または病理診断で再不貧または再不貧相当疾患と診断された例の一致率
(日本剖検検体庫 1966～1970)

病理診断		再 不 貧	汎 骨 髄 病	純 赤 血 球 性 貧 血	骨 髄 低 形 成 萎 縮	脂 肪 腫	無 顆 粒 球 症	白 血 病	がん・ 肉腫	貧 血 の 血 液 の 他 病	肝 疾 患	そ の 他	計
臨床診断													
男	再 不 貧	258	25	1	27	5	1	18	22	4	4	12	377
	汎 骨 髄 病		7		2	8			5				22
	純 赤 血 球 性 貧 血			2					2				4
	無 顆 粒 球 症	5	7		5	1							18
	粒 球 減 少 症	3	2										5
	白 血 病	12	7										19
	が ん ・ 肉 腫	10	3										13
	貧 血 其 他 血 液 病	2	1										3
	肝 疾 患	2											2
	そ の 他	5											5
計		297	52	3	34	14	1	18	29	4	4	12	468
女	再 不 貧	253	22		24	4		16	12	6	3	11	331
	汎 骨 髄 病	4	7		3		1	2				1	18
	純 赤 血 球 性 貧 血	2		3	1								6
	無 顆 粒 球 症	4	2		3								9
	粒 球 減 少 症	4											4
	白 血 病	10	2		1								13
	その 他 の が ん ・ 肉 腫	5	4										9
	貧 血 其 他 血 液 病	2	1										3
	肝 疾 患	0	1										1
	そ の 他	7	2										9
計		291	41	3	32	4	1	18	12	6	3	12	423

注(1) 再生不良性貧血のこと。(2) 人口動態統計の中で公表されている。

(3) 日本で病理解剖された全例が報告されている。

(4) 病理診断では、所見ともいうべき判断が列挙されていることもあって、死因分類(ICD)とは単純には対応していない。

表2. 死亡統計における頻度

		臨床診断での死因						計
		B_0	B_1	...	B_j	...	B_a	
真 の 死 因	A_0	n_{00}	n_{01}	...	n_{0j}	...	n_{0a}	n_0
	A_1	n_{10}	n_{11}	...	n_{1j}	...	n_{1a}	n_1
	$\vdots$	$\vdots$	$\vdots$		$\vdots$		$\vdots$	$\vdots$
	A_i	n_{i0}	n_{i1}	...	n_{ij}	...	n_{ia}	n_i
	$\vdots$	$\vdots$	$\vdots$		$\vdots$		$\vdots$	$\vdots$
	A_a	n_{a0}	n_{a1}	...	n_{aj}	...	n_{aa}	n_a
計		y_0	y_1	...	y_j	...	y_a	

表3. 剖検例における頻度

		病理診断での死因						計
		A_0	A_1	---	A_i	---	A_a	
臨床 診断 での 死 因	B_0	X_{00}	X_{10}	---	X_{i0}	---	X_{a0}	n_0
	B_1	X_{01}	X_{11}	---	X_{i1}	---	X_{a1}	n_1
	$\vdots$	$\vdots$	$\vdots$		$\vdots$		$\vdots$	$\vdots$
	B_j	X_{0j}	X_{1j}	---	X_{ij}	---	X_{aj}	n_j
	$\vdots$	$\vdots$	$\vdots$		$\vdots$		$\vdots$	$\vdots$
	B_a	X_{0a}	X_{1a}	---	X_{ia}	---	X_{aa}	n_a
計		X_0	X_1	---	X_i	---	X_a	n

る。表3の各変数の値はすべて観測できる。

データ解析法を考
えるとき、何を既知
定数、何を未知母数、
何を確率変数とする
かは立場あるいは視
点による。ここでは
英小文字で表わした
ものを既知定数、ギ
リシヤ文字で表わし
たものを未知母数、
英大文字で表したも
のを確率変数とみな
す。

医学上では、論議
のあるところである
が、ここでは、“真
の死因”は“病理診
断での死因”である

と仮定する。

上記各変数および母数は、いずれも統計上で、年齢、性、年度などの関数として定まるものであるが、ここでは簡単のため、年度 t のみの関数として考える。

普通、医学の研究者が知りたいのは $\mu_0(t)$ であるが、それは直接観測できないので、 $y_0(t)$ をもって $\mu_0(t)$ とみなし、諸種の解析を行っている。そこで、 $y_0(t)$ と $\mu_0(t)$ について、以下のような推定や検定の問題が生じる。

(i) $\mu_0(t)$ の推定 観測変数は異なる t の間で独立であるから、推定は各 t ごとに行えばよいが、後の検定の場合のように、 $\mu_0(t)$ を t の関数として推定しなければならない場合もある。

(ii) $\delta(t) \equiv \{y_0(t) - \mu_0(t)\} / \mu_0(t) = y_0(t) / \mu_0(t) - 1$ の推定 $\delta(t)$ は $\mu_0(t)$ を $y_0(t)$ で代用したときの誤差率というべき量である。

(iii) $H_0: \mu_0(t) = y_0(t)$ の検定 H_0 が棄却される場合には補正を考慮しなければならない。各 t ごとに考える場合には対立仮説として $\mu_0(t) \neq y_0(t)$ を考えることになるが、時には、次の仮説 H_1 において $\mu = 1$ というのがこの H_0 だと考えたり、あるいは H_2 において $\delta(t; t) \equiv 1$ としたのがこの H_0 だと考えたりすべきであらう。

(i) $H_1: \pi_0(t) = \gamma y_0(t)$, ただし γ は t に無関係な未知定数の検定
 医学では、しばしば、ある因子の経時変化と $\pi_0(t)$ の変化との相関を調べて因果関係の追究を行おうとする。そして実際には、 $\pi_0(t)$ の代りに $y_0(t)$ を用いて相関係数を計算する。このとき、もし H_1 が成立していれば、この代用は結果を誤らせないが、成立していなければ、相関係数は正しい値を示さない。対立仮説としては次の H_2 も考えられる。

(ii) $H_2: \pi_0(t) = f(t; \theta) y_0(t)$, ただし f の関数形は既知で θ が未知, の検定 $f(t; \theta)$ としては $t > \theta$ において $f(t; \theta) = 1$ となる階段関数, あるいは、 $\gamma(t)$ が指数関数になるようなものが考えられる。

以下では、モデルを想定して、これらの問題に対する推測の方式を考える。

3. ランダム抜取モデル

まず、年 t を固定する。したがって各変数を t で t を省略する。

大きさが y_j で、カテゴリー $(A_0, A_1, \dots, A_a)$ の構成が $(\pi_{j0}, \pi_{j1}, \dots, \pi_{ja})$, $\sum_i \pi_{ji} = y_j$, である集団 B_j からランダムに抜取られた大きさ n_j の標本における $(A_0, A_1, \dots, A_a)$ の観測

結果が $X_j = (X_{0j}, X_{1j}, \dots, X_{aj})$ であると想定する。 $X_0, X_1, \dots, X_a$ は互に独立とする。このとき X_j は多変量超幾何分布に従う。 $\tau_0 = \tau_{00} + \tau_{01} + \dots + \tau_{0a}$ であるから、 τ_0 の推定は、 $X_{00}, X_{01}, \dots, X_{0a}$ に基づくのが自然であろう。

$$E\{X_{0j}\} = \frac{n_j \tau_{0j}}{y_j}, \quad V\{X_{0j}\} = \frac{y_j - n_j}{y_j - 1} \frac{n_j \tau_{0j} (y_j - \tau_{0j})}{y_j^2} \quad (1)$$

であるから、 τ_0 の不偏推定量として

$$\hat{\tau}_0 = \sum_{j=0}^a (y_j / n_j) X_{0j} \quad (2)$$

が自然であり、このとき、一般に $y_j \gg 1$ だから

$$V\{\hat{\tau}_0\} = \sum_{j=0}^a \left(1 - \frac{n_j}{y_j}\right) \frac{\tau_{0j} (y_j - \tau_{0j})}{n_j} \quad (3)$$

したがって、 $H_0: \tau_0 = y_0$ を検定するには、(3) の τ_{0j} をその推定量でおきかえて

$$T = \frac{\left(\sum_{j=0}^a y_j X_{0j} / n_j - y_0\right)^2}{\sum \left(1 - \frac{n_j}{y_j}\right) \frac{y_j X_{0j}}{n_j^2} \left(y_j - \frac{y_j X_{0j}}{n_j}\right)} \quad (4)$$

を検定統計量とし、 $T > c$ という形の棄却域を作ればよい。

然し、 y_0 がいずれも大きければ、 H_0 のもとでの T の分布は自由度 1 のカイ二乗分布で近似できるが、その近似の精度や改善についてはまだ明瞭な結論が得られていない。

$\zeta = (y_0 - \tau_0) / \tau_0$ の推定には $\hat{\zeta} = y_0 / \hat{\tau}_0 - 1$ を用いるの

が自然であろうが、これは不偏でない。一般に $(\hat{\tau}_0 - \tau_0)/\tau_0$ は 1 にくらべてかなり小さいから、展開

$$\hat{\zeta} = \zeta - \frac{y_0}{\tau_0} \left\{ \frac{\hat{\tau}_0 - \tau_0}{\tau_0} - \left(\frac{\hat{\tau}_0 - \tau_0}{\tau_0} \right)^2 + \left(\frac{\hat{\tau}_0 - \tau_0}{\tau_0} \right)^3 - \dots \right\} \quad (5)$$

の第 2 項までをとると、

$$E\{\hat{\zeta}\} = \zeta + \frac{y_0}{\tau_0^3} \sum_{j=0}^a \left(1 - \frac{n_j}{y_j}\right) \frac{\tau_{0j}(y_j - \tau_{0j})}{n_j} \quad (6)$$

右辺第 2 項は正だから、この偏りを修正して

$$\hat{\zeta} = \zeta - y \sum_{j=0}^a \left(1 - \frac{n_j}{y_j}\right) \frac{y_j \tau_{0j}}{n_j^2} \left(y_j - \frac{y_j \tau_{0j}}{n_j}\right) / \hat{\tau}_0^3 \quad (7)$$

を推定量とすることもできる。(5) 式の展開の第 1 項までをとると、 $V\{\hat{\zeta}\} = \frac{y_0^2}{\tau_0^4} \sum_{j=0}^a \left(1 - \frac{n_j}{y_j}\right) \frac{\tau_{0j}(y_j - \tau_{0j})}{n_j}$ (8)

となる。 $V\{\hat{\zeta}\}$ の式は省略する。

ここで年々を考慮にいれて H_1 や H_2 を検定することを考えよう。局外母数が非常に多いので精密標本論にもとづく議論は困難である。そこで、 $\hat{\tau}_0(t)$ の分布を $N(\tau_0(t), \sigma^2(t))$,

$$\text{ただし, } \sigma^2(t) = \sum_{j=0}^a \left(1 - \frac{n_j}{y_j}\right) \frac{\tau_{0j}(y_j - \tau_{0j})}{n_j} \quad (9)$$

で近似できるものと仮定しよう。このとき、問題は、分散が異なる正規変量についての回帰の問題となる。 $\sigma^2(t)$ について

は、(9) 式中の τ_{0j} をその推定量で置きかえ、分散既知として検定を行うのがよいと思われるが、そのロバストネスなどについては、まだ明瞭な結論が得られていない。

4. 優先採取モデル

一般に解剖は、ランダムに行なわれるものではない。症状から見て死因の判別しにくいものが優先的に選ばれて剖検に付されることが多い。したがって前節のランダム採取モデルは必しも現実的な想定ではない。そこで、集団 B_j から大きさ n_j の標本を採取するとき、“設って B_j に判定されていた個体が採取される確率は、正しく B_j に判定されていた個体が採取される確率よりも大きい”というモデルを考えよう。すなわち、集団 B_j を構成する n_j 個の個体のうち、正しい判定の τ_{ji} 個は、選ばれる確率が π_j ($0 < \pi_j < 1$) だけ小さいものとする。

再び τ を固定し、省略して考える。厳密な分布はかなり複雑なものになるが、 n_j が十分大きく、 τ_{ji} がそれほど小さくなくて、多項分布による近似が可能とすると、 $X_j = (X_{0j}, X_{1j}, \dots, X_{ij}, \dots, X_{aj})$ の分布は次のようになる。

$$P\{X_{ij} = x_{ij}, i=0, 1, \dots, a\} = \frac{n_j!}{x_{0j}! x_{1j}! \dots x_{aj}!} \left(\frac{\tau_{0j}}{\tau_j}\right)^{x_{0j}} \left(\frac{\tau_{1j}}{\tau_j}\right)^{x_{1j}} \dots \left(\frac{\tau_{ij}}{\tau_j}\right)^{x_{ij}} \dots \left(\frac{\tau_{aj}}{\tau_j}\right)^{x_{aj}} \quad (10)$$

ただし, $\sum_{i=0}^a x_{ij} = n_j$, $\tau_j = \sum_{k \neq j}^a \tau_{kj} + \pi_j \tau_{jj}$, $\sum_{k=0}^a \tau_{kj} = y_j$.

ここで $\pi_j; j=0, 1, \dots, a$, は他の情報により値が定まるものと仮定しよう。そうでない場合には変数の次元より母数の次元が大きくなり推測が不可能になる。 π_j が与えられたとすると, $(\tau_{0j}, \tau_{1j}, \dots, \tau_{aj})$ に関する対数尤度 $\log L$ は次式になる。

$$\log L = \text{const} + \sum_{i=0}^a x_{ij} \log \tau_{ij} - n_j \log \tau_j \quad (11)$$

最尤推定量は

$$\hat{\tau}_{0j} = \pi_j y_j x_{0j} / \{ \pi_j n_j + (1 - \pi_j) x_{jj} \}, \quad j=1, 2, \dots, a \quad (12)$$

$$\hat{\tau}_{jj} = y_j x_{jj} / \{ \pi_j n_j + (1 - \pi_j) x_{jj} \}, \quad j=0, 1, 2, \dots, a \quad (13)$$

$$\begin{aligned} \hat{\tau}_0 &= \sum_{j=0}^a \hat{\tau}_{0j} \\ &= \sum_{j=0}^a \frac{y_j (x_{0j}/n_j)}{1 + \frac{1-\pi_j}{\pi_j} \frac{x_{jj}}{n_j}} + \frac{1-\pi_0}{\pi_0} \frac{y_0 (x_{00}/n_0)}{1 + \frac{1-\pi_0}{\pi_0} \frac{x_{00}}{n_0}} \end{aligned} \quad (14)$$

すべての π_j が 1 であればこれは前節のモデルに一致する。 π_j が 0 に近くなると x_{jj}/n_j も 0 に近くなり、 $\pi_j = 0$ では確率 1 で $x_{jj} = 0$ となる。一般にこのようなことはない。 $\hat{\tau}_0$ を推定値として計算するには、 $y_j, n_j, x_{0j}/n_j, x_{jj}/n_j, j=0, 1, \dots, a$ が求まっていればよい。 $x_{ij}, i \neq 0$, は計算に不要である。

(5) 式におけると同様を展開を行って 推定量の期待値や分散を求めると近似的に次のようになる。

$$\left. \begin{aligned}
 E\{\hat{\gamma}_j\} &\doteq \gamma_j + (1-\pi_j) \gamma_j \gamma_{jj} \gamma_j / (\pi_j n_j \gamma_j^2), \quad j=1, 2, \dots, a \\
 E\{\hat{\gamma}_{00}\} &\doteq \gamma_{00} + (1-\pi_0) \gamma_{00} (\pi_0 \gamma_{00} - \gamma_{00}) \gamma_{00} / (\pi_0 n_0 \gamma_{00}^2) \\
 E\{\hat{\gamma}_0\} &\doteq \gamma_0 + \sum_{j=1}^a \frac{(1-\pi_j) \gamma_j \gamma_{jj} \gamma_j}{\pi_j n_j \gamma_j^2} + \frac{(1-\pi_0) \gamma_{00} (\pi_0 \gamma_{00} - \gamma_{00}) \gamma_{00}}{\pi_0 n_0 \gamma_{00}^2} \\
 V\{\hat{\gamma}_j\} &\doteq \frac{\gamma_{jj} (\gamma_j - \gamma_{00})}{n_j} + \frac{(1-\pi_j) \gamma_{jj}^2 \gamma_{jj} (\gamma_j + \pi_j \gamma_j)}{\pi_j n_j \gamma_j^2} \\
 V\{\hat{\gamma}_{00}\} &\doteq \gamma_{00}^2 \gamma_{00} (\gamma_{00} - \pi_0 \gamma_{00}) / (\pi_0 n_0 \gamma_{00}^2) \\
 V\{\hat{\gamma}_0\} &\doteq \sum_{j=1}^a \frac{\gamma_{jj} (\gamma_j - \gamma_{00})}{n_j} + \sum_{j=1}^a \frac{(1-\pi_j) \gamma_{jj}^2 \gamma_{jj} (\gamma_j + \pi_j \gamma_j)}{\pi_j n_j \gamma_j^2} + \frac{\gamma_{00}^2 \gamma_{00} (\gamma_{00} - \pi_0 \gamma_{00})}{\pi_0 n_0 \gamma_{00}^2}
 \end{aligned} \right\} (15)$$

$H_0: \gamma_0 = \gamma_{00}$ の検定には、まず尤度比検定が考えられるが、 H_0 のもとでの最尤推定量を求めるのは困難である。むしろ、

$$T = (\hat{\gamma}_0 - \gamma_{00})^2 / \hat{V}(\hat{\gamma}_0) > c \quad (16)$$

ただし、 $\hat{V}(\hat{\gamma}_0)$ は $V\{\hat{\gamma}_0\}$ において母数をその推定量で置きかえたもの、

という形の棄却域のほうが実用的であろう。 γ'_{00} , n'_{00} がいずれも十分大きければ、この検定統計量の H_0 のもとでの分布は自由度 1 のカイ二乗分布で近似できるが、どの程度の大きさの γ'_{00} や n'_{00} でその近似が十分かは明らかでない。

$\gamma = \gamma_0 / \gamma_0 - 1$ の推定や H_1 , H_2 の検定についても、前節と全く同様な議論が成立するが、ここでは省略する。

5. モデル適用上の注意点

前節で論じたモデルにもとづいて誤診が死亡数統計にもたらす偏りを評価しようとするとき、最大の困難は、割引き係数 π の決定である。まずこれについて考えよう。

ある死亡例が剖検に付されるかどうかは多くの要因によって支配される。第一に、病理解剖を許されている機関と医師が限られていて、大学病院や大都市の大きな病院など、ある医療機関が関係する場合にのみ剖検の機会が大きい。そこに来る患者は何らかの偏った傾向を持つことが十分考えられる。

したがって、ここで考えているモデルは机上実験あるいは標本調査で意識的に導入され保証されるようなランダム性に裏付けられたものではなく、かなり巨視的な近似として想定されたものである。それゆえ π の決定の際には、特定の医療機関での慣習に立脚するのではなく、日本全体で、おおむねどのような傾向があるかを評価勘案すべきである。

第二に、剖検は、(i) 病変の本態が研究の焦点になっているときその手がかりを得るため、(ii) 症状から死因を判定するのが困難なとき死因を確かめるため、(iii) 治療の効果を確

認するため、などの目的でなされる。これらの目的のうち、
 (i) が大きな要因として働らくときは、その疾病に関して剖検率が高くなりかつ π も1に近い値になるが、(ii) が大きな要因として働らくときは、その疾病の鑑別の難易性にもよるが一般に剖検率が低くなりかつ π も小さい値になる。したがって巨視的に見た場合、剖検率が大なる疾病に対して π の値を大きくし、剖検率と単調な関係を持つように π を決定するのが妥当な近似と考えられる。

以上のように考察のもとで、次節の数値例では、 π を

$$\pi_i = \min_k (y_k/n_k) \pi_i / y_i \quad (17)$$

と設定した。これは剖検率が最大の疾病に対して π を1とし、剖検率に比例させて π を小さくするというものである。ただし、剖検総数が、剖検機関の増加などにもよってこの20年間で2~3倍に増加したことからわかるように、 $\min_k (y_k/n_k)$ は各年ごとに求めなければならぬ。

応用の際に生じるもう一つの問題点は疾病の分類併合である。死亡統計における死因について一般に使われている分類は、国際疾病分類(ICD)であるが、これを細分類のままを用いるのは現実的でない。その一つの理由は、剖検輯報での病理学的分類が必ずしもこれに一致していないことであり、他のもう一つの理由は、ICDの細分類をそのまま使うと、ほ

とんどすべての疾病 B_j に対して A_0 の設診数 X_{0j} が 0 となり、 n_j も非常に小さくなり、前節で検討した推測方式を無効にしていることである。

このようなわけで、 A_0 と鑑別が困難な疾病についてはある程度細かい分類を用い、他はいくつかの大分類に“その他”にカテゴリーを併合させることが必要になる。そこで問題になるのは正診数 X_{jj} の定義である。たとえば“悪性新生物”というカテゴリーを作ったとき、胃癌と脾臓癌との差を正診とするか設診とするか、もし設診とするならば、さらなどの程度の差までを設診とするかである。本報告における次第の数値例では、原則として ICD の分類で整数部分の番号が臨床診断と病理診断で一致しているものは正診として X_{jj} に繰り入れた。割引係数 π の役割から考えてこれが妥当と判断したわけである。このような取扱によらず、カテゴリーを併合して諸種の統計量を計算することは、併合された B_j について

$$\sum_j \frac{n_j X_{0j}}{n_j + \frac{1-\pi_j}{\pi_j} X_{jj}} \quad \text{を} \quad \frac{\sum_j n_j \sum_j X_{0j}}{\sum_j n_j + \frac{1-\pi}{\pi} \sum_j X_{jj}},$$

ただし、 π は $\min_j \pi_j < \pi < \max_j \pi_j$ なるある値、

で近似することを意味する。 n_j, n_{jj}, n_j が共通であれば問題はないが、一般には近似の設差が多少入り得る。

6. 数値例

剖検報は1958年分より継続して発行されている。ここでは1959年のデータを例示する。

$A \equiv B_0$ として, 1959年当時のICD No. 292 をとり, 他の疾病は次のカテゴリーに併合した。 $a=7$ 。

B_0 : No. 292 その他の型の貧血(溶血性, 再生不良性など)

B_1 : No. 293 詳細不明の貧血

B_2 : No. 295, 296 血友病, 紫斑病その他の出血病

B_3 : No. 290, 291, 294, 297, 298, 299 悪性貧血, 鉄欠乏性貧血, 多血球血症, 無顆粒球症, 脾臓疾患, その他の血液および造血器疾患。

B_4 : No. 200~203, 205 リンパ肉腫, ホジキン病, 多発骨髄腫, 菌状息肉腫

B_5 : No. 204 白血病

B_6 : B_4, B_5 以外の新生物

B_7 : 上記以外のすべて。

併合において考慮したことは次のとおり。各カテゴリーの剖検数が10をこえること。再生不良性貧血は当時、比較的新しく認識され原因や機序も明らかでなかったため詳細不明とされやすかったこと(B_1 を独立)。出血性の症状が目立つこと(出血性疾患 B_2 を独立)。血液病であることは認めやす

かったこと(残りの血液病を B_3 に一括)。白血病が当時急に
 になっていたためとの混同が考えられたこと (B_5 を独立)、
 骨髄腫やリンパ肉腫との混同が考えられたこと (B_4 を独立)、
 などである。

この分類にもとづいて必要な統計量を調べた結果と、それ
 にもとづく推定値とが表4のようになった。割検率最大なのは
 B_5 白血病の .1323 だから $\pi_5 = 1$ とし、以下

$$\pi_j = n_j / y_j \div 0.1323$$

として割引係数を定めてある。

ランダムモデルにもとづくと、

$$\hat{\eta}_0 = 950 \quad \hat{\eta}_0 - y_0 = 240 \quad \hat{v}(\hat{\eta}_0) = 25870$$

$$T = 2.23$$

優先モデルにもとづくと、

$$\hat{\eta}_0 = 738 \quad \hat{\eta}_0 - y_0 = 28 \quad \hat{v}(\hat{\eta}_0) = 3551$$

$$T = 0.22$$

という値が得られる。いずれのモデルでも H_0 は棄却されない。

両者の違いを表4で見ると、総数が大きくて割検率の小さい B_6, B_7 で、ランダムモデルでの推定値と推定分散が大きくなっている。経験的な実感においては、優先モデルのほうが現実に近いように思われる。

表4 1959年の剖検結果に基づく諸統計量の値

カテゴリー	X_{0j}	正診数 X_{1j}	剖検数 n_j	死亡数 y_j	剖検率 n_j/y_j	剖検係数 π_j	ランダムモデル $\hat{e}_{0j}$	$\hat{v}(\hat{e}_{0j})$	優先モデル $\hat{e}_{0j}$	$\hat{v}(\hat{e}_{0j})$
B_0	44	左に同じ	62	710	.0873	.660	504	1528	559	1107
B_1	2	1	12	317	.0379	.286	53	1122	44	911
B_2	3	9	21	745	.0282	.213	106	3134	41	622
B_3	4	26	47	909	.0517	.391	77	1293	42	440
B_4	0	44	195	1492	.1307	.988	0	0	0	0
B_5	3	250	329	2486	<u>.1323</u>	1.000	23	149	23	172
B_6	2	(2549)	(3357)	91806	.0366	.277	55	1448	18	165
B_7	1	(2654)	(4505)	591494	.0076	.057	132	17196	11	134
計			8528	689959			950	25890	738	3551

表中 y_j は人口動態統計より求めた。剖検総数の中には、警察医
勢院での事例、詳細不明と死産が含まれていない。

B_6, B_7 の X_{0j} と n_j は推定値である。すなわち、全剖検機関より
偏りが生じるように約半数を抽出し、そこで剖検例 5047 に
おける正診数 (B_6 で 1636, B_7 で 1704) と剖検数 (B_6 で
2155, B_7 で 2892) を求め、 $n_6 + n_7 = 8528 - 62 - 12 - \dots - 329 = 7862$
より $\hat{n}_6 = 7862 \times 2155 / 5047 \div 3357$, $\hat{n}_7 = 7862 \times 2892 / 5047 \div 4505$
 $\hat{X}_{66} = 3357 \times 1636 / 2155 \div 2549$, $\hat{X}_{77} = 4505 \times 1704 / 2892 \div 2654$
として求めた値である。

7. まとめと考察

病理解剖された事例の記録を、正剖と誤剖という点に着目して分割表形式に整理したとき、そのデータをどう解析すべきかについて、ランダム採取モデルと優先採取モデルとを想定し、検討を行い、解析法を提案した。解析法の精度や正当性については、統計学上と医学上とから検討しなければならない不明点が多く残されているが、一つの解析手段として実用可能であることは数値例を通して示されたと言えよう。解析結果もそれほど的確はすれではないと思われる。

数値例がわずか1年分にも満たなかったのは、剖検報告が事例の列挙に留まっていた統計処理の困難を形式的に示していることによる。日本の全部検例が生データとして公表されていることは高く評価すべきことであり、それに費される関係者の労力にも敬意を表さざるを得ないが、その貴重な資料が有効に利用されるために、たとえば臨床診断での死因がICD分類のどの項に対応するかの記事がほしいと感じられる。そうすれば死亡統計と対比がずうと容易となる。年間3万例にも及ぶ現在では、おそらく編集においてもカードやファイルが利用されているから、分類集計という統計処理をもう少し充実させることも、それほど難しいことではないのではあるか。

本報告で提案した解析法で経年変化の解析を行なおうとするならば、労力の関係で、剖検報中の事例の標本調査が不可避となる。その際、どんな標本抽出がよいか、解析法をどう修正すればよいかは今後の研究課題である。

今回は X のみ確率的変動を考えたが、立場や対象疾病の性質によっては、 π , n , Y 等についても確率的変動を考慮しなければならないこともある。これも今後の問題である。